
Apache ShardingSphere document

v5.2.0

Apache ShardingSphere

2022 年 09 月 08 日

1	概览	1
1.1	什么是 ShardingSphere	1
1.1.1	介绍	1
	ShardingSphere-JDBC	1
	ShardingSphere-Proxy	1
1.1.2	产品功能	2
1.1.3	产品优势	2
1.1.4	线路规划	3
1.1.5	如何参与	3
1.2	设计哲学	3
1.2.1	连接：打造数据库上层标准	4
1.2.2	增强：数据库计算增强引擎	4
1.2.3	可插拔：构建数据库功能生态	5
	L1 内核层	5
	L2 功能层	5
	L3 生态层	6
1.3	部署形态	6
1.3.1	部署形态	6
	ShardingSphere-JDBC 独立部署	6
	ShardingSphere-Proxy 独立部署	7
	混合部署架构	8
1.3.2	运行模式	9
	单机模式	9
	集群模式	9
2	快速入门	10
2.1	ShardingSphere-JDBC	10
2.1.1	应用场景	10
2.1.2	使用限制	10
2.1.3	前提条件	10
2.1.4	操作步骤	10

2.2	ShardingSphere-Proxy	12
2.2.1	应用场景	12
2.2.2	使用限制	12
2.2.3	前提条件	12
2.2.4	操作步骤	13
3	功能	14
3.1	数据分片	14
3.1.1	背景	14
	垂直分片	15
	水平分片	16
3.1.2	挑战	16
3.1.3	目标	17
3.1.4	应用场景	17
	海量数据高并发的 OLTP 场景	17
	海量数据实时分析 OLAP 场景	17
3.1.5	相关参考	17
3.1.6	核心概念	17
	表	17
	数据节点	19
	分片	20
	行表达式	21
	分布式主键	21
3.1.7	使用限制	22
	稳定支持	22
	实验性支持	23
	不支持	24
3.1.8	附录	25
3.2	分布式事务	25
3.2.1	背景	25
3.2.2	挑战	26
3.2.3	目标	26
3.2.4	原理介绍	26
	LOCAL 事务	26
	XA 事务	26
	BASE 事务	27
3.2.5	应用场景	28
	ShardingSphere XA 事务使用场景	28
	ShardingSphere BASE 事务使用场景	28
	ShardingSphere LOCAL 事务使用场景	28
3.2.6	相关参考	28
3.2.7	核心概念	29
	XA 协议	29
3.2.8	使用限制	29
	LOCAL 事务	29

	XA 事务	29
	BASE 事务	30
3.2.9	附录	30
3.3	读写分离	30
3.3.1	背景	30
3.3.2	挑战	31
3.3.3	目标	32
3.3.4	应用场景	32
	复杂的主从数据库架构	32
3.3.5	相关参考	32
3.3.6	核心概念	33
	主库	33
	从库	33
	主从同步	33
	负载均衡策略	33
3.3.7	使用限制	33
3.4	高可用	33
3.4.1	背景	33
3.4.2	挑战	34
3.4.3	目标	34
3.4.4	应用场景	34
3.4.5	相关参考	35
3.4.6	核心概念	35
	高可用类型	35
	动态读写分离	35
3.4.7	使用限制	35
	支持项	35
	不支持项	35
3.5	数据库网关	35
3.5.1	背景	35
3.5.2	挑战	36
3.5.3	目标	36
3.5.4	应用场景	36
3.5.5	核心概念	36
	SQL 方言	36
3.5.6	使用限制	36
3.6	流量治理	36
3.6.1	背景	36
3.6.2	挑战	37
3.6.3	目标	37
3.6.4	应用场景	37
	计算节点过载保护	37
	存储节点限流	37
3.6.5	核心概念	37
	熔断	37

	限流	37
3.7	数据迁移	38
3.7.1	背景	38
3.7.2	挑战	38
3.7.3	目标	38
3.7.4	应用场景	38
3.7.5	相关参考	38
3.7.6	核心概念	39
	节点	39
	集群	39
	源端	39
	目标端	39
	数据迁移作业	39
	存量数据	39
	增量数据	39
3.7.7	使用限制	39
	支持项	39
	不支持项	40
3.8	数据加密	40
3.8.1	背景	40
3.8.2	挑战	40
3.8.3	目标	41
3.8.4	应用场景	41
	新上线业务	41
	成熟业务	41
3.8.5	相关参考	41
3.8.6	核心概念	41
	逻辑列	41
	密文列	41
	查询辅助列	42
	明文列	42
3.8.7	使用限制	42
3.8.8	附录	42
3.9	影子库	42
3.9.1	背景	42
3.9.2	挑战	42
3.9.3	目标	43
3.9.4	应用场景	43
3.9.5	相关参考	43
3.9.6	核心概念	43
	生产库	43
	影子库	43
	影子算法	43
3.9.7	使用限制	44
	基于 Hint 的影子算法	44

基于列的影子算法	44
3.10 可观察性	44
3.10.1 背景	44
3.10.2 挑战	46
3.10.3 目标	46
3.10.4 应用场景	47
监控仪表盘	47
应用性能监控	47
应用链路追踪	47
3.10.5 相关参考	47
3.10.6 核心概念	47
Agent	47
APM	48
Tracing	48
Metrics	48
Logging	48
4 用户手册	49
4.1 ShardingSphere-JDBC	49
4.1.1 YAML 配置	49
简介	49
使用步骤	50
语法说明	51
模式配置	51
数据源配置	53
规则配置	54
算法配置	72
JDBC 驱动	74
4.1.2 Java API	76
简介	76
使用步骤	76
模式配置	77
数据源配置	80
规则配置	81
算法配置	101
4.1.3 Spring Boot Starter	103
简介	103
使用步骤	103
模式配置	104
数据源配置	105
规则配置	107
算法配置	123
4.1.4 Spring 命名空间	125
简介	125
使用步骤	125

	模式配置	126
	数据源配置	129
	规则配置	130
	算法配置	150
4.1.5	特殊 API	152
	数据分片	152
	读写分离	154
	分布式事务	156
4.1.6	不支持项	167
	DataSource 接口	167
	Connection 接口	167
	Statement 和 PreparedStatement 接口	167
	ResultSet 接口	167
	JDBC 4.1	168
4.2	ShardingSphere-Proxy	168
4.2.1	启动手册	168
	使用二进制发布包	168
	使用 Docker	170
	使用 Helm	171
4.2.2	YAML 配置	178
	权限	179
	属性配置	180
	规则配置	181
4.2.3	DistSQL	181
	定义	181
	相关概念	181
	对系统的影响	182
	使用限制	183
	原理介绍	183
	相关参考	184
	语法	184
	使用	225
4.2.4	数据迁移	232
	简介	232
	运行部署	232
	使用手册	236
4.2.5	可观察性	248
	源码编译	248
	agent 配置	248
	ShardingSphere-Proxy 中使用	251
4.3	通用配置	252
4.3.1	属性配置	252
	背景信息	252
	参数解释	252
	操作步骤	252

	配置示例	252
4.3.2	内置算法	252
	简介	252
	使用方式	253
	元数据持久化仓库	253
	分片算法	255
	分布式序列算法	259
	负载均衡算法	262
	加密算法	264
	影子算法	266
	SQL 翻译	268
4.4	错误码	269
4.4.1	SQL 错误码	269
4.4.2	服务器错误码	271
5	开发者手册	272
5.1	运行模式	272
5.1.1	StandalonePersistRepository	272
	全限定类名	272
	定义	272
	已知实现	272
5.1.2	ClusterPersistRepository	273
	全限定类名	273
	定义	273
	已知实现	273
5.1.3	GovernanceWatcher	273
	全限定类名	273
	定义	273
	已知实现	274
5.2	配置	275
5.2.1	RuleBuilder	275
	全限定类名	275
	定义	275
	已知实现	276
5.2.2	YamlRuleConfigurationSwapper	277
	全限定类名	277
	定义	277
	已知实现	278
5.2.3	ShardingSphereYamlConstruct	279
	全限定类名	279
	定义	279
	已知实现	279
5.3	内核	279
5.3.1	SQLRouter	279
	全限定类名	279

	定义	279
	已知实现	280
5.3.2	SQLRewriteContextDecorator	280
	全限定类名	280
	定义	280
	已知实现	280
5.3.3	SQLExecutionHook	280
	全限定类名	280
	定义	281
	已知实现	281
5.3.4	ResultProcessEngine	281
	全限定类名	281
	定义	281
	已知实现	281
5.4	数据源	281
5.4.1	DatabaseType	281
	全限定类名	281
	定义	282
	已知实现	282
5.4.2	DialectSchemaMetaDataLoader	282
	全限定类名	282
	定义	282
	已知实现	283
5.4.3	DataSourcePoolMetaData	283
	全限定类名	283
	定义	283
	已知实现	283
5.4.4	DataSourcePoolActiveDetector	284
	全限定类名	284
	定义	284
	已知实现	284
5.5	SQL 解析	284
5.5.1	DatabaseTypedSQLParserFacade	284
	全限定类名	284
	定义	284
	已知实现	285
5.5.2	SQLVisitorFacade	285
	全限定类名	285
	定义	285
	已知实现	286
5.6	代理端	286
5.6.1	DatabaseProtocolFrontendEngine	286
	全限定类名	286
	定义	286
	已知实现	286

5.6.2	AuthorityProvideAlgorithm	287
	全限定类名	287
	定义	287
	已知实现	287
5.7	数据分片	287
5.7.1	ShardingAlgorithm	287
	全限定类名	287
	定义	287
	已知实现	288
5.7.2	KeyGenerateAlgorithm	289
	全限定类名	289
	定义	289
	已知实现	289
5.7.3	ShardingAuditAlgorithm	289
	全限定类名	289
	定义	289
	已知实现	289
5.7.4	DatetimeService	290
	全限定类名	290
	定义	290
	已知实现	290
5.8	读写分离	290
5.8.1	ReadQueryLoadBalanceAlgorithm	290
	全限定类名	290
	定义	290
	已知实现	291
5.9	高可用	292
5.9.1	DatabaseDiscoveryProviderAlgorithm	292
	全限定类名	292
	定义	292
	已知实现	292
5.10	分布式事务	292
5.10.1	ShardingSphereTransactionManager	292
	全限定类名	292
	定义	292
	已知实现	293
5.10.2	XATransactionManagerProvider	293
	全限定类名	293
	定义	293
	已知实现	293
5.10.3	XADatasourceDefinition	294
	全限定类名	294
	定义	294
	已知实现	294
5.10.4	DatasourcePropertyProvider	294

	全限定类名	294
	定义	295
	已知实现	295
5.11	SQL 检查	295
5.11.1	全限定类名	295
5.11.2	定义	295
5.11.3	已知实现	295
5.12	数据加密	295
5.12.1	EncryptAlgorithm	295
	全限定类名	295
	定义	296
	已知实现	296
5.13	影子库	296
5.13.1	ShadowAlgorithm	296
	全限定类名	296
	定义	296
	已知实现	297
5.14	可观察性	297
5.14.1	PluginBootService	297
	全限定类名	297
	定义	297
	已知实现	297
5.14.2	PluginDefinitionService	298
	全限定类名	298
	定义	298
	已知实现	298
6	测试手册	299
6.1	整合测试	299
6.2	模块测试	299
6.3	性能测试	299
6.4	集成测试	299
6.4.1	设计	299
	测试用例	300
	测试环境	300
	测试引擎	300
6.4.2	使用指南	301
	测试用例配置	301
	环境配置	302
	运行测试引擎	302
6.5	性能测试	304
6.5.1	Sysbench ShardingSphere Proxy 空 Rules 性能测试	304
	测试目的	304
	测试环境搭建	304
	测试阶段	306

6.5.2	BenchmarkSQL ShardingSphere Proxy 分片性能测试	307
	测试目的	307
	测试方法	307
	测试工具微调	308
	压测环境或参数建议	308
	附录	310
	BenchmarkSQL 5.0 PostgreSQL 语句列表	312
6.6	模块测试	321
6.6.1	SQL 解析测试	321
	数据准备	321
6.6.2	SQL 改写测试	322
	目标	322
6.7	Scaling 集成测试	324
6.7.1	测试目的	324
6.7.2	测试环境	324
6.7.3	使用指南	324
	环境配置	324
	测试用例	324
	运行测试用例	325
7	技术参考	326
7.1	数据兼容性	326
7.2	数据库网关	327
7.3	管控	327
7.3.1	注册中心数据结构	327
	/rules	328
	/props	329
	/metadata/databaseName/versions/{versionNumber}/dataSources	329
	/metadata/databaseName/versions/{versionNumber}/rules	330
	/metadata/databaseName/schemas/{schemaName}/tables	330
	/nodes/compute_nodes	330
	/nodes/storage_nodes	331
7.4	数据分片	331
7.4.1	SQL 解析	332
7.4.2	SQL 路由	332
7.4.3	SQL 改写	332
7.4.4	SQL 执行	332
7.4.5	结果归并	332
7.4.6	查询优化	332
7.4.7	解析引擎	332
	抽象语法树	332
	SQL 解析引擎	333
7.4.8	路由引擎	337
	分片路由	337
	广播路由	339

7.4.9	改写引擎	341
	正确性改写	341
	优化改写	346
7.4.10	执行引擎	347
	连接模式	347
	自动化执行引擎	348
7.4.11	归并引擎	351
	遍历归并	352
	排序归并	352
	分组归并	353
	聚合归并	356
	分页归并	356
7.5	分布式事务	357
7.5.1	导览	357
7.5.2	XA 事务	357
	开启全局事务	358
	执行真实分片 SQL	358
	提交或回滚事务	359
7.5.3	Seata 柔性事务	359
	引擎初始化	360
	开启全局事务	360
	执行真实分片 SQL	360
	提交或回滚事务	361
7.6	数据迁移	361
7.6.1	原理说明	361
7.6.2	执行阶段说明	362
	准备阶段	362
	存量数据迁移阶段	362
	增量数据同步阶段	362
	流量切换阶段	362
7.6.3	相关参考	362
7.7	数据加密	362
7.7.1	处理流程详解	362
	整体架构	363
	加密规则	363
	加密处理过程	365
7.7.2	解决方案详解	367
	新上线业务	367
	已上线业务改造	368
7.7.3	中间件加密服务优势	371
7.7.4	加密算法解析	372
	EncryptAlgorithm	372
7.8	影子库	372
7.8.1	原理介绍	372
	DML 语句	373

DDL 语句	373
7.8.2 相关参考	374
7.9 可观察性	374
7.9.1 原理说明	374
7.10 DistSQL	375
7.10.1 语法	375
RDL 语法	375
RQL 语法	406
RAL 语法	418
保留字	418
7.11 基础架构	420
8 FAQ	421
8.1 JDBC	421
8.1.1 JDBC 为什么配置了某个数据连接池的 spring-boot-starter (比如 druid) 和 shardingsphere-jdbc-spring-boot-starter 时, 系统启动会报错?	421
8.1.2 JDBC 使用 Spring 命名空间时找不到 xsd?	421
8.1.3 JDBC 引入 shardingsphere-transaction-xa-core 后, 如何避免 spring-boot 自动加载默认的 JtaTransactionManager?	421
8.2 Proxy	422
8.2.1 Proxy Windows 环境下, 运行 ShardingSphere-Proxy, 找不到或无法加载主类 org.apache.shardingsphere.proxy.Bootstrap, 如何解决?	422
8.2.2 Proxy 在使用 ShardingSphere-Proxy 的时候, 如何动态在添加新的逻辑库?	422
8.2.3 Proxy 在使用 ShardingSphere-Proxy 时, 怎么使用合适的工具连接到 ShardingSphere-Proxy?	423
8.2.4 Proxy 使用 Navicat 等第三方数据库工具连接 ShardingSphere-Proxy 时, 如果 ShardingSphere-Proxy 没有创建 Database 或者没有添加 Resource, 连接失败?	423
8.3 分片	423
8.3.1 分片 Cloud not resolve placeholder ...in string value ...异常的解决方法?	423
8.3.2 分片 inline 表达式返回结果为何出现浮点数?	424
8.3.3 分片如果只有部分数据库分库分表, 是否需要将不分库分表的表也配置在分片规则中?	424
8.3.4 分片指定了泛型为 Long 的 SingleKeyTableShardingAlgorithm, 遇到 ClassCastException: Integer can not cast to Long?	424
8.3.5 [分片、PROXY] 实现 StandardShardingAlgorithm 自定义算法时, 指定了 Comparable 的具体类型为 Long, 且数据库表中字段类型为 bigint, 出现 ClassCastException: Integer can not cast to Long 异常。	424
8.3.6 分片 ShardingSphere 提供的默认分布式自增主键策略为什么是不连续的, 且尾数大多为偶数?	424
8.3.7 分片如何在 inline 分表策略时, 允许执行范围查询操作 (BETWEEN AND、>、<、>=、<=)?	425
8.3.8 分片为什么我实现了 KeyGenerateAlgorithm 接口, 也配置了 Type, 但是自定义的分布式主键依然不生效?	425

8.3.9	分片 ShardingSphere 除了支持自带的分布式自增主键之外, 还能否支持原生的自增主键?	425
8.4	数据加密	425
8.4.1	数据加密 JPA 和数据加密无法一起使用, 如何解决?	425
8.5	DistSQL	426
8.5.1	DistSQL 使用 DistSQL 添加数据源时, 如何设置自定义的 JDBC 连接参数或连接池属性?	426
8.5.2	DistSQL 使用 DistSQL 删除资源时, 出现 Resource [xxx] is still used by [SingleTableRule]。	426
8.5.3	DistSQL 使用 DistSQL 添加资源时, 出现 Failed to get driver instance for jdbcURL=xxx。	426
8.6	其他	426
8.6.1	其他如果 SQL 在 ShardingSphere 中执行不正确, 该如何调试?	426
8.6.2	其他阅读源码时为什么会出现编译错误? IDEA 不索引生成的代码?	427
8.6.3	其他使用 SQLSever 和 PostgreSQL 时, 聚合列不加别名会抛异常?	427
8.6.4	其他 Oracle 数据库使用 Timestamp 类型的 Order By 语句抛出异常提示 “Order by value must implements Comparable” ?	427
8.6.5	其他 Windows 环境下, 通过 Git 克隆 ShardingSphere 源码时为什么提示文件名过长, 如何解决?	428
8.6.6	其他 Type is required 异常的解决方法?	429
8.6.7	其他服务启动时如何加快 metadata 加载速度?	429
8.6.8	其他 ANTLR 插件在 src 同级目录下生成代码, 容易误提交, 如何避免?	429
8.6.9	其他使用 Proxool 时分库结果不正确?	430
8.6.10	其他使用 Spring Boot 2.x 集成 ShardingSphere 时, 配置文件中的属性设置不生效?	430
9	下载	433
9.1	最新版本	433
9.1.1	Apache ShardingSphere - 版本: 5.2.0 (发布日期: Sept 5th, 2022)	433
9.2	全部版本	433
9.3	校验版本	433

本章介绍 Apache ShardingSphere 的定义，设计哲学和部署形态。

更多常见问题，请参考 [FAQ](#)。

1.1 什么是 ShardingSphere

1.1.1 介绍

Apache ShardingSphere 是一款开源的分布式数据库生态项目，由 JDBC 和 Proxy 两款产品组成。其核心采用微内核 + 可插拔架构，通过插件开放扩展功能。它提供多源异构数据库增强平台，进而围绕其上层构建生态。

Apache ShardingSphere 设计哲学为 Database Plus，旨在构建异构数据库上层的标准和生态。它关注如何充分合理地利用数据库的计算和存储能力，而并非实现一个全新的数据库。它站在数据库的上层视角，关注它们之间的协作多于数据库自身。

ShardingSphere-JDBC

ShardingSphere-JDBC 定位为轻量级 Java 框架，在 Java 的 JDBC 层提供的额外服务。

ShardingSphere-Proxy

ShardingSphere-Proxy 定位为透明化的数据库代理端，通过实现数据库二进制协议，对异构语言提供支持。

1.1.2 产品功能

特性	定义
数据分片	数据分片，是应对海量数据存储与计算的有效手段。ShardingSphere 提供基于底层数据库之上，可计算与存储水平扩展的分布式数据库解决方案。
分布式事务	事务能力，是保障数据库完整、安全的关键技术，也是数据库的核心技术之一。ShardingSphere 提供在单机数据库之上的分布式事务能力，可实现跨底层数据源的数据安全。
读写分离	读写分离，是应对高压力业务访问的手段之一。ShardingSphere 基于对 SQL 语义理解及底层数据库拓扑感知能力，提供灵活、安全的读写分离能力，且可实现读访问的负载均衡。
高可用	高可用，是对数据存储计算平台的基本要求。ShardingSphere 基于无状态服务，提供高可用计算服务访问；同时可感知并利用底层数据库自身高可用实现整体的高可用能力。
数据迁移	数据迁移，是打通数据生态的关键能力。ShardingSphere 提供基于数据全场景的迁移能力，可应对业务数据量激增的场景。
联邦查询	联邦查询，是面对复杂数据环境下利用数据的有效手段之一。ShardingSphere 提供跨数据源的复杂数据查询分析能力，简化并提升数据使用体验。
数据加密	数据加密，是保证数据安全的基本手段。ShardingSphere 提供一套完整的、透明化、安全的、低改造成本的数据加密解决方案。
影子库	在全链路压测场景下，ShardingSphere 通过影子库功能支持在复杂压测场景下数据隔离，压测获得测试结果可准确反应系统真实容量和性能水平。

1.1.3 产品优势

- 极致性能

驱动程序端历经长年打磨，效率接近原生 JDBC，性能极致。

- 生态兼容

代理端支持任何通过 MySQL/PostgreSQL 协议的应用访问，驱动程序端可对接任意实现 JDBC 规范的数据库。

- 业务零侵入

面对数据库替换场景，ShardingSphere 可满足业务无需改造，实现平滑业务迁移。

- 运维低成本

在保留原技术栈不变前提下，对 DBA 学习、管理成本低，交互友好。

- 安全稳定

基于成熟数据库底座之上提供增量能力，兼顾安全性及稳定性。

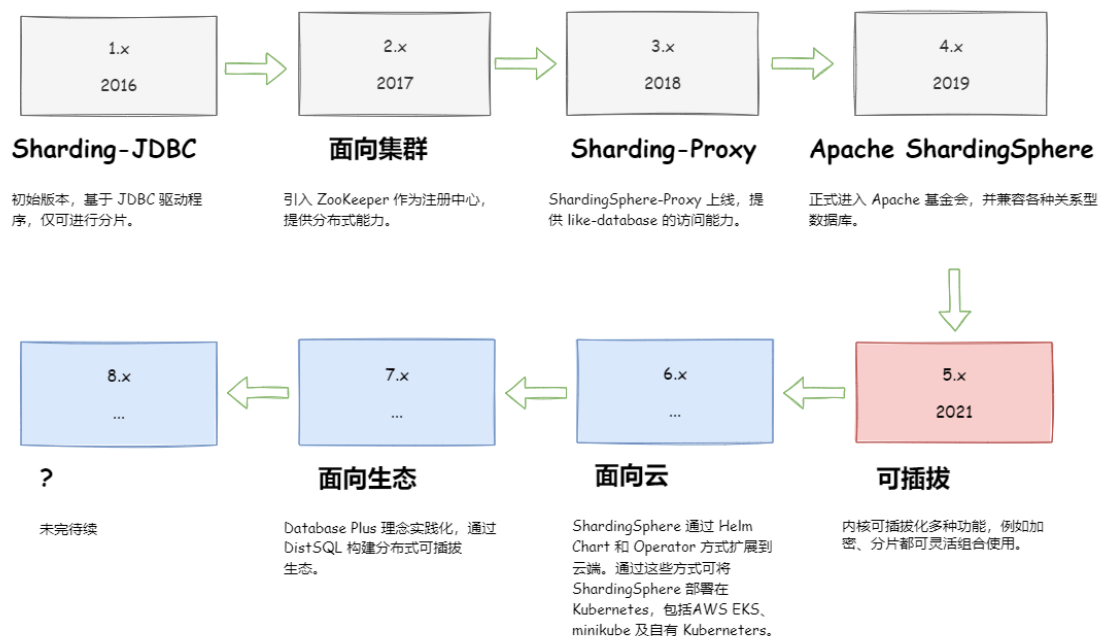
- 弹性扩展

具备计算、存储平滑在线扩展能力，可满足业务多变的需求。

- 开放生态

通过多层次（内核、功能、生态）插件化能力，为用户提供可定制满足自身特殊需求的独有系统。

1.1.4 线路规划



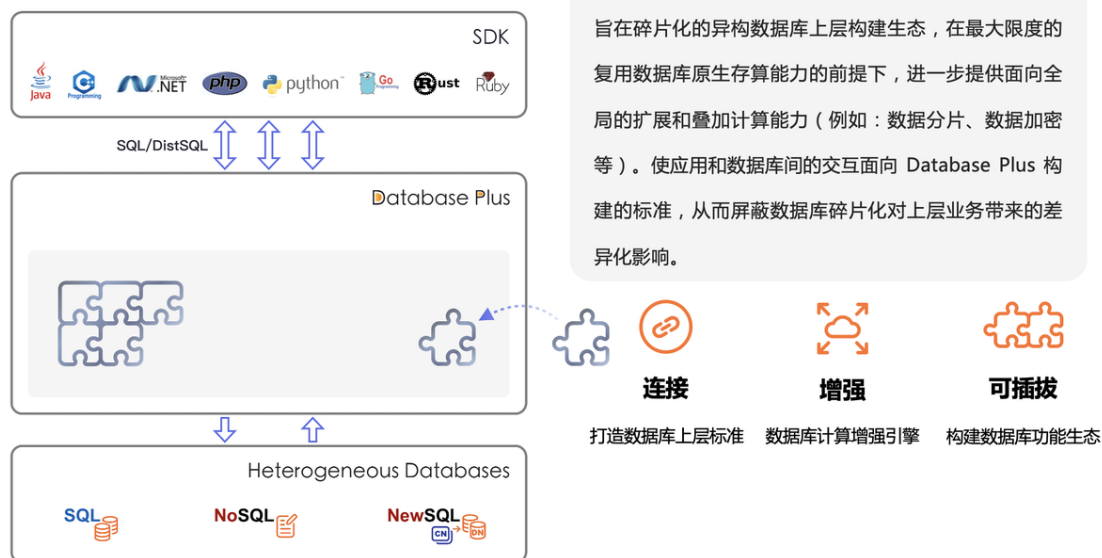
1.1.5 如何参与

ShardingSphere 已于 2020 年 4 月 16 日成为 Apache 软件基金会的顶级项目。欢迎通过[邮件列表](#)参与讨论。

1.2 设计哲学

ShardingSphere 采用 Database Plus 设计哲学，该理念致力于构建数据库上层的标准和生态，在生态中补充数据库所缺失的能力。

设计哲学：Database Plus



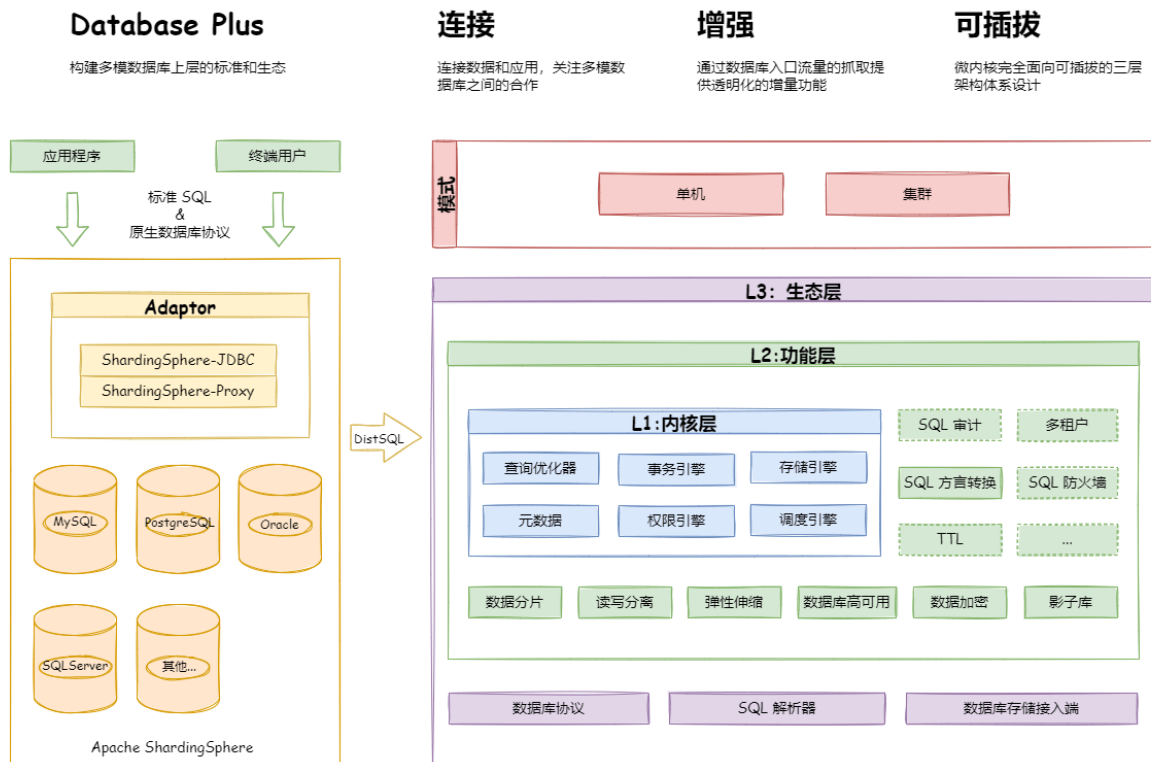
1.2.1 连接：打造数据库上层标准

通过对数据库协议、SQL 方言以及数据库存储的灵活适配，快速构建多模异构数据库上层的标准，同时通过内置 DistSQL 为应用提供标准化的连接方式。

1.2.2 增强：数据库计算增强引擎

在原生数据库基础能力之上，提供分布式及流量增强方面的能力。前者可突破底层数据库在计算与存储上的瓶颈，后者通过对流量的变形、重定向、治理、鉴权及分析能力提供更为丰富的数据应用增强能力。

1.2.3 可插拔：构建数据库功能生态



Apache ShardingSphere 的可插拔架构划分为 3 层，它们是：L1 内核层、L2 功能层、L3 生态层。

L1 内核层

是数据库基本能力的抽象，其所有组件均必须存在，但具体实现方式可通过可插拔的方式更换。主要包括查询优化器、分布式事务引擎、分布式执行引擎、权限引擎和调度引擎等。

L2 功能层

用于提供增量能力，其所有组件均是可选的，可以包含零至多个组件。组件之间完全隔离，互无感知，多组件可通过叠加的方式相互配合使用。主要包括数据分片、读写分离、数据库高可用、数据加密、影子库等。用户自定义功能可全面面向 Apache ShardingSphere 定义的顶层接口进行定制化扩展，而无需改动内核代码。

L3 生态层

用于对接和融入现有数据库生态, 包括数据库协议、SQL 解析器和存储适配器, 分别对应于 Apache ShardingSphere 以数据库协议提供服务的方式、SQL 方言操作数据的方式以及对接存储节点的数据库类型。

1.3 部署形态

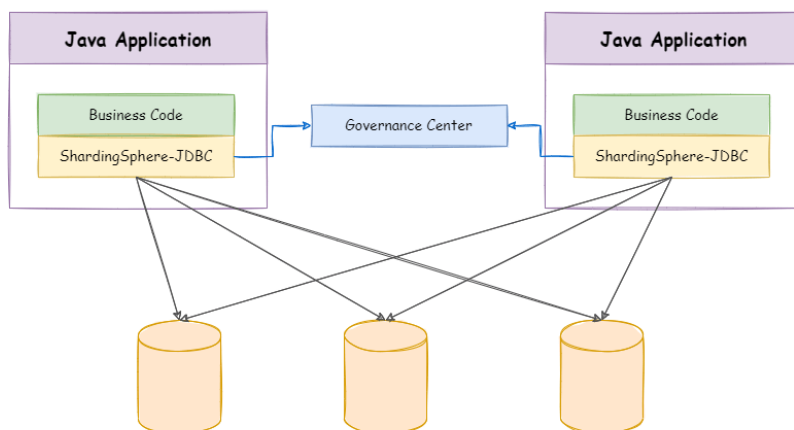
1.3.1 部署形态

Apache ShardingSphere 由 ShardingSphere-JDBC 和 ShardingSphere-Proxy 这 2 款既能够独立部署, 又支持混合部署配合使用的产品组成。它们均提供标准化的基于数据库作为存储节点的增量功能, 可适用于如 Java 同构、异构语言、云原生等各种多样化的应用场景。

ShardingSphere-JDBC 独立部署

ShardingSphere-JDBC 定位为轻量级 Java 框架, 在 Java 的 JDBC 层提供的额外服务。它使用客户端直连数据库, 以 jar 包形式提供服务, 无需额外部署和依赖, 可理解为增强版的 JDBC 驱动, 完全兼容 JDBC 和各种 ORM 框架。

- 适用于任何基于 JDBC 的 ORM 框架, 如: JPA, Hibernate, Mybatis, Spring JDBC Template 或直接使用 JDBC;
- 支持任何第三方的数据库连接池, 如: DBCP, C3P0, BoneCP, HikariCP 等;
- 支持任意实现 JDBC 规范的数据库, 目前支持 MySQL, PostgreSQL, Oracle, SQLServer 以及任何可使用 JDBC 访问的数据库。

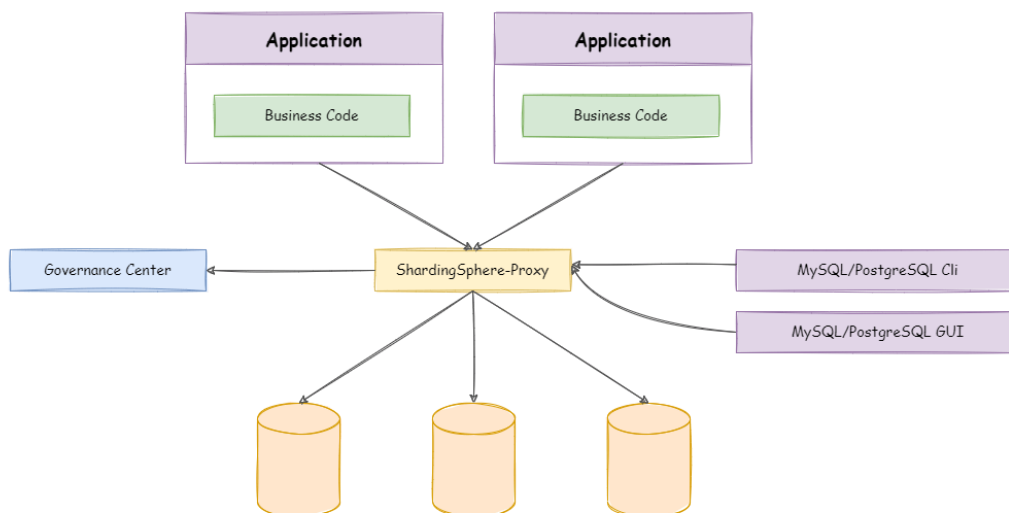


	ShardingSphere-JDBC	ShardingSphere-Proxy
数据库	任意	MySQL/PostgreSQL
连接消耗数	高	低
异构语言	仅 Java	任意
性能	损耗低	损耗略高
无中心化	是	否
静态入口	无	有

ShardingSphere-Proxy 独立部署

ShardingSphere-Proxy 定位为透明化的数据库代理端，通过实现数据库二进制协议，对异构语言提供支持。目前提供 MySQL 和 PostgreSQL 协议，透明化数据库操作，对 DBA 更加友好。

- 向应用程序完全透明，可直接当做 MySQL/PostgreSQL 使用；
- 兼容 MariaDB 等基于 MySQL 协议的数据库，以及 openGauss 等基于 PostgreSQL 协议的数据库；
- 适用于任何兼容 MySQL/PostgreSQL 协议的客户端，如：MySQL Command Client, MySQL Workbench, Navicat 等。

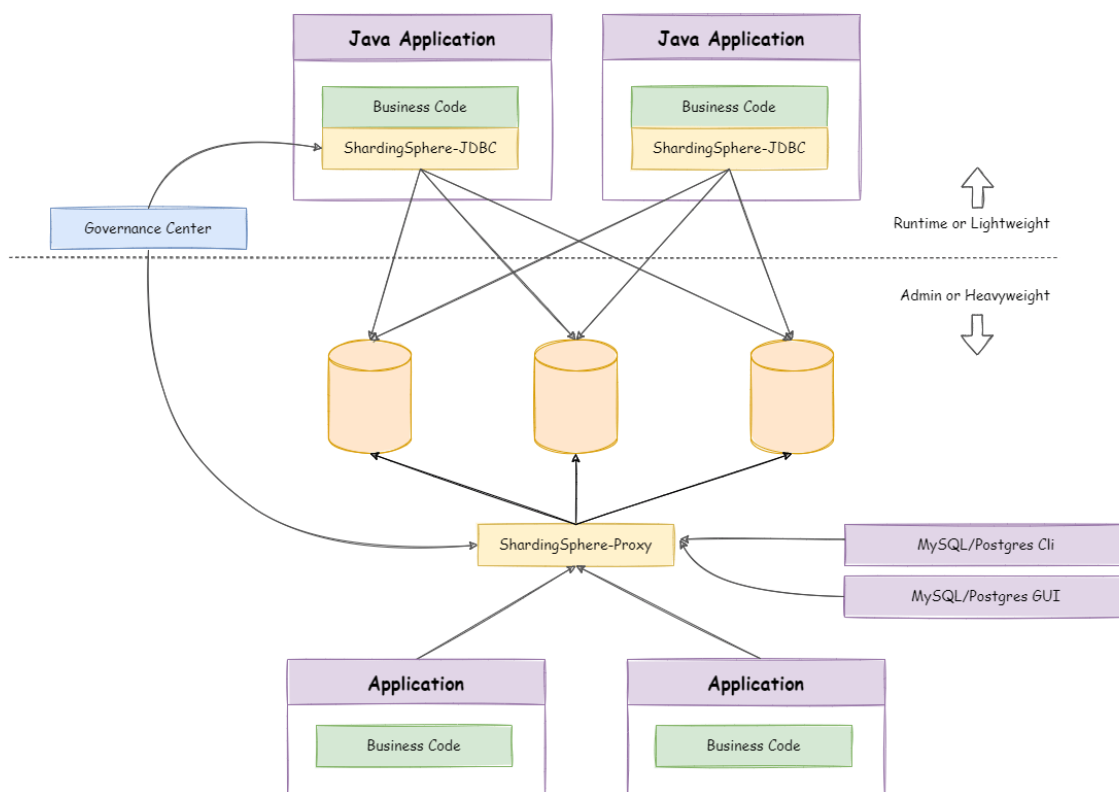


	ShardingSphere-JDBC	ShardingSphere-Proxy
数据库	任意	MySQL/PostgreSQL
连接消耗数	高	低
异构语言	仅 Java	任意
性能	损耗低	损耗略高
无中心化	是	否
静态入口	无	有

混合部署架构

ShardingSphere-JDBC 采用无中心化架构，与应用程序共享资源，适用于 Java 开发的高性能的轻量级 OLTP 应用；ShardingSphere-Proxy 提供静态入口以及异构语言的支持，独立于应用程序部署，适用于 OLAP 应用以及对分片数据库进行管理和运维的场景。

Apache ShardingSphere 是多接入端共同组成的生态圈。通过混合使用 ShardingSphere-JDBC 和 ShardingSphere-Proxy，并采用同一注册中心统一配置分片策略，能够灵活的搭建适用于各种场景的应用系统，使得架构师更加自由地调整适合于当前业务的最佳系统架构。



1.3.2 运行模式

Apache ShardingSphere 提供了两种运行模式，分别是单机模式和集群模式。

单机模式

能够将数据源和规则等元数据信息持久化，但无法将元数据同步至多个 Apache ShardingSphere 实例，无法在集群环境中相互感知。通过某一实例更新元数据之后，会导致其他实例由于获取不到最新的元数据而产生不一致的错误。

适用于工程师在本地搭建 Apache ShardingSphere 环境。

集群模式

提供了多个 Apache ShardingSphere 实例之间的元数据共享和分布式场景下状态协调的能力。它能够提供计算能力水平扩展和高可用等分布式系统必备的能力，集群环境需要通过独立部署的注册中心来存储元数据和协调节点状态。

在生产环境建议使用集群模式。

本章节以尽量短的时间，为使用者提供最简单的 Apache ShardingSphere 的快速入门。

示例代码：<https://github.com/apache/shardingsphere/tree/master/examples>

2.1 ShardingSphere-JDBC

2.1.1 应用场景

Apache ShardingSphere-JDBC 可以通过 Java, YAML, Spring 命名空间和 Spring Boot Starter 这 4 种方式进行配置，开发者可根据场景选择适合的配置方式。

2.1.2 使用限制

目前仅支持 JAVA 语言

2.1.3 前提条件

开发环境需要具备 Java JRE 8 或更高版本。

2.1.4 操作步骤

1. 规则配置。

详情请参见用户手册。

2. 引入 maven 依赖。

```
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-jdbc-core-spring-boot-starter</artifactId>
  <version>${latest.release.version}</version>
</dependency>
```

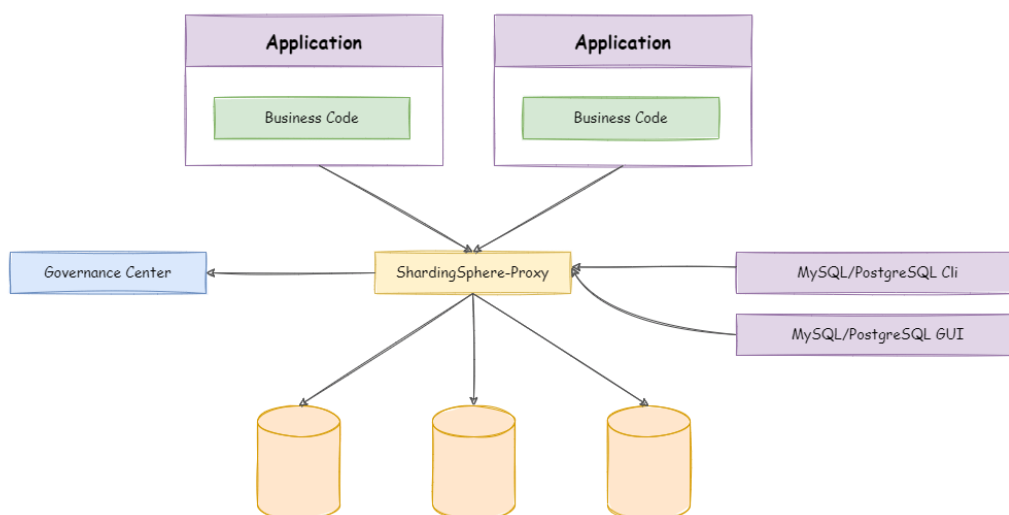
注意：请将 `${latest.release.version}` 更改为实际的版本号。

3. 编辑 application.yml。

```
spring:
  shardingsphere:
    datasource:
      names: ds_0, ds_1
      ds_0:
        type: com.zaxxer.hikari.HikariDataSource
        driverClassName: com.mysql.cj.jdbc.Driver
        jdbcUrl: jdbc:mysql://localhost:3306/demo_ds_0?serverTimezone=UTC&
useSSL=false&useUnicode=true&characterEncoding=UTF-8
        username: root
        password:
      ds_1:
        type: com.zaxxer.hikari.HikariDataSource
        driverClassName: com.mysql.cj.jdbc.Driver
        jdbcUrl: jdbc:mysql://localhost:3306/demo_ds_1?serverTimezone=UTC&
useSSL=false&useUnicode=true&characterEncoding=UTF-8
        username: root
        password:
    rules:
      sharding:
        tables:
          ...
```

2.2 ShardingSphere-Proxy

2.2.1 应用场景



ShardingSphere-Proxy 的定位为透明化的数据库代理，理论上支持任何使用 MySQL、PostgreSQL、open-Gauss 协议的客户端操作数据，对异构语言、运维场景更友好。

2.2.2 使用限制

ShardingSphere-Proxy 对系统库/表（如 `information_schema`、`pg_catalog`）支持有限，通过部分图形化数据库客户端连接 Proxy 时，可能客户端或 Proxy 会有错误提示。可以使用命令行客户端（`mysql`、`psql`、`gsqsl` 等）连接 Proxy 验证功能。

2.2.3 前提条件

使用 Docker 启动 ShardingSphere-Proxy 无须额外依赖。使用二进制分发启动 Proxy，需要环境具备 Java JRE 8 或更高版本。

2.2.4 操作步骤

1. 获取 ShardingSphere-Proxy

目前 ShardingSphere-Proxy 可以通过以下方式：- 二进制发布包 - Docker - Helm

2. 规则配置

编辑 %SHARDINGSPHERE_PROXY_HOME%/conf/server.yaml。

编辑 %SHARDINGSPHERE_PROXY_HOME%/conf/config-xxx.yaml。

%SHARDINGSPHERE_PROXY_HOME% 为 Proxy 解压后的路径，例：/opt/shardingsphere-proxy-bin/

详情请参见 [配置手册](#)。

3. 引入依赖

如果后端连接 PostgreSQL 或 openGauss 数据库，不需要引入额外依赖。

如果后端连接 MySQL 数据库，请下载 [mysql-connector-java-5.1.47.jar](#) 或者 [mysql-connector-java-8.0.11.jar](#)，并将其放入 %SHARDINGSPHERE_PROXY_HOME%/ext-lib 目录。

4. 启动服务

- 使用默认配置项

```
sh %SHARDINGSPHERE_PROXY_HOME%/bin/start.sh
```

默认启动端口为 3307，默认配置文件目录为：%SHARDINGSPHERE_PROXY_HOME%/conf/。

- 自定义端口和配置文件目录

```
sh %SHARDINGSPHERE_PROXY_HOME%/bin/start.sh ${proxy_port} ${proxy_conf_directory}
```

5. 使用 ShardingSphere-Proxy

执行 MySQL / PostgreSQL / openGauss 的客户端命令直接操作 ShardingSphere-Proxy 即可。

使用 MySQL 客户端连接 ShardingSphere-Proxy：

```
mysql -h${proxy_host} -P${proxy_port} -u${proxy_username} -p${proxy_password}
```

使用 PostgreSQL 客户端连接 ShardingSphere-Proxy：

```
psql -h ${proxy_host} -p ${proxy_port} -U ${proxy_username}
```

使用 openGauss 客户端连接 ShardingSphere-Proxy：

```
gsql -r -h ${proxy_host} -p ${proxy_port} -U ${proxy_username} -W ${proxy_password}
```

Apache ShardingSphere 提供了多样化的功能，涵盖范围从数据库内核、数据库分布式到贴近数据库上层的应用，为用户提供了大量的功能池。

功能并无边界，只要满足数据库服务和生态的共性需求即可，期待更多的开源工程师参与 Apache ShardingSphere 社区，提供新颖思路和令人兴奋的功能。

3.1 数据分片

3.1.1 背景

传统的将数据集中存储至单一节点的解决方案，在性能、可用性和运维成本这三方面已经难于满足海量数据的场景。

从性能方面来说，由于关系型数据库大多采用 B+ 树类型的索引，在数据量超过阈值的情况下，索引深度的增加也将使得磁盘访问的 IO 次数增加，进而导致查询性能的下降；同时，高并发访问请求也使得集中式数据库成为系统的最大瓶颈。

从可用性的方面来讲，服务化的无状态性，能够达到较小成本的随意扩容，这必然导致系统的最终压力都落在数据库之上。而单一的数据节点，或者简单的主从架构，已经越来越难以承担。数据库的可用性，已成为整个系统的关键。

从运维成本方面考虑，当一个数据库实例中的数据达到阈值以上，对于 DBA 的运维压力就会增大。数据备份和恢复的时间成本都将随着数据量的大小而愈发不可控。一般来讲，单一数据库实例的数据的阈值在 1TB 之内，是比较合理的范围。

在传统的关系型数据库无法满足互联网场景需要的情况下，将数据存储至原生支持分布式的 NoSQL 的尝试越来越多。但 NoSQL 对 SQL 的不兼容性以及生态圈的不完善，使得它们在与关系型数据库的博弈中始终无法完成致命一击，而关系型数据库的地位却依然不可撼动。

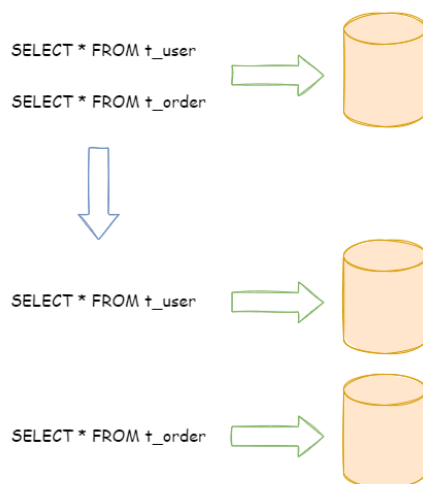
数据分片指按照某个维度将存放在单一数据库中的数据分散地存放至多个数据库或表中以达到提升性能瓶颈以及可用性的效果。数据分片的有效手段是对关系型数据库进行分库和分表。分库和分表均可以有效避免由数据量超过可承受阈值而产生的查询瓶颈。除此之外，分库还能够用于有效的分散对数据库单点的访问量；分表虽然无法缓解数据库压力，但却能够提供尽量将分布式事务转化为本地事务的可能，

一旦涉及到跨库的更新操作，分布式事务往往会使问题变得复杂。使用多主多从的分片方式，可以有效的避免数据单点，从而提升数据架构的可用性。

通过分库和分表进行数据的拆分来使得各个表的数据量保持在阈值以下，以及对流量进行疏导应对高访问量，是应对高并发和海量数据系统的有效手段。数据分片的拆分方式又分为垂直分片和水平分片。

