

基于大数据技术的资源发现平台构建 ——以国家图书馆“文津搜索”系统为例

张红

(国家图书馆, 北京 100081)

摘要: 以国家图书馆“文津搜索”系统为例, 阐述应用Hadoop分布式系统架构和各类NoSQL数据库等大数据技术构建资源发现平台的策略和方法, 并对基于大数据技术构建的“文津搜索”系统的技术优势、实际应用成效以及技术局限性进行分析, 为大数据技术在图书馆数字资源服务领域的应用提供借鉴与参考。

关键词: 大数据技术; 资源发现系统; 国家图书馆; 文津搜索

中图分类号: G250.74

DOI: 10.3772/j.issn.1673-2286.2016.1.009

1 引言

随着信息技术的发展和互联网应用的普及, 大数据时代已然来临。大数据技术突破了传统数据管理技术的瓶颈, 对于种类多样、增长迅速、蕴藏巨大价值的海量数据, 在数据采集、存储、处理、分析和挖掘方面具有一定的技术优势。

在数字图书馆建设深入开展的今天, 数字资源急剧增长, 越来越多的图书馆推出了资源发现服务平台, 为读者提供统一的检索途径, 使读者能够方便、快捷地发现并获取图书馆的数字资源。目前, 国际市场上主流的统一资源发现系统有Serials Solutions公司的Summon系统、Ex Libris公司的Primo Central系统、EBSCO公司的EBSCO DiscoveryService (简称EDS) 系统、OCLC的WorldCat Local系统以及Innovative Interfaces公司的Encore系统^[1], 这些系统既有内容提供商推出的, 也有系统提供商研制的, 均被全球众多图书馆所引进。国内的图书馆有些引进了Summon、Primo等国外的产品, 还有些采用了南京大学数图实验室和EBSCO公司联合研发的Find+知识发现平台, 也有些采用超星公司的超星发现系统^[2]等。

2012年年底, 国家图书馆推出了自己的资源发现系统——“文津搜索”系统, 作为国家数字图书馆的资源检索门户正式向公众提供服务。考虑到项目的建设目

标、个性化需求以及建设成果的推广复用等方面因素, 国家图书馆没有采用商业化的成品软件, 而是联合软件开发商自主研发了“文津搜索”系统。不同于采用传统关系型数据库技术构建的资源发现与获取系统, “文津搜索”系统引入了Hadoop^[3]分布式系统架构和各类NoSQL^[4]数据库等大数据技术, 从而实践了大数据技术在图书馆资源服务领域的应用。

由于各图书馆资源发现服务平台的建设目标和设计理念不尽相同, 各平台采用的系统架构和数据管理技术也各具特点, 呈现出的服务效果也各有所长。任何计算机技术都有其适用领域, 源自开源模式的大数据技术亦是如此, 该技术在其适用领域彰显出技术优势的同时, 在某些方面也存在一定的技术局限性。本文以“文津搜索”系统为例, 探讨大数据技术在资源发现平台中的应用方法, 并结合平台实际运行情况, 分析基于大数据技术构建的资源发现平台所具有的技术优势、应用成效以及受到的制约, 为大数据技术在此领域的应用提供借鉴与参考。

2 建设需求与构建策略

2.1 建设需求

国家图书馆在构建资源发现平台时, 在软件方面的

功能需求主要包括:提供一站式数字资源和实体文献检索服务,支持分类检索、高级检索、全文检索和二次检索,支持中英文检索词翻译、简繁体通检以及同义词扩展检索;对检索结果可通过多种途径的分类和排序方式进行过滤、聚合与导引,方便读者快速定位所需信息。

除上述对软件功能的需求外,国家图书馆对其资源发现平台在数据整合与挖掘、系统性能与架构、与其他系统的衔接以及推广复用等方面还提出了一些个性化需求。

在数据整合与挖掘方面,需对国家图书馆通过自建、购买、征集、采集等方式所获取的各种类型的数字资源和传统文献资源元数据进行整合,实现元数据的本地存储,并构建元数据的集中索引库,以实现高效的元数据检索服务;需采集包括互联网资源在内的相关数据以提高检索和服务质量,并基于统计分析和数据挖掘技术,向用户提供高质量的检索结果排名、相关检索和检索建议等。

