

Named Data Networking Strategies for Improving Large Scientific Data Transfers

Susmit Shannigrahi
Colorado State University
susmit@cs.colostate.edu

Chengyu Fan
Colorado State University
chengyu.fan@colostate.edu

Christos Papadopoulos
Colorado State University
christos@colostate.edu

Abstract—Current scientific workflows such as Climate Science and High Energy Particle Physics (HEP), routinely generate and use large volumes of observed or simulated data. Users are often geographically dispersed and need to transfer large volumes of data over the network for replication, archiving, or local analysis. Scientific communities have built sophisticated applications and dedicated networks to facilitate such data transfers, and yet, users continue to experience failures, delay, and unpredictable transfer latency [1].

Named Data Networking (NDN) is a new Internet architecture that provides a more flexible and intelligent network layer, suitable for large data transfers. In this work, we use a real scientific data flow to demonstrate NDN's flexibility and versatility that makes it a suitable choice for large-data workflows. We use deadline-based data transfers as our driving example since HEP communities widely use them [2] and discuss several NDN forwarding strategies that can help such flows. In addition to using typical forwarding strategies, we propose, at a high level, a bandwidth reservation protocol for NDN and an on-demand high-speed path creation mechanism. Using these as building blocks, we create a deadline-based data transfer protocol and show how NDN can simplify and improve scientific data distribution. Finally, we use a week-long HEP data log to evaluate our protocol analytically.

I. INTRODUCTION

Data-intensive science has transformed modern scientific research. Scientists now use observed and simulated data to translate abstract ideas into conclusive findings and concrete solutions. While large datasets benefit modern scientific research immensely, the ever-increasing size of these datasets creates a considerable data management burden. Since the data volume is very high, for example, the Large Hadron Collider (LHC) generates petabytes of data per year, accommodating all computation and storage needs at the data generation site is not feasible. Consequently, many scientific workflows routinely transfer a substantial amount of data for remote storage, replication, or local analysis that range anywhere from tens of gigabytes to terabytes [3]. While available bandwidth in scientific networks is significant, it is still insufficient to handle the aggregate load. Therefore, such transfers must complete before a deadline to free up network resources for subsequent requests. Today this is often accomplished using intelligent but complex applications or manually orchestrated high-speed paths.

Even with these intelligent applications completing transfers within a deadline is challenging. The inherent limitations of

TCP/IP networking [1] can slow down transfers and waste valuable network resources. The inability of the network to use multiple replicas, failure to reuse in-network data and lack of access to in-network state are some of the limitations that make big data retrieval inefficient. Scientific communities have tried several approaches to solve this problem; modified congestion control algorithms, smart applications and bandwidth reservations over dedicated links are a few examples. However, these solutions are complex, often domain specific and lack support from the underlying TCP/IP network. Though previous work has looked at mathematically optimizing deadline-based data transfers [4] [2] [5], these solutions are hard to deploy due to the inflexibility of the TCP/IP network layer.

Named Data Networking (NDN) [6] is a new Internet architecture that directly addresses content instead of end-hosts. It offers several optimizations such as request aggregation, in-network caching, and multipath retrieval that can significantly enhance large-scale data distribution. Additionally, a forwarding strategy layer in NDN can actively measure network conditions, adapt to changes without involving user applications and provide intelligent request forwarding. This new networking model removes a substantial burden from applications and makes them more straightforward to implement.

In this work, we use High Energy Particle Physics (HEP) data transfers as the driving example to demonstrate NDN's flexibility and versatility at the network layer. While NDN provides other benefits to scientific data such as provenance and content-centric security, we omit their discussion for brevity and concentrate solely on data transfer. First, we propose at a high level, two protocols for NDN-based bandwidth reservation and on-demand path creation. We demonstrate how NDN can dynamically create strategically placed, in-network caches that can reduce hot spots and network resource consumption. We use these newly-proposed protocols along with two NDN strategies to build a deadline based data transfer solution. Finally, we analyze a one-week long HEP data access log to demonstrate how NDN can lower bandwidth consumption by orders of magnitude. In addition to providing a real use case to the NDN community, our work informs the HEP community about data distribution problems.

