

行云(Cirro Data)产品白皮书

东方国信分布式数据库

Version 2.6

北京东方国信科技股份有限公司 | 北京朝阳区创达三路一号院1号楼东方国信大厦

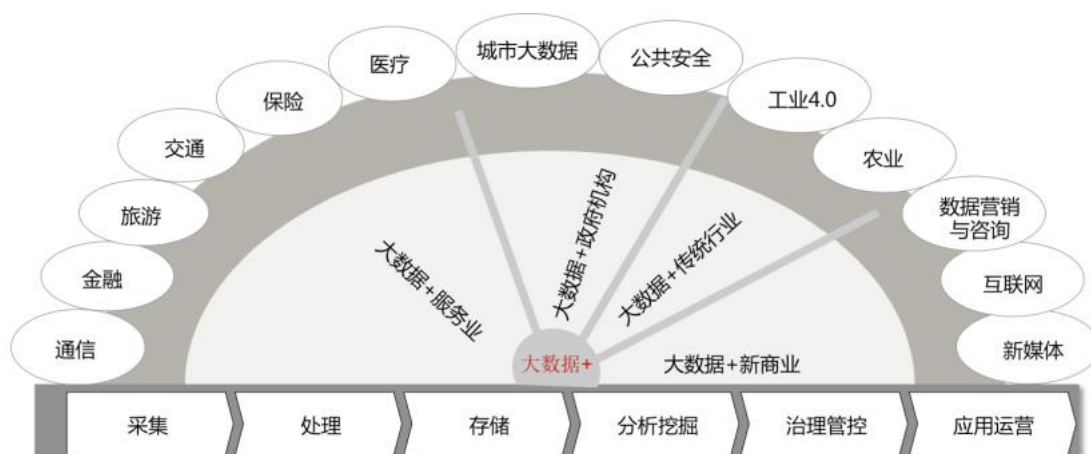
目录

1.	公司概况.....	2
2.	行云特性.....	3
2.1	产品概述.....	3
2.2	系统架构.....	4
2.3	特点描述.....	5
3.	核心功能.....	6
3.1	PL/SQL标准服务	6
3.2	数据类型.....	6
3.3	SQL语法支持.....	6
3.4	函数.....	7
3.5	数据加载导出.....	8
3.6	资源管理.....	8
3.7	安全管理.....	8
3.8	应用接口.....	9
4.	关键技术.....	10
4.1	高性能核心计算引擎.....	10
4.2	分布式动态调配与弹性计算.....	10
4.3	Master对等技术	12
4.4	分布式元数据架构.....	13
4.5	统一的集群事务管理.....	13
4.6	透明高效的行列混合分布式存储引擎.....	15
4.7	基于规则或代价的性能优化.....	15
4.8	多租户技术实现.....	18
4.9	分布式存储过程执行引擎.....	19
4.10	跨地域数据中心架构.....	20
4.11	数据联邦-整合多种异构数据源	20
5.	管理工具.....	22
5.1	可视化集群管理工具.....	22
5.2	可视化 SQL分析工具.....	23
6.	性能指标.....	25
7.	运行环境.....	26

1.公司概况

东方国信成立于 1997年，是中国领先的大数据上市科技公司（股票代码 300166）。自成立以来，东方国信就专注于大数据领域，紧跟全球大数据技术的发展趋势，通过自主研发，打造了面向大数据采集、汇聚、处理、存储、分析、挖掘、应用、管控为一体的大数据核心能力，构建了云化架构的大数据产品体系，对标国外优秀软件产品形成了端到端的软硬件相结合的大数据解决方案，并打造了业内领先的大数据能力开放平台解决方案，成为国内民族软件的第一品牌。

基于大数据的核心能力，东方国信以“大数据+”为战略，紧锣密鼓的加快战略布局，以领先的大数据解决方案服务于通信、金融、智慧城市、公共安全、智慧旅游、工业、农业、医疗、媒体、大数据运营等行业和业务领域，帮助客户从数据中获得价值，得到行业与客户的广泛、高度认可，也铸就了东方国信大数据龙头企业的行业地位。



东方国信目前在全国 31个省及直辖市都设有分支机构或项目实施团队，构建了贴身式服务的大数据落地应用体系。

“让数据改变工作与生活”是东方国信的企业愿景，“专注、智慧、自省、包容”是东方国信人的行为准则，东方国信已经构建起自有特色的企业文化，培养了大批的行业优秀人才，致力于打造员工满意的雇主品牌。

作为国内大数据龙头企业，东方国信发挥大数据行业资源与能力优势，引入专业投资管理机构和其他社会资本共同设立大数据产业基金，承担大数据企业责任，聚合数据价值，推动大数据+战略落地，引领产业升级，全面优化产业生态链，让大数据的价值更好的服务于社会、企业和民生。

2.行云特性

2.1产品概述

行云（Cirro Data）分布式数据库是北京东方国信科技股份有限公司面向海量数据分析型应用领域，完美融合 Hadoop平台和 MPP架构的各自优势，充分利用列存储和行存储的特点，以分布式存储和高效压缩技术为基础，动态计算资源调配，完全自主研发的一款新型分布式高性能数据库产品。

行云分布式数据库能满足 Terabyte到 Petabyte级别的海量数据存储和分析，这些数据可以分布在数百台普通服务器上，并且能够被大量并发用户高速访问，可用于满足各个数据密集型行业日益增大的海量数据分析，数据挖掘，数据备份和即席查询的需求。