垂直分片

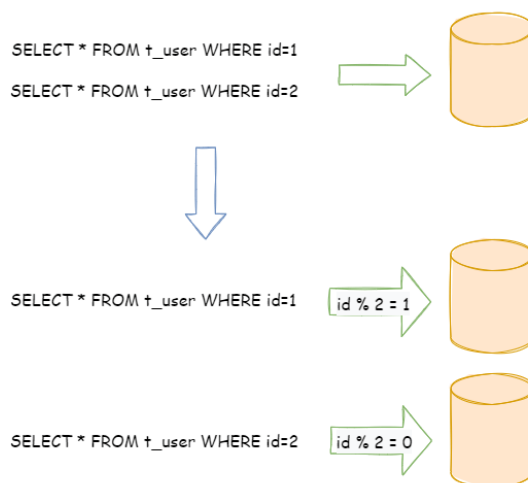
按照业务拆分的方式称为垂直分片，又称为纵向拆分，它的核心理念是专库专用。在拆分之前，一个数据库由多个数据表构成，每个表对应着不同的业务。而拆分之后，则是按照业务将表进行归类，分布到不同的数据库中，从而将压力分散至不同的数据库。下图展示了根据业务需要，将用户表和订单表垂直分片到不同的数据库的方案。



垂直分片往往需要对架构和设计进行调整。通常来讲，是来不及应对互联网业务需求快速变化的；而且，它也并无法真正的解决单点瓶颈。垂直拆分可以缓解数据量和访问量带来的问题，但无法根治。如果垂直拆分之后，表中的数据量依然超过单节点所能承载的阈值，则需要水平分片来进一步处理。

水平分片

水平分片又称为横向拆分。相对于垂直分片，它不再将数据根据业务逻辑分类，而是通过某个字段（或某几个字段），根据某种规则将数据分散至多个库或表中，每个分片仅包含数据的一部分。例如：根据主键分片，偶数主键的记录放入 0 库（或表），奇数主键的记录放入 1 库（或表），如下图所示。



水平分片从理论上突破了单机数据量处理的瓶颈，并且扩展相对自由，是数据分片的标准解决方案。

3.1.2 挑战

虽然数据分片解决了性能、可用性以及单点备份恢复等问题，但分布式的架构在获得了收益的同时，也引入了新的问题。

面对如此散乱的分片之后的数据，应用开发工程师和数据库管理员对数据库的操作变得异常繁重就是其中的重要挑战之一。他们需要知道数据需要从哪个具体的数据库的子表中获取。

另一个挑战则是，能够正确的运行在单节点数据库中的 SQL，在分片之后的数据库中并不一定能够正确运行。例如，分表导致表名称的修改，或者分页、排序、聚合分组等操作的不正确处理。

跨库事务也是分布式的数据库集群要面对的棘手事情。合理采用分表，可以在降低单表数据量的情况下，尽量使用本地事务，善于使用同库不同表可有效避免分布式事务带来的麻烦。在不能避免跨库事务的场景，有些业务仍然需要保持事务的一致性。而基于 XA 的分布式事务由于在并发度高的场景中性能无法满足需要，并未被互联网巨头大规模使用，他们大多采用最终一致性的柔性事务代替强一致事务。

3.1.3 目标

尽量透明化分库分表所带来的影响，让使用方尽量像使用一个数据库一样使用水平分片之后的数据库集群，是 Apache ShardingSphere 数据分片模块的主要设计目标。

3.1.4 应用场景

海量数据高并发的 OLTP 场景

由于关系型数据库大多采用 B+ 树类型的索引，在数据量超过阈值的情况下，索引深度的增加也将使得磁盘访问的 IO 次数增加，进而导致查询性能的下降。通过 ShardingSphere 数据分片，按照某个业务维度，将存放在单一数据库中的数据分散地存放至多个数据库或表中，可以达到提升性能的效果。通过使用 ShardingSphere-JDBC 接入端，可以满足高并发的 OLTP 场景下的性能要求。

海量数据实时分析 OLAP 场景

在传统的数据库架构中，如果用户想要进行数据分析，需要先使用 ETL 工具，将数据同步至数据平台中，然后再进行数据分析，使用 ETL 工具会导致数据分析的实效性大打折扣。ShardingSphere-Proxy 提供静态入口以及异构语言的支持，独立于应用程序部署，适用于实时分析的 OLAP 场景。

3.1.5 相关参考

- [数据分片的配置](#)
- [数据分片的开发者指南](#)

3.1.6 核心概念

表

表是透明化数据分片的关键概念。Apache ShardingSphere 通过提供多样化的表类型，适配不同场景下的数据分片需求。

逻辑表

相同结构的水平拆分数据库（表）的逻辑名称，是 SQL 中表的逻辑标识。例：订单数据根据主键尾数拆分为 10 张表，分别是 `t_order_0` 到 `t_order_9`，他们的逻辑表名为 `t_order`。

真实表

在水平拆分的数据库中真实存在的物理表。即上个示例中的 `t_order_0` 到 `t_order_9`。

绑定表

指分片规则一致的一组分片表。使用绑定表进行多表关联查询时,必须使用分片键进行关联,否则会出现笛卡尔积关联或跨库关联,从而影响查询效率。例如: `t_order` 表和 `t_order_item` 表,均按照 `order_id` 分片,并且使用 `order_id` 进行关联,则此两张表互为绑定表关系。绑定表之间的多表关联查询不会出现笛卡尔积关联,关联查询效率将大大提升。举例说明,如果 SQL 为:

```
SELECT i.* FROM t_order o JOIN t_order_item i ON o.order_id=i.order_id WHERE o.
order_id in (10, 11);
```

在不配置绑定表关系时,假设分片键 `order_id` 将数值 10 路由至第 0 片,将数值 11 路由至第 1 片,那么路由后的 SQL 应该为 4 条,它们呈现为笛卡尔积:

```
SELECT i.* FROM t_order_0 o JOIN t_order_item_0 i ON o.order_id=i.order_id WHERE o.
order_id in (10, 11);

SELECT i.* FROM t_order_0 o JOIN t_order_item_1 i ON o.order_id=i.order_id WHERE o.
order_id in (10, 11);

SELECT i.* FROM t_order_1 o JOIN t_order_item_0 i ON o.order_id=i.order_id WHERE o.
order_id in (10, 11);

SELECT i.* FROM t_order_1 o JOIN t_order_item_1 i ON o.order_id=i.order_id WHERE o.
order_id in (10, 11);
```

在配置绑定表关系,并且使用 `order_id` 进行关联后,路由的 SQL 应该为 2 条:

```
SELECT i.* FROM t_order_0 o JOIN t_order_item_0 i ON o.order_id=i.order_id WHERE o.
order_id in (10, 11);

SELECT i.* FROM t_order_1 o JOIN t_order_item_1 i ON o.order_id=i.order_id WHERE o.
order_id in (10, 11);
```

其中 `t_order` 表由于指定了分片条件,ShardingSphere 将会以它作为整个绑定表的主表。所有路由计算将会只使用主表的策略,那么 `t_order_item` 表的分片计算将会使用 `t_order` 的条件。

广播表

指所有的分片数据源中都存在的表，表结构及其数据在每个数据库中均完全一致。适用于数据量不大且需要与海量数据的表进行关联查询的场景，例如：字典表。

单表

指所有的分片数据源中仅唯一存在的表。适用于数据量不大且无需分片的表。

数据节点

数据分片的最小单元，由数据源名称和真实表组成。例：ds_0.t_order_0。逻辑表与真实表的映射关系，可分为均匀分布和自定义分布两种形式。

均匀分布

指数据表在每个数据源内呈现均匀分布的态势，例如：

```
db0
├─ t_order0
└─ t_order1
db1
├─ t_order0
└─ t_order1
```

数据节点的配置如下：

```
db0.t_order0, db0.t_order1, db1.t_order0, db1.t_order1
```

自定义分布

指数据表呈现有特定规则的分布，例如：

```
db0
├─ t_order0
└─ t_order1
db1
├─ t_order2
├─ t_order3
└─ t_order4
```

数据节点的配置如下：

```
db0.t_order0, db0.t_order1, db1.t_order2, db1.t_order3, db1.t_order4
```

分片

分片键

用于将数据库（表）水平拆分的数据库字段。例：将订单表中的订单主键的尾数取模分片，则订单主键为分片字段。SQL 中如果无分片字段，将执行全路由，性能较差。除了对单分片字段的支持，Apache ShardingSphere 也支持根据多个字段进行分片。

分片算法

用于将数据分片的算法，支持 =、>=、<=、>、<、BETWEEN 和 IN 进行分片。分片算法可由开发者自行实现，也可使用 Apache ShardingSphere 内置的分片算法语法糖，灵活度非常高。

自动化分片算法

分片算法语法糖，用于便捷的托管所有数据节点，使用者无需关注真实表的物理分布。包括取模、哈希、范围、时间等常用分片算法的实现。

自定义分片算法

提供接口让应用开发者自行实现与业务实现紧密相关的分片算法，并允许使用者自行管理真实表的物理分布。自定义分片算法又分为：

- 标准分片算法

用于处理使用单一键作为分片键的 =、IN、BETWEEN AND、>、<、>=、<= 进行分片的场景。

- 复合分片算法

用于处理使用多键作为分片键进行分片的场景，包含多个分片键的逻辑较复杂，需要应用开发者自行处理其中的复杂度。

- Hint 分片算法

用于处理使用 Hint 行分片的场景。

分片策略

包含分片键和分片算法，由于分片算法的独立性，将其独立抽离。真正可用于分片操作的是分片键 + 分片算法，也就是分片策略。

强制分片路由

对于分片字段并非由 SQL 而是其他外置条件决定的场景，可使用 SQL Hint 注入分片值。例：按照员工登录主键分库，而数据库中并无此字段。SQL Hint 支持通过 Java API 和 SQL 注释两种方式使用。详情请参见强制分片路由。

行表达式

行表达式是为了解决配置的简化与一体化这两个主要问题。在繁琐的数据分片规则配置中，随着数据节点的增多，大量的重复配置使得配置本身不易被维护。通过行表达式可以有效地简化数据节点配置工作量。

对于常见的分片算法，使用 Java 代码实现并不有助于配置的统一管理。通过行表达式书写分片算法，可以有效地将规则配置一同存放，更加易于浏览与存储。

行表达式的使用非常直观，只需要在配置中使用 `${ expression }` 或 `$->{ expression }` 标识行表达式即可。目前支持数据节点和分片算法这两个部分的配置。行表达式的内容使用的是 Groovy 的语法，Groovy 能够支持的所有操作，行表达式均能够支持。例如：

`${begin..end}` 表示范围区间 `${[unit1, unit2, unit_x]}` 表示枚举值

行表达式中如果出现连续多个 `${ expression }` 或 `$->{ expression }` 表达式，整个表达式最终的结果将会根据每个子表达式的结果进行笛卡尔组合。

例如，以下行表达式：

```
${['online', 'offline']}_table${1..3}
```

最终会解析为：

```
online_table1, online_table2, online_table3, offline_table1, offline_table2,
offline_table3
```

分布式主键

传统数据库软件开发中，主键自动生成技术是基本需求。而各个数据库对于该需求也提供了相应的支持，比如 MySQL 的自增键，Oracle 的自增序列等。数据分片后，不同数据节点生成全局唯一主键是非常棘手的问题。同一个逻辑表内的不同实际表之间的自增键由于无法互相感知而产生重复主键。虽然可通过约束自增主键初始值和步长的方式避免碰撞，但需引入额外的运维规则，使解决方案缺乏完整性和可扩展性。

目前有许多第三方解决方案可以完美解决这个问题，如 UUID 等依靠特定算法自生成不重复键，或者通过引入主键生成服务等。为了方便用户使用、满足不同用户不同使用场景的需求，Apache ShardingSphere 不仅提供了内置的分布式主键生成器，例如 UUID、SNOWFLAKE，还抽离出分布式主键生成器的接口，方便用户自行实现自定义的自增主键生成器。

3.1.7 使用限制

兼容全部常用的路由至单数据节点的 SQL；路由至多数据节点的 SQL 由于场景复杂，分为稳定支持、实验性支持和不支持这三种情况。

稳定支持

全面支持 DML、DDL、DCL、TCL 和常用 DAL。支持分页、去重、排序、分组、聚合、表关联等复杂查询。支持 PostgreSQL 和 openGauss 数据库 SCHEMA DDL 和 DML 语句。

常规查询

- SELECT 主语句

```
SELECT select_expr [, select_expr ...] FROM table_reference [, table_reference ...]
[WHERE predicates]
[GROUP BY {col_name | position} [ASC | DESC], ...]
[ORDER BY {col_name | position} [ASC | DESC], ...]
[LIMIT {[offset,] row_count | row_count OFFSET offset}]
```

- select_expr

```
* |
[DISTINCT] COLUMN_NAME [AS] [alias] |
(MAX | MIN | SUM | AVG)(COLUMN_NAME | alias) [AS] [alias] |
COUNT(* | COLUMN_NAME | alias) [AS] [alias]
```

- table_reference

```
tbl_name [AS] alias [index_hint_list]
| table_reference ([INNER] | {LEFT|RIGHT} [OUTER]) JOIN table_factor [JOIN ON
conditional_expr | USING (column_list)]
```

子查询

子查询和外层查询同时指定分片键，且分片键的值保持一致时，由内核提供稳定支持。

例如：

```
SELECT * FROM (SELECT * FROM t_order WHERE order_id = 1) o WHERE o.order_id = 1;
```

用于分页的子查询，由内核提供稳定支持。

例如：

```
SELECT * FROM (SELECT row_.*, rownum rownum_ FROM (SELECT * FROM t_order) row_
WHERE rownum <= ?) WHERE rownum > ?;
```

分页查询

完全支持 MySQL、PostgreSQL、openGauss、Oracle 和 SQLServer 由于分页查询较为复杂，仅部分支持。Oracle 和 SQLServer 的分页都需要通过子查询来处理，ShardingSphere 支持分页相关的子查询。

- Oracle

支持使用 rownum 进行分页：

```
SELECT * FROM (SELECT row_.*, rownum rownum_ FROM (SELECT o.order_id as order_id
FROM t_order o JOIN t_order_item i ON o.order_id = i.order_id) row_ WHERE rownum <=
?) WHERE rownum > ?
```

- SQLServer

支持使用 TOP + ROW_NUMBER() OVER 配合进行分页：

```
SELECT * FROM (SELECT TOP (?) ROW_NUMBER() OVER (ORDER BY o.order_id DESC) AS
rownum, * FROM t_order o) AS temp WHERE temp.rownum > ? ORDER BY temp.order_id
```

支持 SQLServer 2012 之后的 OFFSET FETCH 的分页方式：

```
SELECT * FROM t_order o ORDER BY id OFFSET ? ROW FETCH NEXT ? ROWS ONLY
```

- MySQL, PostgreSQL 和 openGauss

MySQL、PostgreSQL 和 openGauss 都支持 LIMIT 分页，无需子查询：

```
SELECT * FROM t_order o ORDER BY id LIMIT ? OFFSET ?
```

运算表达式中包含分片键

当分片键处于运算表达式中时，无法通过 SQL 字面提取用于分片的值，将导致全路由。例如，假设 create_time 为分片键：

```
SELECT * FROM t_order WHERE to_date(create_time, 'yyyy-mm-dd') = '2019-01-01';
```

实验性支持

实验性支持特指使用 Federation 执行引擎提供支持。该引擎处于快速开发中，用户虽基本可用，但仍需大量优化，是实验性产品。

子查询

子查询和外层查询未同时指定分片键，或分片键的值不一致时，由 Federation 执行引擎提供支持。

例如：

```
SELECT * FROM (SELECT * FROM t_order) o;

SELECT * FROM (SELECT * FROM t_order) o WHERE o.order_id = 1;

SELECT * FROM (SELECT * FROM t_order WHERE order_id = 1) o;

SELECT * FROM (SELECT * FROM t_order WHERE order_id = 1) o WHERE o.order_id = 2;
```

跨库关联查询

当关联查询中的多个表分布在不同的数据库实例上时，由 Federation 执行引擎提供支持。假设 `t_order` 和 `t_order_item` 是多数据节点的分片表，并且未配置绑定表规则，`t_user` 和 `t_user_role` 是分布在不同的数据库实例上的单表，那么 Federation 执行引擎能够支持如下常用的关联查询：

```
SELECT * FROM t_order o INNER JOIN t_order_item i ON o.order_id = i.order_id WHERE o.order_id = 1;

SELECT * FROM t_order o INNER JOIN t_user u ON o.user_id = u.user_id WHERE o.user_id = 1;

SELECT * FROM t_order o LEFT JOIN t_user_role r ON o.user_id = r.user_id WHERE o.user_id = 1;

SELECT * FROM t_order_item i LEFT JOIN t_user u ON i.user_id = u.user_id WHERE i.user_id = 1;

SELECT * FROM t_order_item i RIGHT JOIN t_user_role r ON i.user_id = r.user_id WHERE i.user_id = 1;

SELECT * FROM t_user u RIGHT JOIN t_user_role r ON u.user_id = r.user_id WHERE u.user_id = 1;
```

不支持

CASE WHEN

以下 CASE WHEN 语句不支持：

- CASE WHEN 中包含子查询
- CASE WHEN 中使用逻辑表名（请使用表别名）

分页查询

Oracle 和 SQLServer 由于分页查询较为复杂，目前有部分分页查询不支持，具体如下：

- Oracle

目前不支持 `rownum + BETWEEN` 的分页方式。

- SQLServer

目前不支持使用 `WITH xxx AS (SELECT ...)` 的方式进行分页。由于 Hibernate 自动生成的 SQLServer 分页语句使用了 `WITH` 语句，因此目前并不支持基于 Hibernate 的 SQLServer 分页。目前也不支持使用两个 `TOP + 子查询` 的方式实现分页。

3.1.8 附录

不支持的 SQL：

- `CASE WHEN` 中包含子查询
- `CASE WHEN` 中使用逻辑表名（请使用表别名）
- `INSERT INTO tbl_name (col1, col2, ...) SELECT * FROM tbl_name WHERE col3 = ?`（`SELECT` 子句不支持 `*` 和内置分布式主键生成器）
- `REPLACE INTO tbl_name (col1, col2, ...) SELECT * FROM tbl_name WHERE col3 = ?`（`SELECT` 子句不支持 `*` 和内置分布式主键生成器）
- `SELECT MAX(tbl_name.col1) FROM tbl_name`（查询列是函数表达式时，查询列前不能使用表名，可以使用表别名）

3.2 分布式事务

3.2.1 背景

数据库事务需要满足 ACID（原子性、一致性、隔离性、持久性）四个特性。

- 原子性（Atomicity）指事务作为整体来执行，要么全部执行，要么全不执行；
- 一致性（Consistency）指事务应确保数据从一个一致的状态转变为另一个一致的状态；
- 隔离性（Isolation）指多个事务并发执行时，一个事务的执行不应影响其他事务的执行；
- 持久性（Durability）指已提交的事务修改数据会被持久保存。

在单一数据节点中，事务仅限于对单一数据库资源的访问控制，称之为本地事务。几乎所有的成熟的关系型数据库都提供了对本地事务的原生支持。但是在基于微服务的分布式应用环境下，越来越多的应用场景要求对多个服务的访问及其相对应的多个数据库资源能纳入到同一个事务当中，分布式事务应运而生。

关系型数据库虽然对本地事务提供了完美的 ACID 原生支持。但在分布式的场景下，它却成为系统性能的桎梏。如何让数据库在分布式场景下满足 ACID 的特性或找寻相应的替代方案，是分布式事务的重点工作。

3.2.2 挑战

由于应用的场景不同，需要开发者能够合理的在性能与功能之间权衡各种分布式事务。

强一致的事务与柔性事务的 API 和功能并不完全相同，在它们之间并不能做到自由的透明切换。在开发决策阶段，就不得不在强一致的事务和柔性事务之间抉择，使得设计和开发成本被大幅增加。

基于 XA 的强一致事务使用相对简单，但是无法很好的应对互联网的高并发或复杂系统的长事务场景；柔性事务则需要开发者对应用进行改造，接入成本非常高，并且需要开发者自行实现资源锁定和反向补偿。

3.2.3 目标

整合现有的成熟事务方案，为本地事务、两阶段事务和柔性事务提供统一的分布式事务接口，并弥补当前方案的不足，提供一站式的分布式事务解决方案是 Apache ShardingSphere 分布式事务模块的主要设计目标。

3.2.4 原理介绍

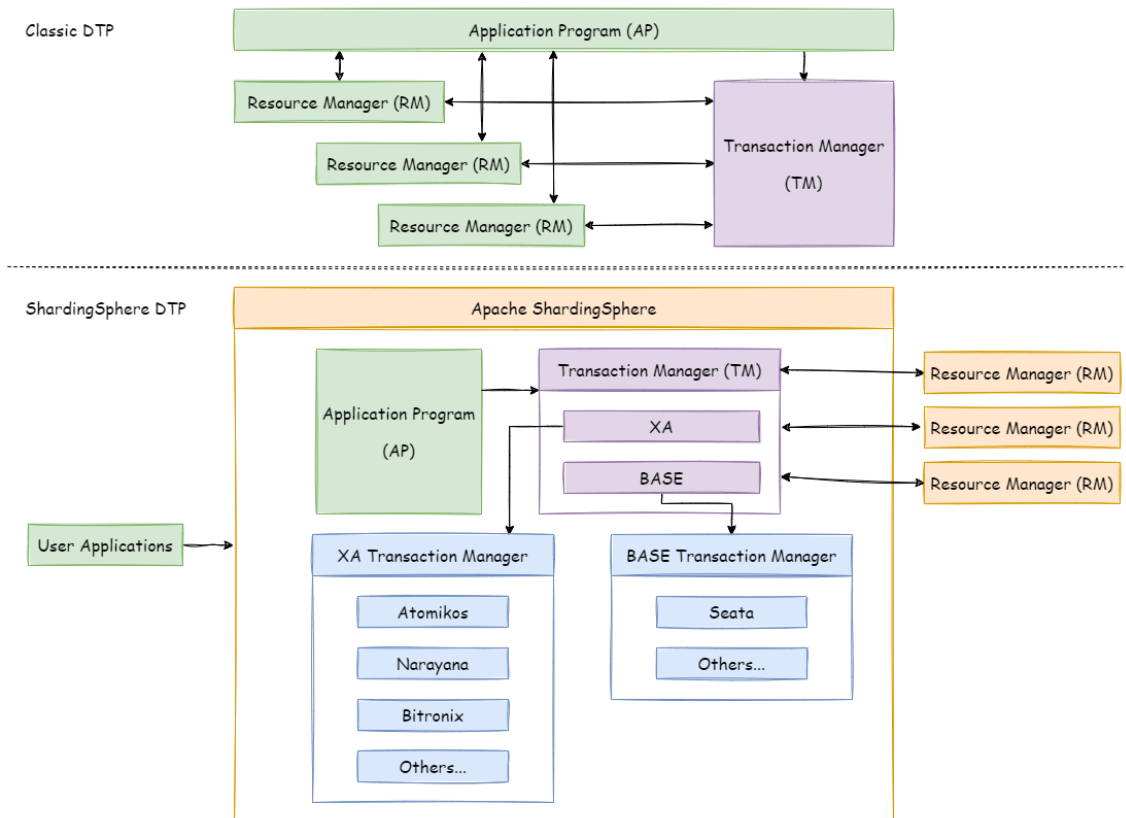
ShardingSphere 对外提供 begin/commit/rollback 传统事务接口，通过 LOCAL，XA，BASE 三种模式提供了分布式事务的能力，

LOCAL 事务

LOCAL 模式基于 ShardingSphere 代理的数据库 begin/commit/rollback 的接口实现，对于一条逻辑 SQL，ShardingSphere 通过 begin 指令在每个被代理的数据库开启事务，并执行实际 SQL，并执行 commit/rollback。由于每个数据节点各自管理自己的事务，它们之间没有协调以及通信的能力，也并不互相知晓其他数据节点事务的成功与否。在性能方面无任何损耗，但在强一致性以及最终一致性方面不能够保证。

XA 事务

XA 事务采用的是 X/OPEN 组织所定义的 [DTP 模型](#) 所抽象的 AP（应用程序），TM（事务管理器）和 RM（资源管理器）概念来保证分布式事务的强一致性。其中 TM 与 RM 间采用 XA 的协议进行双向通信，通过两阶段提交实现。与传统的本地事务相比，XA 事务增加了准备阶段，数据库除了被动接受提交指令外，还可以反向通知调用方事务是否可以被提交。TM 可以收集所有分支事务的准备结果，并于最后进行原子提交，以保证事务的强一致性。



XA 事务建立在 ShardingSphere 代理的数据库 xa start/end/prepare/commit/rollback/recover 的接口上。

对于一条逻辑 SQL，ShardingSphere 通过 xa begin 指令在每个被代理的数据库开启事务，内部集成 TM，用于协调各分支事务，并执行 xa commit/rollback。

基于 XA 协议实现的分布式事务，由于在执行的过程中需要对所需资源进行锁定，它更加适用于执行时间确定的短事务。对于长事务来说，整个事务进行期间对数据的独占，将会对并发场景下的性能产生一定的影响。

BASE 事务

如果将实现了 ACID 的事务要素的事务称为刚性事务的话，那么基于 BASE 事务要素的事务则称为柔性事务。BASE 是基本可用、柔性状态和最终一致性这三个要素的缩写。

- 基本可用（Basically Available）保证分布式事务参与方不一定同时在线；
- 柔性状态（Soft state）则允许系统状态更新有一定的延时，这个延时对客户来说不一定能够察觉；
- 最终一致性（Eventually consistent）通常是通过消息传递的方式保证系统的最终一致性。

在 ACID 事务中对隔离性的要求很高，在事务执行过程中，必须将所有的资源锁定。柔性事务的理念则是通过业务逻辑将互斥锁操作从资源层面上移至业务层面。通过放宽对强一致性要求，来换取系统吞吐量的提升。

基于 ACID 的强一致性事务和基于 BASE 的最终一致性事务都不是银弹，只有在最适合的场景中才能发挥它们的最大长处。Apache ShardingSphere 集成了 SEATA 作为柔性事务的使用方案。可通过下表详细对比它们之间的区别，以帮助开发者进行技术选型。

	LOCAL	XA	BASE
业务改造	无	无	需要 seata server
一致性	不支持	支持	最终一致
隔离性	不支持	支持	业务方保证
并发性能	无影响	严重衰退	略微衰退
适合场景	业务方处理不一致	短事务 & 低并发	长事务 & 高并发

3.2.5 应用场景

在单机应用场景中，依赖数据库提供的事务即可满足业务上对事务 ACID 的需求。但是在分布式场景下，传统数据库解决方案缺乏对全局事务的管控能力，用户在使用过程中可能遇到多个数据库节点上出现数据不一致的问题。

ShardingSphere 分布式事务，为用户屏蔽了分布式事务处理的复杂性，提供了灵活多样的分布式事务解决方案，用户可以根据自己的业务场景在 LOCAL，XA，BASE 三种模式中，选择适合自己的分布式事务解决方案。

ShardingSphere XA 事务使用场景

对于 XA 事务，提供了分布式环境下，对数据强一致性的保证。但是由于存在同步阻塞问题，对性能会有一定影响。适用于对数据一致性要求非常高且对并发性能要求不是很高的业务场景。

ShardingSphere BASE 事务使用场景

对于 BASE 事务，提供了分布式环境下，对数据最终一致性的保证。由于在整个事务过程中，不会像 XA 事务那样全程锁定资源，所以性能较好。适用于对并发性能要求很高并且允许出现短暂数据不一致的业务场景。

ShardingSphere LOCAL 事务使用场景

对于 LOCAL 事务，在分布式环境下，不保证各个数据库节点之间数据的一致性和隔离性，需要业务方自行处理可能出现的不一致问题。适用于用户希望自行处理分布式环境下数据一致性问题的业务场景。

3.2.6 相关参考

- 分布式事务的 [YAML 配置](#)

3.2.7 核心概念

XA 协议

XA 协议最早的分布式事务模型是由 X/Open 国际联盟提出的 X/Open Distributed Transaction Processing (DTP) 模型，简称 XA 协议。

3.2.8 使用限制

虽然 Apache ShardingSphere 希望能够完全兼容所有的分布式事务场景，并在性能上达到最优，但在 CAP 定理所指导下，分布式事务必然有所取舍。Apache ShardingSphere 希望能够将分布式事务的选择权交给使用者，在不同的场景使用最适合的分布式事务解决方案。

LOCAL 事务

支持项

- 完全支持非跨库事务，例如：仅分表，或分库但是路由的结果在单库中；
- 完全支持因逻辑异常导致的跨库事务。例如：同一事务中，跨两个库更新。更新完毕后，抛出空指针，则两个库的内容都能够回滚。

不支持项

- 不支持因网络、硬件异常导致的跨库事务。例如：同一事务中，跨两个库更新，更新完毕后、未提交之前，第一个库宕机，则只有第二个库数据提交，且无法回滚。

XA 事务

支持项

- 支持 Savepoint 嵌套事务；
- PostgreSQL/OpenGauss 事务块内，SQL 执行出现异常，执行 Commit，事务自动回滚；
- 支持数据分片后的跨库事务；
- 两阶段提交保证操作的原子性和数据的强一致性；
- 服务宕机重启后，提交/回滚中的事务可自动恢复；
- 支持同时使用 XA 和非 XA 的连接池。

不支持项

- 服务宕机后，在其它机器上恢复提交/回滚中的数据；
- MySQL 事务块内，SQL 执行出现异常，执行 Commit，数据保持一致。

BASE 事务

支持项

- 支持数据分片后的跨库事务；
- 支持 RC 隔离级别；
- 通过 undo 快照进行事务回滚；
- 支持服务宕机后的，自动恢复提交中的事务。

不支持项

- 不支持除 RC 之外的隔离级别。

待优化项

- Apache ShardingSphere 和 SEATA 重复 SQL 解析。

3.2.9 附录

不支持的 SQL：

- 事务中使用 DistSQL 里的 RAL、RDL 操作；
- XA 事务中使用 DDL 语句。

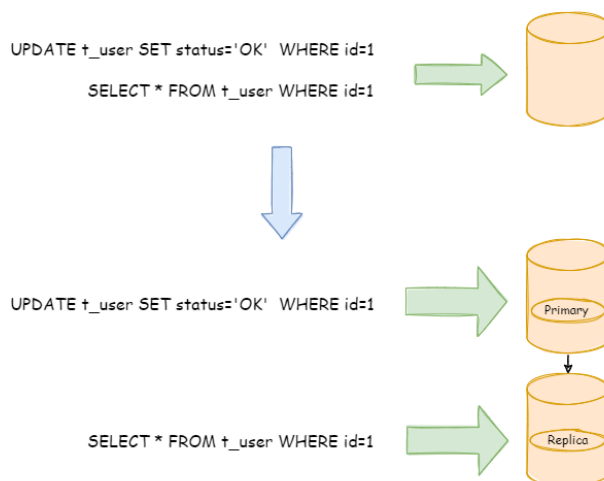
3.3 读写分离

3.3.1 背景

面对日益增加的系统访问量，数据库的吞吐量面临着巨大瓶颈。对于同一时刻有大量并发读操作和较少写操作类型的应用系统来说，将数据库拆分为主库和从库，主库负责处理事务性的增删改操作，从库负责处理查询操作，能够有效的避免由数据更新导致的行锁，使得整个系统的查询性能得到极大的改善。

通过一主多从的配置方式，可以将查询请求均匀的分散到多个数据副本，能够进一步的提升系统的处理能力。使用多主多从的方式，不但能够提升系统的吞吐量，还能够提升系统的可用性，可以达到在任何一数据库宕机，甚至磁盘物理损坏的情况下仍然不影响系统的正常运行。

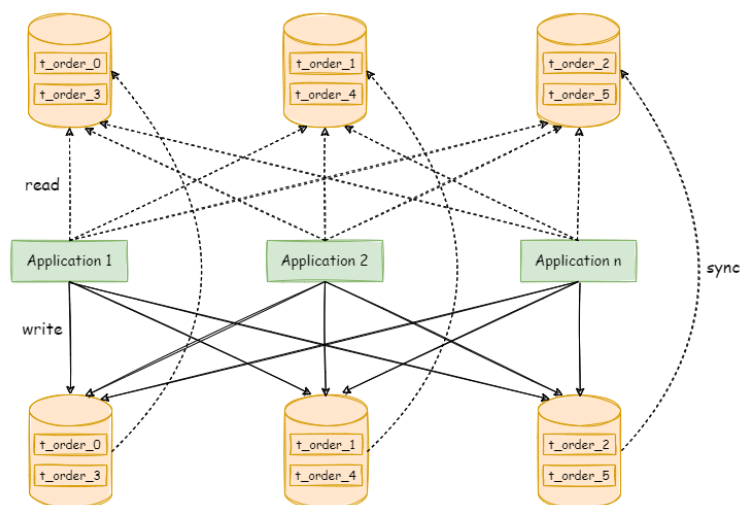
与将数据根据分片键打散至各个数据节点的水平分片不同，读写分离则是根据 SQL 语义的分析，将读操作和写操作分别路由至主库与从库。



读写分离的数据节点中的数据内容是一致的，而水平分片的每个数据节点的数据内容却并不相同。将水平分片和读写分离联合使用，能够更加有效的提升系统性能。

3.3.2 挑战

读写分离虽然可以提升系统的吞吐量和可用性，但同时也带来了数据不一致的问题。这包括多个主库之间的数据一致性，以及主库与从库之间的数据一致性的问题。并且，读写分离也带来了与数据分片同样的问题，它同样会使得应用开发和运维人员对数据库的操作和运维变得更加复杂。下图展现了将数据分片与读写分离一同使用时，应用程序与数据库集群之间的复杂拓扑关系。



3.3.3 目标

透明化读写分离所带来的影响，让使用方尽量像使用一个数据库一样使用主从数据库集群，是 Apache ShardingSphere 读写分离模块的主要设计目标。

3.3.4 应用场景

复杂的主从数据库架构

许多系统通过采用主从数据库架构的配置来提高整个系统的吞吐量，但是主从的配置也给业务的使用带来了一定的复杂性。接入 ShardingSphere，可以利用读写分离功能管理主从数据库，实现透明化的读写分离功能，让用户像使用一个数据库一样使用主从架构的数据库。

3.3.5 相关参考

Java API

YAML 配置

Spring Boot Starter

Spring 命名空间

3.3.6 核心概念

主库

添加、更新以及删除数据操作所使用的数据库，目前仅支持单主库。

从库

查询数据操作所使用的数据库，可支持多从库。

主从同步

将主库的数据异步的同步到从库的操作。由于主从同步的异步性，从库与主库的数据会短时间内不一致。

负载均衡策略

通过负载均衡策略将查询请求疏导至不同从库。

3.3.7 使用限制

- 不处理主库和从库的数据同步
- 不处理主库和从库的数据同步延迟导致的数据不一致
- 不支持主库多写
- 不处理主从库间的事务一致性。主从模型中，事务中的数据读写均用主库。

3.4 高可用

3.4.1 背景

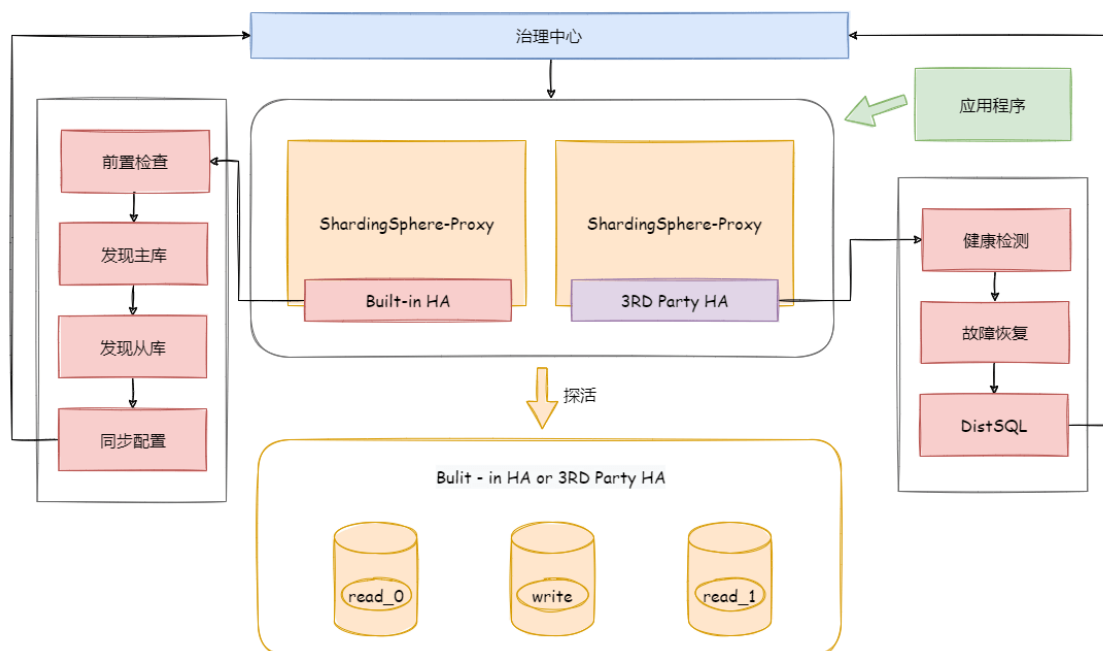
高可用是现代系统的最基本诉求，作为系统基石的数据库，对于高可用的要求也是必不可少的。

在存算分离的分布式数据库体系中，存储节点和计算节点的高可用方案是不同的。对于有状态的存储节点来说，需要其自身具备数据一致性同步、探活、主节点选举等能力；对于无状态的计算节点来说，需要感知存储节点的变化，还需要独立架设负载均衡器，并具备服务发现和请求分发的能力。

Apache ShardingSphere 自身提供计算节点，并通过数据库作为存储节点。因此，它采用的高可用方案是利用数据库自身的高可用方案做存储节点高可用，并自动识别其变化。

3.4.2 挑战

Apache ShardingSphere 需要自动感知多样化的存储节点高可用方案的同时，也能够动态集成对读写分离方案，是实现的主要挑战。



3.4.3 目标

尽可能的保证 7*24 小时不间断的数据库服务，是 Apache ShardingSphere 高可用模块的主要设计目标。

3.4.4 应用场景

在大多数情况下，高可用搭配读写分离功能一起使用。当用户写库或读库关系发生变化时，ShardingSphere 可动态的感知并纠正内部的主从关系，进而保证读流量和写流量的正确路由。同时当从库宕机时，ShardingSphere 也可动态纠正存储节点的状态，保证读流量分发正确。

3.4.5 相关参考

Java API

YAML 配置

Spring Boot Starter

Spring 命名空间

3.4.6 核心概念

高可用类型

Apache ShardingSphere 不提供数据库高可用的能力，它通过第三方提供的高可用方案感知数据库主从关系的切换。确切来说，Apache ShardingSphere 提供数据库发现的能力，自动感知数据库主从关系，并修正计算节点对数据库的连接。

动态读写分离

高可用和读写分离一起使用时，读写分离无需配置具体的主库和从库。高可用的数据源会动态的修正读写分离的主从关系，并正确地疏导读写流量。

3.4.7 使用限制

支持项

- MySQL MGR 单主模式。
- MySQL 主从复制模式。
- openGauss 主从复制模式。

不支持项

- MySQL MGR 多主模式。

3.5 数据库网关

3.5.1 背景

随着数据库碎片化趋势的不可逆转，多种类型数据库的共存已渐成常态。使用一种 SQL 方言访问异构数据库的场景在不断增加。

3.5.2 挑战

多样化的数据库的存在，使访问数据库的 SQL 方言难于标准化，工程师需要针对不同种类的数据库使用不同的方言，缺乏统一化的查询平台。

将不同类型的数据库方言自动翻译为后端数据库所使用的方言，让工程师可以使用任意一种数据库方言访问所有的后端异构数据库，可以极大的降低开发和维护成本。

3.5.3 目标

SQL 方言的自动翻译，是 Apache ShardingSphere 数据库网关希望达成的主要目标。

3.5.4 应用场景

随着业务场景的多元化，企业内部的数据库产品也呈现多元化的趋势，业务应用与不同数据库产品的对接也变得异常复杂，ShardingSphere 数据库网关可以屏蔽业务应用与底层多元化数据库之间连接，同时为不同的业务场景提供统一的访问协议和语法体系，能够帮助企业快速打造统一的数据访问平台。

3.5.5 核心概念

SQL 方言

SQL 方言也就是数据库方言，指的是某些数据库产品除了支持 SQL 之外，还会有一些自己独有的语法，这就称之为方言，不同的数据库产品，也可能会有不同的 SQL 方言。

3.5.6 使用限制

Apache ShardingSphere 的 SQL 方言翻译处于实验阶段。

目前仅支持 MySQL/PostgreSQL 的方言自动翻译，工程师可以使用 MySQL 的方言和协议，访问 PostgreSQL 数据库，反之亦然。

3.6 流量治理

3.6.1 背景

随着数据规模的不断膨胀，使用多节点集群的分布式方式逐渐成为趋势。对集群整体视角的统一管理能力，和针对单独组件细粒度的控制能力，是基于存算分离的现代数据库体系中不可或缺的功能。

3.6.2 挑战

管控的挑战，在于对集群的集中化管理的统一管理能力以及在单点出现故障时精细化的操作能力。

集中化管理的挑战体现在将包括数据库存储节点和中间件计算节点的状态统一管理，并且能够实时的探测到分布式环境下最新的变动情况，进一步为集群的控制和调度提供依据。

面对超负荷的流量下，针对某一节点进行熔断和限流，以保证整个数据库集群得以继续运行，是分布式系统下对单一节点控制能力的挑战。

3.6.3 目标

实现从数据库到计算节点打通的一体化管理能力，在故障中为组件提供细粒度的控制能力，并尽可能的提供自愈的可能，是 Apache ShardingSphere 管控模块的主要设计目标。

3.6.4 应用场景

计算节点过载保护

当 ShardingSphere 集群内某个计算节点超过负载后，通过熔断功能，阻断应用到该计算节点的流量，保证整个集群继续提供稳定服务。

存储节点限流

在读写分离的场景下，当 ShardingSphere 集群内某个负责读流量的存储节点承接超负荷的请求时，通过限流功能，阻断集群内计算节点到该存储节点的流量，以保证存储节点集群正常响应。

3.6.5 核心概念

熔断

阻断 Apache ShardingSphere 和数据库的连接。当某个 Apache ShardingSphere 节点超过负载后，停止该节点对数据库的访问，使数据库能够保证足够的资源为其他节点提供服务。

限流

面对超负荷的请求开启限流，以保护部分请求可以得以高质量的响应。

3.7 数据迁移

3.7.1 背景

当业务持续发展，数据量和并发量达到一定程度，传统单体数据库可能面临性能、可扩展性、可用性问题；

业界曾提出 NoSQL 解决方案，通过数据分片和水平扩容解决以上问题，但是 NoSQL 数据库通常不支持事务和 SQL；

ShardingSphere 也支持数据分片和水平扩容，可以解决以上问题，同时还支持分布式事务和 SQL；

ShardingSphere 提供的数据库迁移方案可以助力传统单体数据库平滑切换到 ShardingSphere。

3.7.2 挑战

在迁移过程中，不应该对正在运行的业务造成影响。尽可能减少迁移时数据不可用的时间窗口，是数据库迁移的第一个挑战；

其次，数据库迁移不应该对现有的数据造成影响，如何保证数据的正确性，是数据库迁移的第二个挑战。

3.7.3 目标

减少数据库迁移时的业务影响，提供一站式的通用数据库迁移解决方案，是 Apache ShardingSphere 数据库迁移的主要设计目标。

3.7.4 应用场景

假如一个应用系统在使用传统单体数据库，单表数据量达到了 1 亿并且还在快速增长，单体数据库负载持续在高位，成为系统瓶颈。一旦数据库成为瓶颈，对应用服务器扩容是无效的，需要对数据库进行扩容。

3.7.5 相关参考

- 数据库迁移的配置
- 数据库迁移的实现原理

3.7.6 核心概念

节点

运行计算层或存储层组件进程的实例，可以是物理机、虚拟机、容器等。

集群

为了提供特定服务而集合在一起的多个节点。

源端

原始数据所在的存储集群。

目标端

原始数据将要迁移的目标存储集群。

数据迁移作业

把数据从某一个存储集群复制到另一个存储集群的完整流程。

存量数据

在数据迁移作业开始前，数据节点中已有的数据。

增量数据

在数据迁移作业执行过程中，业务系统所产生的新数据。

3.7.7 使用限制

支持项

- 将外围数据迁移至 Apache ShardingSphere 所管理的数据库；
- 整型或字符串主键表迁移。

不支持项

- 无主键表迁移;
- 复合主键表迁移;
- 不支持在当前存储节点之上做迁移, 需要准备一个全新的数据库集群作为迁移目标库。

3.8 数据加密

3.8.1 背景

安全控制一直是治理的重要环节, 数据加密属于安全控制的范畴。无论对互联网公司还是传统行业来说, 数据安全一直是极为重视和敏感的话题。数据加密是指对某些敏感信息通过加密规则进行数据的变形, 实现敏感隐私数据的可靠保护。涉及客户安全数据或者一些商业性敏感数据, 如身份证号、手机号、卡号、客户号等个人信息按照相关部门规定, 都需要进行数据加密。

对于数据加密的需求, 在现实的业务场景中一般分为两种情况:

1. 新业务上线, 安全部门规定需将涉及用户敏感信息, 例如银行、手机号码等进行加密后存储到数据库, 在使用的时候再进行解密处理。因为是全新系统, 因而没有存量数据清洗问题, 所以实现相对简单。
2. 已上线业务, 之前一直将明文存储在数据库中。相关部门突然需要对已上线业务进行加密整改。这种场景一般需要处理 3 个问题:
 - 历史数据需要如何进行加密处理, 即洗数。
 - 如何能在不改动业务 SQL 和逻辑情况下, 将新增数据进行加密处理, 并存储到数据库; 在使用时, 再进行解密取出。
 - 如何较为安全、无缝、透明化地实现业务系统在明文与密文数据间的迁移。

3.8.2 挑战

在真实业务场景中, 相关业务开发团队则往往需要针对公司安全部门需求, 自行实行并维护一套加解密系统。而当加密场景发生改变时, 自行维护的加密系统往往又面临着重构或修改风险。此外, 对于已经上线的业务, 在不修改业务逻辑和 SQL 的情况下, 透明化、安全低风险地实现无缝进行加密改造也相对复杂。

3.8.3 目标

根据业界对加密的需求及业务改造痛点，提供了一套完整、安全、透明化、低改造成本的数据加密整合解决方案，是 Apache ShardingSphere 数据加密模块的主要设计目标。

3.8.4 应用场景

新上线业务

对于想要快速上线新业务，同时又需要完成安全部门的加密规定的场景，接入 ShardingSphere encrypt 功能，可以快速完成数据的合规化加密，客户无需自行开发复杂的加密系统，同时 ShardingSphere encrypt 的灵活性，也能够帮助客户避免加密场景变更带来的复杂重构和修改风险。

成熟业务

对于已经上线的成熟业务，用户不仅需要考虑历史数据的清洗，还需要考虑新旧功能的切换。接入 ShardingSphere encrypt，用户就可以方便地完成系统的加密改造，它还能够帮助用户安全快速地切换新旧功能。用户无需改动任何业务逻辑和 SQL 就能够透明化地使用加解密功能。

3.8.5 相关参考

- [配置：数据加密](#)
- [开发者指南：数据加密](#)

3.8.6 核心概念

逻辑列

用于计算加解密列的逻辑名称，是 SQL 中列的逻辑标识。逻辑列包含密文列（必须）、查询辅助列（可选）和明文列（可选）。

密文列

加密后的数据列。

查询辅助列

用于查询的辅助列。对于一些安全级别更高的非幂等加密算法，提供不可逆的幂等列用于查询。

明文列

存储明文的列，用于在加密数据迁移过程中仍旧提供服务。在洗数结束后可以删除。

3.8.7 使用限制

- 需自行处理数据库中原始的存量数据；
- 加密字段无法支持查询不区分大小写功能；
- 加密字段无法支持比较操作，如：大于、小于、ORDER BY、BETWEEN、LIKE 等；
- 加密字段无法支持计算操作，如：AVG、SUM 以及计算表达式。

3.8.8 附录

不支持的 SQL：

- 加密字段无法支持查询不区分大小写功能；
- 加密字段无法支持比较操作，如：大于、小于、ORDER BY、BETWEEN、LIKE 等；
- 加密字段无法支持计算操作，如：AVG、SUM 以及计算表达式。

3.9 影子库

3.9.1 背景

在基于微服务的分布式应用架构下，业务需要多个服务是通过一系列的服务、中间件的调用来完成，所以单个服务的压力测试已无法代表真实场景。在测试环境中，如果重新搭建一整套与生产环境类似的压测环境，成本过高，并且往往无法模拟线上环境的复杂度以及流量。因此，业内通常选择全链路压测的方式，即在生产环境进行压测，这样所获得的测试结果能够准确地反应系统真实容量和性能水平。

3.9.2 挑战

全链路压测是一项复杂而庞大的工作。需要各个微服务、中间件之间配合与调整，以应对不同流量以及压测标识的透传。通常会搭建一整套压测平台以适用不同测试计划。在数据库层面需要做好数据隔离，为了保证生产数据的可靠性与完整性，需要将压测产生的数据路由到压测环境数据库，防止压测数据对生产数据库中真实数据造成污染。这就要求业务应用在执行 SQL 前，能够根据透传的压测标识，做好数据分类，将相应的 SQL 路由到与之对应的数据源。

3.9.3 目标

Apache ShardingSphere 关注于全链路压测场景下，数据库层面的解决方案。将压测数据自动路由至用户指定的数据库，是 Apache ShardingSphere 影子库模块的主要设计目标。

3.9.4 应用场景

在基于微服务的分布式应用架构下，为了提升系统压力测试的准确性，降低测试成本。通常选择在生产环境进行压力测试。测试中风险也会大大提高。通过 ShardingSphere 影子库功能，结合影子算法灵活的配置。可以解决数据污染，数据库性能等问题，满足复杂业务场景的在线压力测试需求。

3.9.5 相关参考

- Java API：影子库
- YAML 配置：影子库
- Spring Boot Starter：影子库
- Spring 命名空间：影子库

3.9.6 核心概念

生产库

生产环境使用的数据库。

影子库

压测数据隔离的影子数据库，与生产数据库应当使用相同的配置。

影子算法

影子算法和业务实现紧密相关，目前提供 2 种类型影子算法。

- 基于列的影子算法通过识别 SQL 中的数据，匹配路由至影子库的场景。适用于由压测数据名单驱动的压测场景。
- 基于 Hint 的影子算法通过识别 SQL 中的注释，匹配路由至影子库的场景。适用于由上游系统透传标识驱动的压测场景。