II. BACKGROUND ON SCIENTIFIC DATA TRANSFERS

Previous work has classified scientific data transfers into two broad categories - bulk data transfer and interactive traffic [7]. In this section, we briefly present these two data transfer modes and discuss why we only consider bulk data transfer. We then point out the significant problems associated with bulk data transfers over TCP/IP networks. Later we use these shortcomings to motivate our NDN based solution.

A. Data Transfer Modes

Bulk data transfers move a considerable amount of data over long distance links. For some workflows, such data transfers can reach more than a Terabyte per day [8]. In addition to transferring data for archiving, replication, or local analysis, researchers also pre-place data copies in caches around the world for efficient, CDN-like access [8]. Data pre-placement has been particularly popular with communities such as the LHC [8], which routinely places a significant amount of data near the users. Bulk data transfers are not overly sensitive to RTT but require a significant amount of bandwidth for a long time and no packet loss. However, a substantial amount of dedicated bandwidth for an extended period is challenging to acquire on public networks.

The other type of scientific traffic is interactive and comes from applications such as data visualization and audio/video conferencing tools. However, such data flows may be short-lived, personalized, and are often generated on-demand, making caching, aggregation, or scheduling less useful for them. We do not consider interactive traffic in this study for these reasons.

B. Problems with Traditional Bulk Data Transfer

Currently, there are two ways to download data over IP networks. When transfers are not very large or does not need to meet a deadline, data is requested using HTTP or FTP over a shared network without any QoS guarantees [1]. When it is critical that requested data reach the requester by a deadline, flows are separated from other traffic using reserved resources. Reservations include router resources such as queue capacity and interfaces, as well as the capacity of network links. We refer this traffic as “resource-reserved traffic” in this work. However, a reservation does not eliminate the underlying architectural shortcomings of TCP/IP. In this section, we discuss these shortcomings in the context of large scientific data flows.

The End-to-end paradigm is unable to use network resources efficiently : In IP all data transfers are between two end hosts. This end-to-end model means two clients must download the same data separately even if they share the same network path, wasting network bandwidth. Though scientific communities use reserved bandwidth channels regularly to avoid congestion and packet loss, such channels are still end-to-end tunnels, and neither the channels nor the data flowing through them are reusable, leading to lower goodput.

Stateless forwarding cannot adapt to network changes: Traditional IP networks do not keep any state in the network

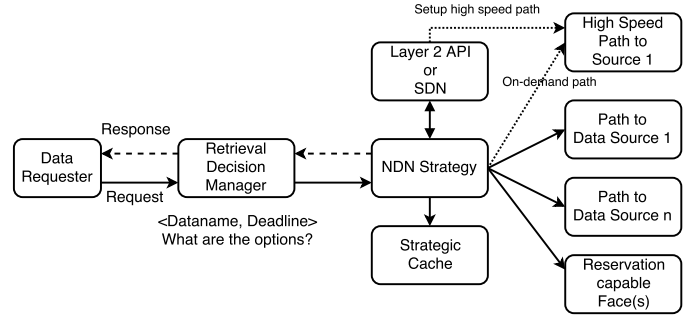


Fig. 1: Overview of a deadline-based data transfer protocol.

which means they are unable to react to changing network conditions. In case of failure or service degradation, data transfers continue to use the same path unless there is some external intervention. Scientific communities have deployed intelligent applications and middleware solutions that keep track of the network performance at the expense of increased application complexity.

Inability to use multiple paths: Even when multiple replicas are available, IP cannot utilize them simultaneously. Though the networking community has proposed workarounds such as multiple connections at the application layer or multipath TCP [9] at the transport layer, these approaches still require knowledge of the underlying network.

TCP congestion control interferes with transfer speed: For a large bandwidth link, losing even one in hundreds of thousands of packets can dramatically reduce transfer speed [10]. To circumvent congestion and packet loss in public networks scientific communities have built dedicated science networks such as the LHC Optical Private Network (LHCOPN) [11] and ESNNet [12] that provide dedicated paths for science data. However, traffic flowing over these networks may still encounter congestion and packet loss, especially when other scientific flows are competing for resources.