产品特性

- 大规模数据处理能力
- 高性能复杂分析能力
- 高性能数据加载能力
- 高并发键值查询能力

- 支持数百节点大集群
- PB 级数据存储与计算
- 线性扩容提升开发数量
- 线性扩容提升计算性能

- 标准 JDBC、ODBC
- 支持 PL/SQL 脚本语言
- 图形化集群管理界面
- 可视化 PL/SQL 开发界面

- DBLink 连接异构数据源
- 多种汇入/汇出格式
- 平台安全验证与加密
- 数据备份、恢复与迁移

三高一低 易管易控

高性能

数据查询性能：可为用户提供秒级的实时查询服务

高并发

单节点最大可以支持上百个用户，并发用户数随节点数线性增长

高可用

企业级发行版产品，能保证长期稳定运行。无单点故障、无单点性能瓶颈

大容量

可以支撑PB级数据的快速加载和查询

低成本

可部署在X86架构的普通服务器上，数据存储在本地磁盘上，高效压缩技术可在同数据规模下有效减少服务器和磁盘数量

易管理

易安装、易管理、部署灵活。提供基于客户端的集群安装和管理界面，系统提供了监控集群节点和增删节点的功能

多功能

提供跨数据库虚拟大表功能，实现数据库复制和备份功能，实现在线平滑扩容

2.2系统架构

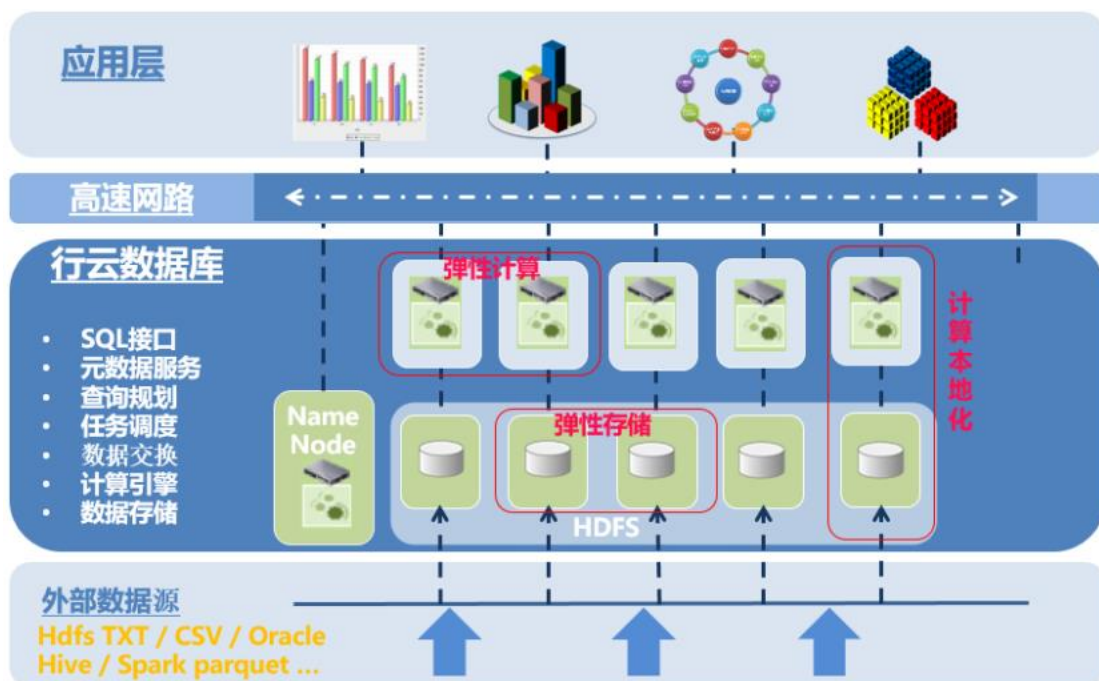


图 1 系统架构图



图 2 产品功能体系

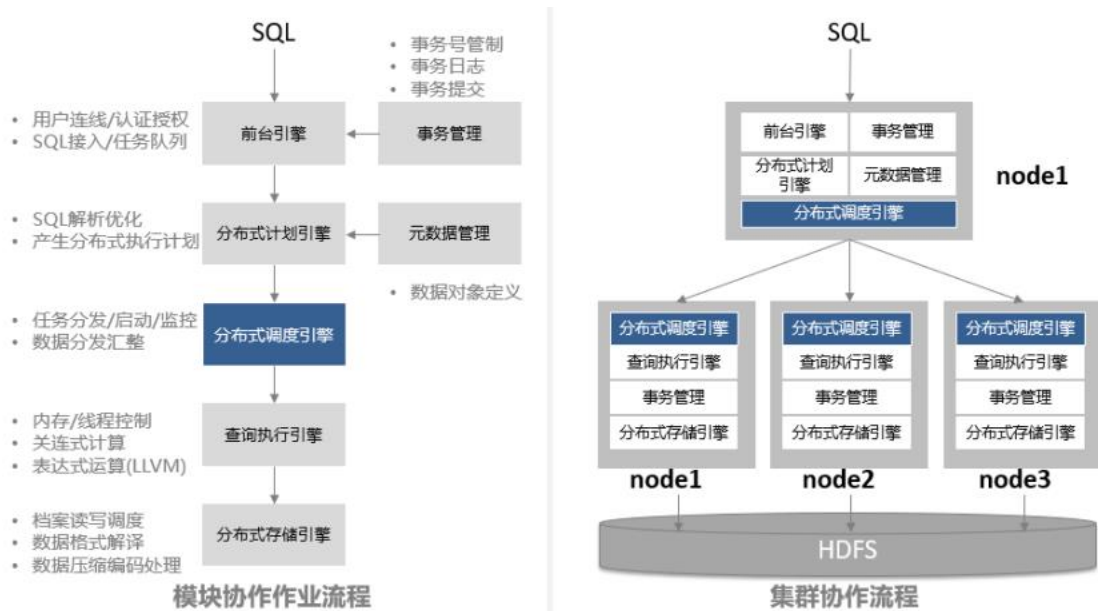


图 3 运行架构

2.3特点描述

- 完整支持 SQL-92标准，兼容 SQL-2003
- 支持 PL/SQL标准存储过程
- 基于 Linux操作系统
- 采用商用 x86-64的标准 PC服务器
- 支持通用 API: JDBC / ODBC
- 支持海量数据的高效粗粒度索引机制，索引占用空间小，膨胀低
- 支持分布式并行处理及内存计算
- 支持多租户资源管理和隔离
- 支持跨域数据计算
- 支持数据联邦计算

关键指标

- 基于行列混合的数据存储结构，压缩比最大可达 1:20
- 可支持 PB规模以上的结构化数据存储
- 高性能数据加载，可达 1GB/sec以上
- 高性能数据统计查询和数据加工处理，亿级数据秒级响应
- 基于索引的全表扫描，十亿级数据毫秒级响应
- 高可扩展性：支持 4~512+节点

3.核心功能

行云分布式数据库能针对海量数据，同时被大量并发用户高速访问。可以满足各个数据密集型行业的用户对日益增大的海量数据科学、便捷和高效分析。

3.1 PL/SQL标准服务

行云分布式数据库提供了与传统关系型数据库高度兼容的 PL/SQL标准服务。同时还为用户提供了清晰、便捷的存储过程图形化集成环境（PL/SQL Developer），用于访问、配置、管理数据库对象和进行 PL/SQL编辑和执行。

因应海量数据时代，批量加工处理任务的高效率需求，行云是首个采用分布式的 P/L SQL调度执行引擎的数据库系统，能有效提高集群的并发能力，缩短任务执行时间，并支持执行引擎集群的线性无限扩展。利用负载均衡技术，根据 P/L SQL执行引擎忙闲状态，灵活调度分配作业，提高资源利用。当执行过程中出现一个执行节点宕机，执行引擎将会分配另一个节点执行任务，确保整个集群的高可用性。

3.2数据类型

行云分布式数据库支持以下数据类型：

- ☐ 数值类型：INT，DOUBLE，LONG，NUMBER，DECIMAL
- ☐ 字符类型：CHAR，VARCHAR，VARCHAR2
- ☐ 日期时间类型：DATE，INTERVAL，TIMESTAMP
- ☐ 特殊类型：NULL

3.3 SQL语法支持

行云分布式数据库支持 SQL 92标准，支持标准 DDL、DCL、DML语法。

- ☐ 数据定义语法：
 - ☐ 支持的 DDL语法包括：
 - CREATE：包含对数据库（SCHEMA），数据表，DBLINK，索引，视图的创建语法。
 - DROP：包含对数据库（SCHEMA），数据表，DBLINK，索引，视图的删除语法。
 - TRUNCATE：包含对数据表的截断语法。

- ALTER: 包含对分区, 列的更改语法。
- 分区: 包含对分区的增加, 删除, 截断语法。
- 列: 包含对列的增加, 删除, 修改和重命名语法。
- 数据操作语法:
 - 包含 SELECT、INSERT、EXPORT、UPDATE、DELETE、MERGE INTO语法等。
 - 支持 group by、order by、case when、limit、in/not in、is null /is not null、where、having, like等。
 - 支持多个事实表之间、事实表与维度表之间的 join, 支持 UNION和子查询（包括相关子查询）。
 - 支持 distinct去重。
 - 支持 compaction数据整合。
- 数据管理语法:

包含用户管理语法、查看语法等。其中用户管理语法包括创建、删除、密码修改、权限授予及权限回收语法等；查看语法包括数据库、表、列查看语法等。

3.4函数

行云分布式数据库函数支持聚合函数、单值函数、字符串函数、日期和时间函数、窗口函数。

□ Math函数

FLOOR, CEIL, ROUND, SIGN, MOD, ABS, TRUNC, LOG, LN, POWER, SQRT, SIN, COS, ATAN, ASIN, ACOS, BITAND, BITOR等

□ String函数

CONCAT, REPLACE, SUBSTR, LENGTH, LOWER, UPPER, LTRIM, RTRIM, TRIM, INITCAP, LEFT, RIGHT, TO_CHAR, REGEXP_LIKE, REGEXP_SUBSTR, TO_SINGLE_BYTE, INSTR, TRANSLATE等