3.9.7 使用限制

基于 Hint 的影子算法

- 无。

基于列的影子算法

- 不支持 DDL;
- 不支持范围、分组和子查询, 如: BETWEEN、GROUP BY ...HAVING 等。SQL 支持列表:

– INSERT

SQL	是否支持
INSERT INTO table (column,...) VALUES (value,...)	支持
INSERT INTO table (column,...) VALUES (value,...),(value,...),...	支持
INSERT INTO table (column,...) SELECT column1 from table1 where column1 = value1	不支持

– SELECT/UPDATE/DELETE

条件类型	SQL	是否支持
=	SELECT/UPDATE/DELETE ...WHERE column = value	支持
LIKE/NOT LIKE	SELECT/UPDATE/DELETE ...WHERE column LIKE/NOT LIKE value	支持
IN/NOT IN	SELECT/UPDATE/DELETE ...WHERE column IN/NOT IN (value1,value2,...)	支持
BETWEEN	SELECT/UPDATE/DELETE ...WHERE column BETWEEN value1 AND value2	不支持
GROUP BY ... HAVING...	SELECT/UPDATE/DELETE ...WHERE ...GROUP BY column HAVING column > value	不支持
子查询	SELECT/UPDATE/DELETE ...WHERE column = (SELECT column FROM table WHERE column = value)	不支持

3.10 可观察性

3.10.1 背景

如何观测集群的运行状态, 使运维人员可以快速掌握当前系统现状, 并进行进一步的维护工作, 是分布式系统的全新挑战。登录到具体服务器的点对点运维方式, 无法适用于面向大量分布式服务器的场景。通

通过对系统观察性数据的遥测是分布式系统推荐的运维方式。Tracing（链路跟踪）、Metrics（指标监控）和 Logging（日志）是系统运行状况的可观察性数据重要的获取手段。

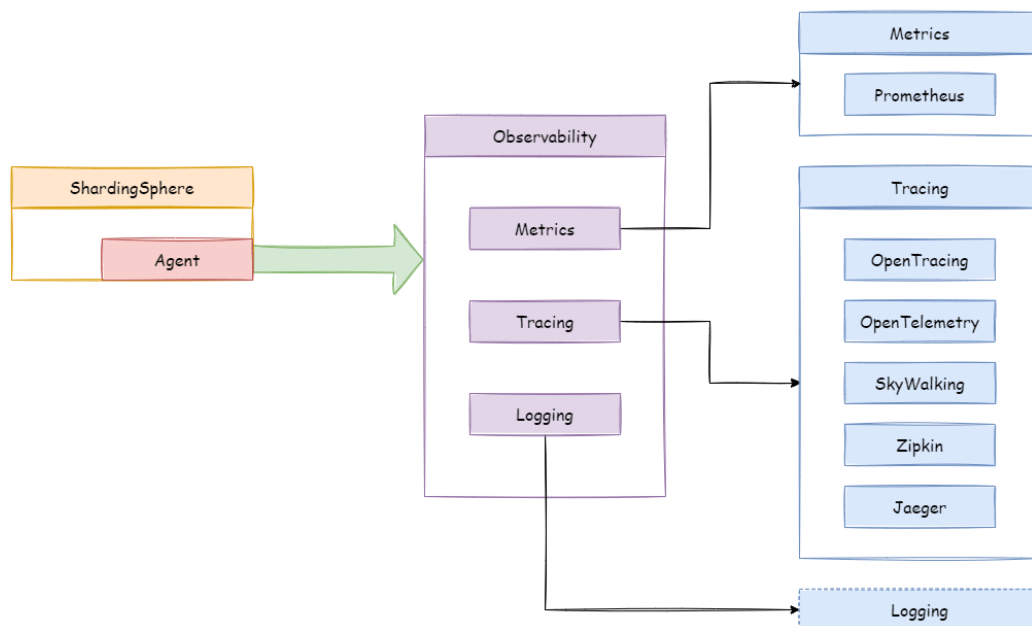
APM（应用性能监控）是通过对系统可观察性数据进行采集、存储和分析，进行系统的性能监控与诊断，主要功能包括性能指标监控、调用链分析，应用拓扑图等。

Apache ShardingSphere 并不负责如何采集、存储以及展示应用性能监控的相关数据，而是为应用监控系统提供必要的指标数据。换句话说，Apache ShardingSphere 仅负责产生具有价值的的数据，并通过标准协议或插件化的方式递交给相关系统。

Tracing 用于获取 SQL 解析与 SQL 执行的链路跟踪信息。Apache ShardingSphere 默认提供了对 SkyWalking, Zipkin, Jaeger 和 OpenTelemetry 的支持，也支持用户通过插件化的方式开发自定义的 Tracing 组件。

- 使用 Zipkin 和 Jaeger 通过在 agent 配置文件中开启对应的插件，并配置好 Zipkin 或者 Jaeger 服务器信息即可。
- 使用 OpenTelemetry OpenTelemetry 在 2019 年由 OpenTracing 和 OpenCensus 合并而来。使用这种方式，只需要在 agent 配置文件中，根据 OpenTelemetry SDK 自动配置说明，填写合适的配置即可。
- 使用 SkyWalking 需要在 agent 配置中配置启用对应插件，并且需要同时配置使用 SkyWalking 的 apm-toolkit 工具。
- 使用 SkyWalking 的内置自动探针 Apache ShardingSphere 团队与 Apache SkyWalking 团队共同合作，在 SkyWalking 中实现了 Apache ShardingSphere 自动探针，可以将相关的应用性能数据自动发送到 SkyWalking 中。注意这种方式的自动探针不能与 Apache ShardingSphere 插件探针同时使用。

Metrics 则用于收集和展示整个集群的统计指标。Apache ShardingSphere 默认提供了对 Prometheus 的支持。



3.10.2 挑战

Tracing 和 Metrics 需要通过埋点来收集系统信息。大量的埋点使项目核心代码支离破碎，难于维护，且不易定制化统计指标。

3.10.3 目标

提供尽量多的性能和统计指标，并隔离核心代码和埋点代码，是 Apache ShardingSphere 可观察性模块的设计目标。

3.10.4 应用场景

ShardingSphere 通过 Agent 模块为应用提供可观察性的能力，可适用于以下场景：

监控仪表盘

将系统静态信息（如应用版本）和动态信息（如线程数、SQL 处理信息）等 Metrics 指标，使用标准接口方式暴露给第三方应用（如 Prometheus），管理员能够通过可视化的方式监控系统实时状态。

应用性能监控

在 ShardingSphere 中，一条 SQL 语句要经历解析、路由、改写、执行、结果归并等流程才能最终执行完成，并输出响应。如果 SQL 语句复杂，整体执行耗时过长，如何知道哪一步存在优化空间呢？

通过 Agent + Tracing，管理员可以了解 SQL 执行过程中每一步的耗时情况，轻松定位性能风险，从而能够有针对性的制定 SQL 优化方案。

应用链路追踪

在分布式应用 + 数据分片的场景下，SQL 语句是哪个节点发出的，最终在哪些数据源执行？这是一个非常棘手的问题。如果 SQL 执行过程中发生异常，如何定位发生异常的节点呢？

Agent + Tracing，能够帮助用户解决以上问题。

通过对 SQL 执行过程的完整链路追踪，用户可以得到“SQL 从哪里来，发到哪里去”这样的完整信息，还能够通过生成的拓扑图来直观的观察 SQL 路由情况，运筹帷幄，同时获得快速定位问题根源的能力。

3.10.5 相关参考

- 可观察性的使用
- 开发者指南：可观察性
- 实现原理

3.10.6 核心概念

Agent

基于字节码增强和插件化设计，以提供 Tracing 和 Metrics 埋点，以及日志输出功能。需要开启 Agent 的插件功能后，才能将监控指标数据输出至第三方 APM 中展示。

APM

APM 是应用性能监控的缩写。着眼于分布式系统的性能诊断，其主要功能包括调用链展示，应用拓扑分析等。

Tracing

链路跟踪，通过探针收集调用链数据，并发送到第三方 APM 系统。

Metrics

系统统计指标，通过探针收集，并且写入到时序数据库，供第三方应用展示。

Logging

日志，通过 Agent 能够方便的扩展日志内容，为分析系统运行状态提供更多信息。

本章节面向 Apache ShardingSphere 的用户，详细阐述项目的使用说明。

4.1 ShardingSphere-JDBC

配置是 ShardingSphere-JDBC 中唯一与应用开发者交互的模块，通过它可以快速清晰的理解 ShardingSphere-JDBC 所提供的功能。

本章节是 ShardingSphere-JDBC 的配置参考手册，需要时可当做字典查阅。

ShardingSphere-JDBC 提供了 4 种配置方式，用于不同的使用场景。通过配置，应用开发者可以灵活的使用数据分片、读写分离、数据加密、影子库等功能，并且能够叠加使用。

混合规则配置与单一规则配置一脉相承，只是从配置单一的规则项到配置多个规则项的异同。

需要注意的是，规则项之间的叠加使用是通过数据源名称和表名称关联的。如果前一个规则是面向数据源聚合的，下一个规则在配置数据源时，则需要使用前一个规则配置的聚合后的逻辑数据源名称；同理，如果前一个规则是面向表聚合的，下一个规则在配置表时，则需要使用前一个规则配置的聚合后的逻辑表名称。

更多使用细节请参见[使用示例](#)。

4.1.1 YAML 配置

简介

YAML 提供通过配置文件的方式与 ShardingSphere-JDBC 交互。配合治理模块一同使用时，持久化在配置中心的配置均为 YAML 格式。

YAML 配置是最常见的配置方式，可以省略编程的复杂度，简化用户配置。

使用步骤

引入 Maven 依赖

```
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-jdbc-core</artifactId>
  <version>${shardingsphere.version}</version>
</dependency>
```

配置 YAML

ShardingSphere-JDBC 的 YAML 配置文件通过 Database 名称、运行模式、数据源集合、规则集合以及属性配置组成。

```
# JDBC 逻辑库名称。在集群模式中，使用该参数来联通 ShardingSphere-JDBC 与 ShardingSphere-Proxy。
# 默认值: logic_db
databaseName (?):

mode:

dataSources:

rules:
- !FOO_XXX
  ...
- !BAR_XXX
  ...

props:
  key_1: value_1
  key_2: value_2
```

模式详情请参见模式配置。

数据源详情请参见数据源配置。

规则详情请参见规则配置。

构建数据源

通过 `YamlShardingSphereDataSourceFactory` 工厂创建的 `ShardingSphereDataSource` 实现自 JDBC 的标准接口 `DataSource`。

```
File yamlFile = // 指定 YAML 文件路径
DataSource dataSource = YamlShardingSphereDataSourceFactory.
createDataSource(yamlFile);
```

使用数据源

使用方式同 Java API。

语法说明

!! 表示实例化该类

! 表示自定义别名

- 表示可以包含一个或多个

[] 表示数组，可以与减号相互替换使用

模式配置

参数解释

```
mode (?): # 不配置则默认单机模式
  type: # 运行模式类型。可选配置: Standalone、Cluster
  repository (?): # 久化仓库配置
    overwrite: # 是否使用本地配置覆盖持久化配置
```

单机模式

```
mode:
  type: Standalone
  repository:
    type: # 持久化仓库类型
    props: # 持久化仓库所需属性
      foo_key: foo_value
      bar_key: bar_value
    overwrite: # 是否使用本地配置覆盖持久化配置
```

集群模式 (推荐)

```
mode:
  type: Cluster
  repository:
    type: # 持久化仓库类型
    props: # 持久化仓库所需属性
      namespace: # 注册中心命名空间
      server-lists: # 注册中心连接地址
      foo_key: foo_value
      bar_key: bar_value
  overwrite: # 是否使用本地配置覆盖持久化配置
```

注意事项

1. 生产环境建议使用集群模式部署。
2. 集群模式部署推荐使用 ZooKeeper 注册中心。

配置示例

单机模式

```
mode:
  type: Standalone
  repository:
    type: File
  overwrite: false
```

集群模式 (推荐)

```
mode:
  type: Cluster
  repository:
    type: ZooKeeper
    props:
      namespace: governance
      server-lists: localhost:2181
      retryIntervalMilliseconds: 500
      timeToLiveSeconds: 60
  overwrite: false
```

相关参考

- [ZooKeeper 注册中心安装与使用](#)
- 持久化仓库类型的详情，请参见[内置持久化仓库类型列表](#)。

数据源配置

背景信息

ShardingSphere-JDBC 支持所有的数据库 JDBC 驱动和连接池。

示例的数据库驱动为 MySQL，连接池为 HikariCP，可以更换为其他数据库驱动和连接池。当使用 ShardingSphere-JDBC 时，JDBC 池的属性名取决于各自 JDBC 池自己的定义，并不由 ShardingSphere 硬定义，相关的处理可以参考类 `org.apache.shardingsphere.infra.datasource.pool.creator.DataSourcePoolCreator`。例如对于 Alibaba Druid 1.2.9 而言，使用 `url` 代替如下示例中的 `jdbcUrl` 是预期行为。

参数解释

```
dataSources: # 数据源配置, 可配置多个 <data-source-name>
  <data-source-name>: # 数据源名称
    dataSourceClassName: # 数据源完整类名
    driverClassName: # 数据库驱动类名, 以数据库连接池自身配置为准
    jdbcUrl: # 数据库 URL 连接, 以数据库连接池自身配置为准
    username: # 数据库用户名, 以数据库连接池自身配置为准
    password: # 数据库密码, 以数据库连接池自身配置为准
    # ... 数据库连接池的其它属性
```

配置示例

```
dataSources:
  ds_1:
    dataSourceClassName: com.zaxxer.hikari.HikariDataSource
    driverClassName: com.mysql.jdbc.Driver
    jdbcUrl: jdbc:mysql://localhost:3306/ds_1
    username: root
    password:
  ds_2:
    dataSourceClassName: com.zaxxer.hikari.HikariDataSource
    driverClassName: com.mysql.jdbc.Driver
    jdbcUrl: jdbc:mysql://localhost:3306/ds_2
    username: root
    password:

# 配置其他数据源
```

规则配置

规则是 Apache ShardingSphere 面向可插拔的一部分。本章节是 ShardingSphere-JDBC 的 YAML 规则配置参考手册。

数据分片

背景信息

数据分片 YAML 配置方式具有非凡的可读性，通过 YAML 格式，能够快速地理解分片规则之间的依赖关系，ShardingSphere 会根据 YAML 配置，自动完成 ShardingSphereDataSource 对象的创建，减少用户不必要的编码工作。

参数解释

```
rules:
- !SHARDING
  tables: # 数据分片规则配置
    <logic-table-name> (+): # 逻辑表名称
      actualDataNodes (?): # 由数据源名 + 表名组成 (参考 Inline 语法规则)
      databaseStrategy (?): # 分库策略, 缺省表示使用默认分库策略, 以下的分片策略只能选其一
        standard: # 用于单分片键的标准分片场景
          shardingColumn: # 分片列名称
          shardingAlgorithmName: # 分片算法名称
        complex: # 用于多分片键的复合分片场景
          shardingColumns: # 分片列名称, 多个列以逗号分隔
          shardingAlgorithmName: # 分片算法名称
      hint: # Hint 分片策略
        shardingAlgorithmName: # 分片算法名称
      none: # 不分片
      tableStrategy: # 分表策略, 同分库策略
      keyGenerateStrategy: # 分布式序列策略
        column: # 自增列名称, 缺省表示不使用自增主键生成器
        keyGeneratorName: # 分布式序列算法名称
  autoTables: # 自动分片表规则配置
    t_order_auto: # 逻辑表名称
      actualDataSources (?): # 数据源名称
      shardingStrategy: # 切分策略
        standard: # 用于单分片键的标准分片场景
          shardingColumn: # 分片列名称
          shardingAlgorithmName: # 自动分片算法名称
  bindingTables (+): # 绑定表规则列表
    - <logic_table_name_1, logic_table_name_2, ...>
    - <logic_table_name_1, logic_table_name_2, ...>
```

```

broadcastTables (+): # 广播表规则列表
  - <table-name>
  - <table-name>
defaultDatabaseStrategy: # 默认数据库分片策略
defaultTableStrategy: # 默认表分片策略
defaultKeyGenerateStrategy: # 默认的分布式序列策略
defaultShardingColumn: # 默认分片列名称

# 分片算法配置
shardingAlgorithms:
  <sharding-algorithm-name> (+): # 分片算法名称
    type: # 分片算法类型
    props: # 分片算法属性配置
    # ...

# 分布式序列算法配置
keyGenerators:
  <key-generate-algorithm-name> (+): # 分布式序列算法名称
    type: # 分布式序列算法类型
    props: # 分布式序列算法属性配置
    # ...

```

操作步骤

1. 在 YAML 文件中配置数据分片规则，包含数据源、分片规则、全局属性等配置项；
2. 调用 `YamlShardingSphereDataSourceFactory` 对象的 `createDataSource` 方法，根据 YAML 文件中的配置信息创建 `ShardingSphereDataSource`。

配置示例

数据分片 YAML 配置示例如下：

```

dataSources:
  ds_0:
    dataSourceClassName: com.zaxxer.hikari.HikariDataSource
    driverClassName: com.mysql.jdbc.Driver
    jdbcUrl: jdbc:mysql://localhost:3306/demo_ds_0?serverTimezone=UTC&useSSL=false&
useUnicode=true&characterEncoding=UTF-8
    username: root
    password:
  ds_1:
    dataSourceClassName: com.zaxxer.hikari.HikariDataSource
    driverClassName: com.mysql.jdbc.Driver
    jdbcUrl: jdbc:mysql://localhost:3306/demo_ds_1?serverTimezone=UTC&useSSL=false&
useUnicode=true&characterEncoding=UTF-8
    username: root

```

```

password:

rules:
- !SHARDING
  tables:
    t_order:
      actualDataNodes: ds_${0..1}.t_order_${0..1}
      tableStrategy:
        standard:
          shardingColumn: order_id
          shardingAlgorithmName: t-order-inline
      keyGenerateStrategy:
        column: order_id
        keyGeneratorName: snowflake
    t_order_item:
      actualDataNodes: ds_${0..1}.t_order_item_${0..1}
      tableStrategy:
        standard:
          shardingColumn: order_id
          shardingAlgorithmName: t_order-item-inline
      keyGenerateStrategy:
        column: order_item_id
        keyGeneratorName: snowflake
    t_account:
      actualDataNodes: ds_${0..1}.t_account_${0..1}
      tableStrategy:
        standard:
          shardingAlgorithmName: t-account-inline
      keyGenerateStrategy:
        column: account_id
        keyGeneratorName: snowflake
  defaultShardingColumn: account_id
  bindingTables:
    - t_order,t_order_item
  broadcastTables:
    - t_address
  defaultDatabaseStrategy:
    standard:
      shardingColumn: user_id
      shardingAlgorithmName: database-inline
  defaultTableStrategy:
    none:

  shardingAlgorithms:
    database-inline:
      type: INLINE
      props:
        algorithm-expression: ds_${user_id % 2}

```

```
t-order-inline:
  type: INLINE
  props:
    algorithm-expression: t_order_${order_id % 2}
t_order-item-inline:
  type: INLINE
  props:
    algorithm-expression: t_order_item_${order_id % 2}
t-account-inline:
  type: INLINE
  props:
    algorithm-expression: t_account_${account_id % 2}
keyGenerators:
  snowflake:
    type: SNOWFLAKE

props:
  sql-show: false
```

通过 `YamlShardingSphereDataSourceFactory` 的 `createDataSource` 方法, 读取 YAML 配置完成数据源的创建。

```
YamlShardingSphereDataSourceFactory.createDataSource(getFile("/META-INF/sharding-
databases-tables.yaml"));
```

相关参考

- 核心特性: 数据分片
- 开发者指南: 数据分片

读写分离

背景信息

读写分离 YAML 配置方式可读性高, 通过 YAML 格式, 能够快速地理解读写分片规则之间的依赖关系, ShardingSphere 会根据 YAML 配置, 自动完成 `ShardingSphereDataSource` 对象的创建, 减少用户不必要的编码工作。

参数解释

静态读写分离

```
rules:
- !READWRITE_SPLITTING
  dataSources:
    <data-source-name> (+): # 读写分离逻辑数据源名称
      static-strategy: # 读写分离类型
      write-data-source-name: # 写库数据源名称
      read-data-source-names: # 读库数据源名称, 多个从数据源用逗号分隔
      loadBalancerName: # 负载均衡算法名称

  # 负载均衡算法配置
  loadBalancers:
    <load-balancer-name> (+): # 负载均衡算法名称
      type: # 负载均衡算法类型
      props: # 负载均衡算法属性配置
      # ...
```

动态读写分离

```
rules:
- !READWRITE_SPLITTING
  dataSources:
    <data-source-name> (+): # 读写分离逻辑数据源名称
      dynamic-strategy: # 读写分离类型
      auto-aware-data-source-name: # 数据库发现逻辑数据源名称
      write-data-source-query-enabled: # 从库全部下线, 主库是否承担读流量
      loadBalancerName: # 负载均衡算法名称

  # 负载均衡算法配置
  loadBalancers:
    <load-balancer-name> (+): # 负载均衡算法名称
      type: # 负载均衡算法类型
      props: # 负载均衡算法属性配置
      # ...
```

算法类型的详情, 请参见内置负载均衡算法列表。查询一致性路由的详情, 请参见核心特性: 读写分离。

操作步骤

1. 添加读写分离数据源
2. 设置负载均衡算法
3. 使用读写分离数据源

配置示例

```
rules:
- !READWRITE_SPLITTING
  dataSources:
    readwrite_ds:
      staticStrategy:
        writeDataSourceName: write_ds
        readDataSourceNames:
          - read_ds_0
          - read_ds_1
        loadBalancerName: random
  loadBalancers:
    random:
      type: RANDOM
```

相关参考

- [核心特性：读写分离](#)
- [Java API：读写分离](#)
- [Spring Boot Starter：读写分离](#)
- [Spring 命名空间：读写分离](#)

分布式事务

背景信息

ShardingSphere 提供了三种模式的分布式事务 LOCAL, XA, BASE。

参数解释

```
rules:
- !TRANSACTION
  defaultType: # 事务模式, 可选值 LOCAL/XA/BASE
  providerType: # 指定模式下的具体实现
```

操作步骤

使用 **LOCAL** 模式

server.yaml 配置文件内容如下:

```
rules:
- !TRANSACTION
  defaultType: LOCAL
```

使用 **XA** 模式

server.yaml 配置文件内容如下:

```
rules:
- !TRANSACTION
  defaultType: XA
  providerType: Narayana/Atomikos
```

手动添加 Narayana 相关依赖:

```
jta-5.12.4.Final.jar
arjuna-5.12.4.Final.jar
common-5.12.4.Final.jar
jboss-connector-api_1.7_spec-1.0.0.Final.jar
jboss-logging-3.2.1.Final.jar
jboss-transaction-api_1.2_spec-1.0.0.Alpha3.jar
jboss-transaction-spi-7.6.0.Final.jar
narayana-jts-integration-5.12.4.Final.jar
shardingsphere-transaction-xa-narayana-x.x.x-SNAPSHOT.jar
```

使用 BASE 模式

server.yaml 配置文件内容如下：

```
rules:
  - !TRANSACTION
    defaultType: BASE
    providerType: Seata
```

搭建 Seata Server，添加相关配置文件，和 Seata 依赖，具体步骤参考 [ShardingSphere 集成 Seata 柔性事务](#)

高可用

背景信息

通过 YAML 格式，ShardingSphere 会根据 YAML 配置，自动完成 ShardingSphereDataSource 对象的创建，减少用户不必要的编码工作。

参数解释

```
rules:
  - !READWRITE_SPLITTING
    dataSources:
      replica_ds:
        dynamicStrategy: Dynamic # 动态读写分离
        autoAwareDataSourceName: # 高可用规则逻辑数据源名称

  - !DB_DISCOVERY
    dataSources:
      <data-source-name> (+): # 逻辑数据源名称
        dataSourceNames: # 数据源名称列表
          - <data-source>
          - <data-source>
        discoveryHeartbeatName: # 检测心跳名称
        discoveryTypeName: # 数据库发现类型名称

    # 心跳检测配置
    discoveryHeartbeats:
      <discovery-heartbeat-name> (+): # 心跳名称
        props:
          keep-alive-cron: # cron 表达式, 如: '0/5 * * * * ?'

    # 数据库发现类型配置
    discoveryTypes:
      <discovery-type-name> (+): # 数据库发现类型名称
```

```

type: # 数据库发现类型, 如: MySQL.MGR
props (?):
  group-name: 92504d5b-6dec-11e8-91ea-246e9612aaf1 # 数据库发现类型必要参数, 如
MGR 的 group-name

```

配置示例

```

databaseName: database_discovery_db

dataSources:
  ds_0:
    url: jdbc:mysql://127.0.0.1:33306/primary_demo_ds?serverTimezone=UTC&
useSSL=false
    username: root
    password:
    connectionTimeoutMilliseconds: 3000
    idleTimeoutMilliseconds: 60000
    maxLifetimeMilliseconds: 1800000
    maxPoolSize: 50
    minPoolSize: 1
  ds_1:
    url: jdbc:mysql://127.0.0.1:33307/primary_demo_ds?serverTimezone=UTC&
useSSL=false
    username: root
    password:
    connectionTimeoutMilliseconds: 3000
    idleTimeoutMilliseconds: 60000
    maxLifetimeMilliseconds: 1800000
    maxPoolSize: 50
    minPoolSize: 1
  ds_2:
    url: jdbc:mysql://127.0.0.1:33308/primary_demo_ds?serverTimezone=UTC&
useSSL=false
    username: root
    password:
    connectionTimeoutMilliseconds: 3000
    idleTimeoutMilliseconds: 60000
    maxLifetimeMilliseconds: 1800000
    maxPoolSize: 50
    minPoolSize: 1

rules:
  - !READWRITE_SPLITTING
    dataSources:
      replica_ds:
        dynamicStrategy:

```

```

    autoAwareDataSourceName: readwrite_ds

- !DB_DISCOVERY
  dataSources:
    readwrite_ds:
      dataSourceNames:
        - ds_0
        - ds_1
        - ds_2
      discoveryHeartbeatName: mgr-heartbeat
      discoveryTypeName: mgr
  discoveryHeartbeats:
    mgr-heartbeat:
      props:
        keep-alive-cron: '0/5 * * * * ?'
  discoveryTypes:
    mgr:
      type: MySQL.MGR
      props:
        group-name: 558edd3c-02ec-11ea-9bb3-080027e39bd2

```

相关参考

- [高可用核心特性](#)
- [JAVA API：高可用配置](#)
- [Spring Boot Starter：高可用配置](#)
- [Spring 命名空间：高可用配置](#)

数据加密

背景信息

数据加密 YAML 配置方式具有非凡的可读性，通过 YAML 格式，能够快速地理解加密规则之间的依赖关系，ShardingSphere 会根据 YAML 配置，自动完成 ShardingSphereDataSource 对象的创建，减少用户不必要的编码工作。

参数解释

```
rules:
- !ENCRYPT
  tables:
    <table-name> (+): # 加密表名称
      columns:
        <column-name> (+): # 加密列名称
          cipherColumn: # 密文列名称
          assistedQueryColumn (?): # 查询辅助列名称
          plainColumn (?): # 原文列名称
          encryptorName: # 加密算法名称
          queryWithCipherColumn(?): # 该表是否使用加密列进行查询

# 加密算法配置
encryptors:
  <encrypt-algorithm-name> (+): # 加解密算法名称
    type: # 加解密算法类型
    props: # 加解密算法属性配置
    # ...

queryWithCipherColumn: # 是否使用加密列进行查询。在有原文列的情况下，可以使用原文列进行查询
```

算法类型的详情，请参见内置加密算法列表。

操作步骤

1. 在 YAML 文件中配置数据加密规则，包含数据源、加密规则、全局属性等配置项；
2. 调用 YamlShardingSphereDataSourceFactory 对象的 createDataSource 方法，根据 YAML 文件中的配置信息创建 ShardingSphereDataSource。

配置示例

数据加密 YAML 配置如下：

```
dataSources:
  unique_ds:
    dataSourceClassName: com.zaxxer.hikari.HikariDataSource
    driverClassName: com.mysql.jdbc.Driver
    jdbcUrl: jdbc:mysql://localhost:3306/demo_ds?serverTimezone=UTC&useSSL=false&
useUnicode=true&characterEncoding=UTF-8
    username: root
    password:

rules:
- !ENCRYPT
```

```

tables:
  t_user:
    columns:
      username:
        plainColumn: username_plain
        cipherColumn: username
        encryptorName: name-encryptor
      pwd:
        cipherColumn: pwd
        assistedQueryColumn: assisted_query_pwd
        encryptorName: pwd_encryptor
    encryptors:
      name-encryptor:
        type: AES
        props:
          aes-key-value: 123456abc
      pwd_encryptor:
        type: assistedTest

```

然后通过 `YamlShardingSphereDataSourceFactory` 的 `createDataSource` 方法创建数据源。

```
YamlShardingSphereDataSourceFactory.createDataSource(getFile());
```

相关参考

- [核心特性：数据加密](#)
- [开发者指南：数据加密](#)

影子库

背景信息

如果您想在 ShardingSphere-Proxy 中使用 ShardingSphere 影子库功能请参考以下配置。

参数解释

```

rules:
- !SHADOW
  dataSources:
    shadowDataSource:
      productionDataSourceName: # 生产数据源名称
      shadowDataSourceName: # 影子数据源名称
  tables:
    <table-name>:

```



```

dataSourceNames: # 影子表关联影子数据源名称列表
  - <shadow-data-source>
shadowAlgorithmNames: # 影子表关联影子算法名称列表
  - <shadow-algorithm-name>
defaultShadowAlgorithmName: # 默认影子算法名称（选配项）
shadowAlgorithms:
  <shadow-algorithm-name> (+): # 影子算法名称
    type: # 影子算法类型
    props: # 影子算法属性配置

```

详情请参见内置影子算法列表

操作步骤

1. 在 YAML 文件中配置影子库规则，包含数据源、影子库规则、全局属性等配置项；
2. 调用 `YamlShardingSphereDataSourceFactory` 对象的 `createDataSource` 方法，根据 YAML 文件中的配置信息创建 `ShardingSphereDataSource`。

配置示例

```

dataSources:
  ds:
    url: jdbc:mysql://127.0.0.1:3306/ds?serverTimezone=UTC&useSSL=false
    username: root
    password:
    connectionTimeoutMilliseconds: 30000
    idleTimeoutMilliseconds: 60000
    maxLifetimeMilliseconds: 1800000
    maxPoolSize: 50
    minPoolSize: 1
  shadow_ds:
    url: jdbc:mysql://127.0.0.1:3306/shadow_ds?serverTimezone=UTC&useSSL=false
    username: root
    password:
    connectionTimeoutMilliseconds: 30000
    idleTimeoutMilliseconds: 60000
    maxLifetimeMilliseconds: 1800000
    maxPoolSize: 50
    minPoolSize: 1
rules:
- !SHADOW
  dataSources:
    shadowDataSource:
      productionDataSourceName: ds
      shadowDataSourceName: shadow_ds

```

```

tables:
  t_order:
    dataSourceNames:
      - shadowDataSource
    shadowAlgorithmNames:
      - user-id-insert-match-algorithm
      - simple-hint-algorithm
    shadowAlgorithms:
      user-id-insert-match-algorithm:
        type: REGEX_MATCH
        props:
          operation: insert
          column: user_id
          regex: "[1]"
      simple-hint-algorithm:
        type: SIMPLE_HINT
        props:
          foo: bar

```

相关参考

- [影子库的核心特性](#)
- [JAVA API：影子库配置](#)
- [Spring Boot Starter：影子库配置](#)
- [Spring 命名空间：影子库配置](#)

SQL 解析

背景信息

SQL 解析 YAML 配置方式具有可读性高，使用简单的特点。通过 YAML 文件的方式，用户可以将代码与配置分离，并且根据需要方便地修改配置文件。

参数解释

```

rules:
- !SQL_PARSER
  sqlCommentParseEnabled: # 是否解析 SQL 注释
  sqlStatementCache: # SQL 语句本地缓存配置项
    initialCapacity: # 本地缓存初始容量
    maximumSize: # 本地缓存最大容量
  parseTreeCache: # 解析树本地缓存配置项

```

```
initialCapacity: # 本地缓存初始容量
maximumSize: # 本地缓存最大容量
```

操作步骤

1. 设置本地缓存配置
2. 设置解析配置
3. 使用解析引擎解析 SQL

配置示例

```
rules:
- !SQL_PARSER
  sqlCommentParseEnabled: true
  sqlStatementCache:
    initialCapacity: 2000
    maximumSize: 65535
  parseTreeCache:
    initialCapacity: 128
    maximumSize: 1024
```

相关参考

- [JAVA API: SQL 解析](#)
- [Spring Boot Starter: SQL 解析](#)
- [Spring 命名空间: SQL 解析](#)

SQL 翻译

配置项说明

```
rules:
- !SQL_TRANSLATOR
  type: # SQL 翻译器类型
  useOriginalSQLWhenTranslatingFailed: # SQL 翻译失败是否使用原始 SQL 继续执行
```

混合规则

背景信息

ShardingSphere 涵盖了很多功能，例如，分库分片、读写分离、高可用、数据脱敏等。这些功能用户可以单独进行使用，也可以配合一起使用，下面是基于 YAML 的参数解释和配置示例。

参数解释

```
rules:
- !SHARDING
  tables:
    <logic-table-name>: # 逻辑表名称:
      actualDataNodes: # 由逻辑数据源名 + 表名组成 (参考 Inline 语法规则)
      tableStrategy: # 分表策略, 同分库策略
        standard:
          shardingColumn: # 分片列名称
          shardingAlgorithmName: # 分片算法名称
      keyGenerateStrategy:
        column: # 自增列名称, 缺省表示不使用自增主键生成器
        keyGeneratorName: # 分布式序列算法名称
      defaultDatabaseStrategy:
        standard:
          shardingColumn: # 分片列名称
          shardingAlgorithmName: # 分片算法名称
      shardingAlgorithms:
        <sharding-algorithm-name>: # 分片算法名称
          type: INLINE
          props:
            algorithm-expression: # INLINE 表达式
          t_order_inline:
            type: INLINE
            props:
              algorithm-expression: # INLINE 表达式
      keyGenerators:
        <key-generate-algorithm-name> (+): # 分布式序列算法名称
          type: # 分布式序列算法类型
          props: # 分布式序列算法属性配置
- !READWRITE_SPLITTING
  dataSources:
    <data-source-name>: # 读写分离逻辑数据源名称
      dynamicStrategy: # 读写分离类型
        autoAwareDataSourceName: # 数据库发现逻辑数据源名称
    <data-source-name>: # 读写分离逻辑数据源名称
      dynamicStrategy: # 读写分离类型
        autoAwareDataSourceName: # 数据库发现逻辑数据源名称
- !DB_DISCOVERY
```

```

dataSources:
  <data-source-name>:
    dataSourceNames: # 数据源名称列表
      - ds_0
      - ds_1
      - ds_2
    discoveryHeartbeatName: # 检测心跳名称
    discoveryTypeName: # 数据库发现类型名称
  <data-source-name>:
    dataSourceNames: # 数据源名称列表
      - ds_3
      - ds_4
      - ds_5
    discoveryHeartbeatName: # 检测心跳名称
    discoveryTypeName: # 数据库发现类型名称
discoveryHeartbeats:
  <discovery-heartbeat-name>: # 心跳名称
  props:
    keep-alive-cron: # cron 表达式, 如: '0/5 * * * * ?'
discoveryTypes:
  <discovery-type-name>: # 数据库发现类型名称
  type: # 数据库发现类型, 如: MySQL.MGR
  props:
    group-name: # 数据库发现类型必要参数, 如 MGR 的 group-name
- !ENCRYPT
encryptors:
  <encrypt-algorithm-name> (+): # 加解密算法名称
  type: # 加解密算法类型
  props: # 加解密算法属性配置
  <encrypt-algorithm-name> (+): # 加解密算法名称
  type: # 加解密算法类型
tables:
  <table-name>: # 加密表名称
  columns:
    <column-name>: # 加密列名称
    plainColumn: # 原文列名称
    cipherColumn: # 密文列名称
    encryptorName: # 加密算法名称
  <column-name>: # 加密列名称
  cipherColumn: # 密文列名称
  encryptorName: # 加密算法名称

```

配置示例

```

rules:
- !SHARDING
  tables:
    t_order:
      actualDataNodes: replica_ds_${0..1}.t_order_${0..1}
      tableStrategy:
        standard:
          shardingColumn: order_id
          shardingAlgorithmName: t_order_inline
      keyGenerateStrategy:
        column: order_id
        keyGeneratorName: snowflake
  defaultDatabaseStrategy:
    standard:
      shardingColumn: user_id
      shardingAlgorithmName: database_inline
  shardingAlgorithms:
    database_inline:
      type: INLINE
      props:
        algorithm-expression: replica_ds_${user_id % 2}
    t_order_inline:
      type: INLINE
      props:
        algorithm-expression: t_order_${order_id % 2}
    t_order_item_inline:
      type: INLINE
      props:
        algorithm-expression: t_order_item_${order_id % 2}
  keyGenerators:
    snowflake:
      type: SNOWFLAKE
- !READWRITE_SPLITTING
  dataSources:
    replica_ds_0:
      dynamicStrategy:
        autoAwareDataSourceName: readwrite_ds_0
    replica_ds_1:
      dynamicStrategy:
        autoAwareDataSourceName: readwrite_ds_1
- !DB_DISCOVERY
  dataSources:
    readwrite_ds_0:
      dataSourceNames:
        - ds_0
        - ds_1

```

```

    - ds_2
    discoveryHeartbeatName: mgr-heartbeat
    discoveryTypeName: mgr
  readwrite_ds_1:
    dataSourceNames:
      - ds_3
      - ds_4
      - ds_5
    discoveryHeartbeatName: mgr-heartbeat
    discoveryTypeName: mgr
  discoveryHeartbeats:
    mgr-heartbeat:
      props:
        keep-alive-cron: '0/5 * * * * ?'
  discoveryTypes:
    mgr:
      type: MySQL.MGR
      props:
        group-name: 558edd3c-02ec-11ea-9bb3-080027e39bd2
- !ENCRYPT
  encryptors:
    aes_encryptor:
      type: AES
      props:
        aes-key-value: 123456abc
    md5_encryptor:
      type: MD5
  tables:
    t_encrypt:
      columns:
        user_id:
          plainColumn: user_plain
          cipherColumn: user_cipher
          encryptorName: aes_encryptor
        order_id:
          cipherColumn: order_cipher
          encryptorName: md5_encryptor

```

算法配置

分片算法

```

shardingAlgorithms:
  # algorithmName 由用户指定, 需要和分片策略中的 shardingAlgorithmName 属性一致
  <algorithmName>:
    # type 和 props, 请参考分片内置算法: https://shardingsphere.apache.org/document/

```

```
current/cn/user-manual/common-config/builtin-algorithm/sharding/
  type: xxx
  props:
    xxx: xxx
```

加密算法

```
encryptors:
  # encryptorName 由用户指定, 需要和加密规则中的 encryptorName 属性一致
  <encryptorName>:
    # type 和 props, 请参考加密内置算法: https://shardingsphere.apache.org/document/
    current/cn/user-manual/common-config/builtin-algorithm/encrypt/
    type: xxx
    props:
      xxx: xxx
```

读写分离负载均衡算法

```
loadBalancers:
  # loadBalancerName 由用户指定, 需要和读写分离规则中的 loadBalancerName 属性一致
  <loadBalancerName>:
    # type 和 props, 请参考读写分离负载均衡内置算法: https://shardingsphere.apache.org/
    document/current/cn/user-manual/common-config/builtin-algorithm/load-balance/
    type: xxx
    props:
      xxx: xxx
```

影子算法

```
loadBalancers:
  # shadowAlgorithmName 由用户指定, 需要和影子库规则中的 shadowAlgorithmNames 属性一致
  <shadowAlgorithmName>:
    # type 和 props, 请参考影子库内置算法: https://shardingsphere.apache.org/document/
    current/cn/user-manual/common-config/builtin-algorithm/shadow/
    type: xxx
    props:
      xxx: xxx
```


高可用

```
discoveryTypes:
  # discoveryTypeName 由用户指定, 需要和数据库发现规则中的 discoveryTypeName 属性一致
  <discoveryTypeName>:
    type: xxx
    props:
      xxx: xxx
```

JDBC 驱动

背景信息

ShardingSphere-JDBC 提供了 JDBC 驱动, 可以仅通过配置变更即可使用, 无需改写代码。

参数解释

驱动类名称

`org.apache.shardingsphere.driver.ShardingSphereDriver`

URL 配置

- 以 `jdbc:shardingsphere:` 为前缀
- 配置文件: `xxx.yaml`, 配置文件格式与 [YAML 配置](#) 一致
- 配置文件加载规则:
 - 无前缀表示从绝对路径加载配置文件
 - `classpath:` 前缀表示从类路径中加载配置文件

操作步骤

1. 引入 Maven 依赖

```
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-jdbc-core</artifactId>
  <version>${shardingsphere.version}</version>
</dependency>
```

2. 使用驱动

- 使用原生驱动:

```

Class.forName("org.apache.shardingsphere.driver.ShardingSphereDriver");
String jdbcUrl = "jdbc:shardingsphere:classpath:config.yaml";

String sql = "SELECT i.* FROM t_order o JOIN t_order_item i ON o.order_id=i.order_id WHERE o.user_id=? AND o.order_id=?";
try (
    Connection conn = DriverManager.getConnection(jdbcUrl);
    PreparedStatement ps = conn.prepareStatement(sql)) {
    ps.setInt(1, 10);
    ps.setInt(2, 1000);
    try (ResultSet rs = preparedStatement.executeQuery()) {
        while(rs.next()) {
            // ...
        }
    }
}

```

- 使用数据库连接池

```

String driverClassName = "org.apache.shardingsphere.driver.ShardingSphereDriver";
String jdbcUrl = "jdbc:shardingsphere:classpath:config.yaml";

// 以 HikariCP 为例
HikariDataSource dataSource = new HikariDataSource();
dataSource.setDriverClassName(driverClassName);
dataSource.setJdbcUrl(jdbcUrl);

String sql = "SELECT i.* FROM t_order o JOIN t_order_item i ON o.order_id=i.order_id WHERE o.user_id=? AND o.order_id=?";
try (
    Connection conn = dataSource.getConnection();
    PreparedStatement ps = conn.prepareStatement(sql)) {
    ps.setInt(1, 10);
    ps.setInt(2, 1000);
    try (ResultSet rs = preparedStatement.executeQuery()) {
        while(rs.next()) {
            // ...
        }
    }
}

```

配置示例

加载 classpath 中 config.yaml 配置文件的 JDBC URL:

```
jdbc:shardingsphere:classpath:config.yaml
```

加载绝对路径中 config.yaml 配置文件的 JDBC URL:

```
jdbc:shardingsphere:/path/to/config.yaml
```

4.1.2 Java API

简介

Java API 是 ShardingSphere-JDBC 中所有配置方式的基础，其他配置最终都将转化成为 Java API 的配置方式。

Java API 是最繁琐也是最灵活的配置方式，适合需要通过编程进行动态配置的场景下使用。

使用步骤

引入 Maven 依赖

```
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-jdbc-core</artifactId>
  <version>${shardingsphere.version}</version>
</dependency>
```

构建数据源

ShardingSphere-JDBC 的 Java API 由 Database 名称、运行模式、数据源集合、规则集合以及属性配置组成。

通过 ShardingSphereDataSourceFactory 工厂创建的 ShardingSphereDataSource 实现自 JDBC 的标准接口 DataSource。

```
String databaseName = "foo_schema"; // 指定逻辑 Database 名称
ModeConfiguration modeConfig = ... // 构建运行模式
Map<String, DataSource> dataSourceMap = ... // 构建真实数据源
Collection<RuleConfiguration> ruleConfigs = ... // 构建具体规则
Properties props = ... // 构建属性配置
DataSource dataSource = ShardingSphereDataSourceFactory.
createDataSource(databaseName, modeConfig, dataSourceMap, ruleConfigs, props);
```

模式详情请参见模式配置。

数据源详情请参见数据源配置。

规则详情请参见规则配置。

使用数据源

可通过 DataSource 选择使用原生 JDBC，或 JPA、Hibernate、MyBatis 等 ORM 框架。

以原生 JDBC 使用方式为例：

```
// 创建 ShardingSphereDataSource
DataSource dataSource = ShardingSphereDataSourceFactory.
createDataSource(databaseName, modeConfig, dataSourceMap, ruleConfigs, props);

String sql = "SELECT i.* FROM t_order o JOIN t_order_item i ON o.order_id=i.order_
id WHERE o.user_id=? AND o.order_id=?";
try (
    Connection conn = dataSource.getConnection();
    PreparedStatement ps = conn.prepareStatement(sql)) {
    ps.setInt(1, 10);
    ps.setInt(2, 1000);
    try (ResultSet rs = preparedStatement.executeQuery()) {
        while(rs.next()) {
            // ...
        }
    }
}
}
```

模式配置

背景信息

通过 Java API 方式构建运行模式。

参数解释

类名称：org.apache.shardingsphere.infra.config.mode.ModeConfiguration

可配置属性：

Standalone 持久化配置

类名称:org.apache.shardingsphere.mode.repository.standalone.StandalonePersistRepositoryConfiguration

可配置属性:

名称	数据类型	说明
type	String	持久化仓库类型
props	Properties	持久化仓库所需属性

Cluster 持久化配置

类名称: org.apache.shardingsphere.mode.repository.cluster.ClusterPersistRepositoryConfiguration

可配置属性:

名称	数据类型	说明
type	String	持久化仓库类型
namespace	String	注册中心命名空间
server-lists	String	注册中心连接地址
props	Properties	持久化仓库所需属性

注意事项

1. 生产环境建议使用集群模式部署。
2. 集群模式部署推荐使用 ZooKeeper 注册中心。

操作步骤

引入 **Maven** 依赖。

```
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-jdbc-core</artifactId>
  <version>${latest.release.version}</version>
</dependency>
```

注意: 请将 `${latest.release.version}` 更改为实际的版本号。

配置示例

Standalone 运行模式

```
ModeConfiguration modeConfig = createModeConfiguration();
Map<String, DataSource> dataSourceMap = ... // 构建真实数据源
Collection<RuleConfiguration> ruleConfigs = ... // 构建具体规则
Properties props = ... // 构建属性配置
DataSource dataSource = ShardingSphereDataSourceFactory.
createDataSource(databaseName, modeConfig, dataSourceMap, ruleConfigs, props);

private ModeConfiguration createModeConfiguration() {
    return new ModeConfiguration("Standalone", new
    StandalonePersistRepositoryConfiguration("H2", new Properties()), true);
}
```

Cluster 运行模式 (推荐)

```
ModeConfiguration modeConfig = createModeConfiguration();
Map<String, DataSource> dataSourceMap = ... // 构建真实数据源
Collection<RuleConfiguration> ruleConfigs = ... // 构建具体规则
Properties props = ... // 构建属性配置
DataSource dataSource = ShardingSphereDataSourceFactory.
createDataSource(databaseName, modeConfig, dataSourceMap, ruleConfigs, props);

private ModeConfiguration createModeConfiguration() {
    return new ModeConfiguration("Cluster", new
    ClusterPersistRepositoryConfiguration("ZooKeeper", "governance-sharding-db",
    "localhost:2181", new Properties()), true);
}
```

相关参考

- [ZooKeeper 注册中心安装与使用](#)
- 持久化仓库类型的详情，请参见[内置持久化仓库类型列表](#)。

数据源配置

背景信息

ShardingSphere-JDBC 支持所有的数据库 JDBC 驱动和连接池。

本节将介绍，通过 JAVA API 的方式配置数据源。

操作步骤

1. 引入 Maven 依赖

```
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-jdbc-core</artifactId>
  <version>${latest.release.version}</version>
</dependency>
```

注意：请将 `${latest.release.version}` 更改为实际的版本号。

配置示例

```
ModeConfiguration modeConfig = // 构建运行模式
Map<String, DataSource> dataSourceMap = createDataSources();
Collection<RuleConfiguration> ruleConfigs = ... // 构建具体规则
Properties props = ... // 构建属性配置
DataSource dataSource = ShardingSphereDataSourceFactory.
createDataSource(databaseName, modeConfig, dataSourceMap, ruleConfigs, props);

private Map<String, DataSource> createDataSources() {
    Map<String, DataSource> dataSourceMap = new HashMap<>();
    // 配置第 1 个数据源
    HikariDataSource dataSource1 = new HikariDataSource();
    dataSource1.setDriverClassName("com.mysql.jdbc.Driver");
    dataSource1.setJdbcUrl("jdbc:mysql://localhost:3306/ds_1");
    dataSource1.setUsername("root");
    dataSource1.setPassword("");
    dataSourceMap.put("ds_1", dataSource1);

    // 配置第 2 个数据源
    HikariDataSource dataSource2 = new HikariDataSource();
    dataSource2.setDriverClassName("com.mysql.jdbc.Driver");
    dataSource2.setJdbcUrl("jdbc:mysql://localhost:3306/ds_2");
    dataSource2.setUsername("root");
    dataSource2.setPassword("");
}
```

```
dataSourceMap.put("ds_2", dataSource2);
}
```

规则配置

规则是 Apache ShardingSphere 面向可插拔的一部分。本章节是 ShardingSphere-JDBC 的 Java 规则配置参考手册。

数据分片

背景信息

数据分片 Java API 规则配置允许用户直接通过编写 Java 代码的方式，完成 ShardingSphereDataSource 对象的创建，Java API 的配置方式非常灵活，不需要依赖额外的 jar 包就能够集成各种类型的业务系统。

参数解释

配置入口

类名称：org.apache.shardingsphere.sharding.api.config.ShardingRuleConfiguration

可配置属性：

分片表配置

类名称：org.apache.shardingsphere.sharding.api.config.ShardingTableRuleConfiguration

可配置属性：

名称	数据类型	说明	默认值
logicTable	String	分片逻辑表名称	.
actualDataNodes (?)	String	由数据源名 + 表名组成，以小数点分隔。多个表以逗号分隔，支持行表达式	使用已知数据源与逻辑表名称生成数据节点，用于广播表或只分库不分表且所有库的表结构完全一致的情况
databaseShardingStrategy (?)	ShardingStrategyConfiguration	分库策略	使用默认分库策略
tableShardingStrategy (?)	ShardingStrategyConfiguration	分表策略	使用默认分表策略
keyGenerateStrategy (?)	KeyGeneratorConfiguration	自增列生成器	使用默认自增主键生成器

