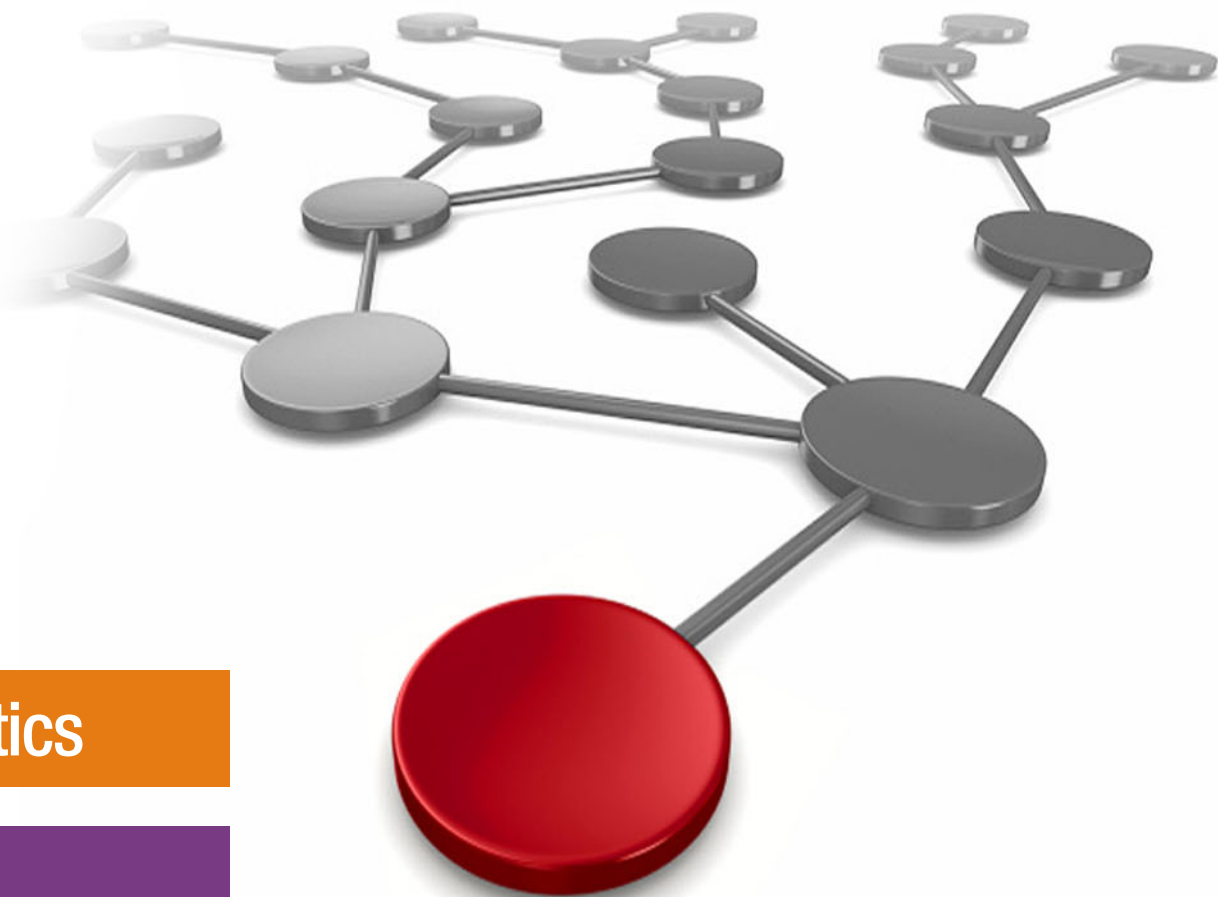


IBM Reference Architecture for High Performance Data and AI in Healthcare and Life Sciences

Dino Quintero

Frank N. Lee, PhD



 **Analytics**

Big Data



International Technical Support Organization

**IBM Reference Architecture for High Performance Data
and AI in Healthcare and Life Sciences**

September 2019

Note: Before using this information and the product it supports, read the information in “Notices” on page vii.

First Edition (September 2019)

© Copyright International Business Machines Corporation 2019. All rights reserved.

Note to U.S. Government Users Restricted Rights -- Use, duplication or disclosure restricted by GSA ADP Schedule Contract with IBM Corp.

Contents

Notices	vii
Trademarks	viii
Preface	ix
Authors	ix
Now you can become a published author, too	xi
Comments welcome	xi
Stay connected to IBM Redbooks	xi
Chapter 1. Trends and challenges for precision medicine	1
1.1 New trend: The era of precision medicine	2
1.2 Challenges	3
1.2.1 Data management challenges	3
1.2.2 Other data challenges	5
Chapter 2. The journey of the reference architecture	9
2.1 The history of IBM Reference Architecture	10
2.1.1 First-generation reference architecture	10
2.1.2 Second-generation reference architecture	11
2.2 Overview of IBM Reference Architecture for High Performance Data and AI	12
2.2.1 Challenges	12
2.2.2 The solution	13
2.2.3 Key values	15
2.3 Datahub for High-Performance Data Analytics	16
2.3.1 Datahub functions	17
2.3.2 Datahub solution and use cases	19
2.4 Orchestrator of High-Performance Data Analytics	20
2.4.1 Orchestrator functions	20
2.4.2 Orchestrator solution and use cases	21
Chapter 3. Deployment model	23
3.1 Composable genomics blueprint	24
3.2 IBM Software-Defined Infrastructure	24
3.3 Multicloud deployment model	25
3.3.1 Clouds over the ocean	25
Chapter 4. Building blocks	29
4.1 IBM Spectrum Storage	30
4.1.1 IBM Spectrum Scale	30
4.1.2 IBM Spectrum Archive	31
4.1.3 IBM Cloud Object Storage	32
4.1.4 IBM Spectrum Discover	34
4.2 IBM Spectrum Computing	34
4.2.1 IBM Spectrum LSF Suite	34
4.2.2 IBM Spectrum Conductor	37
4.3 IBM Power System AC922 for HPC	39
4.3.1 Accelerated computing with IBM POWER9 processor-based systems	39
4.3.2 OpenPOWER Foundation	39
4.3.3 OpenPOWER processors	39

4.3.4 Recent advancements	40
4.3.5 Applications	40
Chapter 5. Use cases	43
5.1 The Broad Institute Genome Analysis Toolkit (GATK)	44
5.2 Expanding IBM Reference Architecture for High-Performance Data Analytics into medical imaging	47
5.2.1 Harnessing AI to transform diagnosis and treatment of brain cancer	47
5.2.2 Pushing the boundaries of traditional medicine	47
5.2.3 Diving into deep learning	48
5.2.4 Giving physicians the tools to excel	49
Chapter 6. Case studies	51
6.1 Sidra Medicine	52
6.1.1 About Sidra	52
6.1.2 The Qatar Genome Project focuses on population health and better treatments	52
6.1.3 Personalized medical advances depend on having a unified view	53
6.1.4 Converging high-performance computing, big data, and cognitive computing	53
6.1.5 Why cognitive computing and IBM	54
6.1.6 A collaboration	54
6.1.7 Software-defined infrastructure for all data and workloads	54
6.1.8 Faster results with scalability, reliability, and speed	55
6.1.9 Adding big data and cognitive computing to high-performance computing	55
6.1.10 Future	56
6.2 Amsterdam UMC	56
6.2.1 Customer background	56
6.2.2 Business challenge	56
6.2.3 Transformation	57
6.2.4 Business benefits	57
6.2.5 Solution components	57
6.3 L7 Informatics	57
6.3.1 Customer background	57
6.3.2 Business challenge	58
6.3.3 Transformation	58
6.3.4 Business benefits	58
6.3.5 Solution components	58
6.4 University of Birmingham	58
6.4.1 Customer background	58
6.4.2 Business challenge	59
6.4.3 Transformation	59
6.4.4 Business benefits	59
6.4.5 Solution components	59
6.5 Thomas Jefferson University	59
6.5.1 Customer background	60
6.5.2 Business challenge	60
6.5.3 Transformation	60
6.5.4 Business benefits	60
6.5.5 Solution components	60
6.6 Biotechnology and Biomedicine Center of the Czech Academy of Sciences and Charles University: BIOCEV	61
6.6.1 Customer background	61
6.6.2 Business challenge	61
6.6.3 Transformation	61

6.6.4 Business benefits	61
6.6.5 Solution components	61
6.7 Washington University St. Louis and Vanderbilt University	62
6.7.1 Customer background	62
6.7.2 Business challenge	62
6.7.3 Transformation	62
6.7.4 Business benefits	62
6.7.5 Solution components	63
Appendix A. Profiling GATK	65
Related publications	71
IBM Redbooks	71
Online resources	71
Help from IBM	72

Notices

This information was developed for products and services offered in the US. This material might be available from IBM in other languages. However, you may be required to own a copy of the product or product version in that language in order to access it.

IBM may not offer the products, services, or features discussed in this document in other countries. Consult your local IBM representative for information on the products and services currently available in your area. Any reference to an IBM product, program, or service is not intended to state or imply that only that IBM product, program, or service may be used. Any functionally equivalent product, program, or service that does not infringe any IBM intellectual property right may be used instead. However, it is the user's responsibility to evaluate and verify the operation of any non-IBM product, program, or service.

IBM may have patents or pending patent applications covering subject matter described in this document. The furnishing of this document does not grant you any license to these patents. You can send license inquiries, in writing, to:

IBM Director of Licensing, IBM Corporation, North Castle Drive, MD-NC119, Armonk, NY 10504-1785, US

INTERNATIONAL BUSINESS MACHINES CORPORATION PROVIDES THIS PUBLICATION "AS IS" WITHOUT WARRANTY OF ANY KIND, EITHER EXPRESS OR IMPLIED, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF NON-INFRINGEMENT, MERCHANTABILITY OR FITNESS FOR A PARTICULAR PURPOSE. Some jurisdictions do not allow disclaimer of express or implied warranties in certain transactions, therefore, this statement may not apply to you.

This information could include technical inaccuracies or typographical errors. Changes are periodically made to the information herein; these changes will be incorporated in new editions of the publication. IBM may make improvements and/or changes in the product(s) and/or the program(s) described in this publication at any time without notice.

Any references in this information to non-IBM websites are provided for convenience only and do not in any manner serve as an endorsement of those websites. The materials at those websites are not part of the materials for this IBM product and use of those websites is at your own risk.

IBM may use or distribute any of the information you provide in any way it believes appropriate without incurring any obligation to you.

The performance data and client examples cited are presented for illustrative purposes only. Actual performance results may vary depending on specific configurations and operating conditions.

Information concerning non-IBM products was obtained from the suppliers of those products, their published announcements or other publicly available sources. IBM has not tested those products and cannot confirm the accuracy of performance, compatibility or any other claims related to non-IBM products. Questions on the capabilities of non-IBM products should be addressed to the suppliers of those products.

Statements regarding IBM's future direction or intent are subject to change or withdrawal without notice, and represent goals and objectives only.

This information contains examples of data and reports used in daily business operations. To illustrate them as completely as possible, the examples include the names of individuals, companies, brands, and products. All of these names are fictitious and any similarity to actual people or business enterprises is entirely coincidental.

COPYRIGHT LICENSE:

This information contains sample application programs in source language, which illustrate programming techniques on various operating platforms. You may copy, modify, and distribute these sample programs in any form without payment to IBM, for the purposes of developing, using, marketing or distributing application programs conforming to the application programming interface for the operating platform for which the sample programs are written. These examples have not been thoroughly tested under all conditions. IBM, therefore, cannot guarantee or imply reliability, serviceability, or function of these programs. The sample programs are provided "AS IS", without warranty of any kind. IBM shall not be liable for any damages arising out of your use of the sample programs.

Trademarks

IBM, the IBM logo, and ibm.com are trademarks or registered trademarks of International Business Machines Corporation, registered in many jurisdictions worldwide. Other product and service names might be trademarks of IBM or other companies. A current list of IBM trademarks is available on the web at “Copyright and trademark information” at <http://www.ibm.com/legal/copytrade.shtml>

The following terms are trademarks or registered trademarks of International Business Machines Corporation, and might also be trademarks or registered trademarks in other countries.

Accesser®	IBM Watson®	Redbooks®
IBM®	LSF®	Redbooks (logo)  ®
IBM Elastic Storage®	POWER®	Slicestor®
IBM Spectrum®	Power Architecture®	Storwize®
IBM Spectrum Conductor®	POWER8®	Tivoli®

The following terms are trademarks of other companies:

Veracity, are trademarks or registered trademarks of Merge Healthcare Inc., an IBM Company.

Linux is a trademark of Linus Torvalds in the United States, other countries, or both.

Microsoft, Windows, and the Windows logo are trademarks of Microsoft Corporation in the United States, other countries, or both.

Java, and all Java-based trademarks and logos are trademarks or registered trademarks of Oracle and/or its affiliates.

Other company, product, or service names may be trademarks or service marks of others.

Preface

This IBM® Redpaper publication provides an update to the original description of IBM Reference Architecture for Genomics. This paper expands the reference architecture to cover all of the major vertical areas of healthcare and life sciences industries, such as genomics, imaging, and clinical and translational research.

The architecture was renamed *IBM Reference Architecture for High Performance Data and AI in Healthcare and Life Sciences* to reflect the fact that it incorporates key building blocks for high-performance computing (HPC) and software-defined storage, and that it supports an expanding infrastructure of leading industry partners, platforms, and frameworks.

The reference architecture defines a highly flexible, scalable, and cost-effective platform for accessing, managing, storing, sharing, integrating, and analyzing big data, which can be deployed on-premises, in the cloud, or as a hybrid of the two. IT organizations can use the reference architecture as a high-level guide for overcoming data management challenges and processing bottlenecks that are frequently encountered in personalized healthcare initiatives, and in compute-intensive and data-intensive biomedical workloads.

This reference architecture also provides a framework and context for modern healthcare and life sciences institutions to adopt cutting-edge technologies, such as cognitive life sciences solutions, machine learning and deep learning, Spark for analytics, and cloud computing. To illustrate these points, this paper includes case studies describing how clients and IBM Business Partners alike used the reference architecture in the deployments of demanding infrastructures for precision medicine.

This publication targets technical professionals (consultants, technical support staff, IT Architects, and IT Specialists) who are responsible for providing life sciences solutions and support.

Authors

This paper was produced by a team of specialists from around the world working at the International Technical Support Organization (ITSO), Poughkeepsie Center.

Dino Quintero is an IT Management Consultant and an IBM Level 3 Senior Certified IT Specialist with the IBM Redbooks® team in Poughkeepsie, New York. Dino shares his technical computing passion and expertise by leading teams developing technical content in the areas of enterprise continuous availability, enterprise systems management, high-performance computing, cloud computing, artificial intelligence (including machine and deep learning), and cognitive solutions. He is also a Certified Open Group Distinguished IT Specialist. Dino holds a Master of Computing Information Systems degree and a Bachelor of Science degree in Computer Science from Marist College.

Frank N. Lee, PhD Dr. Frank Lee is the healthcare and life sciences industry leader for IBM Systems Group with over twenty years' experience in scientific research and information technology. His work includes the creation of industry reference architecture and its implementation as HPC, cloud, big data, and AI platforms for dozens of clients and IBM Business Partners worldwide. As an advocate for the transformation of the industry towards precision medicine, Frank has spoken in dozens of conferences and published in IBM Systems Journals, IBM Redbooks publications, research papers, HPCwire editorials, and HIMSS reports. When encountering gaps in technologies, Frank led charges of innovation with inventions in metadata and provenance management as a co-inventor of IBM Spectrum® Discover software and underlying technologies. Frank also provides subject matter expertise on genomics, an experience that includes participation in the Human Genome Project as a research associate and his training as a molecular biologist at Washington University.

Special acknowledgment to the following for their contributions to this project:

We thank the Sidra technical team for their leadership and contribution to the IBM-Sidra collaboration. We also thank Sidra Communications and the IBM Qatar team for producing and publishing reference materials, including press releases, videos, and media stories. We thank other clients and partners from Washington University Genome Center, Northwestern University Medical Center, and others, for their collaboration in and contribution to development for the reference architecture.

Thanks to the following people for their contributions to this project:

Wade Wallace

International Technical Support Organization, Poughkeepsie Center

Linda Cham, Ruzhu Chen, Jeff Hong, Denise Ruffner, David Wohlford, Joanna Wong,
Jane Yu

IBM US

Bill McMillan, Richard Wale

IBM UK

Jeff Karmioli, Gabor Samu

IBM Canada

Yael Shani

IBM Israel

Now you can become a published author, too

Here's an opportunity to spotlight your skills, grow your career, and become a published author—all at the same time. Join an ITSO residency project and help write a book in your area of expertise, while honing your experience using leading-edge technologies. Your efforts will help to increase product acceptance and customer satisfaction, as you expand your network of technical contacts and relationships. Residencies run from two to six weeks in length, and you can participate either in person or as a remote resident working from your home base.

Find out more about the residency program, browse the residency index, and apply online:

ibm.com/redbooks/residencies.html

Comments welcome

Your comments are important to us.

We want our papers to be as helpful as possible. Send us your comments about this paper or other IBM Redbooks publications in one of the following ways:

- Use the online **Contact us** review Redbooks form:

ibm.com/redbooks

- Send your comments in an email:

redbooks@us.ibm.com

- Mail your comments:

IBM Corporation, International Technical Support Organization
Dept. HYTD Mail Station P099
2455 South Road
Poughkeepsie, NY 12601-5400

Stay connected to IBM Redbooks

- Find us on Facebook:

<http://www.facebook.com/IBMRedbooks>

- Follow us on Twitter:

<http://twitter.com/ibmredbooks>

- Look for us on LinkedIn:

<http://www.linkedin.com/groups?home=&gid=2130806>

- Explore new Redbooks publications, residencies, and workshops with the IBM Redbooks weekly newsletter:

<https://www.redbooks.ibm.com/Redbooks.nsf/subscribe?OpenForm>

- Stay current on recent Redbooks publications with RSS Feeds:

<http://www.redbooks.ibm.com/rss.html>



Trends and challenges for precision medicine

Accelerating personalized healthcare and other biomedical workloads requires the adoption of cost-effective, high-performance infrastructure for big data analytics and artificial intelligence. Healthcare and life sciences organizations worldwide must manage, access, store, share, and analyze an explosive amount of data within the constraints of their IT budgets. IBM reference architecture for high-performance data analytics for healthcare and life sciences defines a platform for delivering the highest levels of performance for big data workloads at the same time lowering the total cost of ownership (TCO) for IT.

Advancements in high-throughput molecular profiling techniques and high-performance computing (HPC) systems ushered in a new era of personalized medicine. In personalized medicine, the treatment and prevention of disease can be tailored to the unique molecular profiles, behavioral characteristics, and environmental exposures of individual patients. Discovering treatment plans that are tailored to specific patient populations requires a clear understanding of the impact that such factors are likely to have on clinical outcomes.

Moreover, the task of delivering those plans in a time-sensitive clinical setting requires technical computing platforms that quickly and accurately classify individual patients into treatment cohorts most likely to achieve favorable outcomes in a timely manner. Such research and clinical tasks require healthcare and life science practitioners to access, process, and analyze various complex, information-rich data sources.

This chapter provides an overview of the IBM Reference Architecture for High Performance Data and AI in Healthcare and Life Sciences.

This chapter contains the following topics:

- ▶ New trend: The era of precision medicine
- ▶ Challenges

1.1 New trend: The era of precision medicine

Advancing the science of medicine by targeting a disease more precisely with treatment that is specific to each patient relies on access to that patient's genomics information and the ability to process massive amounts of genomics data quickly. The following trends are factors for this era of precision medicine:

- Data-driven

Advances in precision medicine, genomics, and imaging, along with widespread adoption of electronic health records and the proliferation of medical internet of things (IoT) and mobile devices, are resulting in an exponential growth of structured and unstructured data:

- By the end of 2020, 25% of data that is used in medical care will be collected and shared with healthcare systems by the patients themselves (“bring your own data”).¹
- Healthcare providers have fully embraced the IoT, with 72.7% of respondents having deployed an IoT solution and the remainder piloting or researching using IoT.²

- Pervasive artificial or augmented intelligence (AI)

To glean actionable insights from these large and complex data sets, healthcare organizations are investing in high-performance systems that support AI workloads. Most of the investment started in the research areas such as bioinformatics, computational chemistry, cellular imaging and natural language processing but started to spread into clinical areas such as medical imaging and informatics.

The AI workflow also requires or intersects with traditional workloads such as high performance and accelerated computing, machine learning, biostatistics, medical analytics and clinical informatics. This trend creates huge challenges on healthcare infrastructure and most institutions cannot keep up with the pace of change and complexities.

- Multicloud

Forward-thinking healthcare organizations are modernizing their infrastructure by deploying data-driven, multicloud storage and software-defined infrastructure because it ensures the highest level of data availability, reliability, and cost-efficiency. The benefits of storage and software-defined infrastructure accrue not only to IT, but to clinicians and researchers through the increased ability to access data and collaborate across the globe.