□ Date函数

CURRENT_DATE, SYSDATE, LAST_DAY, TRUNC, ROUND, ADD_MONTHS, MONTHS_BETWEEN, TO_DATE, TO_CHAR, YEAR, MONTH, DAY, QUARTER等

□ Interval函数

NUMTOYMINTERVAL, NUMTODSINTERVAL, TO_YMINTERVAL, TO_DSINTERVAL, TO_CHAR等

□ Timestamp函数

CURRENT_TIMESTAMP, LOCALTIMESTAMP, SYSTIMESTAMP, TO_CHAR, TO_TIMESTAMP, QUARTER等

□ 聚合函数

sum、count、max、min、avg、STDDEV_POP等

- 窗口函数

row_number、rank、dense_rank、ntile等

- 分层查询

- 其他函数

TO_NUMBER、CAST、EXTRACT、DBMS_RANDOM.STRING等

3.5数据加载导出

- 数据加载：

支持从远端数据库（DBLINK FROM ORACLE，支持多并发）、本地数据库、本地文本文件、本地 HDFS文本文件、外部 HDFS文本文件导入数据。

- 数据导出：

支持从行云分布式数据库导出到本地文本文件、HDFS文本文件。

3.6资源管理

- 全局资源管理：

- 设置数据库级、用户级、队列级的资源配额
- 动态调整配置参数

- 多租户资源管理

- 设置进程级的资源管理和资源隔离
- 资源合理与数据隔离区分

3.7安全管理

- 用户管理：

提供行云用户登录验证功能（创建、删除、更改用户密码）。

- 权限管理（对象授权、权限回收、系统授权、回收）

- 根据数据安全策略，区分用户级别，管理用户权限，保证数据的安全性。

- 支持对 SCHEAM、TABLE对象的权限管理。

3.8应用接口

□ 行云分布式数据库支持 JDBC接口

Cirro Data提供基于 **JDBC**的的数据访问接口，任何 **JAVA**程序都可透过 **ODBC**驱动来访问行云数据库。**Cirro Data JDBC**是遵循 **JDBC 4.0**标准规范，用户端的 **JRE**必须在 **jdk1.6.037**以上版本。

□ 行云分布式数据库支持 ODBC接口

Cirro Data提供基于 **ODBC**的数据访问接口，是基于 **ODBC 3.0**的接口规范，支持 **Windows**平台上任何支持 **ODBC**驱动的开发工具透过 **ODBC**接口来对行云数据库进行操作。

4.关键技术

4.1高性能核心计算引擎

行云是基于 C++高性能原生技术开发的企业级分布式分析型数据库。随着海量数据问题的出现，海量的数据管理能力和单机性能的不足，已经不能通过多上几个节点就可以解决。C++面向对象的、简洁、灵活性实用、数据类型丰富、执行效率高等优势，在实际开发中数据表明相对于 java，原生 C语言代码更能充分发挥硬件的计算能力，从根本入手并解决问题。

行云利用 C++的高性能与高结构化特性实现了以下优势：

- 集群调度最佳化
 - 集群任务负载均衡。
 - 高 K值聚合算法调度优化，基于并行计算来进行算法提速。
 - 大表连接算法调度优化，基于数据重分发技术进行算法提速
 - 极值/区间索引过滤优化。
- 内存计算最佳化
 - 核心引擎使用 C语言开发，保障硬件计算能力的高效运用。
 - 极致的代码优化设计，保障每个运算节点的最大吞吐能力，具体优化的包含内存布局优化、物件池技术、内存配置器技术、C++编译器优化、模板元编程技术、分支预测技术、数据预取技术、无锁技术、流水线技术等。
 - 采用表达式计算代码动态生成技术，保障海量数据进行过滤/表达式运算的最佳计算性能。
- 数据读写最佳化
 - 组合数据压缩算法，有效提高不同类型数据的实际压缩率。
 - 数据文件自适应整并，降低存储文件对整个存储系统的压力。
 - 高效的列存储设计，能够有效的支持结构化和非结构化数据，以及灵活的支持各种高性能压缩算法。

4.2分布式动态调配与弹性计算

行云是率先采用了数据分片存储与计算本地化架构，结合动态数据分片以及动态计算资源分配等技术，实现了 SQL级别的弹性计算能力的分布式数据库产品。

透过行云的智能弹性计算引擎，分别在查询规划时期与查询执行的进行计算资源的优化与配置。可同时兼顾了在集群负载不高时，尽量利用空闲资源提高查询性能；与高并发时酌量限制计算节点个数以提高并发数的负载均衡需

求。更因此能实现透过动态的横向扩展节点数量，即可同时提高单一 SQL 的运行效率又能提高集群的整体并发能力。

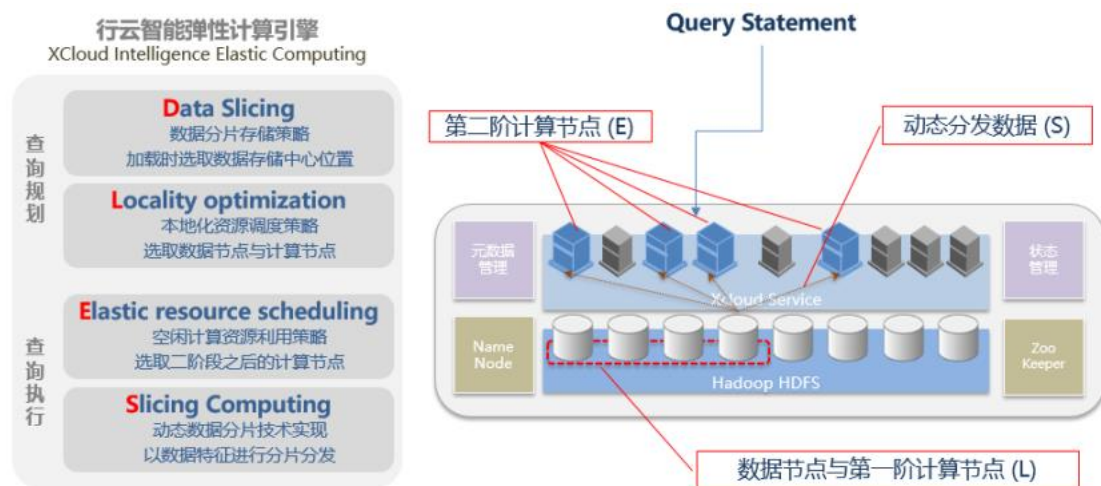


图 5 弹性计算与动态分布

- 数据分片存储策略（Data Slicing）
加载时动态选取数据存储中心位置，以数据特性分布存放不同的数据节点上。
- 本地化资源调度策略（Locality optimization）
依据数据存储位置与计算资源负载情况，选取最优的数据节点与计算节点。
- 空闲计算资源利用策略（Elastic resource scheduling）
依据集群内 CPU 资源的负载情况，选取二阶段之后的计算节点
- 动态数据分片技术实现（Slicing Computing）
以数据特征进行分片分发至指派的二阶段计算节点上进行数据运算。