自动分片表配置

类名称: `org.apache.shardingsphere.sharding.api.config.ShardingAutoTableRuleConfiguration`

可配置属性:

名称	数据类型	说明	默认值
<code>logicTable</code>	<code>String</code>	分片逻辑表名称	<code>.</code>
<code>actualDataSources (?)</code>	<code>String</code>	数据源名称, 多个数据源以逗号分隔	使用全部配置的数据源
<code>shardingStrategy (?)</code>	<code>ShardingStrategyConfiguration</code>	分片策略	使用默认分片策略
<code>keyGenerateStrategy (?)</code>	<code>KeyGeneratorConfiguration</code>	自增列生成器	使用默认自增主键生成器

分片策略配置

标准分片策略配置

类名称: `org.apache.shardingsphere.sharding.api.config.strategy.sharding.StandardShardingStrategyConfiguration`

可配置属性:

名称	数据类型	说明
<code>shardingColumn</code>	<code>String</code>	分片列名称
<code>shardingAlgorithmName</code>	<code>String</code>	分片算法名称

复合分片策略配置

类名称: `org.apache.shardingsphere.sharding.api.config.strategy.sharding.ComplexShardingStrategyConfiguration`

可配置属性:

名称	数据类型	说明
<code>shardingColumns</code>	<code>String</code>	分片列名称, 多个列以逗号分隔
<code>shardingAlgorithmName</code>	<code>String</code>	分片算法名称

Hint 分片策略配置

类名称:org.apache.shardingsphere.sharding.api.config.strategy.sharding.HintShardingStrategyConfiguration

可配置属性:

名称	数据类型	说明
shardingAlgorithmName	String	分片算法名称

不分片策略配置

类名称:org.apache.shardingsphere.sharding.api.config.strategy.sharding.NoneShardingStrategyConfiguration

可配置属性: 无

算法类型的详情, 请参见[内置分片算法列表](#)。

分布式序列策略配置

类名称:org.apache.shardingsphere.sharding.api.config.strategy.keygen.KeyGenerateStrategyConfiguration

可配置属性:

名称	数据类型	说明
column	String	分布式序列列名称
keyGeneratorName	String	分布式序列算法名称

算法类型的详情, 请参见[内置分布式序列算法列表](#)。

操作步骤

1. 创建真实数据源映射关系, key 为数据源逻辑名称, value 为 DataSource 对象;
2. 创建分片规则对象 ShardingRuleConfiguration, 并初始化对象中的分片表对象 ShardingTableRuleConfiguration、绑定表集合、广播表集合, 以及数据分片所依赖的分库策略和分表策略等参数;
3. 调用 ShardingSphereDataSourceFactory 对象的 createDataSource 方法, 创建 ShardingSphereDataSource。

配置示例

```

public final class ShardingDatabasesAndTablesConfigurationPrecise implements
ExampleConfiguration {

    @Override
    public DataSource getDataSource() throws SQLException {
        return ShardingSphereDataSourceFactory.
createDataSource(createDataSourceMap(), Collections.
singleton(createShardingRuleConfiguration()), new Properties());
    }

    private ShardingRuleConfiguration createShardingRuleConfiguration() {
        ShardingRuleConfiguration result = new ShardingRuleConfiguration();
        result.getTables().add(getOrderTableRuleConfiguration());
        result.getTables().add(getOrderItemTableRuleConfiguration());
        result.getBindingTableGroups().add("t_order, t_order_item");
        result.getBroadcastTables().add("t_address");
        result.setDefaultDatabaseShardingStrategy(new
StandardShardingStrategyConfiguration("user_id", "inline"));
        result.setDefaultTableShardingStrategy(new
StandardShardingStrategyConfiguration("order_id", "standard_test_tbl"));
        Properties props = new Properties();
        props.setProperty("algorithm-expression", "demo_ds_${user_id % 2}");
        result.getShardingAlgorithms().put("inline", new AlgorithmConfiguration(
"INLINE", props));
        result.getShardingAlgorithms().put("standard_test_tbl", new
AlgorithmConfiguration("STANDARD_TEST_TBL", new Properties()));
        result.getKeyGenerators().put("snowflake", new AlgorithmConfiguration(
"SNOWFLAKE", new Properties()));
        return result;
    }

    private ShardingTableRuleConfiguration getOrderTableRuleConfiguration() {
        ShardingTableRuleConfiguration result = new ShardingTableRuleConfiguration(
"t_order", "demo_ds_${0..1}.t_order_${[0, 1]}");
        result.setKeyGenerateStrategy(new KeyGenerateStrategyConfiguration("order_
id", "snowflake"));
        return result;
    }

    private ShardingTableRuleConfiguration getOrderItemTableRuleConfiguration() {
        ShardingTableRuleConfiguration result = new ShardingTableRuleConfiguration(
"t_order_item", "demo_ds_${0..1}.t_order_item_${[0, 1]}");
        result.setKeyGenerateStrategy(new KeyGenerateStrategyConfiguration("order_
item_id", "snowflake"));
        return result;
    }
}

```

```
private Map<String, DataSource> createDataSourceMap() {  
    Map<String, DataSource> result = new HashMap<>();  
    result.put("demo_ds_0", DataSourceUtil.createDataSource("demo_ds_0"));  
    result.put("demo_ds_1", DataSourceUtil.createDataSource("demo_ds_1"));  
    return result;  
}  
}
```

相关参考

- 核心特性：数据分片
- 开发者指南：数据分片

读写分离

背景信息

Java API 形式配置的读写分离可以方便的适用于各种场景，不依赖额外的 jar 包，用户只需要通过 java 代码构造读写分离数据源便可以使用读写分离功能。

参数解释

配置入口

类名称：org.apache.shardingsphere.readwritesplitting.api.ReadwriteSplittingRuleConfiguration

可配置属性：

名称	数据类型	说明
dataSources (+)	Collection<ReadwriteSplittingDataSourceRuleConfiguration>	读写数据源配置
loadBalancers (*)	Map<String, AlgorithmConfiguration>	从库负载均衡算法配置

主从数据源配置

类名称:org.apache.shardingsphere.readwritesplitting.api.rule.ReadwriteSplittingDataSourceRuleConfiguration

可配置属性:

名称	数据类型	说明	默认值
name	String	读写分离数据源名称	.
staticStrategy	StaticReadwriteSplittingStrategyConfiguration	静态读写分离配置	.
dynamicStrategy	DynamicReadwriteSplittingStrategyConfiguration	动态读写分离配置	.
loadBalancerName (?)	String	读库负载均衡算法名称	轮询负载均衡算法

类名称:org.apache.shardingsphere.readwritesplitting.api.strategy.StaticReadwriteSplittingStrategyConfiguration

可配置属性:

名称	数据类型	说明
writeDataSourceName	String	写库数据源名称
readDataSourceNames	List<String>	读库数据源列表

类名称:org.apache.shardingsphere.readwritesplitting.api.strategy.DynamicReadwriteSplittingStrategyConfiguration

可配置属性:

名称	数据类型	说明	默认值
autoAwareDataSourceName	String	数据库发现的逻辑数据源名称	.
writeDataSourceQueryEnabled (?)	String	读库全部下线, 主库是否承担读流量	true

算法类型的详情, 请参见[内置负载均衡算法列表](#)。查询一致性路由的详情, 请参见[核心特性: 读写分离](#)。

操作步骤

1. 添加读写分离数据源
2. 设置负载均衡算法
3. 使用读写分离数据源

配置示例

```

public DataSource getDataSource() throws SQLException {
    ReadwriteSplittingDataSourceRuleConfiguration dataSourceConfig = new
ReadwriteSplittingDataSourceRuleConfiguration(
        "demo_read_query_ds", new
StaticReadwriteSplittingStrategyConfiguration("demo_write_ds",
        Arrays.asList("demo_read_ds_0", "demo_read_ds_1")), null, "demo_
weight_lb");
    Properties algorithmProps = new Properties();
    algorithmProps.setProperty("demo_read_ds_0", "2");
    algorithmProps.setProperty("demo_read_ds_1", "1");
    Map<String, AlgorithmConfiguration> algorithmConfigMap = new HashMap<>(1);
    algorithmConfigMap.put("demo_weight_lb", new AlgorithmConfiguration("WEIGHT
", algorithmProps));
    ReadwriteSplittingRuleConfiguration ruleConfig = new
ReadwriteSplittingRuleConfiguration(Collections.singleton(dataSourceConfig),
algorithmConfigMap);
    Properties props = new Properties();
    props.setProperty("sql-show", Boolean.TRUE.toString());
    return ShardingSphereDataSourceFactory.
createDataSource(createDataSourceMap(), Collections.singleton(ruleConfig), props);
}

private Map<String, DataSource> createDataSourceMap() {
    Map<String, DataSource> result = new HashMap<>(3, 1);
    result.put("demo_write_ds", DataSourceUtil.createDataSource("demo_write_ds
"));
    result.put("demo_read_ds_0", DataSourceUtil.createDataSource("demo_read_ds_
0"));
    result.put("demo_read_ds_1", DataSourceUtil.createDataSource("demo_read_ds_
1"));
    return result;
}

```

相关参考

- [核心特性：读写分离](#)
- [YAML 配置：读写分离](#)
- [Spring Boot Starter：读写分离](#)
- [Spring 命名空间：读写分离](#)

分布式事务

配置入口

`org.apache.shardingsphere.transaction.config.TransactionRuleConfiguration`

可配置属性：

名称	数据类型	说明
<code>defaultType</code>	String	默认事务类型
<code>providerType (?)</code>	String	事务提供者类型
<code>props (?)</code>	Properties	事务属性配置

高可用

背景信息

通过 Java API 方式构建高可用规则配置。

参数解释

配置入口

类名称：`org.apache.shardingsphere.dbdiscovery.api.config.DatabaseDiscoveryRuleConfiguration` 可配置属性：

名称	数据类型	说明
<code>dataSources (+)</code>	<code>Collection<DatabaseDiscoveryDataSourceRuleConfiguration></code>	数据源配置
<code>discoveryHeartbeats (+)</code>	<code>Map<String, DatabaseDiscoveryHeartBeatConfiguration></code>	监听心跳配置
<code>discoveryTypes (+)</code>	<code>Map<String, AlgorithmConfiguration></code>	数据库发现类型配置

数据源配置

类名称：`org.apache.shardingsphere.dbdiscovery.api.config.rule.DatabaseDiscoveryDataSourceRuleConfiguration`
可配置属性：

名称	数据类型	说明
groupName (+)	String	数据库发现组名称
dataSourceNames (+)	Collection<String>	数据源名称,多个数据源用逗号分隔如:ds_0,ds_1
discoveryHeartbeatName (+)	String	监听心跳名称
discoveryTypeName (+)	String	数据库发现类型名称

监听心跳配置

类名称:org.apache.shardingsphere.dbdiscovery.api.config.rule.DatabaseDiscoveryHeartBeatConfiguration

可配置属性:

名称	数据类型	说明	默认值 *
properties (+)	Properties	监听心跳属性配置, keep-alive-cron 属性配置 cron 表达式, 如: '0/5 * * * * ?'	.

数据库发现类型配置

类名称: org.apache.shardingsphere.infra.config.algorithm.AlgorithmConfiguration

名称	数据类型	说明
type (+)	String	数据库发现类型, 如: MySQL.MGR
props (?)	Properties	数据库发现类型配置, 如 MGR 的 group-name 属性配置

操作步骤

1. 引入 Maven 依赖

```
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-jdbc-core</artifactId>
  <version>${latest.release.version}</version>
</dependency>
```

注意: 请将 `${latest.release.version}` 更改为实际的版本号。

配置示例

```
// 构建数据源 ds_0, ds_1, ds_2
Map<String, DataSource> dataSourceMap = new HashMap<>(3, 1);
dataSourceMap.put("ds_0", createDataSource1("primary_demo_ds"));
dataSourceMap.put("ds_1", createDataSource2("primary_demo_ds"));
dataSourceMap.put("ds_2", createDataSource3("primary_demo_ds"));

DataSource dataSource = ShardingSphereDataSourceFactory.createDataSource("database_
discovery_db", dataSourceMap, Arrays.asList(createDatabaseDiscoveryConfiguration(),
createReadwriteSplittingConfiguration()), null);

private static DatabaseDiscoveryRuleConfiguration
createDatabaseDiscoveryConfiguration() {
    DatabaseDiscoveryDataSourceRuleConfiguration dataSourceRuleConfiguration = new
DatabaseDiscoveryDataSourceRuleConfiguration("readwrite_ds", Arrays.asList("ds_0,
ds_1, ds_2"), "mgr-heartbeat", "mgr");
    return new DatabaseDiscoveryRuleConfiguration(Collections.
singleton(dataSourceRuleConfiguration), createDiscoveryHeartbeats(),
createDiscoveryTypes());
}

private static ReadwriteSplittingRuleConfiguration
createReadwriteSplittingConfiguration() {
    ReadwriteSplittingDataSourceRuleConfiguration dataSourceConfiguration1 = new
ReadwriteSplittingDataSourceRuleConfiguration("replica_ds", new
DynamicReadwriteSplittingStrategyConfiguration("readwrite_ds", true), "");
    return new ReadwriteSplittingRuleConfiguration(Arrays.
asList(dataSourceConfiguration1), Collections.emptyMap());
}

private static Map<String, AlgorithmConfiguration> createDiscoveryTypes() {
    Map<String, AlgorithmConfiguration> discoveryTypes = new HashMap<>(1, 1);
    Properties props = new Properties();
    props.put("group-name", "558edd3c-02ec-11ea-9bb3-080027e39bd2");
    discoveryTypes.put("mgr", new AlgorithmConfiguration("MGR", props));
    return discoveryTypes;
}

private static Map<String, DatabaseDiscoveryHeartBeatConfiguration>
createDiscoveryHeartbeats() {
    Map<String, DatabaseDiscoveryHeartBeatConfiguration>
discoveryHeartBeatConfiguration = new HashMap<>(1, 1);
    Properties props = new Properties();
    props.put("keep-alive-cron", "0/5 * * * * ?");
    discoveryHeartBeatConfiguration.put("mgr-heartbeat", new
DatabaseDiscoveryHeartBeatConfiguration(props));
    return discoveryHeartBeatConfiguration;
}
```

}

相关参考

- [高可用核心特性](#)
- [YAML 配置：高可用配置](#)
- [Spring Boot Starter：高可用配置](#)
- [Spring 命名空间：高可用配置](#)

数据加密

背景信息

数据加密 Java API 规则配置允许用户直接通过编写 Java 代码的方式，完成 ShardingSphereDataSource 对象的创建，Java API 的配置方式非常灵活，不需要依赖额外的 jar 包就能够集成各种类型的业务系统。

参数解释

配置入口

类名称：org.apache.shardingsphere.encrypt.api.config.EncryptRuleConfiguration

可配置属性：

名称	数据类型	说明	默认值
tables (+)	Collection<EncryptTableRuleConfiguration>	加密表规则配置	
encryptors (+)	Map<String, AlgorithmConfiguration>	加解密算法名称和配置	
queryWithCipherColumn (?)	boolean	是否使用加密列进行查询。在有原文列的情况下，可以使用原文列进行查询	true

加密表规则配置

类名称：org.apache.shardingsphere.encrypt.api.config.rule.EncryptTableRuleConfiguration

可配置属性：

名称	数据类型	说明
name	String	表名称
columns (+)	Collection<EncryptColumnRuleConfiguration>	加密列规则配置列表
queryWithCipherColumn (?)	boolean	该表是否使用加密列进行查询

加密列规则配置

类名称：org.apache.shardingsphere.encrypt.api.config.rule.EncryptColumnRuleConfiguration

可配置属性：

名称	数据类型	说明
logicColumn	String	逻辑列名称
cipherColumn	String	密文列名称
assistedQueryColumn (?)	String	查询辅助列名称
plainColumn (?)	String	原文列名称
encryptorName	String	密文列加密算法名称
assistedQueryEncryptorName	String	查询辅助列加密算法名称
queryWithCipherColumn (?)	boolean	该列是否使用加密列进行查询

加解密算法配置

类名称：org.apache.shardingsphere.infra.config.algorithm.AlgorithmConfiguration

可配置属性：

名称	数据类型	说明
name	String	加解密算法名称
type	String	加解密算法类型
properties	Properties	加解密算法属性配置

算法类型的详情，请参见[内置加密算法列表](#)。

操作步骤

1. 创建真实数据源映射关系，key 为数据源逻辑名称，value 为 DataSource 对象；
2. 创建加密规则对象 EncryptRuleConfiguration，并初始化对象中的加密表对象 EncryptTableRuleConfiguration、加密算法等参数；
3. 调用 ShardingSphereDataSourceFactory 对象的 createDataSource 方法，创建 ShardingSphereDataSource。

配置示例

```
public final class EncryptDatabasesConfiguration implements ExampleConfiguration {

    @Override
    public DataSource getDataSource() {
        Properties props = new Properties();
        props.setProperty("aes-key-value", "123456");
        EncryptColumnRuleConfiguration columnConfigAes = new
EncryptColumnRuleConfiguration("username", "username", "", "username_plain", "name_
encryptor", null);
        EncryptColumnRuleConfiguration columnConfigTest = new
EncryptColumnRuleConfiguration("pwd", "pwd", "assisted_query_pwd", "", "pwd_
encryptor", null);
        EncryptTableRuleConfiguration encryptTableRuleConfig = new
EncryptTableRuleConfiguration("t_user", Arrays.asList(columnConfigAes,
columnConfigTest), null);
        Map<String, AlgorithmConfiguration> encryptAlgorithmConfigs = new
LinkedHashMap<>(2, 1);
        encryptAlgorithmConfigs.put("name_encryptor", new AlgorithmConfiguration(
"AES", props));
        encryptAlgorithmConfigs.put("pwd_encryptor", new AlgorithmConfiguration(
"assistedTest", props));
        EncryptRuleConfiguration encryptRuleConfig = new
EncryptRuleConfiguration(Collections.singleton(encryptTableRuleConfig),
encryptAlgorithmConfigs);
        try {
            return ShardingSphereDataSourceFactory.createDataSource(DataSourceUtil.
createDataSource("demo_ds"), Collections.singleton(encryptRuleConfig), props);
        } catch (final SQLException ex) {
            ex.printStackTrace();
            return null;
        }
    }
}
```

相关参考

- [数据加密的核心特性](#)
- [数据加密的开发者指南](#)

影子库

背景信息

如果您只想使用 Java API 方式配置使用 ShardingSphere 影子库功能请参考以下配置。

参数解释

配置入口

类名称: `org.apache.shardingsphere.shadow.api.config.ShadowRuleConfiguration`

可配置属性:

名称	数据类型	说明
<code>dataSources</code>	<code>Map<String, ShadowData SourceConfiguration></code>	影子数据源映射名称和配置
<code>tables</code>	<code>Map<String, ShadowTableConfiguration></code>	影子表名称和配置
<code>shadowAlgorithms</code>	<code>Map<String, AlgorithmConfiguration></code>	影子算法名称和配置
<code>defaultShadowAlgorithmName</code>	<code>String</code>	默认影子算法名称

影子数据源配置

类名称: `org.apache.shardingsphere.shadow.api.config.datasource.ShadowDataSourceConfiguration`

可配置属性:

名称	数据类型	说明
<code>productionDataSourceName</code>	<code>String</code>	生产数据源名称
<code>shadowDataSourceName</code>	<code>String</code>	影子数据源名称

影子表配置

类名称: org.apache.shardingsphere.shadow.api.config.table.ShadowTableConfiguration

可配置属性:

名称	数据类型	说明
dataSourceNames	Collection<String>	影子表关联影子数据源映射名称列表
shadowAlgorithmNames	Collection<String>	影子表关联影子算法名称列表

影子算法配置

类名称: org.apache.shardingsphere.infra.config.algorithm.AlgorithmConfiguration

可配置属性:

名称	数据类型	说明
type	String	影子算法类型
props	Properties	影子算法配置

算法类型的详情, 请参见[内置影子算法列表](#)。

操作步骤

1. 创建生产和影子数据源。
2. 配置影子规则
 - 配置影子数据源
 - 配置影子表
 - 配置影子算法

配置示例

```
public final class ShadowConfiguration {

    @Override
    public DataSource getDataSource() throws SQLException {
        Map<String, DataSource> dataSourceMap = createDataSourceMap();
        return ShardingSphereDataSourceFactory.createDataSource(dataSourceMap,
            createRuleConfigurations(), createShardingSphereProps());
    }

    private Map<String, DataSource> createDataSourceMap() {
        Map<String, DataSource> result = new LinkedHashMap<>();
    }
}
```

```

        result.put("ds", DataSourceUtil.createDataSource("demo_ds"));
        result.put("ds_shadow", DataSourceUtil.createDataSource("shadow_demo_ds"));
        return result;
    }

    private Collection<RuleConfiguration> createRuleConfigurations() {
        Collection<RuleConfiguration> result = new LinkedList<>();
        ShadowRuleConfiguration shadowRule = new ShadowRuleConfiguration();
        shadowRule.setDataSources(createShadowDataSources());
        shadowRule.setTables(createShadowTables());
        shadowRule.setShadowAlgorithms(createShadowAlgorithmConfigurations());
        result.add(shadowRule);
        return result;
    }

    private Map<String, ShadowDataSourceConfiguration> createShadowDataSources() {
        Map<String, ShadowDataSourceConfiguration> result = new LinkedHashMap<>();
        result.put("shadow-data-source", new ShadowDataSourceConfiguration("ds",
"ds_shadow"));
        return result;
    }

    private Map<String, ShadowTableConfiguration> createShadowTables() {
        Map<String, ShadowTableConfiguration> result = new LinkedHashMap<>();
        result.put("t_user", new ShadowTableConfiguration(Collections.
singletonList("shadow-data-source"), createShadowAlgorithmNames()));
        return result;
    }

    private Collection<String> createShadowAlgorithmNames() {
        Collection<String> result = new LinkedList<>();
        result.add("user-id-insert-match-algorithm");
        result.add("simple-hint-algorithm");
        return result;
    }

    private Map<String, AlgorithmConfiguration>
createShadowAlgorithmConfigurations() {
        Map<String, AlgorithmConfiguration> result = new LinkedHashMap<>();
        Properties userIdInsertProps = new Properties();
        userIdInsertProps.setProperty("operation", "insert");
        userIdInsertProps.setProperty("column", "user_type");
        userIdInsertProps.setProperty("value", "1");
        result.put("user-id-insert-match-algorithm", new AlgorithmConfiguration(
"VALUE_MATCH", userIdInsertProps));
        return result;
    }
}

```

相关参考

影子库的特性描述

SQL 解析

背景信息

SQL 是使用者与数据库交流的标准语言。SQL 解析引擎负责将 SQL 字符串解析为抽象语法树，供 Apache ShardingSphere 理解并实现其增量功能。目前支持 MySQL, PostgreSQL, SQLServer, Oracle, openGauss 以及符合 SQL92 规范的 SQL 方言。由于 SQL 语法的复杂性，目前仍然存在少量不支持的 SQL。通过 Java API 形式使用 SQL 解析，可以方便得集成进入各种系统，灵活定制用户需求。

参数解释

类名称：org.apache.shardingsphere.parser.config.SQLParserRuleConfiguration

可配置属性：

名称	数据类型	说明
sqlCommentParseEnabled (?)	boolean	是否解析 SQL 注释
parseTreeCache (?)	CacheOption	解析语法树本地缓存配置
sqlStatementCache (?)	CacheOption	SQL 语句本地缓存配置

本地缓存配置

类名称：org.apache.shardingsphere.sql.parser.api.CacheOption

可配置属性：

名称	数据类型	说明	默认值
initialCapacity	int	本地缓存初始容量	语法树本地缓存默认值 128, SQL 语句缓存默认值 2000
maximumSize	long	本地缓存最大容量	语法树本地缓存默认值 1024, SQL 语句缓存默认值 65535

操作步骤

1. 设置本地缓存配置
2. 设置解析配置
3. 使用解析引擎解析 SQL

配置示例

```
CacheOption cacheOption = new CacheOption(128, 1024L);
SQLParserEngine parserEngine = new SQLParserEngine("MySQL", cacheOption);
ParseASTNode parseASTNode = parserEngine.parse("SELECT t.id, t.name, t.age FROM table1 AS t ORDER BY t.id DESC;", false);
SQLVisitorEngine visitorEngine = new SQLVisitorEngine("MySQL", "STATEMENT", false, new Properties());
MySQLStatement sqlStatement = visitorEngine.visit(parseASTNode);
System.out.println(sqlStatement.toString());
```

相关参考

- [YAML 配置：SQL 解析](#)
- [Spring Boot Starter：SQL 解析](#)
- [Spring 命名空间：SQL 解析](#)

SQL 翻译

配置入口

类名称：org.apache.shardingsphere.sqltranslator.api.config.SQLTranslatorRuleConfiguration

可配置属性：

名称	数据类型	说明
type	String	SQL 翻译器类型
useOriginalSQLWhenTranslatingFailed (?)	boolean	SQL 翻译失败是否使用原始 SQL 继续执行

混合规则

背景信息

ShardingSphere 涵盖了很多功能，例如，分库分片、读写分离、高可用、数据脱敏等。这些功能用户可以单独进行使用，也可以配合一起使用，下面是基于 JAVA API 的配置示例。

配置示例

```
// 分片配置
private ShardingRuleConfiguration createShardingRuleConfiguration() {
    ShardingRuleConfiguration result = new ShardingRuleConfiguration();
    result.getTables().add(getOrderTableRuleConfiguration());
    result.setDefaultDatabaseShardingStrategy(new
StandardShardingStrategyConfiguration("user_id", "inline"));
    result.setDefaultTableShardingStrategy(new
StandardShardingStrategyConfiguration("order_id", "standard_test_tbl"));
    Properties props = new Properties();
    props.setProperty("algorithm-expression", "demo_ds_${user_id % 2}");
    result.getShardingAlgorithms().put("inline", new AlgorithmConfiguration("INLINE
", props));
    result.getShardingAlgorithms().put("standard_test_tbl", new
AlgorithmConfiguration("STANDARD_TEST_TBL", new Properties()));
    result.getKeyGenerators().put("snowflake", new AlgorithmConfiguration(
"SNOWFLAKE", new Properties()));
    return result;
}

private ShardingTableRuleConfiguration getOrderTableRuleConfiguration() {
    ShardingTableRuleConfiguration result = new ShardingTableRuleConfiguration("t_
order", "demo_ds_${0..1}.t_order_${[0, 1]}");
    result.setKeyGenerateStrategy(new KeyGenerateStrategyConfiguration("order_id",
"snowflake"));
    return result;
}

// 动态读写分离配置
private static ReadwriteSplittingRuleConfiguration
createReadwriteSplittingConfiguration() {
    ReadwriteSplittingDataSourceRuleConfiguration dataSourceConfiguration1 = new
ReadwriteSplittingDataSourceRuleConfiguration("replica_ds_0", new
DynamicReadwriteSplittingStrategyConfiguration("readwrite_ds_0", true), "");
    ReadwriteSplittingDataSourceRuleConfiguration dataSourceConfiguration2 = new
ReadwriteSplittingDataSourceRuleConfiguration("replica_ds_1", new
DynamicReadwriteSplittingStrategyConfiguration("readwrite_ds_1", true), "");
    Collection<ReadwriteSplittingDataSourceRuleConfiguration> dataSources = new
LinkedList<>();
```

```

        dataSources.add(dataSourceRuleConfiguration1);
        dataSources.add(dataSourceRuleConfiguration2);
        return new ReadwriteSplittingRuleConfiguration(dataSources, Collections.
emptyMap());
    }

// 数据库发现配置
private static DatabaseDiscoveryRuleConfiguration
createDatabaseDiscoveryConfiguration() {
    DatabaseDiscoveryDataSourceRuleConfiguration dataSourceRuleConfiguration1 = new
DatabaseDiscoveryDataSourceRuleConfiguration("readwrite_ds_0", Arrays.asList("ds_0,
ds_1, ds_2"), "mgr-heartbeat", "mgr");
    DatabaseDiscoveryDataSourceRuleConfiguration dataSourceRuleConfiguration2 = new
DatabaseDiscoveryDataSourceRuleConfiguration("readwrite_ds_1", Arrays.asList("ds_3,
ds_4, ds_5"), "mgr-heartbeat", "mgr");
    Collection<DatabaseDiscoveryDataSourceRuleConfiguration> dataSources = new
LinkedList<>();
    dataSources.add(dataSourceRuleConfiguration1);
    dataSources.add(dataSourceRuleConfiguration2);
    return new DatabaseDiscoveryRuleConfiguration(configs,
createDiscoveryHeartbeats(), createDiscoveryTypes());
}

private static DatabaseDiscoveryRuleConfiguration
createDatabaseDiscoveryConfiguration() {
    DatabaseDiscoveryDataSourceRuleConfiguration dataSourceRuleConfiguration = new
DatabaseDiscoveryDataSourceRuleConfiguration("readwrite_ds_1", Arrays.asList("ds_3,
ds_4, ds_5"), "mgr-heartbeat", "mgr");
    return new DatabaseDiscoveryRuleConfiguration(Collections.
singleton(dataSourceRuleConfiguration), createDiscoveryHeartbeats(),
createDiscoveryTypes());
}

private static Map<String, AlgorithmConfiguration> createDiscoveryTypes() {
    Map<String, AlgorithmConfiguration> result = new HashMap<>(1, 1);
    Properties props = new Properties();
    props.put("group-name", "558edd3c-02ec-11ea-9bb3-080027e39bd2");
    discoveryTypes.put("mgr", new AlgorithmConfiguration("MGR", props));
    return result;
}

private static Map<String, DatabaseDiscoveryHeartBeatConfiguration>
createDiscoveryHeartbeats() {
    Map<String, DatabaseDiscoveryHeartBeatConfiguration> result = new HashMap<>(1,
1);
    Properties props = new Properties();
    props.put("keep-alive-cron", "0/5 * * * * ?");
    discoveryHeartBeatConfiguration.put("mgr-heartbeat", new

```

```

DatabaseDiscoveryHeartBeatConfiguration(props));
    return result;
}

// 数据脱敏配置
public EncryptRuleConfiguration createEncryptRuleConfiguration() {
    Properties props = new Properties();
    props.setProperty("aes-key-value", "123456");
    EncryptColumnRuleConfiguration columnConfigAes = new
EncryptColumnRuleConfiguration("username", "username", "", "username_plain", "name_
encryptor", null);
    EncryptColumnRuleConfiguration columnConfigTest = new
EncryptColumnRuleConfiguration("pwd", "pwd", "assisted_query_pwd", "", "pwd_
encryptor", null);
    EncryptTableRuleConfiguration encryptTableRuleConfig = new
EncryptTableRuleConfiguration("t_user", Arrays.asList(columnConfigAes,
columnConfigTest), null);
    Map<String, AlgorithmConfiguration> encryptAlgorithmConfigs = new LinkedHashMap
<>(2, 1);
    encryptAlgorithmConfigs.put("name_encryptor", new AlgorithmConfiguration("AES",
props));
    encryptAlgorithmConfigs.put("pwd_encryptor", new AlgorithmConfiguration(
"assistedTest", props));
    EncryptRuleConfiguration result = new EncryptRuleConfiguration(Collections.
singleton(encryptTableRuleConfig), encryptAlgorithmConfigs);
    return result;
}

```

算法配置

分片算法

```

ShardingRuleConfiguration ruleConfiguration = new ShardingRuleConfiguration();
// algorithmName 由用户指定, 需要和分片策略中的分片算法一致
// type 和 props, 请参考分片内置算法: https://shardingsphere.apache.org/document/current/cn/user-manual/common-config/builtin-algorithm/sharding/
ruleConfiguration.getShardingAlgorithms().put("algorithmName", new
AlgorithmConfiguration("xxx", new Properties()));

```

加密算法

```
// encryptorName 由用户指定, 需要和加密规则中的 encryptorName 属性一致
// type 和 props, 请参考加密内置算法: https://shardingsphere.apache.org/document/current/cn/user-manual/common-config/builtin-algorithm/encrypt/
Map<String, AlgorithmConfiguration> algorithmConfigs = new LinkedHashMap<>(1, 1);
algorithmConfigs.put("encryptorName", new AlgorithmConfiguration("xxx", new Properties()));
```

读写分离负载均衡算法

```
// loadBalancerName 由用户指定, 需要和读写分离规则中的 loadBalancerName 属性一致
// type 和 props, 请参考读写分离负载均衡内置算法: https://shardingsphere.apache.org/document/current/cn/user-manual/common-config/builtin-algorithm/load-balance/
Map<String, AlgorithmConfiguration> algorithmConfigs = new HashMap<>(1, 1);
algorithmConfigs.put("loadBalancerName", new AlgorithmConfiguration("xxx", new Properties()));
```

影子算法

```
// shadowAlgorithmName 由用户指定, 需要和影子库规则中的 shadowAlgorithmNames 属性一致
// type 和 props, 请参考影子库内置算法: https://shardingsphere.apache.org/document/current/cn/user-manual/common-config/builtin-algorithm/shadow/
Map<String, AlgorithmConfiguration> algorithmConfigs = new HashMap<>(1, 1);
algorithmConfigs.put("shadowAlgorithmName", new AlgorithmConfiguration("xxx", new Properties()));
```

高可用

```
// discoveryTypeName 由用户指定, 需要和数据库发现规则中的 discoveryTypeName 属性一致
Map<String, AlgorithmConfiguration> algorithmConfigs = new HashMap<>(1, 1);
algorithmConfigs.put("discoveryTypeName", new AlgorithmConfiguration("xxx", new Properties()));
```

4.1.3 Spring Boot Starter

简介

ShardingSphere-JDBC 提供官方的 Spring Boot Starter，使开发者可以非常便捷的整合 ShardingSphere-JDBC 和 Spring Boot。

使用步骤

引入 Maven 依赖

```
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-jdbc-core-spring-boot-starter</artifactId>
  <version>${shardingsphere.version}</version>
</dependency>
```

配置 Spring Boot 属性

ShardingSphere-JDBC 的 Spring Boot 属性配置由 Database 名称、运行模式、数据源集合、规则集合以及属性配置组成。

```
# JDBC 逻辑库名称。在集群模式中，使用该参数来联通 ShardingSphere-JDBC 与 ShardingSphere-Proxy。
spring.shardingsphere.database.name= # 逻辑库名称，默认值: logic_db
spring.shardingsphere.mode.xxx= # 运行模式
spring.shardingsphere.dataSource.xxx= # 数据源集合
spring.shardingsphere.rules.xxx= # 规则集合
spring.shardingsphere.props= # 属性配置
```

模式详情请参见模式配置。

数据源详情请参见数据源配置。

规则详情请参见规则配置。

使用数据源

直接通过注入的方式即可使用 ShardingSphereDataSource；或者将 ShardingSphereDataSource 配置在 JPA、Hibernate、MyBatis 等 ORM 框架中配合使用。

```
@Resource
private DataSource dataSource;
```

模式配置

参数解释

```
spring.shardingsphere.mode.type= # 运行模式类型。可选配置: Standalone、Cluster
spring.shardingsphere.mode.repository= # 持久化仓库配置。
spring.shardingsphere.mode.override= # 是否使用本地配置覆盖持久化配置
```

单机模式

```
spring.shardingsphere.mode.type=Standalone
spring.shardingsphere.mode.repository.type= # 持久化仓库类型
spring.shardingsphere.mode.repository.props.<key>= # 持久化仓库所需属性
spring.shardingsphere.mode.override= # 是否使用本地配置覆盖持久化配置
```

集群模式 (推荐)

```
spring.shardingsphere.mode.type=Cluster
spring.shardingsphere.mode.repository.type= # 持久化仓库类型
spring.shardingsphere.mode.repository.props.namespace= # 注册中心命名空间
spring.shardingsphere.mode.repository.props.server-lists= # 注册中心连接地址
spring.shardingsphere.mode.repository.props.<key>= # 持久化仓库所需属性
spring.shardingsphere.mode.override= # 是否使用本地配置覆盖持久化配置
```

注意事项

1. 生产环境建议使用集群模式部署。
2. 集群模式部署推荐使用 ZooKeeper 注册中心。

操作步骤

1. 引入 MAVEN 依赖

```
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-jdbc-core-spring-boot-starter</artifactId>
  <version>${latest.release.version}</version>
</dependency>
```

注意: 请将 `${latest.release.version}` 更改为实际的版本号。

配置示例

单机模式

```
spring.shardingsphere.mode.type=Standalone
spring.shardingsphere.mode.repository.type=File
spring.shardingsphere.mode.repository.props.path=.shardingsphere
spring.shardingsphere.mode.override=false
```

集群模式 (推荐)

```
spring.shardingsphere.mode.type=Cluster
spring.shardingsphere.mode.repository.type=ZooKeeper
spring.shardingsphere.mode.repository.props.namespace=governance
spring.shardingsphere.mode.repository.props.server-lists=localhost:2181
spring.shardingsphere.mode.repository.props.retryIntervalMilliseconds=500
spring.shardingsphere.mode.repository.props.timeToLiveSeconds=60
spring.shardingsphere.mode.override=false
```

相关参考

- [ZooKeeper 注册中心安装与使用](#)
- – 持久化仓库类型的详情，请参见[内置持久化仓库类型列表](#)。

数据源配置

背景信息

使用本地数据源

示例的数据库驱动为 MySQL，连接池为 HikariCP，可以更换为其他数据库驱动和连接池。当使用 ShardingSphere JDBC 时，JDBC 池的属性名取决于各自 JDBC 池自己的定义，并不由 ShardingSphere 硬定义，相关的处理可以参考类 `org.apache.shardingsphere.infra.datasource.pool.creator.DataSourcePoolCreator`。例如对于 Alibaba Druid 1.2.9 而言，使用 `url` 代替如下示例中的 `jdbc-url` 是预期行为。

使用 JNDI 数据源

如果计划使用 JNDI 配置数据库，在应用容器（如 Tomcat）中使用 ShardingSphere-JDBC 时，可使用 `spring.shardingsphere.datasource.${datasourceName}.jndiName` 来代替数据源的一系列配置。

参数解释

使用本地数据源

```
spring.shardingsphere.datasource.names= # 真实数据源名称，多个数据源用逗号区分

# <actual-data-source-name> 表示真实数据源名称
spring.shardingsphere.datasource.<actual-data-source-name>.type= # 数据库连接池全类名
spring.shardingsphere.datasource.<actual-data-source-name>.driver-class-name= # 数据库驱动类名，以数据库连接池自身配置为准
spring.shardingsphere.datasource.<actual-data-source-name>.jdbc-url= # 数据库 URL 连接，以数据库连接池自身配置为准
spring.shardingsphere.datasource.<actual-data-source-name>.username= # 数据库用户名，以数据库连接池自身配置为准
spring.shardingsphere.datasource.<actual-data-source-name>.password= # 数据库密码，以数据库连接池自身配置为准
spring.shardingsphere.datasource.<actual-data-source-name>.<xxx>= # ... 数据库连接池的其它属性
```

使用 JNDI 数据源

```
spring.shardingsphere.datasource.names= # 真实数据源名称，多个数据源用逗号区分
# <actual-data-source-name> 表示真实数据源名称
spring.shardingsphere.datasource.<actual-data-source-name>.jndi-name= # 数据源 JNDI
```

操作步骤

1. 引入 MAVEN 依赖

```
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-jdbc-core-spring-boot-starter</artifactId>
  <version>${latest.release.version}</version>
</dependency>
```

注意：请将 `${latest.release.version}` 更改为实际的版本号。

配置示例

使用本地数据源

```
# 配置真实数据源
spring.shardingsphere.datasource.names=ds1,ds2

# 配置第 1 个数据源
spring.shardingsphere.datasource.ds1.type=com.zaxxer.hikari.HikariDataSource
spring.shardingsphere.datasource.ds1.driver-class-name=com.mysql.jdbc.Driver
spring.shardingsphere.datasource.ds1.jdbc-url=jdbc:mysql://localhost:3306/ds1
spring.shardingsphere.datasource.ds1.username=root
spring.shardingsphere.datasource.ds1.password=

# 配置第 2 个数据源
spring.shardingsphere.datasource.ds2.type=com.zaxxer.hikari.HikariDataSource
spring.shardingsphere.datasource.ds2.driver-class-name=com.mysql.jdbc.Driver
spring.shardingsphere.datasource.ds2.jdbc-url=jdbc:mysql://localhost:3306/ds2
spring.shardingsphere.datasource.ds2.username=root
spring.shardingsphere.datasource.ds2.password=
```

使用 JNDI 数据源

```
# 配置真实数据源
spring.shardingsphere.datasource.names=ds1,ds2

# 配置第 1 个数据源
spring.shardingsphere.datasource.ds1.jndi-name=java:comp/env/jdbc/ds1

# 配置第 2 个数据源
spring.shardingsphere.datasource.ds2.jndi-name=java:comp/env/jdbc/ds2
```

规则配置

规则是 Apache ShardingSphere 面向可插拔的一部分。本章节是 ShardingSphere-JDBC 的 Spring Boot Starter 规则配置参考手册。

数据分片

背景信息

数据分片 Spring Boot Starter 配置方式适用于使用 SpringBoot 的业务场景，能够最大程度地利用 Spring-Boot 配置初始化及 Bean 管理的能力，完成 ShardingSphereDataSource 对象的创建，减少不必要的编码工作。

参数解释

```

spring.shardingsphere.datasource.names= # 省略数据源配置, 请参考使用手册

# 标准分片表配置
spring.shardingsphere.rules.sharding.tables.<table-name>.actual-data-nodes= # 由数据源名 + 表名组成, 以小数点分隔。多个表以逗号分隔, 支持 inline 表达式。缺省表示使用已知数据源与逻辑表名称生成数据节点, 用于广播表 (即每个库中都需要一个同样的表用于关联查询, 多为字典表) 或只分库不分表且所有库的表结构完全一致的情况

# 分库策略, 缺省表示使用默认分库策略, 以下的分片策略只能选其一

# 用于单分片键的标准分片场景
spring.shardingsphere.rules.sharding.tables.<table-name>.database-strategy.standard.sharding-column= # 分片列名称
spring.shardingsphere.rules.sharding.tables.<table-name>.database-strategy.standard.sharding-algorithm-name= # 分片算法名称

# 用于多分片键的复合分片场景
spring.shardingsphere.rules.sharding.tables.<table-name>.database-strategy.complex.sharding-columns= # 分片列名称, 多个列以逗号分隔
spring.shardingsphere.rules.sharding.tables.<table-name>.database-strategy.complex.sharding-algorithm-name= # 分片算法名称

# 用于 Hint 的分片策略
spring.shardingsphere.rules.sharding.tables.<table-name>.database-strategy.hint.sharding-algorithm-name= # 分片算法名称

# 分表策略, 同分库策略
spring.shardingsphere.rules.sharding.tables.<table-name>.table-strategy.xxx= # 省略

# 自动分片表配置
spring.shardingsphere.rules.sharding.auto-tables.<auto-table-name>.actual-data-sources= # 数据源名

spring.shardingsphere.rules.sharding.auto-tables.<auto-table-name>.sharding-strategy.standard.sharding-column= # 分片列名称
spring.shardingsphere.rules.sharding.auto-tables.<auto-table-name>.sharding-strategy.standard.sharding-algorithm-name= # 自动分片算法名称

# 分布式序列策略配置
spring.shardingsphere.rules.sharding.tables.<table-name>.key-generate-strategy.column= # 分布式序列列名称
spring.shardingsphere.rules.sharding.tables.<table-name>.key-generate-strategy.key-generator-name= # 分布式序列算法名称

spring.shardingsphere.rules.sharding.binding-tables[0]= # 绑定表规则列表
spring.shardingsphere.rules.sharding.binding-tables[1]= # 绑定表规则列表

```

```

spring.shardingsphere.rules.sharding.binding-tables[x]= # 绑定表规则列表

spring.shardingsphere.rules.sharding.broadcast-tables[0]= # 广播表规则列表
spring.shardingsphere.rules.sharding.broadcast-tables[1]= # 广播表规则列表
spring.shardingsphere.rules.sharding.broadcast-tables[x]= # 广播表规则列表

spring.shardingsphere.rules.sharding.default-database-strategy.xxx= # 默认数据库分片策略
spring.shardingsphere.rules.sharding.default-table-strategy.xxx= # 默认表分片策略
spring.shardingsphere.rules.sharding.default-key-generate-strategy.xxx= # 默认分布式序列策略
spring.shardingsphere.rules.sharding.default-sharding-column= # 默认分片列名称

# 分片算法配置
spring.shardingsphere.rules.sharding.sharding-algorithms.<sharding-algorithm-name>.
type= # 分片算法类型
spring.shardingsphere.rules.sharding.sharding-algorithms.<sharding-algorithm-name>.
props.xxx= # 分片算法属性配置

# 分布式序列算法配置
spring.shardingsphere.rules.sharding.key-generators.<key-generate-algorithm-name>.
type= # 分布式序列算法类型
spring.shardingsphere.rules.sharding.key-generators.<key-generate-algorithm-name>.
props.xxx= # 分布式序列算法属性配置

```

算法类型的详情，请参见[内置分片算法列表](#)和[内置分布式序列算法列表](#)。

注意事项：行表达式标识符可以使用 `${...}` 或 `$->{...}`，但前者与 Spring 本身的属性文件占位符冲突，因此在 Spring 环境中使用行表达式标识符建议使用 `$->{...}`。

操作步骤

1. 在 SpringBoot 文件中配置数据分片规则，包含数据源、分片规则、全局属性等配置项；
2. 启动 SpringBoot 程序，会自动加载配置，并初始化 ShardingSphereDataSource。

配置示例

```

spring.shardingsphere.mode.type=Standalone
spring.shardingsphere.mode.repository.type=File
spring.shardingsphere.mode.overwrite=true

spring.shardingsphere.datasource.names=ds-0,ds-1

spring.shardingsphere.datasource.ds-0.jdbc-url=jdbc:mysql://localhost:3306/demo_ds_0?serverTimezone=UTC&useSSL=false&useUnicode=true&characterEncoding=UTF-8
spring.shardingsphere.datasource.ds-0.type=com.zaxxer.hikari.HikariDataSource

```

```

spring.shardingsphere.datasource.ds-0.driver-class-name=com.mysql.jdbc.Driver
spring.shardingsphere.datasource.ds-0.username=root
spring.shardingsphere.datasource.ds-0.password=

spring.shardingsphere.datasource.ds-1.jdbc-url=jdbc:mysql://localhost:3306/demo_ds_
1?serverTimezone=UTC&useSSL=false&useUnicode=true&characterEncoding=UTF-8
spring.shardingsphere.datasource.ds-1.type=com.zaxxer.hikari.HikariDataSource
spring.shardingsphere.datasource.ds-1.driver-class-name=com.mysql.jdbc.Driver
spring.shardingsphere.datasource.ds-1.username=root
spring.shardingsphere.datasource.ds-1.password=

spring.shardingsphere.rules.sharding.default-database-strategy.standard.sharding-
column=user_id
spring.shardingsphere.rules.sharding.default-database-strategy.standard.sharding-
algorithm-name=database-inline
spring.shardingsphere.rules.sharding.binding-tables[0]=t_order,t_order_item
spring.shardingsphere.rules.sharding.broadcast-tables=t_address

spring.shardingsphere.rules.sharding.tables.t_order.actual-data-nodes=ds-${0..1}.
t_order_${0..1}
spring.shardingsphere.rules.sharding.tables.t_order.table-strategy.standard.
sharding-column=order_id
spring.shardingsphere.rules.sharding.tables.t_order.table-strategy.standard.
sharding-algorithm-name=t-order-inline

spring.shardingsphere.rules.sharding.tables.t_order.key-generate-strategy.
column=order_id
spring.shardingsphere.rules.sharding.tables.t_order.key-generate-strategy.key-
generator-name=snowflake

spring.shardingsphere.rules.sharding.tables.t_order_item.actual-data-nodes=ds-${0..1}.t_order_item_${0..1}
spring.shardingsphere.rules.sharding.tables.t_order_item.table-strategy.standard.
sharding-column=order_id
spring.shardingsphere.rules.sharding.tables.t_order_item.table-strategy.standard.
sharding-algorithm-name=t-order-item-inline

spring.shardingsphere.rules.sharding.tables.t_order_item.key-generate-strategy.
column=order_item_id
spring.shardingsphere.rules.sharding.tables.t_order_item.key-generate-strategy.key-
generator-name=snowflake

spring.shardingsphere.rules.sharding.sharding-algorithms.database-inline.
type=INLINE
spring.shardingsphere.rules.sharding.sharding-algorithms.database-inline.props.
algorithm-expression=ds-${0..1}{user_id % 2}
spring.shardingsphere.rules.sharding.sharding-algorithms.t-order-inline.type=INLINE
spring.shardingsphere.rules.sharding.sharding-algorithms.t-order-inline.props.

