

Toward Fine-Grained, Privacy-Preserving, Efficient Multi-Domain Network Resource Discovery

Qiao Xiang, Jingxuan Jensen Zhang, Xin Tony Wang, Yang Jace Liu, Chin Guok, Franck Le, John MacAuley, Harvey Newman, and Y. Richard Yang

Abstract—Multi-domain network resource reservation systems are being deployed, driven by the demand and substantial benefits of providing predictable network resources. However, a major lack of existing systems is their coarse granularity, due to the participating networks' concern of revealing sensitive information, which can result in substantial inefficiencies. This paper presents Mercator, a novel multi-domain network resource discovery system to provide fine-grained, global network resource information, for collaborative sciences. The foundation of Mercator is a resource abstraction through algebraic-expression enumeration (*i.e.*, linear inequalities/equations), as a compact representation of multiple properties of network resources (*e.g.*, bandwidth, delay, and loss rate) in multi-domain networks. In addition, we develop an obfuscating protocol, to address the privacy concerns by ensuring that no participant can associate the algebraic expressions with the corresponding member networks. We also introduce a super-set projection technique to increase Mercator's scalability. We implement a prototype Mercator and deploy it in a small federation network. We also evaluate the performance of Mercator through extensive experiments using real topologies and traces. Results show that Mercator 1) efficiently discovers available networking resources in collaborative networks on average four orders of magnitude faster, and allows fairer allocations of network resources; 2) preserves the member networks' privacy with little overhead; and 3) scales to a collaborative network of 200 member networks.

Index Terms—Multi-domain networks, resource discovery, privacy preserving.

I. INTRODUCTION

MANY of today's premier science experiments, such as the Large Hadron Collider (LHC) [2], the Square Kilometre Array (SKA) [3], and the Linac Coherent Light Source (LCLS) [4], rely on finely-tuned workflows that coordinate geographically distributed resources (*e.g.*, instrument, compute, storage) to enable scientific discoveries. An example of this is the movement of LHC data from Tier 0 (*i.e.*, the data center at European Organization for Nuclear Research, known as CERN) to Tier 1 (*i.e.*, national laboratories) storage sites around the world. This requires deadline scheduling to keep up with the amount of information that is continually generated by instruments when they are online. Another example is the "superfacility" model being developed by LCLS to allow streaming of data from instruments, across the Wide-Area Network (WAN), directly into supercomputers' burst buffers for near real-time analysis. The key to supporting these distributed resource workflows is the ability to reserve and guarantee network resources (*e.g.*, bandwidth) across multiple network domains to facilitate predictable end-to-end network connectivity. As such, several Research and Education (R&E) networks have deployed inter-domain circuit reservation systems. For example, the Energy Sciences Network (ESnet), a network supporting the LHC experiments, has deployed an On-Demand Secure Circuits and Advance Reservation System called OSCARS [5].

However, due to networks' concern of revealing sensitive information, existing systems do not provide a network interface for users to access network resource information (*e.g.*, network capabilities). Instead, they only allow users to submit requests for reserving a specific amount of resources (*e.g.*, a circuit providing a certain amount of bandwidth and delay), and return either success or failure [5]–[12]. This approach, which we call "probe requests" in the rest of this paper, often results in poor performance and fairness. Specifically, while solutions for reserving resources within a single member network, can be very efficient, solutions for discovering and reserving resources for correlated and concurrent flows across multiple member networks face unique challenges. In particular, solutions to reserving resources within a single administrative domain (*e.g.*, NetStitcher [13], SWAN [14] and B4 [15]) are often provided with the network's topology, and links' availability. In contrast, in a network with multiple administrative domains, because this information is typically considered sensitive, member networks do not reveal internal

Manuscript received December 15, 2018; revised June 27, 2019; accepted June 28, 2019. Date of publication July 5, 2019; date of current version August 6, 2019. This research was sponsored by the U.S. Army Research Laboratory and the U.K. Ministry of Defence under Agreement Number W911NF-16-3-0001. The views and conclusions contained in this document are those of the authors and should not be interpreted as representing the official policies, either expressed or implied, of the U.S. Army Research Laboratory, the U.S. Government, the U.K. Ministry of Defence or the U.K. Government. The U.S. and U.K. Governments are authorized to reproduce and distribute reprints for Government purposes notwithstanding any copyright notation hereon. This work was also supported in part by NSF under Award 1440745, Award 1246133, Award 1341024, Award 1120138, and Award 1659403, in part by Department of Energy (DOE) under Award DE-AC02-07CH11359, in part by DOE/Advanced Scientific Computing Research (ASCR) Project under Grant 000219898, in part by DOE/ASCR under Award DE-SC00155527 and Award DE-SC0015528, and in part by the Google Research Award. (*Corresponding authors: Qiao Xiang; Y. Richard Yang.*)

Q. Xiang, J. J. Zhang, X. T. Wang, and Y. R. Yang are with the Department of Computer Science, Yale University, New Haven, CT 06511 USA (e-mail: qiao.xiang@cs.yale.edu; jingxuan.zhang@yale.edu; xin.wang@yale.edu; yry@cs.yale.edu).

Y. L. Liu is with the Computer Science Department, University of Calgary, Calgary, AB T2N 1N4, Canada (e-mail: yang.liu5@ucalgary.ca).

C. Guok and J. MacAuley are with the Lawrence Berkeley National Laboratory, Berkeley, CA 94720 USA (e-mail: chin@es.net; macauley@es.net).

F. Le is with the IBM Thomas J. Watson Research Center, Yorktown Heights, NY 10598 USA (e-mail: fle@us.ibm.com).

H. Newman is with the Division of Physics, Mathematics and Astronomy, California Institute of Technology, Pasadena, CA 91125 USA (e-mail: newman@hep.caltech.edu).

Color versions of one or more of the figures in this paper are available online at <http://ieeexplore.ieee.org>.

Digital Object Identifier 10.1109/JSAC.2019.2927073

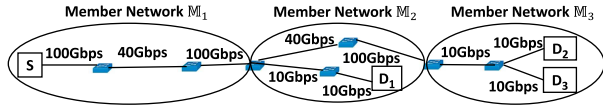


Fig. 1. A motivating example where a user wants to reserve bandwidth for three source-destination pairs: (S, D_1) , (S, D_2) and (S, D_3) , across 3 member networks M_1 , M_2 and M_3 .

network details to external parties. As a result, existing multi-domain reservation systems treat each member network as a black box, probe their available resources by submitting varied circuit reservation requests, and receive boolean responses. In other words, current solutions perform a depth-first search on all member networks, and rely on a trial and error approach: to reserve bandwidth, repeated, and varied attempts may have to be submitted until success.

To illustrate the limitations of existing systems, we consider a collaboration network composed of three member networks running OSCARS [5], as shown in Fig. 1. A user may submit a request to reserve bandwidth for three circuits, from source host S to destination hosts D_1 , D_2 and D_3 . Given the capabilities of the source host (*e.g.*, the source host may have a 100 Gbps network card), and to ensure fairness across the circuits, the user may request 33.33 Gbps for each circuit. Upon receiving this request, OSCARS processes the circuits sequentially, for example, in the order of (S, D_1) , (S, D_2) and (S, D_3) . For each circuit, it uses a depth-first search approach to probe if each member network can provide the requested bandwidth. In this example, there is no path with 33.33 Gbps of bandwidth from S to D_1 , and hence OSCARS notifies the user that this request fails.

The user can then adjust the requested bandwidth. However, with the limited feedback in OSCARS, the user does not know the amount of available bandwidth from S to D_1 . Consequently, the user may use a cut-to-half-until-reserved search strategy. As a result, after 12 attempts, the networks allocate 8.33 Gbps ($33.33 \rightarrow 16.67 \rightarrow 8.33$) for (S, D_1) , 8.33 Gbps ($33.33 \rightarrow 16.67 \rightarrow 8.33$) for (S, D_2) and 1.04 Gbps ($33.33 \rightarrow 16.67 \rightarrow 8.33 \rightarrow 4.17 \rightarrow 2.08 \rightarrow 1.04$) for (S, D_3) . In addition to requiring a large number of search attempts, the approach may obtain a bandwidth allocation that is far from optimal. For example, given the links' capacities and availability, a fair optimal bandwidth allocation is actually 5 Gbps for each circuit. Without a network interface to provide network resource information, designing an algorithm using existing systems to identify this solution can lead to substantially more complexity and churns.

Driven by the benefits of providing network resource information to users to improve the performance of network resource reservation systems, industry and academia have spent substantial efforts on designing network resource discovery systems to provide such information (*e.g.*, [16]–[25]). However, designing a network resource discovery system is a non-trivial task that requires addressing a series of challenges. First, the interface should provide a unified, accurate representation of the availability and sharing of multiple properties of network resources (*e.g.*, bandwidth, delay and loss rate) from multiple networks. Second, it should protect the

$\begin{aligned} M_1: \\ x_1^b + x_2^b + x_3^b &\leq 100, \\ x_1^b + x_2^b + x_3^b &\leq 40, \\ x_1^b + x_2^b + x_3^b &\leq 100, \end{aligned}$	$\begin{aligned} M_2: \\ x_2^b + x_3^b &\leq 40, \quad x_1^b \leq 10, \\ x_2^b + x_3^b &\leq 100, \quad x_1^b \leq 10, \end{aligned}$	$\begin{aligned} M_3: \\ x_2^b + x_3^b &\leq 10, \\ x_2^b &\leq 10, \\ x_3^b &\leq 10, \end{aligned}$
---	---	---

Fig. 2. Illustration of resource abstraction for the reservation request from Fig. 1.

privacy of networks by not exposing their private information (*e.g.*, topology, policy and capacity region). Third, it should not introduce too much computation and communication overhead to networks, and should scale to large multi-domain networks.

Existing resource discovery systems do not fully address these challenges. For example, resource discovery systems in grid-computing [16]–[23] only focus on the discovery of end-point resources (*i.e.*, computation and storage resources) and their availability for different services. Resource discovery systems in cloud computing (*e.g.*, CloudMirror [26], Pretium [27] and Amoeba [28]) adopt a network-does-all approach, in which users are provided with a more expressive interface for specifying requirements on data transfers and the network orchestrates resources between different user requests. Though this approach protects the privacy of the network, the network can only provide elastic resource reservations for user requests (*i.e.*, some requests may be preempted or rejected). Some recent systems (*e.g.*, the ALTO protocol [24], [25], [29] and the SENSE project [30], [31]) provide users the information of certain properties of network resources using the one-big-switch abstraction. While this approach protects the privacy of network, it cannot provide the accurate information of network resource sharing between flows (*e.g.*, bandwidth), which is critical for optimizing the emerging use cases (*e.g.*, large-scale collaborative data sciences).

In this paper, we present Mercator, a novel multi-domain network resource discovery system designed to address the limitations of current reservation systems and optimize multi-domain workflows. Mercator copes with the three aforementioned challenges for providing network resource information through three main components. The first and core component of Mercator is a resource abstraction through algebraic-expression enumeration (*i.e.*, linear inequalities and equations), which provides a compact, unifying representation of multiple properties (*e.g.*, bandwidth, delay and loss rate) of multi-domain network resources. For example, considering the same example of Fig. 1, the resource abstraction captures the constraints of bandwidths from all networks using the set of linear inequalities depicted in Fig. 2. Specifically, the variables x_1^b , x_2^b , x_3^b represent the available bandwidth that can be reserved for (S, D_1) , (S, D_2) and (S, D_3) , respectively. Each linear inequality represents a constraint on the reservable bandwidths over different shared resources by the three circuits. For example, the inequality $x_1^b + x_2^b + x_3^b \leq 100$ indicates that all three circuits share a common resource and that the sum of their bandwidths can not exceed 100 Gbps. With this set of linear inequalities, the user does not need to repeatedly probe the domains, but can immediately derive the bandwidth allocation to satisfy its own objective (*e.g.*, same rate for each transfer, different ratios according to demand ratios, or a fairness allocation such as max-min fairness).

Second, Mercator introduces a resource abstraction obfuscating protocol to ensure that member networks and other external parties cannot associate an algebraic expression with a corresponding member network, leading to a complete unified aggregation of multiple domains, appearing as much as possible as a single (virtual) network. Although such complete integration may not be needed in all settings, it can be highly beneficial in settings with higher privacy or security concerns. For example, in the scenario of Fig. 1, this protocol ensures that (1) the user cannot infer that the constraint $x_2^b + x_3^b \leq 10$ comes from network $\mathbb{M}_3$, and (2) that neither network $\mathbb{M}_1$ nor $\mathbb{M}_2$ knows the existence of this constraint. Finally, Mercator also introduces a super-set projection technique, which substantially improves the scalability and performance of Mercator through pre-computation and projection.