Multicloud storage architecture enables workloads to be deployed to the appropriate environment (private, public, dedicated, or hybrid cloud), increasing the speed to value and insights.

Private and public cloud is not an either or situation. There continues to be value in clients' existing landscapes, but more and more workloads which benefit from the value of private cloud. Think of cloud as a capability and not a location. The ability to deliver agility and composable services does not mean it has to be outside of the firewall; nor it is suggested that everything must be behind that firewall either.

There are applications that can be on public clouds and others that are better off on private clouds. A hybrid cloud configuration gives you the best value because it enables you to put the workloads where they make the most sense.

¹ Source: IDC 2018. From our 2018 FutureScape for Health

² Source: Global IoT Survey, IDC, September 2017

1.2 Challenges

For many biomedical research and clinical organizations, large-scale initiatives in personalized medicine are technically daunting for the reasons that are outlined in this section.

The following video provides insights about the challenges for managing the explosive growth in healthcare data that is discussed throughout this chapter:

<https://bit.ly/2IhSVq0>

1.2.1 Data management challenges

The four Vs that define big data also describe genomic data management challenges:

Volume	Very large data size and capacity.
Velocity	Demanding input/output (I/O) speed and throughput.
Variety	Fast-evolving data types and analytical methods.
Veracity®	The capability to share and explore a large volume of data with context and confidence.

In this case, the challenges are exacerbated by extra requirements, such as regulatory compliance (patient data privacy and protection), provenance management (full versioning and audit trail), and workflow orchestration.

Data volume

Genomic data volume is surging as the cost of sequencing drops precipitously. It is common for an academic medical research center that is equipped with next-generation sequencing (NGS) technologies to double its data storage capacity every 6 - 12 months. Consider a leading academic medical research center in New York City (NYC) that started 2013 with 300 TB of data storage. By the end of 2013, the data storage volume surpassed 1 PB (1000 TB), more than tripling the amount from 12 months before.

What made it even more astonishing is that the speed of growth has been accelerating and continues still today. For some of the world's leading genomic medicine projects, such as Genome England (UK), Genome Arabia (Qatar), Million Veteran Project (US), and China National GeneBank, the starting points or baselines for data volume are no longer measured in terabytes but tens and hundreds of petabytes.

Data velocity

Data velocity in a genomic platform can be demanding because of three divergent requirements:

- Very large files. These files are used to store genomic information from study subjects, which can be a single patient or a group of patients. There are two main types of such files: Genomic sequence alignment (Binary Alignment/Map (BAM)) and genetic variants (Variant Call File (VCF)).

These files are often larger than 1 TB and can take up half of the total storage volume for a typical genomic data repository. Additionally, these files are quickly growing larger, often the result of condensing more genomic information from higher resolution coverage (for example, from 30x to 100x for a full genome) or a larger study size.

As the genomic research evolves from Rare Variant study (variant calling from a single patient) to the Common Variant study, there is an emerging need to make joint variant calling from thousands of patient samples. Consider a hypothetical case that is provided by the Broad Institute: For 57,000 samples to be jointly called, the input BAM file will be 1.4 PB, and the output VCF file will be 2.35 TB, both extremely large in today's standards, but potentially a common standard soon.

- Many small files. These files are used to store raw or temporary genomic information, such as output from sequencers (for example, the BCL file format from Illumina). They are often smaller than 64 KB and can take up half of the total file objects for a typical genomic data repository.

Because each file I/O requires two operations for data and metadata, workload that generates or requires access to many files creates a distinct challenge that is different from that of large files. In this case, the velocity can be measured in I/O operations per second (IOPS), and typically reaches millions of IOPS for the underlying storage system.

Consider a storage infrastructure at a San Diego-based academic medical research center that was not optimized for massive small file operation. A workload, such as BCL conversion (for example, CASAVA, from Illumina), stalls because compute servers are constrained by limited I/O capability, especially IOPS.

A benchmark has confirmed that the central processing unit (CPU) efficiency drops to single digits because the computing power is being wasted waiting for data to be served. To alleviate this computational bottleneck, IBM researchers developed a data caching method and tool to move I/O operation from disk into memory.

- Parallel and workflow operation. To scale performance and speed time to results, genomics computing is often run as an orchestrated workflow in batch mode. This parallel operation is essential to deliver fast turnaround as more workloads evolve from small-scale targeted sequencing to large-scale full-genome sequencing. With hundreds to thousands of diverse workloads running concurrently in such a parallel computational environment, the requirement for storage velocity as measured in I/O bandwidth and IOPS is aggregated and climbs exponentially.

Consider a bioinformatics application from the NYC academic medical research center. This application can be run in parallel on 2,500 compute cores, each writing the output to the disk at a rate of 1 file per second. Collectively this app creates millions of data objects, either 2,500 folders each with 2,500 files or 14 million files in one directory. This workload is one of many that contributes to a data repository with 600 million objects, including 9 million directories that each contain only one file.

Because of the massive amount of metadata, the IOPS load was constraining the overall performance so much that even a file system command to list files (ls in Linux) took several minutes to complete. A parallel application, such as the Broad Institute Genome Analysis Toolkit (GATK) Queue, also suffered from poor performance.

In early 2014, the file system was overhauled with emphasis on improving the metadata infrastructure. As a result, both bandwidth and IOPS performance were improved, and the benchmark showed a 10x speedup of a gene-disease app without any application tuning.

Data variety

There are also many types of data formats to be handled in terms of storage and access. The data formats range from intermediary files that are created during a multistep workflow, to output files that contain vital genomic information, to reference data sets that must be carefully versioned. The common approach today is to store all of this data in online or nearline disks in one storage tier, despite the expense of this approach.

One practical constraint is the lack of lifecycle management capability for the massive amount of data. If it takes the genomic data repository a long time to scan the file system for candidate files for migration or backup, it becomes impossible to complete this task in a timely fashion.

Consider a large US genome center that is struggling to manage its fast-growing data as it adopts an Illumina X10 sequencer for full genome sequencing. To complete a scan of the entire file system, it takes up to four days, making daily or even longer backup windows impossible. As a result, data is piling up quickly in the single-tier storage and slowing down the metadata scan and performance even further, which causes a vicious cycle for data management.

Another emerging challenge for data management is created by the spatial varieties of data. As inter-institutional collaboration becomes more common and a large amount of data must be shared or federated, the locality becomes an indispensable character of data.

The same data set, especially reference data or output data, can exist in multiple copies in different locations. Alternatively, the same data set can exist in duplicates in the same location because of regulatory compliance requirements (for example, physically isolating a clinical sequencing platform from one for research). In this case, managing the metadata efficiently to reduce data movement or copying reduces the cost due to extra storage, and minimizes problems due to versioning and synchronization.

Data veracity

The multi-factorial nature of many complex disorders, such as diabetes, obesity, heart disease, Alzheimer's, and Autism Spectrum Disorder (ASD), requires sophisticated computational capabilities. You use these capabilities to aggregate and analyze large streams of data (genomics, proteomics, and imaging) and observation points (clinical, behavioral, environmental, and actual evidence) from a wide range of sources.

The development of databases and file repositories that are interconnected based on global data sharing and federated networks bring the promise of innovative and smarter approaches to access and analyze data in unprecedented scale and dimensions. It is in this context that the veracity (trustworthiness) of data enters the equation as an essential element.

For example, clinical data (genomics and imaging) must be properly and completely de-identified to protect the confidentiality of the study subject. Genomic data must have end-to-end provenance tracking from bench to bedside, to provide a full audit trail and reproducibility. The authorship and ownership of data must be properly represented in a multi-tenancy and collaborative infrastructure. With the built-in capability to handle data veracity, a genomic computing infrastructure must enable the researchers and data scientists to share and explore a large volume of data with context and confidence.

1.2.2 Other data challenges

This section describes further technical data management challenges.

Data Gravity

Aside from regulatory or compliance reasons that can dictate where data must be located, there are three factors that you need to consider in terms of making the public, private, or hybrid cloud decision:

1. Duration of use and the cost effectiveness of each is the first factor. The more you use something, the more cost effective it is to own it rather than rent it.

2. The less stable or predictable the workload is, the more likely you want a dynamic environment in which you can add any amount of capacity, and you can benefit from not having to pay for that capacity when it is not being used.
3. Finally, data gravity is likely the biggest determination of where applications can be located. The further the applications are from the data, the greater the latency and the lower the throughput. This can dictate application collocation with data.

Big data

Biomedical data that is required to support personalized medicine initiatives is large, varied, and often unstructured. In addition, the amount of data being collected is growing exponentially, and it must be archived for extended periods, and sometimes for decades.

Personalized medicine applications can include the following common data sources:

- ▶ Whole genome sequences
- ▶ Biomedical imaging from clinical and laboratory instrumentation
- ▶ Electronic medical records systems
- ▶ Physiological sensors or wearables
- ▶ Collections of curated scientific and clinical literature

Technical computing systems must be able to process and store this data at low cost as volumes continue to grow.

Data silos

Personalized healthcare requires the aggregation of information that can provide a full view of the biological traits, behaviors, and environmental exposures of each patient. However, data for a single patient is usually captured and scattered across heterogeneous storage silos within health systems and biomedical research institutions, and they must be integrated into a common database before analysis can begin.

Compute and data intensive workloads

Analytics workloads can be compute and data intensive. Common examples include I/O-intensive analysis pipelines that transform raw, next-generation sequence data into the following output types:

- ▶ Genomic variant files
- ▶ Deep learning techniques for discovering patterns within complex biomedical data sets
- ▶ Large-scale data mining of clinical and scientific documents

Such workloads might take hours, or even days, to complete on existing technical computing platforms.

Evolving applications and frameworks

Personalized healthcare must often support hundreds of different applications at any given time, including those that are related to medical informatics, genomics, image analysis, and deep learning.

Such applications are often built on frameworks and databases that are continually evolving, including Spark, Hadoop, TensorFlow, Caffe, Docker, MongoDB, and HBase. Biomedical research organizations often have difficulty supporting multiple versions of these application frameworks and databases, which are proliferating and frequently evolving, sometimes two or more times per year.

Collaboration

Data sharing across institutions and across geographic boundaries is a growing necessity in the study of rare diseases and complex disease mechanisms. International scientific consortia consisting of academic, commercial, non-profit, and government entities are rapidly emerging and sharing biomedical data and related analysis. Collaborating partners often lack solutions for sharing their data sets rapidly and cost-effectively without compromising protected health information and intellectual property rights.

Many organizations worldwide are finding it difficult to overcome these challenges, especially within the constraints of their IT budgets. Clinical and scientific research data must be accessed, stored, analyzed, shared, and archived in a time-efficient and cost-effective manner.

However, for many healthcare organizations, biomedical research institutions, and pharmaceutical companies today, data is collected in very large volumes. These organizations can no longer process, properly store, or transmit this amount of data over regular communication lines in a timely manner.

For many organizations, compute and storage silos are proliferating across clinical and research groups as analysts collect increasing volumes of data and use that data in complex analytical workloads. To move data across long physical distances, organizations often resort to disk drive and shipping companies to transfer raw data to external computing centers for processing and storage, thus hindering speedy access and data analysis.

To overcome the technical challenges that industry practitioners face in the era of personalized healthcare, IBM Systems has created a reference architecture for healthcare and life sciences. This reference architecture, which is built on the IBM history of delivering preferred practices in HPC, can make it possible for healthcare and life science organizations to easily scale compute and storage resources as demand grows.

In addition, this architecture enables organizations to support the wide range of development frameworks and applications that are required for industry innovation, all without unnecessary reinvestments in technology.

Workload management in genomics

Genomics workloads can be complex. There are a growing number of genomic applications with varying degrees of maturity and types of programming models. Many are single-threaded (for example, R) or perfectly parallel (for example, Burrows-Wheeler Aligner (BWA)), and a few others are multi-threaded or MPI-enabled (MPI BLAST). However, all applications must work together in a high-throughput and high-performance mode to generate final results.



The journey of the reference architecture

This chapter describes the journey of the reference architecture for genomics.

This chapter contains the following topics:

- ▶ The history of IBM Reference Architecture
- ▶ Overview of IBM Reference Architecture for High Performance Data and AI
- ▶ Datahub for High-Performance Data Analytics
- ▶ Orchestrator of High-Performance Data Analytics

2.1 The history of IBM Reference Architecture

Key stakeholders in personalized medicine initiatives depend on a reliable and flexible platform that meets their diverse data and analytics needs:

- ▶ Clinical researchers looking for the actionable biomarkers
- ▶ Scientists in pharma research and development organizations progressing potential drugs through clinical trials
- ▶ Physicians in hospitals delivering precision medicine treatments that give patients the best outcomes

IBM Reference Architecture for High Performance Data and AI in Healthcare and Life Sciences (formerly IBM Reference Architecture for Genomics) is an end-to-end architecture that is designed to address the common requirements of organizations pursuing genomics and personalized medicine initiatives in biomedical research.

2.1.1 First-generation reference architecture

In 2013, IBM designed the first enterprise architecture for high-performance computing (HPC) in support of the emerging genomics initiative at MD Anderson Cancer Center. The architecture addressed the critical requirement to provide a high-performance and scalable platform for computational chemistry, bioinformatics, imaging, and other analytical workloads.

Because there were multiple systems involved, we needed to make the storage infrastructure *globally and continuously available* to serve all of the diverse data needs currently and for years to come. Along with the mission-critical need to orchestrate a large compute environment and complicated workloads, we captured and addressed these requirements into a blueprint that we later called [Reference Architecture for Genomics](#) in an IBM Redpaper published in 2015.

In this first-generation architecture for genomics, IBM designed the Datahub as an abstraction layer for handling demanding genomics requirements, such as high-throughput data landing, information lifecycle management, and global name space regardless of sharing protocol. These requirements can sometimes be met easily on a single workstation or small cluster, but the capability to handle hundreds of servers and petabytes of data is what made Datahub unique and essential.

What made Datahub even more valuable was its intrinsic scalability to start small (or big) and grow and scale out rapidly based on the workloads. This turned out to be the most important aspect of precision medicine and genomics workloads: as the next-generation sequencing technologies were rapidly advancing, the data and workloads can grow at a rate of 100% every six months.

Datahub fulfills these requirements through software in concert with storage building blocks (flash, disk, and tape library). At this point, we were already able to take advantage of best practices from other institutions (for example, Mount Sinai) that used tiered storage building blocks to land and then move data based on its attributes, such as size and “temperature” (time since last access or creation).

We also designed the Orchestrator as the second abstraction layer for handling application requirements, and mapped it towards the computing building blocks. It has specified functions, such as parallel computing and workflow automation, which can be fulfilled by software in concert with other computing resources, such as an HPC cluster or virtual machine farms.

This software-defined blueprint was essential to future-proof the infrastructure and sustain the usability of applications. This blueprint is designed so that the hardware building blocks can be expanded or replaced without impacting the operation of the system, the running of the application, and ultimately the user experience.

Between 2013 and 2014, we also tested the reference architecture of genomics at Sidra Medicine in Qatar as part of the Qatar Genome Project. This was when we discovered more strategic values of a referenceable architecture. As illustrated in the case study with Sidra later, the HPC infrastructure started small but grew rapidly as the Qatar Genome project got funded and took off on scale. The architecture served as a roadmap that guided every phase of organic growth and infrastructure expansion, at the same time maintaining the integrity and sustainability of the applications, analytics, workflow and databases.

The architecture also ensured interoperability of the applications. When Sidra decided to collaborate with MD Anderson on a research project/publication, they shared a genomic pipeline based on the Orchestrator and underlying software tools. Given the common architecture patterns and building blocks, the two institutions were able to quickly rebuild the workflow and co-publish the results of genomic analysis using GATK best practice pipeline.

2.1.2 Second-generation reference architecture

Since the launch of the reference architecture for genomics, IBM continued to evolve and optimize the reference architecture to reflect the changing needs and challenges for the healthcare and life sciences industries. We harnessed lessons learned and best practices from our clients, industry collaborators, and IBM Business Partners worldwide. Many of these clients were early adopters of precision medicine, battling in the frontiers of medicine, biotechnology, and pharma to discover and advance care for cancers and rare diseases.

We also conducted survey among users from various industry groups and consortiums. In the survey we inquired about the most challenging aspects of adopting and using infrastructure and informatics tools in support of their precision medicine undertakings. We also asked about the fields of research or clinical practice these users represent.

From the results of the survey we can see that although many researchers are in the fields of -omics (genomics, metabolomics, proteomics, and so on), there is a trend of increasing representation from non-omics fields such as medical imaging, clinical analytics, biostatistics, and even real-world evidence (RWE) and internet of things (IoT).

The need to handle and analyze explosive amount of data and collaborate across fields started to bring more and more users to work together and share their experiences. We are fortunate to witness and document these challenges and needs from people who are in the frontiers of precision medicine.

So what are some of their top challenges?

The reference architecture must meet the following requirements:

- High-performance. Users want fast and faster time-to-results. The results can be a genomics analysis of clinical variants for patients, or an AI model developed for diagnosing Alzheimer's Disease, or new biomarkers for cancer. In all of these cases, traditional desktop computing can no longer keep up with the workloads or the data storage. Often, users have to wait for days or even weeks for data to be transferred and loaded, then an even longer time for processing and analysis.

- **Low-cost.** Many of the users we surveyed emphasized the need to start small in terms of the initial platform, and to be able to control the cost during an expansion. Because most projects were funded by external grants, or are internal projects subject to annual assessment, having low or controlled costs can become a show stopper for many projects.

Because the data needs to be stored and accessed almost forever, this requirement adds a baseline cost that is guaranteed to grow over time. In addition, the location of the data often determines the computational needs, so the choice of data storage can be more critical than most users anticipate at the beginning of the projects.

- **Easy-to-use.** Some users reflected in the survey that they learned about the modern big data and AI toolings, such as Spark, Hadoop, and Tensorflow, but most of these tools require a rather steep learning curve or high threshold of adoption. Users want to spend more time analyzing and thinking about the research problem than learning to program, debug a file system, or operate a complex IT platform.
- **Collaborative.** The users desire collaboration across institutional or geographic boundaries without compromising protected healthcare and private, identifiable information. As parallel discovery based on an ever-growing data set is becoming a norm, more users are experimenting with the public cloud. However, not all of their data or metadata can be shared outside their institutions.

Even within the same institution or team, collaboration for sharing resources or data was lacking. As an example, one customer told us that they used a Slack channel to share the usage of a GPU server for deep learning workloads. “Wouldn’t it be great if a team of data scientists could share an ‘unlimited’ pool of resources for compute and data?” they asked during the review of the survey.

With this user feedback and field input in mind, we made a large overhaul of the reference architecture in the last two years in three main areas:

- Data-driven
- Cloud-ready
- AI-capable

We also renamed the reference architecture to *High Performance Data and AI* to reflect these changes and focuses.

2.2 Overview of IBM Reference Architecture for High Performance Data and AI

Scaling high-performance technical platforms to support growing data volumes and diverse applications, and at the same time continuing to accelerate workloads and minimizing IT costs, requires a flexible yet tightly coordinated framework for data access, computing, and storage. This section describes how the IBM Reference Architecture for High Performance Data and AI helps solve the challenges by providing a solution to benefit users and IT providers.

2.2.1 Challenges

With the arrival of precision medicine and clinical genomics, biomedical research institutes and healthcare providers, such as hospitals, cancer centers, genome centers, pharma R & D and biotech companies are dealing with enormous data growth. This data, mainly unstructured, is flowing at a rate of terabytes per day, or even per hour, from fast-growing sources of instruments, devices, and digital platforms.

This data needs to be captured, labeled, cleaned, stored, managed, analyzed, and archived. The disparate file types generated by different research tools and environments create silos that impede data access, drive down efficiency, drive up costs, and slow times to insight.

The volume and complexity of data also drives the adoption of modern analytical frameworks, such as big data (Hadoop and Spark) and AI (machine learning and deep learning), for thousands of research and business applications (for example, genomics, bioinformatics, imaging, translational, and clinical). The collaborative nature of the biomedical research also facilitates global data sharing in a multicloud environment.

Researchers, clinicians and data scientists struggle to deal with the “ocean” of data and application juggling. Because of this struggle, it is imperative for the infrastructure and underlying IT architecture to be transformed and become agile, data-driven, and application-optimized. The architecture must become data and application ready.

2.2.2 The solution

The IBM High Performance Data and AI (HPDA) architecture for healthcare and life sciences, shown in Figure 2-1 on page 14, is built on the IBM history of delivering best practices in HPC. In fact, the basic HPDA framework and building blocks were used to construct Summit and Sierra, which are currently two of the world’s most powerful supercomputers for data and AI.

The architecture is designed to help healthcare and life science organizations easily scale and expand compute and storage resources independently as demand grows, to ensure maximum performance and business continuity. It supports the wide range of development frameworks and applications required for industry innovation with optimized hardware as a foundation, without unnecessary reinvestments in technology.