```

```

algorithm-expression=t_order_${order_id % 2}
spring.shardingsphere.rules.sharding.sharding-algorithms.t-order-item-inline.
type=INLINE
spring.shardingsphere.rules.sharding.sharding-algorithms.t-order-item-inline.props.
algorithm-expression=t_order_item_${order_id % 2}

spring.shardingsphere.rules.sharding.key-generators.snowflake.type=SNOWFLAKE

```

相关参考

- 核心特性：数据分片
- 开发者指南：数据分片

读写分离

背景信息

读写分离 Spring Boot Starter 配置方式适用于使用 SpringBoot 的业务场景，能够最大程度地利用 Spring Boot 配置初始化及 Bean 管理的能力，完成 ShardingSphereDataSource 对象的创建，减少不必要的编码工作。

参数解释

静态读写分离

```

spring.shardingsphere.datasource.names= # 省略数据源配置，请参考使用手册

spring.shardingsphere.rules.readwrite-splitting.data-sources.<readwrite-splitting-
data-source-name>.static-strategy.write-data-source-name= # 写库数据源名称
spring.shardingsphere.rules.readwrite-splitting.data-sources.<readwrite-splitting-
data-source-name>.static-strategy.read-data-source-names= # 读库数据源列表，多个从数据源
用逗号分隔
spring.shardingsphere.rules.readwrite-splitting.data-sources.<readwrite-splitting-
data-source-name>.load-balancer-name= # 负载均衡算法名称

# 负载均衡算法配置
spring.shardingsphere.rules.readwrite-splitting.load-balancers.<load-balance-
algorithm-name>.type= # 负载均衡算法类型
spring.shardingsphere.rules.readwrite-splitting.load-balancers.<load-balance-
algorithm-name>.props.xxx= # 负载均衡算法属性配置

```

动态读写分离

```
spring.shardingsphere.datasource.names= # 省略数据源配置, 请参考使用手册

spring.shardingsphere.rules.readwrite-splitting.data-sources.<readwrite-splitting-
data-source-name>.dynamic-strategy.auto-aware-data-source-name= # 数据库发现逻辑数据源
名称
spring.shardingsphere.rules.readwrite-splitting.data-sources.<readwrite-splitting-
data-source-name>.dynamic-strategy.write-data-source-query-enabled= # 读库全部下线, 主
库是否承担读流量
spring.shardingsphere.rules.readwrite-splitting.data-sources.<readwrite-splitting-
data-source-name>.load-balancer-name= # 负载均衡算法名称

# 负载均衡算法配置
spring.shardingsphere.rules.readwrite-splitting.load-balancers.<load-balance-
algorithm-name>.type= # 负载均衡算法类型
spring.shardingsphere.rules.readwrite-splitting.load-balancers.<load-balance-
algorithm-name>.props.xxx= # 负载均衡算法属性配置
```

算法类型的详情, 请参见[内置负载均衡算法列表](#)。查询一致性路由的详情, 请参见[核心特性: 读写分离](#)。

操作步骤

1. 添加读写分离数据源
2. 设置负载均衡算法
3. 使用读写分离数据源

配置示例

```
spring.shardingsphere.rules.readwrite-splitting.data-sources.readwrite_ds.static-
strategy.write-data-source-name=write-ds
spring.shardingsphere.rules.readwrite-splitting.data-sources.readwrite_ds.static-
strategy.read-data-source-names=read-ds-0,read-ds-1
spring.shardingsphere.rules.readwrite-splitting.data-sources.readwrite_ds.load-
balancer-name=round_robin
spring.shardingsphere.rules.readwrite-splitting.load-balancers.round_robin.
type=ROUND_ROBIN
```

相关参考

- 核心特性：读写分离
- Java API：读写分离
- YAML 配置：读写分离
- Spring 命名空间：读写分离

高可用

背景信息

Spring Boot Starter 配置方式适用于使用 SpringBoot 的业务场景,能够最大程度地利用 SpringBoot 配置初始化及 Bean 管理的能力, 自动完成 ShardingSphereDataSource 对象的创建。

参数解释

```
spring.shardingsphere.datasource.names= # 省略数据源配置, 请参考使用手册

spring.shardingsphere.rules.readwrite-splitting.data-sources.<readwrite-splitting-
data-source-name>.dynamic-strategy.auto-aware-data-source-name= # 数据库发现的逻辑数据
源名称

spring.shardingsphere.rules.database-discovery.data-sources.<database-discovery-
data-source-name>.data-source-names= # 数据源名称, 多个数据源用逗号分隔 如: ds_0, ds_1
spring.shardingsphere.rules.database-discovery.data-sources.<database-discovery-
data-source-name>.discovery-heartbeat-name= # 检测心跳名称
spring.shardingsphere.rules.database-discovery.data-sources.<database-discovery-
data-source-name>.discovery-type-name= # 数据库发现类型名称
spring.shardingsphere.rules.database-discovery.discovery-heartbeats.<discovery-
heartbeat-name>.props.keep-alive-cron= # cron 表达式, 如: '0/5 * * * * ?'
spring.shardingsphere.rules.database-discovery.discovery-types.<discovery-type-
name>.type= # 数据库发现类型, 如: MySQL.MGR
spring.shardingsphere.rules.database-discovery.discovery-types.<discovery-type-
name>.props.group-name= # 数据库发现类型必要参数, 如 MGR 的 group-name
```

操作步骤

1. 引入 MAVEN 依赖

```
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-jdbc-core-spring-boot-starter</artifactId>
  <version>${latest.release.version}</version>
</dependency>
```


注意：请将 `${latest.release.version}` 更改为实际的版本号。

配置示例

```
spring.shardingsphere.datasource.names=ds-0,ds-1,ds-2
spring.shardingsphere.datasource.ds-0.jdbc-url = jdbc:mysql://127.0.0.1:13306/
primary_demo_ds?serverTimezone=UTC&useSSL=false
spring.shardingsphere.datasource.ds-0.username=root
spring.shardingsphere.datasource.ds-0.password=
spring.shardingsphere.datasource.ds-0.type=com.zaxxer.hikari.HikariDataSource
spring.shardingsphere.datasource.ds-0.driver-class-name=com.mysql.cj.jdbc.Driver

spring.shardingsphere.datasource.ds-1.jdbc-url = jdbc:mysql://127.0.0.1:13307/
primary_demo_ds?serverTimezone=UTC&useSSL=false
spring.shardingsphere.datasource.ds-1.username=root
spring.shardingsphere.datasource.ds-1.password=
spring.shardingsphere.datasource.ds-1.type=com.zaxxer.hikari.HikariDataSource
spring.shardingsphere.datasource.ds-1.driver-class-name=com.mysql.cj.jdbc.Driver

spring.shardingsphere.datasource.ds-2.jdbc-url = jdbc:mysql://127.0.0.1:13308/
primary_demo_ds?serverTimezone=UTC&useSSL=false
spring.shardingsphere.datasource.ds-2.username=root
spring.shardingsphere.datasource.ds-2.password=
spring.shardingsphere.datasource.ds-2.type=com.zaxxer.hikari.HikariDataSource
spring.shardingsphere.datasource.ds-2.driver-class-name=com.mysql.cj.jdbc.Driver

spring.shardingsphere.rules.readwrite-splitting.data-sources.replica_ds.dynamic-
strategy.auto-aware-data-source-name=readwrite_ds

spring.shardingsphere.rules.database-discovery.data-sources.readwrite_ds.data-
source-names=ds-0, ds-1, ds-2
spring.shardingsphere.rules.database-discovery.data-sources.readwrite_ds.discovery-
heartbeat-name=mgr-heartbeat
spring.shardingsphere.rules.database-discovery.data-sources.readwrite_ds.discovery-
type-name=mgr
spring.shardingsphere.rules.database-discovery.discovery-heartbeats.mgr-heartbeat.
props.keep-alive-cron=0/5 * * * * ?
spring.shardingsphere.rules.database-discovery.discovery-types.mgr.type=MGR
spring.shardingsphere.rules.database-discovery.discovery-types.mgr.props.
groupName=b13df29e-90b6-11e8-8d1b-525400fc3996
```

相关参考

- 高可用核心特性
- JAVA API：高可用配置
- YAML 配置：高可用配置
- Spring 命名空间：高可用配置

数据加密

背景信息

数据加密 Spring Boot Starter 配置方式适用于使用 SpringBoot 的业务场景，能够最大程度地利用 Spring-Boot 配置初始化及 Bean 管理的能力，完成 ShardingSphereDataSource 对象的创建，减少不必要的编码工作。

参数解释

```
spring.shardingsphere.datasource.names= # 省略数据源配置，请参考使用手册

spring.shardingsphere.rules.encrypt.tables.<table-name>.query-with-cipher-column= #
该表是否使用加密列进行查询
spring.shardingsphere.rules.encrypt.tables.<table-name>.columns.<column-name>.
cipher-column= # 加密列名称
spring.shardingsphere.rules.encrypt.tables.<table-name>.columns.<column-name>.
assisted-query-column= # 查询列名称
spring.shardingsphere.rules.encrypt.tables.<table-name>.columns.<column-name>.
plain-column= # 原文列名称
spring.shardingsphere.rules.encrypt.tables.<table-name>.columns.<column-name>.
encryptor-name= # 加密算法名称

# 加密算法配置
spring.shardingsphere.rules.encrypt.encryptors.<encrypt-algorithm-name>.type= # 加密
算法类型
spring.shardingsphere.rules.encrypt.encryptors.<encrypt-algorithm-name>.props.xxx=
# 加密算法属性配置

spring.shardingsphere.rules.encrypt.queryWithCipherColumn= # 是否使用加密列进行查询。在
有原文列的情况下，可以使用原文列进行查询
```

算法类型的详情，请参见[内置加密算法列表](#)。

操作步骤

1. 在 SpringBoot 文件中配置数据加密规则，包含数据源、加密规则、全局属性等配置项；
2. 启动 SpringBoot 程序，会自动加载配置，并初始化 ShardingSphereDataSource。

配置示例

```
spring.shardingsphere.datasource.names=ds

spring.shardingsphere.datasource.ds.type=com.zaxxer.hikari.HikariDataSource
spring.shardingsphere.datasource.ds.driver-class-name=com.mysql.jdbc.Driver
spring.shardingsphere.datasource.ds.jdbc-url=jdbc:mysql://localhost:3306/demo_ds?
serverTimezone=UTC&useSSL=false&useUnicode=true&characterEncoding=UTF-8
spring.shardingsphere.datasource.ds.username=root
spring.shardingsphere.datasource.ds.password=

spring.shardingsphere.rules.encrypt.encryptors.name-encryptor.type=AES
spring.shardingsphere.rules.encrypt.encryptors.name-encryptor.props.aes-key-
value=123456abc
spring.shardingsphere.rules.encrypt.encryptors.pwd-encryptor.type=AES
spring.shardingsphere.rules.encrypt.encryptors.pwd-encryptor.props.aes-key-
value=123456abc

spring.shardingsphere.rules.encrypt.tables.t_user.columns.username.cipher-
column=username
spring.shardingsphere.rules.encrypt.tables.t_user.columns.username.encryptor-
name=name-encryptor
spring.shardingsphere.rules.encrypt.tables.t_user.columns.pwd.cipher-column=pwd
spring.shardingsphere.rules.encrypt.tables.t_user.columns.pwd.encryptor-name=pwd-
encryptor

spring.shardingsphere.props.query-with-cipher-column=true
spring.shardingsphere.props.sql-show=true
```

相关参考

- 核心特性：数据加密
- 开发者指南：数据加密

影子库

背景信息

如果您想在 Spring Boot 环境中使用 ShardingSphere 影子库功能请参考以下配置。

参数解释

```
spring.shardingsphere.rules.shadow.data-sources.shadow-data-source.production-data-source-name= # 生产数据源名称
spring.shardingsphere.rules.shadow.data-sources.shadow-data-source.shadow-data-source-name= # 影子数据源名称

spring.shardingsphere.rules.shadow.tables.<table-name>.data-source-names= # 影子表关联影子数据源名称列表（多个值用"," 隔开）
spring.shardingsphere.rules.shadow.tables.<table-name>.shadow-algorithm-names= # 影子表关联影子算法名称列表（多个值用"," 隔开）

spring.shardingsphere.rules.shadow.defaultShadowAlgorithmName= # 默认影子算法名称，选配项

spring.shardingsphere.rules.shadow.shadow-algorithms.<shadow-algorithm-name>.type= # 影子算法类型
spring.shardingsphere.rules.shadow.shadow-algorithms.<shadow-algorithm-name>.props.xxx= # 影子算法属性配置
```

详情请参见[内置影子算法列表](#)

操作步骤

1. 在 SpringBoot 文件中配置影子库规则，包含数据源、影子规则、全局属性等配置项。
2. 启动 SpringBoot 程序，会自动加载配置，并初始化 ShardingSphereDataSource。

配置示例

```

spring.shardingsphere.datasource.names=ds,shadow-ds

spring.shardingsphere.datasource.ds.jdbc-url=jdbc:mysql://localhost:3306/ds?
serverTimezone=UTC&useSSL=false&useUnicode=true&characterEncoding=UTF-8
spring.shardingsphere.datasource.ds.type=com.zaxxer.hikari.HikariDataSource
spring.shardingsphere.datasource.ds.driver-class-name=com.mysql.jdbc.Driver
spring.shardingsphere.datasource.ds.username=root
spring.shardingsphere.datasource.ds.password=

spring.shardingsphere.datasource.shadow-ds.jdbc-url=jdbc:mysql://localhost:3306/
shadow_ds?serverTimezone=UTC&useSSL=false&useUnicode=true&characterEncoding=UTF-8
spring.shardingsphere.datasource.shadow-ds.type=com.zaxxer.hikari.HikariDataSource
spring.shardingsphere.datasource.shadow-ds.driver-class-name=com.mysql.jdbc.Driver
spring.shardingsphere.datasource.shadow-ds.username=root
spring.shardingsphere.datasource.shadow-ds.password=

spring.shardingsphere.rules.shadow.data-sources.shadow-data-source.production-data-
source-name=ds
spring.shardingsphere.rules.shadow.data-sources.shadow-data-source.shadow-data-
source-name=shadow-ds

spring.shardingsphere.rules.shadow.tables.t_user.data-source-names=shadow-data-
source
spring.shardingsphere.rules.shadow.tables.t_user.shadow-algorithm-names=user-id-
insert-match-algorithm,simple-hint-algorithm

spring.shardingsphere.rules.shadow.shadow-algorithms.user-id-insert-match-
algorithm.type=VALUE_MATCH
spring.shardingsphere.rules.shadow.shadow-algorithms.user-id-insert-match-
algorithm.props.operation=insert
spring.shardingsphere.rules.shadow.shadow-algorithms.user-id-insert-match-
algorithm.props.column=user_id
spring.shardingsphere.rules.shadow.shadow-algorithms.user-id-insert-match-
algorithm.props.value=1

spring.shardingsphere.rules.shadow.shadow-algorithms.simple-hint-algorithm.
type=SIMPLE_HINT
spring.shardingsphere.rules.shadow.shadow-algorithms.simple-hint-algorithm.props.
shadow=true
spring.shardingsphere.rules.shadow.shadow-algorithms.simple-hint-algorithm.props.
foo=bar

```

相关参考

- [影子库的特性描述](#)
- [JAVA API：影子库的配置](#)
- [YAML 配置：影子库的配置](#)
- [Spring 命名空间：影子库的配置](#)
- [开发者指南：影子库的接口和示例](#)

SQL 解析

背景信息

Spring Boot Starter 的配置方式适用于使用 SpringBoot 的业务场景。使用这种方式，能够最大程度地利用 SpringBoot 配置初始化以及 Bean 管理的能力，从而达到简化代码开发的目的。

参数解释

```
spring.shardingsphere.rules.sql-parser.sql-comment-parse-enabled= # 是否解析 SQL 注释  
  
spring.shardingsphere.rules.sql-parser.sql-statement-cache.initial-capacity= # SQL  
语句本地缓存初始容量  
spring.shardingsphere.rules.sql-parser.sql-statement-cache.maximum-size= # SQL 语句  
本地缓存最大容量  
  
spring.shardingsphere.rules.sql-parser.parse-tree-cache.initial-capacity= # 解析树本  
地缓存初始容量  
spring.shardingsphere.rules.sql-parser.parse-tree-cache.maximum-size= # 解析树本地缓存  
最大容量
```

操作步骤

1. 设置本地缓存配置
2. 设置解析配置
3. 使用解析引擎解析 SQL

配置示例

```
spring.shardingsphere.rules.sql-parser.sql-comment-parse-enabled=true

spring.shardingsphere.rules.sql-parser.sql-statement-cache.initial-capacity=2000
spring.shardingsphere.rules.sql-parser.sql-statement-cache.maximum-size=65535

spring.shardingsphere.rules.sql-parser.parse-tree-cache.initial-capacity=128
spring.shardingsphere.rules.sql-parser.parse-tree-cache.maximum-size=1024
```

相关参考

- [JAVA API: SQL 解析](#)
- [YAML 配置: SQL 解析](#)
- [Spring 命名空间: SQL 解析](#)

混合规则

背景信息

ShardingSphere 涵盖了很多功能，例如，分库分片、读写分离、高可用、数据脱敏等。这些功能用户可以单独进行使用，也可以配合一起使用，下面是基于 SpringBoot Starter 的参数解释和配置示例。

参数解释

```
spring.shardingsphere.datasource.names= # 省略数据源配置，请参考使用手册
# 标准分片表配置
spring.shardingsphere.rules.sharding.tables.<table-name>.actual-data-nodes= # 由数据源名 + 表名组成，以小数点分隔。多个表以逗号分隔，支持 inline 表达式。缺省表示使用已知数据源与逻辑表名称生成数据节点，用于广播表（即每个库中都需要一个同样的表用于关联查询，多为字典表）或只分库不分表且所有库的表结构完全一致的情况
# 用于单分片键的标准分片场景
spring.shardingsphere.rules.sharding.tables.<table-name>.database-strategy.standard.sharding-column= # 分片列名称
spring.shardingsphere.rules.sharding.tables.<table-name>.database-strategy.standard.sharding-algorithm-name= # 分片算法名称
# 分表策略，同分库策略
spring.shardingsphere.rules.sharding.tables.<table-name>.table-strategy.xxx= # 省略
# 分布式序列策略配置
spring.shardingsphere.rules.sharding.tables.<table-name>.key-generate-strategy.column= # 分布式序列列名称
spring.shardingsphere.rules.sharding.tables.<table-name>.key-generate-strategy.key-generator-name= # 分布式序列算法名称
# 分片算法配置
```

```

spring.shardingsphere.rules.sharding.sharding-algorithms.<sharding-algorithm-name>.
type= # 分片算法类型
spring.shardingsphere.rules.sharding.sharding-algorithms.<sharding-algorithm-name>.
props.xxx= # 分片算法属性配置
# 分布式序列算法配置
spring.shardingsphere.rules.sharding.key-generators.<key-generate-algorithm-name>.
type= # 分布式序列算法类型
spring.shardingsphere.rules.sharding.key-generators.<key-generate-algorithm-name>.
props.xxx= # 分布式序列算法属性配置
# 动态读写分离配置
spring.shardingsphere.rules.readwrite-splitting.data-sources.<readwrite-splitting-
data-source-name>.dynamic-strategy.auto-aware-data-source-name= # 数据库发现逻辑数据源
名称
spring.shardingsphere.rules.readwrite-splitting.data-sources.<readwrite-splitting-
data-source-name>.dynamic-strategy.write-data-source-query-enabled= # 读库全部下线，主
库是否承担读流量
spring.shardingsphere.rules.readwrite-splitting.data-sources.<readwrite-splitting-
data-source-name>.load-balancer-name= # 负载均衡算法名称
# 数据库发现配置
spring.shardingsphere.rules.database-discovery.data-sources.<database-discovery-
data-source-name>.data-source-names= # 数据源名称，多个数据源用逗号分隔 如: ds_0, ds_1
spring.shardingsphere.rules.database-discovery.data-sources.<database-discovery-
data-source-name>.discovery-heartbeat-name= # 检测心跳名称
spring.shardingsphere.rules.database-discovery.data-sources.<database-discovery-
data-source-name>.discovery-type-name= # 数据库发现类型名称
spring.shardingsphere.rules.database-discovery.discovery-heartbeats.<discovery-
heartbeat-name>.props.keep-alive-cron= # cron 表达式，如: '0/5 * * * * ?'
spring.shardingsphere.rules.database-discovery.discovery-types.<discovery-type-
name>.type= # 数据库发现类型，如: MySQL.MGR
spring.shardingsphere.rules.database-discovery.discovery-types.<discovery-type-
name>.props.group-name= # 数据库发现类型必要参数，如 MGR 的 group-name
# 数据脱敏配置
spring.shardingsphere.rules.encrypt.tables.<table-name>.query-with-cipher-column= #
该表是否使用加密列进行查询
spring.shardingsphere.rules.encrypt.tables.<table-name>.columns.<column-name>.
cipher-column= # 加密列名称
spring.shardingsphere.rules.encrypt.tables.<table-name>.columns.<column-name>.
assisted-query-column= # 查询列名称
spring.shardingsphere.rules.encrypt.tables.<table-name>.columns.<column-name>.
plain-column= # 原文列名称
spring.shardingsphere.rules.encrypt.tables.<table-name>.columns.<column-name>.
encryptor-name= # 加密算法名称
# 加密算法配置
spring.shardingsphere.rules.encrypt.encryptors.<encrypt-algorithm-name>.type= # 加密
算法类型
spring.shardingsphere.rules.encrypt.encryptors.<encrypt-algorithm-name>.props.xxx=
# 加密算法属性配置
spring.shardingsphere.rules.encrypt.queryWithCipherColumn= # 是否使用加密列进行查询。在

```


有原文列的情况下，可以使用原文列进行查询

配置示例

```
# 分片配置
spring.shardingsphere.rules.sharding.tables.t_order.actual-data-nodes=replica-ds-${0..1}.t_order_${0..1}
spring.shardingsphere.rules.sharding.tables.t_order.table-strategy.standard.sharding-column=order_id
spring.shardingsphere.rules.sharding.tables.t_order.table-strategy.standard.sharding-algorithm-name=t-order-inline
spring.shardingsphere.rules.sharding.tables.t_order.key-generate-strategy.column=order_id
spring.shardingsphere.rules.sharding.tables.t_order.key-generate-strategy.key-generator-name=snowflake
spring.shardingsphere.rules.sharding.tables.t_order_item.actual-data-nodes=replica-ds-${0..1}.t_order_item_${0..1}
spring.shardingsphere.rules.sharding.tables.t_order_item.table-strategy.standard.sharding-column=order_id
spring.shardingsphere.rules.sharding.sharding-algorithms.database-inline.type=INLINE
spring.shardingsphere.rules.sharding.sharding-algorithms.database-inline.props.algorithm-expression=replica-ds-${user_id % 2}
spring.shardingsphere.rules.sharding.sharding-algorithms.t-order-inline.type=INLINE
spring.shardingsphere.rules.sharding.sharding-algorithms.t-order-inline.props.algorithm-expression=t_order_${order_id % 2}
spring.shardingsphere.rules.sharding.key-generators.snowflake.type=SNOWFLAKE
# 动态读写分离配置
spring.shardingsphere.rules.readwrite-splitting.data-sources.replica-ds-0.dynamic-strategy.auto-aware-data-source-name=readwrite-ds-0
spring.shardingsphere.rules.readwrite-splitting.data-sources.replica-ds-1.dynamic-strategy.auto-aware-data-source-name=readwrite-ds-1
# 数据库发现配置
spring.shardingsphere.rules.database-discovery.data-sources.readwrite-ds-0.data-source-names=ds-0, ds-1, ds-2
spring.shardingsphere.rules.database-discovery.data-sources.readwrite-ds-0.discovery-heartbeat-name=mgr-heartbeat
spring.shardingsphere.rules.database-discovery.data-sources.readwrite-ds-0.discovery-type-name=mgr
spring.shardingsphere.rules.database-discovery.data-sources.readwrite-ds-1.data-source-names=ds-3, ds-4, ds-5
spring.shardingsphere.rules.database-discovery.data-sources.readwrite-ds-1.discovery-heartbeat-name=mgr-heartbeat
spring.shardingsphere.rules.database-discovery.data-sources.readwrite-ds-1.discovery-type-name=mgr
spring.shardingsphere.rules.database-discovery.discovery-heartbeats.mgr-heartbeat.
```

```

props.keep-alive-cron=0/5 * * * * ?
spring.shardingsphere.rules.database-discovery.discovery-types.mgr.type=MGR
spring.shardingsphere.rules.database-discovery.discovery-types.mgr.props.
groupName=b13df29e-90b6-11e8-8d1b-525400fc3996
# 数据脱敏配置
spring.shardingsphere.rules.encrypt.encryptors.name-encryptor.type=AES
spring.shardingsphere.rules.encrypt.encryptors.name-encryptor.props.aes-key-
value=123456abc
spring.shardingsphere.rules.encrypt.encryptors.pwd-encryptor.type=AES
spring.shardingsphere.rules.encrypt.encryptors.pwd-encryptor.props.aes-key-
value=123456abc
spring.shardingsphere.rules.encrypt.tables.t_user.columns.username.cipher-
column=username
spring.shardingsphere.rules.encrypt.tables.t_user.columns.username.encryptor-
name=name-encryptor
spring.shardingsphere.rules.encrypt.tables.t_user.columns.pwd.cipher-column=pwd
spring.shardingsphere.rules.encrypt.tables.t_user.columns.pwd.encryptor-name=pwd-
encryptor
spring.shardingsphere.props.query-with-cipher-column=true
spring.shardingsphere.props.sql-show=true

```

算法配置

分片算法

```

# sharding-algorithm-name 由用户指定, 需要和分片策略中的 sharding-algorithm-name 属性一致
# type 和 props, 请参考分片内置算法: https://shardingsphere.apache.org/document/current/cn/user-manual/common-config/builtin-algorithm/sharding/
spring.shardingsphere.rules.sharding.sharding-algorithms.<sharding-algorithm-name>.
type=xxx
spring.shardingsphere.rules.sharding.sharding-algorithms.<sharding-algorithm-name>.
props.xxx=xxx

```

加密算法

```

# encrypt-algorithm-name 由用户指定, 需要和加密规则中的 encryptor-name 属性一致
# type 和 props, 请参考加密内置算法: https://shardingsphere.apache.org/document/current/cn/user-manual/common-config/builtin-algorithm/encrypt/
spring.shardingsphere.rules.encrypt.encryptors.<encrypt-algorithm-name>.type=xxx
spring.shardingsphere.rules.encrypt.encryptors.<encrypt-algorithm-name>.props.
xxx=xxx

```

读写分离负载均衡算法

```
# load-balance-algorithm-name 由用户指定，需要和读写分离规则中的 load-balancer-name 属性一致
# type 和 props，请参考读写分离负载均衡内置算法: https://shardingsphere.apache.org/document/current/cn/user-manual/common-config/builtin-algorithm/load-balance/
spring.shardingsphere.rules.readwrite-splitting.load-balancers.<load-balance-algorithm-name>.type=xxx
spring.shardingsphere.rules.readwrite-splitting.load-balancers.<load-balance-algorithm-name>.props.xxx=xxx
```

影子算法

```
# shadow-algorithm-name 由用户指定，需要和影子库规则中的 shadow-algorithm-names 属性一致
# type 和 props，请参考影子库内置算法: https://shardingsphere.apache.org/document/current/cn/user-manual/common-config/builtin-algorithm/shadow/
spring.shardingsphere.rules.shadow.shadow-algorithms.<shadow-algorithm-name>.type=xxx
spring.shardingsphere.rules.shadow.shadow-algorithms.<shadow-algorithm-name>.props.xxx=xxx
```

高可用

```
# discovery-type-name 由用户指定，需要和数据库发现规则中的 discovery-type-name 属性一致
spring.shardingsphere.rules.database-discovery.discovery-types.<discovery-type-name>.type=xxx
spring.shardingsphere.rules.database-discovery.discovery-types.<discovery-type-name>.props.xxx=xxx
```

4.1.4 Spring 命名空间

简介

ShardingSphere-JDBC 提供官方的 Spring 命名空间，使开发者可以非常便捷的整合 ShardingSphere-JDBC 和 Spring。

使用步骤

引入 Maven 依赖

```
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-jdbc-core-spring-namespace</artifactId>
  <version>${shardingsphere.version}</version>
</dependency>
```

配置 Spring Bean

配置项说明

命名空间：<http://shardingsphere.apache.org/schema/shardingsphere/datasource/datasource-5.2.0.xsd>

<shardingsphere:data-source />

名称	类型	说明
id	属性	Spring Bean Id
database-name (?)	属性	JDBC 数据源别名
data-source-names	标签	数据源名称，多个数据源以逗号分隔
rule-refs	标签	规则名称，多个规则以逗号分隔
mode (?)	标签	运行模式配置
props (?)	标签	属性配置，详情请参见属性配置

配置示例

```
<beans xmlns="http://www.springframework.org/schema/beans"
  xmlns:xsi="http://www.w3.org/2001/XMLSchema-instance"
  xmlns:shardingsphere="http://shardingsphere.apache.org/schema/shardingsphere/datasource"
  xsi:schemaLocation="http://www.springframework.org/schema/beans
    http://www.springframework.org/schema/beans/spring-beans.xsd
    http://shardingsphere.apache.org/schema/shardingsphere/
```

```
datasource
    http://shardingsphere.apache.org/schema/shardingsphere/
datasource/datasource.xsd
">
    <shardingsphere:data-source id="ds" database-name="foo_schema" data-source-
names="..." rule-refs="...">
        <shardingsphere:mode type="..." />
        <props>
            <prop key="xxx.xxx">${xxx.xxx}</prop>
        </props>
    </shardingsphere:data-source>
</beans>
```

使用数据源

使用方式同 Spring Boot Starter。

模式配置

背景信息

缺省配置为使用单机模式。

参数解释

单机模式

命名空间：<http://shardingsphere.apache.org/schema/shardingsphere/mode-repository/standalone/repository-5.1.1.xsd>

名称	类型	说明
id	属性	持久化仓库 Bean 名称
type	属性	持久化仓库类型
props (?)	标签	持久化仓库所需属性

集群模式 (推荐)

命名空间: <http://shardingsphere.apache.org/schema/shardingsphere/mode-repository/cluster/repository-5.1.1.xsd>

名称	类型	说明
id	属性	持久化仓库 Bean 名称
type	属性	持久化仓库类型
namespace	属性	注册中心命名空间
server-lists	属性	注册中心连接地址
props (?)	标签	持久化仓库所需属性

注意事项

1. 生产环境建议使用集群模式部署。
2. 集群模式部署推荐使用 ZooKeeper 注册中心。

操作步骤

引入 MAVEN 依赖

```
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-jdbc-core-spring-namespace</artifactId>
  <version>${latest.release.version}</version>
</dependency>
```

注意: 请将 `${latest.release.version}` 更改为实际的版本号。

配置示例

单机模式

```
<?xml version="1.0" encoding="UTF-8"?>
<beans xmlns="http://www.springframework.org/schema/beans"
  xmlns:xsi="http://www.w3.org/2001/XMLSchema-instance"
  xmlns:shardingsphere="http://shardingsphere.apache.org/schema/
shardingsphere/datasource"
  xmlns:standalone="http://shardingsphere.apache.org/schema/shardingsphere/
mode-repository/standalone"
  xsi:schemaLocation="http://www.springframework.org/schema/beans
  http://www.springframework.org/schema/beans/spring-
beans.xsd
  http://shardingsphere.apache.org/schema/shardingsphere/
```

```

datasource
    http://shardingsphere.apache.org/schema/shardingsphere/
datasource/datasource.xsd
    http://shardingsphere.apache.org/schema/shardingsphere/
mode-repository/standalone
    http://shardingsphere.apache.org/schema/shardingsphere/
mode-repository/standalone/repository.xsd">
    <standalone:repository id="standaloneRepository" type="File">
        <props>
            <prop key="path">.shardingsphere</prop>
        </props>
    </standalone:repository>

    <shardingsphere:data-source id="ds" database-name="foo_db" data-source-names="
    .." rule-refs="..." >
        <shardingsphere:mode type="Standalone" repository-ref="standaloneRepository
    " overwrite="false" />
    </shardingsphere:data-source>
</beans>

```

集群模式

```

<?xml version="1.0" encoding="UTF-8"?>
<beans xmlns="http://www.springframework.org/schema/beans"
    xmlns:xsi="http://www.w3.org/2001/XMLSchema-instance"
    xmlns:shardingsphere="http://shardingsphere.apache.org/schema/
shardingsphere/datasource"
    xmlns:cluster="http://shardingsphere.apache.org/schema/shardingsphere/mode-
repository/cluster"
    xsi:schemaLocation="http://www.springframework.org/schema/beans
        http://www.springframework.org/schema/beans/spring-
beans.xsd
        http://shardingsphere.apache.org/schema/shardingsphere/
datasource
        http://shardingsphere.apache.org/schema/shardingsphere/
datasource/datasource.xsd
        http://shardingsphere.apache.org/schema/shardingsphere/
mode-repository/cluster
        http://shardingsphere.apache.org/schema/shardingsphere/
mode-repository/cluster/repository.xsd">
    <cluster:repository id="clusterRepository" type="Zookeeper" namespace=
    "regCenter" server-lists="localhost:3182">
        <props>
            <prop key="max-retries">3</prop>
            <prop key="operation-timeout-milliseconds">1000</prop>
        </props>
    </cluster:repository>

```

```
</cluster:repository>

<shardingsphere:data-source id="ds" database-name="foo_db" data-source-names="
.." rule-refs="...">
    <shardingsphere:mode type="Cluster" repository-ref="clusterRepository"
overwrite="false" />
</shardingsphere:data-source>
</beans>
```

相关参考

- [ZooKeeper 注册中心安装与使用](#)
- 持久化仓库类型的详情，请参见[内置持久化仓库类型列表](#)。

数据源配置

背景信息

任何配置成为 Spring Bean 的数据源对象即可与 ShardingSphere-JDBC 的 Spring 命名空间配合使用。

配置示例的数据库驱动为 MySQL，连接池为 HikariCP，可以更换为其他数据库驱动和连接池。当使用 ShardingSphere JDBC 时，JDBC 池的属性名取决于各自 JDBC 池自己的定义，并不由 ShardingSphere 硬定义，相关的处理可以参考类 `org.apache.shardingsphere.infra.datasource.pool.creator.DataSourcePoolCreator`。例如对于 Alibaba Druid 1.2.9 而言，使用 `url` 代替如下示例中的 `jdbcUrl` 是预期行为。

操作步骤

1. 引入 MAVEN 依赖

```
<dependency>
    <groupId>org.apache.shardingsphere</groupId>
    <artifactId>shardingsphere-jdbc-core-spring-namespace</artifactId>
    <version>${latest.release.version}</version>
</dependency>
```

注意：请将 `${latest.release.version}` 更改为实际的版本号。

配置示例

```

<beans xmlns="http://www.springframework.org/schema/beans"
        xmlns:xsi="http://www.w3.org/2001/XMLSchema-instance"
        xmlns:shardingsphere="http://shardingsphere.apache.org/schema/
shardingsphere/datasource"
        xsi:schemaLocation="http://www.springframework.org/schema/beans
                            http://www.springframework.org/schema/beans/spring-
beans.xsd
                            http://shardingsphere.apache.org/schema/shardingsphere/
datasource
                            http://shardingsphere.apache.org/schema/shardingsphere/
datasource/datasource.xsd"
        ">
    <bean id="ds1" class="com.zaxxer.hikari.HikariDataSource" destroy-method="close"
">
        <property name="driverClassName" value="com.mysql.jdbc.Driver" />
        <property name="jdbcUrl" value="jdbc:mysql://localhost:3306/ds1" />
        <property name="username" value="root" />
        <property name="password" value="" />
    </bean>

    <bean id="ds2" class="com.zaxxer.hikari.HikariDataSource" destroy-method="close"
">
        <property name="driverClassName" value="com.mysql.jdbc.Driver" />
        <property name="jdbcUrl" value="jdbc:mysql://localhost:3306/ds2" />
        <property name="username" value="root" />
        <property name="password" value="" />
    </bean>

    <shardingsphere:data-source id="ds" database-name="foo_schema" data-source-
names="ds1,ds2" rule-refs="..." />
</beans>

```

规则配置

规则是 Apache ShardingSphere 面向可插拔的一部分。本章节是 ShardingSphere-JDBC 的 Spring 命名空间规则配置参考手册。

数据分片

背景信息

数据分片 Spring 命名空间的配置方式，适用于传统的 Spring 项目，通过命名空间 xml 配置文件的方式配置分片规则和属性，由 Spring 完成 ShardingSphereDataSource 对象的创建和管理，避免额外的编码工作。

参数解释

命名空间：<http://shardingsphere.apache.org/schema/shardingsphere/sharding/sharding-5.2.0.xsd>

<sharding:rule />

名称	类型	说明
id	属性	Spring Bean Id
table-rules (?)	标签	分片表规则配置
auto-table-rules (?)	标签	自动分片表规则配置
binding-table-rules (?)	标签	绑定表规则配置
broadcast-table-rules (?)	标签	广播表规则配置
default-database-strategy-ref (?)	属性	默认分库策略名称
default-table-strategy-ref (?)	属性	默认分表策略名称
default-key-generate-strategy-ref (?)	属性	默认分布式序列策略名称
default-sharding-column (?)	属性	默认分片列名称

<sharding:table-rule />

名称	• 类型 *	说明
logic-table	属性	逻辑表名称
actual-data-nodes	属性	由数据源名 + 表名组成，以小数点分隔。多个表以逗号分隔，支持 inline 表达式。缺省表示使用已知数据源与逻辑表名称生成数据节点，用于广播表（即每个库中都需要一个同样的表用于关联查询，多为字典表）或只分库不分表且所有库的表结构完全一致的情况
actual-data-sources	属性	自动分片表数据源名
database-strategy-ref	属性	标准分片表分库策略名称
table-strategy-ref	属性	标准分片表分表策略名称
sharding-strategy-ref	属性	自动分片表策略名称
key-generate-strategy-ref	属性	分布式序列策略名称

<sharding:binding-table-rules />

名称	类型	说明
binding-table-rule (+)	标签	绑定表规则配置

<sharding:binding-table-rule />

名称	类型	说明
logic-tables	属性	绑定表名称，多个表以逗号分隔

<sharding:broadcast-table-rules />

名称	类型	说明
broadcast-table-rule (+)	标签	广播表规则配置

<sharding:broadcast-table-rule />

名称	类型	说明
table	属性	广播表名称

<sharding:standard-strategy />

名称	类型	说明
id	属性	标准分片策略名称
sharding-column	属性	分片列名称
algorithm-ref	属性	分片算法名称

<sharding:complex-strategy />

名称	类型	说明
id	属性	复合分片策略名称
sharding-columns	属性	分片列名称，多个列以逗号分隔
algorithm-ref	属性	分片算法名称

<sharding:hint-strategy />

名称	类型	说明
id	属性	Hint 分片策略名称
algorithm-ref	属性	分片算法名称

<sharding:none-strategy />

名称	类型	说明
id	属性	分片策略名称

<sharding:key-generate-strategy />

名称	类型	说明
id	属性	分布式序列策略名称
column	属性	分布式序列列名称
algorithm-ref	属性	分布式序列算法名称

<sharding:sharding-algorithm />

名称	类型	说明
id	属性	分片算法名称
type	属性	分片算法类型
props (?)	标签	分片算法属性配置

<sharding:key-generate-algorithm />

名称	类型	说明
id	属性	分布式序列算法名称
type	属性	分布式序列算法类型
props (?)	标签	分布式序列算法属性配置

算法类型的详情，请参见[内置分片算法列表](#)和[内置分布式序列算法列表](#)。

注意事项：行表达式标识符可以使用 `${...}` 或 `$->{...}`，但前者与 Spring 本身的属性文件占位符冲突，因此在 Spring 环境中使用行表达式标识符建议使用 `$->{...}`。

操作步骤

1. 在 Spring 命名空间配置文件中配置数据分片规则，包含数据源、分片规则、全局属性等配置项；
2. 启动 Spring 程序，会自动加载配置，并初始化 ShardingSphereDataSource。

配置示例

```
<beans xmlns="http://www.springframework.org/schema/beans"
  xmlns:xsi="http://www.w3.org/2001/XMLSchema-instance"
  xmlns:context="http://www.springframework.org/schema/context"
  xmlns:tx="http://www.springframework.org/schema/tx"
  xmlns:shardingsphere="http://shardingsphere.apache.org/schema/shardingsphere/datasource"
  xmlns:sharding="http://shardingsphere.apache.org/schema/shardingsphere/sharding"
  xsi:schemaLocation="http://www.springframework.org/schema/beans
    http://www.springframework.org/schema/beans/spring-beans.xsd"
```

```

        http://www.springframework.org/schema/tx
        http://www.springframework.org/schema/tx/spring-tx.xsd
        http://www.springframework.org/schema/context
        http://www.springframework.org/schema/context/spring-
context.xsd
        http://shardingsphere.apache.org/schema/shardingsphere/
datasource
        http://shardingsphere.apache.org/schema/shardingsphere/
datasource/datasource.xsd
        http://shardingsphere.apache.org/schema/shardingsphere/
sharding
        http://shardingsphere.apache.org/schema/shardingsphere/
sharding/sharding.xsd
    ">
    <context:component-scan base-package="org.apache.shardingsphere.example.core.
mybatis" />

    <bean id="demo_ds_0" class="com.zaxxer.hikari.HikariDataSource" destroy-method=
"close">
        <property name="driverClassName" value="com.mysql.jdbc.Driver"/>
        <property name="jdbcUrl" value="jdbc:mysql://localhost:3306/demo_ds_0?
serverTimezone=UTC&useSSL=false&useUnicode=true&characterEncoding=UTF-8
"/>
        <property name="username" value="root"/>
        <property name="password" value=""/>
    </bean>

    <bean id="demo_ds_1" class="com.zaxxer.hikari.HikariDataSource" destroy-method=
"close">
        <property name="driverClassName" value="com.mysql.jdbc.Driver"/>
        <property name="jdbcUrl" value="jdbc:mysql://localhost:3306/demo_ds_1?
serverTimezone=UTC&useSSL=false&useUnicode=true&characterEncoding=UTF-8
"/>
        <property name="username" value="root"/>
        <property name="password" value=""/>
    </bean>

    <sharding:standard-strategy id="databaseStrategy" sharding-column="user_id"
algorithm-ref="inlineStrategyShardingAlgorithm" />

    <sharding:sharding-algorithm id="inlineStrategyShardingAlgorithm" type="INLINE
">
        <props>
            <prop key="algorithm-expression">demo_ds_${user_id % 2}</prop>
        </props>
    </sharding:sharding-algorithm>

    <sharding:key-generate-algorithm id="snowflakeAlgorithm" type="SNOWFLAKE">

```

```

    </sharding:key-generate-algorithm>

    <sharding:key-generate-strategy id="orderKeyGenerator" column="order_id"
algorithm-ref="snowflakeAlgorithm" />
    <sharding:key-generate-strategy id="itemKeyGenerator" column="order_item_id"
algorithm-ref="snowflakeAlgorithm" />

    <sharding:rule id="shardingRule">
        <sharding:table-rules>
            <sharding:table-rule logic-table="t_order" database-strategy-ref=
"databaseStrategy" key-generate-strategy-ref="orderKeyGenerator" />
            <sharding:table-rule logic-table="t_order_item" database-strategy-ref=
"databaseStrategy" key-generate-strategy-ref="itemKeyGenerator" />
        </sharding:table-rules>
        <sharding:binding-table-rules>
            <sharding:binding-table-rule logic-tables="t_order,t_order_item"/>
        </sharding:binding-table-rules>
        <sharding:broadcast-table-rules>
            <sharding:broadcast-table-rule table="t_address"/>
        </sharding:broadcast-table-rules>
    </sharding:rule>

    <shardingsphere:data-source id="shardingDataSource" database-name="sharding-
databases" data-source-names="demo_ds_0, demo_ds_1" rule-refs="shardingRule" />

    <bean id="transactionManager" class="org.springframework.jdbc.datasource.
DataSourceTransactionManager">
        <property name="dataSource" ref="shardingDataSource" />
    </bean>
    <tx:annotation-driven />

    <bean id="sqlSessionFactory" class="org.mybatis.spring.SqlSessionFactoryBean">
        <property name="dataSource" ref="shardingDataSource"/>
        <property name="mapperLocations" value="classpath*:META-INF/mappers/*.xml"/
>
    </bean>

    <bean class="org.mybatis.spring.mapper.MapperScannerConfigurer">
        <property name="basePackage" value="org.apache.shardingsphere.example.core.
mybatis.repository"/>
        <property name="sqlSessionFactoryBeanName" value="sqlSessionFactory"/>
    </bean>
</beans>

```

相关参考

- 核心特性：数据分片
- 开发者指南：数据分片

读写分离

背景信息

读写分离 Spring 命名空间的配置方式，适用于传统的 Spring 项目，通过命名空间 xml 配置文件的方式配置分片规则和属性，由 Spring 完成 ShardingSphereDataSource 对象的创建和管理，避免额外的编码工作。

参数解释

命名空间：<http://shardingsphere.apache.org/schema/shardingsphere/readwrite-splitting/readwrite-splitting-5.2.0.xsd>

<readwrite-splitting:rule />

名称	类型	说明
id	属性	Spring Bean Id
data-source-rule (+)	标签	读写分离数据源规则配置

<readwrite-splitting:data-source-rule />

名称	类型	说明
id	属性	读写分离数据源规则名称
static-strategy	标签	静态读写分离类型
dynamic-strategy	标签	动态读写分离类型
load-balance-algorithm-ref	属性	负载均衡算法名称

<readwrite-splitting:static-strategy />

名称	类型	说明
id	属性	静态读写分离名称
write-data-source-name	属性	写库数据源名称
read-data-source-names	属性	读库数据源列表，多个从数据源用逗号分隔
load-balance-algorithm-ref	属性	负载均衡算法名称

<readwrite-splitting:dynamic-strategy />

名称	类型	说明
id	属性	动态读写分离名称
auto-aware-data-source-name	属性	数据库发现逻辑数据源名称
write-data-source-query-enabled	属性	读库全部下线，主库是否承担读流量
load-balance-algorithm-ref	属性	负载均衡算法名称

<readwrite-splitting:load-balance-algorithm />

名称	类型	说明
id	属性	负载均衡算法名称
type	属性	负载均衡算法类型
props (?)	标签	负载均衡算法属性配置

算法类型的详情，请参见[内置负载均衡算法列表](#)。查询一致性路由的详情，请参见[核心特性：读写分离](#)。

操作步骤

1. 添加读写分离数据源
2. 设置负载均衡算法
3. 使用读写分离数据源

配置示例

```
<readwrite-splitting:load-balance-algorithm id="randomStrategy" type="RANDOM" />

<readwrite-splitting:rule id="readWriteSplittingRule">
  <readwrite-splitting:data-source-rule id="demo_ds" load-balance-algorithm-ref=
    "randomStrategy">
    <readwrite-splitting:static-strategy id="staticStrategy" write-data-source-
      name="demo_write_ds" read-data-source-names="demo_read_ds_0, demo_read_ds_1"/>
    </readwrite-splitting:data-source-rule>
  </readwrite-splitting:rule>

<shardingsphere:data-source id="readWriteSplittingDataSource" data-source-names=
  "demo_write_ds, demo_read_ds_0, demo_read_ds_1" rule-refs="readWriteSplittingRule"
/>
```


相关参考

- 核心特性：读写分离
- Java API：读写分离
- YAML 配置：读写分离
- Spring Boot Starter：读写分离

高可用

背景信息

Spring 命名空间的配置方式，适用于传统的 Spring 项目，通过命名空间 xml 配置文件的方式配置高可用规则，由 Spring 完成 ShardingSphereDataSource 对象的创建和管理。

参数解释

命名空间：<http://shardingsphere.apache.org/schema/shardingsphere/database-discovery/database-discovery-5.1.1.xsd>

<database-discovery:rule />

名称	类型	说明
id	属性	Spring Bean Id
data-source-rule (+)	标签	数据源规则配置
discovery-heartbeat (+)	标签	检测心跳规则配置

<database-discovery:data-source-rule />

名称	类型	说明
id	属性	数据源规则名称
data-source-names	属性	数据源名称，多个数据源用逗号分隔如：ds_0, ds_1
discovery-heartbeat-name	属性	检测心跳名称
discovery-type-name	属性	数据库发现类型名称

<database-discovery:discovery-heartbeat />

名称	类型	说明
id	属性	监听心跳名称
props	标签	监听心跳属性配置，keep-alive-cron 属性配置 cron 表达式，如：‘0/5 * * * * ?’