The main contributions of this paper are as follows:

- We identify the fundamental reason of the poor performance of current reservation systems for multi-domain data transfers as the lack of visibility of network information (e.g., topology and link availability) of each member network, and design Mercator, a novel multi-domain network resource discovery system, to address this issue;
- In Mercator, we propose a novel, compact resource abstraction to represent the network resource availability and sharing (e.g., bandwidth, delay and loss rate) among virtual circuit requests through algebraic-expression enumeration;
- We design a resource abstraction obfuscating protocol to prevent the user from associating the received algebraic expressions with their corresponding member networks;
- We develop a super-set projection technique to substantially improve the scalability of Mercator;
- We fully implement Mercator, deploy it in a small federation network, and also conduct extensive experiments using real network topologies and traces. Results show that Mercator (1) efficiently discovers available networking resources in collaborative networks on average six orders of magnitude faster, and allows fairer allocations of network resources; (2) preserves the member networks' privacy with little overhead; and (3) scales to a collaborative network of 200 member networks.

The remaining of this paper is organized as follows. We give an overview of Mercator in Section II. We give the details of the algebraic-expression-based resource abstraction in Section III. We discuss the resource abstraction obfuscating protocol and the super-set projection technique in Section IV and Section V, respectively. We introduce the implementation and deployment of Mercator in Section VI. We present the evaluation results of Mercator in Section VII. We discuss the related work in Section VIII and conclude the paper in Section IX.

II. MERCATOR OVERVIEW

This section presents the basic workflow and the architecture of Mercator, and a brief overview of its three main components: the resource abstraction through algebraic-expression enumeration, the resource abstraction obfuscating protocol and the super-set projection technique.

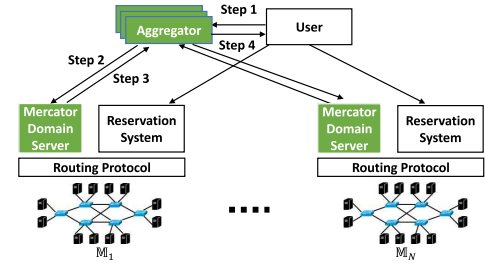


Fig. 3. The architecture and basic workflow of Mercator.

A. Basic Workflow

Mercator introduces and relies on a logically centralized aggregator, and a Mercator domain server in each member network. Consider a multi-domain network of N member networks $\mathbb{M}_i$, where $i = 1, \dots, N$ (Fig. 3). The basic workflow of Mercator to discover the multi-domain network bandwidth availability and sharing for a set of requested circuits is:

- *Step 1:* A user (e.g., an application) submits a resource discovery request for a set of circuits to the aggregator by specifying the source and destination endpoints of each circuit, and what properties of network resources he/she wishes to discover (e.g., bandwidth, delay and loss rate).
- *Step 2:* After authenticating and verifying the authorization of the request, the aggregator determines the member networks that the circuits traverse, and queries the Mercator domain server in each of these member networks to discover their resource abstractions. The determination of the relevant member networks for the aggregator to contact is further described in Section II-B.
- *Step 3:* Upon receiving the query from the aggregator, each Mercator domain server computes the resource abstraction (Section II-C, Section III) of the corresponding member network, and executes an obfuscating protocol (Section II-C, Section IV) to send the obfuscated resource abstraction to the aggregator.
- *Step 4:* The aggregator collects the obfuscated resource abstractions from the relevant member networks, and derives the original resource abstractions to present to the user. Based on the received information, the user determines the bandwidth allocation for each circuit, and sends a reservation request to the underlying reservation system.

The above workflow illustrates the main steps for a user to discover the available network bandwidth and properties for a set of circuits traversing multiple member networks. To further improve the scalability of Mercator, Section V introduces the super-set projection technique. It allows the aggregator to proactively discover the resource abstractions for a set of circuits between every pair of source and destination member networks, and project the pre-computed result to get the resource abstraction when receiving actual requests from users. The super-set projection technique can significantly reduce the delay, as well as number of messages, of resource discovery, and allows the aggregator to process multiple requests concurrently.

B. Architecture

This section describes the roles of the aggregator and Mercator domain servers in further details (Fig. 3).

1) *Aggregator*: The aggregator is the main interface of Mercator. It is responsible for authenticating and verifying the authorization of users' resource discovery requests (*e.g.*, through PKI [32]), querying Mercator domain servers in member networks to discover network resource information, and returning the collected abstractions to users. Depending on the specific requirements of different multi-domain networks, Mercator may adopt different authentication/authorization systems, *e.g.*, OpenID [33] and SAML [34]. We leave the detailed investigation of this issue in Mercator as future work.

The aggregator has connections to Mercator domain servers in all member networks. It also acts as a Border Gateway Protocol (BGP) [35] speaker, and has BGP sessions to all member networks. Consequently, given a request for a set of circuits F , the aggregator can infer the member-network path for each circuit, *i.e.*, the list of member networks a circuit will traverse, and the ingress points of the circuits to each member network¹ (as described in Step 1 of workflow). As such, for this request, the aggregator can also infer the set of circuits traversing and consuming resources in each $\mathbb{M}_i$, denoted as F_i . It can then queries the Mercator domain servers at each $\mathbb{M}_i$ by providing F_i and their ingress points to enter $\mathbb{M}_i$.

2) *Mercator Domain Server*: Given a Mercator domain server in member network $\mathbb{M}_i$, its primary role is to compute the resource abstraction of $\mathbb{M}_i$. To achieve it, Mercator follows the layering design principle to separate the routing protocol and the available network resources. In this way, given a set of circuits sent by the aggregator, their routes in $\mathbb{M}_i$ are computed and provided by the routing protocol in $\mathbb{M}_i$. The Mercator domain server in $\mathbb{M}_i$ takes these routes as inputs, and derives the available bandwidth and shared properties for the requested flows along those routes. After computing the abstraction, the Mercator domain server executes an obfuscating protocol to send the obfuscated resource abstraction to the aggregator, which addresses member networks' privacy concern.

C. Key Design Points

Having illustrated the high-level workflow of Mercator, we next give a brief overview on its key design points.

1) *Resource Abstraction Through Algebraic-Expression Enumeration (Section III)*: Mercator follows two important principles in human-computer interaction, *familiarity* and *uniformity*, to design a unifying abstraction that captures the properties (*e.g.*, available bandwidth, delay and loss rate) of resources shared – within and between member networks – by a set of requested circuits. This novel, compact resource abstraction is the core component of Mercator, and relies

on algebraic expressions (*i.e.*, linear inequalities/equations), a concept familiar to scientists and network engineers [37], to express the available bandwidth sharing for a set of requested circuits to be reserved.

Existing resource abstractions, including graph-based abstractions [38], [39] and the one-big-switch abstractions [24], [25], either fail to protect the private, sensitive information of each member network, or fail to capture the accurate resource availability and sharing between virtual circuit requests. In contrast, the resource abstraction of Mercator, expressed through algebraic-expression enumeration, naturally and accurately captures different properties (*e.g.*, bandwidth, delay and loss rate) of shared resources of a set of circuits without requiring member networks to reveal their network topology. Compared with the Boolean response of current resource reservation systems such as OSCARS, the user receives the complete resource feasible region of the collaboration networks for the requested circuits represented through algebraic expressions. A point in that feasible region represents a feasible allocation of resources for the different circuits in the request. In other words, the user can choose any point in the returned region as the parameters for the circuits to be reserved, using his own resource allocation strategy (*e.g.*, max-min fairness), and get predictable performance guarantee (*e.g.*, bandwidth, delay and loss rate).

2) *Resource Abstraction Obfuscating Protocol (Section IV)*: The algebraic-expression-based abstraction provides a compact, unifying representation of the multi-domain network resource information. It does not require member networks to reveal their network topologies and link availabilities. However, it does expose the resource feasible region of each member network (illustrated by the examples in Section I and Section III). Some member networks might prefer not to expose such information, as malicious parties may use it to identify links where to launch attacks (*e.g.*, DDoS). To address this issue, we develop a resource abstraction obfuscating protocol, which prevents the resource discovery aggregator from identifying the source of each received resource constraint. Specifically, the key idea consists of having each Mercator domain server obfuscate its own set of linear inequalities as a set of linear equations through a private random matrix of its own and a couple of random matrices shared with few other Mercator domain servers from other member networks (*e.g.*, through a consensus protocol), and then sends the obfuscated set of linear equations back to the aggregator using symmetric-key encryption, *e.g.*, Advanced Encryption Standard (AES) [40]. We demonstrate that from the received obfuscated equations, the aggregator can retrieve the actual resource feasible region for the circuits across member networks, but cannot associate any linear inequality with its corresponding member network. As a result, even if a malicious party obtains the resource feasible region across member networks, launching attacks to all member networks is much harder than attacking a particular member network.

3) *Super-Set Projection (Section V)*: To improve the scalability of Mercator, we introduce the super-set projection technique. The main idea consists of having the aggregator periodically query Mercator domain servers to discover the

¹In BGP glossary, such a path is also called an autonomous-system-path, or an AS-path, which is announced in BGP update messages along BGP sessions. The Route View Project [36] relies on a similar architecture with BGP speakers establishing sessions with hundreds of peering networks to collect BGP updates, and provides a real time monitoring infrastructure. In particular, we observe that the AS path for each destination prefix is currently already collected and made publicly available. As such, Mercator does not introduce additional privacy issues.

resource abstraction for a set of circuits between every pair of source and destination member networks. With these pre-computed abstractions, when a user submits a resource discovery request, the aggregator does not need to query the Mercator domain servers to compute the abstraction for each received request. Instead, the aggregator performs a projection on the precomputed abstractions based on the source and destination member networks of each circuit in the actual user request, to get the abstraction for this request. For example, consider a network of 2 member networks M_1 and M_2 . Using super-set projection, the aggregator queries the Mercator domain servers at both member networks about the bandwidth properties for a set of 2 circuits, one from M_1 to M_2 and the other from M_2 to M_1 , and gets a set of linear inequalities $\{x_{12}^b + x_{21}^b \leq 100, x_{12}^b \leq 50\}$. Suppose later a user submits a request for 1 circuit, with the source being an endpoint in M_2 and the destination being an endpoint in M_1 , to the aggregator. The aggregator projects the precomputed set of linear inequalities by removing all variables that are not x_{21}^b , and returns the result $\{x_{21}^b \leq 100\}$ to the user.

Such projection is much more efficient than having Mercator domain servers compute the abstraction for each received circuit request. With this technique, when a user submits a resource discovery request to the aggregator, the aggregator does not need to query Mercator domain servers (Step 2 in Section II-A), and the Mercator domain servers do not need to compute and obfuscate the resource abstraction for the request (Step 3 in Section II-A). Only when the user fails to reserve the resource based on the projected abstraction will the aggregator query the Mercator domain servers to obtain an up-to-date abstraction for the user. As such, servers in the aggregator pool can process requests concurrently (e.g., using optimistic concurrency control), significantly improving the scalability, fault-tolerance, and performance of Mercator.

After an overview of the key design points in Mercator, we discuss these designs in detail in the next few sections.

III. RESOURCE ABSTRACTION THROUGH ALGEBRAIC-EXPRESSION ENUMERATION

In this section, we give the details of the resource abstraction through algebraic-expression enumeration, the core component of Mercator. We first discuss the limitations of existing design options. Next we give the specifications of this abstraction, and how it handles important use cases, e.g., multicast, multi-path routing and load balancing, using the bandwidth property as an example. Then we discuss how the resource abstraction can represent other important properties of network resources (e.g., delay and loss rate), and how the resource abstractions from different member networks are aggregated to provide a unified representation of network resources.

A. Basic Issue

As illustrated by the example in Section I, the fundamental reason for the poor performance of existing circuit reservation systems is they are lack of the visibility of properties, e.g., bandwidth, of shared network resources for a set of circuits to be reserved. One may think of a strawman to let each

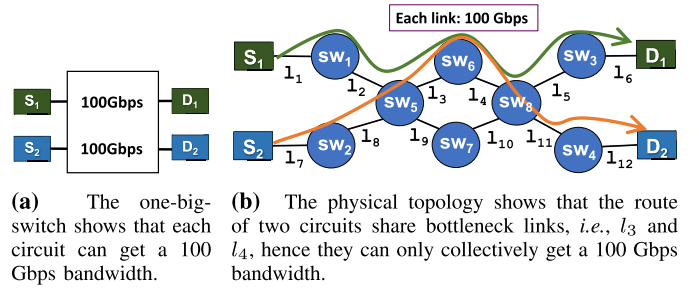


Fig. 4. A running example for illustrating the inefficiency of one-big-switch abstraction and the basic idea of resource abstraction through algebraic-expression enumeration, where two circuits (S_1, D_1) and (S_2, D_2) need to be reserved.

member network provide the full topology information to the aggregator in a graph-based abstraction [38], [39]. This design, however, exposes all the sensitive, private information of each member network, i.e., network topology and links' availability, to external parties, leading to security breaches.