在系统性能与架构方面,需支持每分钟10万次并发和亚秒级响应,满足大并发、低功耗以及动态可伸缩的要求,可随着资源数据量和用户访问量的增加方便地进行扩展。

在与其他系统的衔接方面,需与国家图书馆统一用户管理系统集成,实现与国图其他读者服务系统之间的单点登录以及资源访问权限控制;需提供到馆内外资源发布系统的链接,实现自动认证并能够直接跳转到资源的详细显示页面,方便读者获取对象数据;对实

体文献,需提供到各馆OPAC的馆藏链接以及到国图馆际互借与文献传递系统的接口。

在推广复用方面,需提供多种复用方式,可作为数字图书馆工程建设成果被推广复用到地方图书馆。

2.2 构建策略

国家图书馆在构建“文津搜索”系统服务平台时,基于实时高效的检索、处理海量数据、便于扩展和推广复用等方面的考虑,选定了能够支持分布式服务和大规模数据处理的系统架构,通过前端服务平台、数据存储与处理平台以及多级互联的搜索集群三个核心组件协同运转,来实现系统的预期目标,如图1所示。

前端服务平台承载着检索服务软件各功能的实现,负责接收读者发送的资源检索、获取资源详细信息、个性化设置等各类请求,并提供合理、灵活的检索结果展示页,方便读者获取所需资源。平台采用多种类型的NoSQL^[4]数据库存储各类需向读者提供的数据以及系统运行中产生的历史数据,采用特定的数据组织方式、算法和软件技术提高检索质量。平台集成国图统一用户管理系统、馆内外资源发布系统以及国图馆际互借与文献传递系统的接口,整合国图和地方馆OPAC的馆藏链接地址数据,可以实现单点登录、用户认证、到对象数据发布系统和各馆OPAC系统的跳转,并可发送馆际互借与文献传递请求。

多级互联的搜索集群接收前端服务平台发来的检

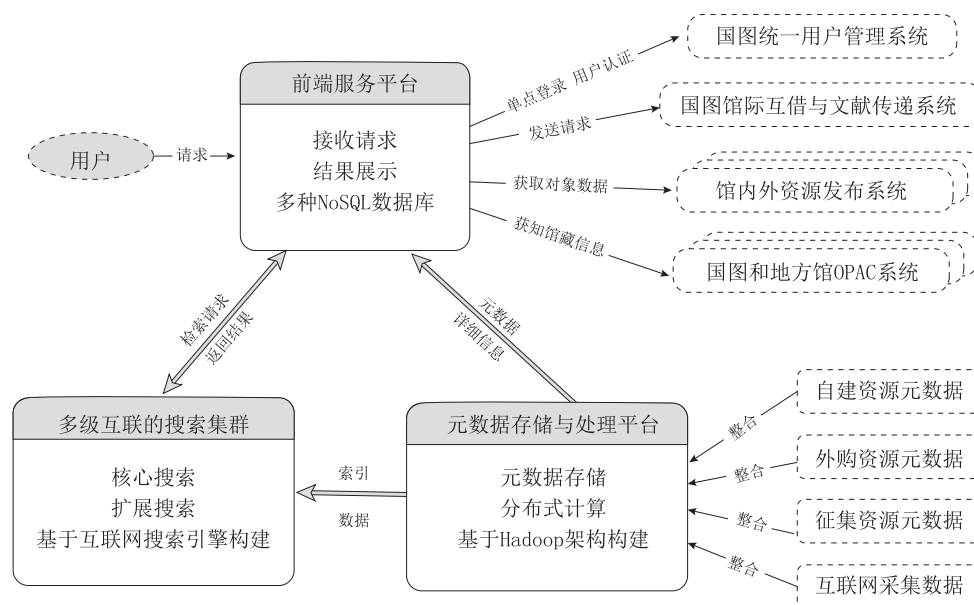


图1 文津搜索核心组件示意图

索请求, 提供元数据核心检索和扩展检索服务, 并将检索结果返回给前端服务平台。搜索集群采用互联网上成熟的搜索引擎技术构建, 通过负载均衡支持大并发, 通过预索引的方式, 将索引文件预先装载到搜索集群内存或SSD闪存中以提高检索速度, 从而满足系统既定的性能指标。