4.3 Master对等技术

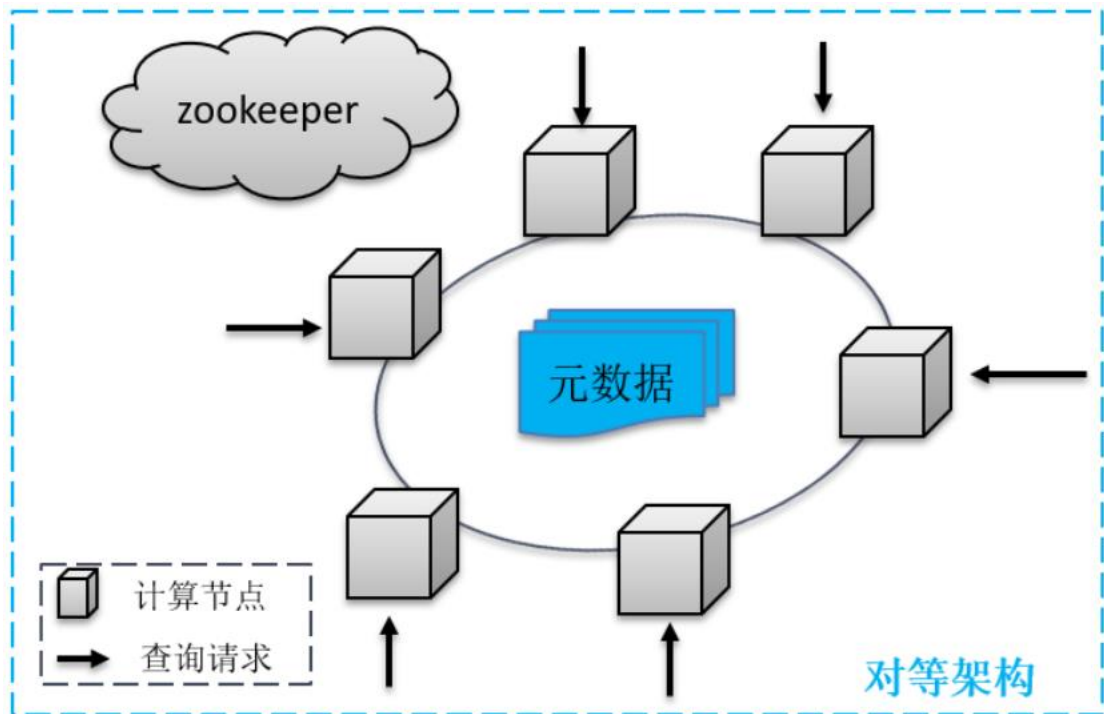


图 6 Master 对等技术的分布式集群架构

不同于传统 MPP架构，行云采用了无专用入口节点的 Master对等技术。集群中的任一节点，都可以做为接受 SQL请求的入口节点。不需要传统的 Master节点 H/A机制，不受 Master节点软硬件性能的限制，可极为有效的提高整体集群的并发负载能力，均匀的分散集群计算压力。