NDN addresses these problems at the network layer. For example, NDN uses caching and Interest aggregation to efficiently use available bandwidth by reducing request duplication; name-based forwarding adapts to network changes immediately; intelligent forwarding strategies use multiple paths, and NDN’s hop-by-hop congestion control mechanism does not need to slow down consumers if another path is available [13].

III. BUILDING AN NDN-BASED LARGE DATA TRANSFER PROTOCOL

This section presents a high-level overview of two new protocols, an NDN-based Bandwidth Reservation Protocol, and an NDN-based circuit-creation protocol. Besides, we discuss two NDN forwarding strategies that we later use. Fig. 1 shows a high-level overview of our protocol. Note that while we describe these strategies separately, they can be integrated into a single larger strategy. Descriptions of these constructs are deliberately high-level since the goal of this paper is not to

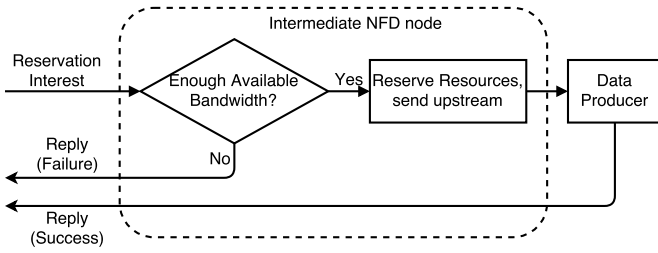


Fig. 2: Reservation with NDN.

present the protocol specifications but to simply demonstrate NDN’s capabilities as the network layer.

A. Bandwidth Reservation Protocol

Distributed, connected systems with many simultaneous users such as Computing Grids routinely encounter resource contention and congestion leading to data transfer delays. Satisfying transfer deadlines in such an environment often requires dedicated per-flow bandwidth allocation [2]. In this work, we propose a protocol to create hop-by-hop reservations in an NDN network. Fig. 2 shows a high-level overview of the reservation protocol; to set up a reservation an NDN node sends a special reservation Interest that is forwarded hop-by-hop upstream. A tuple containing **<data name, requested bandwidth, start_time, deadline>** represents the reservation request. Each node has a reservation manager that checks if the requested bandwidth is available during the requested period. If the request is successful, the reservation manager forwards the Interest upstream. Upon reaching a data producer or a repository, the Interest brings back a reply with a success message. The reservation manager keeps track of the reservation using a reservation table similar to TABLE I.

While similar to RSVP [14], there are two important differences between our reservation protocol and RSVP: (a) Our reservation is per name prefix, not end-to-end. Per-prefix reservation means transfers can share the reservation as long as they share some of the network path, and (b) the NDN reservation manager can aggregate requests for the same content if they fall within a common deadline, a feature not available in RSVP.

B. Integrating Dynamic Path Creation with NDN Strategies

Sometimes existing bandwidth is simply not enough to satisfy a request deadline without a reservation. At the same time, creating permanent, high-bandwidth paths between all sites is not economically feasible. Moreover, since scientific data flows are often bursty [15], creating permanent paths is far from optimal. The scientific communities address this short-term bandwidth shortage problem by creating temporary, high-bandwidth paths for large data transfers.

ESnet’s On-Demand Secure Circuits and Advance Reservation System (OSCARs) [16] is a service that allows users to create such guaranteed bandwidth-reserved paths. However, users are still responsible for knowing the endpoints, creating paths, and scheduling transfers. Arbitrary creation of reserved

TABLE I:
RESERVATION SCHEDULING TABLE

ReqID	Prefix	Requested StartTime	Deadline	BW
1	/xrootd	1463330393	1463355592	1Gbps
2	/xrootd	1463330519	1463355623	1Gbps

paths may create conflict between users, and the network resources may not be optimally utilized. We argue that the network, not the users should be in charge of creating network paths. Our protocol allows NDN to provide a network-driven approach to path creation. In our implementation, an NDN strategy invokes OSCAR’s path creation mechanism when the available bandwidth on existing paths are not sufficient to meet a deadline. If the path creation succeeds, the strategy adds a new route to the FIB. Since NDN consolidates requests for the same data, a high-speed path can potentially speed up other best-effort traffic if such flows are temporally close to the high-speed flow. Note that strategies are not constrained to using a specific lower layer protocol but can interact with any protocol or an SDN controller to create similar high-speed paths.

