

HiBase: 一种基于分层式索引的高效 HBase 查询技术与系统

葛微^{1,3)}, 罗圣美²⁾, 周文辉^{1,3)}, 赵頔^{1,3)}, 唐云^{1,3)}, 周娟^{1,3)}, 曲文武²⁾, 袁春风^{1,3)}, 黄宜华^{1,3)}

¹⁾ 南京大学 计算机软件新技术国家重点实验室 南京 中国 210046

²⁾ 中兴通讯股份有限公司 南京市雨花台区软件大道 50 号 南京 中国 210012

³⁾ 江苏省软件新技术与产业化协同创新中心 南京 中国 210046

摘 要 Hadoop HBase 系统为大数据的存储管理提供了一种具有高可扩展性的技术方法和系统平台。然而 HBase 不支持非主键索引,导致 HBase 的数据查询效率较低,难以满足数据实时/准实时查询需求。本文研究提出了一种基于分层式 HBase 非主键索引的查询模型和方法,包括基于 HBase 的持久性索引、基于分布式内存的索引热点数据缓存技术和高效的热度累积缓存替换策略,并实现于分层式索引和查询系统 HiBase。在千万至十亿条记录规模数据集上的测试结果表明,HiBase 总体查询性能比标准 HBase 快 300 多倍(大结果集)到 1.7 万倍(小结果集),比开源的 Hindex 系统快 5~20 倍。

关键词 HBase; 非主键索引; 查询处理; 索引热点缓存; 缓存替换策略

中图法分类号 TP311

HiBase: A Hierarchical Indexing Mechanism and System for Efficient HBase Query

GE Wei^{1,3)}, LUO Sheng-Mei²⁾, ZHOU Wen-Hui^{1,3)}, ZHAO Di^{1,3)}, TANG Yun^{1,3)}, ZHOU Juan^{1,3)}, QU Wen-Wu²⁾,
YUAN Chun-Feng^{1,3)}, HUANG Yi-Hua^{1,3)}

¹⁾(State Key Laboratory for Novel Software Technology at Nanjing University, Nanjing 210046)

²⁾(ZTE Corporation, NO. 50, Software Avenue, Yuhuatai District, Nanjing 210012)

³⁾(Collaborative Innovation Center for Novel Software Technology and Industry of Jiangsu Province, Nanjing 210046)

Abstract Hadoop HBase provides a technical method and system with excellent scalability for storing and querying big data. However, HBase only provides the row key indexing and does not support non-key indexing, which makes it not efficient enough to meet the need of realtime or near-realtime applications. In this paper, we propose a hierarchical secondary indexing model and method for HBase. It builds the secondary index for non-key columns in HBase table to speed up the query process. Furthermore, based on the distributed memory, we present a hot index caching method for the secondary index in HBase, plus an efficient cache replacement policy, the Hotscore Algorithm, to reduce the disk access overhead for index data. Based on the above techniques, we implemented a hierarchical indexing system HiBase. The results from experiments on datasets

收稿日期: 年-月-日*投稿时不填写此项*; 最终修改稿收到日期: 年-月-日 *投稿时不填写此项*. 本课题得到国家自然科学基金优秀重点实验室专项基金项目(No. 61223003)资助, 同时得到中兴通讯研发基金资助. 葛微, 女, 1979年生, 博士生, 计算机学会(CCF)会员(E200035709M), 主要研究领域为查询处理、查询优化、分布式和并行计算, E-mail: gloria.w.ge@gmail.com. 罗圣美, 男, 1971年生, 硕士学位, 高级工程师, 计算机学会(CCF)会员(E2000160404M), 主要研究领域为云计算、云存储、大数据等技术领域, E-mail: luo.shengmei@zte.com.cn. 周文辉, 男, 1987年生, 硕士, 主要研究领域为并行计算、大数据分布式一致性, E-mail: cnzhwh@qq.com. 袁春风, 女, 1963年生, 教授, 计算机学会(CCF)高级会员, 主要研究领域为体系结构与并行计算、多媒体文档处理、Web信息抽取与集成等, Email: cfyuan@nju.edu.cn. 黄宜华, 通讯作者, 男, 1962年生, 教授, 计算机学会(CCF)高级会员(14302S), 中国计算机学会大数据专家委员会委员, 主要研究领域为大数据处理技术、计算机体系结构、分布式与并行计算, E-mail: yhuang@nju.edu.cn.

第1作者手机号码: 18260089153, E-mail: gloria.w.ge@gmail.com

ranging from 10 million to one billion records show that, the HiBase cold query (the cache-missed query) outperforms native HBase by 65 times (for large returned result sets) to more than 3000 times (for small returned result sets) respectively. Further, the HiBase hot query (the cache-hit query) after adopting the Hotscore Algorithm-based index caching mechanism can achieve extra 5-20 times speedup compared to the HiBase cold query, making the overall performance speedup more than 300 times (for large returned result sets) to 17,000 times (for small returned result sets) compared to the native HBase and also speedup 5-20 times compared to the open-source Hindex secondary indexing system.

Key words HBase; Secondary Index; Query Processing; Hierarchical Index; Cache

1 引言

继 Google 在 2006 年发表论文阐述 BigTable[1] 数据模型后, 为了有效应对海量数据的存储与查询管理, 人们越来越多地用 NoSQL 分布式数据存储系统替代关系数据库管理系统(RDBMS)。在迫切的企业大数据应用需求推动下, 目前出现了一些 NoSQL 非关系数据存储系统, 例如, Apache 社区的 HBase[2], Facebook 的 Cassandra[3], Amazon 的 Dynamo[4], 以及支持高效数据查询的内存数据存储系统 Redis[5]等等。这些数据存储都采用了 key-value 数据模型。key-value 数据存储系统中, HBase 的使用最为广泛。

HBase 在主键上建立了类 B+树索引, 可以高效地支持基于主键的快速数据查询。然而, 由于 HBase 缺少非主键索引能力, 在面对以非主键为条件的查询请求时, 需要对全表进行扫描, 导致查询效率低下, 难以满足需要快速响应的数据查询和统计分析场景。在传统关系数据库领域, 面对这样的问题, 通常会借助索引来解决。为了提高非主键查询响应时间, 减少查询开销, 本文研究提出了一种基于 HBase 的非主键索引和查询模型与技术。

随着内存性能的持续提升和价格的不断下降, 内存在支持实时大数据处理的商业需求中扮演着越来越重要的角色。目前, 用内存计算提升大数据处理性能已成为普遍公认的发展方向。在大数据查询分析过程中, 数据访问通常具有一定的 80/20 分布特征: 即 20% 的数据通常会被频繁访问, 我们把这些数据称之为“热数据”; 而 80% 的数据很少被访问, 我们称之为“冷数据”。为了提高查询性能, 低延迟、高带宽的内存被普遍用做热数据缓存来填补 CPU 和磁盘之间的性能鸿沟。