The HPDA is based on software-defined infrastructure solutions that offer advanced policy-driven data and resource management capabilities. It has two key layers, one each for managing storage resources and compute resources. Currently, it can support major computing paradigms, such as traditional HPC, data analytics, cognitive computing, machine learning, and deep learning.

These capabilities then become the infrastructure and informatics foundation for developing and deploying applications for various fields, such as genomics, imaging, clinical, real-world evidence (RWE), and IoT. The HPDA architecture can be implemented on-premise in a local data center, or off-premise in a private or public cloud. We have also demonstrated advanced use cases and platforms that can be deployed as a hybrid cloud.

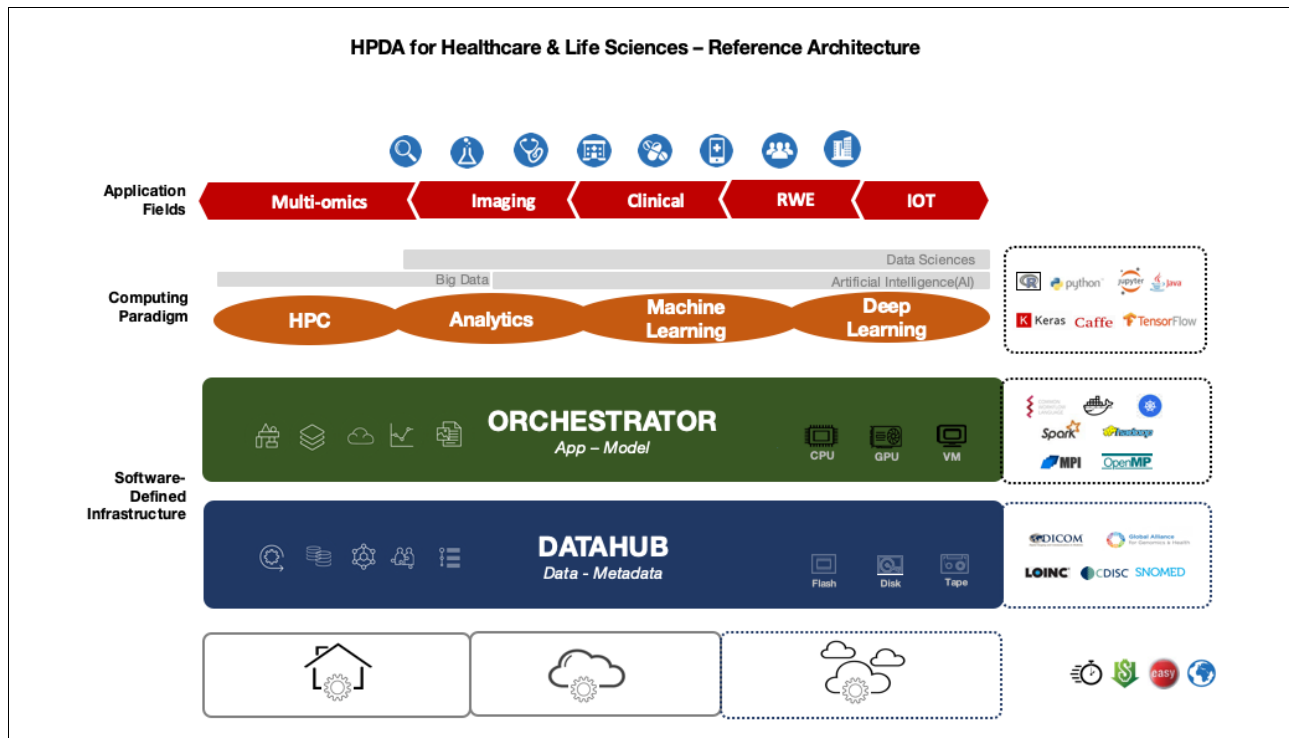


Figure 2-1 IBM Reference Architecture for High-Performance Data Analytics and AI

Datahub

The HPDA has a *Datahub* layer designed to manage the deluge of unstructured data that is pouring into and siloed in disparate systems and locations. We define five key functions in handling the full lifecycle of data and metadata:

- ▶ High-performance ingesting
- ▶ Policy-based auto tiering
- ▶ Multiprotocol sharing
- ▶ Active-active peering
- ▶ Metadata cataloging

These five mission-critical functions anchor infrastructural capabilities for data to be captured rapidly, stored safely, accessed transparently, and shared globally, *wherever and whenever*.

The data ingest function is the most basic yet important one: a large amount of raw genomic, imaging, and sensor data needs to be quickly ingested into the infrastructure from the various data sources. These sources include genomic sequencers, high content screening scanners, microscopes, and IoT devices.

One essential requirement for high-performance data ingestion is the ability to load data in parallel, so that a large file can be split into many blocks and written into the target storage device using a parallel file system. The file system must also be able to handle many thousands or even millions of files concurrently. IBM Spectrum Scale is one such file system that meets the requirement for high performance data ingestion.

As a reference architecture, the HPDA Datahub can be implemented as a software-defined storage infrastructure on-premise, in a private cloud, or in a public cloud. The infrastructure building blocks include low-latency Flash/NVME devices, large-capacity and high-performance disk/file system appliances (for example, IBM Elastic Storage® Server), low-cost tape libraries, and cloud object stores.

Datahub can also be deployed as a software-only solution using customers' existing public cloud infrastructure.

Based on the requirements for capacity, performance, and projected future growth, an HPDA Datahub-based infrastructure can be designed with various building blocks of different sizes and price-performance profiles. The management software (for example, IBM Spectrum Scale, IBM Spectrum Discover, and Cloud Object Store) then works together to connect storage hardware into a global and extensible name space for data services.

Orchestrator

The HPDA has an *Orchestrator* layer designed to manage myriads of applications and workloads, ranging from a high-throughput genomics (DNAseq, RNAseq) pipeline running on a large cluster to a medical imaging deep learning training job running on a multi-GPU system. There are thousands of applications, tools, frameworks, and workflows that are available for use, with many of them still being actively developed.

The updated versions of the applications often have newer requirements and dependencies for infrastructure (operating system, drivers, libraries, configuration, and so on) that can conflict with older versions or other packages on the system. Because some next-gen applications are developed as *Cloud-ready*, they can take use modern technologies, such as containers, to gain mobility and elasticity. This technology use makes it necessary to orchestrate and manage containers so that they can share resources among themselves and with non-containerized workloads.

Applications can and do consume computational resources in many types of ways. From highly parallel applications hosted across thousands of nodes, to genomics applications running as single-threaded processes with large memory requirements. To handle the diversity of the workloads and deliver a consistently high-performance computational infrastructure, the first function of Orchestrator is to manage the infrastructure building blocks and turn them into usable resources by way of policy-based allocation and job scheduling.

The system administrator can set the policies that prioritize the placement and execution of workloads, while users can submit jobs to the scheduler through scripts or the graphical user interface.

The agility, elasticity, and flexibility of the compute infrastructure can be accomplished through various functions, such as building platform as a service and cloud computing. The workload isolation into containers, automation by pipelining, and sharing through the catalog makes efficient use of the resources possible.

2.2.3 Key values

The Data Hub and Orchestrator were designed as two separate abstraction layers that can work together to manage data and orchestrate workloads on any supporting storage and computational building blocks. The resulting infrastructure is a true data-driven, cloud-ready, AI-capable platform that is able to handle both complex data at scale and the most demanding analytics and AI workloads.

All of the applications and use cases developed for the architecture are based on deep industry experience, collaboration, and feedback from leading organizations that are at the forefront of precision medicine.

Users and infrastructure providers are achieving valuable results and significant benefits from the HPDA solution, as described in the following sections.

Key values for users

The HPDA solution provides users the following benefits:

- ▶ **Ease-of-use.** It provides a self-service App Center with a GUI based on an advanced catalog and search engines for users that enables them to manage the data easily, in real-time, with maximum flexibility.
- ▶ **High-performance.** Cloud-scale data management and multicloud workload orchestration locate data where it makes sense, and provision the required environment for peak demand periods in the cloud, dynamically and automatically, only for as long as needed, to maximize performance.
- ▶ **Low cost.** It includes policy-based data management that can reduce storage costs up to 90% by automatically migrating file and object medical data to the optimal storage tier, based on data value and performance criteria.
- ▶ **Global collaboration.** This functionality facilitates multi-tenant access and data sharing that spans across storage systems and geographic locations, enabling many research initiatives around the globe that use a common reference architecture to establish strategic partnerships and collaboration.

Key values for IT providers

The reference architecture offers the following values to IT providers:

- ▶ **Easy to install.** A blueprint compiles best practices and enables IT architects to quickly deploy an end-to-end solution architecture, which is designed and tuned specifically to match different use cases and requirements from different business and research disciplines.
- ▶ **Fully tested.** IT architecture based on a solid roadmap of future-ready proven infrastructure can easily be integrated into the existing environment protecting already made investments, especially the hardware purchase and cloud services.
- ▶ **Global Industry Ecosystem.** A wide ecosystem aligns with the latest technologies for hybrid multicloud, big data analytics, and AI to optimize data for cost, compliance, and performance that is needed and expected by end users for better services and patient care.

Note: The reference architecture is based on collaborative work with leading institutions, such as UPMC and Sidra. IBM is expanding it to new domains of imaging, clinical work, and AI. This [YouTube video](#) illustrates the challenges and their solution.

2.3 Datahub for High-Performance Data Analytics

Data management is the most fundamental capability for genomics platforms, because a huge amount of data must be processed at the correct time and place at a feasible cost. The temporal factors can range from hours (analyzing data in an HPC system) to years (when data must be recalled from a storage archive for reanalysis). The spatial aspect can span a local infrastructure that provides nearline storage capability to a cloud-based, remote cold archive.

2.3.1 Datahub functions

To address the challenges of data management for genomics, define an enterprise capability that functions as a scalable and extensible layer for serving data and metadata to all workloads. Name this layer *Datahub* to reflect its critical role as a central hub for data: storing, moving, sharing, and indexing a massive amount of genomic raw and processed data. The Datahub also manages the underlying heterogeneous storage infrastructure from solid-state disk (SSD) and flash storage to disk to tape to cloud (Figure 2-2).

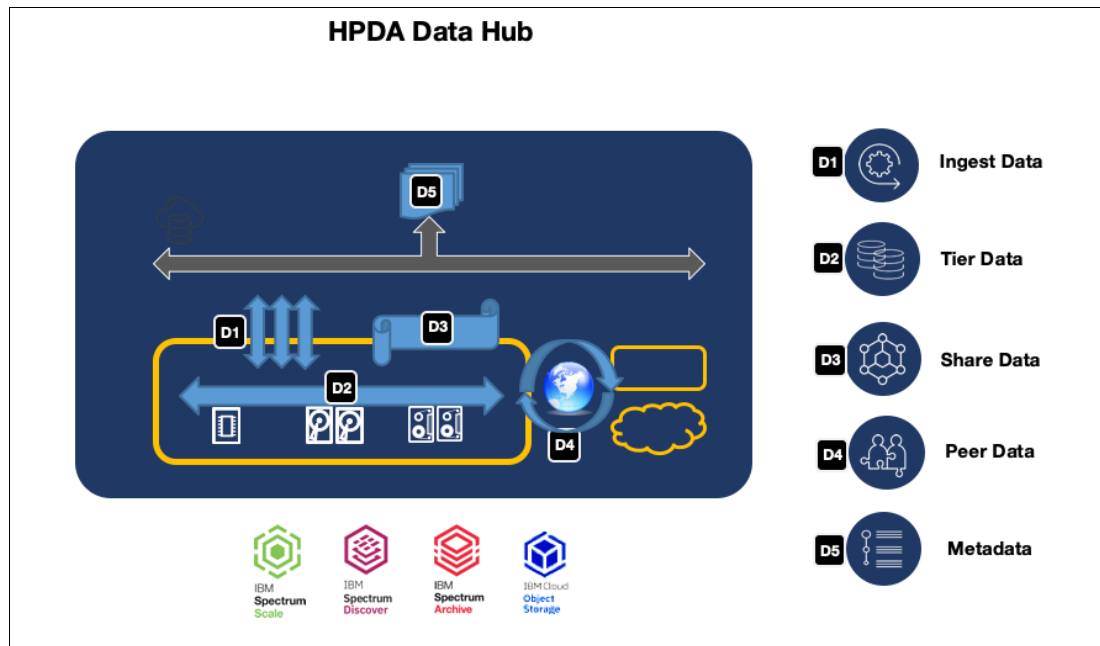


Figure 2-2 Overview of Datahub

The Datahub is the enterprise capability for serving data and metadata to all of the workloads (Figure 2-2). It defines a scalable and extensible layer that virtualizes and globalizes all storage resources under a global name space. Datahub is designed to provide the following key functions:

- ▶ High-performance data I/O
- ▶ Policy-driven information lifecycle management (ILM)
- ▶ Efficient data sharing through caching and necessary replication
- ▶ Large-scale metadata management

For physical deployment, the Datahub must support an increasing number of storage technologies as modular building blocks, including the following components:

- ▶ SSD and flash storage
- ▶ High-performance fast disks
- ▶ Large-capacity slow disks (4 TB per drive)
- ▶ High-density and low-cost tape library
- ▶ External storage cache that can be locally or globally distributed
- ▶ Big data storage that is based on Hadoop
- ▶ Cloud-based external storage

The following key functions are mapped to the Datahub:

- I/O management. This Datahub function addresses the need for large and scalable I/O capability. There are two dimensions to the capability: I/O bandwidth for serving large-size files, such as Binary Alignment/Map (BAM), and I/O operations per second (IOPS) for serving many small files, such as BCL and FASTQ. Because of these divergent needs, the traditional one-size-fit-all architecture struggles to deliver performance and scalability.

Datahub I/O management solves this challenge by introducing the pooling concept to separate the I/O operations for metadata and small files from those for the large files. These storage pools, although mapped to different underlying hardware to deliver optimal storage performance, are still unified at the file system level to provide a single global name space for all of the data and metadata, and are transparent to users.

- Lifecycle management. This Datahub function addresses the need for managing the lifecycle of data from creation to deletion or preservation. We use the analogy of *temperature* to describe the stages and timeliness in which data needs to be captured, processed, moved, and archived. When raw data is captured from some instruments, such as high-throughput sequencers, that data is the *hottest* in temperature and must be processed by an HPC cluster with robust I/O performance (so-called *scratch* storage).

After initial processing, the raw and processed data becomes *warm* in temperature, because it takes a policy-based process to determine the final outcome: deletion, preservation in a long-term storage pool, or being archived. The process accounts for file type, size, usage (for example, the last time that it was accessed by a user), and system utilization information.

Any files that meet the requirement for action are either deleted or migrated from one storage pool to another, typically one that has a larger capacity, slower performance, and much lower cost. One such target tier can be a *tape library*. Coupled with storage pooling and low-cost media, such as tape, this function enables the efficient use of underlying storage hardware and drastically lowers the total cost of ownership (TCO) for a Datahub-based solution.

- Sharing management. This Datahub function addresses the need for data sharing within and across logical domains of a storage infrastructure. As genomic sample and reference data sets grow larger (in some cases exceeding 1 PB per workload), it is increasingly difficult to move and duplicate data for sharing and collaboration purposes.

To minimize the impact of data duplication and at the same time enable data sharing, Datahub introduces the following elements under sharing management:

- Storage multi-clustering. One compute cluster can access a remote system directly, and pull data/storage only on demand.
- Cloud data caching. The metadata index and full data set for a specific data repository (host) can be selectively and asynchronously cached on a remote (client) system for fast local access.
- Database federation. This enables secured federation among distributed databases.

In all of these functions, the data sharing and movement can occur over a private high-performance network, or over a wide area network (such as the internet). The technology accounts for the security and fault tolerance.

- Metadata management. This Datahub function provides a foundation for I/O, lifecycle, and sharing management. The ability to store, manage, and analyze billions of metadata objects is necessary for any data repository that scales beyond petabytes of size, which is increasingly the case for genomics infrastructures. The metadata includes system metadata, such as file name, path, size, pool name, creation time, and modified or access time. It can also include custom metadata in the form of key value pairs that applications, workflow, or users can create to associate with the files of interest.

This metadata must be efficiently used to accomplish the following goals:

- Facilitate the I/O management by placing and moving files based on file size, type, or usage.
- Enable policy-based lifecycle management of data based on information that is collected from lightning-fast metadata scans.
- Enable data-caching so that the distribution of metadata can be lightweight and depend less on networking.

2.3.2 Datahub solution and use cases

Datahub is based on IBM Spectrum Scale (formerly IBM GPFS) which provides a file system with high-performance, scalability, and extensibility characteristics.

Initially developed and optimized as an HPC parallel file system, IBM Spectrum Scale manages to serve large volumes of data at a high bandwidth and in parallel to all the compute nodes in the computing system. Because genomic pipelines can consist of hundreds of applications that are engaged in concurrent data processing on many files, this capability is critical in feeding data to the computational genomics workflow.

The traditional genomics pipeline generates a huge amount of metadata and data. The file system, which is a system pool that is built on SSD and flash storage with high-IOPS capability, can be dedicated to store metadata for files and directories, and in some cases small-size files directly. This situation drastically improves file system performance and responsiveness to metadata-heavy operations, such as listing all files in a directory.

As a file system with a connector to MapReduce, Datahub can also serve MapReduce and big data jobs on the same set as compute nodes, eliminating the need for (and complexity of) Hadoop Distributed File System (HDFS).

The policy-based data lifecycle management capability enables Datahub to move data from one storage pool to others, maximizing I/O performance and storage utilization, and minimizing operational cost. These storage pools can range from high-I/O flash storage or a high-capacity storage appliance, to low-cost tape media through integration with a tape management solution, such as IBM Spectrum Protect (formerly IBM Tivoli® Storage Manager) and IBM Spectrum Archive (formerly IBM Linear Tape File System).

The increasingly distributed nature of genomics infrastructure also requires data management on a larger and global scale. Data must be moved or shared across different sites, and its movement or sharing must be coordinated with computational workload and workflow.

To achieve this goal, Datahub relies on a sharing function that is based on the Active File Management (AFM) feature of IBM Spectrum Scale. AFM enables the Datahub to extend the global name space to multiple sites, enabling them to share a common metadata catalog and a cache copy of home data for a remote client site to access locally.

For example, a genomic center can own, operate, and control the versions of all the reference databases or data sets during the time the affiliated or partner sites or centers can access the reference data set through this sharing function. When the centralized copy of a database is updated, so are the cache copies of the other sites.

With Datahub, a system-wide metadata engine can also be built to index and search all the genomic and clinical data, enabling powerful downstream analytics and translational research.

2.4 Orchestrator of High-Performance Data Analytics

This section describes the challenges in workload management with genomics, and using *Orchestrator* to help minimize workload management challenges.

2.4.1 Orchestrator functions

Through the Orchestrator, the IBM Reference Architecture for Genomics defines the capability to orchestrate resources, workload, and workflow, as shown in Figure 2-3. This functionality, a unique combination of the workload manager and workflow engine, links and coordinates a spectrum of computational and analytical jobs into fully automated pipelines that can be easily built, customized, shared, and run on a shared platform.

The workload manager allows and enables multiple platforms. This setup provides the necessary abstraction of applications from the underlying infrastructure, such as an HPC cluster with a graphical processor unit (GPU) or a big data cluster in the cloud.

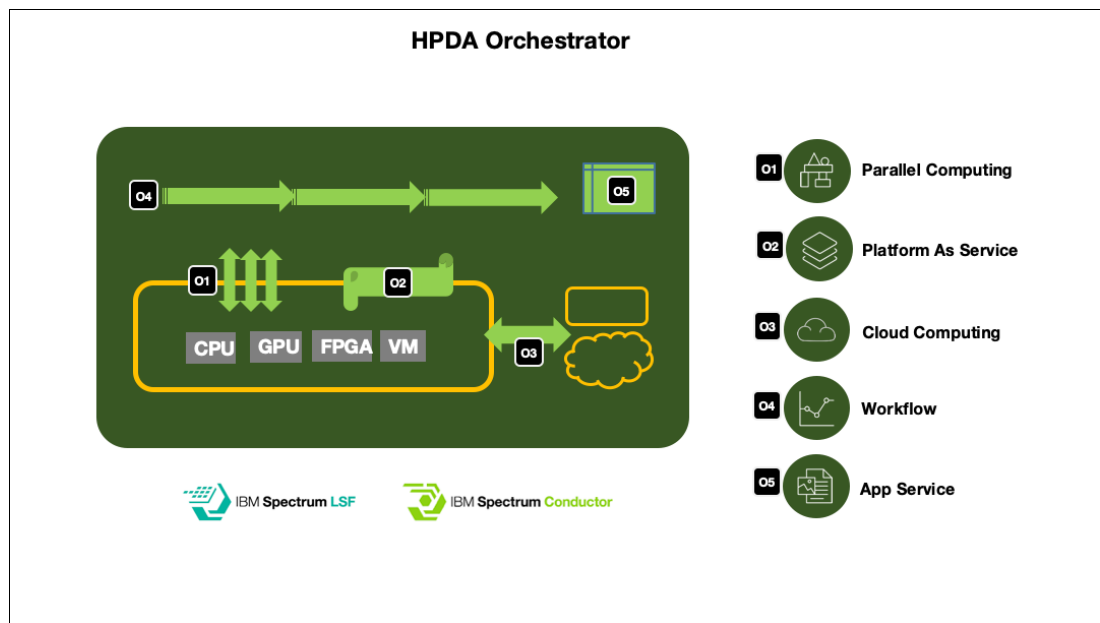


Figure 2-3 Orchestrator overview

The Orchestrator is the enterprise capability for orchestrating resources, workloads, and managing provenance, as shown in Figure 2-3. It is designed to provide the following key functions:

- ▶ Resource management, by allocating infrastructure to computational requirements dynamically and elastically.
- ▶ Workload management, by efficient placement of jobs onto various computational resources, such as local or remote clusters.
- ▶ Workflow management, by linking applications into logical and automated pipelines.
- ▶ Provenance management, by recording and saving metadata that is associated with the workload and workflow.