数据存储与处理平台是基于Hadoop架构构建的,负责将收集到的各种来源、不同格式的元数据按照映射规则整合成统一格式的元数据,实现各类数字资源元数据的本地存储,并接收网络资源采集系统采集到的互联网上开放的相关数据以丰富、优化检索结果。该平台同时承担着大规模分布式计算任务,负责索引构建以及资源重要性计算,并能够对收集的日志进行分析、统计和挖掘,其结果用于改善读者服务。该平台将整合好的元数据详细信息转换成特定格式提供给前端服务平台用于详细信息展示,将抽取的索引数据提供给搜索集群用于元数据检索。

设计三种复用模式,即前端接口方式、缩小版整体部署方式和云平台应用方式用于“文津搜索”系统在地方馆的推广复用。

3 大数据技术及应用

资源发现系统的建设涉及硬件平台、软件环境、网络设施、数据整合规则、系统功能、安全策略等各方面技术,本文仅就与大数据相关技术的应用进行阐述。

国家图书馆的“文津搜索”系统服务平台采用分布式系统架构,用近400台PC服务器组成服务器集群,在各核心组件上部署了定制开发的应用软件以实现特定功能。其系统结构如图2所示。

3.1 基于NoSQL数据库的前端服务平台

(1) 前端服务平台构建

前端反向代理是前端服务平台的入口,负责接收读者发来的请求,并向读者返回请求结果,通过Nginx^[5]中间件进行反向代理和负载均衡,同时提供静态文件的缓存。

Web Server和Media Server采用多个Tomcat^[6]中间件,接收前端反向代理发来的请求,并返回请求结果,实现高可用负载均衡。Web Server将检索请求发