为了优化 HBase 上非主键属性的查询性能, 本

文研究实现了基于 HBase 的分层式索引和查询系统 HiBase (Hierarchical-Indexed HBase)。在索引存储模型上, HiBase 分为两层: 1) 基于 HBase 的持久化存储层。用来存储 HBase 用户表对应的所有非主键索引数据; 2) 基于分布式内存的缓存层。为了进一步提高索引访问的速度, 本文研究提出了一种基于分布式内存的索引热点数据缓存替换策略和技术, 将部分索引热点数据缓存在内存中, 大幅提高数据查询访问的效率。在 HiBase 中, 我们进一步研究实现了热度累积的缓存替换策略, 它克服了 LRU 只关注不同数据块最近一次访问时间、不关注不同数据块访问频繁程度不同的算法局限性, 考虑数据访问的累积热度, 并对早前时间的数据热度累积进行指数衰减, 从而更准确地捕获数据访问的特征。

本文以下内容的组织结构如下: 第 2 节介绍相关工作, 第 3 节给出了 HiBase 的非主键索引及其存储模型, 第 4 节提出了 HiBase 索引热点数据缓存替换策略, 第 5 节阐述了 HiBase 的系统设计与实现, 第 6 节对实验结果进行了分析评估, 第 7 节是总结与展望。

2 相关工作

通常, 检索 HBase 表中的数据有三种方式: 指定单个行键查询、指定行键的范围查询、以及扫描(Scan)。HBase 以字节数组的字典序对行键进行排序(这种存储格式更适合于记录的顺序读取, 充分利用磁盘传输带宽), 支持高效的指定行键的单个查询和指定行键范围的范围查询。扫描操作主要用于对非主键属性的查询, 它允许指定扫描的数据范围, 可以在查询过程中指定限定条件来过滤目标结果集。扫描操作允许对任意一个数据属性进行查询, 因而比基于行键的查询灵活, 但是对非主键属性扫描操作的代价很大。基于行键检索的时间复杂性是 $O(\log N)$, 甚至如果使用 Bloom

Filter 可以达到 $O(1)$ ，而扫描操作的时间复杂性是 $O(N)$ 。大数据应用的数据规模在千万至亿级以上，对全表进行非主键扫描的时间开销是不可接受的，这导致了 HBase 在很多实际应用中难以满足高实时性的查询需求。为此，一些研究工作致力于改善 HBase 非主键查询的性能。

开源系统 Hindex[6] 是华为公司研发的基于 HBase 的非主键索引机制。Hindex 通过对 HBase 表的每个 Region 建立单独的索引表来提高非主键查询的效率。查询请求发送到每个 Region 服务器，然后通过 Region 上的索引表将结果集过滤并返回。Hindex 查询需要访问所有的 Region，由于绝大多数检索任务的目标记录数量相对较少，在分布式集群中并行地执行该任务往往造成很多未存储任何目标记录的计算节点也触发了检索过程，而最终却返回空集。在检索任务频繁的情况下，这一并行执行过程将会耗费大量不必要的计算资源，最终将降低系统的并发查询吞吐量和系统的查询效率。

NGDATA 公司的 HBase-indexer[7] 也提供了 HBase 上的非主键索引。HBase-indexer 将 HBase 的更新数据异步发送到索引服务器上，在索引服务器上分析数据并生成对应的索引数据。索引服务器会定期将索引数据推送到 SolrCloud[8] 服务集群上。查询则通过访问 Solr 服务来定位 HBase 上的内容。这种索引机制周期性地对索引进行异步更新，索引的时效性稍差，在面向实时应用时难以有效地满足需求。

Interval Index[9] 是希腊 Patras 大学提出的基于 HBase 的非主键范围查询索引，其索引通过 MapReduce 构造 Segment Tree 索引结构并保存在内存中，同时将索引树中存储的每段范围的上界保存在 HBase 表中，以支持有效的范围查询。Interval Index 的可扩展性好，但是对单点查询，线段树退化成二叉树，索引的空间开销和维护代价都很大。

在大数据的内存缓存研究上，TBF[10] 在 HBase 数据存储上提出了基于固态硬盘(Solid State Disk)的缓存替换策略。TBF 算法结合 CLOCK 和 Bloom Filter 的优点，充分利用 HBase 数据存储行键有序的特点，有效地降低缓存替换算法的两个元数据（“数据记录是否被缓存”和“最近访问信息”）的空间开销。TBF 算法的空间效率高，但它无法避免 CLOCK 算法固有的缺陷，不能量化数据

在最近被访问的频率，并且不能保存一个 CLOCK 周期以外的数据访问信息。

微软在文献[11]中提出了如何区分持续数据访问中的冷热数据。通过指数平滑算法，将数据记录的访问频率量化。由于针对的是离线日志的数据访问分析，所以微软在这篇文章里提出了回溯(backward)算法和并行回溯算法，能够通过海量的数据访问日志高效地区分出冷热数据。但此回溯算法只适用于日志数据的离线分析，因为只有这样才能高效地从新时间到旧时间逆向分析日志。

3 分层式索引存储机制

3.1 非主键持久化索引存储模型

为了避免对 HBase 非主键查询时的全表扫描，提供快速的非主键查询能力，HiBase 为保存在 HBase 用户表中的非主键属性建立索引表，并将索引表保存在 HBase 中，借助 HBase 的特性获得良好的可扩展性和容错性。每个 HiBase 索引表用来存储管理用户表中的某个待查询非主键属性的索引。我们为用户表中待建立索引的非主键属性定义如下格式的索引表主键：

<用户表索引列名（简短别名），用户表索引列值，用户表主键>

其中，用户表索引列名是用户表中被索引属性的名称。通过将列名映射到一个简短的别名，如“Age”映射到“a”，可减少索引表主键存储空间的开销。在索引表主键中存储用户表主键有两个作用：一是保证了索引表主键的唯一性；二是提供 HBase 用户表中被索引记录的地址，通过用户表主键，可快速获得用户表中被索引的记录。

表 1 HiBase 的用户表和索引表结构

用户表			索引表	
Name	Age	Income	key	value
Bob	21	3000	a,12,Tom	{Income:500}
Jerry	30	10000	a,21,Bob	{Income:3000}
Ron	36	20000	a,30,Jerry	{Income:10000}
Simon	42	22000	a,36,Ron	{Income:20000}
Tom	12	500	a,42,Simon	{Income:22000}

表 1 给出了一个 HBase 用户表及以 Age 属性为索引的索引模型和索引表示例。该示例中，索引表的主键形如“a, 12, Tom”，其中，a 为 Age 属性简短别名；12 是用户表数据记录 Tom 的 Age 值，也就是索引列值；Tom 是该记录在用户表中对应