By mapping and distributing workloads to elastic and heterogeneous resources (HPC, Hadoop, Spark, Openstack/Docker, and cloud) based on workflow logic and application requirements (for example, architecture, CPU, memory, or I/O), the Orchestrator defines the abstraction layer between the diverse computing infrastructure and the fast-growing array of genomic workloads.

The IBM reference architecture reflects work that is continually underway within IBM Systems to integrate elements of our compute and storage platforms. This integration delivers the highest levels of performance for big data, while also lowering the total cost of IT ownership.

Many healthcare organizations are optimizing their systems with a high-performance analytics framework to improve business results, as shown in this [High-Performance Data Architecture for Healthcare](#) YouTube video.

2.4.2 Orchestrator solution and use cases

The IBM Spectrum LSF® family of products is one such solution to the requirements of the Orchestrator. It is used in business-critical environments ranging from semiconductor design to aircraft design to financial services.

Genomics pipelines consist of hundreds of applications generating tens of thousands to millions of individual jobs. Being able to scale reliably to these volumes is critical to delivering a production-quality, personalized medicine solution. Telling a patient that their tests are delayed because the IT can't handle the load is just not acceptable.

A powerful policy engine enables arbitration between different users and applications, with time-critical patient analyses being automatically prioritized over research tasks. However, fine-grained role-based access control (RBAC) ensures that users only see what they are allowed to see, protecting confidentiality and intellectual property.

IBM Spectrum LSF also supports powerful multicloud capabilities, not only enabling bursting or flexing if additional capacity is required, but supporting data affinity. Therefore, particular analyses and pipelines can be run where the data resides, rather than incurring the resource use required by replicating remote data on-site.

IBM Spectrum LSF Application Manager and IBM Spectrum LSF Process Manager enable custom workload portals and workflows to be readily built. These portals and workflows enable the healthcare professional to learn new tools faster and with fewer errors. More importantly, the professionals can keep their focus on healthcare and not IT.

IBM Spectrum Conductor®, IBM Deep Learning Impact, and IBM Watson® Machine Learning Accelerator for Enterprise AI can all use IBM Spectrum LSF for orchestration, allowing a broad range of analyses to be supported on a single orchestration platform.

For more information about how to get your data ready for precision medicine and experience new records for speed and scale through tips that are based on real-world use cases of high-performance genomics and imaging, see *The smart tips guide to high performance data and AI architecture* at the following website:

<https://ibm.co/2r3trFc>



Deployment model

The IBM reference architecture reflects the work that is continually underway within IBM Systems to integrate elements of the IBM compute and storage portfolio such that they deliver high levels of performance for big data, while at the same time lowering the total cost of IT ownership.

This chapter contains the following topics:

- ▶ Composable genomics blueprint
- ▶ IBM Software-Defined Infrastructure
- ▶ Multicloud deployment model

3.1 Composable genomics blueprint

A composable, building-block-based solution for genomics, such as the IBM Reference Architecture for High-Performance Data Analytics and AI, addresses the most complex aspect of data management. By doing so, it enables organizations to store, access, manage, and share huge volumes of genomic sequencing data. In addition, the solution addresses workload management challenges to enable IT to build, refine, submit, and orchestrate computational jobs with maximum resource utilization.

In Figure 2-1 on page 14, the solution components are represented graphically by the following layers:

- ▶ The top two layers represent the applications, databases, and frameworks that researchers and clinicians are using (red and orange boxes).
- ▶ The bottom layer represents virtual or physical high-performance compute and storage servers where data is processed and stored (white boxes).
- ▶ The two middle layers (green and blue boxes) contain software that enables a diverse set of biomedical applications to run on shared compute and storage servers efficiently and cost-effectively.

The two middle layers have the following functionality:

- ▶ The workload management layer (*Orchestrator*) makes it possible to distribute thousands of computational workflows, in parallel, across enterprise compute servers in a way that maximizes utilization of those servers.
- ▶ The data management layer (*Datahub*) enables low-latency data access, convergence of heterogeneous data silos, metadata collection, and automated information lifecycle management (ILM).

These two layers give organizations the ability to rapidly scale compute and storage capacity upwards, or shrink capacity downward as workloads demand.

For more information about running genomics workloads, see *IBM Spectrum Scale Best Practices for Genomics Medicine Workloads*, REDP-5479 at the following website:

<http://www.redbooks.ibm.com/abstracts/redp5479.html>

3.2 IBM Software-Defined Infrastructure

Using the IBM Software-Defined Infrastructure, the solution provides healthcare and life sciences leaders with a set of options that includes the following features:

- ▶ A fully qualified and tested stack for next-generation sequencing (NGS) workloads with crisp configuration templates, including tuning and optimization of the blocks to efficiently run complete workflows from the Broad Institute Genome Analysis Toolkit (GATK) and Burrows-Wheeler Aligner (BWA).
- ▶ Easy and rapid assembly of a tested, composable infrastructure to create a scalable, HPC-like cluster to analyze genomics data.
- ▶ Flexible deployment models that are derived from the composable design, enabling simple and easy scalability of the storage, compute, and network building elements. These models include preferred practices for implementing genomics clusters in different configuration sizes based on actual customer needs.
- ▶ Neutral compute platforms.

- An innovative GUI for submitting and managing compute jobs and for viewing cluster status and utilization.
- User nodes to access interactive applications.

The IBM solution can enable data scientists to efficiently deliver results to physicians, meet physicians' on-demand requests for analysis, and improve the underlying infrastructure to meet core objectives. Using this solution, IT architects and IT administrators can easily design, install, and manage deployment in a timely manner without being overwhelmed.

3.3 Multicloud deployment model

You can learn how your organization can take advantage of a cloud with software-defined infrastructure to advance precision medicine. You can start getting answers fast and save costs on infrastructure, as shown in the following video:

<https://bit.ly/2F7Wb85>

3.3.1 Clouds over the ocean

Advances in precision medicine, genomics, and imaging; the widespread adoption of electronic health records; and the proliferation of medical internet of things (IoT) and mobile devices are resulting in an explosion of structured and unstructured healthcare-related data. Industry analysts expect that by 2020, the amount of medical data in the world will double every 73 days. In addition, a typical healthcare consumer in developed countries will generate 1,200 terabytes of data in a lifetime.¹

A substantial amount of this healthcare data deluge serves to advance precision medicine, such as medical research, which drives the need for HPC environments to support big data demands. For example, sequencing an individual's entire genome, a task growing ever more common, requires the same amount of data storage as 100 feature-length movies.²

However, data volume is not the only problem faced by medical research. Disparate file types generated by different research tools and environments create silos that impede data access, drive down efficiency, drive up costs, and slow times to insight. To address the challenges posed by both the volume and variety of medical research data, world-class healthcare organizations are building data "oceans".³

To construct data oceans, software-defined infrastructure solutions are deployed as foundations to manage and run rapidly evolving healthcare and life sciences applications (for example genomics, imaging, and clinical). These software-defined infrastructure solutions enhance the HPC platform with analytics open frameworks, such as Hadoop and Spark, and consolidate disparate data stores.

Behind the scenes, the software-defined infrastructure architecture creates a Databus to manage the ocean of data, orchestrate the different applications, and provide intelligent workload and policy-driven resource management. Putting all the data together into one coherent data resource to be analyzed, and making it available to all users anywhere-anytime, is key to facilitating research and accelerating time to insights.

¹ IDC InfoBrief: Modernizing Healthcare IT for the Data-driven Cognitive Era, April 2018

(<https://www-01.ibm.com/common/ssi/cgi-bin/ssialias?htmlfid=41015041USEN&>)

² Workgroup for Electronic Data Exchange (WEDI)

(<https://www.wedi.org/docs/publications/a-white-paper-by-the-genomics-workgroup.pdf?sfvrsn=0>)

³ IBM video: High-Performance Data Architecture for Healthcare

(https://www.youtube.com/watch?v=Fgrq8nyihY&list=PLS7mekU2kxDowLMKIpbA_ZQb5Fipyv4zR&index=1)

The benefits are substantial. The ability to automatically migrate medical data to the optimal storage tier can substantially reduce costs.

Eliminating the need for separate processing platforms for different data types dramatically increases resource utilization. Massive parallel processing and enhanced application and data portability accelerate time to insights. However, it is not entirely serene “sailing” across HPC data oceans. They do not solve every data processing problem. Medical research, like many other HPC environments, generates peaks and valleys of resource demand.

The efficiency and high utilization rates that data oceans are designed to produce can work against them when demand peaks beyond infrastructure capabilities. To accommodate these spikes in demand, traditional HPC environments often divide jobs and stretch out scheduling, lengthening time to insight. However, the same software-defined infrastructure solutions used to create data oceans can also address this challenge by adopting a hybrid cloud.

IBM Spectrum Scale and IBM Spectrum Computing family members, such as IBM Spectrum LSF, have long and successful histories of providing solutions to the full landscape of HPC challenges.

For example, by using IBM Spectrum LSF, healthcare researchers can determine through an advanced reporting functionality where the bottlenecks are that cause jobs to run slower. Then IBM Spectrum LSF can move targeted jobs to the cloud:

- ▶ Does an HPC job require more memory? Run it on servers in the cloud with more memory.
- ▶ Does it need faster access to the underlying data? Provision a massively parallel IBM Spectrum Scale file system for the fastest access on the planet.

Whatever resources are needed, with IBM Software-Defined Infrastructure solutions, healthcare researchers can provision the required system for peak demand periods in the cloud, dynamically and automatically, only for as long as needed, resulting in faster insights for a fraction of the cost of building on-premises solutions.

The question becomes, can every high-performance data architecture provide the infrastructure support, flexibility, and agility needed to meet highly demanding and unpredictable healthcare HPC requirements? Again, IBM offers solutions for cloud-scale data management and multicloud workload orchestration based on a Reference Architecture for High Performance Data and AI Platforms (HPDA).

Teaming with L7 Informatics (L7), the two solution-providers have built a cloud-based HPC environment that enables scientists to process and analyze huge volumes of genomics data up to 96% faster. Built on the IBM Cloud platform, the L7 Genomic Cloud uses IBM Spectrum Scale and IBM Spectrum LSF to support rapid data processing and analysis.⁴ For more information, see 6.3, “L7 Informatics” on page 57.

Chris Mueller, Founder of L7 Informatics, explains: *“IBM Spectrum Scale provides high-performance data storage that we can scale quickly and easily. Built-in tiering capabilities allow a lot of flexibility in how we move data around, enabling customers to seamlessly migrate data from lab instruments up to the cloud for analysis and long-term storage. IBM Spectrum LSF, meanwhile, offers everything we need for HPC workload management in a single package, from job scheduling tools to resource management capabilities. It gives us the tools to manage the L7 Genomic Cloud as a complete HPC environment rather than just as a virtual machine and associated storage layer, providing intelligent, policy-driven scheduling and improved visibility to increase throughput.”*

⁴ IBM Marketplace: L7 Informatics, (<https://www.ibm.com/case-studies/l7-informatics-systems-spectrum-hpc>)

Oceans and clouds. For millennia these natural systems have nourished and supported humanity. Perhaps it is not as surprising as it might seem that digital versions of them are now helping to accelerate medical advancements that offer great human benefit.

Follow the links in the case study to learn more about how you can build software-defined data oceans and agile hybrid clouds that help your organization lower costs, and at the same time gain precious insights faster.



Building blocks

This chapter describes the building blocks for the solution, and contains the following topics:

- ▶ IBM Spectrum Storage
- ▶ IBM Spectrum Computing
- ▶ IBM Power System AC922 for HPC

4.1 IBM Spectrum Storage

The life science industry's explosive data growth in genomics, translational, and personalized medicine sectors have created massive challenges for their IT organizations. Managing this magnitude of data requires high-performance technical environments to process the enormous amount of unstructured data and support the increasingly sophisticated simulations and analyses workloads.

Organizations looking for root causes and cure for diseases need their data infrastructure to seamlessly produce the speed, accessibility, reliability, and security to increase productivity at lower costs, foster innovation and collaboration, and compete more effectively.

These critical data infrastructure capabilities can be addressed by the IBM Spectrum Storage family of products. IBM Spectrum Storage delivers proven technology for software-defined storage that can dynamically and flexibly store data at optimal cost, helping maximize performance and ensure data protection.

The IBM Spectrum Storage family includes IBM Spectrum Control and IBM Spectrum Protect for simplified management, IBM Spectrum Archive and IBM Spectrum Virtualize for increased efficiency, and IBM Spectrum Accelerate and IBM Spectrum Scale for the agility to meet changing needs.

IBM Spectrum Scale fits well for the workload this book focuses on, which are distributed and cloud-ready applications.

4.1.1 IBM Spectrum Scale

In the IBM Reference Architecture for Genomics, data management for computational workloads is enabled by IBM Spectrum Scale. IBM Spectrum Scale is a proven, scalable, and high-performance data and file management solution that provides world-class storage management with scalability. Leading genomics centers, medical institutions, and pharmaceutical companies worldwide are already investing in IBM Spectrum Scale.

These organizations use IBM Spectrum Scale to store, archive, process, and manage a vast amount of structured and unstructured data, including genomic sequences, biomedical images, and electronic medical records. IBM Spectrum Scale is well-suited for managing biomedical data and related analytics because it addresses the following challenges:

- ▶ Rapid growth of data volumes.
- ▶ I/O-intensive workloads, which are often required to analyze raw genomic sequences.
- ▶ Heterogeneous file and object storage clusters that use different operating systems, storage access protocols, storage media, and storage hardware.
- ▶ Data sharing across globally distributed projects.
- ▶ Metadata identification, collection, and search, which are required for scientific and clinic repeatability, validation, and long-term archiving.
- ▶ Requirements for secure storage and secure deletion of data containing sensitive information.

Moreover, IBM Spectrum Scale enables policy-driven information lifecycle management (ILM) across multiple storage tiers that are built on flash storage, disk, and tape media across local or remote locations. Automated policies make it possible for administrators to define where, when, and on what media data (or metadata) will be stored to maximize workload performance and minimize overall storage costs.

For example, low-latency flash storage systems might provide the best price performance for small volume and highly used data in I/O-intensive workloads. In contrast, Linear Tape File System (LTFS) tape offers the best price performance for large volume and less-used data sets that are ready for long-term archive.

4.1.2 IBM Spectrum Archive

IBM Spectrum Archive, a member of the IBM Spectrum Storage family, enables direct, intuitive, and graphical access to data that is stored in IBM tape drives and libraries by incorporating the LTFS format standard for reading, writing, and exchanging descriptive metadata on formatted tape cartridges. IBM Spectrum Archive eliminates the need for extra tape management and software to access data.

IBM Spectrum Archive offers several software solutions for managing your digital files with the LTFS format:

- ▶ Single Drive Edition (SDE)
- ▶ Library Edition (LE)
- ▶ Enterprise Edition (EE)

This book focuses on IBM Spectrum Archive EE.

IBM Spectrum Archive Enterprise Edition provides seamless integration of LTFS with IBM Spectrum Scale, which is another member of the IBM Spectrum Storage family, by creating a tape-based storage tier. You can run any application that is designed for disk files on tape by using IBM Spectrum Archive EE because it is fully transparent and integrates into the IBM Spectrum Scale file system. IBM Spectrum Archive EE can play a major role in reducing the cost of storage for data that does not need the access performance of a primary disk.

With IBM Spectrum Archive EE, you can enable the use of LTFS for the policy management of tape as a storage tier in an IBM Spectrum Scale environment, and use tape as a critical tier in the storage environment.

The use of IBM Spectrum Archive EE to replace online disk storage with tape in tier 2 and tier 3 storage can improve data access compared to other storage solutions because it improves efficiency and streamlines management for files on tape. IBM Spectrum Archive EE simplifies the use of physical tape by making it transparent to the user, and manageable by the administrator, under a single infrastructure.

IBM Spectrum Archive EE uses an enhanced version of the IBM Spectrum Archive LE, which is referred to as *the IBM Spectrum Archive LE+ component*, for the movement of files to and from tape devices. The scale-out architecture of IBM Spectrum Archive EE can add nodes and tape devices as needed to satisfy bandwidth requirements between IBM Spectrum Scale and the IBM Spectrum Archive EE tape tier.

Low-cost storage tier, data migration, and archive needs that are described in the following use cases can benefit from IBM Spectrum Archive EE:

- ▶ Operational storage
 - Provides a low-cost, scalable tape storage tier.
- ▶ Active archive
 - A local or remote IBM Spectrum Archive EE node that serves as a migration target for IBM Spectrum Scale and transparently archives data to tape based on policies that are set by the user.

The following IBM Spectrum Archive EE characteristics cover a broad base of integrated storage management software with leading tape technology and the highly scalable IBM tape libraries:

- ▶ Integrates with IBM Spectrum Scale by supporting file-level migration and recall with an innovative database-less storage of metadata.
- ▶ Provides a scale-out architecture that supports multiple IBM Spectrum Archive EE nodes that share tape inventory with load balancing over multiple tape drives and nodes.
- ▶ Enables tape cartridge pooling and data exchange for IBM Spectrum Archive EE tape tier management:
 - Tape cartridge pooling enables the user to group data on sets of tape cartridges.
 - Multiple copies of files can be written on different tape cartridge pools, including different tape libraries in different locations.
 - Supports tape cartridge export with and without the removal of file metadata from IBM Spectrum Scale.
 - Supports tape cartridge import with pre-population of file metadata in IBM Spectrum Scale.

Furthermore, IBM Spectrum Archive EE provides the following key benefits:

- ▶ A low-cost storage tier in an IBM Spectrum Scale environment.
- ▶ An active archive or big data repository for long-term storage of data that requires file system access to that content.
- ▶ File-based storage in the LTFS tape format that is open, self-describing, portable, and interchangeable across platforms.
- ▶ Lowers capital expenditure and operational expenditure costs by using cost-effective and energy-efficient tape media without dependencies on external server hardware or software.
- ▶ Enables the retention of data on tape media for long-term preservation (10+ years).
- ▶ Provides portability for large amounts of data by bulk transfer of tape cartridges between sites for disaster recovery, and the initial synchronization of two IBM Spectrum Scale sites by using open-format, portable, and self-describing tapes.
- ▶ Migration of data to newer tape or newer technology that is managed by IBM Spectrum Scale.
- ▶ Provides ease of management for operational and active archive storage.
- ▶ Expand archive capacity by adding and provisioning media without affecting the availability of data in the pool.

4.1.3 IBM Cloud Object Storage

As hybrid cloud adoption continues to accelerate in the healthcare and life sciences market, customers are incorporating more object storage in their data infrastructure for new AI and Big Data requirements. IBM Cloud Object Storage is a marketing leading object storage solution in the data center that meets many of these workload requirements, and can also be used as a server in the IBM Cloud.

The IBM system provides data that is accessible from any location with concurrent parallel access and the simplicity to lower cost for growing and modern data requirements. With proven scalability to exabyte (EB) capacity and the ability to start as small as 72 TB, this solution provides a strong foundation to keep data safe and available for many years.

IBM Cloud Object Storage delivers the following key benefits:

- ▶ Flexibility. Security and control in your data center, or offloaded operational costs using the public cloud
- ▶ Security. Built in air gap, data encryption, and lockable WORM (write once, read many) data
- ▶ Data Protection. A geo-dispersed protection layer for demanding SLAs
- ▶ Lower Costs. Designed to easily scale-up and scale-out with the simplicity and investment protection to minimize costly upgrades
- ▶ Simplicity. Proven scalability from TB to EB with simplicity to scale online

With IBM industry-leading object storage platform, you can select the best configuration and approach to address the unique application, data, and workload requirements for your business. Start with as few as 3 nodes and grow to 1000s of share-nothing nodes seamlessly, without any downtime and with investment protection.

With a unique access layer for increased throughput and performance, the system can be configured from multiple demanding workloads and configurations. The platform uses current and new modern applications, such as containers, micro services, and cloud native apps.