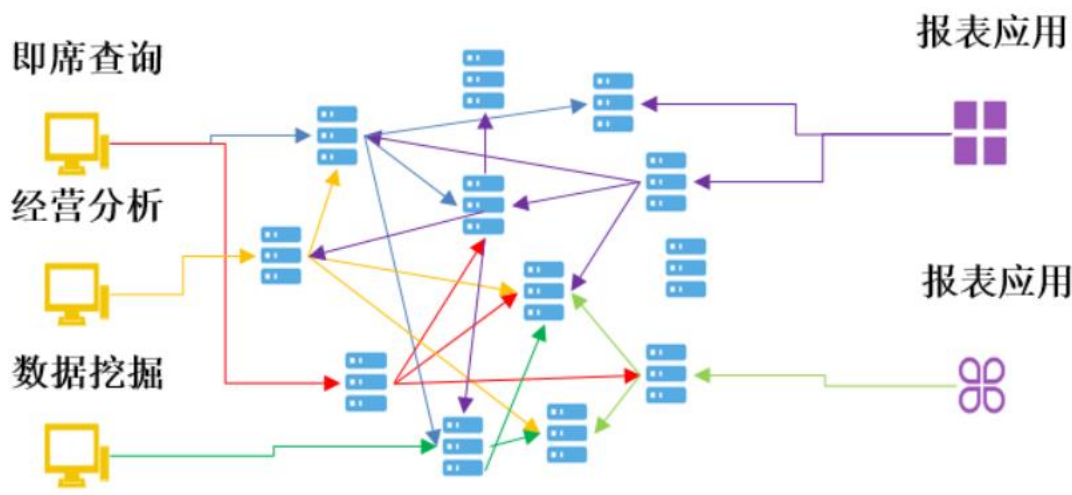


图 7 Master 对等技术的使用场景

4.4 分布式元数据架构

行云采用的分布式元数据管理技术。

支持分布式元数据管理

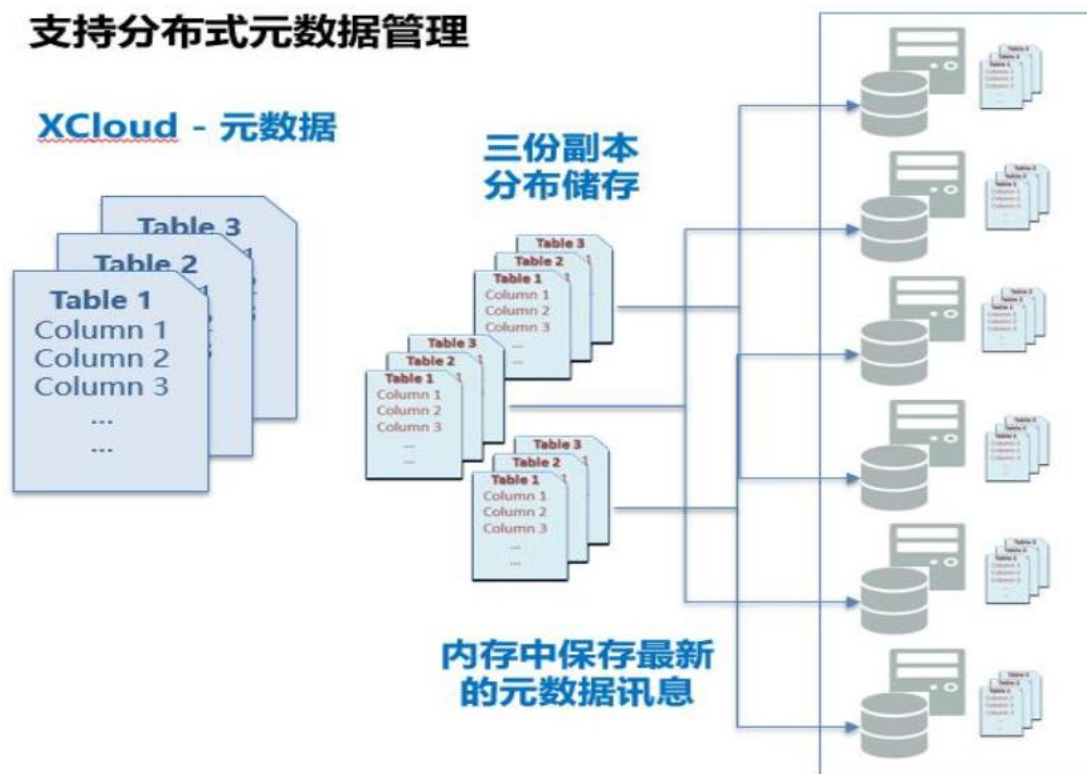


图 8 分布式元数据管理

□ 高可用性

元数据管理采用无集中服务器方式，利用 HDFS 以节点冗余存放技术，解决元数据读写瓶颈和元数据服务器单一节点失效问题，具有极高的可靠性和容错性

□ 高性能

各节点以内存缓存常用的元数据，极大的提高 SQL 运行效率

□ 高扩展性

元数据管理支持平滑扩展，支持动态增删节点，具有良好的伸缩性

4.5 统一的集群事务管理

□ 支持强一致性事务处理能力

相对于其他开源 Hadoop 方案在事务支持上的不足，行云支持强一致性事务处理能力，满足企业级数据库一致性的需求

□ 支持高并发的数据读取

支持事务的 RCSI隔离级别，并行读事务完全不受影响，支持高并发的数据读取，满足用户对数据处理的高并发需求

□ 高可靠性

基于 HDFS的分布式事务日志，确保在任何条件下数据库状态的持久性

□ DDL/DML操作快速响应

批次数据清理机制，保证 DDL/DML操作的快速响应

□ 完整的 ACID实现

实现基于分布式环境下基于 MVCC的 ACID事务处理机制

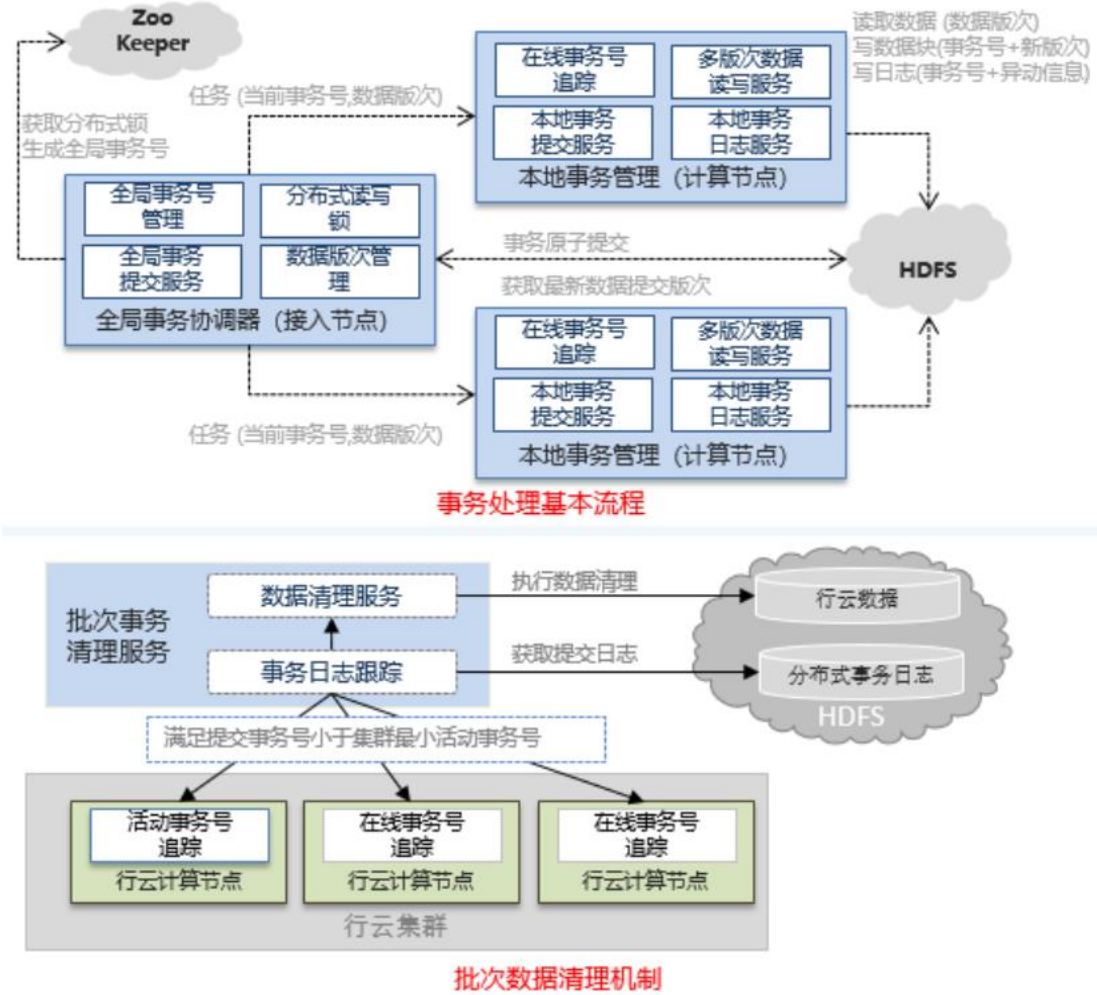


图 9 集群事务管理

4.6透明高效的行列混合分布式存储引擎

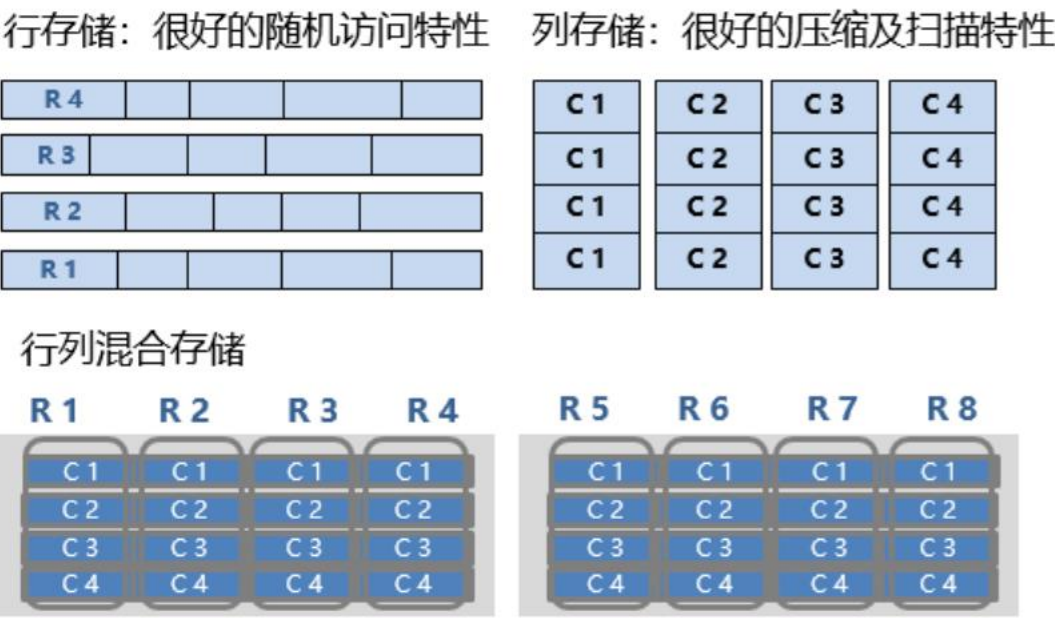


图 10 集群事务管理

- 5-20倍高压缩比例
列存储在压缩方面比传统的行存储数据库更加有效。由于同一列中的所有数据具有相同的数据类型，连续存储的数据具有很大的相似性，数据压缩效率比不同数据类型字段连续存储的行存储表更高
- 毫秒级高效检索
每个区描述项中都存储了区数据的最大值和最小值，且每个列的数据是在段中连续存储的，相当于对每个列都有分段的范围索引，能大大提高数据检索的效率

4.7基于规则或代价的性能优化

- 基于规则的优化（RBO），使用 hint的方式触发当前支持如下：
 - /*+SLICE_SORT*/
适用语句
SELECT hint ... WHERE in_query_clause
意义
建议优化器使用排序算法计算 in条件。仅在 in_query_clause的两个表都是有序表、slice个数相同且 in的列是 slice key时生效。

□ **/*+ORDERED_JOIN*/**

适用语句

SELECT hint ... JOIN ...

意义

建议优化器使用 **sort join**算法。仅在 JOIN的两个表都是有序表、**slice**个数相同且 **join key**是 **slice key**时生效。

□ **/*+MULTI_AGG*/**

适用语句

SELECT hint ... GROUP BY ...

意义

建议优化器使用 **multiple agg**算法。这种 **agg**算法适用于分组列的 **distinct**值很高的情况。

□ **/*+SEMI_AGG*/**

适用语句

SELECT hint ... GROUP BY ...

意义

建议优化器使用 **semi agg**算法。这种 **agg**算法适用于分组列的 **distinct**值不很高也不很低的情况。这种算法是优化器的默认选项。

□ **/*+NORMAL_AGG*/**

适用语句

SELECT hint ... GROUP BY ...