<database-discovery:discovery-type />

名称	类型	说明
id	属性	数据库发现类型名称
type	属性	数据库发现类型，如：MySQL.MGR
props (?)	标签	数据库发现类型配置，如 MGR 的 group-name 属性配置

操作步骤

1. 引入 MAVEN 依赖

```
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-jdbc-core-spring-namespace</artifactId>
  <version>${latest.release.version}</version>
</dependency>
```

配置示例

```
<beans xmlns="http://www.springframework.org/schema/beans"
  xmlns:xsi="http://www.w3.org/2001/XMLSchema-instance"
  xmlns:cluster="http://shardingsphere.apache.org/schema/shardingsphere/mode-
repository/cluster"
  xmlns:shardingsphere="http://shardingsphere.apache.org/schema/
shardingsphere/datasource"
  xmlns:database-discovery="http://shardingsphere.apache.org/schema/
shardingsphere/database-discovery"
  xmlns:readwrite-splitting="http://shardingsphere.apache.org/schema/
shardingsphere/readwrite-splitting"
  xsi:schemaLocation="http://www.springframework.org/schema/beans
    http://www.springframework.org/schema/beans/spring-
beans.xsd
    http://shardingsphere.apache.org/schema/shardingsphere/
database-discovery
    http://shardingsphere.apache.org/schema/shardingsphere/
database-discovery/database-discovery.xsd
    http://shardingsphere.apache.org/schema/shardingsphere/
readwrite-splitting
    http://shardingsphere.apache.org/schema/shardingsphere/
readwrite-splitting/readwrite-splitting.xsd
    http://shardingsphere.apache.org/schema/shardingsphere/
mode-repository/cluster
    http://shardingsphere.apache.org/schema/shardingsphere/
mode-repository/cluster/repository.xsd
    http://shardingsphere.apache.org/schema/shardingsphere/
datasource
```

```

        http://shardingsphere.apache.org/schema/shardingsphere/
datasource/datasource.xsd
    ">
    <bean id="ds_0" class="com.zaxxer.hikari.HikariDataSource" destroy-method=
"close">
        <property name="driverClassName" value="com.mysql.jdbc.Driver" />
        <property name="jdbcUrl" value="jdbc:mysql://127.0.0.1:33306/primary_demo_
ds?serverTimezone=UTC&useSSL=false&useUnicode=true&
characterEncoding=UTF-8" />
        <property name="username" value="root" />
        <property name="password" value="" />
    </bean>
    <bean id="ds_1" class="com.zaxxer.hikari.HikariDataSource" destroy-method=
"close">
        <property name="driverClassName" value="com.mysql.jdbc.Driver"/>
        <property name="jdbcUrl" value="jdbc:mysql://127.0.0.1:33307/primary_demo_
ds?serverTimezone=UTC&useSSL=false&useUnicode=true&
characterEncoding=UTF-8" />
        <property name="username" value="root" />
        <property name="password" value="" />
    </bean>
    <bean id="ds_2" class="com.zaxxer.hikari.HikariDataSource" destroy-method=
"close">
        <property name="driverClassName" value="com.mysql.jdbc.Driver"/>
        <property name="jdbcUrl" value="jdbc:mysql://127.0.0.1:33308/primary_demo_
ds?useSSL=false"/>
        <property name="username" value="root"/>
        <property name="password" value="" />
    </bean>
    <cluster:repository id="clusterRepository" type="ZooKeeper" namespace=
"governance" server-lists="localhost:2181">
        <props>
            <prop key="max-retries">3</prop>
            <prop key="operation-timeout-milliseconds">3000</prop>
        </props>
    </cluster:repository>
    <readwrite-splitting:rule id="readWriteSplittingRule">
        <readwrite-splitting:data-source-rule id="replica_ds">
            <readwrite-splitting:dynamic-strategy id="dynamicStrategy" auto-aware-
data-source-name="readwrite_ds" />
        </readwrite-splitting:data-source-rule>
    </readwrite-splitting:rule>
    <database-discovery:rule id="mgrDatabaseDiscoveryRule">
        <database-discovery:data-source-rule id="readwrite_ds" data-source-names=
"ds_0,ds_1,ds_2" discovery-heartbeat-name="mgr-heartbeat" discovery-type-name="mgr"
/>
        <database-discovery:discovery-heartbeat id="mgr-heartbeat">
            <props>

```

```

        <prop key="keep-alive-cron" >0/5 * * * * ?</prop>
    </props>
</database-discovery:discovery-heartbeat>
</database-discovery:rule>
<database-discovery:discovery-type id="mgr" type="MySQL.MGR">
    <props>
        <prop key="group-name">558edd3c-02ec-11ea-9bb3-080027e39bd2</prop>
    </props>
</database-discovery:discovery-type>
<shardingsphere:data-source id="databaseDiscoveryDataSource" schema-name=
"database-discovery-db" data-source-names="ds_0, ds_1, ds_2" rule-refs=
"readWriteSplittingRule, mgrDatabaseDiscoveryRule">
    <shardingsphere:mode repository-ref="clusterRepository" type="Cluster" />
</shardingsphere:data-source>
</beans>

```

相关参考

- 高可用核心特性
- JAVA API: 高可用配置
- YAML 配置: 高可用配置
- Spring Boot Starter: 高可用配置

数据加密

背景信息

数据加密 Spring 命名空间的配置方式，适用于传统的 Spring 项目，通过命名空间 xml 配置文件的方式配置分片规则和属性，由 Spring 完成 ShardingSphereDataSource 对象的创建和管理，避免额外的编码工作。

参数解释

命名空间: <http://shardingsphere.apache.org/schema/shardingsphere/encrypt/encrypt-5.2.0.xsd>

<encrypt:rule />

名称	类型	说明	默认值
id	属性	Spring Bean Id	
queryWithCipherColumn (?)	属性	是否使用加密列进行查询。在有原文列的情况下，可以使用原文列进行查询	true
table (+)	标签	加密表配置	

<encrypt:table />

名称	类型	说明
name	属性	加密表名称
column (+)	标签	加密列配置
query-with-cipher-column(?)	属性	该表是否使用加密列进行查询。在有原文列的情况下，可以使用原文列进行查询

<encrypt:column />

名称	类型	说明
logic-column	属性	加密列逻辑名称
cipher-column	属性	加密列名称
assisted-query-column (?)	属性	查询辅助列名称
plain-column (?)	属性	原文列名称
encrypt-algorithm-ref	属性	加密算法名称

<encrypt:encrypt-algorithm />

名称	类型	说明
id	属性	加密算法名称
type	属性	加密算法类型
props (?)	标签	加密算法属性配置

算法类型的详情，请参见[内置加密算法列表](#)。

操作步骤

1. 在 Spring 命名空间配置文件中配置数据加密规则，包含数据源、加密规则、全局属性等配置项；
2. 启动 Spring 程序，会自动加载配置，并初始化 ShardingSphereDataSource。

配置示例

```
<beans xmlns="http://www.springframework.org/schema/beans"
       xmlns:xsi="http://www.w3.org/2001/XMLSchema-instance"
       xmlns:shardingsphere="http://shardingsphere.apache.org/schema/
shardingsphere/datasource"
       xmlns:context="http://www.springframework.org/schema/context"
       xmlns:tx="http://www.springframework.org/schema/tx"
       xmlns:encrypt="http://shardingsphere.apache.org/schema/shardingsphere/
encrypt"
       xsi:schemaLocation="http://www.springframework.org/schema/beans
                           http://www.springframework.org/schema/beans/spring-
beans.xsd
                           http://www.springframework.org/schema/tx
                           http://www.springframework.org/schema/tx/spring-tx.xsd
                           http://www.springframework.org/schema/context
                           http://www.springframework.org/schema/context/spring-
context.xsd
                           http://shardingsphere.apache.org/schema/shardingsphere/
datasource
                           http://shardingsphere.apache.org/schema/shardingsphere/
datasource/datasource.xsd
                           http://shardingsphere.apache.org/schema/shardingsphere/
encrypt
                           http://shardingsphere.apache.org/schema/shardingsphere/
encrypt/encrypt.xsd"
       ">
    <context:component-scan base-package="org.apache.shardingsphere.example.core.
mybatis" />

    <bean id="ds" class="com.zaxxer.hikari.HikariDataSource" destroy-method="close"
">
        <property name="driverClassName" value="com.mysql.jdbc.Driver"/>
        <property name="jdbcUrl" value="jdbc:mysql://localhost:3306/demo_ds?
serverTimezone=UTC&useSSL=false&useUnicode=true&characterEncoding=UTF-8"
"/>
        <property name="username" value="root"/>
        <property name="password" value=""/>
    </bean>

    <encrypt:encrypt-algorithm id="name_encryptor" type="AES">
        <props>
```

```

        <prop key="aes-key-value">123456</prop>
    </props>
</encrypt:encrypt-algorithm>
<encrypt:encrypt-algorithm id="pwd_encryptor" type="assistedTest" />

<encrypt:rule id="encryptRule">
    <encrypt:table name="t_user">
        <encrypt:column logic-column="username" cipher-column="username" plain-
column="username_plain" encrypt-algorithm-ref="name_encryptor" />
        <encrypt:column logic-column="pwd" cipher-column="pwd" assisted-query-
column="assisted_query_pwd" encrypt-algorithm-ref="pwd_encryptor" />
    </encrypt:table>
</encrypt:rule>

<shardingsphere:data-source id="encryptDataSource" data-source-names="ds" rule-
refs="encryptRule" />

<bean id="transactionManager" class="org.springframework.jdbc.datasource.
DataSourceTransactionManager">
    <property name="dataSource" ref="encryptDataSource" />
</bean>
<tx:annotation-driven />

<bean id="sqlSessionFactory" class="org.mybatis.spring.SqlSessionFactoryBean">
    <property name="dataSource" ref="encryptDataSource"/>
    <property name="mapperLocations" value="classpath*:META-INF/mappers/*.xml"/
>
</bean>

<bean class="org.mybatis.spring.mapper.MapperScannerConfigurer">
    <property name="basePackage" value="org.apache.shardingsphere.example.core.
mybatis.repository"/>
    <property name="sqlSessionFactoryBeanName" value="sqlSessionFactory"/>
</bean>
</beans>

```

相关参考

- 核心特性：数据加密
- 开发者指南：数据加密

影子库

背景信息

如果您只想使用 XML 配置文件方式配置使用 ShardingSphere 影子库功能请参考以下配置。

参数解释

配置入口

```
<shadow:rule />
```

可配置属性：

名称	类型	说明
id	属性	Spring Bean Id
data-source(?)	标签	影子库数据源映射配置
shadow-table(?)	标签	影子表配置
shadow-algorithm(?)	标签	影子表配置
default-shadow-algorithm-name(?)	标签	默认影子算法名称

影子数据源配置：

```
<shadow:data-source />
```

名称	类型	说明
id	属性	Spring Bean Id
production-data-source-name	属性	生产数据源名称
shadow-data-source-name	属性	影子数据源名称

影子表配置：

```
<shadow:shadow-table />
```

名称	类型	说明
name	属性	影子表名称
data-sources	属性	影子表关联影子数据源名称列表（多个值用” ，“隔开）
algorithm (?)	标签	影子表关联影子算法配置


```
<shadow:algorithm />
```

名称	类型	说明
shadow-algorithm-ref	属性	影子表关联影子算法名称

影子算法配置：

```
<shadow:shadow-algorithm />
```

名称	类型	说明
id	属性	影子算法名称
type	属性	影子算法类型
props (?)	标签	影子算法属性配置

详情请参见[内置影子算法列表](#)

操作步骤

1. 创建生产和影子数据源。
2. 配置影子规则
 - 配置影子数据源
 - 配置影子表
 - 配置影子算法

配置示例

```
<beans xmlns="http://www.springframework.org/schema/beans" xmlns:xsi="http://www.
w3.org/2001/XMLSchema-instance" xmlns:shadow="http://shardingsphere.apache.org/
schema/shardingsphere/shadow" xsi:schemaLocation="http://www.springframework.org/
schema/beans
                                http://www.springframework.org/schema/beans/spring-
beans.xsd
                                http://shardingsphere.apache.org/schema/shardingsphere/
shadow
                                http://shardingsphere.apache.org/schema/shardingsphere/
shadow/shadow.xsd
                                ">
    <shadow:shadow-algorithm id="user-id-insert-match-algorithm" type="VALUE_MATCH
">
        <props>
```

```

        <prop key="operation">insert</prop>
        <prop key="column">user_id</prop>
        <prop key="value">1</prop>
    </props>
</shadow:shadow-algorithm>

    <shadow:rule id="shadowRule">
        <shadow:data-source id="shadow-data-source" production-data-source-name="ds
" shadow-data-source-name="ds_shadow"/>
        <shadow:shadow-table name="t_user" data-sources="shadow-data-source">
            <shadow:algorithm shadow-algorithm-ref="user-id-insert-match-algorithm" />
        </shadow:shadow-table>
    </shadow:rule>
</beans>

```

相关参考

- 影子库的特性描述
- JAVA API：影子库的配置
- YAML 配置：影子库的配置
- Spring 命名空间：影子库的配置
- 开发者指南：影子库的接口和示例

SQL 解析

背景信息

Spring 命名空间的配置方式，适用于传统的 Spring 项目，它通过命名空间 xml 配置文件的方式配置 SQL 解析规则和属性。

参数解释

命名空间：<http://shardingsphere.apache.org/schema/shardingsphere/sql-parser/sql-parser-5.2.0.xsd>

<sql-parser:rule />

名称	类型	说明
id	属性	Spring Bean Id
sql-comment-parse-enable	属性	是否解析 SQL 注释
parse-tree-cache-ref	属性	解析树本地缓存名称
sql-statement-cache-ref	属性	SQL 语句本地缓存名称

<sql-parser:cache-option />

名称	类型	说明
id	属性	本地缓存配置项名称
initial-capacity	属性	本地缓存初始容量
maximum-size	属性	本地缓存最大容量

操作步骤

1. 设置本地缓存配置
2. 设置解析配置
3. 使用解析引擎解析 SQL

配置示例

```
<sql-parser:rule id="sqlParseRule" sql-comment-parse-enable="true" parse-tree-cache-ref="parseTreeCache" sql-statement-cache-ref="sqlStatementCache" />
<sql-parser:cache-option id="sqlStatementCache" initial-capacity="1024" maximum-size="1024"/>
<sql-parser:cache-option id="parseTreeCache" initial-capacity="1024" maximum-size="1024"/>
```

相关参考

- [JAVA API: SQL 解析](#)
- [YAML 配置: SQL 解析](#)
- [Spring Boot Starter: SQL 解析](#)

混合规则

背景信息

ShardingSphere 涵盖了很多功能，例如，分库分片、读写分离、高可用、数据脱敏等。这些功能用户可以单独进行使用，也可以配合一起使用，下面是基于 Spring 命名空间配置示例。

配置示例

```

<!-- 分片配置 -->
<sharding:standard-strategy id="databaseStrategy" sharding-column="user_id"
algorithm-ref="inlineStrategyShardingAlgorithm" />
<sharding:sharding-algorithm id="inlineStrategyShardingAlgorithm" type="INLINE">
  <props>
    <prop key="algorithm-expression">replica_ds_${user_id % 2}</prop>
  </props>
</sharding:sharding-algorithm>
<sharding:key-generate-algorithm id="snowflakeAlgorithm" type="SNOWFLAKE">
</sharding:key-generate-algorithm>
<sharding:key-generate-strategy id="orderKeyGenerator" column="order_id" algorithm-
ref="snowflakeAlgorithm" />
<sharding:rule id="shardingRule">
  <sharding:table-rules>
    <sharding:table-rule logic-table="t_order" database-strategy-ref=
"databaseStrategy" key-generate-strategy-ref="orderKeyGenerator" />
  </sharding:table-rules>
</sharding:rule>

<!-- 动态读写分离配置 -->
<readwrite-splitting:rule id="readWriteSplittingRule">
  <readwrite-splitting:data-source-rule id="replica_ds_0">
    <readwrite-splitting:dynamic-strategy id="dynamicStrategy" auto-aware-data-
source-name="readwrite_ds_0" />
  </readwrite-splitting:data-source-rule>
  <readwrite-splitting:data-source-rule id="replica_ds_1">
    <readwrite-splitting:dynamic-strategy id="dynamicStrategy" auto-aware-data-
source-name="readwrite_ds_1" />
  </readwrite-splitting:data-source-rule>
</readwrite-splitting:rule>

<!-- 数据库发现配置 -->
<database-discovery:rule id="mgrDatabaseDiscoveryRule">
  <database-discovery:data-source-rule id="readwrite_ds_0" data-source-names="ds_
0,ds_1,ds_2" discovery-heartbeat-name="mgr-heartbeat" discovery-type-name="mgr" />
  <database-discovery:data-source-rule id="readwrite_ds_1" data-source-names="ds_
3,ds_4,ds_5" discovery-heartbeat-name="mgr-heartbeat" discovery-type-name="mgr" />
  <database-discovery:discovery-heartbeat id="mgr-heartbeat">
    <props>
      <prop key="keep-alive-cron">0/5 * * * * ?</prop>
    </props>
  </database-discovery:discovery-heartbeat>
</database-discovery:rule>
<database-discovery:discovery-type id="mgr" type="MySQL.MGR">
  <props>
    <prop key="group-name">558edd3c-02ec-11ea-9bb3-080027e39bd2</prop>
  </props>
</database-discovery:discovery-type>

```

```

    </props>
</database-discovery:discovery-type>

<!-- 数据脱敏配置 -->
<encrypt:encrypt-algorithm id="name_encryptor" type="AES">
    <props>
        <prop key="aes-key-value">123456</prop>
    </props>
</encrypt:encrypt-algorithm>
<encrypt:encrypt-algorithm id="pwd_encryptor" type="assistedTest" />

<encrypt:rule id="encryptRule">
    <encrypt:table name="t_user">
        <encrypt:column logic-column="username" cipher-column="username" plain-
column="username_plain" encrypt-algorithm-ref="name_encryptor" />
        <encrypt:column logic-column="pwd" cipher-column="pwd" assisted-query-
column="assisted_query_pwd" encrypt-algorithm-ref="pwd_encryptor" />
    </encrypt:table>
</encrypt:rule>

```

算法配置

分片算法

```

<!-- algorithmName 由用户指定，需要和分片策略中的 algorithm-ref 属性一致 -->
<!-- type 和 props，请参考分片内置算法: https://shardingsphere.apache.org/document/current/cn/user-manual/common-config/builtin-algorithm/sharding/ -->
<sharding:sharding-algorithm id="algorithmName" type="xxx">
    <props>
        <prop key="xxx">xxx</prop>
    </props>
</sharding:sharding-algorithm>

```

加密算法

```

<!-- encryptorName 由用户指定，需要和加密规则中的 encrypt-algorithm-ref 属性一致 -->
<!-- type 和 props，请参考加密内置算法: https://shardingsphere.apache.org/document/current/cn/user-manual/common-config/builtin-algorithm/encrypt/ -->
<encrypt:encrypt-algorithm id="encryptorName" type="xxx">
    <props>
        <prop key="xxx">xxx</prop>
    </props>
</encrypt:encrypt-algorithm>

```

读写分离负载均衡算法

```
<!-- loadBalancerName 由用户指定, 需要和读写分离规则中的 load-balance-algorithm-ref 属性一致 -->
<!-- type 和 props, 请参考读写分离负载均衡内置算法: https://shardingsphere.apache.org/document/current/cn/user-manual/common-config/builtin-algorithm/load-balance/ -->
<readwrite-splitting:load-balance-algorithm id="loadBalancerName" type="xxx">
  <props>
    <prop key="xxx">xxx</prop>
  </props>
</readwrite-splitting:load-balance-algorithm>
```

影子算法

```
<!-- shadowAlgorithmName 由用户指定, 需要和影子库规则中的 shadow-algorithm-ref 属性一致 -->
<!-- type 和 props, 请参考影子库内置算法: https://shardingsphere.apache.org/document/current/cn/user-manual/common-config/builtin-algorithm/shadow/ -->
<shadow:shadow-algorithm id="shadowAlgorithmName" type="xxx">
  <props>
    <prop key="xxx">xxx</prop>
  </props>
</shadow:shadow-algorithm>
```

高可用

```
<!-- discoveryTypeName 由用户指定, 需要和数据库发现规则中的 discovery-type-name 属性一致 -->
<database-discovery:discovery-type id="discoveryTypeName" type="xxx">
  <props>
    <prop key="xxx">xxx</prop>
  </props>
</database-discovery:discovery-type>
```

4.1.5 特殊 API

本章节将介绍 ShardingSphere-JDBC 的特殊场景 API。

数据分片

本章节将介绍 ShardingSphere-JDBC 的分片场景 API。

强制路由

背景信息

Apache ShardingSphere 使用 ThreadLocal 管理分片键值进行强制路由。可以通过编程的方式向 HintManager 中添加分片值，该分片值仅在当前线程内生效。Apache ShardingSphere 还可以通过 SQL 中增加注释的方式进行强制路由。

Hint 的主要使用场景：- 分片字段不存在 SQL 和数据库表结构中，而存在于外部业务逻辑。- 强制在指定数据库进行某些数据操作。

操作步骤

1. 调用 HintManager.getInstance() 获取 HintManager 实例；
2. 调用 HintManager.addDatabaseShardingValue, HintManager.addTableShardingValue 方法设置分片键值；
3. 执行 SQL 语句完成路由和执行；
4. 调用 HintManager.close 清理 ThreadLocal 中的内容。

配置示例

使用 Hint 分片

规则配置

Hint 分片算法需要用户实现 org.apache.shardingsphere.sharding.api.sharding.hint.HintShardingAlgorithm 接口。Apache ShardingSphere 在进行路由时，将会从 HintManager 中获取分片值进行路由操作。

参考配置如下：

```
rules:
- !SHARDING
  tables:
    t_order:
      actualDataNodes: demo_ds_${0..1}.t_order_${0..1}
```

```

    databaseStrategy:
      hint:
        algorithmClassName: xxx.xxx.xxx.HintXXXAlgorithm
    tableStrategy:
      hint:
        algorithmClassName: xxx.xxx.xxx.HintXXXAlgorithm
    defaultTableStrategy:
      none:
    defaultKeyGenerateStrategy:
      type: SNOWFLAKE
      column: order_id

props:
  sql-show: true

```

获取 HintManager

```
HintManager hintManager = HintManager.getInstance();
```

添加分片键值

- 使用 `hintManager.addDatabaseShardingValue` 来添加数据源分片键值。
- 使用 `hintManager.addTableShardingValue` 来添加表分片键值。

分库不分表情况下，强制路由至某一个分库时，可使用 `hintManager.setDatabaseShardingValue` 方式添加分片。

清除分片键值

分片键值保存在 `ThreadLocal` 中，所以需要在操作结束时调用 `hintManager.close()` 来清除 `ThreadLocal` 中的内容。

hintManager 实现了 **AutoCloseable** 接口，可推荐使用 **try with resource** 自动关闭。

完整代码示例

```

// Sharding database and table with using HintManager
String sql = "SELECT * FROM t_order";
try (HintManager hintManager = HintManager.getInstance();
     Connection conn = dataSource.getConnection();
     PreparedStatement preparedStatement = conn.prepareStatement(sql)) {
    hintManager.addDatabaseShardingValue("t_order", 1);
    hintManager.addTableShardingValue("t_order", 2);
    try (ResultSet rs = preparedStatement.executeQuery()) {

```



```
        while (rs.next()) {
            // ...
        }
    }
}

// Sharding database without sharding table and routing to only one database with
// using HintManager
String sql = "SELECT * FROM t_order";
try (HintManager hintManager = HintManager.getInstance();
     Connection conn = dataSource.getConnection();
     PreparedStatement preparedStatement = conn.prepareStatement(sql)) {
    hintManager.setDatabaseShardingValue(3);
    try (ResultSet rs = preparedStatement.executeQuery()) {
        while (rs.next()) {
            // ...
        }
    }
}
```

相关参考

- 核心特性：数据分片
- 开发者指南：数据分片

读写分离

本章节将介绍 ShardingSphere-JDBC 的读写分离场景 API。

强制路由

背景信息

Apache ShardingSphere 使用 ThreadLocal 管理主库路由标记进行强制路由。可以通过编程的方式向 HintManager 中添加主库路由标记，该值仅在当前线程内生效。Apache ShardingSphere 还可以通过 SQL 中增加注释的方式进行主库路由。

Hint 在读写分离场景下，主要用于强制在主库进行某些数据操作。

操作步骤

1. 调用 `HintManager.getInstance()` 获取 `HintManager` 实例;
2. 调用 `HintManager.setWriteRouteOnly()` 方法设置主库路由标记;
3. 执行 SQL 语句完成路由和执行;
4. 调用 `HintManager.close()` 清理 `ThreadLocal` 中的内容。

配置示例

使用 **Hint** 强制主库路由

使用手动编程的方式

获取 **HintManager**

与基于 `Hint` 的数据分片相同。

设置主库路由

使用 `hintManager.setWriteRouteOnly` 设置主库路由。

清除分片键值

与基于 `Hint` 的数据分片相同。

完整代码示例

```
String sql = "SELECT * FROM t_order";
try (HintManager hintManager = HintManager.getInstance();
    Connection conn = dataSource.getConnection();
    PreparedStatement preparedStatement = conn.prepareStatement(sql)) {
    hintManager.setWriteRouteOnly();
    try (ResultSet rs = preparedStatement.executeQuery()) {
        while (rs.next()) {
            // ...
        }
    }
}
```

使用 SQL 注释的方式

使用规范

SQL Hint 功能需要用户提前开启解析注释的配置，设置 `sqlCommentParseEnabled` 为 `true`。注释格式暂时只支持 `/* */`，内容需要以 `ShardingSphere hint:` 开始，属性名为 `writeRouteOnly`。

完整示例

```
/* ShardingSphere hint: writeRouteOnly=true */  
SELECT * FROM t_order;
```

- 核心特性：读写分离
- 开发者指南：读写分离

分布式事务

通过 Apache ShardingSphere 使用分布式事务，与本地事务并无区别。除了透明化分布式事务的使用之外，Apache ShardingSphere 还能够在每次数据库访问时切换分布式事务类型。支持的事务类型包括：本地事务、XA 事务和柔性事务。可在创建数据库连接之前设置，缺省为 Apache ShardingSphere 启动时的默认事务类型。

使用 Java API

背景信息

使用 ShardingSphere-JDBC 时，可以通过 API 的方式使用 XA 和 BASE 模式的事务。

前提条件

引入 Maven 依赖

```
<dependency>  
  <groupId>org.apache.shardingsphere</groupId>  
  <artifactId>shardingsphere-jdbc-core</artifactId>  
  <version>${shardingsphere.version}</version>  
</dependency>  
  
<!-- 使用 XA 事务时，需要引入此模块 -->  
<dependency>  
  <groupId>org.apache.shardingsphere</groupId>  
  <artifactId>shardingsphere-transaction-xa-core</artifactId>  
  <version>${shardingsphere.version}</version>  
</dependency>
```

```

<!-- 使用 XA 的 Narayana 模式时, 需要引入此模块 -->
<dependency>
    <groupId>org.apache.shardingsphere</groupId>
    <artifactId>shardingsphere-transaction-xa-narayana</artifactId>
    <version>${project.version}</version>
</dependency>

<!-- 使用 BASE 事务时, 需要引入此模块 -->
<dependency>
    <groupId>org.apache.shardingsphere</groupId>
    <artifactId>shardingsphere-transaction-base-seata-at</artifactId>
    <version>${shardingsphere.version}</version>
</dependency>

```

操作步骤

1. 设置事务类型
2. 执行业务逻辑

配置示例

```

TransactionTypeHolder.set(TransactionType.XA); // 支持 TransactionType.LOCAL,
TransactionType.XA, TransactionType.BASE
    try (Connection conn = dataSource.getConnection()) { // 使用
ShardingSphereDataSource
        conn.setAutoCommit(false);
        PreparedStatement ps = conn.prepareStatement("INSERT INTO t_order (user_id,
status) VALUES (?, ?)");
        ps.setObject(1, 1000);
        ps.setObject(2, "init");
        ps.executeUpdate();
        conn.commit();
    }

```

使用 Spring Boot Starter

背景信息

使用 ShardingSphere-JDBC 时, 可以通过 spring boot starter 的方式使用。## 前提条件

引入 Maven 依赖

```

<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-jdbc-core-spring-boot-starter</artifactId>
  <version>${shardingsphere.version}</version>
</dependency>

<!-- 使用 XA 事务时, 需要引入此模块 -->
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-transaction-xa-core</artifactId>
  <version>${shardingsphere.version}</version>
</dependency>

<!-- 使用 XA 的 Narayana 模式时, 需要引入此模块 -->
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-transaction-xa-narayana</artifactId>
  <version>${project.version}</version>
</dependency>

<!-- 使用 BASE 事务时, 需要引入此模块 -->
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-transaction-base-seata-at</artifactId>
  <version>${shardingsphere.version}</version>
</dependency>

```

操作步骤

1. 配置事务类型
2. 使用分布式事务

配置示例

配置事务类型

```

@Configuration
@EnableTransactionManagement
public class TransactionConfiguration {

    @Bean
    public PlatformTransactionManager txManager(final DataSource dataSource) {
        return new DataSourceTransactionManager(dataSource);
    }
}

```

```

@Bean
public JdbcTemplate jdbcTemplate(final DataSource dataSource) {
    return new JdbcTemplate(dataSource);
}
}

```

使用分布式事务

```

@Transactional
@ShardingSphereTransactionType(TransactionType.XA) // 支持 TransactionType.LOCAL,
TransactionType.XA, TransactionType.BASE
public void insert() {
    jdbcTemplate.execute("INSERT INTO t_order (user_id, status) VALUES (?, ?)",
        (PreparedStatementCallback<Object>) ps -> {
            ps.setObject(1, i);
            ps.setObject(2, "init");
            ps.executeUpdate();
        });
}

```

使用 Spring 命名空间

背景信息

使用 ShardingSphere-JDBC 时，可以通过 spring namespace 的方式使用。## 前提条件

引入 Maven 依赖

```

<dependency>
    <groupId>org.apache.shardingsphere</groupId>
    <artifactId>shardingsphere-jdbc-core-spring-namespaces</artifactId>
    <version>${shardingsphere.version}</version>
</dependency>

<!-- 使用 XA 事务时，需要引入此模块 -->
<dependency>
    <groupId>org.apache.shardingsphere</groupId>
    <artifactId>shardingsphere-transaction-xa-core</artifactId>
    <version>${shardingsphere.version}</version>
</dependency>

<!-- 使用 XA 的 Narayana 模式时，需要引入此模块 -->
<dependency>
    <groupId>org.apache.shardingsphere</groupId>
    <artifactId>shardingsphere-transaction-xa-narayana</artifactId>
    <version>${project.version}</version>

```

```

</dependency>

<!-- 使用 BASE 事务时, 需要引入此模块 -->
<dependency>
    <groupId>org.apache.shardingsphere</groupId>
    <artifactId>shardingsphere-transaction-base-seata-at</artifactId>
    <version>${shardingsphere.version}</version>
</dependency>

```

操作步骤

1. 配置事务管理器
2. 使用分布式事务

配置示例

配置事务管理器

```

<!-- ShardingDataSource 的相关配置 -->
<!-- ... -->

<bean id="transactionManager" class="org.springframework.jdbc.datasource.
DataSourceTransactionManager">
    <property name="dataSource" ref="shardingDataSource" />
</bean>
<bean id="jdbcTemplate" class="org.springframework.jdbc.core.JdbcTemplate">
    <property name="dataSource" ref="shardingDataSource" />
</bean>
<tx:annotation-driven />

<!-- 开启自动扫描 @ShardingSphereTransactionType 注解, 使用 Spring 原生的 AOP 在类和方法上
进行增强 -->
<sharding:tx-type-annotation-driven />

```

使用分布式事务

```

@Transactional
@ShardingSphereTransactionType(TransactionType.XA) // 支持 TransactionType.LOCAL,
TransactionType.XA, TransactionType.BASE
public void insert() {
    jdbcTemplate.execute("INSERT INTO t_order (user_id, status) VALUES (?, ?)",
(PreparedStatementCallback<Object>) ps -> {
        ps.setObject(1, i);
    }
}

```

```
ps.setObject(2, "init");
ps.executeUpdate();
});
}
```

Atomikos 事务

背景信息

Apache ShardingSphere 提供 XA 事务，默认的 XA 事务实现为 Atomikos。## 操作步骤

1. 配置事务类型
2. 配置 Atomikos

配置示例

配置事务类型

Yaml:

```
- !TRANSACTION
  defaultType: XA
  providerType: Atomikos
```

SpringBoot:

```
spring:
  shardingsphere:
    props:
      xa-transaction-manager-type: Atomikos
```

Spring Namespace:

```
<shardingsphere:data-source id="xxx" data-source-names="xxx" rule-refs="xxx">
  <props>
    <prop key="xa-transaction-manager-type">Atomikos</prop>
  </props>
</shardingsphere:data-source>
```


配置 Atomikos

可以通过在项目的 classpath 中添加 `jta.properties` 来定制化 Atomikos 配置项。

详情请参见 [Atomikos 官方文档](#)。

数据恢复

在项目的 `logs` 目录中会生成 `xa_tx.log`, 这是 XA 崩溃恢复时所需的日志, 请勿删除。

Narayana 事务

背景信息

Apache ShardingSphere 提供 XA 事务, 集成了 Narayana 的实现。

前提条件

引入 Maven 依赖

```
<properties>
  <narayana.version>5.12.4.Final</narayana.version>
  <jboss-transaction-spi.version>7.6.0.Final</jboss-transaction-spi.version>
  <jboss-logging.version>3.2.1.Final</jboss-logging.version>
</properties>

<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-jdbc-core</artifactId>
  <version>${shardingsphere.version}</version>
</dependency>

<!-- 使用 XA 事务时, 需要引入此模块 -->
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-transaction-xa-core</artifactId>
  <version>${shardingsphere.version}</version>
</dependency>

<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-transaction-xa-narayana</artifactId>
  <version>${shardingsphere.version}</version>
</dependency>
<dependency>
  <groupId>org.jboss.narayana.jta</groupId>
```

```

        <artifactId>jta</artifactId>
        <version>${narayana.version}</version>
</dependency>
<dependency>
    <groupId>org.jboss.narayana.jts</groupId>
    <artifactId>narayana-jts-integration</artifactId>
    <version>${narayana.version}</version>
</dependency>
<dependency>
    <groupId>org.jboss</groupId>
    <artifactId>jboss-transaction-spi</artifactId>
    <version>${jboss-transaction-spi.version}</version>
</dependency>
<dependency>
    <groupId>org.jboss.logging</groupId>
    <artifactId>jboss-logging</artifactId>
    <version>${jboss-logging.version}</version>
</dependency>

```

操作步骤

1. 配置 Narayana
2. 设置 XA 事务类型

配置示例

配置 Narayana

可以通过在项目的 classpath 中添加 jbossts-properties.xml 来定制化 Narayana 配置项。
详情请参见 [Narayana 官方文档](#)。

设置 XA 事务类型

Yaml:

```

- !TRANSACTION
  defaultType: XA
  providerType: Narayana

```

SpringBoot:

```

spring:
  shardingsphere:

```

```
props:
  xa-transaction-manager-type: Narayana
```

Spring Namespace:

```
<shardingsphere:data-source id="xxx" data-source-names="xxx" rule-refs="xxx">
  <props>
    <prop key="xa-transaction-manager-type">Narayana</prop>
  </props>
</shardingsphere:data-source>
```

Bitronix 事务

背景信息

Apache ShardingSphere 提供 XA 事务，集成了 Bitronix 的实现。

前提条件

引入 Maven 依赖

```
<properties>
  <btm.version>2.1.3</btm.version>
</properties>

<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-jdbc-core</artifactId>
  <version>${shardingsphere.version}</version>
</dependency>

<!-- 使用 XA 事务时，需要引入此模块 -->
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-transaction-xa-core</artifactId>
  <version>${shardingsphere.version}</version>
</dependency>

<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-transaction-xa-bitronix</artifactId>
  <version>${shardingsphere.version}</version>
</dependency>

<dependency>
  <groupId>org.codehaus.btm</groupId>
```

```
<artifactId>btm</artifactId>
<version>${btm.version}</version>
</dependency>
```

操作步骤

1. 配置 XA 事务类型
2. 配置 Bitronix

配置示例

配置 XA 事务类型

Yaml:

```
- !TRANSACTION
  defaultType: XA
  providerType: Bitronix
```

SpringBoot:

```
spring:
  shardingsphere:
    props:
      xa-transaction-manager-type: Bitronix
```

Spring Namespace:

```
<shardingsphere:data-source id="xxx" data-source-names="xxx" rule-refs="xxx">
  <props>
    <prop key="xa-transaction-manager-type">Bitronix</prop>
  </props>
</shardingsphere:data-source>
```

配置 Bitronix（可省略）

详情请参见 [Bitronix 官方文档](#)。

Seata 事务

背景信息

Apache ShardingSphere 提供 BASE 事务，集成了 Seata 的实现。

操作步骤

1. 启动 Seata Server
2. 创建日志表
3. 添加 Seata 配置

配置示例

启动 Seata Server

按照 [seata-work-shop](#) 中的步骤，下载并启动 Seata 服务器。

创建 undo_log 表

在每一个分片数据库实例中执创建 undo_log 表（以 MySQL 为例）。

```
CREATE TABLE IF NOT EXISTS `undo_log`
(
  `id`          BIGINT(20)  NOT NULL AUTO_INCREMENT COMMENT 'increment id',
  `branch_id`   BIGINT(20)  NOT NULL COMMENT 'branch transaction id',
  `xid`         VARCHAR(100) NOT NULL COMMENT 'global transaction id',
  `context`     VARCHAR(128) NOT NULL COMMENT 'undo_log context,such as
serialization',
  `rollback_info` LONGBLOB  NOT NULL COMMENT 'rollback info',
  `log_status`   INT(11)     NOT NULL COMMENT '0:normal status,1:defense status',
  `log_created`  DATETIME    NOT NULL COMMENT 'create datetime',
  `log_modified` DATETIME    NOT NULL COMMENT 'modify datetime',
  PRIMARY KEY (`id`),
  UNIQUE KEY `ux_undo_log` (`xid`, `branch_id`)
) ENGINE = InnoDB
  AUTO_INCREMENT = 1
  DEFAULT CHARSET = utf8 COMMENT ='AT transaction mode undo table';
```

修改配置

在 classpath 中增加 seata.conf 文件。

```
client {  
    application.id = example    ## 应用唯一主键  
    transaction.service.group = my_test_tx_group    ## 所属事务组  
}
```

根据实际场景修改 Seata 的 file.conf 和 registry.conf 文件。

4.1.6 不支持项

DataSource 接口

- 不支持 timeout 相关操作。

Connection 接口

- 不支持存储过程，函数，游标的操作；
- 不支持执行 native SQL；
- 不支持 savepoint 相关操作；
- 不支持 Schema/Catalog 的操作；
- 不支持自定义类型映射。

Statement 和 PreparedStatement 接口

- 不支持返回多结果集的语句（即存储过程，非 SELECT 多条数据）；
- 不支持国际化字符的操作。

ResultSet 接口

- 不支持对于结果集指针位置判断；
- 不支持通过非 next 方法改变结果指针位置；
- 不支持修改结果集内容；
- 不支持获取国际化字符；
- 不支持获取 Array。

JDBC 4.1

- 不支持 JDBC 4.1 接口新功能。

查询所有未支持方法，请阅读 `org.apache.shardingsphere.driver.jdbc.unsupported` 包。

4.2 ShardingSphere-Proxy

配置是 ShardingSphere-Proxy 中唯一与开发者交互的模块，通过它可以快速清晰的理解 ShardingSphere-Proxy 所提供的功能。

本章节是 ShardingSphere-Proxy 的配置参考手册，需要时可当做字典查阅。

ShardingSphere-Proxy 提供基于 YAML 的配置方式，并使用 DistSQL 进行交互。通过配置，应用开发者可以灵活的使用数据分片、读写分离、数据加密、影子库等功能，并且能够叠加使用。

规则配置部分与 ShardingSphere-JDBC 的 YAML 配置完全一致。DistSQL 与 YAML 配置能够相互取代。

更多使用细节请参见[使用示例](#)。

4.2.1 启动手册

本章节将介绍 ShardingSphere-Proxy 相关部署和启动等相关操作。

使用二进制发布包

背景信息

本节主要介绍如何通过二进制发布包启动 ShardingSphere-Proxy。

前提条件

使用二进制发布包启动 Proxy，需要环境具备 Java JRE 8 或更高版本。

操作步骤

1. 获取 ShardingSphere-Proxy 二进制发布包

在[下载页面](#)获取。

2. 配置 `conf/server.yaml`

ShardingSphere-Proxy 运行模式在 `server.yaml` 中配置，配置格式与 ShardingSphere-JDBC 一致，请参考[模式配置](#)。

其他配置项请参考：[* 权限配置](#) [* 属性配置](#)

3. 配置 `conf/config-*.yaml`

修改 `conf` 目录下以 `config-` 前缀开头的文件，如：`conf/config-sharding.yaml` 文件，进行分片规则、读写分离规则配置。配置方式请参考[配置手册](#)。`config-*.yaml` 文件的 `*` 部分可以任意命名。ShardingSphere-Proxy 支持配置多个逻辑数据源，每个以 `config-` 前缀命名的 YAML 配置文件，即为一个逻辑数据源。

4. （可选）引入数据库驱动

如果后端连接 PostgreSQL 或 openGauss 数据库，不需要引入额外依赖。

如果后端连接 MySQL 数据库，请下载 `mysql-connector-java-5.1.47.jar` 或者 `mysql-connector-java-8.0.11.jar`，并将其放入 `ext-lib` 目录。

5. （可选）引入集群模式所需依赖

ShardingSphere-Proxy 默认集成 ZooKeeper Curator 客户端，集群模式使用 ZooKeeper 无须引入其他依赖。

如果集群模式使用 Etcd，需要将 Etcd 的客户端驱动程序 `jetcd-core 0.5.0` 复制至目录 `ext-lib`。

6. （可选）引入分布式事务所需依赖

与 ShardingSphere-JDBC 使用方式相同。具体可参考[分布式事务](#)。

7. （可选）引入自定义算法

当用户需要使用自定义的算法类时，可通过以下方式配置使用自定义算法，以分片为例：

1. 实现 ``ShardingAlgorithm`` 接口定义的算法实现类。
2. 在项目 ``resources`` 目录下创建 ``META-INF/services`` 目录。
3. 在 ``META-INF/services`` 目录下新建文件 ``org.apache.shardingsphere.sharding.spi.ShardingAlgorithm``
4. 将实现类的全限定类名写入至文件 ``org.apache.shardingsphere.sharding.spi.ShardingAlgorithm``
5. 将上述 Java 文件打包成 jar 包。
6. 将上述 jar 包拷贝至 ``ext-lib`` 目录。
7. 将上述自定义算法实现类的 Java 文件引用配置在 YAML 文件中，具体可参考 [\[配置规则\]](https://shardingsphere.apache.org/document/current/cn/user-manual/shardingsphere-proxy/yaml-config/) (<https://shardingsphere.apache.org/document/current/cn/user-manual/shardingsphere-proxy/yaml-config/>)。

8. 启动 ShardingSphere-Proxy

Linux/macOS 操作系统请运行 `bin/start.sh`，Windows 操作系统请运行 `bin/start.bat` 启动 ShardingSphere-Proxy。默认监听端口 3307，默认配置目录为 Proxy 内的 `conf` 目录。启动脚本可以指定监听端口、配置文件所在目录，命令如下：

```
bin/start.sh [port] [/path/to/conf]
```

9. 使用客户端连接 ShardingSphere-Proxy

执行 MySQL / PostgreSQL / openGauss 的客户端命令直接操作 ShardingSphere-Proxy 即可。

使用 MySQL 客户端连接 ShardingSphere-Proxy：

```
mysql -h${proxy_host} -P${proxy_port} -u${proxy_username} -p${proxy_password}
```


使用 PostgreSQL 客户端连接 ShardingSphere-Proxy:

```
psql -h ${proxy_host} -p ${proxy_port} -U ${proxy_username}
```

使用 openGauss 客户端连接 ShardingSphere-Proxy:

```
gsql -r -h ${proxy_host} -p ${proxy_port} -U ${proxy_username} -W ${proxy_password}
```

配置示例

完整配置请参考 ShardingSphere 仓库中的示例: <https://github.com/apache/shardingsphere/tree/master/examples/shardingsphere-proxy-example>

使用 Docker

背景信息

本节主要介绍如何通过 Docker 启动 ShardingSphere-Proxy。

注意事项

使用 Docker 启动 ShardingSphere-Proxy 无须额外依赖。

操作步骤

1. 获取 Docker 镜像

- 方式一（推荐）：从 DockerHub 获取

```
docker pull apache/shardingsphere-proxy
```

- 方式二：获取 master 分支最新镜像: <https://github.com/apache/shardingsphere/pkgs/container/shardingsphere-proxy>
- 方式三：自行构建镜像

```
git clone https://github.com/apache/shardingsphere
mvn clean install
cd shardingsphere-distribution/shardingsphere-proxy-distribution
mvn clean package -Prelease,docker
```

如果遇到以下问题，请确保 Docker daemon 进程已经运行。

```
I/O exception (java.io.IOException) caught when processing request to {}->unix://
localhost:80: Connection refused?
```

2. 配置 conf/server.yaml 和 conf/config-*.yaml

可以从 Docker 容器中获取配置文件模板，拷贝到宿主机任意目录中：

```
docker run -d --name tmp --entrypoint=bash apache/shardingsphere-proxy
docker cp tmp:/opt/shardingsphere-proxy/conf /host/path/to/conf
docker rm tmp
```

由于容器内的网络环境可能与宿主机的网络环境有差异，如果启动时报无法连接到数据库错误等错误，请确保 conf/config-*.yaml 配置文件中指定的数据库的 IP 可以被 Docker 容器内部访问到。

具体配置请参考 [ShardingSphere-Proxy 启动手册 - 使用二进制发布包](#)。

3. （可选）引入第三方依赖或自定义算法

如果存在以下任意需求：* ShardingSphere-Proxy 后端使用 MySQL 数据库；* 使用自定义算法；* 使用 Etcd 作为集群模式的注册中心。

请在宿主机中任意位置创建 ext-lib 目录，并参考 [ShardingSphere-Proxy 启动手册 - 使用二进制发布包](#)中的对应步骤。

4. 启动 ShardingSphere-Proxy 容器

将宿主机中的 conf 与 ext-lib 目录挂载到容器中，启动容器：

```
docker run -d \
  -v /host/path/to/conf:/opt/shardingsphere-proxy/conf \
  -v /host/path/to/ext-lib:/opt/shardingsphere-proxy/ext-lib \
  -e PORT=3308 -p13308:3308 apache/shardingsphere-proxy:latest
```

其中，ext-lib 非必需，用户可按需挂载。ShardingSphere-Proxy 默认端口 3307，可以通过环境变量 -e PORT 指定。自定义 JVM 相关参数可通过环境变量 JVM_OPTS 设置。

5. 使用客户端连接 ShardingSphere-Proxy

请参考 [ShardingSphere-Proxy 启动手册 - 使用二进制发布包](#)。

配置示例

完整配置请参考 ShardingSphere 仓库中的示例：<https://github.com/apache/shardingsphere/tree/master/examples/shardingsphere-proxy-example>

使用 Helm

背景信息

使用 Helm 在 Kubernetes 集群中引导 ShardingSphere-Proxy 实例进行安装。

前提条件

- kubernetes 1.18+
- kubectl
- helm 3.3.0+
- 可以动态申请 PV(Persistent Volumes) 的 StorageClass 用于持久化数据。(可选)

操作步骤

在线安装

1. 将 ShardingSphere-Proxy 添加到 Helm 本地仓库:

```
helm repo add shardingsphere https://shardingsphere.apache.org/charts
```

1. 以 ShardingSphere-Proxy 命名安装 charts:

```
helm install shardingsphere-proxy shardingsphere/shardingsphere-proxy
```

源码安装

1. 执行下述命令以执行默认配置进行安装。

```
cd shardingsphere-proxy/charts/governance
helm dependency build
cd ../../
helm dependency build
cd ..
helm install shardingsphere-proxy shardingsphere-proxy
```

1. 其他的配置详见下方的配置列表。
2. 执行 `helm list` 获取所有安装的 release。

卸载

1. 默认删除所有发布记录, 增加 `--keep-history` 参数保留发布记录。

```
helm uninstall shardingsphere-proxy
```

参数解释

治理节点配置项

配置项	描述	值
governance.enabled	用来切换是否使用治理节点的 chart	true

治理节点 ZooKeeper 配置项

配置项	描述	值
governance.zookeeper.enabled	用来切换是否使用 ZooKeeper 的 chart	“true”
governance.zookeeper.replicaCount	ZooKeeper 节点数量	1
governance.zookeeper.persistence.enabled	标识 ZooKeeper 是否使用持久卷申领 (PersistentVolumeClaim) 用来申请持久卷 (PersistentVolume)	false
governance.zookeeper.persistence.storageClass	持久卷 (PersistentVolume) 的存储类 (Storage-Class)	""
governance.zookeeper.persistence.accessModes	持久卷 (PersistentVolume) 的访问模式	["ReadWriteOnce"]
governance.zookeeper.persistence.size	持久卷 (PersistentVolume) 大小	“8Gi”
governance.zookeeper.resources.limits	ZooKeeper 容器的资源限制	{}
governance.zookeeper.resources.requests.memory	ZooKeeper 容器申请的内存	256Mi
governance.zookeeper.resources.requests.cpu	ZooKeeper 容器申请的 cpu 核数	“250m”

计算节点 ShardingSphere-Proxy 配置项

配置项	描述	值
compute.image.repository	ShardingSphere-Proxy 的镜像名	apache/shardingsphere-proxy
compute.image.pullPolicy	ShardingSphere-Proxy 镜像拉取策略	IfNotPresent
compute.image.tag	ShardingSphere-Proxy 镜像标签	5.1.2
compute.imagePullSecrets	拉取私有仓库的凭证	[]
compute.resources.limits	ShardingSphere-Proxy 容器的资源限制	{}
compute.resources.requests.memory	ShardingSphere-Proxy 容器申请的内存	2Gi
compute.resources.requests.cpu	ShardingSphere-Proxy 容器申请的 cpu 核数	200m
compute.replicas	ShardingSphere-Proxy 节点个数	3
compute.service.type	ShardingSphere-Proxy 网络模式	ClusterIP
compute.service.port	ShardingSphere-Proxy 暴露端口	3307
compute.mysqlConnector.version	MySQL 驱动版本	5.1.49
compute.startPort	ShardingSphere-Proxy 启动端口	3307
compute.serverConfig	ShardingSphere-Proxy 模式配置文件	""

配置示例

```
#
# Licensed to the Apache Software Foundation (ASF) under one or more
# contributor license agreements. See the NOTICE file distributed with
# this work for additional information regarding copyright ownership.
# The ASF licenses this file to You under the Apache License, Version 2.0
# (the "License"); you may not use this file except in compliance with
# the License. You may obtain a copy of the License at
#
#     http://www.apache.org/licenses/LICENSE-2.0
#
# Unless required by applicable law or agreed to in writing, software
# distributed under the License is distributed on an "AS IS" BASIS,
# WITHOUT WARRANTIES OR CONDITIONS OF ANY KIND, either express or implied.
# See the License for the specific language governing permissions and
# limitations under the License.
#
## @section Governance-Node parameters
```

```

## @param governance.enabled Switch to enable or disable the governance helm chart
##
governance:
  enabled: true
  ## @section Governance-Node ZooKeeper parameters
  zookeeper:
    ## @param governance.zookeeper.enabled Switch to enable or disable the
    ZooKeeper helm chart
    ##
    enabled: true
    ## @param governance.zookeeper.replicaCount Number of ZooKeeper nodes
    ##
    replicaCount: 1
    ## ZooKeeper Persistence parameters
    ## ref: https://kubernetes.io/docs/user-guide/persistent-volumes/
    ## @param governance.zookeeper.persistence.enabled Enable persistence on
    ZooKeeper using PVC(s)
    ## @param governance.zookeeper.persistence.storageClass Persistent Volume
    storage class
    ## @param governance.zookeeper.persistence.accessModes Persistent Volume access
    modes
    ## @param governance.zookeeper.persistence.size Persistent Volume size
    ##
    persistence:
      enabled: false
      storageClass: ""
      accessModes:
        - ReadWriteOnce
      size: 8Gi
    ## ZooKeeper's resource requests and limits
    ## ref: https://kubernetes.io/docs/user-guide/compute-resources/
    ## @param governance.zookeeper.resources.limits The resources limits for the
    ZooKeeper containers
    ## @param governance.zookeeper.resources.requests.memory The requested memory
    for the ZooKeeper containers
    ## @param governance.zookeeper.resources.requests.cpu The requested cpu for the
    ZooKeeper containers
    ##
    resources:
      limits: {}
      requests:
        memory: 256Mi
        cpu: 250m

## @section Compute-Node parameters
##
compute:
  ## @section Compute-Node ShardingSphere-Proxy parameters

