

Huge Data: A Computing, Networking, and Distributed Systems Perspective Workshop Report

The Networking & Information Technology R&D Program
Large Scale Networking
Interagency Working Group

January 2021



Contents

Executive Summary	1
Introduction	2
Important Issues in Huge Data Research	2
Huge Data Challenges and Opportunities	3
Data Generation and Storage	3
Data Movement.....	4
Data Processing and Security.....	5
How to Enable and Sustain Big Data	5
New Areas of Huge Data Research	5
New Types of Huge Data and Obstacles to Accessing Them	6
Cross-Disciplinary Collaboration.....	6
Critical Research Infrastructure Needs.....	6
Conclusion	7
Abbreviations.....	8
About the Authors	8
Acknowledgments	8

Copyright Notice: This document is a work of the United States Government and is in the public domain (see 17 U.S.C. §105). It may be freely distributed and copied with acknowledgment to the NITRD Program. This and other NITRD documents are available online at <https://www.nitrd.gov/publications/>. Published in the United States of America, 2020.

Executive Summary

On April 13-14, 2020, the Large Scale Networking Interagency Working Group of the Networking and Information Technology Research and Development Program held a workshop¹ to explore new paradigms to address the challenges and requirements of “huge data” science and engineering research. The workshop brought together domain scientists, network and systems researchers, and infrastructure providers to address the problems associated with processing, storing, and transferring huge data. This document summarizes the workshop discussions.

There is an ever-increasing demand in science and engineering for the creation, analysis, archiving, and sharing of extremely large datasets, or “big data.” Increasing evidence suggests that these demands may soon outgrow current big data methods and infrastructure capacity, and therefore, the community has begun using the term “huge data.” The dataset from a blackhole image collected over a period of 7 days by the Event Horizon Telescope is approaching 5 petabytes of data. The Large Hadron Collider generates 2000 petabytes of data over a typical 12-hour run. These datasets can grow continuously without bounds. They are often collected from distributed devices (e.g., sensors), processed on-site or at distributed clouds, and placed or duplicated in distributed sites for reasons of reliability, scalability, and availability. Measurement, generation, and transformation of data over distributed locations is stressing the computing paradigm; however, efficient processing, persistent availability, and timely delivery (especially over wide areas) of huge data is critically important to the success of scientific research.

The size of the data collected today has surpassed the maximum size assumed in prior research. Today, most computing systems and applications operate based on clear delineation of data movement and data computing. Data is moved from one or more data stores to a computing system where computation is performed locally. This framework requires significant storage capacity at each computing system as well as significant time for data transfer before and after the computation. Looking forward, researchers see potential benefits in a completely new paradigm that supports *in situ* computation, at unprecedented scales, across distributed computing systems that are interconnected by high-speed networks. These high-performance data transfer functions will be closely integrated in software (e.g., operating systems) and hardware infrastructures, and they will avoid bottlenecks for scientific discoveries and engineering innovations through much faster, more efficient, and more scalable computation across a globally distributed, highly interconnected, and vast collection of data and computation infrastructure.

The Huge Data workshop included 120 participants and the presentation of 40 white papers in four topical areas: data generation, data storage, data movement, and data processing and security. One key observation of participants was that huge data problems often arise from continuously generated, domain-specific data that is extremely big, does not warrant long-term storage, and tends to “break” when using existing networking and computing infrastructure and methods. To enable and sustain the effective use of huge data going forward, focused research is needed. Suggestions from the workshop included new research on data infrastructure and methods, systems research on new architecture and integration methods, increased cross-disciplinary collaboration, and new research infrastructure.

¹ <https://www.nitrd.gov/nitrdgroups/index.php?title=LSN-Huge-Data-2020>

Introduction

On April 13-14, 2020, the Large Scale Networking (LSN) Interagency Working Group (IWG) of the Networking and Information Technology Research and Development (NITRD) Program held its workshop,² “Huge Data: A Computing, Networking, and Distributed Systems Perspective.” The purpose of the workshop was to engage members from stakeholder communities to discuss important issues in huge data research, including how to define huge data’s nature and scope, identify the challenges and opportunities it presents, and suggest how to enable and sustain it. Workshop attendees included university and national lab researchers; operators of regional, national, and international research and education (R&E) networks; and Federal agency representatives.

This report provides input for the LSN IWG and its member agencies to consider opportunities and strategies to enable and sustain huge data research. Following the keynote presentation, workshop presentations and discussions were organized around four topical areas: data generation, data storage, data movement, and data processing and security. (Please note in this report data generation and storage were combined into one section.) Follow-on discussions focused on new areas of huge data research, new types of huge data, cross-disciplinary collaboration, and critical research infrastructure needs.