的主键。查询时通过比较索引列值定位记录。与关系数据库中的部分索引 (partial index) 一样，HiBase 以查询属性为索引表主键，索引表并不是包含用户表的全部字段，通常只将查询中可能需要辅助性访问的字段存放在索引表的非主键属性中。这是为了方便对查询中需要访问的辅助性字段提供快读访问，避免再次访问用户表而导致的二次磁盘访问。本例中，索引表中的 value 部分所保存的如 {Income:500}，表示查询 Age 时可能需要访问 Income 列。

3.2 分层式索引存储模型与索引缓存

上述索引表将为 HBase 表实现索引数据的持久化存储。由于索引数据是存放在 HBase 中，每次查询访问 HBase 表会涉及到很多磁盘访问，我们进一步考虑把索引中那些访问频度高的索引数据作为热点数据缓存在内存中，形成基于 HBase 和分布式内存的分层式索引存储和查询机制，进一步提高索引表查询速度。

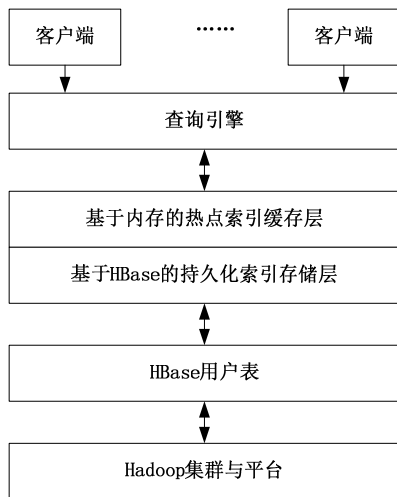


图 1 HiBase 分层式索引存储模型

HiBase 中的内存热点索引数据缓存格式不同于持久化存储中的索引格式，内存缓存索引的主键格式为：

<用户表索引列名（简短别名），用户表索引列值>

其中，用户表索引列名和用户表索引列值的含义与持久化索引存储层中的相同。内存索引构建的基本思路类似于倒排索引，内存索引缓存层中的每个索引主键对应着一个具有相同索引列值的索引记录集合，该集合包含了与该索引值对应的所有索引表数据记录。与持久化索引存储层一

样，集合中也包含了可能需要访问的其他非主键属性。因此，完整的内存索引数据格式如下：

索引主键：<用户表索引列名（简短别名），用户表索引列值>

索引集合：{<用户表主键，{<频繁访问列名，频繁访问列值>}>}

图 2 是内存缓存层的索引结构示例，该例中缓存了 Age 属性下的两个年龄值“21”和“30”，分别对应于索引内存缓存层的主键“a, 21”和“a, 30”，其中，a 是 Age 的简短别名。当以年龄为 30 进行查询时，根据索引内存缓存层主键“a, 30”的哈希值找到对应索引在内存中的存储地址，可以得到该年龄值所对应的索引集合 {<Jerry, {Income:10000}>, <Ron, {Income:20000}>}，其中包括了该年龄记录对应的用户名字和收入。由该集合对应的用户表主键到用户表中可获得相应的用户数据记录。在实现上索引内存缓存层数据存储在 Redis 内存数据库中，并由 Redis 自动完成上述哈希和快速查询过程，这一层全内存化的查询和更新都相当高效。

分层式索引存储模型基本的查询过程是：首先到索引内存缓存层查询热点索引数据，若缓存中不存在该记录，则将查询转发到索引持久化存储层进行检索。可以看出，通过将索引热点数据缓存在内存中，部分查询可以直接在内存中命中结果集，从而降低了磁盘访问开销，提高整体查询性能，这对于具有倾斜的数据访问分布特性的应用来说尤为有效。

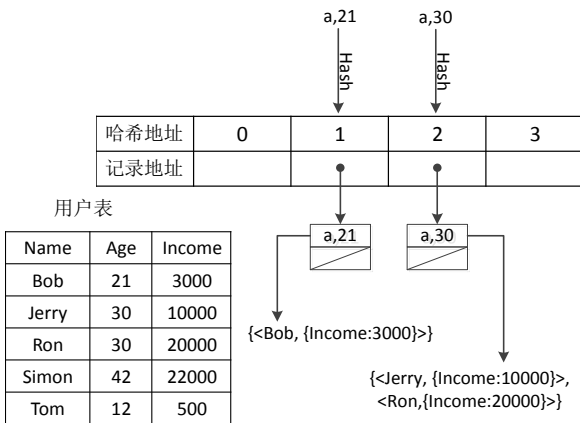


图 2 内存缓存层的索引结构示意图

3.3 基于一致性哈希的分布式内存缓存

HiBase 利用 HBase 服务器节点上的分布式内存来存储管理所有的索引热点数据。为了在分布式内存环境下提供有效的索引存储管理，我们引

入了一致性哈希[12]来完成索引热点数据在分布式内存中的存储管理。一致性哈希在许多分布式系统中得到了广泛的应用，如 Dynamo[4]，Cassandra[3]。

在 HiBase 分布式内存缓存中，使用一致性哈希来确定数据所在的服务器节点。在节点发生变化时（如节点失效或节点加入），只有和变化节点相邻的节点数据需要迁移，从而可减少节点的加入和退出带来的计算和数据传输开销。例如，当图 3(a)中存储节点 2 出现故障时，原先映射到该节点的数据会映射到顺时针方向遇到的第一个存储节点上，如图 3(b)所示部分数据被重新映射到存储节点 3。而增加存储节点的时候，数据映射关系的变化与上述的过程正好相反，新节点将会负责管理其哈希地址到它逆时针方向上第一个存储节点之间的所有数据。

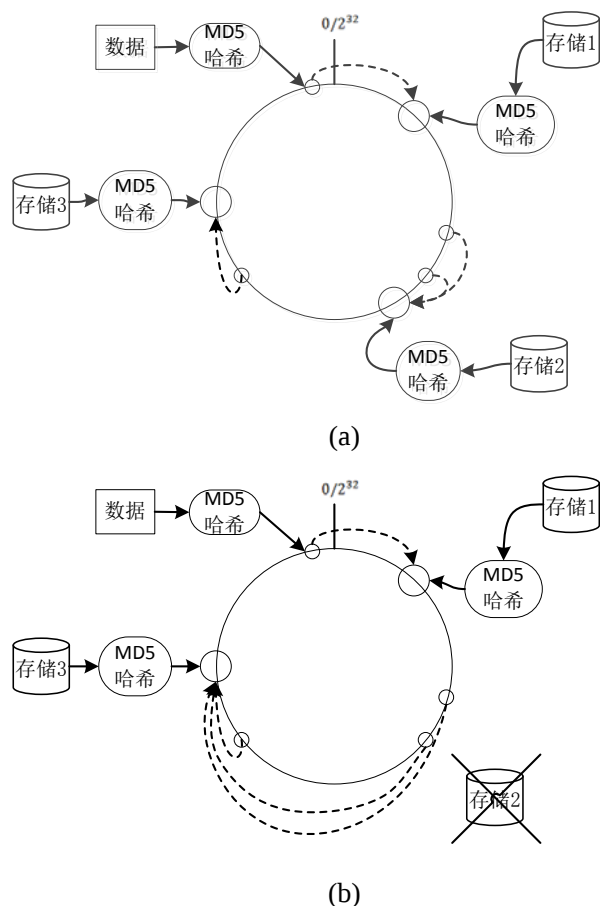


图 3 基于一致性哈希的分布式索引内存缓存存储机制

可以看出，当存储节点发生变化时，只需要迁移小部分的数据即可。而且通过将存储节点映射到圆环的伪随机分布点上，能够有效地保证各存储节点之间的负载均衡。

HiBase 的内存缓存需要通过两次哈希来找到对应索引记录的真实位置：第一次通过图 3 所示的一致性哈希找到索引数据所在的服务器节点；第二次则通过 Redis 的哈希机制找到节点内的索引数据地址。