```

```

## ref: https://kubernetes.io/docs/concepts/containers/images/
## @param compute.image.repository Image name of ShardingSphere-Proxy.
## @param compute.image.pullPolicy The policy for pulling ShardingSphere-Proxy
image
## @param compute.image.tag ShardingSphere-Proxy image tag
##
image:
  repository: "apache/shardingsphere-proxy"
  pullPolicy: IfNotPresent
  ## Overrides the image tag whose default is the chart appVersion.
  ##
  tag: "5.1.2"
## @param compute.imagePullSecrets Specify docker-registry secret names as an
array
## e.g:
## imagePullSecrets:
##   - name: myRegistryKeySecretName
##
imagePullSecrets: []
## ShardingSphere-Proxy resource requests and limits
## ref: https://kubernetes.io/docs/concepts/configuration/manage-resources-
containers/
## @param compute.resources.limits The resources limits for the ShardingSphere-
Proxy containers
## @param compute.resources.requests.memory The requested memory for the
ShardingSphere-Proxy containers
## @param compute.resources.requests.cpu The requested cpu for the
ShardingSphere-Proxy containers
##
resources:
  limits: {}
  requests:
    memory: 2Gi
    cpu: 200m
## ShardingSphere-Proxy Deployment Configuration
## ref: https://kubernetes.io/docs/concepts/workloads/controllers/deployment/
## ref: https://kubernetes.io/docs/concepts/services-networking/service/
## @param compute.replicas Number of cluster replicas
##
replicas: 3
## @param compute.service.type ShardingSphere-Proxy network mode
## @param compute.service.port ShardingSphere-Proxy expose port
##
service:
  type: ClusterIP
  port: 3307
## MySQL connector Configuration
## ref: https://shardingsphere.apache.org/document/current/en/quick-start/

```

```

shardingsphere-proxy-quick-start/
  ## @param compute.mysqlConnector.version MySQL connector version
  ##
  mysqlConnector:
    version: "5.1.49"
  ## @param compute.startPort ShardingSphere-Proxy start port
  ## ShardingSphere-Proxy start port
  ## ref: https://shardingsphere.apache.org/document/current/en/user-manual/shardingsphere-proxy/startup/docker/
  ##
  startPort: 3307
  ## @section Compute-Node ShardingSphere-Proxy ServerConfiguration parameters
  ## NOTE: If you use the sub-charts to deploy Zookeeper, the server-lists field must be "{{ printf \"%s-zookeeper.%s:2181\" .Release.Name .Release.Namespace }}",
  ## otherwise please fill in the correct zookeeper address
  ## The server.yaml is auto-generated based on this parameter.
  ## If it is empty, the server.yaml is also empty.
  ## ref: https://shardingsphere.apache.org/document/current/en/user-manual/shardingsphere-jdbc/yaml-config/mode/
  ## ref: https://shardingsphere.apache.org/document/current/en/user-manual/common-config/builtin-algorithm/metadata-repository/
  ##
  serverConfig:
    ## @section Compute-Node ShardingSphere-Proxy ServerConfiguration authority parameters
    ## NOTE: It is used to set up initial user to login compute node, and authority data of storage node.
    ## ref: https://shardingsphere.apache.org/document/current/en/user-manual/shardingsphere-proxy/yaml-config/authentication/
    ## @param compute.serverConfig.authority.privilege.type authority provider for storage node, the default value is ALL_PERMITTED
    ## @param compute.serverConfig.authority.users[0].password Password for compute node.
    ## @param compute.serverConfig.authority.users[0].user Username,authorized host for compute node. Format: <username>@<hostname> hostname is % or empty string means do not care about authorized host
    ##
    authority:
      privilege:
        type: ALL_PRIVILEGES_PERMITTED
      users:
        - password: root
          user: root@%
    ## @section Compute-Node ShardingSphere-Proxy ServerConfiguration mode Configuration parameters
    ## @param compute.serverConfig.mode.type Type of mode configuration. Now only support Cluster mode
    ## @param compute.serverConfig.mode.repository.props.namespace Namespace of

```



```

registry center
  ## @param compute.serverConfig.mode.repository.props.server-lists Server lists
of registry center
  ## @param compute.serverConfig.mode.repository.props.maxRetries Max retries of
client connection
  ## @param compute.serverConfig.mode.repository.props.
operationTimeoutMilliseconds Milliseconds of operation timeout
  ## @param compute.serverConfig.mode.repository.props.retryIntervalMilliseconds
Milliseconds of retry interval
  ## @param compute.serverConfig.mode.repository.props.timeToLiveSeconds Seconds
of ephemeral data live
  ## @param compute.serverConfig.mode.repository.type Type of persist repository.
Now only support ZooKeeper
  ## @param compute.serverConfig.mode.override Whether overwrite persistent
configuration with local configuration
  ##
  mode:
    type: Cluster
  repository:
    type: ZooKeeper
  props:
    maxRetries: 3
    namespace: governance_ds
    operationTimeoutMilliseconds: 5000
    retryIntervalMilliseconds: 500
    server-lists: "{{ printf \"%s-zookeeper.%s:2181\" .Release.Name .Release.
Namespace }}"
    timeToLiveSeconds: 60
  overwrite: true

```

4.2.2 YAML 配置

ShardingSphere-JDBC 的 YAML 配置是 ShardingSphere-Proxy 的子集。在 `server.yaml` 文件中, ShardingSphere-Proxy 能够额外配置权限功能和更多的 Proxy 专有属性。

本章节将介绍 ShardingSphere-Proxy 的 YAML 额外配置。

权限

权限配置用于设置能够连接到 ShardingSphere-Proxy 的用户，并可以为他们授予不同的权限。

背景信息

在 ShardingSphere-Proxy 中，通过全局规则 Authority Rule（标识为!AUTHORITY）来配置用户和授权信息。

得益于 ShardingSphere 的可插拔架构，Proxy 提供了两种级别的权限提供者，分别是：

- ALL_PERMITTED：授予所有权限，不鉴权；
- DATABASE_PERMITTED：为用户授予指定逻辑库的权限，通过 user-database-mappings 进行映射。

在配置 Authority Rule 时，管理员可根据需要选择使用哪一种权限提供者。

参数解释

```
rules:
- !AUTHORITY
  users:
    - # 用于登录计算节点的用户名，授权主机和密码的组合。格式：<username>@<hostname>:
      <password>, hostname 为 % 或空字符串表示不限制授权主机
  provider:
    type: # 存储节点数据授权的权限提供者类型，缺省值为 ALL_PERMITTED
```

配置示例

ALL_PERMITTED

```
rules:
- !AUTHORITY
  users:
    - root@localhost:root
    - my_user@:pwd
  provider:
    type: ALL_PERMITTED
```

以上配置表示：- 用户 root，仅可从 localhost 连接 Proxy，密码为 root；- 用户 my_user，可以从任意主机连接 Proxy，密码为 pwd；- provider 类型为 ALL_PERMITTED，表示对用户授予所有权限，不鉴权。

DATABASE_PERMITTED

```
rules:
- !AUTHORITY
  users:
    - root@localhost:root
    - my_user@:pwd
  provider:
    type: DATABASE_PERMITTED
    props:
      user-database-mappings: root@localhost=sharding_db, root@localhost=test_db,
my_user@=sharding_db
```

以上配置表示：

- provider 类型为 DATABASE_PERMITTED，表示对用户授予库级别权限，需要配置；
- 用户 root 仅可从 localhost 主机连接，可访问 sharding_db 和 test_db；
- 用户 my_user 可从任意主机连接，可访问 sharding_db。

相关参考

权限提供者具体实现可以参考 [权限提供者](#)。

属性配置

背景信息

Apache ShardingSphere 提供属性配置的方式配置系统级配置。本节介绍 server.yaml 中的配置项。

参数解释

属性配置可以通过 `DistSQL#RAL` 修改。支持动态修改的属性可以立即生效，不支持动态修改的属性需要重启后生效。

配置示例

完整配置示例请参考 ShardingSphere 仓库内的 server.yaml：<https://github.com/apache/shardingsphere/blob/aac0d3026e00575114701be603ec189a02a45747/shardingsphere-proxy/shardingsphere-proxy-bootstrap/src/main/resources/conf/server.yaml#L71-L93>

规则配置

背景信息

本节介绍如何进行 ShardingSphere-Proxy 的规则配置。

参数解释

ShardingSphere-Proxy 的规则配置与 ShardingSphere-JDBC 一致，具体规则请参考 [ShardingSphere-JDBC 规则配置](#)。

注意事项

与 ShardingSphere-JDBC 不同的是，以下规则需要配置在 ShardingSphere-Proxy 的 `server.yaml` 中：

- SQL 解析
- 分布式事务
- SQL 翻译

4.2.3 DistSQL

本章节将介绍 DistSQL 的详细语法。

定义

DistSQL (Distributed SQL) 是 Apache ShardingSphere 特有的操作语言。它与标准 SQL 的使用方式完全一致，用于提供增量功能的 SQL 级别操作能力。

灵活的规则配置和资源管控能力是 Apache ShardingSphere 的特点之一。

在使用 4.x 及其之前版本时，开发者虽然可以像使用原生数据库一样操作数据，但却需要通过本地文件或注册中心配置资源和规则。然而，操作习惯变更，对于运维工程师并不友好。

从 5.x 版本开始，DistSQL (Distributed SQL) 让用户可以像操作数据库一样操作 Apache ShardingSphere，使其从面向开发人员的框架和中间件转变为面向运维人员的数据库产品。

相关概念

DistSQL 细分为 RDL、RQL、RAL 和 RUL 四种类型。

RDL

Resource & Rule Definition Language, 负责资源和规则的创建、修改和删除。

RQL

Resource & Rule Query Language, 负责资源和规则的查询和展现。

RAL

Resource & Rule Administration Language, 负责强制路由、熔断、配置导入导出、数据迁移控制等管理功能。

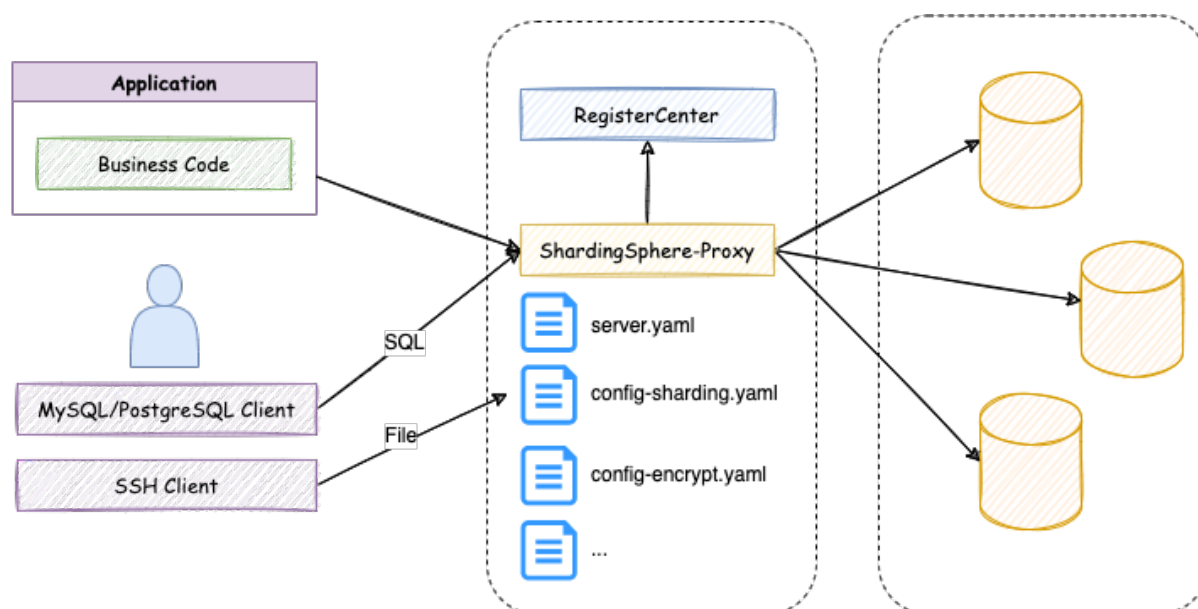
RUL

Resource & Rule Utility Language, 负责 SQL 解析、SQL 格式化、执行计划预览等功能。

对系统的影响

之前

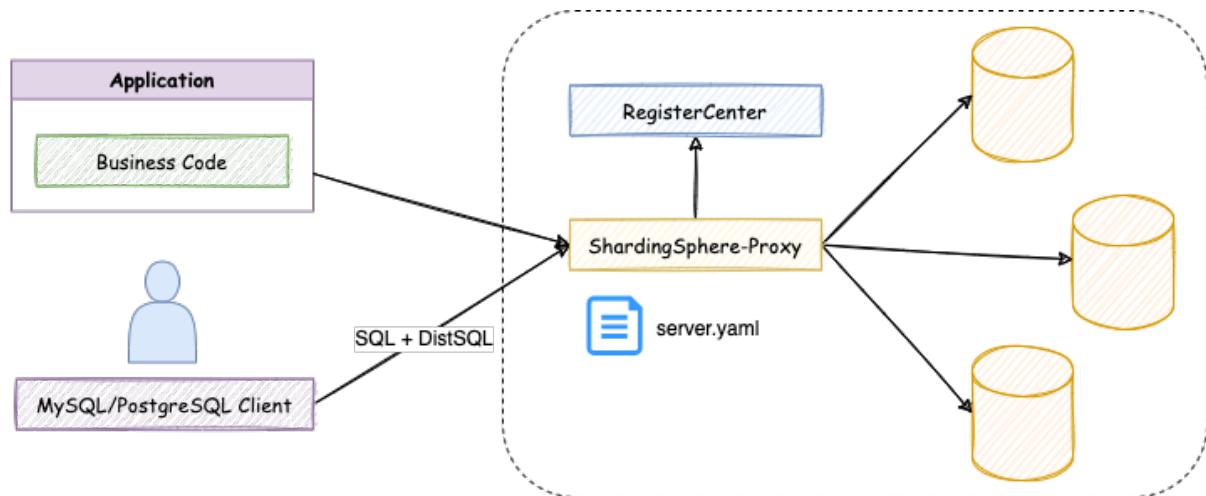
在拥有 DistSQL 以前, 用户一边使用 SQL 语句操作数据, 一边使用 YAML 文件来管理 ShardingSphere 的配置, 如下图:



这时用户不得不面对以下几个问题: - 需要通过不同类型的客户端来操作数据和管理 ShardingSphere 规则; - 多个逻辑库需要多个 YAML 文件; - 修改 YAML 需要文件的编辑权限; - 修改 YAML 后需要重启 ShardingSphere。

之后

随着 DistSQL 的出现，对 ShardingSphere 的操作方式也得到了改变：



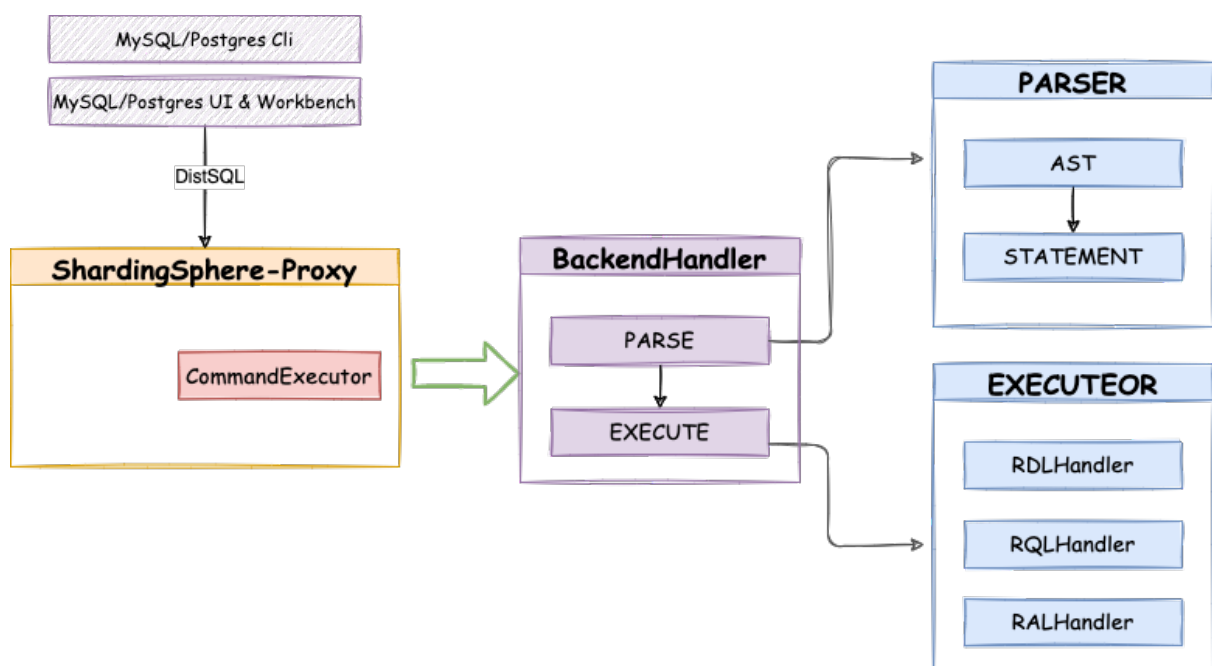
现在，用户的使用体验得到了巨大改善：- 使用相同的客户端来管理数据和 ShardingSphere 配置；- 不再额外创建 YAML 文件，通过 DistSQL 管理逻辑库；- 不再需要文件的编辑权限，通过 DistSQL 来管理配置；- 配置的变更实时生效，无需重启 ShardingSphere。

使用限制

DistSQL 只能用于 ShardingSphere-Proxy，ShardingSphere-JDBC 暂不提供。

原理介绍

与标准 SQL 一样，DistSQL 由 ShardingSphere 的解析引擎进行识别，将输入语句转换为抽象语法树，进而生成各个语法对应的 Statement，最后由合适的 Handler 进行业务处理。整体流程如下图所示：



相关参考

用户手册：DistSQL

语法

本章节将对 DistSQL 的语法进行详细说明，并以实际的例子介绍 DistSQL 的使用。

语法规则

在 DistSQL 语句中，除关键字外，其余元素的输入格式应符合以下规则。

标识符

1. 标识符代表 SQL 语句中的一个对象，包括：
 - 数据库名称
 - 表名
 - 列名
 - 索引名称
 - 资源名称
 - 规则名称
 - 算法名称
2. 标识符中允许使用的字符有：[a-z,A-Z,0-9, _]（字母、数字、下划线），且应以字母开头。
3. 当标识符中出现关键字或特殊字符时，使用反引号 (`)。

字面量

字面量包括字符串、整数值和布尔值：

- 字符串：是由单引号 (`) 或双引号 (") 括起来的字符序列；
- 整数值：一般为正整数，如 0-9；

说明：部分 DistSQL 语法允许负值，此时可在数字前加负号 (-)，如 -1。

- 布尔值：TRUE 或 FALSE，大小写不敏感。

RDL 语法

RDL (Resource & Rule Definition Language) 为 Apache ShardingSphere 的资源 and 规则定义语言。

资源定义

语法说明

```
ADD RESOURCE resourceDefinition [, resourceDefinition] ...

ALTER RESOURCE resourceDefinition [, resourceDefinition] ...

DROP RESOURCE resourceName [, resourceName] ... [ignore single tables]

resourceDefinition:
    simpleSource | urlSource

simpleSource:
    resourceName(HOST=hostname,PORT=port,DB=dbName,USER=user [,PASSWORD=password]
[,PROPERTIES(property [,property]) ...])

urlSource:
    resourceName(URL=url,USER=user [,PASSWORD=password] [,PROPERTIES(property [,
property]) ...])

property:
    key=value
```

参数解释

名称	数据类型	说明
resourceName	IDENTIFIER	资源名称
hostname	STRING	数据源地址
port	INT	数据源端口
dbName	STRING	物理库名称
url	STRING	URL 地址
user	STRING	用户名
password	STRING	密码

注意事项

- 添加资源前请确认已经创建分布式数据库，并执行 `use` 命令成功选择一个数据库；
- 确认将要添加或修改的资源是可以正常连接的，否则将不能操作成功；
- 不允许重复的 `resourceName`；
- `PROPERTIES` 用于自定义连接池参数，`key` 和 `value` 均为 `STRING` 类型；
- `ALTER RESOURCE` 修改资源时不允许改变该资源关联的真实数据源；
- `ALTER RESOURCE` 修改资源时会发生连接池的切换，此操作可能对进行中的业务造成影响，请谨慎使用；
- `DROP RESOURCE` 只会删除逻辑资源，不会删除真实的数据源；
- 被规则引用的资源将无法被删除；
- 若资源只被 `single table rule` 引用，且用户确认可以忽略该限制，则可以添加可选参数 `ignore single tables` 进行强制删除。

示例

```
ADD RESOURCE resource_0 (
    HOST="127.0.0.1",
    PORT=3306,
    DB="db0",
    USER="root",
    PASSWORD="root"
),resource_1 (
    HOST="127.0.0.1",
    PORT=3306,
    DB="db1",
    USER="root"
),resource_2 (
    HOST="127.0.0.1",
    PORT=3306,
    DB="db2",
    USER="root",
    PROPERTIES("maximumPoolSize"="10")
),resource_3 (
    URL="jdbc:mysql://127.0.0.1:3306/db3?serverTimezone=UTC&useSSL=false",
    USER="root",
    PASSWORD="root",
    PROPERTIES("maximumPoolSize"="10","idleTimeout"="30000")
);

ALTER RESOURCE resource_0 (
    HOST="127.0.0.1",
    PORT=3309,
```

```

    DB="db0",
    USER="root",
    PASSWORD="root"
),resource_1 (
    URL="jdbc:mysql://127.0.0.1:3309/db1?serverTimezone=UTC&useSSL=false",
    USER="root",
    PASSWORD="root",
    PROPERTIES("maximumPoolSize"="10","idleTimeout"="30000")
);

DROP RESOURCE resource_0, resource_1;
DROP RESOURCE resource_2, resource_3 ignore single tables;

```

规则定义

本章节将对规则定义的语法进行详细说明。

数据分片

语法说明

Sharding Table Rule

```

CREATE SHARDING TABLE RULE shardingTableRuleDefinition [,
shardingTableRuleDefinition] ...

ALTER SHARDING TABLE RULE shardingTableRuleDefinition [,
shardingTableRuleDefinition] ...

DROP SHARDING TABLE RULE tableName [, tableName] ...

CREATE DEFAULT SHARDING shardingScope STRATEGY (shardingStrategy)

ALTER DEFAULT SHARDING shardingScope STRATEGY (shardingStrategy)

DROP DEFAULT SHARDING shardingScope STRATEGY;

CREATE SHARDING ALGORITHM shardingAlgorithmDefinition [,
shardingAlgorithmDefinition] ...

ALTER SHARDING ALGORITHM shardingAlgorithmDefinition [,
shardingAlgorithmDefinition] ...

DROP SHARDING ALGORITHM algorithmName [, algorithmName] ...

```

```

CREATE SHARDING KEY GENERATOR keyGeneratorDefinition [, keyGeneratorDefinition] ...

ALTER SHARDING KEY GENERATOR keyGeneratorDefinition [, keyGeneratorDefinition] ...

DROP SHARDING KEY GENERATOR keyGeneratorName [, keyGeneratorName] ...

shardingTableRuleDefinition:
    shardingAutoTableRule | shardingTableRule

shardingAutoTableRule:
    tableName(resources, shardingColumn, algorithmDefinition [,
keyGenerateDeclaration])

shardingTableRule:
    tableName(dataNodes [, databaseStrategy] [, tableStrategy] [,
keyGenerateDeclaration])

resources:
    RESOURCES(resource [, resource] ...)

dataNodes:
    DATANODES(dataNode [, dataNode] ...)

resource:
    resourceName | inlineExpression

dataNode:
    resourceName | inlineExpression

shardingColumn:
    SHARDING_COLUMN=columnName

algorithmDefinition:
    TYPE(NAME=shardingAlgorithmType [, PROPERTIES([algorithmProperties]))

keyGenerateDeclaration:
    keyGenerateDefinition | keyGenerateConstruction

keyGenerateDefinition:
    KEY_GENERATE_STRATEGY(COLUMN=columnName, strategyDefinition)

shardingScope:
    DATABASE | TABLE

databaseStrategy:
    DATABASE_STRATEGY(shardingStrategy)

tableStrategy:

```

```

TABLE_STRATEGY(shardingStrategy)

keyGenerateConstruction
  KEY_GENERATE_STRATEGY(COLUMN=columnName, KEY_
GENERATOR=keyGenerateAlgorithmName)

shardingStrategy:
  TYPE=strategyType, shardingColumn, shardingAlgorithm

shardingAlgorithm:
  existingAlgorithm | autoCreativeAlgorithm

existingAlgorithm:
  SHARDING_ALGORITHM=shardingAlgorithmName

autoCreativeAlgorithm:
  SHARDING_ALGORITHM(algorithmDefinition)

strategyDefinition:
  TYPE(NAME=keyGenerateStrategyType [, PROPERTIES([algorithmProperties]))

shardingAlgorithmDefinition:
  shardingAlgorithmName(algorithmDefinition)

algorithmProperties:
  algorithmProperty [, algorithmProperty] ...

algorithmProperty:
  key=value

keyGeneratorDefinition:
  keyGeneratorName (algorithmDefinition)

```

- RESOURCES 需使用 RDL 管理的数据源资源；
- shardingAlgorithmType 指定自动分片算法类型，请参考 [自动分片算法](#)；
- keyGenerateStrategyType 指定分布式主键生成策略，请参考 [分布式主键](#)；
- 重复的 tableName 将无法被创建；
- shardingAlgorithm 能够被不同的 Sharding Table Rule 复用，因此在执行 DROP SHARDING TABLE RULE 时，对应的 shardingAlgorithm 不会被移除；
- 如需移除 shardingAlgorithm，请执行 DROP SHARDING ALGORITHM；
- strategyType 指定分片策略，请参考[分片策略](#)；
- Sharding Table Rule 同时支持 Auto Table 和 Table 两种类型，两者在语法上有所差异，对应配置文件请参考 [数据分片](#)；
- 使用 autoCreativeAlgorithm 方式指定 shardingStrategy 时，将会自动创建新的分

片算法, 算法命名规则为 `tableName_strategyType_shardingAlgorithmType`, 如 `t_order_database_inline`。

Sharding Binding Table Rule

```
CREATE SHARDING BINDING TABLE RULES bindTableRulesDefinition [,
bindTableRulesDefinition] ...

ALTER SHARDING BINDING TABLE RULES bindTableRulesDefinition [,
bindTableRulesDefinition] ...

DROP SHARDING BINDING TABLE RULES bindTableRulesDefinition [,
bindTableRulesDefinition] ...

bindTableRulesDefinition:
    (tableName [, tableName] ... )
```

- ALTER 会使用新的配置直接覆盖数据库内的绑定表配置

Sharding Broadcast Table Rule

```
CREATE SHARDING BROADCAST TABLE RULES (tableName [, tableName] ...)

ALTER SHARDING BROADCAST TABLE RULES (tableName [, tableName] ...)

DROP SHARDING BROADCAST TABLE RULES (tableName [, tableName] ...)
```

- ALTER 会使用新的配置直接覆盖数据库内的广播表配置

示例

Sharding Table Rule

Key Generator

```
CREATE SHARDING KEY GENERATOR snowflake_key_generator (
TYPE(NAME="SNOWFLAKE")
);

ALTER SHARDING KEY GENERATOR snowflake_key_generator (
TYPE(NAME="SNOWFLAKE")
);

DROP SHARDING KEY GENERATOR snowflake_key_generator;
```

Auto Table

```
CREATE SHARDING TABLE RULE t_order (
  RESOURCES(resource_0,resource_1),
  SHARDING_COLUMN=order_id,TYPE(NAME="hash_mod",PROPERTIES("sharding-count"="4")),
  KEY_GENERATE_STRATEGY(COLUMN=another_id,TYPE(NAME="snowflake"))
);

ALTER SHARDING TABLE RULE t_order (
  RESOURCES(resource_0,resource_1,resource_2,resource_3),
  SHARDING_COLUMN=order_id,TYPE(NAME="hash_mod",PROPERTIES("sharding-count"="16")),
  KEY_GENERATE_STRATEGY(COLUMN=another_id,TYPE(NAME="snowflake"))
);

DROP SHARDING TABLE RULE t_order;

DROP SHARDING ALGORITHM t_order_hash_mod;
```

Table

```
CREATE SHARDING ALGORITHM table_inline (
  TYPE(NAME="inline",PROPERTIES("algorithm-expression"="t_order_item_${order_id % 2}"))
);

CREATE SHARDING TABLE RULE t_order_item (
  DATANODES("resource_${0..1}.t_order_item_${0..1}"),
  DATABASE_STRATEGY(TYPE="standard",SHARDING_COLUMN=user_id,SHARDING_
  ALGORITHM(TYPE(NAME="inline",PROPERTIES("algorithm-expression"="resource_${user_id
  % 2}")))),
  TABLE_STRATEGY(TYPE="standard",SHARDING_COLUMN=order_id,SHARDING_ALGORITHM=table_
  inline),
  KEY_GENERATE_STRATEGY(COLUMN=another_id,KEY_GENERATOR=snowflake_key_generator)
);

ALTER SHARDING ALGORITHM database_inline (
  TYPE(NAME="inline",PROPERTIES("algorithm-expression"="resource_${user_id % 4}"))
),table_inline (
  TYPE(NAME="inline",PROPERTIES("algorithm-expression"="t_order_item_${order_id % 4}"))
);

ALTER SHARDING TABLE RULE t_order_item (
  DATANODES("resource_${0..3}.t_order_item_${0..3}"),
  DATABASE_STRATEGY(TYPE="standard",SHARDING_COLUMN=user_id,SHARDING_
  ALGORITHM=database_inline),
  TABLE_STRATEGY(TYPE="standard",SHARDING_COLUMN=order_id,SHARDING_ALGORITHM=table_
  inline),
  KEY_GENERATE_STRATEGY(COLUMN=another_id,KEY_GENERATOR=snowflake_key_generator)
);
```

```

DROP SHARDING TABLE RULE t_order_item;

DROP SHARDING ALGORITHM database_inline;

CREATE DEFAULT SHARDING DATABASE STRATEGY (
TYPE="standard",SHARDING_COLUMN=order_id,SHARDING_ALGORITHM=database_inline
);

ALTER DEFAULT SHARDING DATABASE STRATEGY (
TYPE="standard",SHARDING_COLUMN=another_id,SHARDING_ALGORITHM=database_inline
);

DROP DEFAULT SHARDING DATABASE STRATEGY;

```

Sharding Binding Table Rule

```

CREATE SHARDING BINDING TABLE RULES (t_order,t_order_item),(t_1,t_2);

ALTER SHARDING BINDING TABLE RULES (t_order,t_order_item);

DROP SHARDING BINDING TABLE RULES;

DROP SHARDING BINDING TABLE RULES (t_order,t_order_item);

```

Sharding Broadcast Table Rule

```

CREATE SHARDING BROADCAST TABLE RULES (t_b,t_a);

ALTER SHARDING BROADCAST TABLE RULES (t_b,t_a,t_3);

DROP SHARDING BROADCAST TABLE RULES;

DROP SHARDING BROADCAST TABLE RULES t_b;

```

单表

定义

```

CREATE DEFAULT SINGLE TABLE RULE singleTableRuleDefinition

ALTER DEFAULT SINGLE TABLE RULE singleTableRuleDefinition

```

```
DROP DEFAULT SINGLE TABLE RULE
```

```
singleTableRuleDefinition:
    RESOURCE = resourceName
```

- RESOURCE 需使用 RDL 管理的数据源资源。

示例

Single Table Rule

```
CREATE DEFAULT SINGLE TABLE RULE RESOURCE = ds_0
```

```
ALTER DEFAULT SINGLE TABLE RULE RESOURCE = ds_1
```

```
DROP DEFAULT SINGLE TABLE RULE
```

读写分离

语法说明

```
CREATE READWRITE_SPLITTING RULE readwriteSplittingRuleDefinition [,
readwriteSplittingRuleDefinition] ...
```

```
ALTER READWRITE_SPLITTING RULE readwriteSplittingRuleDefinition [,
readwriteSplittingRuleDefinition] ...
```

```
DROP READWRITE_SPLITTING RULE ruleName [, ruleName] ...
```

```
readwriteSplittingRuleDefinition:
    ruleName ([staticReadwriteSplittingRuleDefinition |
dynamicReadwriteSplittingRuleDefinition]
            [, loadBalancerDefinition])
```

```
staticReadwriteSplittingRuleDefinition:
    WRITE_RESOURCE=writeResourceName, READ_RESOURCES(readResourceName [,
readResourceName] ... )
```

```
dynamicReadwriteSplittingRuleDefinition:
    AUTO_AWARE_RESOURCE=autoAwareResourceName [, WRITE_DATA_SOURCE_QUERY_
ENABLED=writeDataSourceQueryEnabled]
```

```
loadBalancerDefinition:
    TYPE(NAME=loadBalancerType [, PROPERTIES([algorithmProperties] )) )
```



```
algorithmProperties:
    algorithmProperty [, algorithmProperty] ...

algorithmProperty:
    key=value

writeDataSourceQueryEnabled:
    TRUE | FALSE
```

参数解释

名称	数据类型	说明
ruleName	IDENTIFIER	规则名称
writeResourceName	IDENTIFIER	写库数据源名称
readResourceName	IDENTIFIER	读库数据源名称
autoAwareResourceName	IDENTIFIER	数据库发现的逻辑数据源名称
writeDataSourceQueryEnabled	BOOLEAN	读库全部下线，主库是否承担读流量
loadBalancerType	STRING	负载均衡算法类型

注意事项

- 支持创建静态读写分离规则和动态读写分离规则；
- 动态读写分离规则依赖于数据库发现规则；
- loadBalancerType 指定负载均衡算法类型，请参考 [负载均衡算法](#)；
- 重复的 ruleName 将无法被创建。

示例

```
// Static
CREATE READWRITE_SPLITTING RULE ms_group_0 (
WRITE_RESOURCE=write_ds,
READ_RESOURCES(read_ds_0,read_ds_1),
TYPE(NAME="random")
);

// Dynamic
CREATE READWRITE_SPLITTING RULE ms_group_1 (
AUTO_AWARE_RESOURCE=group_0,
WRITE_DATA_SOURCE_QUERY_ENABLED=false,
TYPE(NAME="random",PROPERTIES("read_weight"="2:1"))
);
```

```
ALTER READWRITE_SPLITTING RULE ms_group_1 (
WRITE_RESOURCE=write_ds,
READ_RESOURCES(read_ds_0,read_ds_1,read_ds_2),
TYPE(NAME="random",PROPERTIES("read_weight"="2:0"))
);

DROP READWRITE_SPLITTING RULE ms_group_1;
```

数据库发现

语法说明

```
CREATE DB_DISCOVERY RULE ruleDefinition [, ruleDefinition] ...

ALTER DB_DISCOVERY RULE ruleDefinition [, ruleDefinition] ...

DROP DB_DISCOVERY RULE ruleName [, ruleName] ...

CREATE DB_DISCOVERY TYPE databaseDiscoveryTypeDefinition [,
databaseDiscoveryTypeDefinition] ...

ALTER DB_DISCOVERY TYPE databaseDiscoveryTypeDefinition [,
databaseDiscoveryTypeDefinition] ...

DROP DB_DISCOVERY TYPE discoveryTypeName [, discoveryTypeName] ...

CREATE DB_DISCOVERY HEARTBEAT databaseDiscoveryHeartbeatDefinition [,
databaseDiscoveryHeartbeatDefinition] ...

ALTER DB_DISCOVERY HEARTBEAT databaseDiscoveryHeartbeatDefinition [,
databaseDiscoveryHeartbeatDefinition] ...

DROP DB_DISCOVERY HEARTBEAT discoveryHeartbeatName [, discoveryHeartbeatName] ...

ruleDefinition:
    (databaseDiscoveryRuleDefinition | databaseDiscoveryRuleConstruction)

databaseDiscoveryRuleDefinition
    ruleName (resources, typeDefinition, heartbeatDefinition)

databaseDiscoveryRuleConstruction
    ruleName (resources, TYPE = discoveryTypeName, HEARTBEAT =
discoveryHeartbeatName)

databaseDiscoveryTypeDefinition
    discoveryTypeName (typeDefinition)
```

```

databaseDiscoveryHeartbeatDefinition
    discoveryHeartbeatName (PROPERTIES (properties))

resources:
    RESOURCES(resourceName [, resourceName] ...)

typeDefinition:
    TYPE(NAME=typeName [, PROPERTIES([properties] )])

heartbeatDefinition
    HEARTBEAT (PROPERTIES (properties))

properties:
    property [, property] ...

property:
    key=value

```

参数解释

名称	数据类型	说明
discoveryTypeName	IDENTIFIER	数据库发现类型名
ruleName	IDENTIFIER	规则名称
discoveryHeartbeatName	IDENTIFIER	监听心跳名称
typeName	STRING	数据库发现类型，如：MySQL.MGR
resourceName	IDENTIFIER	资源名称

注意事项

- discoveryType 指定数据库发现服务类型，ShardingSphere 内置支持 MySQL.MGR；
- 重复的 ruleName 将无法被创建；
- 正在被使用的 discoveryType 和 discoveryHeartbeat 无法被删除；
- 带有 - 的命名在改动时需要使用 " "；
- 移除 discoveryRule 时不会移除被该 discoveryRule 使用的 discoveryType 和 discoveryHeartbeat。

示例

创建 **discoveryRule** 时同时创建 **discoveryType** 和 **discoveryHeartbeat**

```
CREATE DB_DISCOVERY RULE db_discovery_group_0 (
  RESOURCES(ds_0, ds_1, ds_2),
  TYPE(NAME='MySQL.MGR',PROPERTIES('group-name'='92504d5b-6dec')),
  HEARTBEAT(PROPERTIES('keep-alive-cron'='0/5 * * * * ?'))
);

ALTER DB_DISCOVERY RULE db_discovery_group_0 (
  RESOURCES(ds_0, ds_1, ds_2),
  TYPE(NAME='MySQL.MGR',PROPERTIES('group-name'='246e9612-aaf1')),
  HEARTBEAT(PROPERTIES('keep-alive-cron'='0/5 * * * * ?'))
);

DROP DB_DISCOVERY RULE db_discovery_group_0;

DROP DB_DISCOVERY TYPE db_discovery_group_0_mgr;

DROP DB_DISCOVERY HEARTBEAT db_discovery_group_0_heartbeat;
```

使用已有的 **discoveryType** 和 **discoveryHeartbeat** 创建 **discoveryRule**

```
CREATE DB_DISCOVERY TYPE db_discovery_group_1_mgr(
  TYPE(NAME='MySQL.MGR',PROPERTIES('group-name'='92504d5b-6dec'))
);

CREATE DB_DISCOVERY HEARTBEAT db_discovery_group_1_heartbeat(
  PROPERTIES('keep-alive-cron'='0/5 * * * * ?')
);

CREATE DB_DISCOVERY RULE db_discovery_group_1 (
  RESOURCES(ds_0, ds_1, ds_2),
  TYPE=db_discovery_group_1_mgr,
  HEARTBEAT=db_discovery_group_1_heartbeat
);

ALTER DB_DISCOVERY TYPE db_discovery_group_1_mgr(
  TYPE(NAME='MySQL.MGR',PROPERTIES('group-name'='246e9612-aaf1'))
);

ALTER DB_DISCOVERY HEARTBEAT db_discovery_group_1_heartbeat(
  PROPERTIES('keep-alive-cron'='0/10 * * * * ?')
);
```

```

ALTER DB_DISCOVERY RULE db_discovery_group_1 (
  RESOURCES(ds_0, ds_1),
  TYPE=db_discovery_group_1_mgr,
  HEARTBEAT=db_discovery_group_1_heartbeat
);

DROP DB_DISCOVERY RULE db_discovery_group_1;

DROP DB_DISCOVERY TYPE db_discovery_group_1_mgr;

DROP DB_DISCOVERY HEARTBEAT db_discovery_group_1_heartbeat;

```

数据加密

语法说明

```

CREATE ENCRYPT RULE encryptRuleDefinition [, encryptRuleDefinition] ...

ALTER ENCRYPT RULE encryptRuleDefinition [, encryptRuleDefinition] ...

DROP ENCRYPT RULE tableName [, tableName] ...

encryptRuleDefinition:
    tableName(COLUMNS(columnDefinition [, columnDefinition] ...), QUERY_WITH_
    CIPHER_COLUMN=queryWithCipherColumn)

columnDefinition:
    (NAME=columnName [, PLAIN=plainColumnName] , CIPHER=cipherColumnName,
    encryptAlgorithm)

encryptAlgorithm:
    TYPE(NAME=encryptAlgorithmType [, PROPERTIES([algorithmProperties] )) )

algorithmProperties:
    algorithmProperty [, algorithmProperty] ...

algorithmProperty:
    key=value

```

参数解释

名称	数据类型	说明
tableName	IDENTIFIER	表名称
columnName	IDENTIFIER	逻辑数据列名称
plainColumnName	IDENTIFIER	明文数据列名称
cipherColumnName	IDENTIFIER	加密数据列名称
encryptAlgorithmType	STRING	加密算法类型名称

注意事项

- PLAIN 指定明文数据列，CIPHER 指定密文数据列；
- encryptAlgorithmType 指定加密算法类型，请参考 [加密算法](#)；
- 重复的 tableName 将无法被创建；
- queryWithCipherColumn 支持大写或小写的 true 或 false。

示例

```
CREATE ENCRYPT RULE t_encrypt (
  COLUMNS(
    (NAME=user_id,PLAIN=user_plain,CIPHER=user_cipher,TYPE(NAME='AES',PROPERTIES('aes-key-value'='123456abc'))),
    (NAME=order_id, CIPHER =order_cipher,TYPE(NAME='MD5'))
  ),QUERY_WITH_CIPHER_COLUMN=true),
t_encrypt_2 (
  COLUMNS(
    (NAME=user_id,PLAIN=user_plain,CIPHER=user_cipher,TYPE(NAME='AES',PROPERTIES('aes-key-value'='123456abc'))),
    (NAME=order_id, CIPHER=order_cipher,TYPE(NAME='MD5'))
  ), QUERY_WITH_CIPHER_COLUMN=FALSE);

ALTER ENCRYPT RULE t_encrypt (
  COLUMNS(
    (NAME=user_id,PLAIN=user_plain,CIPHER=user_cipher,TYPE(NAME='AES',PROPERTIES('aes-key-value'='123456abc'))),
    (NAME=order_id,CIPHER=order_cipher,TYPE(NAME='MD5'))
  ), QUERY_WITH_CIPHER_COLUMN=TRUE);

DROP ENCRYPT RULE t_encrypt,t_encrypt_2;
```

影子库压测

语法说明

```
CREATE SHADOW RULE shadowRuleDefinition [, shadowRuleDefinition] ...

ALTER SHADOW RULE shadowRuleDefinition [, shadowRuleDefinition] ...

CREATE SHADOW ALGORITHM shadowAlgorithm [, shadowAlgorithm] ...

ALTER SHADOW ALGORITHM shadowAlgorithm [, shadowAlgorithm] ...

DROP SHADOW RULE ruleName [, ruleName] ...

DROP SHADOW ALGORITHM algorithmName [, algorithmName] ...

CREATE DEFAULT SHADOW ALGORITHM NAME = algorithmName

shadowRuleDefinition: ruleName(resourceMapping, shadowTableRule [, shadowTableRule]
...)

resourceMapping: SOURCE=resourceName, SHADOW=resourceName

shadowTableRule: tableName(shadowAlgorithm [, shadowAlgorithm] ...)

shadowAlgorithm: ([algorithmName, ] TYPE(NAME=shadowAlgorithmType,
PROPERTIES([algorithmProperties] ...)))

algorithmProperties: algorithmProperty [, algorithmProperty] ...

algorithmProperty: key=value
```

参数解释

名称	数据类型	说明
ruleName	IDENTIFIER	规则名称
resourceName	IDENTIFIER	数据库名称
tableName	IDENTIFIER	影子表名称
algorithmName	IDENTIFIER	影子算法名称
shadowAlgorithmType	STRING	影子算法类型

注意事项

- 重复的 ruleName 无法被创建;
- resourceMapping 指定源数据库和影子库的映射关系, 需使用 RDL 管理的 resource, 请参考 [数据源资源](#);
- shadowAlgorithm 可同时作用于多个 shadowTableRule;
- algorithmName 未指定时会根据 ruleName、tableName 和 shadowAlgorithmType 自动生成;
- shadowAlgorithmType 目前支持 VALUE_MATCH、REGEX_MATCH 和 SIMPLE_HINT;
- shadowTableRule 能够被不同的 shadowRuleDefinition 复用, 因此在执行 DROP SHADOW RULE 时, 对应的 shadowTableRule 不会被移除;
- shadowAlgorithm 能够被不同的 shadowTableRule 复用, 因此在执行 ALTER SHADOW RULE 时, 对应的 shadowAlgorithm 不会被移除。

示例

```
CREATE SHADOW RULE shadow_rule(
SOURCE=demo_ds,
SHADOW=demo_ds_shadow,
t_order((simple_hint_algorithm, TYPE(NAME="SIMPLE_HINT", PROPERTIES("shadow"="true", "foo"="bar"))), (TYPE(NAME="REGEX_MATCH", PROPERTIES("operation"="insert", "column"="user_id", "regex"='[1]')))),
t_order_item((TYPE(NAME="VALUE_MATCH", PROPERTIES("operation"="insert", "column"="user_id", "value"='1')))));

ALTER SHADOW RULE shadow_rule(
SOURCE=demo_ds,
SHADOW=demo_ds_shadow,
t_order((simple_hint_algorithm, TYPE(NAME="SIMPLE_HINT", PROPERTIES("shadow"="true", "foo"="bar"))), (TYPE(NAME="REGEX_MATCH", PROPERTIES("operation"="insert", "column"="user_id", "regex"='[1]')))),
t_order_item((TYPE(NAME="VALUE_MATCH", PROPERTIES("operation"="insert", "column"="user_id", "value"='1')))));

CREATE SHADOW ALGORITHM
(simple_hint_algorithm, TYPE(NAME="SIMPLE_HINT", PROPERTIES("shadow"="true", "foo"="bar"))),
(user_id_match_algorithm, TYPE(NAME="REGEX_MATCH", PROPERTIES("operation"="insert", "column"="user_id", "regex"='[1]')));

ALTER SHADOW ALGORITHM
(simple_hint_algorithm, TYPE(NAME="SIMPLE_HINT", PROPERTIES("shadow"="false", "foo"="bar"))),
```



```
(user_id_match_algorithm, TYPE(NAME="VALUE_MATCH",PROPERTIES("operation"="insert",
"column"="user_id", "value"='1')));

DROP SHADOW RULE shadow_rule;

DROP SHADOW ALGORITHM simple_hint_algorithm;

CREATE DEFAULT SHADOW ALGORITHM NAME = simple_hint_algorithm;
```

RQL 语法

RQL (Resource & Rule Query Language) 为 Apache ShardingSphere 的资源 and 规则查询语言。

资源查询

语法说明

```
SHOW DATABASE RESOURCES [FROM databaseName]
```

返回值说明

列	说明
name	数据源名称
type	数据源类型
host	数据源地址
port	数据源端口
db	数据库名称
attribute	数据源参数

示例

```
mysql> SHOW DATABASE RESOURCES;
+-----+-----+-----+-----+-----+-----+-----+
|      |      |      |      |      |      |      |
+-----+-----+-----+-----+-----+-----+
|      |      |      |      |      |      |      |
+-----+-----+-----+-----+-----+-----+
|      |      |      |      |      |      |      |
+-----+-----+-----+-----+-----+-----+
|      |      |      |      |      |      |      |
+-----+-----+-----+-----+-----+-----+
|      |      |      |      |      |      |      |
+-----+-----+-----+-----+-----+-----+
|      |      |      |      |      |      |      |
+-----+-----+-----+-----+-----+-----+
```

```

-----+
| name | type | host      | port | db   | connection_timeout_milliseconds | idle_
timeout_milliseconds | max_lifetime_milliseconds | max_pool_size | min_pool_size |
read_only | other_attributes
|
+-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
| ds_0 | MySQL | 127.0.0.1 | 3306 | db_0 | 30000 | 60000
| 1800000 | 50 | 1 |
false | {"dataSourceProperties":{"cacheServerConfiguration":"true",
"elideSetAutoCommits":"true","useServerPrepStmts":"true","cachePrepStmts":"true",
"rewriteBatchedStatements":"true","cacheResultSetMetadata":"false",
"useLocalSessionState":"true","maintainTimeStats":"false","prepStmtCacheSize":"8192
","tinyInt1isBit":"false","prepStmtCacheSqlLimit":"2048",
"netTimeoutForStreamingResults":"0","zeroDateTimeBehavior":"round"},
"healthCheckProperties":{"},"initializationFailTimeout":1,"validationTimeout":5000,
"leakDetectionThreshold":0,"poolName":"HikariPool-1","registerMbeans":false,
"allowPoolSuspension":false,"autoCommit":true,"isolateInternalQueries":false} |
| ds_1 | MySQL | 127.0.0.1 | 3306 | db_1 | 30000 | 60000
| 1800000 | 50 | 1 |
false | {"dataSourceProperties":{"cacheServerConfiguration":"true",
"elideSetAutoCommits":"true","useServerPrepStmts":"true","cachePrepStmts":"true",
"rewriteBatchedStatements":"true","cacheResultSetMetadata":"false",
"useLocalSessionState":"true","maintainTimeStats":"false","prepStmtCacheSize":"8192
","tinyInt1isBit":"false","prepStmtCacheSqlLimit":"2048",
"netTimeoutForStreamingResults":"0","zeroDateTimeBehavior":"round"},
"healthCheckProperties":{"},"initializationFailTimeout":1,"validationTimeout":5000,
"leakDetectionThreshold":0,"poolName":"HikariPool-2","registerMbeans":false,
"allowPoolSuspension":false,"autoCommit":true,"isolateInternalQueries":false} |
+-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+

```

```

-----
-----
-----
-----
-----+
2 rows in set (0.84 sec)

```

规则查询

本章节将对规则查询的语法进行详细说明。

数据分片

语法说明

Sharding Table Rule

```

SHOW SHARDING TABLE tableRule | RULES [FROM databaseName]

SHOW SHARDING ALGORITHMS [FROM databaseName]

SHOW UNUSED SHARDING ALGORITHMS [FROM databaseName]

SHOW SHARDING AUDITORS [FROM databaseName]

SHOW SHARDING TABLE RULES USED ALGORITHM shardingAlgorithmName [FROM databaseName]

SHOW SHARDING KEY GENERATORS [FROM databaseName]

SHOW UNUSED SHARDING KEY GENERATORS [FROM databaseName]

SHOW SHARDING TABLE RULES USED KEY GENERATOR keyGeneratorName [FROM databaseName]

SHOW DEFAULT SHARDING STRATEGY

SHOW SHARDING TABLE NODES

tableRule:
    RULE tableName