意义

建议优化器使用 **sort agg**算法。这种 **agg**算法适用于分组列的 **distinct**值很低的情况。

□ **/*+CONCURRENT_TASK_RATIO(n)*/**

适用语句

SELECT hint

INSERT INTO hint

UPDATE hint

DELETE hint

意义

指示优化器对此条件产生的并发执行计划的倍数。**n**为倍数，有效值为**[0,100]**内的整数。此 **hint**的优先级高于 **concurrent_task_ratio**配置参数，本语

句以 hint为准。

□ `/*+BROADCAST_TABLE*/`

适用语句

```
SELECT .. FROM hint table_or_subquery JOIN table_or_subquery
SELECT .. FROM table_or_subquery JOIN hint table_or_subquery
SELECT .. FROM table_or_subquery WHERE expr [NOT] IN (SELECT ... FROM
hint table_or_subquery)
SELECT .. FROM table_or_subquery WHERE [NOT] EXISTS (SELECT ... FROM
hint table_or_subquery)
```

意义

指示优化器对此 hint后的表或子查询进行广播分发，可以用于 JOIN、IN/NOT IN和 EXISTS/NOT EXISTS语句中。此算法适用于两表数据相差很大的情况，对数据量小的表使用广播分发。

□ 基于代价的优化（CBO）

统计信息是描述数据存储的元数据。优化器通过统计信息来选择最优的执行计划。主要包括扫描方式的选择和 join路径的优化。

计算统计信息会花费时间和资源，所以 CirroData对大表进行分析时会进行取样。

统计信息包括：

reltuples：表的行数，是个估算值。

stanullfrac：此列为 null的比例。

stawidth：此列非 null值的平均宽度，单位是字节。

stadistinct：此列的 distinct情况，正数表示估算的 distinct值，负数表示 distinct的比例，0表示不确定 distinct情况。

比如，值-1.0表示此列是唯一值列。平均宽度大于 1024的列被认为是唯一值列。

stakindN：一个编码值，表明第 N个 slot中的统计信息的类型。

staop：用于描述第 N个 slot中的统计信息的操作符。比如，直方图 slot此列为<操作符，定义数据的排列顺序。

stanumbers：包含第 N个 slot中适当的数值类型的数组，如果 slot类型不涉及数值类型，则为 NULL。

stavalues：包含第 N个 slot适当的列的数据的数组，如果 slot类型没有存储任何数据，则为 NULL。

最常见值（MCV）

□ slot

staop, =操作符。

stavalues, K个最常见的非 NULL值。

stanumbers, stavalues中最常见的非 NULL值的比例。

stavalues中值是按照出现的比例排序的, K是统计收集器选择的一个值。stavalues中的值必须出现不止一次, 唯一值列没有此 slot。

□直方图 Slot。

staop, <操作符。

stavalues, M个非 NULL值, 这 M个值把总体数据近似等分为 M-1份。第一个值最小, 最后一个值最大。

stanumbers, 未使用, NULL。

如果同时存在 MCV slot, 则直方图 Slot描述的是将 MSV slot中描述的值删除后的数据分布(压缩的直方图)。当列的数据只有比较少值时, 用 MCV slot就能很好的描述数据分布, 直方图 slot应该省略。

□相关性 slot。描述此列数据的物理顺序和排序顺序之间的相关性。

staop, <操作符。

stavalues, 未使用, NULL。

stanumbers, 单值, 相关系数, 值的范围+1到-1。

4.8多租户技术实现

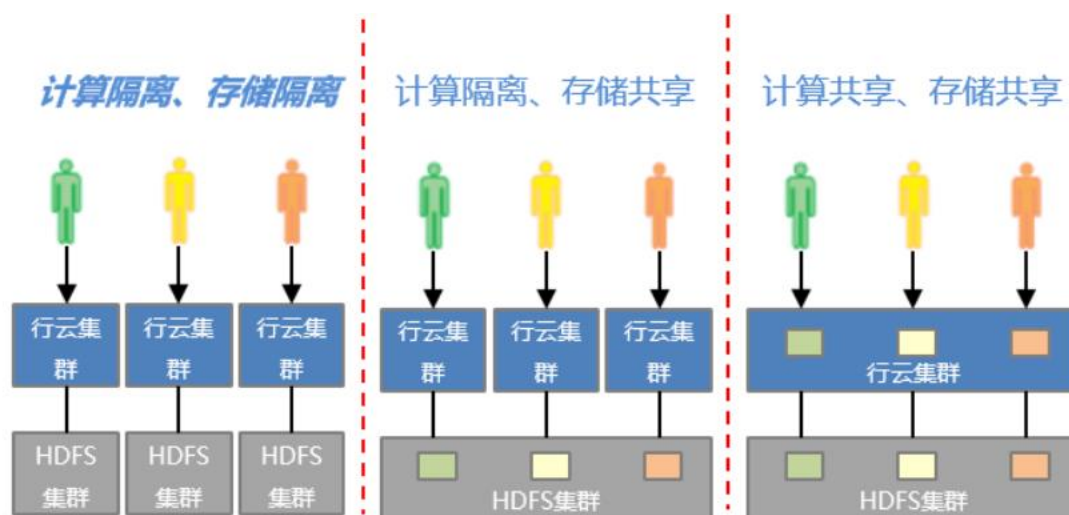


图 11 多租户资源管理

行云实现了多种级别的多租户隔离技术。

□ 技术特点:

- 基于 HDFS支持存储配额分配、文件数限制, 实现租户的存储资源管理,透过服务端配置 HDFS目录权限及租户存储配额, 限制租户 HDFS使用量, 达到租户存储资源分配隔离, 以及基于用户组的存储资源

共享

- 基于 MESOS服务的 CPU、内存分配, 实现租户的计算资源管理, 通过队列限制用户使用集群资源的比率
 - 基于租户共享元数据来实现数据共享机制
- 功能特性:
- 支持多种租户管理机制, 计算、存储隔离与共享
 - 支持租户动态扩容, 无须停机, 降低对线上服务的影响, 服务灵活平滑运行
 - 支持弹性计算, 保证执行效率
 - 统一管理介面, 整合数据库用户权限和资源配额管理机制, 降低操作复杂度, 实现统一对多租户资源分配能力

4.9 分布式存储过程执行引擎

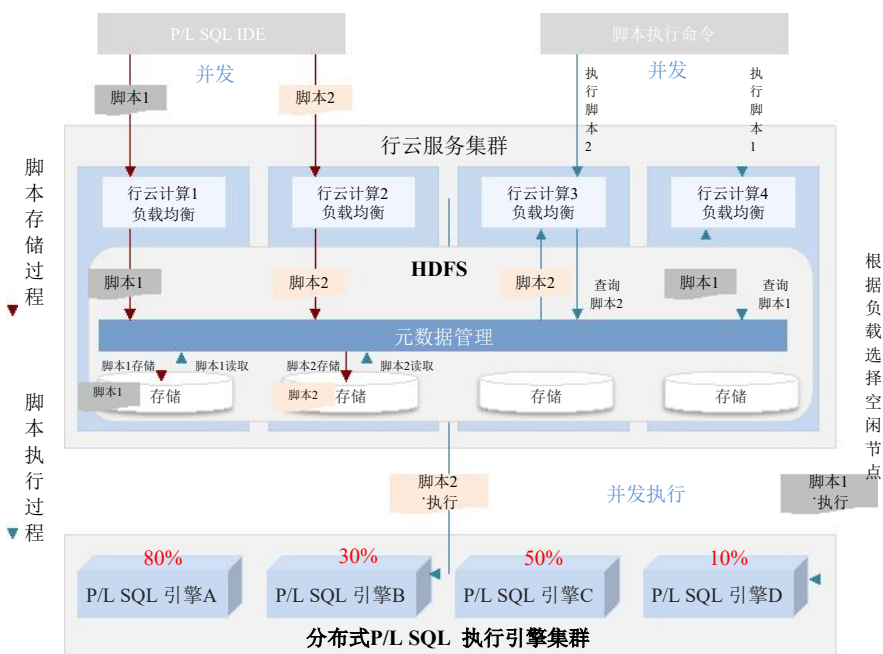


图 12 分布式存储过程执行引擎

针对海量数据批处理加工的任务需求，行云实现了分布式的 PL/SQL 服务。透过使用统一的 PL/SQL 引擎来完成用户的 SQL 请求的接入、解析与优化，以及后续执行计划的生成与分发，保证 SQL 支持的高度一致性与完整性。