A second strawman is to use a one-big-switch abstraction to provide simplified views of network information [24], [25], which protects the privacy of each member network. However, this abstraction fails to capture the information of shared resource among virtual circuit requests and thus is inaccurate. Consider the example in Fig. 4, where the user wants to reserve two circuits from S_1 to D_1 and S_2 to D_2 , respectively. Using the one-big-switch abstraction in the P4P system [25], the user will get the information that each circuit can reserve a bandwidth up to 100 Gbps (Fig. 4a). However, the routes for the two circuits – computed by the underlying routing protocol – share common links l_3 and l_4 (Fig. 4b), making it infeasible for both circuits to each reserve a 100 Gbps bandwidth.

In some recent studies [41], [42], a variation of the one-big-switch abstraction was proposed to define the resource sharing among different traffic flows as operations defined in different algebra fields. However, this abstraction is too complex and can only handle single-path routing policies.

B. Basic Idea

Different from the graph-based abstraction and the one-big-switch abstraction, the basic idea of the resource abstraction in Mercator is simple yet powerful: *given a set of requested circuits to be reserved, capture the properties (e.g., available bandwidth) of relevant shared resources, through a set of algebraic expressions.*

Specifically, suppose the Mercator domain server at a member network receives the resource discovery request of a set of circuits F entering this member network. For each circuit $f_j \in F$, we use x_j^b to denote the available bandwidth the user can reserve for this circuit. Upon receiving this request, the Mercator domain server first checks the intradomain route of each circuit f_j . Then the server enumerates all the links in the member network. For each link l_u , it generates a linear inequality:

$$\sum x_j^b \leq l_u.\text{bandwidth}, \quad \forall f_j \text{ that uses link } l_u \text{ in its route.}$$

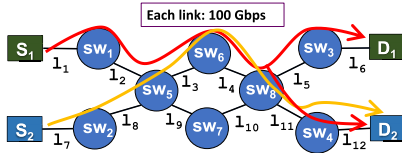


Fig. 5. A running example illustrating how the resource abstraction handles multicast through algebraic-expression enumeration, where two circuits $(S_1, \{D_1, D_2\})$ and (S_2, D_2) need to be reserved.

Revisit the example in Fig. 4, the Mercator domain server will generate the following set of linear inequalities $\Pi(F)$:

$$\begin{aligned} x_1^b &\leq 100 \quad \forall l_u \in \{l_1, l_2, l_5, l_6\}, \\ x_2^b &\leq 100 \quad \forall l_u \in \{l_7, l_8, l_{11}, l_{12}\}, \\ x_1^b + x_2^b &\leq 100 \quad \forall l_u \in \{l_3, l_4\}, \end{aligned} \quad (1)$$

which accurately captures the bandwidth sharing among two circuits' routes.

C. Removing Redundant Linear Inequalities

Observe the set of linear inequalities in the above example. One may realize that this set has redundancies, *e.g.*, there are 4 same inequalities $x_1^b \leq 100$ in this set. Given $\Pi(F)$, a linear inequality $y \in \Pi(F)$ is redundant if and only if the optimal solution of any optimization problem with $\Pi(F)$ as the constraint is the same as that with $\Pi(F) - \{y\}$ as the constraint. In our system, the Mercator domain server adopts a classic compression algorithm [43] to remove the redundant linear inequalities. In this example, the compressed $\Pi(F)$ will only contain one inequality, *i.e.*, $x_1^b + x_2^b \leq 100$.

Through algebraic-expression enumeration, the resource abstraction can handle not only unicast, as shown above, but many other settings. Below we show how resource abstraction handles three important use cases in collaborative data sciences.

D. Use Case 1 - Multicast

Consider the example in Fig. 5, where the first circuit is a multicast circuit from S_1 to D_1 and D_2 , and the second one is a unicast circuit from S_2 to D_2 . The routes for these circuits, computed by the underlying routing protocol, are marked in red and yellow, respectively. The resource abstraction captures the bandwidth sharing between these two circuits by introducing auxiliary variables x_{11}^b and x_{12}^b for the multicast circuit. Because the traffic duplication for the first circuit happens at switch 8, we use x_{11}^b to represent the traffic from switch 8 to D_1 , and x_{12}^b to represent the traffic from switch 8 to D_2 . In this way, the Mercator domain server will generate the following set of linear inequalities:

$$\begin{aligned} x_{11}^b &= x_1^b, \quad x_{12}^b = x_1^b, \\ x_1^b &\leq 100 \quad \forall l_u \in \{l_1, l_2\}, \\ x_{11}^b &\leq 100 \quad \forall l_u \in \{l_5, l_6\}, \\ x_2^b &\leq 100 \quad \forall l_u \in \{l_7, l_8\}, \\ x_1^b + x_2^b &\leq 100 \quad \forall l_u \in \{l_3, l_4\}, \\ x_{12}^b + x_2^b &\leq 100 \quad \forall l_u \in \{l_{11}, l_{12}\}, \end{aligned} \quad (2)$$

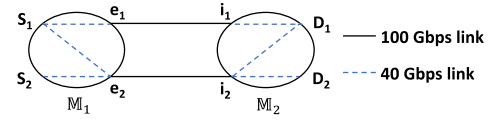


Fig. 6. A running example illustrating how resource abstraction handles complex routing and traffic engineering policies through algebraic-expression enumeration and how resource abstractions from different member networks are stitched, where two circuits (S_1, D_1) and (S_2, D_2) need to be reserved.

E. Use Case 2 - Multi-Path Routing

Consider the example in Fig. 6, where the user wants to discover the bandwidth sharing for two circuits $f_1 : (S_1, D_1)$ and $f_2 : (S_2, D_2)$, and M_1 uses multi-path routing for the circuit f_1 , *i.e.*, routing to two egresses e_1, e_2 .

In particular, the Mercator domain server at M_1 introduces variables x_{11} and x_{12} to represent the available bandwidth from S to egresses e_1 and e_2 , respectively, and share the introduction of these variables to M_2 . Then M_1 independently adds an equation $x_1 = x_{11} + x_{12}$ into its set of linear inequalities $\Pi_1(F)$. The resulting resource abstraction at both member networks are then expressed as

$$\begin{aligned} \Pi_1(F) : x_1^b &= x_{11}^b + x_{12}^b, & \Pi_2(F) : x_{11}^b &\leq 40, \\ x_{11}^b &\leq 40, & x_{12}^b &\leq 40, \\ x_{12}^b &\leq 40, & x_2^b &\leq 40, \\ x_2^b &\leq 40, & & \\ x_{11}^b &\leq 100, & & \\ x_{12}^b + x_2^b &\leq 100. & & \end{aligned} \quad (3)$$

Using $\Pi_1(F)$ and $\Pi_2(F)$ as the constraint, the user can then make reservation requests based on the optimization of her own objective function. For example, to achieve the max-min fairness between two circuits, the user will reserve $x_1^b = 80$ Gbps for (S_1, D_1) and $x_2^b = 40$ Gbps for (S_2, D_2) , where internally M_1 can allocate $x_{11}^b = x_{12}^b = 40$ Gbps.

F. Use Case 3 - Load-Balancing

In the same example in Fig. 6, assume M_1 uses weighted-cost-multi-path (WCMP) and has an internal policy to allocate bandwidth for the circuit (S_1, D_1) along two path $S_1 \rightarrow e_1$ and $S_1 \rightarrow e_2$ in a ratio of 1:2. With this policy, the previous reservation request with $x_1^b = 80$ Gbps and $x_2^b = 40$ Gbps is no longer valid as x_{11}^b and x_{12}^b cannot reach 40 Gbps simultaneously. To capture this policy so that the user does not make the invalid reservation request, the Mercator domain server at M_1 introduces an additional equation $x_{12}^b = 2x_{11}^b$ into $\Pi_1(F)$ and sends to the user. And the user can compute the valid, optimal reservation decisions, *e.g.*, $x_1^b = 60$ Gbps and $x_2^b = 40$ Gbps, to achieve max-min fairness.

G. Resource Abstraction for Other Properties

The algebraic-expression-based resource abstraction provides a generic representation for different properties of network resources. We now illustrate this generality by showing how it can represent two other important properties of network resources, delay and loss rate.

Specifically, suppose the Mercator domain server at a member network receives the resource discovery request of the delay and loss rate of a set of circuits F entering this network. For each circuit $f_j \in F$, we use x_j^d to denote the delay of this circuit in the network. The Mercator domain server checks the intradomain route of each f_j , and generates the following linear expression:

$$x_j^d = \sum_u l_u \cdot \text{delay}, \quad \forall l_u \text{ in the route of } f_j.$$

Similarly, for the property of loss rate, we use x_j^r to denote the loss rate of circuit f_j , and the Mercator domain server generates the following linear expression:

$$x_j^r = 1 - \prod (1 - l_u \cdot \text{lossrate}), \quad \forall l_u \text{ in the route of } f_j.$$

As such, the algebraic-expression-based resource abstraction is a unified representation of different properties of network resources for a set of circuits. Given $\Pi_i(F)$, the resource abstraction of $\mathbb{M}_i$ for a set of F circuits, from the geometric perspective, represents the resource feasible region of $\mathbb{M}_i$ for providing bandwidths, delays and loss rates to this set of circuits.

H. Aggregation of Multi-Domain Resource Abstraction

Given a set of F circuits spanning over N member networks, the resource abstractions $\Pi_i(F_i)$ from all the Mercator domain servers in the member networks can be aggregated into a unified, aggregated representation of multi-domain network resources $\oplus \Pi_i(F_i)$, where $\oplus$ is a property-specific operator.

Specifically, for the bandwidth property, the $\oplus$ operator is $\cup$, i.e., the union of multiple sets of linear inequalities. Geometrically speaking, $\cup \Pi_i^b(F_i)$ represents the intersection of the bandwidth feasible region of all member networks. For the delay property, the $\oplus$ operator is $\sum_d$, i.e., the sum of delays from different networks. In $\sum_d \Pi_i^d(F_i)$, for each circuit f_j and all member networks $\mathbb{M}_i$, $x_j^d = \sum \text{delay}_{ji}$, where delay_{ji} is the delay of circuit f_j in network $\mathbb{M}_i$. For the loss rate property, the $\oplus$ operator is $\sum_r$. In $\sum_r \Pi_i^r(F_i)$, for each circuit f_j and all member networks i , $x_j^r = 1 - \prod (1 - \text{lossrate}_{ji})$, where lossrate_{ji} is the loss rate of circuit f_j in network $\mathbb{M}_i$.

IV. PRIVACY-PRESERVING RESOURCE ABSTRACTION

Given a member network, the algebraic-expression-based resource abstraction specified in Section III accurately captures different properties of the available network resources among virtual circuits without exposing its network topology and links' availability. However, a resource abstraction still represents the resource feasible region of the corresponding member network for a set of circuits. Such information is still private and sensitive, and a malicious party who acquires it may use it to launch attacks to the corresponding member network. To address the privacy challenge for network resource discovery and preserve the privacy of resource feasible region of member networks while still providing the accurate network resource information for circuits, we extend the base resource abstraction to develop an obfuscating protocol in Mercator.

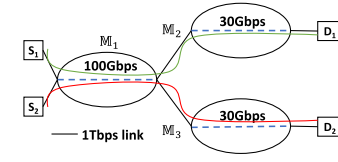


Fig. 7. A running example to illustrate the resource abstraction obfuscating.

In this section, we first formally define the privacy-preserving resource abstraction problem. Next, we present the details of our protocol and conduct a rigorous analysis.

A. Privacy-Preserving Resource Abstraction Problem

1) *Basic Issue*: We use the example in Fig. 7 to illustrate the privacy concern of the resource abstraction, where Mercator tries to discover the shared bandwidth of two virtual circuits (S_1, D_1) and (S_2, D_2) across 3 member networks. In this example, all links in black line are 1 Tbps aggregating links. The inter-member-network-paths of two circuits are $[\mathbb{M}_1, \mathbb{M}_2]$ and $[\mathbb{M}_1, \mathbb{M}_3]$, respectively. And two circuits share the same intra-domain path in $\mathbb{M}_1$.

When receiving the resource discovery request, the Mercator domain server at each member network will abstract the bandwidth sharing of both circuits into a set of linear inequalities. After removing the redundant inequalities of each member network, the resource abstraction of each member network's bandwidth is:

$$\begin{aligned} \Pi_1^b(F_1) &: \{x_1^b + x_2^b \leq 100\} \\ \Pi_2^b(F_2) &: \{x_1^b \leq 30\} \\ \Pi_3^b(F_3) &: \{x_2^b \leq 30\}. \end{aligned} \quad (4)$$

If each Mercator domain server directly sends its own resource abstraction to the aggregator, the aggregator will have the knowledge of the resource feasible region of each individual member network. This makes the whole collaboration network vulnerable because the aggregator is a single point of failure possessing the private information of all member networks. In other words, if an attacker gains the control to the aggregator, he can leverage such specific information to attack any member network.