C. Delay-based Forwarding Strategy

NDN supports forwarding strategies that record RTT on each outgoing link. On receiving the first Interest for a namespace (e.g., **</xrootd/data1>**), the strategy sends it over all matching interfaces. Once data comes back, it records the RTT of each incoming Data packet before forwarding it downstream. It then ranks the faces based on RTT and uses that ranking to forward subsequent Interests. Periodically the strategy tries out other, lower ranked interfaces, and as Interest/Data exchange continues the strategy adjusts the ranking based on new observed RTTs. At any given time this strategy chooses the lowest latency path for data retrieval.

D. Multipath Strategy

If multiple producers are reachable from a node, NDN can route packets simultaneously for data retrieval. In our example, we created a strategy that can send Interests over multiple links. This strategy can also forward Interests based on the Interface ranking. For example, if two routes are available, NFD might forward 75% of the Interests to route one and the rest to route two based on their ranking. We use available bandwidth and RTT as the metrics for initial Interface ranking and then adjust based on constant measurement. Unlike the previous strategy, this strategy utilizes multiple links for data transfer.

E. Namespace

Hierarchical NDN names align well with existing scientific namespaces. We have shown in the previous work [17] [18] that many scientific domains already use hierarchical names that we can use unmodified or with minor changes. In this work, we use xrootd [19], a data management tool for High Energy Physics datasets (HEP) as an example. We assume data is published under a root prefix, e.g., **</xrootd>**. Different sub-namespaces under the root prefix identify data and various

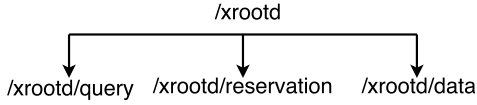


Fig. 3: Namespace Design.

services. Fig. 3 shows such an example; actual data is served under `</xrootd/data>` while two other services are advertised under `</xrootd/query>` and `</xrootd/reservation>`. The first service allows applications to query the current network state for the `</xrootd>` prefix. The second provides a reservation service that a strategy can use to set up a reserved path for `</xrootd>`.

IV. A DEADLINE-BASED DATA TRANSFER PROTOCOL: DESIGN AND IMPLEMENTATION

In this section, we use the NDN-based primitives we discussed previously to design a deadline-based data transfer protocol for both best-effort and reserved bandwidth data transfers.

The reference implementation of our protocol has three main components; a data requester/client, a per-node retrieval decision manager and custom NDN strategies. The forwarding strategy controls intelligent Interest forwarding decisions and when necessary, reserves bandwidth, creates strategically placed in-network caches and interacts with the upper/lower layer protocols for dynamic path creation. The retrieval manager acts as an intermediary between the client and the strategy. In addition to communicating with strategies and the clients, the retrieval manager works as a policy module to enforce retrieval or reservation quotas. Policies are needed to ensure applications do not force the network to use dedicated paths for all transfers by setting impossible deadlines.

A. Component Interaction

In our protocol, the client notifies the network of its requirements by sending an Interest packet to the retrieval manager. The Interest takes the following form: `<data name, deadline, dataset size, hard/soft deadline flag>`. The name of the dataset defines the requested dataset; the retrieval deadline denotes the latest acceptable time for data retrieval; a “hard deadline” flag means the deadline is non-negotiable while a “soft deadline” flag denotes best effort traffic. While transfer time for requests with soft deadlines is not guaranteed, the strategy still may use intelligent forwarding, e.g., multipath forwarding, to fulfill the request within a reasonable time. The Interest also tells the retrieval manager the size of the data. While estimating dataset sizes is not easy for general Internet traffic, scientific data sizes are usually recorded in a catalog, and therefore relatively easy to estimate. We have described such a distributed scientific catalog in our previous work [17].