- 可视化的 P/L SQL IDE环境

支持图形和文本两种方式建立 PL/SQL存储过程（完全支持 SQL92标准），并实现即时编译，同时为数据人员提供过程调试环境。同时支持 ORACLE、DB2 等异构数据库存储过程的无缝迁移。

- ## □ 高效的分布式执行引擎

首个采用分布式的 P/L SQL调度执行引擎，高并发能力，缩短执行时间，并支持线性无限扩展。

□ 动态的负载均衡技术

利用负载均衡技术，根据 P/L SQL执行引擎忙闲状态，灵活调度分配作业，提高资源利用。

□ 良好的容错机制

执行过程中出现一个执行节点宕机，执行引擎将会分配另一个节点执行任务。

4.10 跨地域数据中心架构

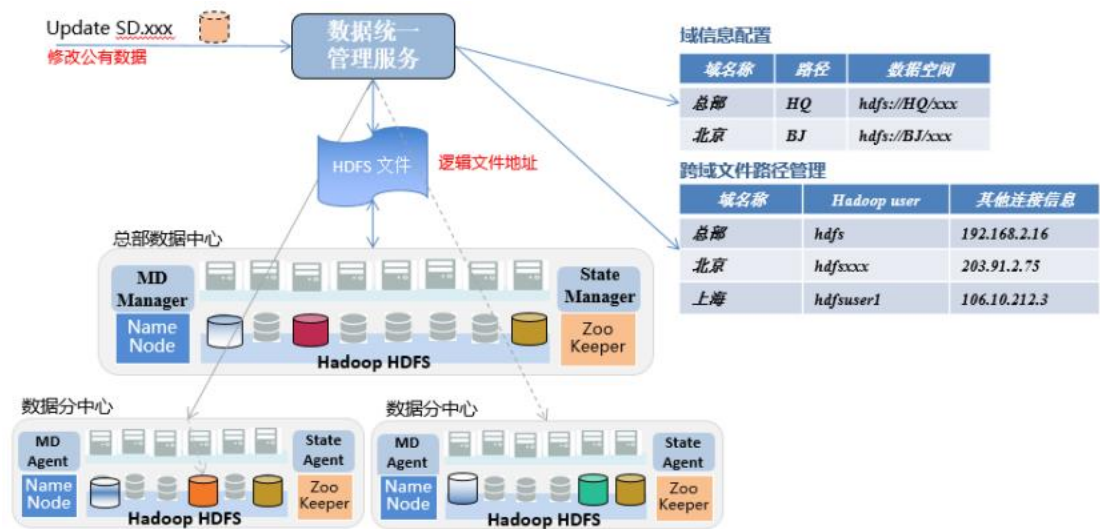


图 12 分布式存储过程执行引擎

行云提供了跨域统一数据管理的服务。

□ 域信息配置管理

为每个域的 HDFS配置信息，包括获取每个域所需的 IP、端口号、域名、访问账号等信息

□ 跨域文件路径管理

对不同域的数据路径进行整合，从而为不同域的数据提供统一的文件路径

□ 统一数据管理

修改公有数据，需通过统一数据管理服务，提交给总部数据中心获取权限许可和数据域信息，根据权限许可进行公有数据的修改；

4.11 数据联邦-整合多种异构数据源

Cirro Data支持多种异构数据源：

□支持 Oracle、MySQL、DB2、Hive等数据库，通过原生接口连接，性能

更高。

- 支持使用 JDBC的通用连接访问异构数据源，便于快速扩展新业务。
- 支持多个 Cirro Data相连组成一个共享的大数据平台。

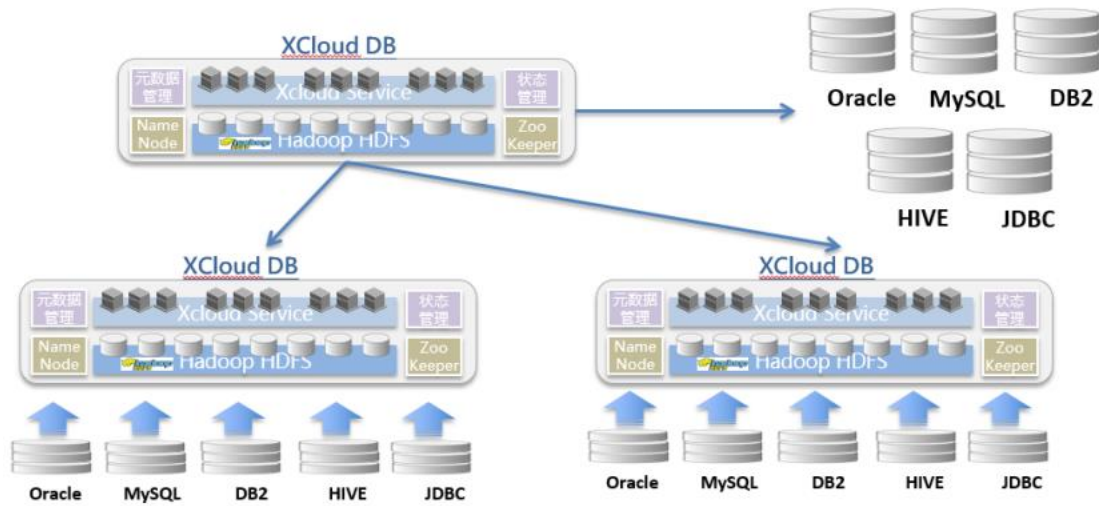


图 13 整合多种异构数据源

5.管理工具

5.1可视化集群管理工具

Cirro Data Enterprise Administrator是部署、监控、配置行云集群的管理控制台、并为 Cirro Data PL/SQL Developer提供代理服务。

XEA主要提供以下功能特点：

- 易部署，易使用。
 - 行云HDFS查询代理组件(HDFSMetaQueryProxy)、行云计算引擎(xpkg)、行云存储过程执行引擎(taskmanager)安装介质管理。
 - 多集群多版本管理。
 - 一键式部署、升级、管理向导。
 - 在线升级、扩容、缩容。
 - 灵活的扩容方式，支持克隆扩容、新增扩容。
 - 支持单节点多进程方式部署。
 - 统一配置管理中心，方便用户统一配置、优化管理各节点配置。
 - 友好的日志交互，帮助用户快速诊断定位问题。
 - 实时监控集群资源、节点状态。
 - 联机监控管理行云会话与SQL状态。
- 为Cirro Data PL/SQL Developer提供安全代理服务。