2) *Problem Definition*: To make Mercator functional and secure, therefore, we need a solution that provides the accurate network resource information for the set of virtual circuits to be reserved, and at the same time protects each member network from exposing its private resource feasible region. To this end, we first give a formal definition of privacy-preserving, equivalent resource abstraction:

Definition 1 (Equivalent, Privacy-Preserving Resource Abstraction): Given a set of circuits F that span over $N > 1$ member networks, the resource abstraction $\Pi_p(F)$ collected by the aggregator is equivalent and privacy-preserving if for all network resource properties (e.g., bandwidth, delay and loss rate), (1) the resource feasible region represented by $\Pi_p(F)$ is the same as that represented by $\oplus \Pi_i(F_i)$ where $i = 1, 2, \dots, N$; and (2) for any linear inequality $c \in \Pi_p(F)$, the aggregator cannot associate it with a particular member network.

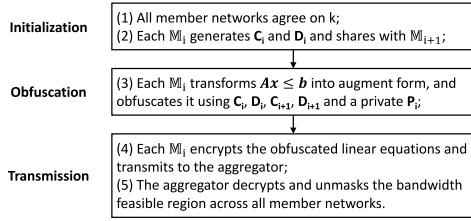


Fig. 8. The resource abstraction obfuscating protocol.

With this definition, we further define the privacy-preserving resource abstraction problem:

Problem 1 (Privacy-Preserving Resource Abstraction Problem): Given a set of circuits F that span over $N > 1$ member networks, design a security protocol in the resource discovery system to ensure that for all network resource properties (e.g., bandwidth, delay and loss rate), (1) the aggregator receives the equivalent, privacy-preserving resource abstraction $\Pi_p(F)$; and (2) for any M_i , it does not know any linear inequality from any other $\Pi_j(F_j)$, where $j \neq i$.

3) *Security Model:* In this paper, we assume a *semi-honest* security model, i.e., the aggregator and all member networks will not deviate from the security protocol, but merely try to gather information during the execution of the protocol [44]. This is sufficient for collaboration science networks where member networks share resources to collaboratively conduct common tasks such as data transfers, storage and analytics.

B. Resource Abstraction Obfuscating Protocol

There are different design options for Problem IV-A.2, e.g., garbled circuit based protocols [45]. However, these designs would incur expensive computation and communication overhead, hence are not suitable for the need of multi-domain resource discovery. In this paper, we tackle this problem by designing a novel resource abstraction obfuscating protocol that only requires simple operations on matrices, i.e., addition and multiplication.

1) *Basic Idea:* Our protocol leverages random matrix theory [46], [47]. In particular, each M_i independently computes and sends to the aggregator a set of disguised *linear equations*, which are derived from the private $\Pi_i(F_i)$, a random matrix P_i known only to M_i , two random matrices C_i and D_i known only to M_i and M_{i-1} , and two random matrices C_{i+1} and D_{i+1} known only to M_i and M_{i+1} .

2) *Protocol:* The protocol is composed of three phases: initialization, obfuscation and transmission, as shown in Fig. 8. For the simplicity of presentation, we let $m_i = |\Pi_i(F_i)|$, i.e., the number of linear inequalities in $\Pi_i(F_i)$ after redundancy removal, and $M_i = \sum_{j=1}^i m_j$. And for each circuit f_j , we also omit the superscript representing different properties in the corresponding x_j . As such, the resource abstraction of a member network i is written as $\Pi_i(F_i) = A_i x \leq b_i$.

During the *initialization phase*, all member networks agree on a common $k > \sum m_i$. For each M_i where

$i = 1, 2, \dots, N-1$, it generates a k -by- $(|F| + m_i + m_{i+1})$ random matrix $C_i = [C_i^{|F|} \ C_i^{m_i} \ C_i^{m_{i+1}}]$, and a k -by-1 random matrix D_i , and sends to M_{i+1} . And we define C_0 , D_0 , C_N and D_N as zero matrices. As we will illustrate in the remaining of this section, these zero matrices are used for presentation completeness and will not affect the correctness of the obfuscating protocol.

During the *obfuscation phase*, each M_i introduces m_i slack variables, denoted by x_i^s , to transform $\Pi_i(F_i) = A_i x \leq b_i$ from the standard form to the augment form and gets the following equivalent linear system:

$$[A_i \ I_{m_i}] [x, x_i^s] = b_i. \quad (5)$$

We then add slack variables introduced by all other member networks with zero coefficients into the linear system in Equation (5) and get the following equivalent linear system:

$$[A_i \ 0_{M_{i-1}} \ I_{m_i} \ 0] [x, x_1^s, \dots, x_i^s, \dots, x_N^s] = b_i. \quad (6)$$

Next, each M_i generates a private random matrix $P_i \in R^{k \times m_i}$, and left-multiplies both sides of Equation (6) to get:

$$[P_i A_i \ 0_{M_{i-1}} \ P_i \ 0] [x, x_1^s, \dots, x_i^s, \dots, x_N^s] = P_i b_i. \quad (7)$$

Then each M_i adds

$$[C_i^{|F|} - C_{i-1}^{|F|} \ 0_{M_{i-2}} \ -C_{i-1}^{m_{i-1}} \ -C_{i-1}^{m_i} + C_i^{m_i} \ C_i^{m_{i+1}} \ 0],$$

to the coefficient matrix of the left-hand-side (LHS) of Equation (7), and adds $-D_{i-1} + D_i$ to its right-hand-side (RHS) to get Equation (8), as shown at the bottom of this page, where it can be observed that for each M_i , the coefficient matrix of LHS of Equation (8) is of dimension k -by- $|F| + M_N$, and the RHS is of dimension k -by-1.

In the *transmission phase*, each M_i encrypts the set of linear equations in Equation (8) using a symmetric-key algorithm, e.g., AES, and sends the cypher text to the aggregator. After collecting the linear equations from all member networks, the aggregator decrypts them and computes the sum of all LHS matrices and RHS matrices of all member networks, respectively. After simple elimination, the LHS sum is expressed as:

$$[\sum P_i A_i \ P_1 \ \dots \ P_N].$$

Similarly, the sum of all RHS matrices of all member networks can be expressed as $\sum P_i b_i$. Denoting $[x_1^s, \dots, x_N^s]$ as x^s , the aggregator can get the privacy-preserving abstraction $\Pi_p(F)$:

$$[\sum P_i A_i \ P_1 \ \dots \ P_N] [x, x^s] = \sum P_i b_i. \quad (9)$$

3) *Example:* We use the example in Fig. 7 to illustrate the resource abstraction obfuscating protocol. For simplicity,

$$[P_i A_i + C_i^{|F|} - C_{i-1}^{|F|} \ 0_{M_{i-2}} \ -C_{i-1}^{m_{i-1}} \ P_i - C_{i-1}^{m_i} + C_i^{m_i} \ C_i^{m_{i+1}} \ 0] \cdot [x, x_1^s, \dots, x_i^s, \dots, x_N^s] = P_i b_i - D_{i-1} + D_i \quad (8)$$

we assume three member networks agree on $k = 4$. The private random matrices $\mathbf{P}_1$, $\mathbf{P}_2$ and $\mathbf{P}_3$ are generated as $\mathbf{P}_1 = [11, 49, 95, 34]$, $\mathbf{P}_2 = [58, 22, 75, 25]$, and $\mathbf{P}_3 = [50, 69, 89, 95]$. The obfuscated resource abstractions computed by each network are:

$$\begin{aligned} 15x_1 + 14x_2 + 15x_{11}^s + 4x_{21}^s + 0x_{31}^s &= 1130, \\ 53x_1 + 50x_2 + 53x_{11}^s + 2x_{21}^s + 0x_{31}^s &= 4910, \\ 96x_1 + 97x_2 + 96x_{11}^s + 4x_{21}^s + 0x_{31}^s &= 9540, \\ 38x_1 + 37x_2 + 38x_{11}^s + 1x_{21}^s + 0x_{31}^s &= 3420, \\ -2x_1 + 47x_2 + 0x_{11}^s + -3x_{21}^s + 47x_{31}^s &= 1470, \\ -4x_1 + 68x_2 + 0x_{11}^s + -4x_{21}^s + 68x_{31}^s &= 2040, \\ -4x_1 + 85x_2 + 0x_{11}^s + -3x_{21}^s + 86x_{31}^s &= 2630, \\ -4x_1 + 91x_2 + 0x_{11}^s + -2x_{21}^s + 94x_{31}^s &= 2810, \end{aligned}$$

and

$$\begin{aligned} 56x_1 + 0x_2 + -4x_{11}^s + 57x_{21}^s + 3x_{31}^s &= 1740, \\ 22x_1 + 0x_2 + -4x_{11}^s + 24x_{21}^s + 1x_{31}^s &= 680, \\ 78x_1 + 2x_2 + -1x_{11}^s + 74x_{21}^s + 3x_{31}^s &= 2250, \\ 25x_1 + 0x_2 + 0x_{11}^s + 25x_{21}^s + 0x_{31}^s &= 770. \end{aligned}$$

Summing these obfuscated resource abstractions together, the resulting resource abstraction $\Pi_p(F)$ collected by the aggregator is:

$$\begin{aligned} 69x_1 + 61x_2 + 11x_{11}^s + 58x_{21}^s + 50x_{31}^s &= 4340, \\ 71x_1 + 118x_2 + 49x_{11}^s + 22x_{21}^s + 69x_{31}^s &= 7630, \\ 170x_1 + 184x_2 + 95x_{11}^s + 75x_{21}^s + 89x_{31}^s &= 14420, \\ 59x_1 + 129x_2 + 34x_{11}^s + 25x_{21}^s + 95x_{31}^s &= 7000, \end{aligned}$$

where x_{11}^s , x_{21}^s and x_{31}^s are slack variables. Assume the user's objective is to maximize the throughput, i.e., $x_1 + x_2$. Using this set of linear inequalities as the constraint, it can get the optimal solution where $x_1 = x_2 = 30$ Gbps, the same as when using Equation (4) as the constraint.

C. Analysis

We conduct rigorous analysis on different properties of the proposed obfuscating protocol.

1) *Correctness*: We first study the correctness of this protocol. In particular, we prove the correctness of this protocol for different properties in the following propositions.

Proposition 1 (Bandwidth Resource Abstraction Equivalence): *If the resource abstraction $\mathbf{Ax} \leq \mathbf{b}$, where $\mathbf{A} = [\mathbf{A}_1, \mathbf{A}_2, \dots, \mathbf{A}_N]$ and $\mathbf{b} = [\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_N]$, represents the bandwidth property for a set of circuits F over N member networks. Using the proposed obfuscating protocol, the bandwidth feasible region of represented by Equation (9) is the same as the bandwidth feasible region represented by $\mathbf{Ax} \leq \mathbf{b}$.*

Proof: To prove this proposition, we first observe that the bandwidth feasible region of $\mathbf{Ax} \leq \mathbf{b}$ is the same as that of

$$[\mathbf{A} \quad \mathbf{I}_{M_N}] [\mathbf{x}, \mathbf{x}^s] = \mathbf{b} \quad (10)$$

Representing $\mathbf{P} = [\mathbf{P}_1, \dots, \mathbf{P}_N] \in R^{k \times M_N}$, we first observe that $[\sum \mathbf{P}_i \mathbf{A}_i \quad \mathbf{P}_1 \dots \mathbf{P}_N] = \mathbf{P} [\mathbf{A} \quad \mathbf{I}_{M_N}]$, and that $\sum \mathbf{P}_i \mathbf{b}_i = \mathbf{Pb}$ [47]. It is easy to see that when $[\mathbf{x} \quad \mathbf{x}^s]$ satisfies Equation (10), it also satisfies Equation (9).

Next, from the results in [46] and that $\mathbf{P} \in R^{k \times M_N}$, we have $\text{rank}(\mathbf{P}) = M_N < k$. As a result, $\mathbf{P}$ has a left inverse matrix $\mathbf{P}_{left}^{-1}$ where $\mathbf{P}_{left}^{-1} \mathbf{P} = \mathbf{I}_{M_N}$. Hence when $[\mathbf{x} \quad \mathbf{x}^s]$ satisfies Equation (9), i.e., $\mathbf{P} [\mathbf{A} \quad \mathbf{I}_{M_N}] [\mathbf{x}, \mathbf{x}^s] = \mathbf{Pb}$, we have

$$\mathbf{P}_{left}^{-1} \mathbf{P} [\mathbf{A} \quad \mathbf{I}_{M_N}] [\mathbf{x}, \mathbf{x}^s] = \mathbf{P}_{left}^{-1} \mathbf{Pb},$$

which then transforms into Equation (10). Therefore, Equations (9) and (10) represent the same bandwidth feasible region, which completes the proof. ■

Next, we give the correctness proof for delay and loss rate resource abstraction equivalence.