If the retrieval manager sees two or more requests with overlapping deadlines, it simply aggregates them and suggests a start time to the clients. Otherwise, the retrieval manager talks to the local NFD about possible retrieval options using a special query namespace. For querying retrieval options for

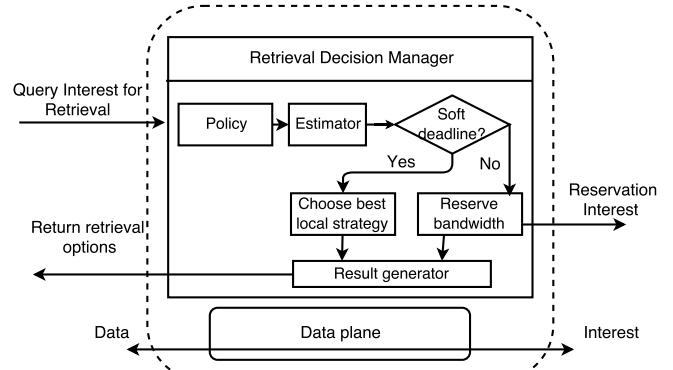


Fig. 4: Strategy decisions on client node's NFD.

`</xrootd>`, the retrieval manager sends an Interest to the NFD with the following structure: `</xrootd/query/</xrootd/data1/start_time/deadline/deadline_type>>`. On receiving this Interest, NFD compiles a list of options and returns it to the retrieval manager.

B. Fulfilling Requests with Soft Deadlines

Fig. 4 shows the decision path for such Interests. If the deadline type is soft, our custom NDN strategy looks up the list of faces in the FIB that can be used for retrieving `</xrootd/data1>`. If the Interest is under a namespace that was not previously used for data retrieval, the strategy fetches a few chunks using each matching face and records the following information for each face: `<FaceID, RTT, Max Bandwidth>`. The strategy then compiles retrieval options and sends it to the retrieval manager, which in-turn notifies the client. Instead of sending binary yes/no response to the client, our protocol replies with more detailed return values along with a suggested start time; for our implementation they are (a) the request can be satisfied, and the client starts retrieval immediately; (b) the retrieval can be satisfied only with an extended deadline; if the new deadline is acceptable, the client adjusts the deadline and requests again; (c) the retrieval is aggregated and the client starts retrieval at the suggested time; and (d) the request can not be satisfied. Such fine-grained information may enable the clients to make intelligent decisions, a feature that is not available today. However, since the network condition may change after the initial reply, there is no guarantee that network will be able to satisfy the soft deadline.

C. Reserved Bandwidth Path for Hard Deadlines and Strategic In-network Caching

While strategies such as multi-path may work well for best-effort traffic, the only way to guarantee timely completion of large data transfers is to create a reserved bandwidth path [1]. In case the deadline is “hard”, the strategy must create a reserved path to a publisher or a cache. In our implementation we send a reservation Interest using a special namespace, `</xrootd/reservation/>`. A reservation Interest looks like `</xrootd/reservation/>`

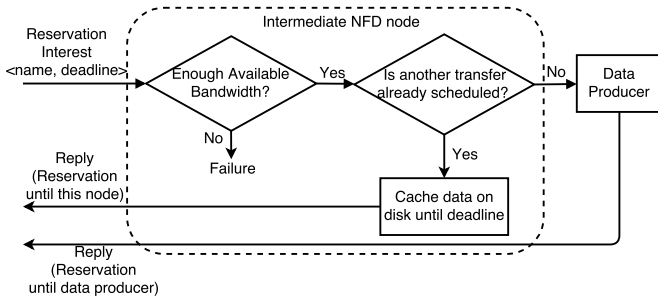


Fig. 5: Using reservation with strategic caching.

`</xrootd/data1/start_time/deadline/bandwidth>>`. On receiving this Interest, a node reserves the appropriate amount of downstream bandwidth for future incoming Data Packets. We assume Interest packets are small and the path can support them without requiring bandwidth reservation. If the reservation is successful, the node forwards the reservation Interest upstream. If the node is the publisher, it returns a success message.