4 索引热点数据缓存替换策略

缓存是数据查询优化常用和行之有效的方法，缓存技术一直是数据存储和查询的研究热点，通过将部分数据保存内存中，可以提高随后对数据的重复访问效率。通常缓存的容量远远小于保存全部数据的磁盘数据库的容量，所以当缓存满了之后需要选择合适的牺牲者淘汰出缓存，这就是缓存替换策略。

HBase 也有自己的缓存策略，即 HBase 的块缓存(BlockCache)。HBase 在读数据时会以块(Block)为单位进行缓存，用来提升读性能，块的默认大小是 64KB。读请求先到写缓存(MemStore)中查询，查询失败则继续到块缓存中查询，再查询失败就会到磁盘上的 HBase 表中读取，同时会把读取的结果放入块缓存。在配置文件中会指定块缓存的容量，当块缓存的容量已经被使用 85%以上的时候开始采用缓存替换策略淘汰数据块，淘汰过程在块缓存容量占用比为 75%的时候停止。然而，HBase 的缓存策略是面向系统的读写性能提升的，而不是面向特定的数据表的。

HBase 采用的是最近最少使用(LRU, Least Recently Used)的缓存替换策略来淘汰块。LRU 算法是基于如下假定：最近被访问的数据在近期最有可能被重复访问，而如果数据块很长时间未被使用，则在最近的未来被访问的概率很低。LRU 算法只关注不同数据块最近一次访问时间的差异，没有关注不同数据块访问的频繁程度不同，也即是说，LRU 算法没有考虑数据的累积热度。

4.1 HiBase索引热点数据缓存基本思想

在 HiBase 中，我们提出了热度累积的缓存替换策略，其基本设计思想是周期性地累积记录被访问的次数。

在内存中缓存的索引热点数据基于 Redis 的集合(Set)存储，Redis Set 也是以<key,value>格式来组织数据。索引热点数据的索引主键（如 3.2 节描述）做 Redis Set 的 key，而索引集合（如 3.2 节描述）作为 Redis Set 的 value 保存在内存缓存中。显然，

具有相同索引列值的记录被绑定在同一个集合中，基于索引列值的查询命中是以集合为单位的。同时，它们也是热度累积的基本单位，每个集合都会累积它在一个计算周期内的访问次数。

热度累积的缓存替换策略基于与 LRU 算法相同的假设：最近被访问的数据在最近的未来最有可能被重复访问。算法周期性地计算集合的累积热度，对所有的记录累积热度排序，选择累积热度 TOP-K 的索引记录缓存到内存中，这就是热度累积的缓存替换策略。

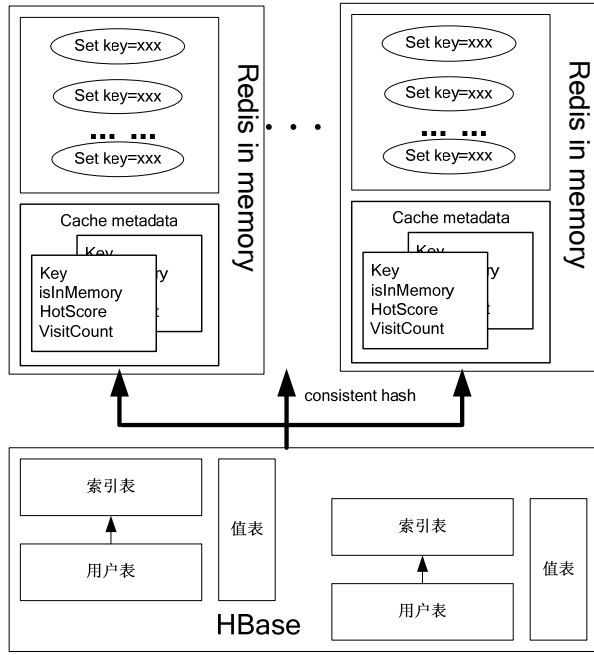


图4 内存中保存的缓存元数据

4.2 热度累积的缓存替换策略

在 HiBase 中，热度累积缓存替换策略的热度计算公式如下：

$$score_n = \alpha \times \frac{visitCount}{countPeriod} + (1 - \alpha) \times score_{n-1}$$

其中 $0 \leq \alpha \leq 1$ 。公式中的 $countPeriod$ 即热度计算周期， $visitCount$ 指当前热度计算周期中，该索引集合被访问的次数。历史热度 $score_{n-1}$ 则反映集合累积的历史热度。参数 α 是衰减系数，用来确定当前周期累积的热度和历史热度在 $score_n$ 中各自所占的权重。 α 越大，则最近的访问在数据访问热度中所占的权重越大，历史访问记录对数据热度的影响越小，反之亦然。集合的历史热度在本计算周期内以系数 $(1 - \alpha)$ 的速率衰减，经过多次迭代，更早计算周期的累积热度经过了更多次

衰减，所以早期的累积热度对数据热度的影响不断降低。

为了降低热度计算带来的计算和更新开销，在执行查询请求时，内存缓存的服务进程将会对访问到的每条索引数据记录本周期内的访问次数，此时并不对内存缓存的数据进行替换。直到查询请求次数达到 $countPeriod$ ，即到达热度计算周期时，服务进程触发缓存的更新替换。按照热度累积公式对所有的记录计算热度，根据热度排序，将热度排序 TOP-K 的集合记录缓存到内存中，集合中包含的记录条数是不固定的，所以选择 TOP-K 时，根据缓存空间能够容纳的记录条数限制计算出热度门限，高于门限的集合被缓存到内存中。

然而，在系统初始阶段，缓存是大量空闲的。LRU 算法在系统初始阶段的命中率提升很快，这是由于 LRU 算法中数据记录是访问即进入缓存，最长时间没有被访问的数据记录会在缓存充满后被淘汰。所以 LRU 可以快速进入稳定状态。而热度累积的访问如果在系统初始阶段通过周期性地计算热度，等被访问数据记录的热度累积到门限时才可以进入缓存的话，初始阶段预热代价大。所以我们的热度累积算法在缓存空闲阶段做了优化，只要缓存有空闲，我们就采用“访问即进入”的策略，将所有访问到的记录都插入缓存。而当缓存充满以后，热度累积的缓存替换策略根据记录的热度累积评分选择“牺牲者”淘汰出内存，选择获得热度高分的记录保存在缓存中。