Proposition 2 (Delay/Loss Rate Resource Abstraction Equivalence): *If the resource abstraction $\mathbf{Ax} \leq \mathbf{b}$, where $\mathbf{A} = [\mathbf{A}_1, \mathbf{A}_2, \dots, \mathbf{A}_N]$ and $\mathbf{b} = [\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_N]$, represents the delay or loss rate property for a set of circuits F over N member networks. The aggregator can compute the aggregated multi-domain delay or loss rate resource abstraction $\oplus \Pi_i(F_i)$, where $\Pi_i(F_i) = \mathbf{A}_i \mathbf{x} \leq \mathbf{b}_i$ using the set of linear equations in Equation (9).*

Proof: From the proof of Proposition 1, we know that $\mathbf{P}$ is full column rank. Observe $[\sum \mathbf{P}_i \mathbf{A}_i \quad \mathbf{P}_1 \dots \mathbf{P}_N]$, the coefficient matrix on the LHS of Equation (10), we can find that each column of $\sum \mathbf{P}_i \mathbf{A}_i$ can be linearly expressed by the columns in $\mathbf{P}$. As a result, we can dissect $[\sum \mathbf{P}_i \mathbf{A}_i \quad \mathbf{P}_1 \dots \mathbf{P}_N]$ into $\mathbf{P} [\mathbf{A} \quad \mathbf{I}_{M_N}]$ through Gaussian Elimination and learn $\mathbf{A}$. Similarly, we can learn $\mathbf{b}$. As such, we can reconstruct $\mathbf{Ax} \leq \mathbf{b}$ and then compute the aggregated multi-domain delay or loss rate resource abstraction $\oplus \Pi_i(F_i)$. ■

2) *Security*: Next, we give the following proposition on the privacy-preserving property of the proposed protocol.

Proposition 3 (Resource Abstraction Privacy-Preserving): *In the semi-honest security model, the proposed resource abstraction obfuscating protocol ensures that (1) the aggregator cannot associate any linear equation it receives in $\Pi_p(F)$ with any particular member network, and (2) for any $\mathbb{M}_i$, it does not know any linear inequality from any other $\mathbb{M}_j(F_j)$ ($j \neq i$).*

Proof: From the description of the resource abstraction obfuscation proof, we see that each $\mathbb{M}_i$ directly sends its own set of disguised linear equations back to the aggregator, hence it does not know any linear inequality from any other member network. Furthermore, even though Proposition 2 shows that it is possible for the aggregator to compute $\mathbf{Ax} \leq \mathbf{b}$, the aggregator cannot associate any $\mathbf{A}_i$ or $\mathbf{b}_i$ to any particular $\mathbb{M}_i$ because $\mathbf{P}_i$ is also disguised by matrices $\mathbf{C}_i$, $\mathbf{C}_{i-1}$, $\mathbf{C}_{i+1}$, $\mathbf{D}_i$, $\mathbf{D}_{i-1}$ and $\mathbf{D}_{i+1}$ before sending back to the aggregator. ■

Even with Proposition 2 and the inter-member-network-path information of each circuit, the aggregator still cannot associate any linear inequality in $\mathbf{Ax} \leq \mathbf{b}$ with the corresponding member network or any networking device (i.e., switch or link). This is because (1) the set of linear equations sent by each member network do not represent its original feasible region, and (2) the inter-member-network-path does not reveal any topology information inside member networks.

With Propositions 1, 2, and 3, we can get the following theorem.

Theorem 1: *Given a set of circuits F that span over N member networks, the proposed resource abstraction obfuscating protocol ensures that the aggregator receives equivalent, privacy-preserving resource abstraction and each member network only knows its own resource feasible region.*

As stated in Section IV-A, the resource abstraction obfuscating protocol was designed for the semi-honest security model. Next, we analyze the privacy-preserveness of our obfuscating protocol in a collusion security model, a more adversarial model in which some member networks may share their resource abstractions with the aggregator. Specifically, we prove the following proposition.

Proposition 4 (Resource Abstraction Privacy-Preserving against Collusion): *Assume some member networks may collude with the aggregator to share their resource abstractions $\Pi_j(F_j)$. Given a non-colluding member network $\mathbb{M}_i$, the resource abstraction obfuscating protocol ensures that the aggregator cannot associate any linear equation it receives in $\Pi_p(F)$ with $\mathbb{M}_i$ unless all other $N - 1$ member networks choose to share their resource abstractions $\Pi_j(F_j)$, where $j \neq i$, with the aggregator.*

Proof: The proof of this proposition follows the proof of Proposition 2 and Proposition 3. Essentially, in order to associate any linear equation in $\Pi_p(F)$ to a non-colluding member network $\mathbb{M}_i$, the aggregator needs to know that this linear equation does not belong to any other member network. Given that the colluding member networks only share their own resource abstractions with the aggregator, this can only be achieved for all other $N - 1$ networks to share their own resource abstractions $\Pi_j(F_j)$, where $j \neq i$, with the aggregator. ■

This proposition indicates that the obfuscating protocol can preserve the privacy of a member network against the collusion between the aggregator and up to $N - 2$ member networks.

3) *Efficiency:* We next analyze the efficiency of our protocol at different phases. During the initialization phase, the main overhead comes from the process each member network agreeing on k , and each $\mathbb{M}_i$ share $\mathbf{C}_i$ and $\mathbf{D}_i$ with $\mathbb{M}_{i+1}$. The first part can be efficiently realized using leader-election algorithms in ring topology or pre-configured. For the second part, it can be efficiently realized by sharing random seeds between $\mathbb{M}_i$ and $\mathbb{M}_{i+1}$. In the obfuscating phase, the computation overhead is also low because it only involves simple, cheap matrix operations, *e.g.*, addition and multiplication.

One may have concern on the transmission overhead of our protocol in the transmission phase because we disguise the set of linear inequalities of each member network into a larger set of linear equations. As such, we quantify the transmission overhead of our obfuscating protocol as follows:

Proposition 5 (Transmission Overhead): *Given a resource discovery procedure for a set of circuits F spanning over N member networks, the transmission overhead of the resource abstraction obfuscating protocol at each member network is $O(k|F|)$, where $k > \sum m_i$.*

Proof: Observing the set of equations sent by each $\mathbb{M}_i$ in Equation (8), we can see that most of the columns of the

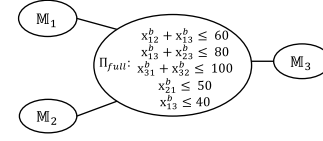


Fig. 9. An illustrating example of super-set projection.

LHS coefficient matrix are zero-columns. Therefore, each $\mathbb{M}_i$ only needs to send nonzero-columns to the aggregator and specifies the indice of these columns. As such, the number of elements to be sent to the aggregator is bounded by $O(k|F|)$, where $k > \sum m_i$. This substantially reduces the amount of data needs to be transmitted from $\mathbb{M}_i$ to the aggregator. ■

V. SUPER-SET RESOURCE ABSTRACTION PROJECTION

As pointed out in Section I, the third challenge for resource discovery is efficiency and scalability, as the number of resource discovery requests may be large in collaboration networks and each request could trigger a resource discovery procedure. This procedure requires the communication between the aggregator and the user, and between the aggregator and every Mercator domain server in member networks. Furthermore, the introduction of resource abstraction obfuscating in Section IV may also increase the communication and computation overhead of resource discovery. To address the efficiency and scalability issue, we develop a novel super-set projection technique, which does not require any change to the resource abstraction design in Section III or the extended obfuscating protocol in Section IV. In this section, we first describe its basic idea, and then give the details of this mechanism.

A. Basic Idea

The intuition of super-set projection is simple: to have the aggregator proactively discover the resource abstraction for a set of circuits between every pair of source and destination member networks, and use these pre-computed abstractions to quickly project to get the resource abstraction for user's requests.

In particular, in a collaboration network of N member networks, the super-set projection technique first simulates the need of $N(N - 1)$ artificial circuits, where each circuit f_{ij} represents an artificial circuit from $\mathbb{M}_i$ to $\mathbb{M}_j$. With this artificial resource discovery request, the aggregator follows the normal resource discovery process to discover the shared bandwidth of all these $N(N - 1)$ circuits across the whole collaboration network, represented by Π_{full} . When a user sends an actual resource discovery request for a set of F circuits, the aggregator checks the source and destination member networks of each circuit, and uses the stored Π_{full} to derive $\Pi(F)$ by removing unrelated inequalities and unrelated artificial circuits, instead of starting a new resource discovery procedure. In this way, the overhead of resource discovery is reduced to a single round of message exchange between the aggregator and the user.

Example: Consider an example of 3 member networks in Fig. 9. With the super-set projection, the aggregator discovers

the bandwidth sharing of all $3 \times 2 = 6$ network-to-network artificial circuits as Π_{full} in the figure. When a user submits a resource discovery request for two circuits (S_1, D_1) and (S_2, D_2) , where S_1 is in $\mathbb{M}_1$, S_2 and D_1 are in $\mathbb{M}_2$ and D_2 is in $\mathbb{M}_3$. The aggregator first maps the (S_1, D_1) to the artificial circuit from $\mathbb{M}_1$ to $\mathbb{M}_2$, and (S_2, D_2) to the artificial circuit from $\mathbb{M}_2$ to $\mathbb{M}_3$. Next, it projects Π_{full} to these two circuits to get the resource abstraction for these two circuits by (1) removing all linear inequalities that do not contain x_{12}^b or x_{23}^b , and (2) for every remaining linear inequality, remove all the items on the LHS that are not x_{12}^b or x_{23}^b . Finally, it returns the resource abstraction: $\{x_{12}^b \leq 60, x_{23}^b \leq 80\}$, to the user.

B. Update of Π_{full}

We ensure the freshness of Π_{full} via two mechanisms. First, the Mercator domain servers at member networks periodically send updated information to the aggregator. Second, when the reservation system receives and successfully executes a resource reservation request from the user, it sends a notification to the aggregator with the reservation details so that the aggregator can update Π_{full} . The aggregator will only query the Mercator domain servers to obtain an up-to-date abstraction for the user when the user fails to reserve the resource based on the projected abstraction.

C. Handling Heterogeneous Flows

One may notice that the super-set projection technique is designed based on the assumption that given a source-destination member network pair, all the traffic flows between these two member networks will be treated homogeneously by all other member networks. In practice, flows between the same source-destination member network pair may be handled differently by other member networks, *i.e.*, they are heterogeneous flows. To address this limitation, we use traffic classes to differentiate heterogeneous flows. In particular, for each source-destination member network pair with G different traffic classes, the super-set projection technique considers these classes as G separate artificial circuits and proactively discovers the bandwidth sharing among these G circuits and other artificial circuits.

VI. IMPLEMENTATION AND DEPLOYMENT

In this section, we describe the implementation of the Mercator and the recent deployment of Mercator in a small federation network to orchestrate large-scale science dataset transfer between two major cities in the United States.

A. Implementation

Figure 10 shows the Mercator domain server implementation, including the Mercator domain server and the aggregator.

1) *Mercator Domain Server*: We build the Mercator domain server on top of the OpenDaylight Software Defined Network controller [48]. Essentially, the Mercator domain server collects the network state information from the OpenDaylight controller, *e.g.*, topology, policy and traffic statistics, processes the collected information into resource abstraction, and sends the abstraction back to the aggregator.

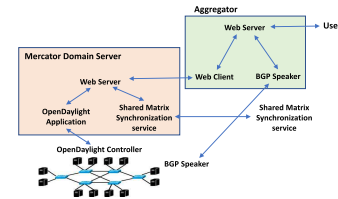


Fig. 10. The Mercator implementation.

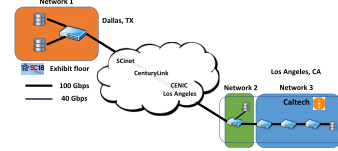


Fig. 11. Deployment of Mercator on a small federation network at Dallas, Texas and Los Angeles, California.

The Mercator domain server has three modules: an OpenDaylight application running in a Karaf container, and a web server accepting the resource discovery from the aggregator and responding with the resource abstraction, and a synchronization service communicating with neighbor domain servers to exchange shared-random matrices that are used for abstraction obfuscating.

2) *Aggregator*: The aggregator has three modules: a web server, a web client and a BGP speaker. The web server provides interfaces for the user to submit a resource discovery request for a set of circuits in a format specified by the ALTO protocol [24]. The web client communicates with Mercator domain servers in different member networks by sending resource discovery requests. In addition, the BGP speaker maintains BGP sessions with the border routers or route servers at member networks to collect inter-member-network paths information.