Important Issues in Huge Data Research

The workshop’s background and purpose were established in the opening address, *[Huge Data is outgrowing the Internet’s file transfer protocols](#)*, by the following statement: Roughly 1 in every 121 huge file transfer delivers bad data.³ With this as the premise, the following observations set the stage for the workshop’s 40 presentations and numerous informative discussions:

- Current transmission control protocol (TCP) checksums are known to be weak for detecting many types of errors, and such errors typically occur on hosts and routers due to memory and data bus timing problems.
- Evidence suggests some middleboxes⁴ (which are prevalent in networks today) overwrite checksums transferred with files with corrupted checksums.
- Increasing speeds on link layer technologies, such as WiFi, are stressing the link layer cyclic redundancy check (CRC-32) checksum algorithm’s ability to catch link-level errors.
- Errors occur often enough that many scientists, as a best practice, retrieve data from multiple sites and compare the copies to detect errors.
- There has not been a major study of errors in networks for 20 years.⁵ Research is needed into new file transfer tools that are developed to handle huge data transfer needs and to incorporate innovative error detection and correction methods.

² <https://www.nitrd.gov/nitrdgroups/images/3/3b/LSN-KC-Wang-06092020.pdf>

³ <https://www.nitrd.gov/nitrdgroups/images/0/0a/Craig-Partridge-042020.pdf>

⁴ A device within a network that enforces transport policy such as a firewall or network address translation.

⁵ <http://conferences.sigcomm.org/sigcomm/2000/conf/paper/sigcomm2000-9-1.pdf>

Huge Data Challenges and Opportunities

Data Generation and Storage

Applications that generate huge data are found in many different research areas including medical imaging, genomics, computational pathology, connected vehicles, and the Internet of Things (IoT). All of these areas have projected petabytes to exabytes of sustained data generation, and they have important differences based on how the data is generated, either from scientific instruments or extremely distributed sources.

- *Data from scientific instruments* are too large to be stored *in situ*. Challenges include the need for:
 - Network-connected, professionally operated, large-scale storage facilities.
 - Processing on large-scale, shared, high-performance computing (HPC) facilities.
 - Preplanned studies using subsets of data (i.e., because saving everything is impractical).
 - Lossy compression⁶ methods that can preserve signals of interest in lieu of storing all the data from huge datasets.
 - Increased numbers of data and metadata management professionals.
 - Centralized and converged compute and storage facilities (i.e., data movement leads to reduced computing throughput and higher cost). However, it was noted that distributed, federated facilities may still be needed.
 - Storage in dedicated and secure storage facilities to ensure data privacy compliance.
- *Data from extremely distributed sources* (e.g., vehicles and IoT devices) are generated at smaller rates per device but from large numbers of devices. Challenges include the need for:
 - Models that support mostly end-system device use cases.
 - Very large, persistent, cloud-bound data flows for some data analytics applications. (However, the majority of such device data can be processed at the edge or on the device.)
 - Trusted execution environment computing for devices that convey private information.
 - Advanced compression methods to address device resource constraints and heightened attacks (i.e., compressed data and models can be executed in the limited on-processor, trusted execution hardware).
 - Data storage.

Workshop participants identified radio astronomy (e.g., at the National Radio Astronomy Observatory [NRAO]); high-energy physics (e.g., at the High-Luminosity Large Hadron Collider [HL-LHC]); and large scholarly research data (e.g., at CiteSeerX⁷) as examples of major projects in distinct domains that illustrate the challenges involved in storage and access solutions for extremely large data. In addition, they identified scalability of data handling at the network level (exposed buffer processing) and the server level (computational antenna analysis) as important topics to be considered.

As examples of the huge data challenges, NRAO's next-generation Very Large Array facility plans to bring data from 244 antenna sites around the globe to a single high-performance processing center for processing at ~60 petaflops and archiving at least 240 petabytes of data per year. The computing requirements can be met by a contemporary HPC center such as the Texas Advanced Computing Center, but a long-term storage solution has not yet been identified. HL-LHC produces high-energy particle

⁶ Lossy compression is a method of data compression in which the size of the file is reduced by eliminating data in the file. In doing so, image quality is sacrificed to decrease file size. Any data that the compression algorithm deems expendable is removed from the image, thereby reducing its size.

⁷ CiteSeerX is a public search engine and digital library for scientific and academic papers.

collision data at ~1 exabytes/year at the European Laboratory for Particle Physics, which then sends the data to member facilities across the globe for processing and storage. The tremendous data volume requires that the data be stored at multiple distributed storage sites, then retrieved and sent to distributed processing centers to be computed, then results returned to the distributed storage sites to be archived. In the NRAO case, a converged datacenter with sufficient compute and storage capacity may be the most efficient solution. In the HL-LHC case, a facility does not exist that has the capacity to store or process all of its data.