热度累积的缓存替换策略不仅考虑了数据的访问时间远近，同时考虑了数据的访问频率，所以比 LRU 更准确。从实验结果看出，热度累积的缓存替换策略明显优于 LRU 算法，和不使用内存缓存策略相比，可以提升 5-15 倍的查询性能。

5 HiBase 的系统设计与实现

5.1 系统总体架构

基于以上的非主键索引模型和技术方法，我们设计并实现了一个基于 HBase 的分层式非主键索引查询系统 HiBase，该系统实现了基于 HBase 的持久化索引存储和基于内存的索引热点数据缓存，并使用热度累积的缓存替换策略。HiBase 提供了基于分层式非主键索引的数据查询方法，并对范围查询进

行了并行化优化。

HiBase 将整个分层式索引存储系统划分为以下几个模块，系统功能模块划分如图 5 所示。

1) 索引构建管理模块。管理索引的元数据(记录用户表对应的索引表名称、索引列等信息)，并实现针对 HBase 的流式数据和静态数据两种不同特性数据的索引构建方法，包括支持索引表和值表的插入、删除、更新操作。

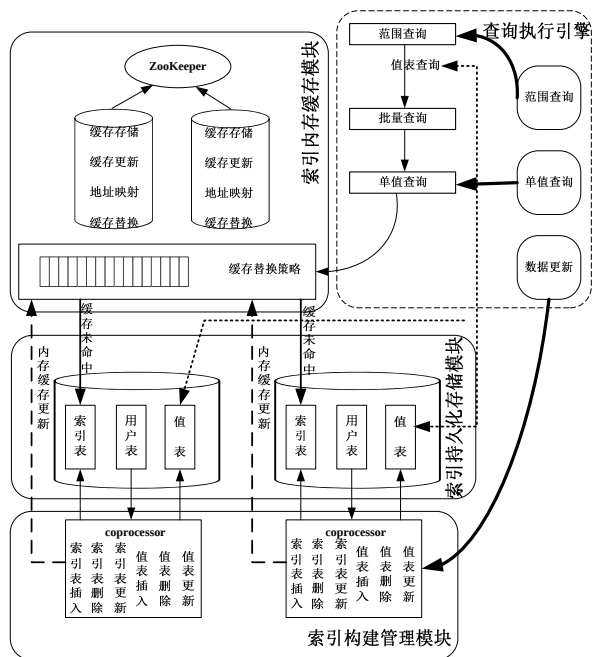


图 5 HiBase 系统架构

2) 持久化存储管理模块。提供索引表和值表的持久化存储，依赖于 HBase 为持久化存储数据提供可扩展性和容错性。

3) 索引内存缓存模块。管理索引热点数据的缓存存储、更新和地址映射，实现热度累积的缓存替换策略，使得最近频繁访问的数据能够缓存到内存中。

4) 查询执行引擎。将用户的查询请求翻译成系统识别的命令，调用对应的方法执行查询，并将查询结果汇总返回给客户端。

为了提供索引内存缓存的高可用性，HiBase 使用了 Hadoop 环境中的分布式协调管理系统 ZooKeeper 来可靠地检测分布式内存节点上服务进程的存活状态。索引内存缓存层每个内存节点服务进程会分别建立与 ZooKeeper 的会话，并创建临时 Znode 来表示自身的存活状态。每个内存节点服务进程从 ZooKeeper 系统镜像中可以观察到其他节点进程的存活状态。系统通过对分布内存节点状态的监测实现内存节点的失效检测和上线处理，

从而实现索引内存缓存层的高可用性。

5.2 索引构建过程

根据数据输入方式的不同，索引构建可分为面向流式数据和面向批处理数据的索引构建。索引创建过程都是读取用户表的一条记录，在非主键属性上生成一条索引记录，如果满足缓存条件，同时生成内存缓存层的索引数据。最后将索引数据分别更新到持久化存储层与内存缓存层，并更新值表。

对于流式数据索引构建方法，HiBase 利用 Observer 类型的 Coprocessor 来构建相关的索引。具体来说是使用 HBase 提供的 RegionObserver 接口的回调函数 prePut，插入一条数据之前会被触发调用。prePut 方法首先根据索引信息对用户发起的 Put 操作进行分析，如果 Put 操作的数据包含有索引列，即包含要索引的数据，则触发索引数据的插入。

由于静态数据一般相对较大，为了能够加快静态数据索引的构建速度，本文利用 Hadoop MapReduce 来并行化执行静态数据索引构建。MapReduce 任务首先得到输入<Row, Result>，其中 Row 为用户表的行键，Result 是通过 Row 获得的 HBase 记录。然后根据索引信息生成其对应的索引数据，并将索引数据插入到分层式索引中。整个过程不需要 Reduce 阶段即可完成，同时由于 HBase 用户表记录之间是相互独立的，所以可以充分利用 MapReduce 提供的并行化处理能力来加速索引构建过程。

5.3 数据查询过程

HiBase 通过建立非主键属性上的索引，支持高效的非主键列上的单值查询和范围查询。

对于单值查询，客户端的基本流程如下：

1) 从配置文件获取 ZooKeeper 的地址，建立与 ZooKeeper 的连接，获取注册的所有服务进程，确定当前提供内存缓存的所有服务进程位置信息。

2) 向内存缓存层的服务进程发起查询请求。若内存缓存层命中，则返回内存缓存层给出的结果，结束查询。

3) 若内存缓存层未命中，则向 HBase 的索引表发起查询请求。获取结果后，返回查询得到的结果，结束查询。

可以看出，如果在内存缓存层命中，整个查询流程都不会访问到磁盘，减少了磁盘访问开销，能

够大幅度提高响应速度。另外，对 ZooKeeper 访问获取的内存缓存层服务进程信息会被缓存在客户端，在后续的访问中会进一步减少数据查询的通信开销。

对于范围查询，内存缓存层通过一致性哈希的方法将数据分布到各个存储节点上来提高查询性能，而哈希破坏了索引列的有序性，所以我们需要记录非主键列上的数据属性值，将属性的所有取值集合保存在值表中。通过该值表，我们可以获得索引列某个范围内所有存在的属性值。

对于范围查询，客户端的基本流程如下：

1) 从配置文件获取 ZooKeeper 的地址，建立与 ZooKeeper 的连接，获取注册的所有服务进程，确定当前提供内存缓存的所有服务进程位置信息。

2) 从 HBase 中的值表获取查询范围内所有的查询属性值。

3) 依次发起单值查询请求。将查询结果汇总并返回。

与单值查询一样，缓存命中率的提高会降低范围查询的响应时间。范围查询中对值表的访问是 HiBase 为哈希索引付出的代价，但是由于每次范围查询只需要访问一次值表，其访问代价占整个范围查询的比重是较小的。