IBM Cloud Object Storage is available in the following modes:

- ▶ Private on-premises object storage
- ▶ Dedicated object storage (single-tenant)
- ▶ Public object storage (multi-tenant)
- ▶ Hybrid object storage (a mix of on-premises, dedicated, or public offerings)

IBM Cloud Object Storage gives you the choice to deploy object storage on-premises, in the public cloud, or both on-premises and in the cloud in a hybrid solution. In addition, public cloud services (Standard Object Storage and Vault Object Storage) can be configured in either a Regional or Cross-Regional model, providing even more choices when it comes to the level of data protection and resiliency that you need for workloads.

IBM Cloud Object Storage is a dispersed storage mechanism that uses a cluster of storage nodes to store pieces of the data across the available nodes. IBM Cloud Object Storage uses an *Information Dispersal Algorithm (IDA)* to break files into unrecognizable slices that are then distributed to the storage nodes.

No single node has all of the data, which makes it safe and less susceptible to data breaches, while at the same time needing only a subset of the storage nodes to be available to retrieve fully the stored data. This ability to reassemble all of the data from a subset of the chunks dramatically increases the tolerance to node and disk failures.

The IBM Cloud Object Storage architecture is composed of three functional components. Each of these components runs ClevOS software that can be deployed on compatible, industry-standard hardware. The three components include:

- ▶ IBM Cloud Object Storage Manager
IBM Cloud Object Storage Manager provides an out-of-band management interface that is used for administrative tasks, such as system configuration, storage provisioning, and monitoring the health and performance of the system.
- ▶ IBM Cloud Object Storage Accesser®
IBM Cloud Object Storage Accesser imports and reads data, encrypting/encoding data on import, and decrypting/decoding data on read. It is a stateless component that presents the storage interfaces to the client applications, and transforms data by using an IDA.

- IBM Cloud Object Storage Slicestor®

The IBM Cloud Object Storage Slicestor node is primarily responsible for storage of the data slices. It receives data from the IBM Cloud Object Storage Accesser on import, and returns data to the IBM Cloud Object Storage Accesser as required by reads.

4.1.4 IBM Spectrum Discover

IBM Spectrum Discover is modern metadata management software that provides data insight for exabyte-scale unstructured storage. IBM Spectrum Discover easily connects to multiple file and object storage systems both on-premises and in the cloud to rapidly ingest, consolidate and index metadata for billions of files and objects, providing a rich metadata layer on top of these storage sources. This metadata enables data scientists, storage administrators, and data stewards to efficiently manage, classify and gain insights from massive amounts of unstructured data.

For more information about Spectrum Discover, please refer to the following website:

<https://www.ibm.com/us-en/marketplace/spectrum-discover>

Challenges associated with processing unstructured data is particular prevalent in the healthcare industry. The following published article discusses the issue and explains how IBM Spectrum Discover is helping with these challenges:

<https://bit.ly/2wMJ4DQ>

4.2 IBM Spectrum Computing

The workload management layer dynamically and elastically allocates computational tasks across compute servers in a manner that is transparent to the user. It consists of multiple coherent workflow schedulers that are coordinated to place diverse compute jobs on local and remote clusters in an efficient and cost-effective way.

IBM resource-aware and policy-based schedulers include industry-leading IBM Spectrum LSF (for HPC batch workloads), IBM Spectrum Conductor (for Spark workloads), IBM Spectrum Conductor Deep Learning Impact, and IBM Watson Machine Learning Accelerator for Enterprise AI (for deep learning workloads). These schedulers are tightly integrated: If one type of workload is using only a few resources in a cluster, then the other workload types can fully use the remaining resources in that cluster.

The flexibility and elasticity of server use across these schedulers eliminates the need for IT organizations to provide dedicated clusters for each of the workload types. When serving a multitenant environment, these schedulers protect individual tenants through secure isolation and role-based access control (RBAC). They make it possible for workloads to be distributed seamlessly across multiple physical and cloud environments, and they support the distribution of workloads that are deployed in Docker and other container technologies.

4.2.1 IBM Spectrum LSF Suite

IBM Spectrum LSF Suite provides a tightly integrated solution for HPC delivering systems and workload management designed to increase the productivity of users and utilization of hardware, at the same time controlling system management costs. It supports high-performance and high-throughput workloads through to big data, GPU machine learning, and containerized workloads, both on-premises and in the cloud.

IBM Spectrum LSF has been selected as the preferred workload management system by large genome analysis organizations for its ability to routinely orchestrate hundreds of thousands of jobs that are submitted in batch, and for its ability to readily scale with growing user demand. Clients worldwide are using technical computing environments that are supported by IBM Spectrum LSF to run hundreds of genomic workloads, including Burrows-Wheeler Aligner (BWA), SAMtools, Picard, Broad Institute Genome Analysis Toolkit (GATK), Isaac, CASAVA, and other frequently used pipelines for genomic analysis.

IBM Spectrum LSF Suite is available in three versions with progressively greater capabilities suitable for single clusters through to the largest supercomputers, as shown in Figure 4-1.

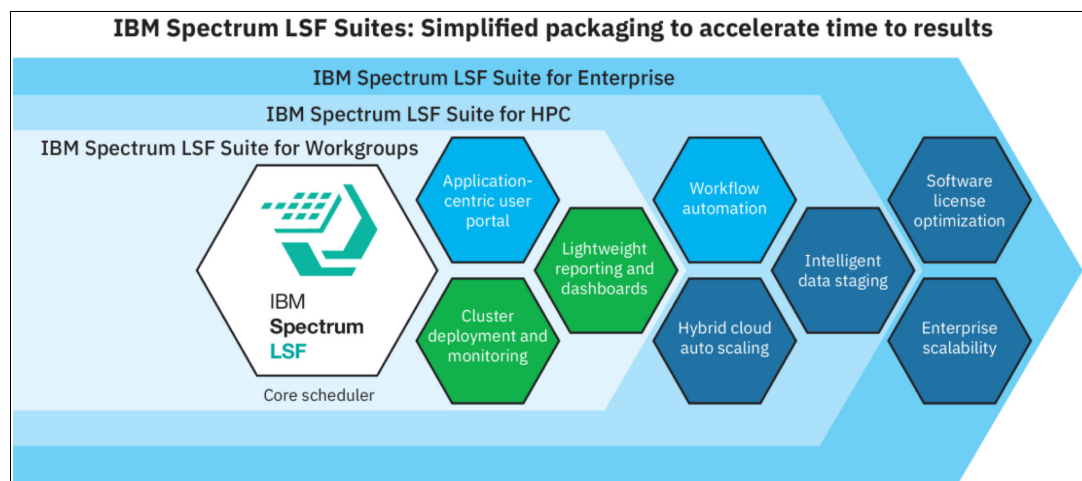


Figure 4-1 IBM Spectrum LSF Suite

User productivity

IBM Spectrum LSF provides a superior user experience through a simple and intuitive web-based interface for submitting, managing, and monitoring jobs. Application-specific templates enable users to run jobs without having to create customer scripts. Additionally, intuitive, self-documenting scripting guidelines help simplify the creation of additional job submission templates, resulting in reduced setup time at the same time minimizing user errors during workload submissions.

Ease of use is further extended by the availability of job-monitoring mobile clients for Google Android and Apple iOS platforms, and as an integrated client for Microsoft Windows environments for job management and a RESTful API. These ease-of-use features enable organizations to drive productivity by hiding complexity, with interfaces that enables domain experts to focus on research rather than IT.

Computational workflows

Clients that are interested in personalized healthcare research have taken advantage of the workflow management capabilities provided by the IBM Spectrum LSF family. These clients use this tool to simplify the process of writing genomic workflow scripts that transform raw data from next-generation sequencers (in the FASTQ format) into variant files (for example, Variant Call File (VCF), SNV, and CV) for downstream analysis. This makes it possible for bioinformaticians to share workflows with selected users who do not have formal experience.

Additionally, organizations using Common Workflow Language (CWL) for workflows can run these workflows on IBM Spectrum LSF with support provided by the CWLEXEC open source project(<https://github.com/IBMSpectrumComputing/cwlexec>). This provides a smooth integration of CWL workflows with IBM Spectrum LSF, benefiting from high efficiency, scalability, self-healing of workflows, and support for user-specified options, while keeping CWL definitions portable.

Containerized workloads

IBM Spectrum LSF provides support for organizations that are using container technologies to help streamline the building, testing, and shipping of applications, which enables an application stack to be consistently deployed both on-premises and in the cloud. A generalized interface provides support for Docker, Shifter, and Singularity container technologies.

Containerized jobs submitted to IBM Spectrum LSF benefit from resource binding, interactive and parallel job support, and reliability from automatically rerunning containers during failures. Additionally, access controls enable administrators to define which container images can be run in the environment.

Hybrid cloud

The majority of HPC environments serve groups of users potentially working on multiple projects with differing priorities and deadlines. Situations occur where multiple users, projects, and applications outstrip the available resources, leading to costly delays in time to results. To address this, organizations are increasingly turning to hybrid-cloud solutions to address these peaks and valleys in demand for computing resources.

IBM Spectrum LSF provides advanced hybrid-cloud capabilities, enabling workloads to be forwarded to multiple clouds with support for different cloud providers, and data to be automatically staged to or from the cloud. By dynamically provisioning external cloud resources in response to workload demands, IBM Spectrum LSF enables organizations to intelligently and transparently control the use of these cloud resources, so that you only pay for what you use.

GPU-aware

GPU-accelerated computing is now commonplace in HPC environments, and GPU support is emerging in an ever-increasing number of genomics and life science applications. As with any other resources in a computing environment, GPUs must be intelligently managed and utilized for maximum effectiveness. IBM Spectrum LSF provides a number of advanced capabilities supporting NVIDIA GPUs:

- ▶ GPUs as a schedulable resource
- ▶ Reporting on GPU utilization
- ▶ CPU-GPU affinity
- ▶ Enforcement of GPU allocations by way of cgroup
- ▶ GPU power management and boost control
- ▶ NVIDIA Multi-Process Service (MPS) support
- ▶ NVIDIA Data Center GPU Manager (DCGM) support for enhanced job accounting and health check
- ▶ Extended syntax for supporting GPU workloads and monitoring

4.2.2 IBM Spectrum Conductor

IBM Spectrum Conductor is an enterprise-class, multitenant solution for Apache Spark and Anaconda/Python. It provides a framework to enable other application integrations, sharing resources dynamically.

It specifically enables your organization to deploy Apache Spark-based and Python-based applications efficiently and effectively, supporting multiple concurrent instances and versions. It can help increase performance and scale, optimize resource usage, and eliminate silos of resources that otherwise be tied to multiple, separate Apache Spark or Anaconda or Python implementations.

Apache Spark

Apache Spark is a common application framework used by individuals conducting big data, AI and deep learning, computational research in personalized healthcare, and other areas of biomedical science. Spark is also a disruptive technology with huge momentum. It dominates all other Apache open source projects in terms of community activity.

Spark is generating significant interest as a big data analytics solution because of its perceived advantages over Hadoop MapReduce. It runs faster than MapReduce (especially when run in-memory) and offers easy-to-use APIs for a variety of programming languages. It also includes a rich set of high-level tools, including Spark SQL, machine learning, graph processing, and stream processing capabilities.

Apache Spark requires a resource manager and can interface with various distributed storage models. IBM Spectrum Conductor addresses both of these requirements:

- ▶ Incorporates a generalized resource manager that provides a shared-service backbone for a broad portfolio of distributed software frameworks
- ▶ Enables a wide variety of applications to share resources and coexist on the same infrastructure

The offering includes a Spark distribution, providing a robust end-to-end solution for organizations that are considering Spark deployment, both for exploratory projects and in-production environments.

Benefits

IBM Spectrum Conductor provides the following benefits:

- ▶ Helps obtain business insights from data while reducing both capital and operational expenses
- ▶ Can lower costs by using a service orchestration framework for distributed Apache Spark workloads to optimize resource utilization
- ▶ Improves performance and efficiency with granular and dynamic resource allocation
- ▶ Simplifies administration with a consolidated framework for Apache Spark deployment, monitoring, and reporting
- ▶ Helps future-proof data centers with multidimensional scaling, including independent scaling of compute and storage infrastructure
- ▶ Enables a wide variety of applications to share resources and coexist on the same infrastructure
- ▶ Simplifies Spark deployment without requiring the complexity of a Hadoop stack
- ▶ Provides flexible, efficient management of data when deployed using IBM Spectrum Scale

Unique differentiators

IBM Spectrum Conductor is designed to enable organizations to deploy Spark efficiently, effectively, and confidently. Unlike other offerings that require piecemeal assembly of components or a full Hadoop stack, IBM Spectrum Conductor is a robust offering that accelerates results with policy-based, workload-aware resource management. The solution provides an IBM-supported framework for workload management, monitoring, reporting, and end-to-end security.

It also includes a generalized resource manager that provides a shared-service backbone for a broad portfolio of distributed software frameworks. IBM Spectrum Conductor, when combined with IBM Spectrum Scale for storage management, is POSIX-compliant, unlike open-source HDFS, and provides significant storage efficiencies compared to HDFS. Nevertheless, IBM Spectrum Conductor also supports HDFS and other storage management technologies for clients who prefer alternative options.

Collaboration

IBM Spectrum Conductor supports and deploys with the Spark instances. Notebooks provide an interactive environment for data analysis, enabling you to explore and visualize data analytics from your browser.

Create shared Spark batch applications, enabling multiple users to submit Spark jobs to an application and use the same Resilient Distributed Datasets (RDDs). The benefit is that, through the use of new Spark jobs, multiple users can analyze the same RDDs without having to recompute them multiple times in new Spark apps. Shared Spark batch apps use an API to create and manage sharable RDDs in a Spark context. The sharable RDD API provides a data caching layer, wherein the shared RDD data is computed after and cached for reuse.

The Spark RDD abstraction is a collection of partitioned data elements that can be operated on in parallel. RDDs work at the application level, wherein each application manages its own RDDs.

Faster and consistent performance

IBM Spectrum Conductor includes an advanced scheduler, which provides high-performance resource and workload managers. The logic behind the scheduler controls multitenancy, enabling multiple concurrent and different versions of an application to exist in a single cluster. It also controls the orchestration and execution of each task, and delivers superior performance compared to Mesos and YARN:¹

- ▶ 5 - 88% higher throughput than Apache Mesos
- ▶ 0 - 224% higher throughput than Apache YARN

Enterprise security

Security in IBM Spectrum Conductor has been proven in the field by being deployed in highly regulated industries, such as financial services and other life science research environments. Security is implemented around the following principles:

- ▶ Authentication. Support for Kerberos, Siteminder, AD/LDAP, and operating system authentication, including Kerberos authentication for HDFS
- ▶ Authorization. Fine-grained access control, ACL RBAC, Spark binary lifecycle, notebook updates, deployments, resource plans, reporting, monitoring, log retrieval, and execution
- ▶ Impersonation. Allow different tenants to define production execution users
- ▶ Encryption. SSL and authentication between all daemons

¹ Source: <https://stacresearch.com/news/2017/05/19/IBM170405>

4.3 IBM Power System AC922 for HPC

The IBM Power System AC922 is the next generation of the IBM POWER® processor-based systems that are specifically designed for DL and AI, HPDA, and HPC.

The system is co-designed with the OpenPOWER Foundation members and delivers the latest technologies available for HPC and improves the movement of data from memory to GPUs and back, which enables faster and lower latency data processing. This massive computing capacity is packed into just 2Us of rack space. To accommodate the thermal challenges inherent with large deployments, IBM offers two models with different cooling implementations: The 8335-GTH is air-cooled, and the 8335-GTX is water-cooled.

4.3.1 Accelerated computing with IBM POWER9 processor-based systems

IBM Systems is committed to providing superior performance for the most challenging computational workloads. The IBM strategy for achieving this goal is based on the observation that next-generation HPC systems will need to support both data-intensive workloads and compute-intensive workloads.

As requirements for traditional HPC and newer big data analytics converge, system throughput depends on I/O performance improvements at the level of the central processing units (CPUs), and on minimizing data movement within the architecture. It also depends on tightening the integration of the CPU with hardware accelerators, such as GPUs and Field Programmable Gate Arrays (FPGAs), which are often used to dramatically speed up user applications.

4.3.2 OpenPOWER Foundation

In 2013, IBM began open source licensing of products that are related to the IBM Power Architecture® to create an alternative to proprietary solutions and spur innovation in computing technology. IBM initiated the creation of the OpenPOWER Foundation, which is a technical community of more than 130 commercial and academic organizations collaborating on the development of IBM POWER processor-based solutions that better meet specific business needs. Innovations arising from OpenPOWER collaborations include custom systems for workload acceleration by using GPUs, FPGAs, and advanced I/O.

4.3.3 OpenPOWER processors

OpenPOWER systems are being designed to support compute and data intensive workloads, such as machine learning and deep neural networks. IBM POWER9 processors have simultaneous multithreading (SMT), which enables up to eight hardware threads from a single physical core. It also employs the most advanced memory subsystem available to achieve leading-edge performance, by using many on- and off-chip memory caches. This processor design reduces memory latency and generates high bandwidth for memory and I/O.

For more information about the IBM POWER9 processors, see *IBM Power System AC922 Introduction and Technical Overview*, REDP-5472, at the following website:

<http://www.redbooks.ibm.com/abstracts/redp5472.html>

4.3.4 Recent advancements

The following features are designed to augment the performance of POWER9 systems:

- ▶ IBM CAPI2 is the evolution of CAPI that defines a coherent accelerator interface structure for attaching special processing devices to the POWER9 processor bus. As with the original CAPI, CAPI2 can attach accelerators that have coherent shared memory access with the processors in the server. CAPI2 can also share full virtual address translation with these processors by using standard PCIe Gen4 buses with twice the bandwidth compared to the previous generation.
- ▶ Applications can have customized functions in Field Programmable Gate Arrays (FPGAs), and queue work requests directly in shared memory queues to the FPGA. Applications can also have customized functions by using the same effective addresses (pointers) that they use for any threads running on a host processor. From a practical perspective, CAPI enables a specialized hardware accelerator to be seen as an extra processor in the system, with access to the main system memory and coherent communication with other processors in the system.
- ▶ NVLink 2.0 is the NVIDIA new generation high-speed interconnect technology for GPU-accelerated computing. Supported on SXM2-based Tesla V100 accelerator system boards, NVLink increases performance for both GPU-to-GPU communications and for GPU access to system memory.
- ▶ Support for the GPU ISA enables programs running on NVLink-connected GPUs to run directly on data in the memory of another GPU and on local memory. GPUs can also perform atomic memory operations on remote GPU memory addresses, enabling much tighter data sharing and improved application scaling.

4.3.5 Applications

POWER9 processor-based servers are supported by standard Linux distributions, making it easy to port existing codes to the platform. Preferred applications in the healthcare and life sciences sector, such as GATK, BWA, SAMtools, BLAST, MuTect2, and tranSMART, already are enabled and optimized on IBM PowerLinux.

Bioinformatics specialists within IBM Systems, along with IBM Business Partners in the healthcare and life science industry, are actively engaged in porting and optimizing the performance of more biomedical research codes on POWER. More than 100 open source applications are enabled on POWER9.

Within research fields that are relevant to personalized healthcare, such as genomics and bioinformatics, adoption of GPU-enabled workloads has been slow; however, these workloads appear to be gaining traction. An increasing number of organizations conducting biomedical research are applying deep learning techniques to uncover predictive patterns within large sets of often unstructured data, such as biomedical images and time-varying physiological signals. In support of such workloads, IBM and NVIDIA are collaborating on IBM Watson Machine Learning Accelerator for Enterprise AI, a new deep learning toolkit.

PowerAI is an easy-to-deploy deep learning platform that delivers popular deep learning frameworks, including Caffe, Torch, and Theano, within the IBM Power Architecture. IBM has optimized each of these deep learning software distributions to take advantage of the high bandwidth that is offered by the IBM POWER9 processor and NVIDIA NVLink 2.0 interconnect. The toolkit also uses NVIDIA GPUDL libraries, including cuDNN, cuBLAS, and NCCL, to deliver multi-GPU acceleration on IBM servers.

IBM and NVIDIA partnered in integrating IBM POWER systems with NVIDIA GPUs and the enablement of GPU-accelerated applications and workloads. The computational capability that is provided by the combination of NVIDIA Tesla GPUs and IBM Power Systems servers enables workloads from scientific, technical, and HPC to run on data center hardware.

This computational capability is built on top of massively parallel and multithreaded cores with NVIDIA Tesla GPUs and IBM POWER architecture processors, where processor-intensive operations are offloaded to GPUs and coupled with the system's high memory-hierarchy bandwidth and I/O throughput.