B. Deployment

We deploy Mercator in a small federation network shown in Figure 11. Specifically, this federation is composed of three member networks. Network 1 is in Dallas, Texas, and Network 2 and network 3 are in Los Angeles, California. Network 1 is connected to network 2 through a layer-2 WAN circuit with a 100 Gbps bandwidth, provisioned by several providers such as SCinet, CenturyLink and CENIC. Network 1 is a temporal science network in the CMS experiment [49], while network 2 and 3 are long-running CMS Tier-2 sites. In this federation, users need to reserve network resources to transfer large-scale science datasets (*e.g.*, with a size of hundreds of PB) between networks.

In our deployment, a Mercator domain server is deployed in each network, and the aggregator is deployed in Dallas. We also deploy SFP, a BGP-compatible routing protocol providing fine-grained routing information [50] in the federation. Upon receiving a user's request, *e.g.*, to discover network resources for circuits from network 1 to network 2 and 3, the aggregator in Mercator contacts the SFP speaker at different networks to discover the interdomain routes for

these circuits, and then sends resource discovery requests to the Mercator domain servers at different networks. After collecting the resource abstraction from Mercator domain servers, the aggregator assembles them and returns to the user. The user then uses such information to compute the optimal amount of network resources to reserve for each circuit and send to the underlying reservation to reserve the resources.

1) *Performance*: We evaluate the accuracy and latency of Mercator for discovering network resources in this network. During our evaluation, Mercator accurately discovers the network resource information for a large amount of circuits reservation requests with a very low discovery latency. Specifically, for all the reservation requests, Mercator always provides the accurate information of available bandwidth sharing in the network (*i.e.*, a 100% accuracy), with an average discovery latency of ~ 100 milliseconds, and a worst latency of less than 1 second. With the discovered network resource information, users can transmit large-scale science datasets at a speed up to 100 Gbps, (*i.e.*, the theoretical maximal throughput). A demonstration of the Mercator deployment in this federation network can be found at [51].

VII. EVALUATION

We implement Mercator on commodity servers (*i.e.*, equipped with Intel(R) Xeon(R) E5-2609 2.50GHz 4-core CPU and 32 GB memory) and evaluate its performance based on a member-network-level topology from a large federation of networks supporting large-scale distributed science collaborations, and using real traffic traces from recent science experiments. After describing our experimental setup, we first demonstrate the benefits of resource abstraction through algebraic-expression enumeration. Second, we demonstrate the efficiency of the proposed resource abstraction obfuscation protocol. Finally, we demonstrate that the super-set projection technique substantially increases the scalability of Mercator.

A. Experimental Setup

We evaluate Mercator on the member-network-level topology from LHC Open Network Environment (LHCONE), a global science network consisting of 62 member networks, where scientists conduct large-scale distributed analytics. Because inter-member-network routing typically is not based on shortest path routing, but follows business relationships (*e.g.*, customer, peer, provider), we label the connections between every pair of connected member networks with their business relationship using the CAIDA network relationships dataset [52], and we compute the inter-member-network paths according to conventional policies for selecting and exporting routes. For member networks' intradomain topologies, we randomly select a topology for each network from the Topology Zoo [53], which provides a collection of real intradomain topologies. The topology of transit member networks varies from 31 switches/routers with 33 links to 49 switches/routers with 85 links. The topology of stub member networks (*e.g.*, campus science networks) ranges from 7 switches/routers with 6 links to 21 switches/routers with 44 links.

B. Benefits of Resource Abstraction Through Algebraic-Expression Enumeration

The first set of experiments demonstrate the benefits of the resource abstraction through algebraic-expression enumeration. We show that this abstraction reduces the time to discover network resources by up to six orders of magnitude, and allows fairer allocations of network resources.

1) *Methodology*: To evaluate the benefits of this resource abstraction, we replay the trace from a large-scale distributed experiment, and submit network resource reservations for the corresponding flows. More specifically, we use the actual trace from the CMS experiment [54], a major scientific experiment in LHC, and a main source of traffic in LHCONE. We extract the traffic flows, with their source member network, destination member network and the time. We focus on the 7-day trace starting from September 30, 2018 to October 6, 2018, and slice the data trace into 24 continuous 2-hour time windows. We apply the resources reservation once every time window. In other words, resources for traffic flows starting at the same time window are reserved in the same request, and we assume all resources will be released in the next time window.

We compare the performance of Mercator with that of existing reservation systems. In particular, for existing systems, we consider one that adopts a probe-requests based approach:

- **Mercator**: As described in Section II, for every resource discovery request, the aggregator queries the relevant member networks for their resource abstraction, and then derives the feasible bandwidth allocation region.

- **Probe requests**: As described in Section I, existing resource reservation systems such as OSCARS process each circuit in the request one at a time and in a sequential order. For each circuit, the resource reservation system initiates a depth-first search to probe if each member network can provide the requested bandwidth. We set the initial requested bandwidth for a circuit as C/N where C is the source host's capacity, and N is the number of flows from that host. In the event of a failure, the resource reservation system performs a binary search of the available bandwidth repeatedly halving the requested bandwidth until success. The process is repeated for each circuit in the request.

2) *Results*: First, we consider that the goal of the resource allocation policy is to maximize the minimum throughput of all the requested flows (max-min fairness). Such a policy is commonly desired as it ensures high throughput and fairness across the circuits. We compare the fairness of the network resource allocations obtained with Mercator to that obtained with the probe-requests based solution. We adopt Jain's fairness index [55] to measure the fairness [56]:

$$J(x_1, x_2, \dots, x_n) = \frac{(\sum_{i=1}^n x_i)^2}{n \cdot \sum_{i=1}^n x_i^2}$$

where x_i is the ratio of the actual allocation and the optimal fair allocation for a single flow. 12a shows that with resource abstraction, Mercator can always compute the optimal max-min fairness allocation. Hence its fairness index is always 1. In contrast, the highest fairness index the probe-requests can get is 0.05, with most of the slots even smaller than 0.01.

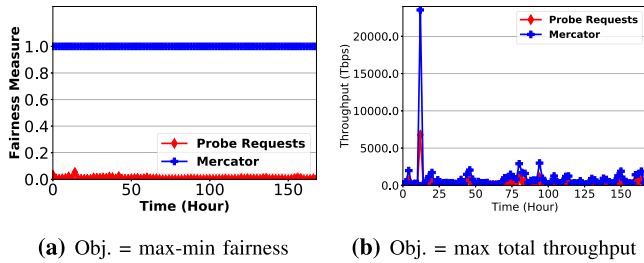


Fig. 12. Comparison of performance between the probe-requests approach and Mercator in different objectives.

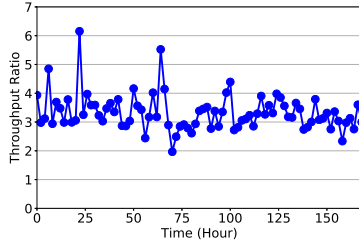


Fig. 13. Ratio of throughput between Mercator and the probe-requests approach when the objective is to maximize the total throughput.

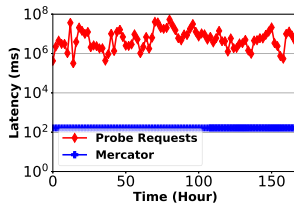


Fig. 14. Resource discovery latency of the probe-requests approach and Mercator.

Second, we consider the case where the objective is to maximize the total throughput. 12b shows that the total throughput of Mercator is larger than that obtained by the probe-requests based solution throughout the whole experiment. Fig. 13 shows that the ratio of throughput of Mercator over that of probe-requests based solution is by 3.47x on average, and up to 6.2x. The results are noteworthy given that Mercator assumes the routes for each circuit to be completely determined by the underlying intradomain routing protocol. In contrast, the probe-requests approach sequentially explores every possible route for each circuit until it finds an available one. In other words, even with much less exploration, Mercator still outperforms the probe requests significantly. Allowing Mercator to consider not only the routes provided by the underlying routing protocols, but also all other available routes, could lead to significant additional improvements. We leave the extension of Mercator to consider all possible routes in the network as future work.

Fig. 14 presents the total resource discovery latency for completing all circuits resource reservations in a time window. We assume the aggregator to be in New York, and consider network latencies as measured in [57]. The figure shows the total resource discovery latency with Mercator can reduce the time to discover network resources by four orders of magnitude on average and up to six orders of magnitude at times.

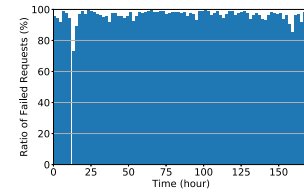


Fig. 15. Ratio of failed requests in the probe-requests approach.

This is because resource abstraction allows users to query the information from different member networks in parallel. In contrast, existing probe-requests based solutions process requests sequentially, and continuously probe to discover the available network resources.

Finally, we highlight that the probe-requests based solution suffers high request failure ratio, *i.e.*, a large number of requests cannot succeed: We define a failure of a request as the inability to reserve resource for the circuit, due to the lack of remaining capacity despite the gradually decreasing requested bandwidth. Fig. 15 shows that during the 7-day period Mercator is running, the probe-requests based solution has an average request failure ratio of 87%. In other words, more than 80% of the circuits cannot reserve network resources. This is because the probe-requests approach processes the request for each circuit sequentially. Therefore, the first few circuits may successfully reserve network resources and saturate the network. As such, the majority of the latter requests may fail as the links do not have any spare resources. In contrast, the request failure ratio of Mercator is null because Mercator returns a feasible region for the set of circuits so that the user can make optimal reservation decisions for all circuits.

C. Efficiency of Resource Abstraction Obfuscating Protocol

This second set of experiments evaluate the performance of the resource abstraction obfuscating protocol. We show that this protocol efficiently scales for collaboration networks of 200 member networks, with a maximal overall latency around 3 seconds and an average data transmission overhead between the aggregator and member networks of only around 180 KB.

1) *Methodology*: We conduct our experiment by using the member-network-level topology from the LHC Open Network Environment (LHCONE). In each round of the experiment, we randomly select a set of member networks from the topology. For each chosen member network, we randomly select a set of m linear inequalities, where m is randomly chosen between 5 and 15, to represent the bandwidth feasible reason for 10 circuits in this member network. For the encryption and decryption operations in the obfuscating protocol, we use the AES algorithm, provided by the Python Cryptography Toolkit (pycrypto) [58]. The parameters k , C_i and D_i are pre-configured as discussed in Section IV-C.

We consider two metrics, *i.e.*, the latency and the data transmission overhead of the resource abstraction obfuscating protocol. First, the overall latency of the protocol is measured from the beginning of the obfuscation phase, when each member network independently starts to obfuscate its own set of linear inequalities, to the end of the transmission process,

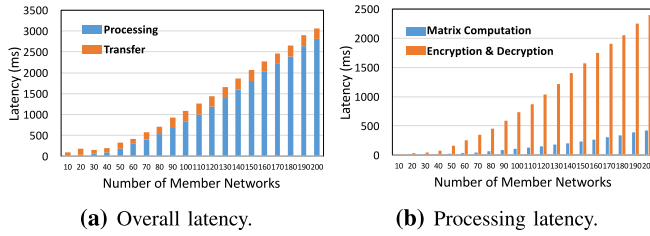


Fig. 16. The latency of the resource abstraction obfuscating protocol.

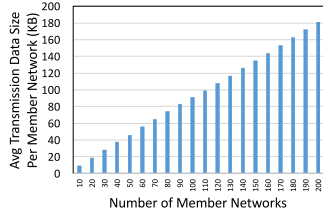


Fig. 17. The data transmission overhead of the resource abstraction obfuscating protocol.

when the aggregator obtains $\sum \mathbf{P}_i \mathbf{A}_i \mathbf{x} = \mathbf{b}$. We use the field statistic results measured in [57] as the communication latencies between the aggregator and the Mercator domain servers at the different member networks. Second, the data transmission overhead is measured as the size of the set of encrypted, obfuscated linear equations transferred from each member network to the aggregator. We vary the number of member networks from 10 to 200, in a step size of 10. For each number of member networks, we repeat the experiment 10 times and measure the average values of these metrics.

2) *Results*: We present the results of our experiments in Fig. 16 and Fig. 17. In particular, Fig. 16a shows the overall latency of the obfuscating protocol under different numbers of member networks, together with a break down on processing delay and transmission latency. We observe that even for a large collaboration network with 200 member networks, which is larger than most existing operational collaboration networks, the overall latency of the resource abstraction obfuscating protocol is only slightly over 3 seconds, which demonstrates that the latency of this protocol is reasonably low. We also observe that the processing latency takes a much higher percentage than the transmission latency and that the processing latency has a linear growth as the number of member networks increases. We further plot the breakdown of the processing latency. Fig. 16b shows that both the cryptography operations of AES and the matrix operations in the resource abstraction obfuscating protocol increases linearly as the number of member networks increases, but the AES encryption and decryption operations are the most expensive operations in the protocol (*i.e.*, up to 2.4 seconds for federations of 200 member networks). More importantly, although the obfuscating protocol may take over 3 seconds for a federation of 200 member networks, we emphasize that with the super-set projection technique, the Mercator domain servers do not need to execute the obfuscating protocol for each individual request.