6 实验与性能分析评估

6.1 实验环境与数据集

我们从数据查询性能、数据扩展性、集群规模扩展性分别测试了 HiBase 的性能。我们的集群测试环境共有 10 个节点，集群的节点配置参见表 2。

表 2 计算节点配置信息

属性	配置信息
CPU	2 路 4 Core Intel Xeon E-5620 2.4GHz
Memory	24 GB
Disk	6 TB 7200RPM SATA II
Network	Bandwidth 1Gbps
OS	Red Hat Enterprise Linux Server 6.0
JVM Version	Java 1.6.0
Hadoop Version	Hadoop 1.2.1
HBase Version	HBase 0.94.14
ZooKeeper Version	ZooKeeper 3.4.5

实验采用的数据集包括布朗大学的大数据基准测试和国内电信运营商的电话用户上网记录数据。

首先，我们用电话用户上网记录数据测试不同系统的查询响应时间，该数据集共 1,072,868,115 条记录，每条数据记录不等长，平均是约 200 个字节，10 个属性。查询需求包括下面四个：

1) 根据用户号码 MDN 查询某一个时间段内的详单列表；

2) 根据 IP 查询某个时间段内使用该 IP 进行上网的用户列表；

3) 根据 IP+PORT 查询某个时间段内使用该 IP+PORT 进行上网的用户列表；

4) 根据 URL 查询某个时间段内登录该 URL 的用户列表；

针对以上查询需求，我们分别建立查询属性上的索引（其中查询需求 1 是在用户表主键上查询，不用建立索引），并进行查询响应时间测试。我们一共设计了 10 个测试用例，每个查询测试用例分别测试缓存未命中的冷查询响应时间和缓存命中的热查询响应时间。

然后，我们再进行布朗大学的大数据基准测试实验。布朗大学的大数据基准测试[13]是布朗大学在文献[14]中为了对比关系数据库和 Hadoop 在大数据上的查询性能而提出的基准测试。我们用该基准测试在 4 个节点上对数据集进行 10,000 次符合齐夫分布的持续查询（包括点查询和范围查询），测试在不同的数据缓存比例下缓存的命中率和查询响应时间。

6.2 实验结果及其分析

1) 非主键索引查询性能对比实验与结果分析

我们用 10 亿条电话用户上网记录数据对 HiBase 实现的非主键索引的性能进行对比测试。对比如下三种情况：(1) 不带任何非主键索引的标准 HBase 表上进行扫描的查询性能；(2) 华为公司的开源 HBase 非主键索引系统 Hindex 的查询性能；(3) HiBase 基于热度累积缓存算法的非主键索引的查询性能。图 6 给出了 4 个测试用例的查询响应时间性能对比。图 6 中的横坐标分别是查询获得的结果集大小，纵坐标是查询响应时间，由于 HBase 和 HiBase 的查询响应时间数量级相差太大，无法用等比例坐标，在这里使用指数指标。

实验结果可以看出，HiBase 的查询响应时间大大优于无非主键索引的标准 HBase 和开源的 Hindex 非主键索引系统的查询性能。对于结果集很小（极端情况下，结果集大小为 0）的查询，HiBase 和无非主键索引的标准 HBase 的扫描方法相比，冷

查询的性能提升达到 116,912 倍，热查询的性能提升可以达到 131,732 倍。对于结果集较小（结果集为 1155）的查询，冷查询的性能提升达到 3082 倍，热查询的性能提升达到 17,751 倍。对于结果集较大（结果集为 97515）的查询，冷查询的性能提升是 65 倍，热查询的性能提升是 343 倍。和开源的 Hindex 非主键索引系统相比，HiBase 热查询的性能提升可以达到 5~20 倍。从 HiBase 和标准 HBase 性能对比来看，结果集小的查询性能提升效果更明显，这是因为对于标准 HBase，不论结果集大小，非主键属性上的查询都需要遍历全部数据，而 HiBase 上的非主键索引可以直接定位到记录，查询响应时间随着结果集的大小而增长。

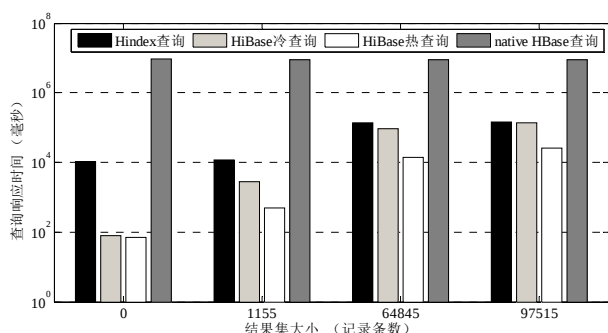


图 6 非主键索引的查询性能对比

2) 缓存命中率查询性能对比实验与结果分析

缓存命中率是指在批量查询中，在缓存中获得结果集的数据查询个数占全部数据查询个数的比率。提高缓存命中率会减少访问磁盘的次数，有效降低查询响应时间，因此命中率是缓存算法最重要的性能评价指标。命中率越高，缓存系统的查询性能越好。我们用布朗大学的大数据基准测试对 HiBase 进行 10,000 次符合齐夫分布的持续查询来测试缓存性能，数据规模为 10,000,000 条记录。图 7 给出了 HiBase 的热度累积缓存算法和 LRU 算法的命中率对比，图 8 给出了查询响应时间对比。

从图 7 中可以看出，缓存空间增加，命中率提高。热度累积缓存算法的命中率高于 LRU，尤其是在数据缓存比例为 0.2 时，热度累积缓存算法的命中率大约比 LRU 的命中率高 16%，因为热度累积缓存算法的热度累积机制能够更精确地记录数据的冷热程度，将热点数据缓存到内存中。对于大数据应用，缓存空间受内存限制，热度累积缓存算法在缓存比例不高时正可以充分发挥优势，提高大数据下的查询性能。

HiBase 是为了实现大数据上的高效非主键索引而实现的系统，我们在 HiBase 上做了无内存缓存、LRU 缓存替换策略、热度累积的缓存替换策略的 10000 次符合齐夫分布的持续查询，并记录查询响应时间，如图 8 所示。可以看出，随着缓存空间的增大，命中率不断提升，使用两种替换策略的索引查询时间都在随之降低。由于热度累积的缓存替换策略在命中率上的优势，查询响应时间明显优于使用 LRU 替换策略的查询性能，提升幅度从缓存比例为 0.2 时的 30% 到缓存比例为 0.8 时的 67%。