Use cases

This chapter describes use cases and contains the following topics:

- ▶ The Broad Institute Genome Analysis Toolkit (GATK)
- ▶ Expanding IBM Reference Architecture for High-Performance Data Analytics into medical imaging

5.1 The Broad Institute Genome Analysis Toolkit (GATK)

The Broad Institute GATK is a set of applications that is used to create multi-step workflows for variant discovery analysis of both germline and somatic genomes. Each step has its own set of tools. The output from each step is the input to the next step. There are a variety of GATK-based workflows that are used in the field. For this paper, the examples profile the workflow that is documented in the Broad Institute GATK best practices (as shown in Figure 5-1).

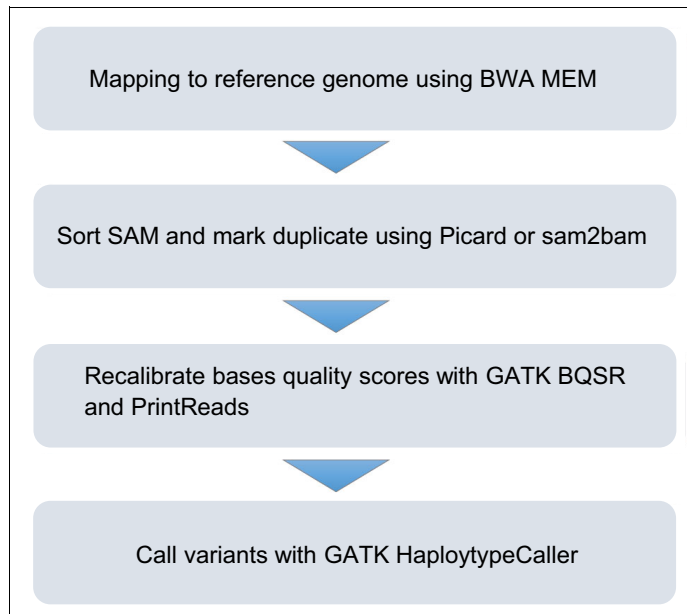


Figure 5-1 GATK based workflows

To profile that workflow, the following environment was used:

- ▶ A single IBM POWER8® node (IBM 8247-22L with SMT8) with 256 GB of memory to execute the whole workflow
- ▶ 1x IBM Elastic Storage Server (ESS) GS4 with SSDs (≥ 23 GBps write bandwidth and ≥ 30 GBps read bandwidth)
- ▶ Dual-rail FDR InfiniBand aggregating to ~ 13 GBps
- ▶ GATK pipeline execution using Whole Genome Sequence (WGS) input data set with B37 reference data set

We set SMT8 on the Compute Nodes based on Table 3 in “A guide to GATK4 best practice pipeline performance and optimization on the IBM OpenPOWER system” (<https://www.ibm.com/downloads/cas/ZJQD0QAL>). Furthermore, BWA-Mem required SMT8 to launch 160 threads on a node.

Note: This profiling setup is different than the example configuration. However, the insights gained by the profiling influenced the example configuration.

Figure 5-1 is a four-step workflow that requires running six applications. The profiling was performed with the Solexa WGS data set provided by the Broad Institute.

Figure 5-2 also shows the run time achieved on the profiling environment.

	Solexa WGS Broad dataset with b37 reference
BWA-Mem	303 min 47 sec
sam2bam (storage mode)	35 min 53 sec
GATK BaseRecalibrator (java setting -Xmn10g -Xms10g -Xmx10g)	87 min 21 sec
GATK PrintReads (java setting -Xmn10g -Xms10g -Xmx10g)	97 min 1 sec
GATK HaplotypeCaller (java setting -Xmn10g -Xms10g -Xmx10g)	261 min 37 sec
GATK mergeVCF (java setting -Xmn10g -Xms10g -Xmx10g)	0 min 51 sec
➔ Execution time was measured on the profiling environment using the Solexa WGS Broad dataset. The actual throughput or performance that any user will experience will vary depending upon many factors.	

Figure 5-2 Execution time on the profiling environment using the Solexa WGS Broad data set

Table 5-1 summarizes the workload profile of each processing step and shows the run time achieved on the profiling environment. The actual throughput or performance that any user experiences can vary, depending upon many factors. See Appendix A, “Profiling GATK” on page 65 for profiling details.

Table 5-1 Workload profile of each processing step

	BWA-Mem	sam2bam (storage mode)	GATK BaseRecalibrator	GATK PrintReads	GATK HaplotypeCaller	GATK mergeVCF
Run Time ^{a b}	303 min 47 sec	35 min 53 sec	87 min 21 sec	97 min 1 sec	261 min 37 sec	0 min 51 sec
CPU	Intensive. Close to 100% CPU utilization	~93% (initial phase) and ~40% in later phases	~70% CPU utilization	~70% CPU utilization	~40% CPU utilization	Less than 1% CPU utilization
Memory	Low memory consumption	Low memory consumption	Total of 18 x Java threads with each thread customized with 10 GB → 180 GB	Total of 18 x Java threads with each thread customized with 10 GB → 180 GB	Not memory intensive	Not memory intensive
File data I/O access pattern	Pattern of writes followed by reads. Predominantly sequential I/O.	Write I/O predominantly sequential I/O. Read I/O is random access in units of 512 KiB.	Predominantly read intensive. Read is mix of sequential and random I/O.	Mix of read and write. Write I/O is mostly 512 KiB with mix of sequential and random. Read is mostly sequential.	Mix of read and write. Write I/O is mix of sequential and random. Read is mostly sequential.	Mix of read and write. Read and write I/O is predominantly sequential I/O.
File I/O bandwidth	<= 200 MBps (read and write)	Write < 2.5 GBps. Sustained read < 300 MBps. High degree of pagepool cache hits during reads (< 36 GBps).	<= 100 MBps (read and write)	Write < 150 MBps and read < 75 MBps	Write < 100 MBps and read < 100 MBps	Write < 1.5 GBps and read < 2 GBps
File Metadata	<=2 inode updates	Initial phase <= 60 inode updates. Later phase, <=2 inode updates.	~24 file open and ~24 file closes.	~24 file open and ~24 file closes.	~20 file open and ~20 file closes.	~2 file open and ~2 file closes.
Output file or files	Single output file (*.sam) <= 380 GB file size	Two output files: ~77 GB (.bam) and ~9 MB (.bam.bai).	Total of 52 files: 26 x *.table.log-4" files (<200 KB) and 26 x *.table" files (< 300 KB)	Total of 78 files: 26 x *.recal_reads*.bam" files (< 15 GB), 26 x *.bai" files (< 750 KB), and 26 x *.recal_reads*.bam.log" files (< 200 KB)	Total of 78 files: 26 x *.raw_variants*.vcf" files (< 6 GB), 26 x *.raw_variants*.vcf.log" files (< 400 KB), and 26 x *.raw_variants*.vcf.idx" files (< 20 KB)	Single output file (*.raw_variants.vcf) with ~66 GiB file size

a. Execution time was measured on the profiling environment using the Solexa WGS Broad data set. The actual throughput or performance that any user will experience will vary depending upon many factors.

b. Java setting -Xmn10g -Xms10g -Xmx10g

The analysis of the GATK workflow of the Solexa WGS Broad data set guided the details of the design of the compute cluster and the /gpfs/data file system that is served by the storage services. In summary, BWA-Mem is CPU-intensive. For optimal performance, run this application on nodes with higher core count and higher clock frequency. The sam2bam is memory-intensive. It can be run in one of two possible modes: memory mode and storage mode.

In our test, sam2bam in storage mode took 35 minutes and 53 seconds to complete. For optimal performance, run sam2bam in memory mode on a node with >= 1 TiB of memory. The GATK Base Recalibrator and PrintRead steps are memory-intensive. To achieve optimal performance, run these application on nodes with >= 512 GiB of memory.

An initial mention of the file systems is provided here because the GATK performance profile influenced the choices. Place the file system metadata and data on separate storage pools. Configure the data storage pool with a larger IBM Spectrum Scale File System block size (8 MiB). Because the IBM Spectrum Scale File Systems are remotely mounted on the compute cluster from the storage services, the IBM Spectrum Scale networking must be over a low-latency and high throughput network interface.

Additional information can be read in the following document:

- *A guide to GATK4 best practice pipeline performance and optimization on the IBM OpenPOWER system*

<https://www.ibm.com/downloads/cas/ZJQDQAL>

5.2 Expanding IBM Reference Architecture for High-Performance Data Analytics into medical imaging

Medical imaging is a vital tool for diagnosis, which is creating major analytical and data challenges as more sophisticated methods are used to extract clinical insights. With deep learning supported by IBM software-defined infrastructure and high-performance computing solutions, researchers are analyzing brain scans to identify brain tumors faster and more accurately, helping physicians improve patient care.

5.2.1 Harnessing AI to transform diagnosis and treatment of brain cancer

Researchers at two universities sought to enhance magnetic resonance imaging (MRI) scan analysis, to enable physicians to use huge amounts of data generated effectively, while at the same time keeping scan times short.

Washington University St. Louis and Vanderbilt University deployed IBM solutions to accelerate the creation and deployment of deep learning models that fill in the gaps in incomplete MRI brain scans:

- Provides 20x faster training of deep learning models than traditional PC environments
- Increases speed and accuracy of diagnosis, enhancing treatments and patient outcomes
- Lowers barriers to deep learning for physicians, addressing the big data skills gap

As stated, the IBM solution trains deep learning models 20 times faster than traditional PC environments, enhancing patient outcomes through accurate, fast diagnosis. It provides physicians with an entry point to deep learning.

5.2.2 Pushing the boundaries of traditional medicine

With medical imaging tools, physicians can see inside patients' bodies without lifting a scalpel. By enabling clinicians to identify evidence of disease and injury non-invasively, it is no surprise that advances in this field have had a revolutionary impact on the ability to diagnose and treat patients.

Recent gains in computing power have made it possible to capture more detailed medical images, but making sense of the growing volumes of data is a challenge. Applying AI capabilities, and in particular deep learning, can potentially overcome these obstacles. However, deep learning is an area where skills are in short supply in every industry.

Also, capturing more detailed images often involves longer scanning times. This can be uncomfortable for patients, and it ties up hospital resources and slows down the delivery of health care. To get around this, clinical engineers have developed techniques to minimize scanning times by generating under-sampled, or incomplete, images. The trade-off is that these images can include distortions that prevent accurate diagnoses.

Researchers at the Washington University St. Louis School of Medicine and the Vanderbilt University Institute of Imaging Science wanted to bring high-performance analytics capabilities to the medical imaging fields, with low barriers to entry for physicians with little to no previous experience in this area.

Yong Wang, PhD, Assistant Professor of Gynecology and Obstetrics, Radiology, and Biomedical Engineering at Washington University St. Louis, and Xiaoyu Jiang, PhD, Research Fellow at Vanderbilt University Institute of Imaging Science, teamed up to extract insights from MRI scans more efficiently. By doing so, they created an effective method of using under-sampled scans for medical imaging.

Yong Wang picks up the story: *"We saw an opportunity to develop a predictive solution to fill in the gaps in incomplete MRI images. Because of the vast amounts of data involved, we knew that AI, and more specifically deep learning, held the key to success."*

To create a practical tool that helps rather than hinder physicians, the research team knew that the solution needed to incorporate sophisticated data analytics technologies alongside a short learning curve.

Yong Wang adds: *"Traditionally, AI innovation has been led by specialist computing engineers. But as relative newcomers to this field, we wanted to find an easy entry point for AI both for us, and for the eventual users of the solution: physicians. In other words, we needed a tool that can be used by someone with no prior coding knowledge to augment their ability to diagnose and treat patients immediately, which can be optimized once they got more familiar with the technology. On top of that, we needed a computing backbone with the processing power to provide results fast."*

5.2.3 Diving into deep learning

First on the agenda for the researchers: find an IT infrastructure that can enable them to fully exploit the value of big data. When analyzing MRI images, both patient outcomes and experiences are at stake, so insights are required as soon as possible.

Xiaoyu Jiang sums up the challenge: *"From an IT perspective, we were seeking technology that can set new records for speed, scalability, flexibility and efficiency. It needed to help us deal with data that is growing fast, in volume, variety and complexity across siloed systems. And to provide usable results in the short timelines demanded by clinical settings, we wanted to be able to apply and automate deep learning workflows."*

Washington University St. Louis School of Medicine and the Vanderbilt University Institute of Imaging Science deployed IBM Spectrum Conductor Deep Learning Impact to get started with deep learning faster. A software-defined infrastructure solution, IBM Spectrum Conductor Deep Learning Impact helps to automate and accelerate system resource management, distributed processing and prototyping. The team also implemented IBM Power System S822LC servers to provide the performance required to support AI workloads.

Xiaoyu Jiang says, “*IBM Spectrum Conductor Deep Learning Impact gives us a full workflow for deep learning, broken down step by step, making it as easy as ordering your shopping online. Even a beginner like me can use it to start exploring data immediately. It simplifies prototyping and hyperparameter tuning so you can get your models ready for production sooner. And IBM Power Architecture is the ideal platform for deep learning, allowing us to experiment at scale.*”

Using IBM Spectrum Conductor Deep Learning Impact, the research team can upload data in multiple file formats. TFRecord from TensorFlow was the researchers’ format of choice, and they uploaded vast numbers of previous brain scan TFRecord image files to the system. The team took advantage of object detection and classification features to prepare the data, before uploading training models built using Python.

“*With IBM Spectrum Conductor Deep Learning Impact, we can monitor and adjust our models in real-time, and optimize them very quickly,*” explains Xiaoyu Jiang. “*We can also keep track of hardware utilization, which enables us to share resources more effectively across the team. Once we’ve finished refining our training models, we can apply them to under-sampled images to predict what the missing parts are, and reconstruct them with a high degree of accuracy.*”

5.2.4 Giving physicians the tools to excel

Using the IBM platform, Washington University St. Louis School of Medicine and the Vanderbilt University Institute of Imaging Science developed an effective deep learning solution for the enhancement of MRI brain scans.

Xiaoyu Jiang comments: “*IBM Spectrum Computing and Power Systems solutions made short work of training our models with 1,300 MRI images, finishing in just two hours. On a traditional PC, we estimate that this would have taken 20 times as long, so the IBM technology is directly responsible for helping us innovate faster.*”

The research team’s solution increased the signal-to-noise ratio and reduced the number of artifacts in MRI images compared to existing under-sampling techniques. As a result, physicians will be able to identify cancerous brain tumors with greater precision. Yong Wang says: “*By filling in the gaps in under-sampled MRI images with help from IBM solutions, we can support better diagnosis and assessment of treatments.*”

The research team expects to reduce time-to-diagnosis for patients using its deep learning models, yielding benefits for both patients and healthcare organizations. Next, the researchers will feed more data into the network, including images featuring a variety of brain structures and tumors.

“*Supported by IBM solutions, we’ve been able to recalibrate under-sampled images so that they match the level of detail you would find in a scan that took twice the time,*” explains Xiaoyu Jiang. “*For patients, this means less time spent in MRI machines with no impact on the care they receive. It frees up healthcare equipment and staff, so they can accommodate more patients.*”

Although the research study at Washington University St. Louis School of Medicine and the Vanderbilt University Institute of Imaging Science was focused on MRI scans, the project has broader relevance. For example, a similar approach can enhance ultrasound scanning to help prevent premature births.

Yong Wang sums up: *“Combined, our research and IBM technology brings cutting-edge deep learning capabilities within reach of physicians so that they can give patients higher-quality care. We are inventing a whole new imaging system that directly enables better diagnosis and disease monitoring, so physicians can select the right treatment plans and evaluate patients’ progress. IBM provided a skilled team that supported us every step of the way, and their backing is essential to researchers like us that deal with data from hundreds of patients each and every day.”*

For more details about the case study, see 6.7, “Washington University St. Louis and Vanderbilt University” on page 62.



Case studies

This chapter illustrates the following case studies:

- ▶ Sidra Medicine
- ▶ Amsterdam UMC
- ▶ L7 Informatics
- ▶ University of Birmingham
- ▶ Thomas Jefferson University
- ▶ Biotechnology and Biomedicine Center of the Czech Academy of Sciences and Charles University: BIOCEV
- ▶ Washington University St. Louis and Vanderbilt University

6.1 Sidra Medicine

Supporting game-changing genomics research to improve the health of a nation.

“Analyzing hundreds of samples in parallel on a regular basis requires a robust HPC system to handle the load properly. From our experience, IBM systems have proven to be reliable in helping us address this technical requirement.”

—**Dr. Mohamed-Ramzi Temanni**, Manager, BioInformatics Technical Group at Sidra Medical and Research Center

Per this [announcement letter](#), IBM is collaborating and providing solutions for a compute and storage infrastructure for Sidra Medical and Research Center (Sidra). The goal of Sidra Medical and Research Center is to be a research and education institution, in addition to a world-class hospital focusing on the health and well-being of children and women.

For more information about the IBM collaboration to provide the platform that is deployed by Sidra to advance Qatar's biomedical research capabilities, see this [press release](#).

This is not the first collaboration between Sidra and IBM. One of the first programs that Sidra used from the IBM technology platform was for the [Qatar Genome Project \(QGP\)](#).

The goal of this section is to describe the Sidra Medical and Research Center Advancing Qatar's Biomedical Research Capabilities with IBM solutions. For more information about this solution, see this [YouTube video](#).

6.1.1 About Sidra

Sidra focuses on three pillars:

- ▶ Healthcare
- ▶ Education
- ▶ Biomedical research

There has been an increase in obesity, diabetes, and cardiovascular-related diseases. Sidra has discovered new insights, which is improving the treatment of patients.

Sidra had to adjust the way it approached healthcare, research, and education to accommodate the need for personalized healthcare. Genomic care is important because it can be tailored to a specific treatment that is based on the genomic signature of the patient. In the past, a treatment was generic; now, a treatment is customized for each patient.

6.1.2 The Qatar Genome Project focuses on population health and better treatments

QGP is the first initiative in the world where a sequence of the entire population is done. Sidra is the provider for the sequencing and bioinformatics analysis for this project. Until this project, only a small subset (of about 350,000 samples total) had been sequenced. Now, you can sequence 18,000 samples per year by using the sequencer at Sidra.

The QGP looks at genomic signatures and the disease that fits with that specific genomic signature to provide a customized treatment.

The goal of the QGP is to help the researcher answer complex research questions. The research starts with a hypothesis where the researcher is trying to understand the importance of the genomic factor in the prevalence of a disease. A blood sample moves from the biological realm to the digital realm through the sequencing step, and then that data is transferred to a HPC cluster. A pipeline and analysis are performed to, for example, specify what mutations are responsible or related to the disease.

6.1.3 Personalized medical advances depend on having a unified view

Genomic is not everything. Genomic is taking a snapshot of your gene at a certain point in time. But beyond the genome, there are other types of data, such as the RNA-seq, which is analyzing the transcriptome (you can also analyze the ribosome). With Ribo-seq, you also have information about the metabolome, which is where you look at the metabolite. This data tells a story about how the human body is functioning and the status of the disease. You combine this different heterogeneous data to gain a better insight about the disease.

When you work with multifactorial diseases, such as obesity, you must account for genomic data and other information, such as phenotypical or environmental data, to better understand the disease. If you do not account for these factors, it is as though you are watching a video without any sound.

For example, if the clinical data is worth one dollar, and the genomic data is worth one dollar, combining the two is worth a thousand dollars because the combined data provides much better insight into the disease than looking at the data separately. You must combine multiple sorts of data, such as combining the phenotypical and genomics data. For example, for a multifactorial disease, such as obesity, there are genetic and environmental factors, so capturing only the genetic information provides only part of the answer.

Scientists are trying to correlate the genome data and clinical data to find abnormal cases that are related to obesity or blood pressure-related activities, which is the reason to study the clinical data and genome data together in a single platform. The goal is to build a unique and integrated platform to help the researchers analyze their bioinformatics data in an easy and efficient way.

6.1.4 Converging high-performance computing, big data, and cognitive computing

The Sidra infrastructure is a national resource for other research institutions. The unified infrastructure is used for genomic workflow, big data analytics, and machine learning in a single platform.

Genomics workloads with high-performance computing

HPC addresses a “one size fits all” approach for many research requirements. An application-driven architecture helps you build genomic workflows on HPC, machine learning algorithms on big data, and image processing on a centralized infrastructure. An application-driven architecture also helps you run multi-disciplinary applications on a single infrastructure.

Genomic workloads deal with large amounts of data. Data analysis can take a couple of weeks, and if the analysis fails, you must redo the entire workload. A failure adds to the number of days to run to complete the job, which adds constraints to meeting deadlines.

There are three key elements that are used to select the best solution for HPC: scalability, support, and flexibility.

Accelerating pipelines with Apache Spark

Apache Spark is ideal for organizing big data and genomic pipelines. There are several bioinformatics tools that are integrated into Spark, such as PacBio, Partic-Floor, and Galaxy, which improve biomedical pipeline development, and also big data tools, such as Spark, that are integrated to help minimize the run time of the pipeline.