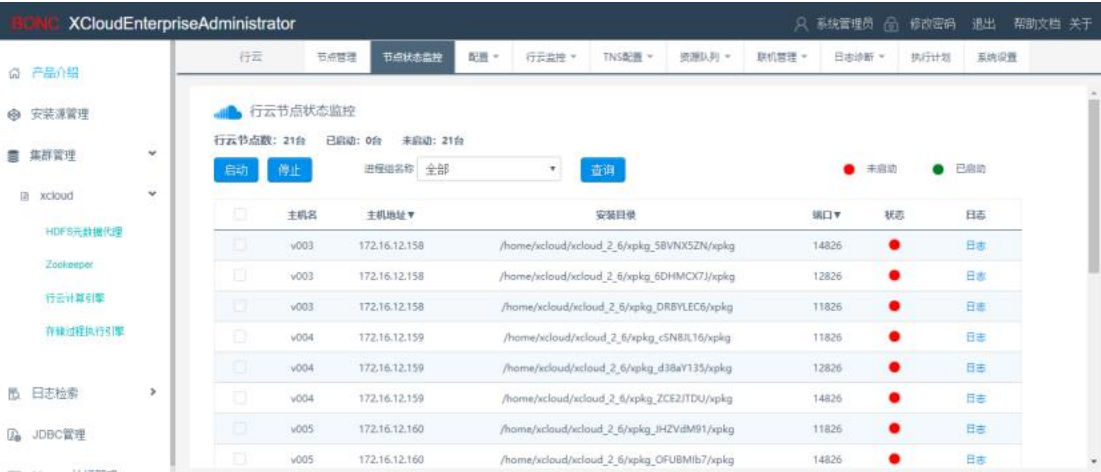


图 14 集群管理



图 15 集群监控

5.2 可视化SQL分析工具

行云提供的 Cirro Data PL/SQL Developer是一款全功能的数据查询及分析工具。支持图形和文本两种方式，完全支持 SQL92标准。

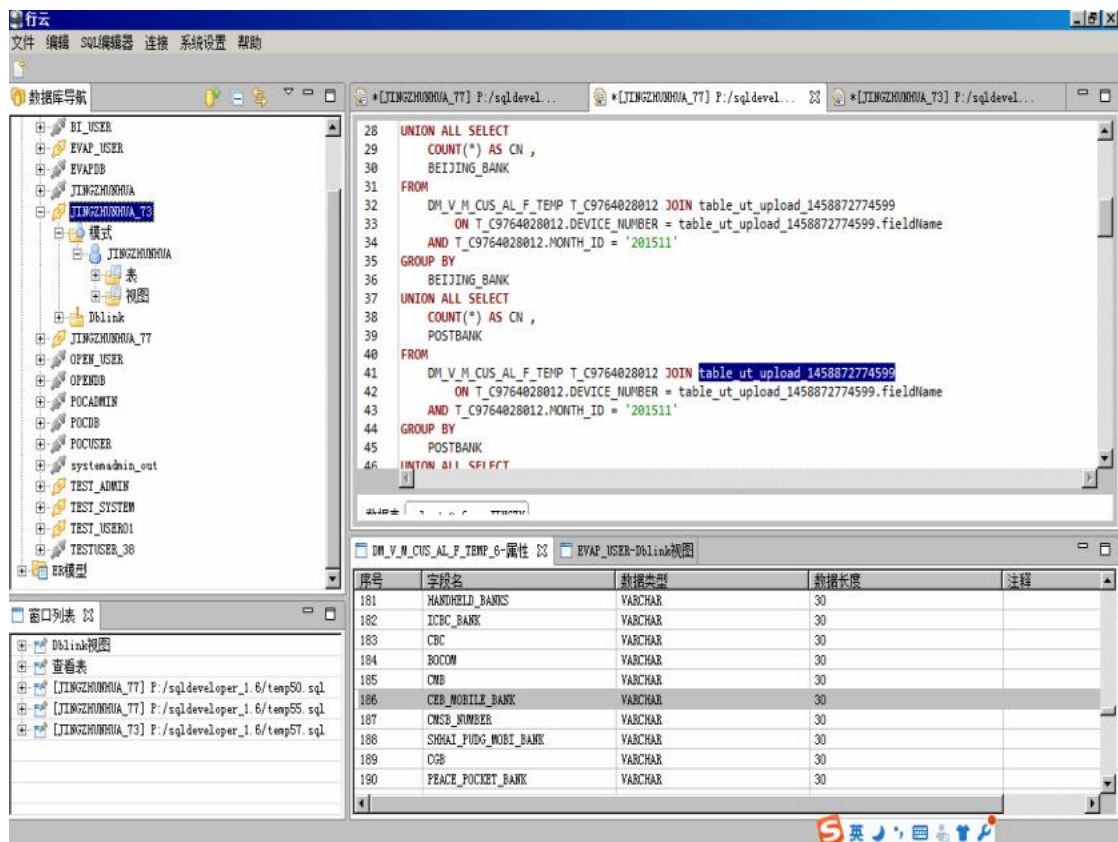


图 16 SQL 分析工具

- ☐ 管理数据库、表、DBLINK、索引等数据对象。
- ☐ 可视化创建用户、赋予用户权限、编辑用户和删除用户。
- ☐ 可视化浏览表中的数据记录。

- 可视化浏览数据表中的 Partition及 Slice。
- SQL脚本编写支持语法高亮显示、错误提示以及自动完成功能。
- 支持自动生成 SQL脚本对应的执行树，并对其进行分析。

同时，该分析工具还提供了存储过程编译、调测、执行功能等功能的可视化界面。

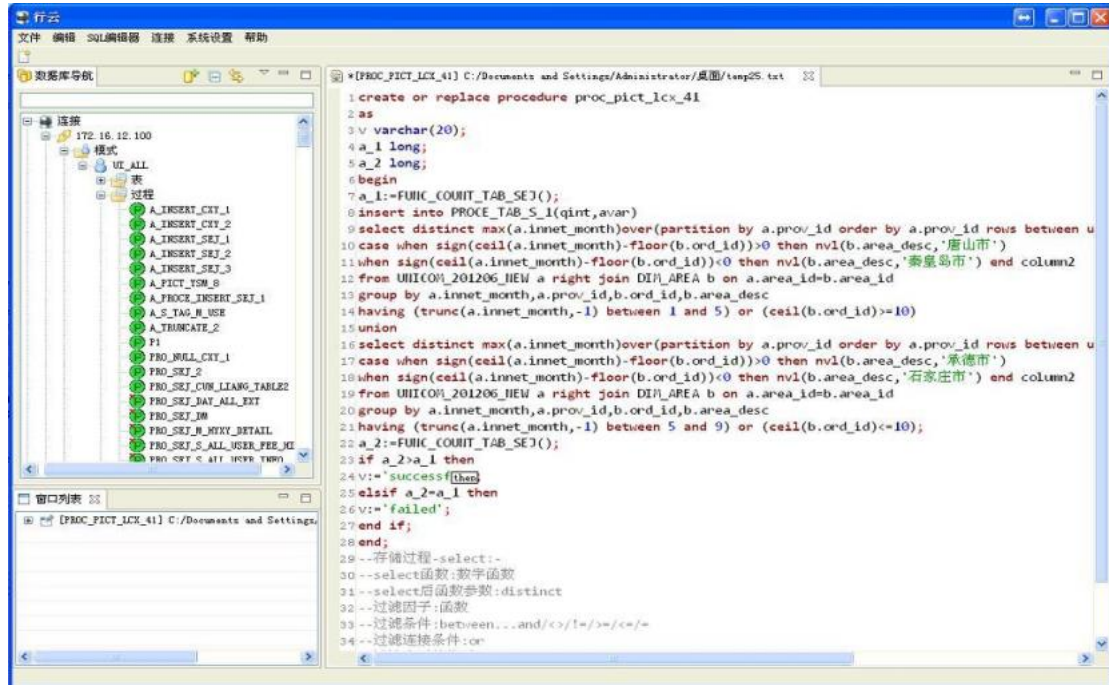


图 17 PL/SQL 编辑器

6.性能指标

针对下列各项主要的分析型应用特性，根据行云数据库在 20个节点集群上的性能测试结果，其性能指标为：

数据导入	每秒 GB级
聚合	千万级数据毫秒级，十亿级数据秒级
剔重	千万级数据毫秒级，亿级数据秒级
多维分析	千万级数据毫秒级，亿级数据秒级
大表关联	千万级数据毫秒级，亿级数据秒级
点查询	十亿级数据毫秒级

7.运行环境

- 硬件平台：普通的 X86服务器即可
- 操作系统：Redhat-6.4/6.5/7.1/7.2/7.3 (64bit)。
- Hadoop环境： hadoop 2.3.0~ Hadoop 2.8.x / CDH 5.0.1~CDH 5.13.x