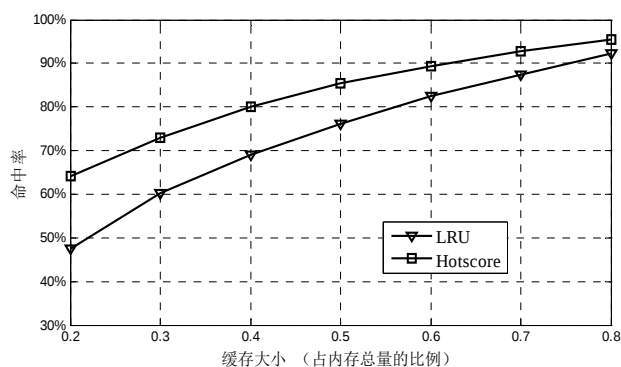


图 7 热度累积缓存算法和 LRU 的命中率对比

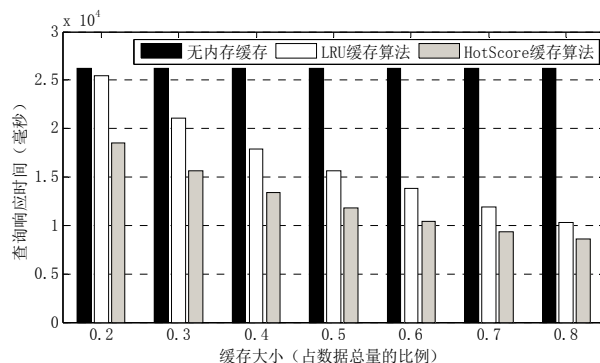


图 8 缓存算法的查询响应时间对比

3) 冷热查询对比实验与结果分析

我们在 10 亿条电话用户上网记录数据上做了 10 个测试用例的冷热查询性能对比实验。每个测试用例测试了第一次冷查询和随后的缓存命中热查询的查询响应时间。性能对比实验结果如图 9 所示，这里我们的查询响应时间是单次查询的性能指标，横坐标是查询得到的结果集大小（结果集包含的记录条数）。

可以看出，每个用例的第一次冷查询的时间都是最长的，第二次热查询由于缓存命中，所以查询响应时间缩短 5-15 倍（注：图 9 中的查询响应时间是指数坐标）。

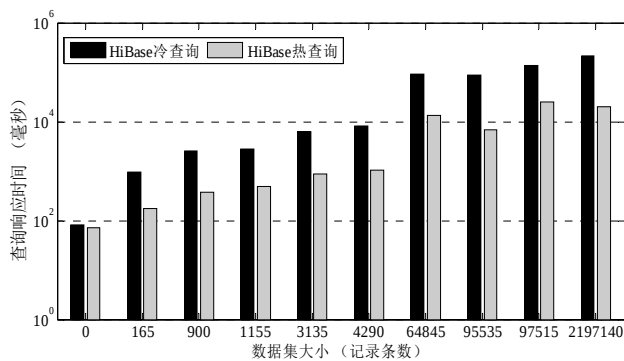


图9 电话用户上网记录数据查询响应时间

4) 扩展性实验与结果分析

我们通过布朗大学的大数据基准测试验证了HiBase在不同数据规模下的查询性能，如图10所示。可以看到，部署在10个节点上的HiBase，在数据规模从10,000,000扩展到50,000,000条数据记录时，查询时间保持线性增长，这验证了HiBase在数据可扩展性方面有着良好的表现。查询响应时间和结果集大小成正比，因为在查询响应时间中，主要开销是对查询结果集的访问。在数据集中，数据是符合均匀分布的，随着数据行总数的增长，查询结果集也随之增加，因此查询响应时间会与数据行总数的大小成正比。

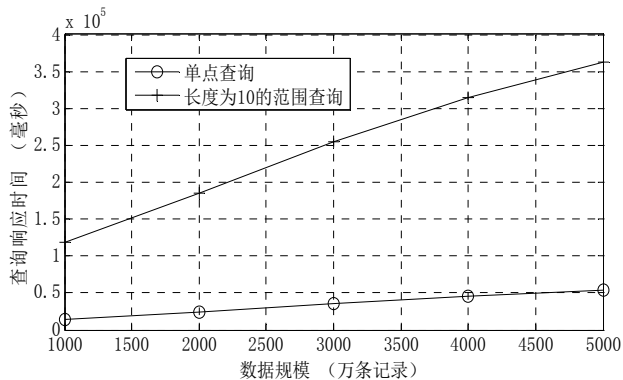


图10 数据规模增大的查询响应时间对比

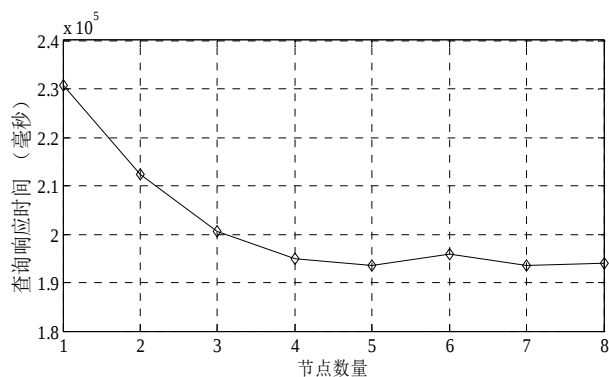


图11 节点数量增大的查询响应时间对比

同时，我们在20,000,000条数据记录上进行了10000次范围长度为10的范围查询可扩展性测试，如图11所示。本实验不使用单值查询的原因是：在进行单值查询时，没有并发的过程，所以节点的变化几乎不影响查询时间。可以看出，随着节点数量的增加，查询响应时间逐渐降低。在范围查询中，查询会根据范围内存在的索引列值向各节点发送查询请求，查询被并行执行后汇总结果，所以查询响应时间会随着节点的增加而降低。另一方面，索引列值是通过一致性哈希算法映射到各个节点上的，我们的测试实例范围固定为10，所以当节点个数增加到一定程度之后，查询响应时间趋于平缓。

5) 索引构建的时间和空间开销分析

HiBase非主键索引系统是在数据导入阶段创建索引，我们在电话用户上网记录数据上测试HiBase系统的数据导入时间来衡量系统的索引构建时间开销。同时我们测试了HBase的数据导入时间作为对比。数据导入时间的对比实验结果如表3所示。

在电话用户上网记录实验数据集上我们分别做1.4GB部分数据集(6,501,311条记录)和225GB全部数据集(1,072,868,115条记录)的数据导入时间对比测试。可以看出，HiBase的数据导入时间分别是HBase的1.03和1.12倍，因为HiBase在将用户表插入到HBase的同时，触发HBase Coprocessor并行地构建索引记录并插入到索引表中。实验结果表明，HiBase为建立索引所带来的时间开销是完全可以接受的。

表3 HiBase和HBase的数据导入时间对比

数据规模 (记录行数)	数据集大小 (GB)	HBase 数据导入时间(ms)	HiBase 数据导入时间(ms)
6,501,311	1.4	113,352	116,768
1,072,868,115	225	120,382,424	135,297,604