Next, we present the average data transmission overhead of the obfuscating protocol at each member network in Fig. 17.

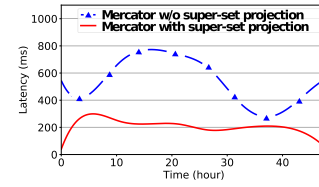


Fig. 18. Comparison of latency between Mercator with and without super-set projection.

We see in this figure that even after the encryption, the size of data to be transmitted from member networks to the aggregator is still very small. For example, for a collaboration network with 200 member networks, the average size of data transmitted from a member network to the aggregator is only 180 KB. As discussed in Section IV-C, this is because most of the columns of the LHS coefficient matrix are zero-columns and each member network only needs to send nonzero-columns to the aggregator. The linear scaling of the data transmission overhead (*i.e.*, the ciphertext) at each member network comes from the linear increase of the number of disguised linear equations (*i.e.*, the plaintext), which is caused by the linear increase of k due to the increased number of member networks. This is consistent with Proposition 5 in Section IV-C.

D. Efficiency of Super-Set Projection

In this experiment, we evaluate the efficiency of the super-set projection technique in improving the scalability of Mercator. We show that this mechanism improves the resource discovery delay of Mercator by 2 times, and that its update latency is within seconds in a collaborative network with 200 member networks.

1) *Methodology*: We conduct our experiments by using the same settings as in Section VII-B.1. We focus on two metrics. The first one is the resource discovery latency. When Mercator uses super-set projection, the resource discovery latency is reduced to only the round-trip time from the user to the Mercator aggregator because the aggregator can derive the resource abstraction for a request from the precomputed Π_{full} .

To have a comprehensive understanding on the scalability of super-set projection, we are also interested in a second metric, the update latency. This is measured as the resource discovery latency of from the time the aggregator starting the artificial resource abstraction discovery procedure to the time the aggregator receives the latest Π_{full} . In particular, we measure this latency under different collaboration scales by varying the number of member networks and the number of stub member networks in the collaborative network. For each setting, we repeat the experiment 10 times and compute the average update latency. In each repetition, we also randomly choose different sizes of intradomain topologies from the Topology Zoo dataset for each member network.

2) *Results*: Fig. 18 compares the resource discovery latency of Mercator with and without super-set projection for a 48-hour period in the LHCONE trace. The results for the whole 7-day period is similar, and hence is omitted.

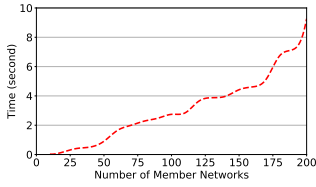


Fig. 19. Update latency of super-set projection.

We observe that the super-set projection technique decreases the average resource discovery latency by around 2 times. Fig. 19 presents the update latency of this mechanism. It shows that even in a collaborative network with 200 member networks, the update latency of Π_{full} is still less than 10 seconds. Most importantly, although computing Π_{full} may take up to ten seconds for a federation of 200 member networks, we emphasize that resource discovery requests do not get blocked at the aggregator because servers from the aggregator pool can still process incoming requests using the previously computed resource abstraction, which is continuously locally updated (*e.g.*, available resources are continuously reduced as incoming requests reserve resources).

VIII. RELATED WORK

A. Resource Reservation

Network resource reservation systems are deployed driven by the demand and substantial benefits of providing predictable network resources [5]–[11], [13]–[15]. Systems running in a single administrative domain (*e.g.*, NetStitcher [13], SWAN [14] and B4 [15]) are often provided with detailed network information, such as the topology and links' availability. Therefore, optimizing resource reservations in a single administrative domain can be very efficient. In contrast, in a multi-domain network (*e.g.*, LHC), due to networks' concern of revealing sensitive information, resource reservation systems only allow users to submit requests for reserving a specific amount of resources (*e.g.*, a circuit providing a certain amount of bandwidth and delay), and return either success or failure [5]–[11]. Without an interface to provide network resource information, optimizing resource reservations in a multi-domain network requires a complex, time-consuming trial-and-error process.

B. Resource Discovery

Multiple multi-domain resource discovery systems (*e.g.*, [16]–[20]) are designed to discover endpoint resources (*i.e.*, computation and storage resources) and their availability for different services across multiple domains. In contrast, there has been little progress on multi-domain network resource discovery systems that provide fine-grained, global network resource information, to support high-performance, collaborative data sciences.

Many cluster/grid resource management systems [38], [39], [59]–[64] adopt a graph-based abstraction to discover and manage network resources. This abstraction is designed for single administrative domains (*e.g.*, a company or a university) to manage their own network, where they do not need to preserve the privacy of network. If this abstraction is directly

ported to a multi-domain collaborative network, it would expose the private information (*e.g.*, the network topology) of member networks, leading to security breaches.

Some systems in cloud computing [26]–[28] adopt a network-does-all approach, in which users are provided with a more expressive interface for specifying requirements on data transfers and the network orchestrates resources between different user requests. Though this approach protects the privacy of the network, the network can only provide elastic resource reservation for user requests (*i.e.*, some requests may be preempted or rejected). Some recent systems (*e.g.*, the ALTO protocol [24], [25], [29] and the SENSE project [30], [31]) provide users the information of certain properties of network resources using the one-big-switch abstraction. While this approach protect the privacy of network, it cannot provide accurate information of network resource sharing between flows (*e.g.*, bandwidth), which is critical for optimizing the emerging use cases (*e.g.*, large-scale collaborative sciences).

Some recent studies [37], [41], [42], [65], [66] propose variations of the one-big-switch abstraction to represent the resource availability and sharing among different data traffic flows using operations defined on different algebra fields. However, this abstraction (1) cannot handle complex routing and traffic engineering policies, *e.g.*, WCMP, and (2) will raise security concern when applied to multi-domain science collaborations. In contrast, Mercator provides fine-grained, global network resource information, to support high-performance, collaborative data sciences, through a unifying representation and composition framework to reveal compact, complete multi-domain network resource information.

IX. CONCLUSION

Existing multi-domain network resource reservation systems often operate on coarse-grained or localized information, resulting in substantial inefficiencies. To address this issue, We present Mercator, a novel multi-domain network resource discovery system to provide fine-grained, global network resource information, to support high-performance, collaborative data sciences. The core of Mercator is a unifying representation resource abstraction using algebraic expressions to represent multi-domain network resources. We develop a resource abstraction obfuscating protocol and a super-set projection technique to ensure the privacy-preserving and the scalability of Mercator. Evaluation using real data shows that Mercator discovers fine-grained network resources by up to six orders of magnitude, allows fairer allocations of network resources, and scales to a collaborative network of 200 member networks.

ACKNOWLEDGMENT

The authors thank Haizhou Du, Kai Gao, Linghe Kong, Geng Li, Yeon-sup Lim, Alan Liu, Ennan Zhai, and Yan Zhu for their help during the preparation of this paper.

REFERENCES

- [1] Q. Xiang *et al.*, "Fine-grained, multi-domain network resource abstraction as a fundamental primitive to enable high-performance, collaborative data sciences," in *Proc. Int. Conf. High Perform. Comput., Netw., Storage Anal.*, 2018, pp. 58–70.