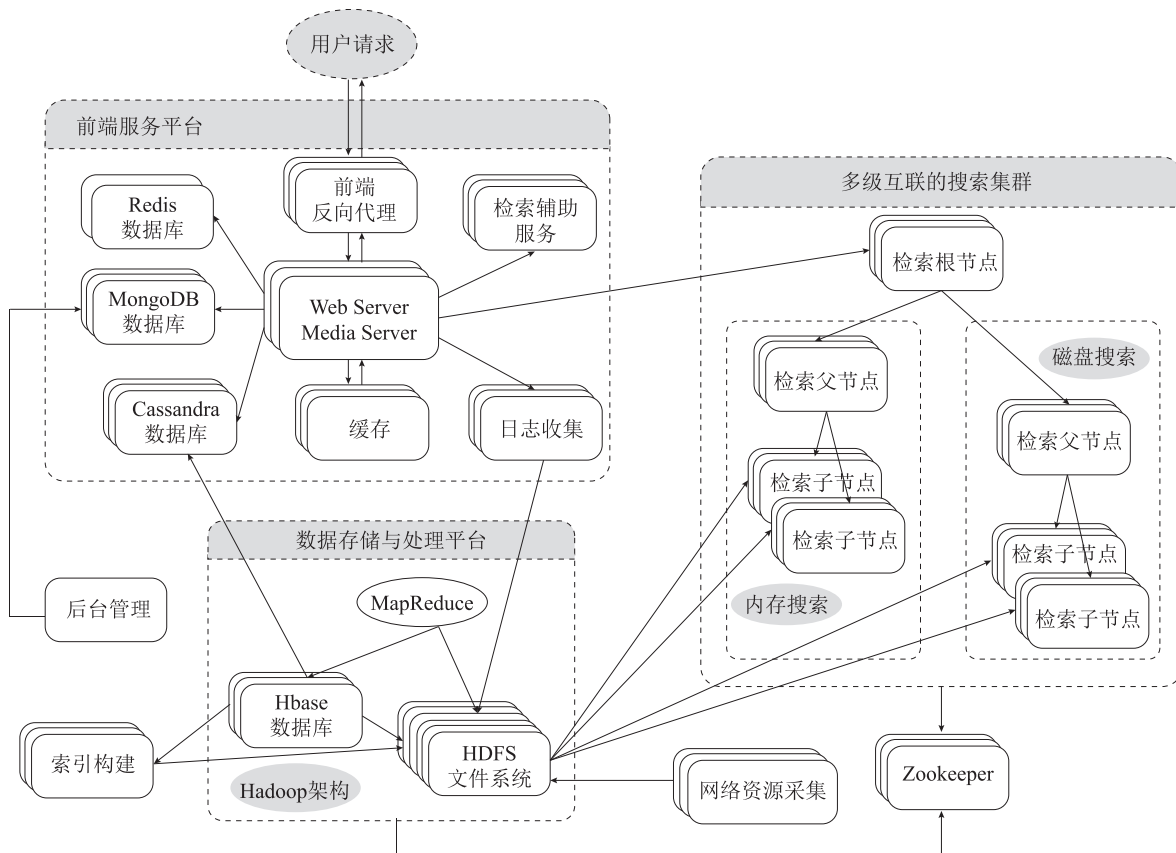


图2 文津搜索系统结构图

送到搜索集群进行检索并获取检索结果,检索辅助服务器为其提供相关检索和检索推荐服务,Redis^[7]、MongoDB^[8]和Cassandra^[9]三种NoSQL数据库负责记录历史信息、保存个人设置信息以及向读者提供检索结果展示数据。

日志收集服务定时收集Web Server上的日志信息,并发送到数据存储与处理平台用于对日志的统计、分析和挖掘。

缓存服务器部署了Memcached^[10]分布式内存对象缓存系统,用来存储数据库检索结果、相关文件和图片等,可向Web Server和Media Server提供缓存的数据,从而提高检索速度。

(2) 前端服务数据管理

前端服务平台采用三种NoSQL数据库存储与读者交互的数据信息。NoSQL数据库是开源型数据库,具有海量数据的处理能力,在低成本、高性能和易于扩展方面具有一定的优势^[11],因而近年来发展非常迅速。

前端服务平台的读者检索历史信息和翻译词库信息保存在Redis数据库中。Redis是一个键值模型的内存数据库,性能非常出色,但因受物理内存的限制,存储的数据量和扩展能力也会受到限制^[12]。由于读者的检索词、检索串信息和翻译词库信息需要快速读写,数据量不大且所需存储空间可以预见,因此采用Redis数据库比较适合。

前端服务平台的检索热词信息、标签云信息以及读者个人设置信息保存在MongoDB数据库中。此外,该数据库还存储了后台管理系统的用户信息、设置信息以及各类统计信息。MongoDB是一种可扩展、高性能、面向文档的数据库,支持松散的数据结构,数据查询功能强大,支持建立索引,能够解决海量数据的访问效率问题^[12]。由于在读者访问时系统需频繁地读写读者信息和标签云信息,管理员通过后台管理系统进行各种操作也需要得到实时响应,因而选择MongoDB数据库存储相关信息可以满足数据读写的高效性。

前端服务平台检索结果的详细显示数据从Cassandra数据库调取。Cassandra是一套高性能、可伸缩的分布式数据管理系统,可实现纯粹意义上的水平扩展和分布式的写操作,具有较强的容灾能力^[12]。Cassandra数据库中预先存储了包含书封和URL信息的元数据、图片数据以及书评信息。基于Cassandra数据库多数据节点的设计,元数据的详细信息被写到多个节点上,在读数据时可路由到其中一个节点进行读取,从

而保证存储的数据安全可靠且能够快速读取。

3.2 基于Hadoop架构的数据存储与处理平台

数据存储与数据处理平台是基于开源Hadoop分布式系统架构构建的。Hadoop是一个能够对大量数据进行分布式处理的软件框架,具有高可靠性、高可扩展性和高容错性的优点,其技术核心是分布式文件系统HDFS和MapReduce引擎^[13]。

数据存储与数据处理平台采用数十台服务器部署了HDFS分布式文件系统,用以存储需进行数据处理的文件 and 数据处理后的结果文件,包括从互联网采集到的信息、收集到的日志文件、元数据整合中的过程文件、生成的引文件等。HDFS包含名字节点和诸多数据节点,名字节点负责管理文件系统命名空间和对文件的访问,数据节点用于管理自身节点上的数据存储。HDFS可跨服务器存储海量文件,它将大文件分成多个同样大小的数据块,并根据策略将数据块复制到相同及不同机架的数据存储节点中,从而保证平台的高容错性和数据读取的高效性。HDFS具有数据的错误检测和快速自动恢复功能,可以在数据节点间动态迁移数据,从而保证数据的高可靠性。HDFS可以实现完全的水平扩展,增加一台新服务器到集群里,就可以增加更多容量,不需要手动迁移任何数据,就可以实现资源的自动分配,从而保证平台的高可扩展性。