空间开销上，索引数据与原始用户表数据的占比不是绝对固定的，而是依赖于用户表的结构以及索引属性列值的占用空间。在本文的电话用户测试数据中，225GB全部数据集的用户表和索引表空间开销总和为575GB，其中225GB文本数据导入到HBase后的用户表空间开销为238GB。针对6.1节中提到的查询需求1~4（其中查询需求1是用户表主键查询），我们在HiBase中为查询需求2~4建立

非主键索引。索引的空间开销总和为 337GB，单个索引的平均开销约为 112.3GB，故单个非主键索引表的空间开销约为用户表的 47%。对于采用廉价的分布式存储、具有优异的存储空间可扩展性的 Hadoop 和 HBase 系统来说，构建一个十多个节点的 Hadoop 集群即可获得数十 TB 的存储容量。因此，通过为每个非主键索引提供约一半用户表的空间开销而达到数十至数百倍以上的查询性能提升是完全值得的。

7 总结和展望

为了提高 HBase 上非主键查询的效率，本文提出了一种基于 HBase 的分层式非主键索引存储模型，并设计实现了分层式索引查询系统 HiBase。HiBase 在 HBase 上持久化存储用户表和索引表，并将索引表的热点数据缓存在内存中。同时我们研究提出并实现了一种基于热度累积的缓存替换策略，通过高效的缓存替换策略，HiBase 获得了优于 LRU 的缓存命中率，进一步提升了 HBase 上非主键查询的性能。

下一步，我们会针对非主键属性上的聚集查询进行研究并提出高效的解决方案。另外，在 HiBase 中为了支持范围查询需要引入值表，而海量数据下的值表开销不容忽视，我们会针对值表做大粒度存储优化，降低值表的空间开销。

参 考 文 献

[1] F. Chang, J. Dean, S. Ghemawat, W. C. Hsieh, D. A. Wallach, M. Burrows, T. Chandra, A. Fikes, and R. E. Gruber. Bigtable: a distributed storage system for structured data. Proceedings of the 7th USENIX Symposium on Operating Systems Design and Implementation (OSDI),



GE Wei, born in 1979. PhD candidate. Member of China Computer Federation (E200035709M). Her research interests include query processing, query optimization, distributed and parallel computing.

LUO Sheng-Mei, born in 1971, M.S. Senior Engineer. Member of China Computer Federation (E2000160404M). His research interests include cloud computing, cloud storage, and big data processing.

Seattle, WA. 2006:205-218

[2] Apache HBase, <http://hbase.apache.org/>

[3] A. Lakshman and P. Malik. Cassandra- a decentralized structured storage system. ACM SIGOPS Operating Systems Review. 2010. 44(2): 35-40.

[4] G. DeCandia, D. Hastorun, M. Jampani, G. Kakulapati, A. Lakshman, A. Pilchin, S. Sivasubramanian, P. Voshall, and W. Vogels. Dynamo: Amazon's highly available key-value store. Proceedings of the 21st ACM Symposium on Operating Systems Principles (SOSP). Stevenson, WA. 2007:205--220

[5] Redis, <http://redis.io/>

[6] Huawei hindex, <https://github.com/Huawei-Hadoop/hindex>.

[7] NGDATA hbase-indexer, <https://github.com/NGDATA/hbase-indexer>.

[8] Apache Solr, <https://lucene.apache.org/solr/>.

[9] George Sfakianakis, Ioannis Patlakas, Nikos Ntarmos, Peter Triantafillou. Interval indexing and querying on key-value cloud stores. Proceedings of the 29th IEEE International Conference on Data Engineering (ICDE). Brisbane, Australia. 2013: 805-816

[10] Cristian Ungureanu, Biplob Debnath, Stephen Rago, Akshat Aranya. TBF: A memory-efficient replacement policy for flash-based caches. Proceedings of the 29th IEEE International Conference on Data Engineering (ICDE). Brisbane, Australia. 2013:1117-1128

[11] Justin J. Levandoski, Per-Ake Larson, Radu Stoica. Identifying Hot and Cold Data in Main-Memory Databases. Proceedings of the 29th IEEE International Conference on Data Engineering (ICDE). Brisbane, Australia. 2013:26-37

[12] David Karger, Eric Lehman, and Tom Leighton. Consistent hashing and random trees: distributed caching protocols for relieving hot spots on the World Wide Web, Proceedings of the 29th annual ACM symposium on Theory of computing, NY, USA. 1997:654-663.

[13] Benchmark of Brown University, <http://database.cs.brown.edu/projects/mapreduce-vs-dbms/>

[14] Andrew Pavlo, Erik Paulson, Alexander Rasin, Daniel J. Abadi, David J. DeWitt, Samuel Madden, Michael Stonebraker. A comparison of approaches to large-scale data analysis. Proceedings of the ACM SIGMOD Conference, Providence, USA. 2009: 165-178

ZHOU Wen-Hui, born in 1987, M.S. His research interests include parallel computing, distributed consistency of big data.

YUAN Chun-Feng, born in 1963, Professor, Senior Member of China Computer Federation. Her research interests include computer system architecture, multimedia document processing, web information extraction and integration.

HUANG Yi-Hua, born in 1962, Professor, Senior Member of China Computer Federation (14302S). His research interests include big data parallel processing and cloud computing, computer architecture, distributed and parallel computing.

Background

HBase is an open-source key-value storing system in Hadoop Ecosystem of Apache Foundation. Data in HBase are composed of much smaller entities in format of key-value pairs, and HBase organizes the small records into large storage files and stores them on HDFS to provide well scalability and fault tolerance. HBase is designed for fast point queries and sequential scans on row keys. Each data split maintains key-value pairs with ranges of rows sorted by row key, so it excels at providing key-based or sequential range access. Point queries are supported by providing a row key to find what you are looking for. When faced to non-key query requirement, HBase takes a measurement of sequential scan, that is, scanning data from start to end one by one. It is suitable for reading adjacent key-value pairs and is optimized for block disk access operations that can make full use of disk transfer channels. Whereas, sequential scan misses the basic features of key-value retrieval, such as selection of row keys based on regular expressions. Thus the approach of non-key query leaves much to be desired. Some research works try to provide a

solution by building secondary index in HBase.

HiBase is our high-performance storing and management system based on HBase to provide efficient query on non-key column. It builds secondary index according to global secondary index model, which outperforms local secondary index model (adopted by Hindex) for it saves much unnecessary computation overhead. In local secondary index query process accesses index table in all Regions, but some Regions returns null result set. Furthermore, HiBase provides efficient memory caching policy, so as to reduce the index's disk access overhead, and as a result, query performance is improved in a large degree. Thus, HiBase offers the realtime or near-realtime capability for big data analysis.

This work is supported by China National Science Foundation of Outstanding Key Laboratory Special Research Grant (61223003) and the ZTE Research Funding.