6.1.5 Why cognitive computing and IBM

Cognitive computing plays a major role in Sidra's application development. There are two cases:

- ▶ The entire querying mechanism to which a scientist provides a natural language query. This mechanism uses IBM cognitive solutions to convert this query into a technical query to the system.
- ▶ IBM cognitive solutions help the user by suggesting the best way to submit a job to the HPC workload, with the goal of efficiently using the resources to the maximum.

When the user submits an inefficient job, the number of resources are used less, which leads to inefficiency. IBM cognitive solutions help resolve those inefficiencies by providing a better way of submitting the jobs.

6.1.6 A collaboration

Sidra scientists decided to work with IBM because IBM has solid experience in engineering systems. There are many complex problems to solve, and there is a good team of scientists with Sidra who know how to deal with those problems. So, the goal of Sidra was to find a partner that collaborated with their scientists to address the problem by providing a robust solution, and by working hand-in-hand to tackle those problems on both sides.

Sidra collaborated with IBM to build centralized natural resources to address the diversified categories of applications, which include a pathogenome project, machine learning, big data analytics, and image processing. All these applications must run in a centralized infrastructure so that data can move in and around for research requirements.

The project started from the ground up. The team built the entire infrastructure in collaboration with IBM. The scientists provided input, and with their experience in combination with IBM knowledge, a new robust infrastructure was built that is now used for many projects at Sidra, and with many business partners. The Sidra infrastructure is also used as nationalized resources for other organizations.

6.1.7 Software-defined infrastructure for all data and workloads

The teams are running diversified projects, such as pathogenome projects, machine learning, image processing, and big data, in a single infrastructure (see Figure 2-1 on page 14). The customized design infrastructure is suitable for machine learning algorithms and image-processing applications.

The sources for image processing are mostly MRI scans and scanning machines, which are processed by open source and MATLAB publications. The image-processing applications are integrated with MATLAB and open source applications, so they can be processed in a single infrastructure.

The data can be moved from HPC to big data analytics. HPC addresses genome data, and big data addresses clinical data.

6.1.8 Faster results with scalability, reliability, and speed

Optimization of the scientific workload is important because pathogenome projects contain much data and many samples. Scalability plays a crucial role in pathogenome projects in terms of computing, storage, and networking.

IBM Spectrum Computing Solutions are flexible, scalable, and expandable. IBM Systems, IBM Storage, and IBM Spectrum Compute Solutions are a key combination to run pipeline and bioinformatics tools optimally. IBM Spectrum Computing Solutions provide end-to-end solutions for your research requirements.

Researchers and scientists must run more than a thousand jobs per day. Intelligent resource management systems provide scalability, quality of service (QoS), and the best turnaround time in the infrastructure. Intelligent resource management systems can be implemented in IBM Spectrum LSF. For example, in the last two years, researchers have run 700,000 genomics jobs in an IBM Spectrum LSF cluster.

By optimizing the different aspects, such as optimizing the population calling and tweaking the parameters of IBM Spectrum LSF, the researchers reduced one of the steps from 30 days to only four days.

The research team has used this solution for the last three years, and never had any failures or outages. The IBM Spectrum LSF cluster is 90 - 100% used, and the number of jobs are increasing day by day.

IBM Spectrum LSF RTM and IBM Spectrum LSF Application Center produce reports on performance metrics and job slot utilization, and many other reports that help management plan the capacity of the HPC cluster. For example, IBM Spectrum LSF RTM helps monitor the jobs across the cluster. IBM Spectrum LSF RTM is a dashboard monitoring system where the user can log in and check their own jobs that are running on the cluster, and can pull the reports, which helps their research.

IBM Spectrum LSF Application Center helps the researchers submit any jobs to the IBM Spectrum LSF cluster through a web interface. IBM Spectrum LSF Application Center is a tool where the user does not need to remember any IBM Spectrum LSF command-line arguments, and it helps any researchers to send a job to the IBM Spectrum LSF cluster.

6.1.9 Adding big data and cognitive computing to high-performance computing

IBM Spectrum Computing products integrate HPC and big data workloads in a single platform. To support the QGP, IBM Spectrum LSF provides cluster integration with current technologies, such as Docker and Open Stack, and includes integration with other big data tools.

IBM already has integrated Spark and Docker containers with IBM Spectrum LSF successfully, and is integrating IBM Spectrum Conductor with a Spark container for IBM Spectrum LSF to optimize the computer sources and applications.

IBM Spectrum Scale helps customize the genomics solution design, HPC capabilities, and cognitive computing, and integrates them in a single infrastructure. IBM Spectrum Scale is useful for data-intensive applications.

The sample solution has about 3 petabytes of data. However, IBM Spectrum Scale is a highly scalable solution, and compared to other file system options, IBM Spectrum Scale has better features and capabilities, is highly integrated, and is more stable.

Additionally, IBM Spectrum Scale RTM helps downsize the resource requirements for applications. IBM provides test fixes in a short period of time.

The genomics solution proof of concept sequenced 3,000 samples, which were successfully analyzed, and needed about 1.5 PB of storage.

6.1.10 Future

The field of precision medicine keeps evolving thanks to the revolution that is happening in biotechnology, and specifically in the world of sequencer technology, where observations show higher data generation, lower cost, and a faster turnaround time to generate the data. For scientists, the expectation of an HPC system is to have an innovative technology platform that can match the incremental data generation in the biotechnology world so that there is rapid and on-time analysis.

The Qatar biomedical informatics division is considered a national resource. This division helps other scientists and researchers from other institutes with their bioinformatics and research computing needs. In addition, all of these collaborations use the HPC system. Scientists and researchers expect HPC systems to provide innovative technology that can help them meet their needs in terms of data analysis and ever-increasing rapid data generation.

The scientists hope that the current analysis that they are doing will identify variants that cause major diseases, which will help them develop personalized and more effective treatments.

6.2 Amsterdam UMC

Enabling ground breaking research with scalable, cost-effective storage for big data.

“Thanks to our work with IBM and E-Storage, we’ve created a secure, scalable storage platform to support stakeholders across the organization.”

—Patrick Dekkers, Storage Specialist, Amsterdam UMC

6.2.1 Customer background

In 2019, VU Medical Center (VUmc) and the Academic Medical Center (AMC) joined forces as Amsterdam UMC - <https://www.vumc.com/>. The two Amsterdam academic hospitals are working together and have the same goals: keep delivering high-quality patient care, conduct ground-breaking scientific research, and provide excellent academic education.

6.2.2 Business challenge

As its unstructured data (including administrative documents, research materials, and medical images) exploded, the customer wanted the security, cost-efficiency, and scalability of a centralized storage platform.

6.2.3 Transformation

By using a centralized, scalable and flexible storage platform for big data, Amsterdam UMC is able to optimize data performance and costs through capacity planning, storage utilization, and data placement based on IBM Spectrum Scale that automatically moves data between storage systems without disrupting users or applications.

This guarantees clinicians, researchers, and users have maximum availability for any type of data based on self-provisioning with multiple levels of service (gold, silver and bronze) depending on how often the data needs to be accessed or used. Today, their archive based on IBM Tapes is spanning more than 100 years' worth of data, protected from disasters and GDPR-compliant.

6.2.4 Business benefits

The following business benefits are described from the case scenario:

- ▶ 99% faster data migrations enable IT to focus on value-added development
- ▶ 7% increase in backup frequency due to reduced complexity and increased efficiency
- ▶ Streamlines governance by migrating the organization to a centralized pool of storage

6.2.5 Solution components

The following list shows the solution components:

- ▶ IBM Spectrum Archive Enterprise Edition
- ▶ IBM Spectrum Scale
- ▶ IBM TS4500 Tape Library

You can read the full story at the following website:

<https://www.ibm.com/case-studies/vu-medical-center-research-spectrum-storage>

You can also watch the video at the following website:

<https://www.youtube.com/watch?v=ISFVscG20xU>

6.3 L7 Informatics

Building a high-performance Genomic Cloud to support ground-breaking research.

"We were able to cut the run time of one standard genome analysis pipeline down from 24 hours to just over an hour—a time saving of 96 percent."

—Chris Mueller, Founder, L7 Informatics

6.3.1 Customer background

L7 Informatics (<https://www.l7informatics.com/>) provides software and services that enable synchronized solutions for science and health. L7's novel Enterprise Science Platform (ESP) is a scientific process and data management (SPDM) solution that enables life science and healthcare companies to connect people, processes, and systems to accelerate discoveries and drive precision healthcare.

6.3.2 Business challenge

To advance our understanding of the human genome, scientists must process vast amounts of data. However, many research centers struggle to manage the immense volume of data that they generate, so that they can put it to its best use.

6.3.3 Transformation

L7 teamed up with IBM to build an HPC environment on the cloud, leveraging IBM Spectrum technology for flexible, highly scalable data storage and user-friendly workload management.

6.3.4 Business benefits

The following business benefits are described from the case scenario:

- ▶ 96% reduction in the run time of a standard genome analysis pipeline
- ▶ 1/3 the price of using commodity solutions to perform the same work at scale
- ▶ 2 weeks from conceptual design to fully-functional IBM HPC environment on the cloud

6.3.5 Solution components

The following list shows the solution components:

- ▶ IBM Cloud
- ▶ IBM Spectrum LSF
- ▶ IBM Spectrum Scale

You can read the full story at the following website:

<https://www.ibm.com/case-studies/l7-informatics-systems-spectrum-hpc>

You can also watch the video at the following website:

https://www.youtube.com/watch?time_continue=46&v=1D8CiPQYRTI

6.4 University of Birmingham

Driving innovative research forward by taking control of data.

“Breakthroughs are happening all the time at the university. Underpinning all of this pioneering innovation, IBM Spectrum Storage solutions make sure that the data is there, whenever our researchers need it.”

—**Simon Thompson**, Research Computing Infrastructure Architect, University of Birmingham

6.4.1 Customer background

Established by Queen Victoria in 1900, the University of Birmingham (<https://www.birmingham.ac.uk/index.aspx>) is one of the largest universities in the UK, serving approximately 34,000 undergraduate and graduate students.

The university's Computer Centre is the centerpiece of the Birmingham Environment for Academic Research (BEAR), a collection of IT resources available without cost to the University of Birmingham community and qualified external researchers.

6.4.2 Business challenge

To maintain its reputation as a premier research institution, the University of Birmingham must ensure that data is always available to a growing number of users running increasingly complex simulations.

6.4.3 Transformation

The university deployed IBM Spectrum Scale and IBM Spectrum Protect, increasing transparency around data's location and who accesses it, and increasing its mobility within a diverse IT environment.

6.4.4 Business benefits

The following features represent a few of the business benefits gained from the implemented solution:

- ▶ Supports compliance with data protection regulations at low cost and without disruption
- ▶ Up to an estimated 2 FTEs savings due to enhanced operational efficiency
- ▶ 5000 researchers supported by infrastructure that helps them find solutions to key issues faster

6.4.5 Solution components

The following list shows the solution components:

- ▶ IBM Spectrum Scale Data Management Edition
- ▶ IBM Spectrum Protect
- ▶ IBM Power Systems AC922
- ▶ IBM PowerAI Enterprise

You can read the full story at the following website:

<https://www.ibm.com/case-studies/university-of-birmingham-systems-software-spectrum-scale>

6.5 Thomas Jefferson University

Deepening the understanding of disease enables radically new approaches to diagnosis and treatment.

"When you let data lead the way, you can entertain bolder journeys that are not limited by what is already known in the literature. High-performance computing is the catalyst that makes such scientific explorations possible."

— **Isidore Rigoutsos, PhD**, Founding Director of the Computational Medicine Center, Thomas Jefferson University

6.5.1 Customer background

Jefferson (Philadelphia University + Thomas Jefferson University: <https://www.jefferson.edu/>) is a distinctive, comprehensive national university setting a new standard for 21st-century professional education. It has 7,800 students, more than 4,000 faculty members, and offers approximately 160 undergraduate and graduate programs on multiple campuses. Its unique Nexus Learning model focuses on collaborative, inter-professional, and trans-disciplinary approaches to learning supported by design and systems thinking, innovation, entrepreneurship, empathy, and the modes of thought central to the liberal arts and scientific inquiry.

6.5.2 Business challenge

What causes some people to develop diseases and not others? The attempt to find an answer is driving ground-breaking research, and leading pioneers to challenge traditional approaches to treatment.

6.5.3 Transformation

The Computational Medicine Center at Jefferson is breaking new ground in the understanding of disease by analyzing huge amounts of biological data with the help of high-performance computing.

6.5.4 Business benefits

The following are a few of the business benefits from the implemented solution:

- ▶ Push the boundaries of knowledge, anticipating new breakthroughs in healthcare
- ▶ Support the development of diagnostics and therapies that boost positive outcomes
- ▶ Remove barriers to scientific exploration through data-driven research

6.5.5 Solution components

The following list shows the solution components:

- ▶ IBM Spectrum Scale
- ▶ IBM Spectrum Protect
- ▶ IBM Storwize® V5030
- ▶ IBM TS3310 Tape Library

You can read the full story at the following website:

<https://www.ibm.com/case-studies/jefferson>

You can also read the story with the following link to the e-book:

<https://www.ibm.com/downloads/cas/ZBXXNGP2>

You can also watch the videos at the following website:

<https://bit.ly/2wBh4TN>

6.6 Biotechnology and Biomedicine Center of the Czech Academy of Sciences and Charles University: BIOCEV

Building research infrastructure with performance, efficiency, and reliability in its DNA. Ł

“As scientists introduce the latest generation of appliances and lab equipment, the data generated by their research activities surges, and the IBM platform ensures we can cope with even the highest demand.”

—Michal Sedláček, IT Architect, BIOCEV

6.6.1 Customer background

Biotechnology and Biomedicine Centre of the Academy of Sciences and Charles University in Vestec (BIOCEV: <https://www.biocev.eu/en>) was founded as a joint initiative from the Academy of Sciences of the Czech Republic and two faculties at Charles University in Prague. The project’s goal is to establish a European center of excellence for biomedicine and biotechnology, with the following aims: detailed study of cellular mechanisms at the molecular level, the research and development of novel therapeutic strategies, early diagnostics, biologically active agents including chemotherapeutic, protein engineering, and other technologies.

6.6.2 Business challenge

Scientific research does not end after experiment results are achieved. Instead, the outcomes must be evaluated and verified, making data storage an essential component of successful innovation. To achieve its goal of becoming a European center of excellence for biomedicine and biotechnology, BIOCEV had to help scientists store huge amounts of research data.

6.6.3 Transformation

With a software-defined storage solution from IBM, BIOCEV gained the high-performance, efficient, and reliable platform needed to support scientific breakthroughs, offering scientists fast, reliable storage and access to data. The automated storage management contributes to low TCO.

6.6.4 Business benefits

The following are a few of the business benefits from the implemented solution:

- ▶ Facilitates research by offering scientists fast, reliable storage and access to data
- ▶ Enables high efficiency through automated storage management
- ▶ Elevates the organization’s reputation by enabling non-stop services

6.6.5 Solution components

The following list shows the solution components:

- ▶ IBM Spectrum Scale
- ▶ IBM Spectrum Protect

- ▶ IBM Storwize V7000 Gen2
- ▶ IBM TS3500 Tape Library

You can read the full story at the following website:

<https://www.ibm.com/case-studies/BIOCEV>

6.7 Washington University St. Louis and Vanderbilt University

Advancing medical imaging research with deep learning.

“IBM did an excellent job to form a very skilled team to support us in a very timely manner.”

—**Dr. Yong Wang, PhD**, Assistant Professor of Gynecology and Obstetrics, Radiology, and Biomedical Engineering, Washington University St. Louis

6.7.1 Customer background

Washington University School of Medicine (<https://medicine.wustl.edu/research/>) in St. Louis is committed to advancing human health throughout the world, and has an outstanding history of biomedical research in an environment that cultivates the best minds in science and medicine. To advance medical imaging with AI, they are collaborating with the Vanderbilt Institute of Imaging Science (<https://vuiis.vumc.org/>), which is a trans-institutional initiative within Vanderbilt University serving physicians, scientists, students, and corporate affiliates.

6.7.2 Business challenge

Recent gains in computing power have made it possible to capture more detailed medical images, but making sense of the growing volumes of data is a challenge. Applying AI capabilities, and in particular, deep learning, can hold the potential to overcome these obstacles and fill in the gaps in incomplete MRI brain scans.

6.7.3 Transformation

With deep learning supported by IBM software-defined infrastructure and high-performance computing solutions, researchers are analyzing brain scans to identify brain tumors faster and more accurately, helping physicians improve patient care. The enhancement of MRI scan analysis enables physicians to effectively use large amounts of generated data, while keeping scan times short.

6.7.4 Business benefits

The following features represent a few of the business benefits gained from the implemented solution:

- ▶ 20x faster training of deep learning models than traditional PC environments
- ▶ Increased speed and accuracy of diagnosis, enhancing treatments and patient outcomes
- ▶ Lower barriers to deep learning for physicians, addressing the big data skills gap

6.7.5 Solution components

The following list shows the solution components:

- ▶ IBM Spectrum Conductor Deep Learning Impact
- ▶ IBM POWER8

For additional information about the solution implemented, and feedback about it, see the following website:

<https://ibm.co/2PsVgFJ>

You can read the full story at the following website:

<https://ibm.co/2EMbz8e>

You can also watch the informative videos at the following websites:

<https://bit.ly/2RAaqco>

<https://bit.ly/2C001zX>



Profiling GATK

This appendix provides details about the profiling of the Genome Analysis Toolkit (GATK) workflow, which is described in 5.1, “The Broad Institute Genome Analysis Toolkit (GATK)” on page 44.

Note: The product release levels indicated in this appendix were the ones used in the lab environment for this paper.

Figure A-1 shows step one of the Application Workflow, *Mapping to reference genome using BWA MEM*.

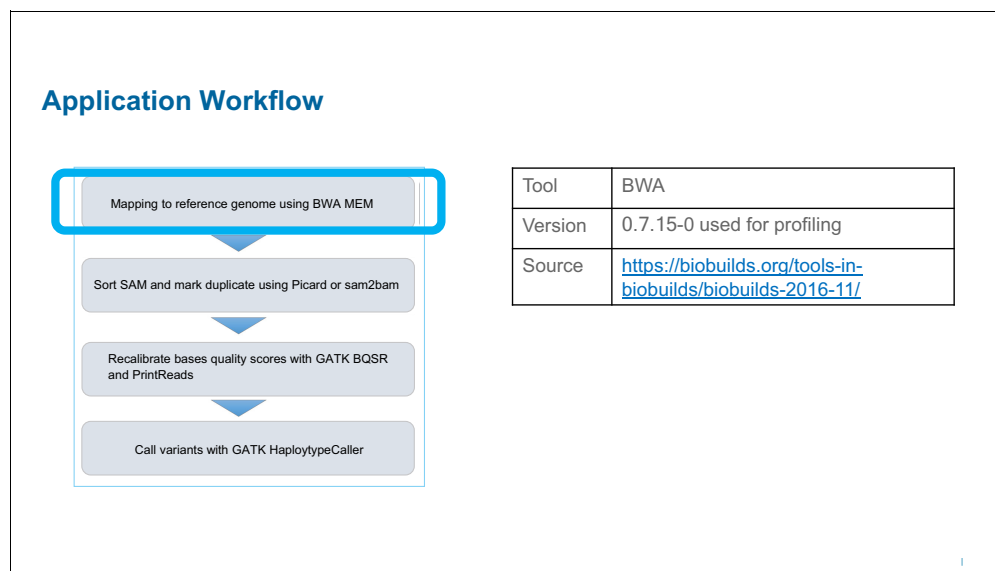


Figure A-1 Application Workflow, Mapping to reference genome using BWA MEM

Figure A-2 shows Application Profiling - BWA MEM.

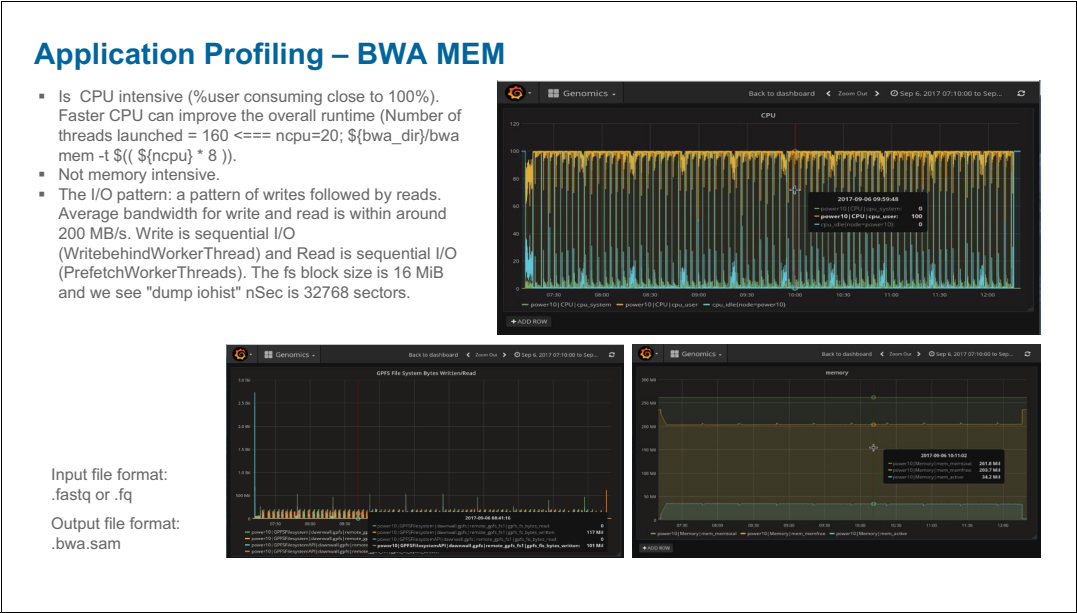


Figure A-2 Application Profiling - BWA MEM

Figure A-3 shows step two of the Application Workflow, *Sort SAM and mark duplicate using Picard or sam2bam*.