Key observations of workshop participants include:

- Storage is not a one-size-fits-all decision. The type of storage depends on how data is produced, when and where it needs to be processed, and the scope of data to be retrieved for research inquiries.
- Storage is a multi-tiered function that reflects how data is handled throughout its life cycle. Storage of different types (volatile vs. non-volatile, fast vs. slow, large vs. small) reside in or on tiers that consist of data sources, processing servers, data pipelines, and long-term storage. The data processing workflows and requirements (e.g., quality, security, privacy, and throughput) that are application- and domain specific directly influence how data storage should be implemented at each tier. Different network architectures (e.g., Internet Protocol, Named Data Networking, and Exposed Buffer Processing) can also influence the choice of storage approaches.

Data Movement

Managers and experts from very large data projects have varied perspectives on research data transfer and access methods. Workshop discussions on movement of huge data centered around the topics of data management; operational support; and protocols, tools, and infrastructure.

Data Management

- The provenance and quality of data (e.g., authenticity, correctness, security, and privacy) is especially important for huge datasets because these datasets require significantly more resources to access the data, and they have potentially more serious implications if there are errors or policy violations.
- Broadening the FAIR (findable, accessible, interoperable, and reusable) data management principles when managing huge data is important to enable efficient search, access, and processing.

Operational Support

- Proactive monitoring and coordinated troubleshooting of high-speed R&E networks is valuable.
- International coordination of network infrastructure is increasingly important because of the global scope of most huge data research efforts.

Protocols, Tools, and Infrastructure

- Advanced networking solutions have seen heavy investment by the most mature large data projects (e.g., high energy physics and astronomy) in order to speed up and streamline data transfer and processing. This includes such technologies as cloud datacenter networking, network function virtualization, programmable wide area networking, traffic engineering, data transfer tools, and data transfer node architecture.
- Data transfer protocols should be enhanced to consider the compression, packet sizing, transfer speed, and latency tradeoffs.
- Network orchestration should be workflow-aware and infrastructure-aware.
- “Reading” a file is inefficient. An approach is needed that is more innovative, efficient, and broadly supported for reading a file (without downloading the whole file).

- Enhanced data repositories with built-in analytic capabilities for a broad range of data types are needed. Taking the computation to the data rather than moving the data to the computation is a very useful function to develop.
- Another digital divide can be avoided by improving smaller institutions' data infrastructure to make huge data research possible.

Data Processing and Security

Huge data presents a range of diverse issues for data processing and security including the need for workflow management, infrastructure, program analysis, platform optimization, and cache management. For example, the increasing importance of cyber-physical systems and IoT has reshaped the computing paradigm—the entire network (with its connected, distributed computing resources) has become a major source of huge data. Needs and suggestions from the discussion included the following.

- Innovative techniques and technologies are needed for:
 - Debugging bad performance.
 - Detecting in-network anomalies using artificial intelligence/machine learning (e.g., for TCP traffic).
 - In-network computing for memory-centric applications and streaming data analysis.
 - Developing policies for interdomain data sharing and resource use.
- Suggested techniques and technologies to consider include:
 - Both application-aware networking (e.g., software-defined networking [SDN]) and network-aware application (e.g., application layer traffic optimization) when designing a system.
 - Error-bounded lossy compression for reducing dataset size where appropriate.
 - Federated machine learning and data partitioning for smart transfers.⁸
 - Named Data Networking for new models for data use and access.

How to Enable and Sustain Big Data

New Areas of Huge Data Research

Huge data tends to involve specific data objects or data streams that are “huge”; this is in contrast to big data that is typically a large collection of small data. Therefore, networking and system researchers need to improve their understanding of huge data scientific workflows that include the “deep integration” of sciences and methodologies across technology domains. New ways are needed of composing workflows as “macro processes” that consider elements beyond a single compute node such as compute, storage, and network resources. Workflow design should include fine-grained resource control, temporal control, and quality control. Other suggestions by workshop participants for R&D needs include:

- Enable light computation on storage.
- Select data worthy of keeping by revisiting pipeline design.
- Question Internet design assumptions, for example:
 - Networking nodes need more robust compute and storage.
 - Clarity of resources across layers needs enhancement to enable design of higher-layer algorithms and applications that optimize the use of lower-layer resources.
- Move the computations to the data (data can be centralized or highly distributed, even on mobile platforms.)

⁸ Simple transfers are one-time transfer sessions that require entering all transfer information. Smart transfers are reusable templates with saved transfer settings.

- Develop tools for data privacy and ownership control.
- Address bottlenecks at all levels: storage, memory, network, and processors.
- Allow for the delivery of research data to citizens, not just scientists.