数据存储与数据处理平台的大规模并行数据处理、数据分析和数据挖掘的功能是基于MapReduce编程模型实现的,主要包括五个方面:①对灌装到Hadoop集群中的原始元数据进行数据清洗、数据查重、数据转换、数据合并、挂接书封目次信息以及建立数据关联等一系列数据处理;②对从互联网采集到的信息进行处理,完成资源重要性计算,以提高检索质量;③在索引构建服务的控制下进行核心信息索引和全文信息索引的构建,并将索引结果文件存入Hadoop的HDFS分布式文件系统中;④对收集到的日志数据进行统计分析生成统计数据;⑤对用户的行为数据进行分析挖掘,作为用户相关推荐的数据支撑。

数据存储与数据处理平台采用HBase^[14]数据库存储灌装后的元数据、整合后标准化的元数据、书封数据以及到数字资源发布系统的URL信息。HBase数据库也是一种NoSQL开源数据库,它是构建在Hadoop上的分布式的、面向列的数据库,有很好的扩展性,可以

存储数以亿计的行,可利用成熟的HDFS文件系统和MapReduce计算框架,高效和高可靠地处理海量数据^[12],其数据库的一致性服务由ZooKeeper^[15]提供。由于存储的元数据囊括各种文献类型,且不同类型的元数据描述项存在很大差异,而HBase数据库是稀疏存储,同一张表里的每一行数据都可以有截然不同的列,因而非常适合存储结构复杂的海量元数据。此外,存储于Hbase中的元数据可以非常方便地通过MapReduce程序进行数据处理,并可以将数据以SSTable格式导出、生成索引数据,导入Cassandra数据库提供显示数据,这些都有效地利用了Hbase数据库的优势。

3.3 基于互联网搜索引擎构建的搜索集群

多级互联的搜索集群是基于互联网成熟的搜索引擎构建的,是一个分布式搜索引擎架构,分为提供核心信息检索服务的内存搜索集群和提供扩展信息(含全文信息)检索服务的SSD磁盘搜索集群两部分,每一部分包含若干组结构相同的服务器,每一组服务器又分为检索根节点、检索父节点和检索子节点三个层级,集群中根节点、父节点和子节点之间的同步服务、配置维护和命名服务也是由ZooKeeper提供的。通过这种多级互联的方式实现负载均衡,保证高并发情况下的检索效率。

以内存搜索集群为例,每一组服务器含有一个检索根节点,每一个检索根节点下联3个检索父节点,每一个检索父节点下联20个检索子节点。数据存储与处理平台HDFS中保存的索引文件被预先分发到各组服务器,在每组检索子节点服务器的硬盘中平均分布,即每组检索子节点的服务器中都保存了一套全量的索引文件。当服务器启动时,索引文件自动加载到内存中提供检索,因而每组检索子节点服务器的内存容量之和也应当能够容纳全量的索引文件。当检索根节点获得检索请求后,会通过检索父节点把检索请求信息发送到每一个检索子节点中进行检索,检索子节点将检索结果及每个结果的重要性分数返回检索父节点,通过检索父节点和检索根节点进行归并和排序,最后由检索根节点合并为一个相关排序表返还给前端展示页面。从上述检索过程可以看出,核心搜索是通过分布式并行搜索技术实现的,而且是在内存中进行检索,加之在各层级上都构建了缓存系统,因而大大提高了检索速度。

4 应用成效与制约

4.1 应用成效分析

国家图书馆在推出“文津搜索”系统之前,曾采用艾利贝斯公司的产品MetaLib构建国图的数字资源门户,该系统运行在小型机上,采用的是联邦检索的模式,其检索速度受各远程资源提供系统以及网络链路的影响很大,服务效果不甚理想。

“文津搜索”系统是基于元数据集中的理念建设的,通过将所有元数据都整合到本地的方式,有效地提高了检索速度,改善了服务效果。在2012年年底上线时,“文津搜索”系统整合了近2亿条书目元数据和目次数据。随着原有数据库的数据更新和新数据库的陆续整合,至2015年6月系统中的元数据量已达到近3亿条,并且还在不断增加。本着优先整合能够提供对象数据的资源库元数据的原则,除数据来自于联合编目系统的传统文献外,系统中绝大部分的数字资源都能提供到对象数据的访问链接。“文津搜索”系统自上线以来,向读者提供了安全、稳定、高效的数字资源统一检索服务,在实际运行过程中,平台采用的大数据技术在以下方面体现出了技术优势,取得了良好的应用成效。