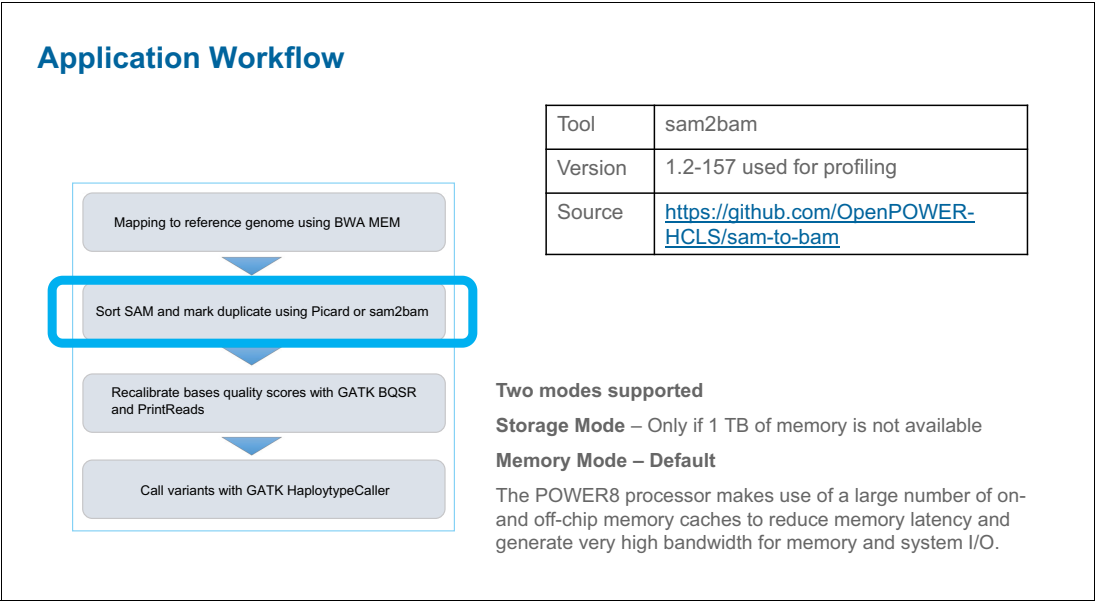


Figure A-3 Application Workflow, Sort SAM and mark duplicate using Picard or sam2bam

Figure A-4 shows Application Profiling - Sam2Bam (Storage Mode).

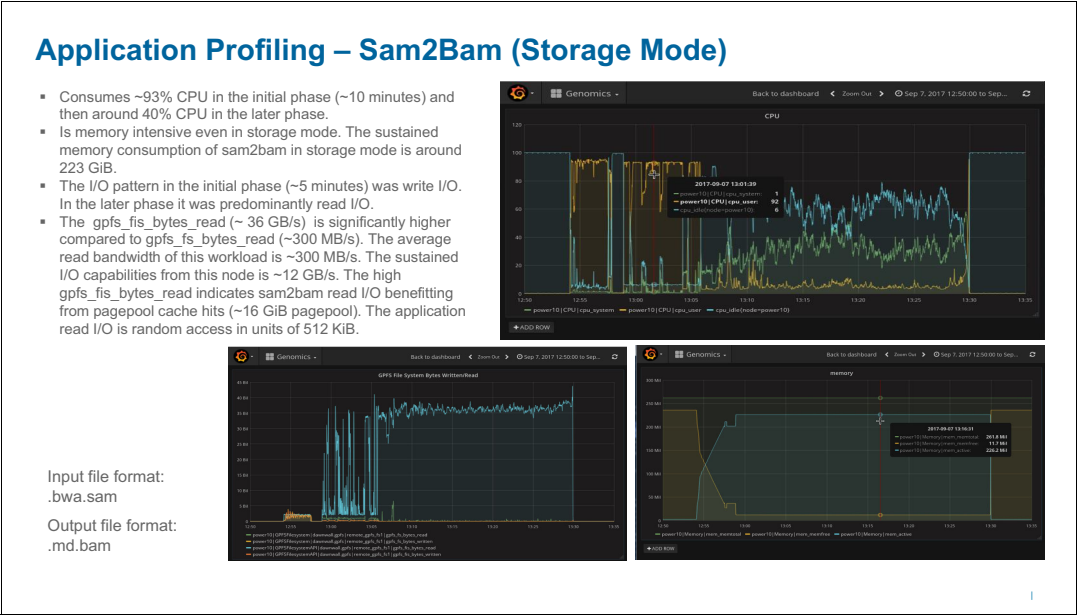


Figure A-4 Application Profiling - Sam2Bam (Storage Mode)

Figure A-5 shows step three of the Application Workflow, *Recalibrate bases quality scores with GATK BSQR and PrintReads*.

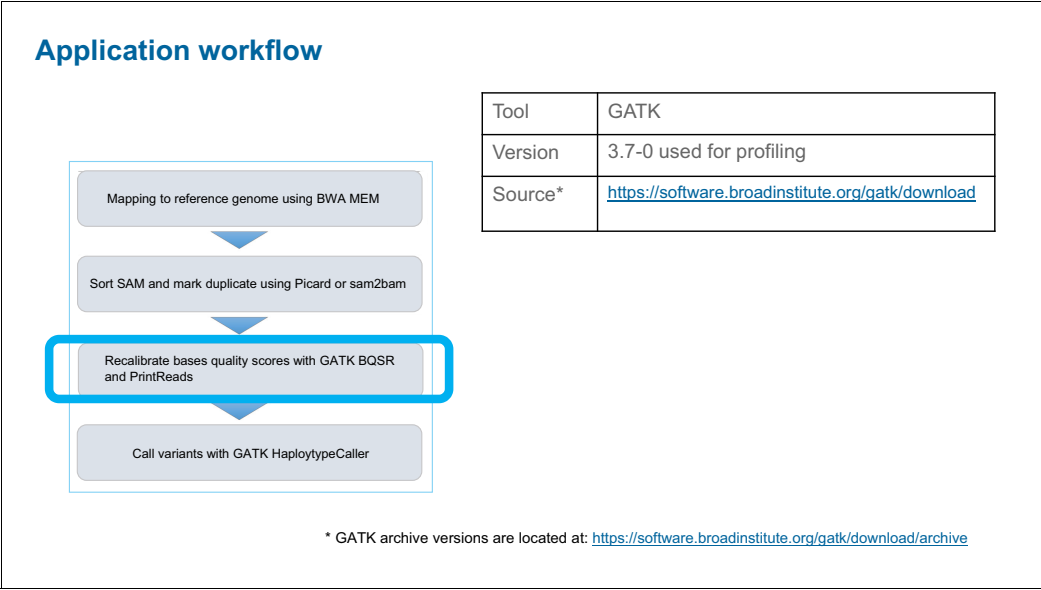


Figure A-6 shows Application Profiling - GATK BQSR, and Figure A-7 shows GATK PrintRead.

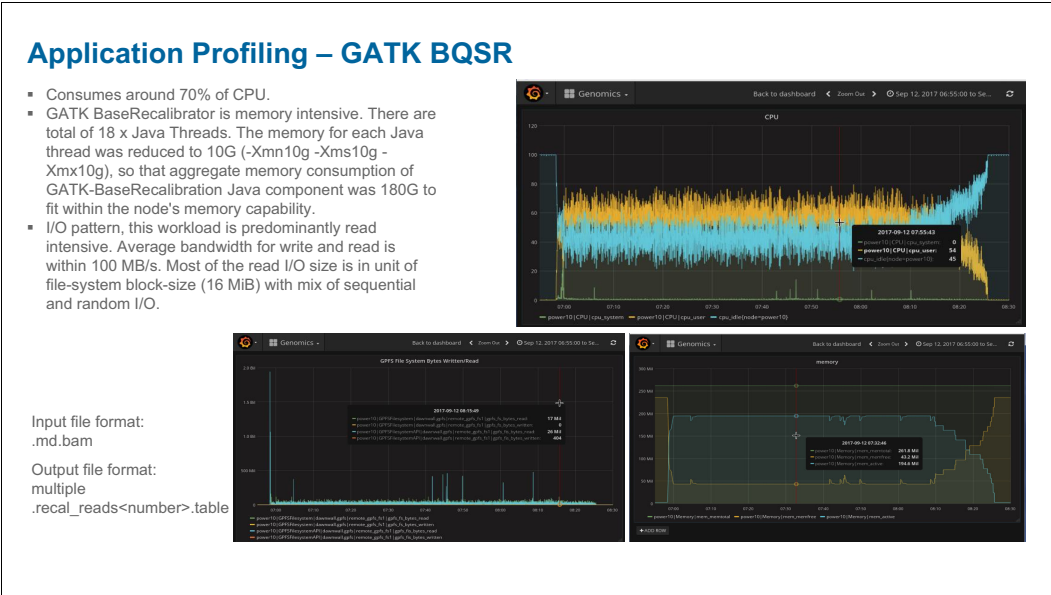


Figure A-6 Application Profiling - GATK BQR

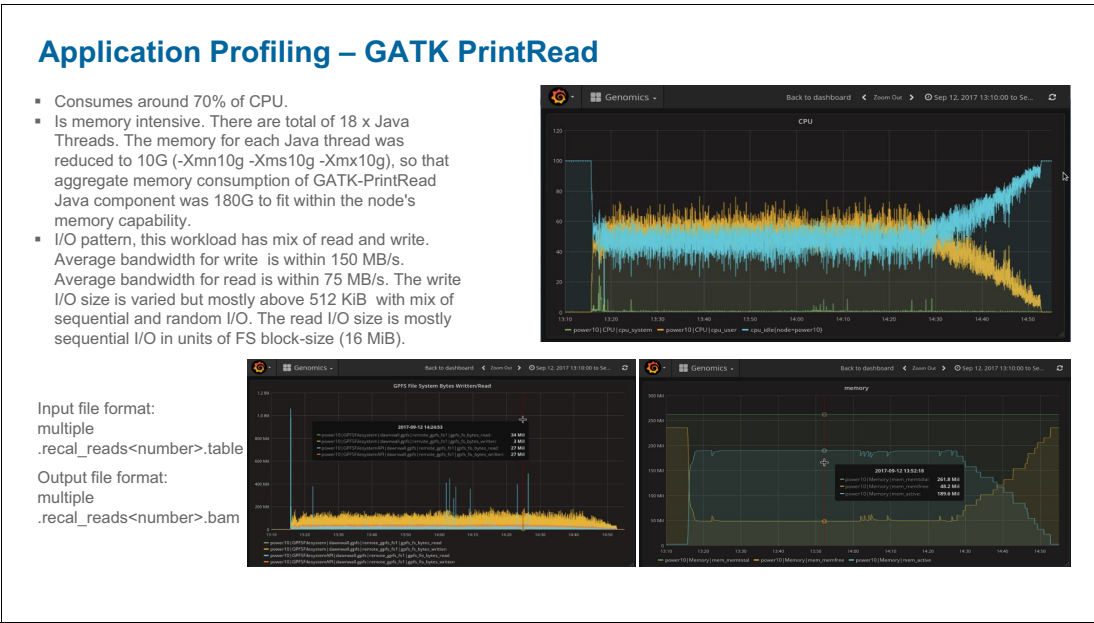


Figure A-7 Application Profiling - GATK PrintRead

Figure A-8 shows step four of the Application Workflow, *Call variants with GATK HaplotypeCaller*.

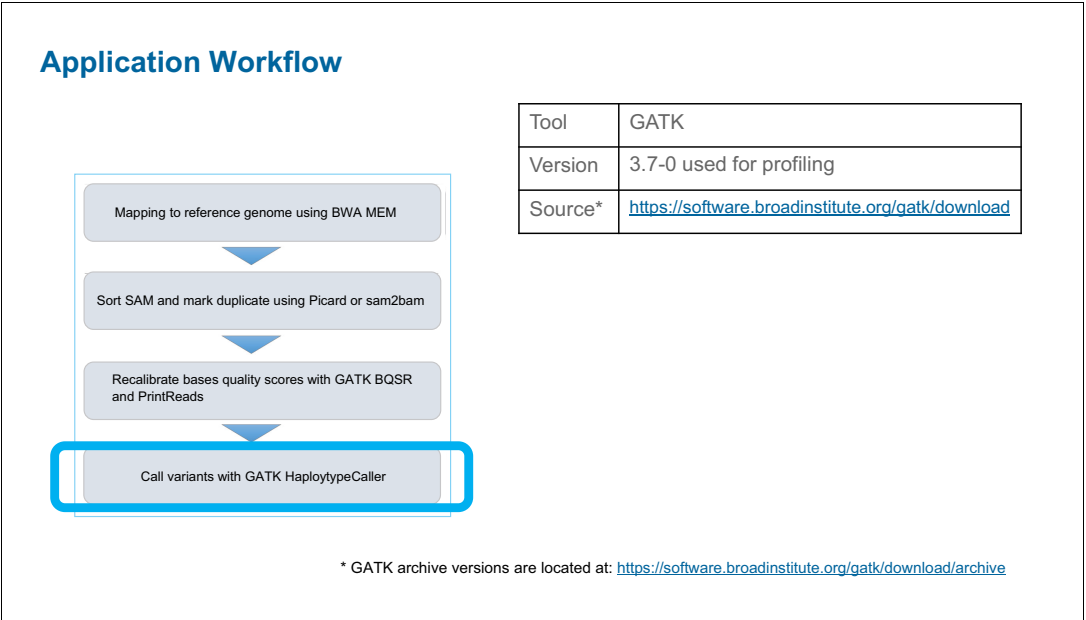
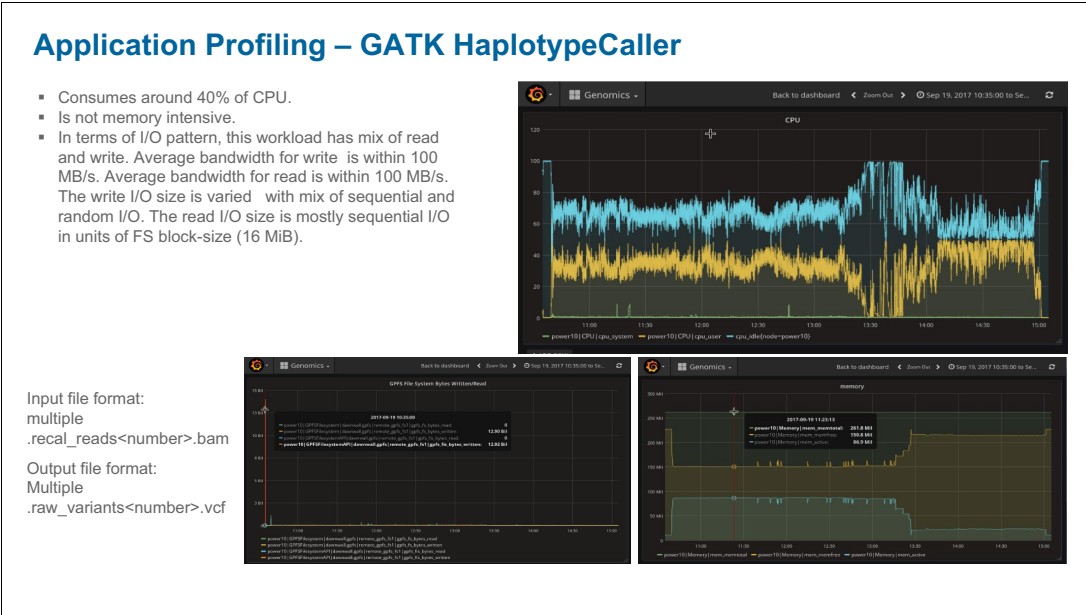


Figure A-8 Application Workflow, Call variants with GATK HaplotypeCaller

Figure A-9 shows Application Profiling - GATK HaplotypeCaller, and Figure A-10 on page 70 shows GATK MergeVCF.



Application Profiling – GATK MergeVCF

- Not CPU intensive.
- Is not memory intensive.
- I/O pattern, this workload has mix of read and write.
Average bandwidth for write is within 1.5 GB/s.
Average bandwidth for read is within 2 GB/s. The read I/O size is mostly sequential I/O in units of FS block-size (16 MiB). The write I/O size is mostly sequential I/O in units of FS block-size (16 MiB).

Input file format:
multiple
.raw_variants<number>.vcf

Output file format:
Single
.raw_variants.vcf file

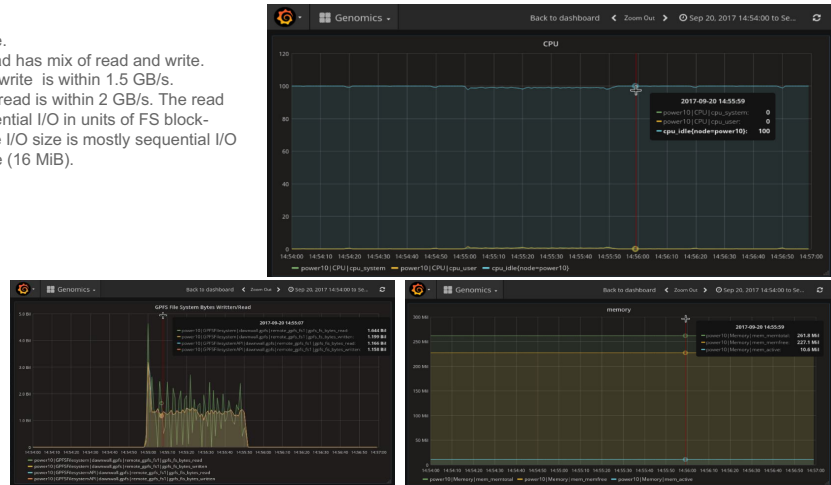


Figure A-10 Application Profiling - GATK MergeVCF

Additional information can be read in the following document:

- *A guide to GATK4 best practice pipeline performance and optimization on the IBM OpenPOWER system*

<https://www.ibm.com/downloads/cas/ZJQDOQAL>

Related publications

The publications that are listed in this section are considered suitable for a more detailed discussion of the topics that are covered in this paper.

IBM Redbooks

The following IBM Redbooks publications provide additional information about the topic in this document. Some publications referenced in this list might be available in softcopy only.

- ▶ *IBM Power System AC922 Introduction and Technical Overview*, REDP-5472
- ▶ *IBM Platform Computing Cloud Services*, REDP-5214
- ▶ *IBM Platform Computing Solutions for High Performance and Technical Computing Workloads*, SG24-8264
- ▶ *IBM Reference Architecture for Genomics, Power Systems Edition*, SG24-8279
- ▶ *IBM Reference Architecture for Genomics: Speed, Scale, Smarts*, REDP-5210
- ▶ *IBM Spectrum Computing Solutions*, SG24-8373
- ▶ *IBM Spectrum Scale Best Practices for Genomics Medicine Workloads*, REDP-5479
- ▶ *IBM Spectrum Scale (formerly GPFS)*, SG24-8254
- ▶ *Implementing IBM FlashSystem 900*, SG24-8271
- ▶ *Implementing IBM Spectrum Scale*, REDP-5254

You can search for, view, download or order these documents and other Redbooks, Redpapers, web docs, drafts, and more materials, at the following website:

ibm.com/redbooks

Online resources

These websites are also relevant as further information sources:

- ▶ A guide to GATK4 best practice pipeline performance and optimization on the IBM OpenPOWER system
<https://www.ibm.com/downloads/cas/ZJQD0QAL>
- ▶ IBM Aspera
<https://www.ibm.com/software/info/aspera/>
- ▶ IBM Cloud
<http://www.softlayer.com/>
- ▶ IBM Cloud Object Storage
<https://www.ibm.com/cloud-computing/products/storage/object-storage/cloud/>
- ▶ IBM Power Systems
<http://www.ibm.com/systems/power/>

- ▶ IBM Spectrum Computing
<http://www.ibm.com/systems/spectrum-computing/>
- ▶ IBM Spectrum Scale
<http://www.ibm.com/systems/storage/spectrum/scale/>

Help from IBM

IBM Support and downloads

ibm.com/support

IBM Global Services

ibm.com/services



REDP-5481-00

ISBN 073845690x

Printed in U.S.A.

Get connected