The reservation protocol can be tweaked to create intelligent data dissemination strategies; Fig. 5 shows an example. If a new request overlaps with an existing one, our strategy merges the requests and sends back a reply indicating success and the time of the reservation. Note that in this case the reserved path is created only between the client and the replying node. In addition to performing this simple aggregation, an intelligent strategy may create a temporary cache in the intermediate nodes. For example, a node may decide to cache data for `</xrootd>` until $t_{deadline}$ if there are n requests scheduled between now and time $t_{deadline}$. Since scientific datasets show a high degree of temporal locality [20], in-network strategic caching is helpful for these data flows.

Our method has two benefits for scientific datasets: unlike today, an end-to-end per-client path reservation is not required which frees up network resources. Second, our strategy can dynamically create in-network caches without requiring prior planning and operator intervention.

D. Interacting with Other Layers

If none of the existing options can meet the deadline, the network must create a new high-speed path. We built an NDN strategy that works with Layer 2 protocols (or SDN) to create such a path. In this work, our prototype implementation interfaces with OSCARS [21] to set up reserved paths using a strategy. Once the strategy on a node decides that a new path is needed it calls the OSCARS API which then creates a new VLAN between the node and a data producer. Once the high-speed path is set up, the strategy uses it as simply another available link.

Our protocol is generic and should work in both NDN-only networks, and NDN overlays over IP. No additional mechanism is required for fulfilling requests with soft deadlines; however, meeting hard deadlines will require QoS guarantees not only from the NDN entities but also from the underlying IP routers.

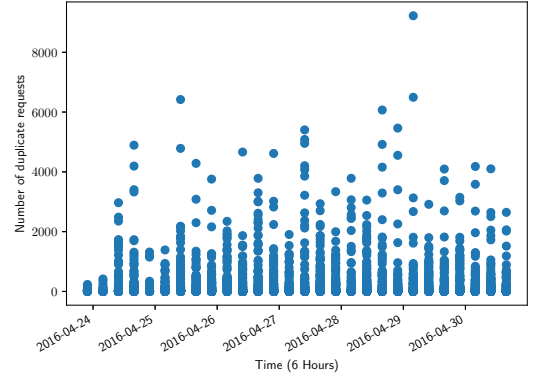


Fig. 6: Duplicate requests for individual datasets over time.

V. EVALUATION

The deadline-based data transfer protocol we described optimizes network usage by aggregating requests, caching, and intelligent request scheduling. We are working on implementing and evaluating the strategy we described in the previous sections. In this section, we analytically explore how much bandwidth our NDN-based protocol can save for a sample HEP data flow. We analyze a xrootd [19] access log recorded from Apr 23th, 2016 to Apr 30, 2016. The logs had 114K unique requests from 267 users, recorded at a ten-minute interval. The access logs showed a high degree of duplicate requests; users requested only 1871 unique datasets over 114K requests; so on average, each dataset was requested sixty times. Duplicate requests can occur in xrootd for popular datasets or if transfers fail, which then automatically triggers another request for the same data. NDN can optimize HEP data flow by de-duplicating requests using our deadline based protocol. For combining the requests, we introduce a “scheduling window”; requests falling within this window can potentially be combined.

To investigate temporal locality within our scheduling window, we first separated the requests in (arbitrary) six-hour bins. The actual window will depend on network capacity, storage, and individual workflows. Fig. 6 shows the number of duplicate requests over time; each dot represents the number of requests for a specific dataset in a six-hour window. We find that many datasets were requested several thousand of times and over 50% of the requests have one or more follow-up request(s) within 10 minutes. Having so many duplicate requests in the log is good news for our strategy since they can be combined efficiently.