(1) 检索速度快。系统在设计时力求从各个环节提高检索速度,从高可用负载均衡及缓存技术,到分布式多级互联搜索架构,从元数据集中预索引,到内存搜索和SSD闪存搜索,共同保证了大并发情况下对海量资源检索的快速响应。系统的设计目标是具备每分钟10万次以及峰值每秒钟1万次的检索请求处理能力,该性能指标在系统上线前项目验收的性能测试以及系统上线后的第三方软件测评机构的性能测试中均得到了验证。在之后的读者服务中,实际的用户访问量远低于每分钟10万次的高并发量,系统在检索性能方面的优势也得到了很好的发挥,能够做到即搜即得。

(2) 动态可伸缩设计,方便扩展。系统充分利用了分布式架构的技术优势,各个核心组件都可以根据需求单独扩展。前端服务平台在16台服务器上部署了64个页面节点,可以根据读者访问量调整页面节点数。内存搜索集群为10组服务器,可以根据检索并发量增加或减少服务器组数,每组服务器中的检索子节点为10台大内存服务器,共同承担全量索引数据的存储,可以根据索引数据量的多少调整服务器的配置。目前,前端服务平台和数据存储与处理平台的硬件资源还很充足,搜

索集群子节点的内存资源已接近饱和,将采取硬件扩充的方式扩展其容量。

(3) 分析挖掘用户行为数据,实现个性化推荐。系统每天将前端服务平台记录的用户行为日志收集到数据存储与处理平台中集中保存,目前已经收集了自上线以来一千多天约为120G的用户行为日志。系统定期对收集到的读者行为日志进行分析挖掘,通过为用户的各种行为设定相应的权重,根据用户的行为轨迹生成用户对资源评价的数据模型,并基于协同过滤推荐机制实现针对用户和针对资源两方面的相关资源推荐。

(4) 数据处理和索引构建不影响读者服务。系统设计为前端服务平台与数据存储与处理平台分离,读者服务与数据处理、索引构建采用不同的核心组件支撑,因而在数据处理和索引构建时不会影响读者服务。在向搜索集群更新构建好的索引文件时,采用分批分组更新的模式,也不会造成服务的中断。

(5) 系统和数据的安全可靠性高,系统运行稳定。在系统部署方面,近400台服务器单独组网,与其他业务系统物理隔离,通过一台跳板机进行内部管理,采用密码口令和服务端密钥文件相结合的方式访问控制,从而增强了系统的安全性。此外,在系统部署时,充分利用了分布式架构的优点,在各个模块的重要节点都部署了冗余服务器,使得个别服务器故障不会影响系统运行,从而增强了系统的可靠性。在数据安全方面,对于元数据、索引数据、页面显示数据以及用户行为数据等各种类型的数据,系统都采用了分布式、冗余的数据存储方案,各类数据都能够自动进行备份,从而保证数据的安全可靠。此外,为应对不断进行数据更新存在的风险,系统保留了历次更新的索引数据,一旦更新数据存在问题,可以快速回退到上一次更新数据,从而保证在短时间内恢复服务。系统的安全可靠性在近3年的实际运行中得到检验,系统自上线以来,在全年开馆日全天候地提供了持续稳定的读者服务。

4.2 存在的制约

(1) 不支持元数据的实时在线修改。由于Hadoop比较适合一次写入和多次读取数据的应用,根据系统的设计理念,基于Hadoop的元数据仓储中的元数据是通过批量灌装的方式导入的,经过一系列数据处理生成统一格式的元数据保存在Hbase数据库中。而提供服务的简要显示元数据来自搜索集群中保存的索引数

据,详细显示元数据来自Cassandra数据库中特有格式的数据,它们分别是Hbase数据库中的元数据经过抽取索引和格式转换得到的。因此,只要修改元数据就需要经过数据导入、数据整合和数据转换全过程,如果修改的是核心数据部分,还需要重抽索引,才能向读者提供更新后的数据。因此,与存储于单一关系型数据库中的数据不同,“文津搜索”系统的数据不支持在线实时修改,只适于批量更新。