New Types of Huge Data and Obstacles to Accessing Them

Increasingly, researchers need to address a breadth and diversity of huge data that includes genomics, IoT and sensors (streaming data), high-resolution images and video, metadata, and traditional sources of big data that are moving to huge data, such as radio astronomy and high energy physics.

Workshop participants discussed a range of obstacles to accessing these types of data:

- Lack of federated user management.
- Restrictive access policies.
- Inability to view “sketches” of the data (to help determine what is needed before accessing the entire dataset).
- Inability to hold all the data, which then requires stream processing.
- Difficulty accessing highly distributed data across many nodes.
- Constraints on network bandwidth.
- High cost of access (e.g., data egress fees).
- Friction between the firewall and the infrastructure.
- Lack of automation to meet identification and protection requirements.
- Changed network infrastructure expectations resulting from SDN.

Cross-Disciplinary Collaboration

Information technology (IT) sits in the intersection of many disciplines, and all Federal research agencies encourage use of multidisciplinary research teams. For those agencies dealing with huge data projects, reaching across multiple domains helps researchers learn from each other and drives innovative thinking. The National Science Foundation has many programs that distinctively encourage multidisciplinary teams, and the NSF Computer and Information Science and Engineering Directorate continues to discuss new programs that leverage these teams. It also takes effort and structure for such teams to be effective. Funding agencies can strongly influence how collaboration occurs.

Domain scientists must take an active role in driving technology solutions. Many disciplines, not all represented at the workshop, could contribute to the success of projects (e.g., economists, political scientists, information theorists, and game theorists). Productive collaboration, which begins by parties patiently listening and understanding each other’s “language”, can take years to build, but it is worth the effort. Providing incentives to collaborate, such as finding and influencing publication venues and reviews, and finding a common architecture, could facilitate the process.

Critical Research Infrastructure Needs

Workshop participants identified the different types of critical infrastructure that is needed to sustain huge data research:

- Data storage that is less expensive and includes solutions for data retention and stability.
- Distributed and federated infrastructure that includes a focus on shared storage, compute, and network functionality.
- Specific research infrastructure that is based on the problem to be solved (i.e., not all huge data research requires all resources at the same time).

- Data movement options that are less expensive. Research networks that contain QoS (quality of service) solutions when demand approaches capacity.

Conclusion

The Huge Data workshop brought together domain scientists, network and systems researchers, and infrastructure providers to explore new paradigms to address the challenges and requirements of “huge data” science and engineering research. Huge data problems often arise from continuously generated, domain-specific data that are both time sensitive and tend to “break” over existing networking and computing infrastructure and methods. Workshop attendees discussed new research on data infrastructure and methods, system architecture and integration strategies, cross-disciplinary collaboration, and the infrastructure to enable and sustain huge data innovation.

Abbreviations

FAIR	findable, accessible, interoperable, and reusable
HL-LHC	High-Luminosity Large Hadron Collider
HPC	high performance computing
IoT	Internet of Things
IT	information technology
IWG	interagency working group
LSN	Large Scale Networking (NITRD IWG)
NITRD	Networking and Information Technology R&D Program
NRAO	National Radio Astronomy Observatory
R&D	research and development
R&E	research and education
SDN	software-defined networking
TCP	transmission control protocol

About the Authors

The NITRD Program is the Nation’s primary source of federally funded coordination of pioneering IT research and development (R&D) in computing, networking, and software. The multiagency NITRD Program, guided by the NITRD Subcommittee of the National Science and Technology Council Committee on Science and Technology Enterprise, seeks to provide the R&D foundations for ensuring continued U.S. technological leadership and meeting the Nation’s needs for advanced IT. More information is available at <https://www.nitrd.gov>.

The NITRD Program’s LSN Interagency Working Group coordinates Federal R&D in leading-edge networking technologies, services, and performance. The advanced Federal research network infrastructure supports national security, commerce, industry, and scientific research, and enables the robust transfer of data among systems on the ground, on the sea, in the air, and in space. LSN- coordinated Federal R&D aims to ensure that the next generation of the Internet will be scalable, flexible, and trustworthy. More information is available at <https://www.nitrd.gov/groups/lsn>.

Acknowledgments

NITRD’s LSN Interagency Working Group gratefully acknowledges the LSN Co-chairs Deep Mehdi (NSF), Bob Bonneau (OSD), Richard Carlson (DOE/SC), and NITRD Technical Coordinator Joyce Lee; and Kuang-Ching Wang (Clemson University), and co-authors Ronald Hutchins (University of Virginia); James Griffioen (University of Kentucky); and Zongming Fei (University of Kentucky), who helped plan and implement the workshop and write and review this report. Also, we gratefully acknowledge the workshop participants for their contributions to the workshop discussions.