We assume each request was for a 2GB file, the average file size in xrootd [22]. Intelligent request scheduling allows NDN to retrieve only one copy of the data that satisfies all requests for the same data. To compare IP’s bandwidth consumption with NDN, we first calculate the amount of bandwidth needed for each request and then calculate the total aggregate bandwidth required to serve all requests over six-hour periods. Fig. 7 shows that the total bandwidth requirement for xrootd is very high, with peaks at approximately 64 Gbps. We also calculated how much bandwidth NDN could save compared

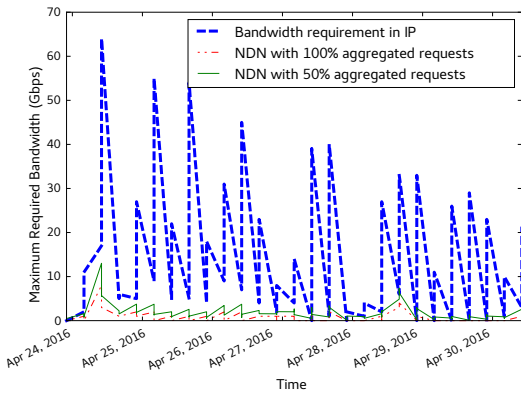


Fig. 7: Reduced bandwidth consumption with NDN.

to IP when requests are aggregated. We take the same requests over 6 hours, de-duplicate the requests and calculate the total bandwidth requirement.

Looking from the servers' point of view, Fig. 7 shows overall bandwidth demand of the system in two scenarios. First one is the best case; if we can aggregate all duplicate requests over a 6-hour period, the max bandwidth requirement at the servers drops from 64 Gbps to around 8.2 Gbps, an 85% reduction. However, we acknowledge that not all requests can be aggregated; some requests might have very tight deadlines and need to be served immediately. Even if we assume only 50% of the duplicate requests can be aggregated using our protocol, the bandwidth requirement comes down to 13.2 Gbps, a 79% reduction. This result shows that even some degree of de-duplication can go a long way in reducing resource consumption for HEP data flows.

While our result demonstrates the improvements an NDN based intelligent network layer can bring, we did not evaluate bandwidth reservation or strategic caching in this work. We are optimistic that incorporating these features will improve our example data flow even more. We will investigate them in future work.

VI. CONCLUSIONS AND FUTURE WORK

In this work, we used deadline-based data transfers as the driving example to demonstrate that an NDN based network layer can be more flexible and versatile compared to the current IP networks. Additionally, we proposed an RSVP-like bandwidth reservation protocol for NDN. We show that NDN does not always require end-to-end reserved paths for high-speed transfers. Instead, a reserved path to the nearest cache or repository can speed up transfers and at the same time, improve network utilization. We also present a mechanism that enables NDN strategies to interact with lower layer protocols to set up dedicated paths for large data transfers. Finally, we used a real HEP access log to demonstrate that our NDN-based protocol can potentially reduce bandwidth consumption by over 75% for this particular example.

We acknowledge that our study is preliminary and can be improved in several ways. We are working on improving the

implementation to include in-network strategic caching and the NDN-based reservation protocol. Once these pieces are implemented, we plan to evaluate our complete protocol using a real topology and more detailed access logs.