(2) URL链接信息未实现参数化配置。为向读者提供到对象数据的链接以及到地方馆OPAC系统的馆藏链接,数据存储与处理平台对这些链接的URL信息进行了整合。为提高服务效率,并考虑到Hadoop能够快速数据处理的优势,数据存储与处理平台将URL信息整合到每条元数据中进行保存,而未采用参数化配置的方式,因而URL信息的修改不够灵活,一旦链接地址发生变化,需对整个资源库的URL数据或该地方馆的馆藏链接数据逐条进行更新。

(3) 不适用于小型图书馆单独部署。“文津搜索”系统是针对提供海量资源检索服务的需求设计的,其分布式集群平台的搭建最少也需要数十台服务器,内存搜索集群等服务器还有大内存的要求,即使采用造价相对低廉的PC服务器,也需要一定的硬件开销。因此,“文津搜索”系统设计的三种推广复用模式中,缩小版整体部署方式不适用于资源有限的小型图书馆。对于这种类型的图书馆,建议采用前端接口方式或云平台应用方式。

5 结语

目前,大数据技术已经被越来越多的行业用来处理海量数据,它能够有效地利用分布式和并行技术,结合各种NoSQL数据库,解决海量数据带来的数据采集、数据存储、数据处理和数据挖掘等问题,不仅展现出高性能和高可用性的优势,而且具有良好的可扩展性和容错性。随着图书馆数字资源量的迅速增长,在采用传统数据管理技术遇到瓶颈时,可以考虑采用大数据技术,它为海量资源的数据管理提供了全新的解决方案。在图书馆资源服务领域应用大数据技术时,既要充分发挥其技术优势,又要针对其技术局限性制定合理的应对方案,必要时还可以与传统数据管理技术相结合,实现对海量数字资源科学有效的管理,力争为广大读者提供优质高效的数字资源服务。

参考文献

- [1] 包凌,蒋颖.图书馆统一资源发现系统的比较研究[J].情报资料工作,2012,(5):67-72.
- [2] 孙宇,张磊,刘伟.图书馆资源发现系统选型研究[J].图书馆杂志,2013,32(12):63-70.
- [3] Hadoop [EB/OL]. [2015-10-01]. <http://hadoop.apache.org/>.
- [4] NoSQL [EB/OL]. [2015-10-01]. <http://nosql-database.org/>.
- [5] Nginx [EB/OL]. [2015-10-01]. <http://nginx.org/en/>.
- [6] Apache Tomcat [EB/OL]. [2015-10-01]. <http://tomcat.apache.org/>.
- [7] Redis [EB/OL]. [2015-10-01]. <http://redis.io/>.
- [8] MongoDB for GIANT Ideas [EB/OL]. [2015-10-01]. <https://www.mongodb.org/>.
- [9] Cassandra [EB/OL]. [2015-10-01]. <http://cassandra.apache.org/>.
- [10] Memcached [EB/OL]. [2015-10-01]. <http://memcached.org/>.
- [11] 李冯筱,罗高松.NoSQL理论体系及应用[J].电信科学,2012,28(12):23-30.
- [12] 陆嘉恒.大数据挑战与NoSQL数据库技术[M].北京:电子工业出版社,2013:49-53,175,293.
- [13] deRoos D, Eaton C, Lapis G, et al. Understanding Big Data: Analytics for Enterprise Class Hadoop and Streaming Data [M]. McGraw Hill, 2011: 54-55.
- [14] Apache HBase [EB/OL]. [2015-10-01]. <http://hbase.apache.org/>.
- [15] Apache ZooKeeper [EB/OL]. [2015-10-01]. <http://zookeeper.apache.org/>.

作者简介

张红,女,1965年生,国家图书馆高级工程师,研究方向:数字图书馆,E-mail: hong@nlc.cn。

Application of Big Data Technology in Building Resource Discovery Platform: Taking Wenjin Retrieval System as an Example

ZHANG Hong
(National Library of China, Beijing 100081, China)

Abstract: This paper illustrates the strategy and methodology of constructing resource discovery system applying big data technology, which involves Hadoop Distributed Framework and several NoSQL databases by taking Wenjin Retrieval System of National Library of China as an example. It analyses the technological advantages, practical appliance effectiveness as well as the limitation of Wenjin Retrieval System which is structured based on big data technology, and provides reference for the further utilization of big data technology in the field of digital resource service of libraries.

Keywords: Big Data; Resource Discovery System; National Library of China; Wenjin Retrieval System

(收稿日期: 2015-12-28)