- [2] *The Large Hadron Collider (LHC) Experiment*. Accessed: Dec. 2018. [Online]. Available: <https://home.cern/topics/large-hadron-collider>
- [3] *The Square Kilometre Array*. Accessed: Dec. 2018. [Online]. Available: <https://www.skatelescope.org/>
- [4] *The Linac Coherent Light Source*. Accessed: Dec. 2018. [Online]. Available: <https://lcls.slac.stanford.edu/>
- [5] *Oscars: On-Demand Secure Circuits and Advance Reservation System*. Accessed: Dec. 2018. [Online]. Available: <https://www.es.net/engineering-services/oscars/>
- [6] M. Campanella *et al.*, "Bandwidth on demand services for European research and education networks," in *Proc. IEEE 1st Int. Workshop Bandwidth Demand*, Nov. 2006, pp. 65–72.
- [7] C. Guok and D. Robertson, "ESnet on-demand secure circuits and advance reservation system (OSCARs)," *Internet2 Joint*, 2006, vol. 92.
- [8] W. Johnston, C. Guok, and E. Chaniotakis, "Motivation, design, deployment and evolution of a guaranteed bandwidth network service," in *Proc. TERENA Netw. Conf.*, 2011, pp. 1–14.
- [9] B. Riddle, "BRUW: A bandwidth reservation system to support end-user work," in *Proc. TERENA Netw. Conf.*, Poznan, Poland, 2005.
- [10] J. Sobieski, T. Lehman, and B. Jabbari, "DRAGON: Dynamic resource allocation via GMPLS optical networks," in *Proc. MCNC Opt. Control Planes Workshop*, Chicago, IL, USA, 2004.
- [11] X. Zheng, M. Veeraraghavan, N. S. V. Rao, Q. Wu, and M. Zhu, "CHEETAH: Circuit-switched high-speed end-to-end transport architecture tested," *IEEE Commun. Mag.*, vol. 43, no. 8, pp. S11–S17, Aug. 2005.
- [12] Q. Xiang, H. Yu, J. Aspnes, F. Le, L. Kong, and Y. R. Yang, "Optimizing in the dark: Learning an optimal solution through a simple interface," in *Proc. AAAI*, Nov. 2018, pp. 1–8.
- [13] N. Laoutaris, M. Sirivianos, X. Yang, and P. Rodriguez, "Inter-datacenter bulk transfers with netstitcher," *ACM SIGCOMM Comput. Commun. Rev.*, vol. 41, no. 4, pp. 74–85, 2011.
- [14] C.-Y. Hong *et al.*, "Achieving high utilization with software-driven WAN," *ACM SIGCOMM Comput. Commun. Rev.*, vol. 43, no. 4, pp. 15–26, 2013.
- [15] S. Jain *et al.*, "B4: Experience with a globally-deployed software defined WAN," *ACM SIGCOMM Comput. Commun. Rev.*, vol. 43, no. 4, pp. 3–14, 2013.
- [16] Y. Deng, F. Wang, and A. Ciura, "Ant colony optimization inspired resource discovery in P2P grid systems," *J. Supercomput.*, vol. 49, no. 1, pp. 4–21, 2009.
- [17] A. Iamnitchi and I. Foster, "A peer-to-peer approach to resource location in grid environments," in *Grid Resource Management*. Springer, 2004, pp. 413–429.
- [18] T. Kocak and D. Lacks, "Design and analysis of a distributed grid resource discovery protocol," *Cluster Comput.*, vol. 15, no. 1, pp. 37–52, 2012.
- [19] I. Sfiligoi, D. C. Bradley, B. Holzman, P. Mhashilkar, S. Padhi, and F. Wurthwein, "The pilot way to grid resources using glideinWMS," in *Proc. IEEE CSIE*, Mar./Apr. 2009, pp. 428–432.
- [20] D. Thain, T. Tannenbaum, and M. Livny, "Distributed computing in practice: The condor experience," *Concurrency Pract. Exper.*, vol. 17, nos. 2–4, pp. 323–356, 2005.
- [21] R. Ahmed, N. Limam, J. Xiao, Y. Iraqi, and R. Boutaba, "Resource and service discovery in large-scale multi-domain networks," *IEEE Commun. Surveys Tuts.*, vol. 9, no. 4, pp. 2–30, 4th Quart., 2007.
- [22] A. Hameurlain, D. Cokuslu, and K. Erciyes, "Resource discovery in grid systems: A survey," *Int. J. Metadata, Semantics Ontologies*, vol. 5, no. 3, pp. 251–263, 2010.
- [23] N. J. Navimipour, A. M. Rahmani, A. H. Navin, and M. Hosseinzadeh, "Resource discovery mechanisms in grid systems: A survey," *J. Netw. Comput. Appl.*, vol. 41, pp. 389–410, May 2014.
- [24] R. Alimi, Y. Yang, and R. Penno, *Application-Layer Traffic Optimization (ALTO) Protocol*, document RFC 7285, 2014.
- [25] H. Xie, Y. R. Yang, A. Krishnamurthy, Y. G. Liu, and A. Silberschatz, "P4P: Provider portal for applications," *ACM SIGCOMM*, vol. 38, no. 4, pp. 351–362, Aug. 2008.
- [26] J. Lee *et al.*, "Application-driven bandwidth guarantees in datacenters," *ACM SIGCOMM Comput. Commun. Rev.*, vol. 44, no. 4, pp. 467–478, 2014.
- [27] H. Zhang *et al.*, "Guaranteeing deadlines for inter-data center transfers," *IEEE/ACM Trans. Netw.*, vol. 25, no. 1, pp. 579–595, Feb. 2017.
- [28] V. Jalaparti, I. Bliznets, S. Kandula, B. Lucier, and I. Menache, "Dynamic pricing and traffic engineering for timely inter-datacenter transfers," in *Proc. ACM SIGCOMM Conf.*, 2016, pp. 73–86.
- [29] D. Perez and C. Rothenberg, "ALTO-based broker-assisted multi-domain orchestration," in *Proc. IETF Draft*, Mar. 2019.
- [30] H. Newman *et al.*, "Next-generation exascale network integrated architecture for global science," *J. Opt. Commun. Netw.*, vol. 9, no. 2, pp. A162–A169, 2017.
- [31] I. Monga *et al.*, "SDN for end-to-end networked science at the exascale (SENSE)," in *Proc. IEEE/ACM INDIS*, Nov. 2018, pp. 33–44.
- [32] V. Welch, M. Thompson, D. E. Engert, S. Tuecke, and L. Pearlman, *Internet X.509 Public Key Infrastructure (PKI) Proxy Certificate Profile*, document RFC 3820, RFC Editor, Jun. 2004, p. 37. [Online]. Available: <https://rfc-editor.org/rfc/rfc3820.txt>. doi: 10.17487/RFC3820.
- [33] N. Sakimura, J. Bradley, M. Jones, B. de Medeiros, and C. Mortimore, "Open ID connect core 1.0," OpenID, San Ramon, CA, USA, Tech. Rep., Nov. 2014.
- [34] OSST Committee. (2012). *Security Assertion Markup Language (SAML) 2.0*. [Online]. Available: http://www.oasis-open.org/committees/tc_home.php
- [35] Y. Rekhter, S. Hares, and D. T. Li, *A Border Gateway Protocol 4 (BGP-4)*, document RFC 4271, Jan. 2006. [Online]. Available: <https://rfc-editor.org/rfc/rfc4271.txt>
- [36] *Route Views Project*. Accessed: Dec. 2018. [Online]. Available: <http://www.routeviews.org/routeviews/>
- [37] V. Heorhiadi, M. K. Reiter, and V. Sekar, "Simplifying software-defined network optimization using SOL," in *Proc. NSDI*, 2016, pp. 223–237.
- [38] B. Hindman *et al.*, "Mesos: A platform for fine-grained resource sharing in the data center," in *Proc. NSDI*, 2011, p. 22.
- [39] A. Verma, L. Pedrosa, M. Korupolu, D. Oppenheimer, E. Tune, and J. Wilkes, "Large-scale cluster management at Google with Borg," in *Proc. ACM EuroSys*, 2015, p. 18.
- [40] W. Stallings, "The advanced encryption standard," *Cryptologia*, vol. 26, no. 3, pp. 165–188, Jul. 2002. doi: 10.1080/0161-110291890876.
- [41] K. Gao, C. Gu, Q. Xiang, X. Wang, Y. R. Yang, and J. Bi, "ORSAP: Abstracting routing state on demand," in *Proc. IEEE ICNP*, Nov. 2016, pp. 1–2.
- [42] K. Gao, Q. Xiang, X. Wang, Y. R. Yang, and J. Bi, "NOVA: Towards on-demand equivalent network view abstraction for network optimization," in *Proc. ACM/IEEE IWQoS*, Jun. 2017, pp. 1–10.
- [43] J. Telgen, "Identifying redundant constraints and implicit equalities in systems of linear constraints," *Manage. Sci.*, vol. 29, no. 10, pp. 1209–1222, 2002.
- [44] M. Raykova, *Secure Computation in Heterogeneous Environments: How to Bring Multiparty Computation Closer to Practice?*. Columbia Univ., 2012.
- [45] A. C.-C. Yao, "How to generate and exchange secrets," in *IEEE FOCS 1986*.
- [46] X. Feng and Z. Zhang, "The rank of a random matrix," *Appl. Math. Comput.*, vol. 185, no. 1, pp. 689–694, Jan. 2007.
- [47] O. L. Mangasarian, "Privacy-preserving horizontally partitioned linear programs," *Optim. Lett.*, vol. 6, no. 3, pp. 431–436, 2012.
- [48] *The OpenDaylight Project*. Accessed: Dec. 2018. [Online]. Available: <https://www.opendaylight.org>
- [49] The CMS Collaboration *et al.*, "The CMS experiment at the CERN LHC," *J. Instrum.*, vol. 3, no. 8, pp. S08004–S08004, 2008.
- [50] Q. Xiang, C. Guok, F. Le, J. MacAuley, H. Newman, and Y. R. Yang, "SFP: Toward interdomain routing for SDN networks," in *Proc. ACM SIGCOMM Conf. Posters Demos*, 2018, pp. 87–89.
- [51] *A Demonstration of Mercator*. Accessed: Dec. 2018. [Online]. Available: <https://youtu.be/kUK78gHIQDI>
- [52] (2016). *The CAIDA AS Relationships Dataset*. [Online]. Available: <http://www.caida.org/data/as-relationships/>
- [53] S. Knight, H. X. Nguyen, N. Falkner, R. Bowden, and M. Roughan, "The Internet topology zoo," *IEEE J. Sel. Areas Commun.*, vol. 29, no. 9, pp. 1765–1775, Oct. 2011.
- [54] *CMS Task Monitoring*. Accessed: Dec. 2018. [Online]. Available: <http://dashb-cms-job.cern.ch/>
- [55] R. Jain, D.-M. Chiu, and W. R. Hawe, "A quantitative measure fairness discrimination for resource allocation shared computer system," *Eastern Res. Lab., Digit. Equip. Corp.*, Hudson, MA, USA, Tech. Rep., 1984, vol. 38.
- [56] J. Y. Boudec, "Rate adaptation, congestion control and fairness: A tutorial," EPFL, Lausanne, Switzerland, Tech. Rep., Oct. 2000.
- [57] (2018). *Global Ping Statistics—WonderNetwork*. [Online]. Available: <https://wondernetwork.com/pings/>
- [58] *Python Cryptography Toolkit*. Accessed: Dec. 2018. [Online]. Available: <https://pypi.python.org/pypi/pycrypt>

- [59] *Under the Hood: Scheduling MapReduce Jobs More Efficiently With Corona*. Accessed: May 9, 2017. [Online]. Available: <http://on.fb.me/TxUsYN>
- [60] E. Boutin *et al.*, "Apollo: Scalable and coordinated scheduling for cloud-scale computing," in *Proc. OSDI*, 2014, pp. 285–300.
- [61] M. Isard, V. Prabhakaran, J. Currey, U. Wieder, K. Talwar, and A. Goldberg, "Quincy: Fair scheduling for distributed computing clusters," in *Proc. IEEE Int. Conf. Recent Trends Inf. Syst.*, 2009, pp. 261–276.
- [62] Q. Pu *et al.*, "Low latency geo-distributed data analytics," in *Proc. ACM SIGCOMM*, 2015, pp. 421–434.
- [63] R. Viswanathan, G. Ananthanarayanan, and A. Akella, "CLARINET: WAN-aware optimization for analytics queries," in *Proc. Usenix Conf. Operating Syst. Design Implement.*, 2016, pp. 435–450.
- [64] A. Vulimiri, C. Curino, B. Godfrey, K. Karanasos, and G. Varghese, "WANalytics: Analytics for a geo-distributed data-intensive world," in *Proc. CIDR*, 2015.
- [65] Q. Xiang *et al.*, "Unicorn: Unified resource orchestration for multi-domain, geo-distributed data analytics," in *Proc. IEEE SmartWorld, DAIS Workshop*, Aug. 2017, pp. 1–6.
- [66] Q. Xiang, X. Wang, J. Zhang, H. Newman, Y. R. Yang, and Y. J. Liu, "Unicorn: Unified resource orchestration for multi-domain, geo-distributed data analytics," in *Proc. IEEE INDIS Workshop*, 2017.



Qiao Xiang received the bachelor's degree in information security and the bachelor's degree in economics from Nankai University in 2007, and the master's and Ph.D. degrees in computer science from Wayne State University in 2012 and 2014, respectively. From 2016 to 2018, he was a Post-Doctoral Fellow with the Department of Computer Science, Yale University. From 2014 to 2015, he was a Post-Doctoral Fellow with the School of Computer Science, McGill University. He is currently an Associate Research Scientist with the Department

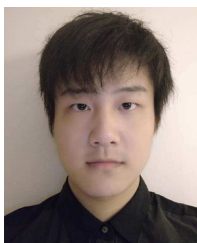
of Computer Science, Yale University. His research interests include software defined networking, resource discovery and orchestration in collaborative data sciences, interdomain routing, and wireless cyber-physical systems.



Jingxuan Jensen Zhang received the bachelor's degree in engineering from the Department of Computer Science and Engineering, Tongji University, in 2015. He is currently Visiting Ph.D. with the Department of Computer Science, Yale University. His research focuses on network resource discovery, abstraction and programming consistency for large-scale data analytics systems. He is also an Active Member of the IETF ALTO Working Group and the OpenDaylight Open Source Community.



Xin Tony Wang received the bachelor's degree in engineering from the Department of Computer Science and Engineering, Tongji University, in 2014. He is currently Visiting Ph.D. with the Department of Computer Science, Yale University. His research interests include software defined networking, interdomain routing, and distributed computing.



Yang Jace Liu received the bachelor's degree in engineering from the Department of Computer Science and Engineering, Tongji University, in 2017. He is currently pursuing the degree with the Computer Science Department, University of Calgary, Canada. His research interests include software defined networking, large-scale data analytics systems, and high-performance computing.



Chin Guok received the B.S. degree in computer science from the University of Pacific in 1991 and the M.S. degree in computer science from The University of Arizona in 1997. He joined ESnet in 1997 as a Network Engineer, where he is focusing primarily on network statistics. He was a Core Engineer in the testing and production deployment of MPLS and QoS (Scavenger Service) within ESnet. He is the Technical Lead of the ESnet On-demand Secure Circuits and Advanced Reservation System (OSCARS) Project which enables end-users to provision guaranteed bandwidth virtual circuits within ESnet. He also serves as a Co-Chair for the Open Grid Forum On-Demand Infrastructure Service Provisioning Working Group. His research interests include high-performance networking and network protocols, dynamic network resource provisioning, network tuning issues, and hybrid network traffic engineering.



Franck Le received the Diplôme d'Ingenieur from the Ecole Nationale Supérieure des Telecommunications de Bretagne in 2000 and the Ph.D. degree from Carnegie Mellon University in 2010. He is currently a Research Staff Member with the IBM Thomas J. Watson Research Center. His current research interests lie at the intersection of Internet of Things, artificial intelligence, and distributed systems & networks.



John MacAuley is currently a Chief Software Architect with Energy Sciences Network. His main research interests include high-speed computer networks and systems, resource discovery, and orchestration in science networks.



Harvey Newman received the Sc.D. degree from MIT in 1974. From 1973 to 1974, he co-lead the team that discovered fourth quark flavor known as charm. He co-lead the MARK J Collaboration that discovered the gluon, the carrier of the strong force in 1979. Since 1982, he has been a Faculty Member with the California Institute of Technology (Caltech), where he is currently the Marvin L. Goldberger Professor of physics. He has been leading a role in originating, developing, and operating state of the art international networks and collaborative systems serving the high energy and nuclear physics communities since 1982. He served on the IETF and the Technical Advisory Group that led to the NSFNet from 1985 to 1986, originated the worldwide LHC Computing Model in 1996, and has been leading the science and network engineering teams defining the state of the art in long distance data transfers since 2002. Since 1994, he has been a member of CMS that discovered the Higgs boson at LHC in 2012.



Y. Richard Yang received the B.E. degree in computer science and technology from Tsinghua University in 1993, and the M.S. and Ph.D. degrees in computer science from The University of Texas at Austin in 1998 and 2001, respectively. He is currently a Professor of computer science and electrical engineering with Yale University. His research is supported by both U.S. government funding agencies and leading industrial corporations, and spans areas including computer networks, mobile computing, wireless networking, and network security. His work

has been implemented/adopted in products/systems of major companies (e.g., AT&T, Alcatel-Lucent, Cisco, Google, Microsoft, and Youku), and featured in mainstream media including Economist, Forbes, Guardian, Chronicle of Higher Education, Information Week, MIT Technology Review, Science Daily, USA Today, Washington Post, and Wired, among others. His awards include a CAREER Award from the National Science Foundation and a Google Faculty Research Award.