REFERENCES

- [1] B. Tierney, E. Kissel, M. Swamy, and E. Pouyoul, "Efficient data transfer protocols for big data," in *E-Science (e-Science), 2012 IEEE 8th International Conference on*. IEEE, 2012, pp. 1–9.
- [2] B. B. Chen and P. V.-B. Primet, "Scheduling deadline-constrained bulk data transfers to minimize network congestion," in *Cluster Computing and the Grid, 2007. CCGRID 2007. Seventh IEEE International Symposium on*. IEEE, 2007, pp. 410–417.
- [3] J. Rehn, T. Barrass, D. Bonacorsi, J. Hernandez, I. Semeniouk, L. Tuura, and Y. Wu, "Phedex high-throughput data transfer management system," *Computing in High Energy and Nuclear Physics (CHEP) 2006*, 2006.
- [4] W. Lu, Z. Zhu, and B. Mukherjee, "Optimizing deadline-driven bulk-data transfer to revitalize spectrum fragments in eons," *J. Opt. Commun. Netw.*, no. 12, pp. B173–B183, 2015.
- [5] S. M. Srinivasan, T. Truong-Huu, and M. Gurusamy, "Flexible bandwidth allocation for big data transfer with deadline constraints," in *Computers and Communications (ISCC), 2017 IEEE Symposium on*. IEEE, 2017, pp. 347–352.
- [6] L. Zhang, A. Afanasyev, J. Burke, V. Jacobson, P. Crowley, C. Papadopoulos, L. Wang, B. Zhang *et al.*, "Named data networking," *ACM SIGCOMM Computer Communication Review*, vol. 44, no. 3, pp. 66–73, 2014.
- [7] I. Foster, M. Fidler, A. Roy, V. Sander, and L. Winkler, "End-to-end quality of service for high-end applications," *Computer Communications*, vol. 27, no. 14, pp. 1375–1388, 2004.
- [8] A. Barczyk, "World-wide networking for lhc data processing," in *National Fiber Optic Engineers Conference*. Optical Society of America, 2012, pp. NTu1E–1.
- [9] S. Barré, C. Paasch, and O. Bonaventure, "Multipath tcp: from theory to practice," *NETWORKING 2011*, pp. 444–457, 2011.
- [10] B. Tierney and J. Metzger, "High Performance Bulk Data Transfer," <https://fasterdata.es.net/assets/fasterdata/JT-201010.pdf>.
- [11] E. Martelli and S. Stancu, "Lhcopn and lhcone: Status and future evolution," in *Journal of Physics: Conference Series*, vol. 664, no. 5. IOP Publishing, 2015, p. 052025.
- [12] ESNet. <https://fasterdata.es.net>.
- [13] K. Schneider, C. Yi, B. Zhang, and L. Zhang, "A practical congestion control scheme for named data networking," in *Proceedings of the 2016 conference on 3rd ACM Conference on Information-Centric Networking*. ACM, 2016, pp. 21–30.
- [14] L. Zhang, S. Deering, D. Estrin, S. Shenker, and D. Zappala, "Rsvp: A new resource reservation protocol," *Network, IEEE*, vol. 7, no. 5, pp. 8–18, 1993.
- [15] A. Shoshani and D. Rotem, *Scientific Data Management: Challenges, Technology, and Deployment*. CRC Press, 2009.
- [16] C. Guok, "A user driven dynamic circuit network implementation," *Lawrence Berkeley National Laboratory*, 2009.
- [17] C. Fan, S. Shannigrahi, S. DiBenedetto, C. Olschanowsky, C. Papadopoulos, and H. Newman, "Managing scientific data with named data networking," in *Proceedings of the Fifth International Workshop on Network-Aware Data Management*. ACM, 2015, p. 1.
- [18] C. Olschanowsky, S. Shannigrahi, and C. Papadopoulos, "Supporting climate research using named data networking," in *Local & Metropolitan Area Networks (LANMAN), 2014 IEEE 20th International Workshop on*. IEEE, 2014, pp. 1–6.
- [19] L. Bauerdick, K. Bloom, B. Bockelman, D. Bradley, S. Dasu, I. Sfiligoi, A. Tadel, M. Tadel, F. Wuerthwein, and A. Yagil, "Xrootd monitoring for the cms experiment," in *Journal of Physics: Conference Series*, vol. 396. IOP Publishing, 2012, p. 042058.
- [20] S. Shannigrahi, C. Fan, and C. Papadopoulos, "Request aggregation, caching, and forwarding strategies for improving large climate data distribution with ndn: A case study," in *To Appear in ICN17*, 2017.
- [21] C. Guok and D. Robertson, "Esnet on-demand secure circuits and advance reservation system (oscars)," *Internet2 Joint*, 2006.
- [22] S. Shannigrahi, C. Papadopoulos, E. Yeh, H. Newman, A. J. Barczyk, R. Liu, A. Sim, A. Mughal, I. Monga, J.-R. Vlimant *et al.*, "Named data networking in climate research and hep applications," in *Journal of Physics: Conference Series*, vol. 664. IOP Publishing, 2015, p. 052033.