```

- 支持查询所有数据分片规则和指定表查询；
- 支持查询所有分片算法；
- 支持查询所有分片审计算法。

Sharding Binding Table Rule

```
SHOW SHARDING BINDING TABLE RULES [FROM databaseName]
```

Sharding Broadcast Table Rule

```
SHOW SHARDING BROADCAST TABLE RULES [FROM databaseName]
```

Sharding Table Rule

列	说明
table	逻辑表名
actual_data_nodes	实际的数据节点
actual_data_sources	实际的数据源（通过 RDL 创建的规则时显示）
database_strategy_type	数据库分片策略类型
database_sharding_column	数据库分片键
database_sharding_algorithm_type	数据库分片算法类型
database_sharding_algorithm_props	数据库分片算法参数
table_strategy_type	表分片策略类型
table_sharding_column	表分片键
table_sharding_algorithm_type	表分片算法类型
table_sharding_algorithm_props	表分片算法参数
key_generate_column	分布式主键生成列
key_generator_type	分布式主键生成器类型
key_generator_props	分布式主键生成器参数

Sharding Algorithms

列	说明
name	分片算法名称
type	分片算法类型
props	分片算法参数

Unused Sharding Algorithms

列	说明
name	分片算法名称
type	分片算法类型
props	分片算法参数

Sharding Auditors

列	说明
name	分片审计算法名称
type	分片审计算法类型
props	分片审计算法参数

Sharding Key Generators

列	说明
name	主键生成器名称
type	主键生成器类型
props	主键生成器参数

Unused Sharding Key Generators

列	说明
name	主键生成器名称
type	主键生成器类型
props	主键生成器参数

Default Sharding Strategy

列	说明
name	策略名称
type	分片策略类型
sharding_column	分片键
sharding_algorithm_name	分片算法名称
sharding_algorithm_type	分片算法类型
sharding_algorithm_props	分片算法参数

Sharding Table Nodes

列	说明
name	分片规则名称
nodes	分片节点

Sharding Binding Table Rule

列	说明
sharding_binding_tables	绑定表名称

Sharding Broadcast Table Rule

列	说明
sharding_broadcast_tables	广播表名称

Sharding Table Rule

SHOW SHARDING TABLE RULES

```
mysql> SHOW SHARDING TABLE RULES;
+-----+-----+-----+-----+-----+-----+-----+-----+-----+-----+
| table          | actual_data_nodes          | actual_data_sources | database_ |
| strategy_type | database_sharding_column | database_sharding_algorithm_type |
| database_sharding_algorithm_props | table_strategy_type | table_sharding_ |
| column | table_sharding_algorithm_type | table_sharding_algorithm_props |
|          | key_generate_column | key_generator_type | key_generator_props |
+-----+-----+-----+-----+-----+-----+-----+-----+-----+-----+
| t_order       | ds_${0..1}.t_order_${0..1} |          | INLINE |
|               | user_id                  |          | INLINE |
|               | algorithm-expression:ds_${user_id % 2} | order_id | INLINE |
|               | algorithm-expression:t_order_${order_id % 2} | order_id |
|               | SNOWFLAKE                |          |
| t_order_item | ds_${0..1}.t_order_item_${0..1} |          | INLINE |
```

```

      | user_id      | INLINE      | algorithm-
expression:ds_${user_id % 2} | INLINE      | order_id      | INLINE
      | algorithm-expression:t_order_item_${order_id % 2} | order_item_id
      | SNOWFLAKE      |              | |
| t2      |              | ds_0,ds_1      |
      |              |              |
      |              | mod      | id      | mod
      | sharding-count:10      |              |
      |              |              |
+-----+-----+-----+-----+
-----+-----+-----+-----+
-----+-----+-----+-----+
-----+-----+-----+-----+
+-----+-----+-----+
3 rows in set (0.02 sec)

```

SHOW SHARDING TABLE RULE tableName

```

mysql> SHOW SHARDING TABLE RULE t_order;
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
| table | actual_data_nodes | actual_data_sources | database_strategy_
type | database_sharding_column | database_sharding_algorithm_type | database_
sharding_algorithm_props | table_strategy_type | table_sharding_column |
table_sharding_algorithm_type | table_sharding_algorithm_props |
key_generate_column | key_generator_type | key_generator_props |
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
| t_order | ds_${0..1}.t_order_${0..1} | | INLINE |
user_id | INLINE | algorithm-expression:ds_${
{user_id % 2} | INLINE | order_id | INLINE
| algorithm-expression:t_order_${order_id % 2} | order_id | SNOWFLAKE
| |
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
1 row in set (0.01 sec)

```

SHOW SHARDING ALGORITHMS

```
mysql> SHOW SHARDING ALGORITHMS;
+-----+-----+-----+
| name          | type   | props          |
+-----+-----+-----+
| t_order_inline | INLINE | algorithm-expression=t_order_${order_id % 2}
| t_order_item_inline | INLINE | algorithm-expression=t_order_item_${order_id % 2}
+-----+-----+-----+
2 row in set (0.01 sec)
```

SHOW UNUSED SHARDING ALGORITHMS

```
mysql> SHOW UNUSED SHARDING ALGORITHMS;
+-----+-----+-----+
| name          | type   | props          |
+-----+-----+-----+
| t1_inline      | INLINE | algorithm-expression=t_order_${order_id % 2}
+-----+-----+-----+
1 row in set (0.01 sec)
```

SHOW SHARDING AUDITORS

```
mysql> SHOW SHARDING AUDITORS;
+-----+-----+-----+
| name          | type          | props          |
+-----+-----+-----+
| dml_audit      | DML_SHARDING_CONDITIONS
+-----+-----+-----+
2 row in set (0.01 sec)
```

SHOW SHARDING TABLE RULES USED ALGORITHM *shardingAlgorithmName*

```
mysql> SHOW SHARDING TABLE RULES USED ALGORITHM t_order_inline;
+-----+-----+
| type | name   |
+-----+-----+
| table | t_order
+-----+-----+
1 row in set (0.01 sec)
```

SHOW SHARDING KEY GENERATORS

```
mysql> SHOW SHARDING KEY GENERATORS;
+-----+-----+-----+
| name          | type          | props          |
+-----+-----+-----+
```


name	type	props
t_order_snowflake	snowflake	
t_order_item_snowflake	snowflake	
uuid_key_generator	uuid	

3 row in set (0.01 sec)

SHOW UNUSED SHARDING KEY GENERATORS

```
mysql> SHOW UNUSED SHARDING KEY GENERATORS;
```

name	type	props
uuid_key_generator	uuid	

1 row in set (0.01 sec)

SHOW SHARDING TABLE RULES USED KEY GENERATOR *keyGeneratorName*

```
mysql> SHOW SHARDING TABLE RULES USED KEY GENERATOR t_order_snowflake;
```

type	name
table	t_order

1 row in set (0.01 sec)

SHOW DEFAULT SHARDING STRATEGY

```
mysql> SHOW DEFAULT SHARDING STRATEGY ;
```

name	type	sharding_column	sharding_algorithm_name	sharding_algorithm_type	sharding_algorithm_props
TABLE	NONE				
DATABASE	STANDARD	order_id	database_inline	INLINE	{algorithm-expression=ds_\${user_id % 2}}

2 rows in set (0.07 sec)

SHOW SHARDING TABLE NODES

```
mysql> SHOW SHARDING TABLE NODES;
+-----+-----+
| name      | nodes                                                                 |
+-----+-----+
| t_order   | ds_0.t_order_0, ds_1.t_order_1, ds_0.t_order_2, ds_1.t_order_3 |
+-----+-----+
1 row in set (0.02 sec)
```

Sharding Binding Table Rule

```
mysql> SHOW SHARDING BINDING TABLE RULES;
+-----+
| sharding_binding_tables |
+-----+
| t_order,t_order_item |
| t1,t2                |
+-----+
2 rows in set (0.00 sec)
```

Sharding Broadcast Table Rule

```
mysql> SHOW SHARDING BROADCAST TABLE RULES;
+-----+
| sharding_broadcast_tables |
+-----+
| t_1                      |
| t_2                      |
+-----+
2 rows in set (0.00 sec)
```

单表

语法说明

```
SHOW SINGLE TABLE (table | RULES) [FROM databaseName]

SHOW SINGLE TABLES

COUNT SINGLE_TABLE RULE [FROM databaseName]

table:
    TABLE tableName
```

返回值说明

Single Table Rule

列	说明
name	规则名称
resource_name	数据源名称

Single Table

列	说明
table_name	单表名称
resource_name	单表所在的数据源名称

Single Table Rule Count

列	说明
rule_name	规则名称
database	单表所在的数据库名称
count	规则个数

示例

SHOW SINGLE TABLES RULES

```
sql> SHOW SINGLE TABLES RULES;
+-----+-----+
| name    | resource_name |
+-----+-----+
| default | ds_1          |
+-----+-----+
1 row in set (0.01 sec)
```

SHOW SINGLE TABLE tableName

```
sql> SHOW SINGLE TABLE t_single_0;
+-----+-----+
| table_name | resource_name |
+-----+-----+
| t_single_0 | ds_0          |
+-----+-----+
1 row in set (0.01 sec)
```

SHOW SINGLE TABLES

```
mysql> SHOW SINGLE TABLES;
+-----+-----+
| table_name | resource_name |
+-----+-----+
| t_single_0 | ds_0          |
| t_single_1 | ds_1          |
+-----+-----+
2 rows in set (0.02 sec)
```

COUNT SINGLE_TABLE RULE

```
mysql> COUNT SINGLE_TABLE RULE;
+-----+-----+-----+
| rule_name | database | count |
+-----+-----+-----+
| t_single_0 | ds       | 2      |
+-----+-----+-----+
1 row in set (0.02 sec)
```

读写分离

语法说明

```
SHOW READWRITE_SPLITTING RULES [FROM databaseName]
```

返回值说明

列	说明
name	规则名称
auto_aware_data_source_name	自动发现数据源名称（配置动态读写分离规则显示）
write_data_source_query_enabled	读库全部下线，主库是否承担读流量
write_data_source_name	写数据源名称
read_data_source_names	读数据源名称列表
load_balancer_type	负载均衡算法类型
load_balancer_props	负载均衡算法参数

示例

静态读写分离规则

```
mysql> SHOW READWRITE_SPLITTING RULES;
+-----+-----+-----+-----+
| name      | auto_aware_data_source_name | write_data_source_name | read_data_source_names | load_balancer_type | load_balancer_props |
+-----+-----+-----+-----+
| ms_group_0 |                               | ds_primary              | ds_slave_0, ds_slave_1 | random              |                      |
+-----+-----+-----+-----+
1 row in set (0.00 sec)
```

动态读写分离规则

```
mysql> SHOW READWRITE_SPLITTING RULES FROM readwrite_splitting_db;
+-----+-----+-----+-----+
| name      | auto_aware_data_source_name | write_data_source_query_enabled | write_data_source_name | read_data_source_names | load_balancer_type | load_balancer_props |
+-----+-----+-----+-----+
| readwrite_ds | ms_group_0                  |                                |                        |                        | random              | read_weight=2:1      |
+-----+-----+-----+-----+
1 row in set (0.01 sec)
```

静态读写分离规则和动态读写分离规则

```
mysql> SHOW READWRITE_SPLITTING RULES FROM readwrite_splitting_db;
+-----+-----+-----+-----+
| name      | auto_aware_data_source_name | write_data_source_query_enabled | write_data_source_name | read_data_source_names | load_balancer_type | load_balancer_props |
+-----+-----+-----+-----+
| readwrite_ds | ms_group_0                  |                                |                        | read_ds_0, read_ds_1 | random              | read_weight=2:1      |
+-----+-----+-----+-----+
```

```
1 row in set (0.00 sec)
```

数据库发现

语法说明

```
SHOW DB_DISCOVERY RULES [FROM databaseName]

SHOW DB_DISCOVERY TYPES [FROM databaseName]

SHOW DB_DISCOVERY HEARTBEATS [FROM databaseName]
```

返回值说明

DB Discovery Rule

列	说明
group_name	规则名称
data_source_names	数据源名称列表
primary_data_source_name	主数据源名称
discovery_type	数据库发现服务类型
discovery_heartbeat	数据库发现服务心跳

DB Discovery Type

列	说明
name	类型名称
type	类型种类
props	类型参数

DB Discovery Heartbeat

列	说明
name	心跳名称
props	心跳参数

示例

DB Discovery Rule

```
mysql> SHOW DB_DISCOVERY RULES;
+-----+-----+-----+-----+
| group_name          | data_source_names | primary_data_source_name | discovery_
type                  |
discovery_heartbeat   |
+-----+-----+-----+-----+
| db_discovery_group_0 | ds_0,ds_1,ds_2    | ds_0                      | {name=db_
discovery_group_0_mgr, type=MySQL.MGR, props={group-name=92504d5b-6dec}} |
{name=db_discovery_group_0_heartbeat, props={keep-alive-cron=0/5 * * * * ?}} |
+-----+-----+-----+-----+
1 row in set (0.20 sec)
```

DB Discovery Type

```
mysql> SHOW DB_DISCOVERY TYPES;
+-----+-----+-----+
| name                  | type          | props                      |
+-----+-----+-----+
| db_discovery_group_0_mgr | MySQL.MGR    | {group-name=92504d5b-6dec} |
+-----+-----+-----+
1 row in set (0.01 sec)
```

DB Discovery Heartbeat

```
mysql> SHOW DB_DISCOVERY HEARTBEATS;
+-----+-----+
| name                  | props                      |
+-----+-----+
| db_discovery_group_0_heartbeat | {keep-alive-cron=0/5 * * * * ?} |
+-----+-----+
1 row in set (0.01 sec)
```

数据加密

语法说明

```
SHOW ENCRYPT RULES [FROM databaseName]
```

```
SHOW ENCRYPT TABLE RULE tableName [FROM databaseName]
```

- 支持查询所有的数据加密规则和指定逻辑表名查询。

返回值说明

列	说明
table	逻辑表名
logic_column	逻辑列名
logic_data_type	逻辑列数据类型
cipher_column	密文列名
cipher_data_type	密文列数据类型
plain_column	明文列名
plain_data_type	明文列数据类型
assisted_query_column	辅助查询列名
assisted_query_data_type	辅助查询列数据类型
encryptor_type	加密算法类型
encryptor_props	加密算法参数
query_with_cipher_column	是否使用加密列进行查询

示例

显示加密规则

```
mysql> SHOW ENCRYPT RULES FROM encrypt_db;
```

```
+-----+-----+-----+-----+-----+
+-----+-----+-----+-----+-----+
+-----+-----+-----+-----+-----+
| table      | logic_column | logic_data_type | cipher_column | cipher_data_type |
plain_column | plain_data_type | assisted_query_column | assisted_query_data_type |
encryptor_type | encryptor_props | query_with_cipher_column |
+-----+-----+-----+-----+-----+
+-----+-----+-----+-----+-----+
+-----+-----+-----+-----+-----+
| t_encrypt  | user_id      |                  | user_cipher   |                  |
user_plain   |              |                  |               |                  |
AES          | aes-key-value=123456abc | true              |               |                  |
| t_encrypt  | order_id     |                  | order_cipher  |                  |
```


MD5				true	
t_encrypt_2	user_id			user_cipher	
user_plain					
AES	aes-key-value=123456abc		false		
t_encrypt_2	order_id			order_cipher	
MD5			false		
+-----+-----+-----+-----+-----+					
+-----+-----+-----+-----+-----+					
+-----+-----+-----+-----+-----+					
4 rows in set (0.78 sec)					

显示加密表规则表名

```
mysql> SHOW ENCRYPT TABLE RULE t_encrypt;
```

+-----+-----+-----+-----+-----+					
+-----+-----+-----+-----+-----+					
+-----+-----+-----+-----+-----+					
table	logic_column	logic_data_type	cipher_column	cipher_data_type	
plain_column	plain_data_type	assisted_query_column	assisted_query_data_type		
encryptor_type	encryptor_props		query_with_cipher_column		
+-----+-----+-----+-----+-----+					
+-----+-----+-----+-----+-----+					
+-----+-----+-----+-----+-----+					
t_encrypt	user_id		user_cipher		
user_plain					
AES	aes-key-value=123456abc		true		
t_encrypt	order_id		order_cipher		
MD5			true		
+-----+-----+-----+-----+-----+					
+-----+-----+-----+-----+-----+					
+-----+-----+-----+-----+-----+					
2 rows in set (0.01 sec)					

```
mysql> SHOW ENCRYPT TABLE RULE t_encrypt FROM encrypt_db;
```

+-----+-----+-----+-----+-----+					
+-----+-----+-----+-----+-----+					
+-----+-----+-----+-----+-----+					
table	logic_column	logic_data_type	cipher_column	cipher_data_type	
plain_column	plain_data_type	assisted_query_column	assisted_query_data_type		
encryptor_type	encryptor_props		query_with_cipher_column		
+-----+-----+-----+-----+-----+					
+-----+-----+-----+-----+-----+					
+-----+-----+-----+-----+-----+					
t_encrypt	user_id		user_cipher		
user_plain					

AES	aes-key-value=123456abc	true		
t_encrypt	order_id		order_cipher	
MD5		true		
+-----+-----+-----+-----+-----+				
+-----+-----+-----+-----+-----+				
+-----+-----+-----+-----+-----+				
2 rows in set (0.01 sec))				

影子库压测

语法说明

SHOW SHADOW shadowRule RULES [FROM databaseName]
SHOW SHADOW TABLE RULES [FROM databaseName]
SHOW SHADOW ALGORITHMS [FROM databaseName]
shadowRule:
RULE ruleName

- 支持查询所有影子规则和指定表查询；
- 支持查询所有表规则；
- 支持查询所有影子算法。

返回值说明

Shadow Rule

列	说明
rule_name	规则名称
source_name	源数据库
shadow_name	影子数据库
shadow_table	影子表

Shadow Table Rule

列	说明
shadow_table	影子表
shadow_algorithm_name	影子算法名称

Shadow Algorithms

列	说明
shadow_algorithm_name	影子算法名称
type	算法类型
props	算法参数
is_default	是否默认

Shadow Rule status

列	说明
status	是否启用

示例

SHOW SHADOW RULES

```
mysql> SHOW SHADOW RULES;
+-----+-----+-----+-----+
| rule_name      | source_name | shadow_name | shadow_table |
+-----+-----+-----+-----+
| shadow_rule_1  | ds_1       | ds_shadow_1 | t_order      |
| shadow_rule_2  | ds_2       | ds_shadow_2 | t_order_item |
+-----+-----+-----+-----+
2 rows in set (0.02 sec)
```

SHOW SHADOW RULE ruleName

```
mysql> SHOW SHADOW RULE shadow_rule_1;
+-----+-----+-----+-----+
| rule_name      | source_name | shadow_name | shadow_table |
+-----+-----+-----+-----+
| shadow_rule_1  | ds_1       | ds_shadow_1 | t_order      |
+-----+-----+-----+-----+
1 rows in set (0.01 sec)
```

SHOW SHADOW TABLE RULES

```
mysql> SHOW SHADOW TABLE RULES;
```

shadow_table	shadow_algorithm_name
t_order_1	user_id_match_algorithm,simple_hint_algorithm_1
1 rows in set (0.01 sec)	

SHOW SHADOW ALGORITHMS

```
mysql> SHOW SHADOW ALGORITHMS;
```

shadow_algorithm_name	type	props
is_default		
user_id_match_algorithm	REGEX_MATCH	operation=insert,column=user_id,regex=[1]
simple_hint_algorithm_1	SIMPLE_HINT	shadow=true,foo=bar
2 rows in set (0.01 sec)		

RAL 语法

RAL (Resource & Rule Administration Language) 为 Apache ShardingSphere 的管理语言，负责强制路由、熔断、配置导入导出、数据迁移控制等管理功能。

强制路由

语句	说明	示例
SET READ- WRITE_SPLITTING HINT SOURCE = [auto / write]	针对当前连接，设置读写分离的路由策略（自动路由或强制到写库）	SET RE AD- WRITE_SPLITTING HINT SOURCE = write
SET SHARDING HINT DATABASE_VALUE = yy	针对当前连接，设置 hint 仅对数据库分片有效，并添加分片值，yy：数据库分片值	SET SHARDING HINT DATABASE_VALUE = 100
ADD SHARDING HINT DATABASE_VALUE xx = yy	针对当前连接，为表 xx 添加分片值 yy，xx：逻辑表名称，yy：数据库分片值	ADD SHARDING HINT DATABASE_VALUE t_order= 100
ADD SHARDING HINT TA- BLE_VALUE xx = yy	针对当前连接，为表 xx 添加分片值 yy，xx：逻辑表名称，yy：表分片值	ADD SHARDING HINT TABLE_VALUE t_order = 100
CLEAR HINT	针对当前连接，清除 hint 所有设置	CLEAR HINT
CLEAR [SHARDING HINT / READWRITE_SPLITTING HINT]	针对当前连接，清除 sharding 或 read-write_splitting 的 hint 设置	CLEAR RE AD- WRITE_SPLITTING HINT
SHOW [SHARDING / R EAD- WRITE_SPLITTING] HINT STATUS	针对当前连接，查询 sharding 或 read-write_splitting 的 hint 设置	SHOW RE AD- WRITE_SPLITTING HINT STATUS

数据迁移

语句	说明	示例
MIGRATE TABLE ds.schema.table INTO table	从源端迁移到目标端	MIGRATE TABLE ds_0.public.t_order INTO t_order
SHOW MIGRATION LIST	查询运行列表	SHOW MIGRATION LIST
SHOW MIGRATION STATUS jobId	查询作业状态	SHOW MIGRATION STATUS 1234
STOP MIGRATION jobId	停止作业	STOP MIGRATION 12345
START MIGRATION jobId	开启停止的作业	START MIGRATION 1234
CHECK MIGRATION jobId	数据一致性校验	CHECK MIGRATION 1234
SHOW MIGRATION CHECK AL- GORITHMS	展示可用的一致性校验算法	SHOW MIGRATION CHECK AL- GORITHMS
CHECK MIGRATION jobId (by type(name=algorithmTypeName)?	数据一致性校验, 使用指定的 校验算法	CHECK MIGRATION 1234 by type(name= "DATA_MATCH")
ROLLBACK MIGRATION jobId	撤销作业。注意：该语句会清 理目标端表，请谨慎操作	ROLLBACK MIGRATION 1234
COMMIT MIGRATION jobId	完成作业	COMMIT MIGRATION 1234

熔断

语句	说明	示例
[ENABLE / DISABLE] READWRITE_SPLITTING (READ)? resourceName [FROM databaseName]	启用 / 禁用读库	ENABLE R EAD- WRITE_SPLITTING READ resource_0
[ENABLE / DISABLE] INSTANCE instanceId	启用 / 禁用 proxy 实例	DISABLE INSTANCE in- stance_1
SHOW INSTANCE LIST	查询 proxy 实例信息	SHOW INSTANCE LIST
SHOW READWRITE_SPLITTING (READ)? resource-Name [FROM databaseName]	查询所有读库的状态	SHOW R EAD- WRITE_SPLITTING READ RESOURCES

全局规则

语句	说明	示例
SHOW AUTHORITY RULE	查询权限规则配置	SHOW AUTHORITY RULE
SHOW TRANSACTION RULE	查询事务规则配置	SHOW TRANSACTION RULE
SHOW SQL_PARSER RULE	查询解析引擎规则配置	SHOW SQL_PARSER RULE
ALTER TRANSACTION RULE(DE FAULT=xx,TYPE(NAME=xxx, PROPERTIES(key1=value1,key2=value2...)))	更新事务规则配置, DE FAULT: 默认事务类型, 支持 LOCAL、XA、BASE; NAME: 事务管理器名称, 支持 Ato mikos、Narayana 和 Bitronix	ALTER TRANSACTION RULE(DEFAULT= "XA" ,TYPE(NAME= "Narayana" , PROPERTIES("databaseName" = "jbossts" , "host" = "127.0.0.1")))
ALTER SQL_PARSER RULE SQL_COMMENT_PARSE_ENABLE=xx, PARSE_TREE_CACHE(INITIAL_CAPACITY=xx, MAXIMUM_SIZE=xx, CONCURRENCY_LEVEL=xx), SQL_STATEMENT_CACHE(INITIAL_CAPACITY=xxx, MAXIMUM_SIZE=xxx, CONCURRENCY_LEVEL=xxx)	更新解析引擎规则配置, SQL_COMMENT_PARSE_ENABLE: 是否解析 SQL 注释, "PARSE_TREE_CACHE": 语法树本地缓存配置, SQL_STATEMENT_CACHE: SQL 语句本地缓存配置项	ALTER SQL_PARSER RULE SQL_COMMENT PARSE_ENABLE=false, PARSE_TREE_CACHE HE(INITIAL_CAPACITY=10, MAXIMUM_SIZE=11, CON- CURRENCY_LEVEL=1), SQL_STATEMENT_CAC HE(INITIAL_CAPACITY=11, MAXIMUM_SIZE=11, CONCUR- RENCY_LEVEL=100)

其他

语句	说明	示例
SHOW INSTANCE INFO	查询当前 proxy 的实例信息	SHOW INSTANCE INFO
SHOW MODE INFO	查询当前 proxy 的 mode 配置	SHOW MODE INFO
COUNT DATABASE RULES [FROM database]	查询 database 中的规则数量	COUNT DATABASE RULES
SET VARIABLE proxy_property_name = xx	proxy_property_name 为 proxy 的属性配置，需使用下划线命名	SET VARIABLE sql_show = true
SET VARIABLE transaction_type = xx	修改当前连接的事务类型，支持 LOCAL, XA, BASE	SET VARIABLE transaction_type = "XA"
SET VARIABLE agent_plugins_enabled = [TRUE / FALSE]	设置 agent 插件的启用状态，默认值 false	SET VARIABLE agent_plugins_enabled = TRUE
SHOW ALL VARIABLES	查询 proxy 所有的属性配置	SHOW ALL VARIABLES
SHOW VARIABLE variable_name	查询 proxy 属性，需使用下划线命名	SHOW VARIABLE sql_show
REFRESH TABLE METADATA	刷新所有表的元数据	REFRESH TABLE METADATA
REFRESH TABLE METADATA tableName	刷新指定表的元数据	REFRESH TABLE METADATA t_order
REFRESH TABLE METADATA tableName FROM RESOURCE resourceName	刷新指定数据源中表的元数据	REFRESH TABLE METADATA t_order FROM RESOURCE ds_1
REFRESH TABLE METADATA FROM RESOURCE resourceName SCHEMA schemaName	刷新指定 schema 中表的元数据，如果 schema 中不存在表，则会删除该 schema	REFRESH TABLE METADATA FROM RESOURCE ds_1 SCHEMA db_schema
SHOW TABLE METADATA tableName [, tableName] ...	查询表的元数据	SHOW TABLE METADATA t_order
EXPORT DATABASE CONFIG [FROM database_name] [, file= "file_path"]	将 database 中的资源和规则配置导出为 YAML 格式	EXPORT DATABASE CONFIG FROM readwrite_splitting_db
IMPORT DATABASE CONFIG FILE= "file_path"	将 YAML 中的配置导入到 database 中，仅支持对空库进行导入操作	IMPORT DATABASE CONFIG FILE = "/xxx/config-sharding.yaml"
SHOW RULES USED RESOURCE resourceName [from database]	查询 database 中使用指定资源的规则	SHOW RULES USED RESOURCE ds_0 FROM databaseName

注意事项

ShardingSphere-Proxy 默认不支持 hint，如需支持，请在 `conf/server.yaml` 中，将 `props` 的属性 `proxy-hint-enabled` 设置为 `true`。

RUL 语法

RUL (Resource Utility Language) 为 Apache ShardingSphere 的工具类语言，提供 SQL 解析、SQL 格式化、执行计划预览等功能。

SQL 工具

语句	说明	示例
PARSE SQL	解析 SQL 并输出抽象语法树	PARSE SELECT * FROM t_order
FORMAT SQL	解析并输出格式化后的 SQL 语句	FORMAT SELECT * FROM t_order
PREVIEW SQL	预览 SQL 执行计划	PREVIEW SELECT * FROM t_order

使用

本章节将结合 DistSQL 的语法，并以实战的形式分别介绍如何使用 DistSQL 管理分布式数据库下的资源和规则。

前置工作

以 MySQL 为例，其他数据库可直接替换。

1. 启动 MySQL 服务；
2. 创建待注册资源的 MySQL 数据库；
3. 在 MySQL 中为 ShardingSphere-Proxy 创建一个拥有创建权限的角色或者用户；
4. 启动 ZooKeeper 服务；
5. 添加 `mode` 和 `authentication` 配置参数到 `server.yaml`；
6. 启动 ShardingSphere-Proxy；
7. 通过应用程序或终端连接到 ShardingSphere-Proxy；

创建数据库

1. 创建逻辑库。

```
CREATE DATABASE foo_db;
```

2. 使用新创建的逻辑库。

```
USE foo_db;
```

资源操作

详见具体规则示例。

规则操作

详见具体规则示例。

注意事项

1. 当前, DROP DATABASE 只会移除 逻辑的分布式数据库, 不会删除用户真实的数据库;
2. DROP TABLE 会将逻辑分片表和数据库中真实的表全部删除;
3. CREATE DATABASE 只会创建 逻辑的分布式数据库, 所以需要用户提前创建好真实的数据库。

数据分片

资源操作

```
ADD RESOURCE ds_0 (  
    HOST="127.0.0.1",  
    PORT=3306,  
    DB="ds_1",  
    USER="root",  
    PASSWORD="root"  
) , ds_1 (  
    HOST="127.0.0.1",  
    PORT=3306,  
    DB="ds_2",  
    USER="root",  
    PASSWORD="root"  
);
```

规则操作

- 创建分片规则

```
CREATE SHARDING TABLE RULE t_order(
  RESOURCES(ds_0,ds_1),
  SHARDING_COLUMN=order_id,
  TYPE(NAME="hash_mod",PROPERTIES("sharding-count"="4")),
  KEY_GENERATE_STRATEGY(COLUMN=order_id,TYPE(NAME="snowflake"))
);
```

- 创建切分表

```
CREATE TABLE `t_order` (
  `order_id` int NOT NULL,
  `user_id` int NOT NULL,
  `status` varchar(45) DEFAULT NULL,
  PRIMARY KEY (`order_id`)
) ENGINE=InnoDB DEFAULT CHARSET=utf8mb4
```

- 删除切分表

```
DROP TABLE t_order;
```

- 删除分片规则

```
DROP SHARDING TABLE RULE t_order;
```

- 删除数据源

```
DROP RESOURCE ds_0, ds_1;
```

- 删除分布式数据库

```
DROP DATABASE foo_db;
```

读写分离

资源操作

```
ADD RESOURCE write_ds (
  HOST="127.0.0.1",
  PORT=3306,
  DB="ds_0",
  USER="root",
  PASSWORD="root"
),read_ds (
```

```

    HOST="127.0.0.1",
    PORT=3307,
    DB="ds_0",
    USER="root",
    PASSWORD="root"
);

```

规则操作

- 创建读写分离规则

```

CREATE READWRITE_SPLITTING RULE group_0 (
WRITE_RESOURCE=write_ds,
READ_RESOURCES(read_ds),
TYPE(NAME="random")
);

```

- 修改读写分离规则

```

ALTER READWRITE_SPLITTING RULE group_0 (
WRITE_RESOURCE=write_ds,
READ_RESOURCES(read_ds),
TYPE(NAME="random",PROPERTIES("read_weight"="2:0"))
);

```

- 删除读写分离规则

```

DROP READWRITE_SPLITTING RULE group_0;

```

- 删除数据源

```

DROP RESOURCE write_ds,read_ds;

```

- 删除分布式数据库

```

DROP DATABASE readwrite_splitting_db;

```

数据加密

资源操作

```

ADD RESOURCE ds_0 (
    HOST="127.0.0.1",
    PORT=3306,
    DB="ds_0",
    USER="root",

```

```
PASSWORD="root"
);
```

规则操作

- 创建加密规则

```
CREATE ENCRYPT RULE t_encrypt (
  COLUMNS(
    (NAME=user_id,PLAIN=user_plain,CIPHER=user_cipher,TYPE(NAME='AES',
  PROPERTIES('aes-key-value'='123456abc'))),
    (NAME=order_id,PLAIN=order_plain,CIPHER =order_cipher,TYPE(NAME='RC4',
  PROPERTIES('rc4-key-value'='123456abc'))))
);
```

- 创建加密表

```
CREATE TABLE `t_encrypt` (
  `id` int(11) NOT NULL,
  `user_id` varchar(45) DEFAULT NULL,
  `order_id` varchar(45) DEFAULT NULL,
  PRIMARY KEY (`id`)
) ENGINE=InnoDB DEFAULT CHARSET=utf8mb4;
```

- 修改加密规则

```
ALTER ENCRYPT RULE t_encrypt (
  COLUMNS(
    (NAME=user_id,PLAIN=user_plain,CIPHER=user_cipher,TYPE(NAME='AES',
  PROPERTIES('aes-key-value'='123456abc'))))
);
```

- 删除加密规则

```
DROP ENCRYPT RULE t_encrypt;
```

- 删除数据源

```
DROP RESOURCE ds_0;
```

- 删除分布式数据库

```
DROP DATABASE encrypt_db;
```

数据库发现

资源操作

```
ADD RESOURCE ds_0 (
    HOST="127.0.0.1",
    PORT=3306,
    DB="ds_0",
    USER="root",
    PASSWORD="root"
),ds_1 (
    HOST="127.0.0.1",
    PORT=3306,
    DB="ds_1",
    USER="root",
    PASSWORD="root"
),ds_2 (
    HOST="127.0.0.1",
    PORT=3306,
    DB="ds_2",
    USER="root",
    PASSWORD="root"
);
```

规则操作

- 创建数据库发现规则

```
CREATE DB_DISCOVERY RULE db_discovery_group_0 (
    RESOURCES(ds_0, ds_1),
    TYPE(NAME='MySQL.MGR',PROPERTIES('group-name'='92504d5b-6dec')),
    HEARTBEAT(PROPERTIES('keep-alive-cron'='0/5 * * * * ?'))
);
```

- 修改数据库发现规则

```
ALTER DB_DISCOVERY RULE db_discovery_group_0 (
    RESOURCES(ds_0, ds_1, ds_2),
    TYPE(NAME='MySQL.MGR',PROPERTIES('group-name'='92504d5b-6dec')),
    HEARTBEAT(PROPERTIES('keep-alive-cron'='0/5 * * * * ?'))
);
```

- 删除数据库发现规则

```
DROP DB_DISCOVERY RULE db_discovery_group_0;
```

- 删除数据库发现类型

```
DROP DB_DISCOVERY TYPE db_discovery_group_0_mgr;
```

- 删除数据库发现心跳

```
DROP DB_DISCOVERY HEARTBEAT db_discovery_group_0_heartbeat;
```

- 删除数据源

```
DROP RESOURCE ds_0,ds_1,ds_2;
```

- 删除分布式数据库

```
DROP DATABASE discovery_db;
```

影子库压测

资源操作

```
ADD RESOURCE ds_0 (  
    HOST="127.0.0.1",  
    PORT=3306,  
    DB="ds_0",  
    USER="root",  
    PASSWORD="root"  
) ,ds_1 (  
    HOST="127.0.0.1",  
    PORT=3306,  
    DB="ds_1",  
    USER="root",  
    PASSWORD="root"  
) ,ds_2 (  
    HOST="127.0.0.1",  
    PORT=3306,  
    DB="ds_2",  
    USER="root",  
    PASSWORD="root"  
);
```

规则操作

- 创建影子库压测规则

```
CREATE SHADOW RULE group_0(
SOURCE=ds_0,
SHADOW=ds_1,
t_order((simple_hint_algorithm, TYPE(NAME="SIMPLE_HINT", PROPERTIES("foo"="bar"))),
(TYPE(NAME="REGEX_MATCH", PROPERTIES("operation"="insert", "column"="user_id",
"regex"='[1]')))),
t_order_item((TYPE(NAME="SIMPLE_HINT", PROPERTIES("foo"="bar")))));
```

- 修改影子库压测规则

```
ALTER SHADOW RULE group_0(
SOURCE=ds_0,
SHADOW=ds_2,
t_order_item((TYPE(NAME="SIMPLE_HINT", PROPERTIES("foo"="bar")))));
```

- 删除影子库压测规则

```
DROP SHADOW RULE group_0;
```

- 删除数据源

```
DROP RESOURCE ds_0,ds_1,ds_2;
```

9. 删除分布式数据库

```
DROP DATABASE foo_db;
```

4.2.4 数据迁移

简介

ShardingSphere 可以提供给用户通用的数据迁移解决方案。

于 **4.1.0** 开始向用户提供，目前仍处于 实验室版本。

运行部署

背景信息

对于使用单数据库运行的系统来说，如何安全简单地将数据迁移至水平分片的数据库上，一直以来都是一个迫切的需求。

前提条件

- Proxy 采用纯 JAVA 开发，JDK 建议 1.8 或以上版本。
- 数据迁移使用集群模式，目前支持 ZooKeeper 作为注册中心。

操作步骤

1. 执行以下命令，编译生成 ShardingSphere-Proxy 二进制包：

```
git clone --depth 1 https://github.com/apache/shardingsphere.git
cd shardingsphere
mvn clean install -Dmaven.javadoc.skip=true -Dcheckstyle.skip=true -Drat.skip=true
-Djacoco.skip=true -DskipITs -DskipTests -Prelease
```

发 布 包：- /shardingsphere-distribution/shardingsphere-proxy-distribution/target/apache-shardingsphere-`{{latest.release.version}}`-shardingsphere-proxy-bin.tar.gz

或者通过[下载页面](#)获取安装包。

2. 解压缩 proxy 发布包，修改配置文件 conf/config-sharding.yaml。详情请参见[proxy 启动手册](#)。
3. 修改配置文件 conf/server.yaml，详情请参见[模式配置](#)。

目前 mode 必须是 Cluster，需要提前启动对应的注册中心。

配置示例：

```
mode:
  type: Cluster
  repository:
    type: ZooKeeper
    props:
      namespace: governance_ds
      server-lists: localhost:2181
      retryIntervalMilliseconds: 500
      timeToLiveSeconds: 60
      maxRetries: 3
      operationTimeoutMilliseconds: 500
  overwrite: false
```

4. 引入 JDBC 驱动。

proxy 已包含 PostgreSQL JDBC 驱动。

如果后端连接以下数据库，请下载相应 JDBC 驱动 jar 包，并将其放入 `/${shardingsphere-proxy}/ext-lib` 目录。

数 据 库	JDBC 驱动	参考
MySQL	<code>'mysql-connector-java-5.1.47.jar < https://repo1.maven.org/maven2/mysql/mysql-connector-java/5.1.47/mysql-connector-java-5.1.47.jar>'_</code>	Connector/J Versions
open-Gauss	<code>opengauss-jdbc-3.0.0.jar</code>	

如果是异构迁移，源端支持范围更广的数据库，比如：Oracle。JDBC 驱动处理方式同上。

5. 启动 ShardingSphere-Proxy:

```
sh bin/start.sh
```

6. 查看 proxy 日志 logs/stdout.log, 看到日志中出现:

```
[INFO ] [main] o.a.s.p.frontend.ShardingSphereProxy - ShardingSphere-Proxy start success
```

确认启动成功。

7. 按需配置迁移

7.1. 查询配置。

```
SHOW MIGRATION PROCESS CONFIGURATION;
```

默认配置如下:

```
+-----+-----+
| read                                     | write
|                                     |
| stream_channel                       |
|                                     |
+-----+-----+
| {"workerThread":40,"batchSize":1000,"shardingSize":10000000} | {"workerThread":40,"batchSize":1000} | {"type":"MEMORY","props":{"block-queue-size":10000}} |
+-----+-----+
```

7.2. 新建配置 (可选)。

不配置的话有默认值。

完整配置 DistSQL 示例:

```
CREATE MIGRATION PROCESS CONFIGURATION (
READ(
  WORKER_THREAD=40,
  BATCH_SIZE=1000,
```

```

    SHARDING_SIZE=10000000,
    RATE_LIMITER (TYPE(NAME='QPS',PROPERTIES('qps'='500')))
),
WRITE(
    WORKER_THREAD=40,
    BATCH_SIZE=1000,
    RATE_LIMITER (TYPE(NAME='TPS',PROPERTIES('tps'='2000')))
),
STREAM_CHANNEL (TYPE(NAME='MEMORY',PROPERTIES('block-queue-size'='10000')))
);

```

配置项说明:

```

CREATE MIGRATION PROCESS CONFIGURATION (
READ( -- 数据读取配置。如果不配置则部分参数默认生效。
    WORKER_THREAD=40, -- 从源端摄取全量数据的线程池大小。如果不配置则使用默认值。
    BATCH_SIZE=1000, -- 一次查询操作返回的最大记录数。如果不配置则使用默认值。
    SHARDING_SIZE=10000000, -- 全量数据分片大小。如果不配置则使用默认值。
    RATE_LIMITER ( -- 限流算法。如果不配置则不限流。
        TYPE( -- 算法类型。可选项: QPS
            NAME='QPS',
            PROPERTIES( -- 算法属性
                'qps'='500'
            )))
),
WRITE( -- 数据写入配置。如果不配置则部分参数默认生效。
    WORKER_THREAD=40, -- 数据写入到目标端的线程池大小。如果不配置则使用默认值。
    BATCH_SIZE=1000, -- 一次批量写入操作的最大记录数。如果不配置则使用默认值。
    RATE_LIMITER ( -- 限流算法。如果不配置则不限流。
        TYPE( -- 算法类型。可选项: TPS
            NAME='TPS',
            PROPERTIES( -- 算法属性
                'tps'='2000'
            )))
),
STREAM_CHANNEL ( -- 数据通道, 连接生产者和消费者, 用于 read 和 write 环节。如果不配置则默认使用 MEMORY 类型。
    TYPE( -- 算法类型。可选项: MEMORY
        NAME='MEMORY',
        PROPERTIES( -- 算法属性
            'block-queue-size'='10000' -- 属性: 阻塞队列大小
        )))
);

```

DistSQL 示例: 配置 READ 限流。

```

CREATE MIGRATION PROCESS CONFIGURATION (
READ(

```

```
RATE_LIMITER (TYPE(NAME='QPS',PROPERTIES('qps'='500'))
)
);
```

配置读取数据限流，其它配置使用默认值。

7.3. 修改配置。

ALTER MIGRATION PROCESS CONFIGURATION, 内部结构和 CREATE MIGRATION PROCESS CONFIGURATION 一致。

DistSQL 示例：调整限流参数

```
ALTER MIGRATION PROCESS CONFIGURATION (
READ(
  RATE_LIMITER (TYPE(NAME='QPS',PROPERTIES('qps'='1000'))
)
);
---
ALTER MIGRATION PROCESS CONFIGURATION (
READ(
  RATE_LIMITER (TYPE(NAME='QPS',PROPERTIES('qps'='1000'))
), WRITE(
  RATE_LIMITER (TYPE(NAME='TPS',PROPERTIES('tps'='1000'))
)
);
```

7.4. 清除配置。

DistSQL 示例：清空 READ 配置、恢复为默认值。

```
DROP MIGRATION PROCESS CONFIGURATION '/READ';
```

DistSQL 示例：清空 READ/RATE_LIMITER 配置。

```
DROP MIGRATION PROCESS CONFIGURATION '/READ/RATE_LIMITER';
```

使用手册

MySQL 使用手册

环境要求

支持的 MySQL 版本：5.1.15 ~ 5.7.x。

权限要求

1. 开启 binlog

MySQL 5.7 my.cnf 示例配置:

```
[mysqld]
server-id=1
log-bin=mysql-bin
binlog-format=row
binlog-row-image=full
max_connections=600
```

执行以下命令，确认是否有开启 binlog:

```
show variables like '%log_bin%';
show variables like '%binlog%';
```

如以下显示，则说明 binlog 已开启

Variable_name	Value
log_bin	ON
binlog_format	ROW
binlog_row_image	FULL

2. 赋予 MySQL 账号 Replication 相关权限。

执行以下命令，查看该用户是否有迁移权限:

```
SHOW GRANTS FOR 'user';
```

示例结果:

Grants for \${username}@\${host}
GRANT REPLICATION SLAVE, REPLICATION CLIENT ON *.* TO \${username}@\${host}
.....

完整流程示例

前提条件

1. 在 MySQL 已准备好源端库、表、数据。

示例:

```
DROP DATABASE IF EXISTS migration_ds_0;
CREATE DATABASE migration_ds_0 DEFAULT CHARSET utf8;

USE migration_ds_0

CREATE TABLE t_order (order_id INT NOT NULL, user_id INT NOT NULL, status
VARCHAR(45) NULL, PRIMARY KEY (order_id));

INSERT INTO t_order (order_id, user_id, status) VALUES (1,2,'ok'),(2,4,'ok'),(3,6,
'ok'),(4,1,'ok'),(5,3,'ok'),(6,5,'ok');
```

2. 在 MySQL 准备目标端库。

示例:

```
DROP DATABASE IF EXISTS migration_ds_10;
CREATE DATABASE migration_ds_10 DEFAULT CHARSET utf8;

DROP DATABASE IF EXISTS migration_ds_11;
CREATE DATABASE migration_ds_11 DEFAULT CHARSET utf8;

DROP DATABASE IF EXISTS migration_ds_12;
CREATE DATABASE migration_ds_12 DEFAULT CHARSET utf8;
```

操作步骤

1. 在 proxy 新建逻辑数据库并配置好资源和规则。

```
CREATE DATABASE sharding_db;

USE sharding_db

ADD RESOURCE ds_2 (
    URL="jdbc:mysql://127.0.0.1:3306/migration_ds_10?serverTimezone=UTC&
useSSL=false",
    USER="root",
    PASSWORD="root",
    PROPERTIES("minPoolSize"="1","maxPoolSize"="20","idleTimeout"="60000")
), ds_3 (
    URL="jdbc:mysql://127.0.0.1:3306/migration_ds_11?serverTimezone=UTC&
```

```

useSSL=false",
    USER="root",
    PASSWORD="root",
    PROPERTIES("minPoolSize"="1","maxPoolSize"="20","idleTimeout"="60000")
), ds_4 (
    URL="jdbc:mysql://127.0.0.1:3306/migration_ds_12?serverTimezone=UTC&
useSSL=false",
    USER="root",
    PASSWORD="root",
    PROPERTIES("minPoolSize"="1","maxPoolSize"="20","idleTimeout"="60000")
);

CREATE SHARDING TABLE RULE t_order(
RESOURCES(ds_2,ds_3,ds_4),
SHARDING_COLUMN=order_id,
TYPE(NAME="hash_mod",PROPERTIES("sharding-count"="6")),
KEY_GENERATE_STRATEGY(COLUMN=order_id,TYPE(NAME="snowflake"))
);

```

如果是迁移到异构数据库，那目前需要在 proxy 执行建表语句。

2. 在 proxy 配置源端资源。

```

ADD MIGRATION SOURCE RESOURCE ds_0 (
    URL="jdbc:mysql://127.0.0.1:3306/migration_ds_0?serverTimezone=UTC&useSSL=false
",
    USER="root",
    PASSWORD="root",
    PROPERTIES("minPoolSize"="1","maxPoolSize"="20","idleTimeout"="60000")
);

```

3. 启动数据迁移。

```
MIGRATE TABLE ds_0.t_order INTO t_order;
```

或者指定目标端逻辑库：

```
MIGRATE TABLE ds_0.t_order INTO sharding_db.t_order;
```

4. 查看数据迁移作业列表。

```
SHOW MIGRATION LIST;
```

示例结果：

```

+-----+-----+-----+-----+-----+
| id | tables | sharding_total_count | active |
| create_time | stop_time |

```

+-----+-----+-----+-----+-----+				
-----+-----+				
j015d4ee1b8a5e7f95df19babb2794395e8	t_order	1	true	
2022-08-22 16:37:01	NULL			
+-----+-----+-----+-----+-----+				
-----+-----+				

5. 查看数据迁移详情。

```
SHOW MIGRATION STATUS 'j015d4ee1b8a5e7f95df19babb2794395e8';
```

+-----+-----+-----+-----+-----+				
-----+-----+				
item	data_source	status	active	inventory_finished_
percentage	incremental_idle_seconds			
+-----+-----+-----+-----+-----+				
-----+-----+				
0	ds_0	EXECUTE_INCREMENTAL_TASK	true	100
141				
+-----+-----+-----+-----+-----+				
-----+-----+				

6. 执行数据一致性校验。

```
CHECK MIGRATION 'j015d4ee1b8a5e7f95df19babb2794395e8' BY TYPE (NAME='CRC32_MATCH');
```

+-----+-----+-----+-----+			
+-----+-----+			
table_name	source_records_count	target_records_count	records_count_matched
records_content_matched			
+-----+-----+-----+-----+			
+-----+-----+			
t_order	6	6	true
true			
+-----+-----+-----+-----+			
+-----+-----+			

数据一致性校验算法类型来自：

SHOW MIGRATION CHECK ALGORITHMS;		
+-----+-----+-----+		
-----+-----+		
type	supported_database_types	
description		
+-----+-----+-----+		
-----+-----+		
CRC32_MATCH	MySQL	
Match CRC32 of records.		
DATA_MATCH	SQL92,MySQL,MariaDB,PostgreSQL,openGauss,Oracle,SQLServer,H2	
Match raw data of records.		
+-----+-----+-----+		

目标端开启数据加密的情况需要使用 `DATA_MATCH`。

异构迁移需要使用 `DATA_MATCH`。

7. 完成作业。

```
COMMIT MIGRATION 'j015d4ee1b8a5e7f95df19babb2794395e8';
```

8. 刷新元数据。

```
REFRESH TABLE METADATA;
```

更多 DistSQL 请参见 [RAL # 数据迁移](#)。

PostgreSQL 使用手册

环境要求

支持的 PostgreSQL 版本：9.4 或以上版本。

权限要求

1. 开启 `test_decoding`。
2. 调整 WAL 配置。

postgresql.conf 示例配置：

```
wal_level = logical
max_wal_senders = 10
max_replication_slots = 10
max_connections = 600
```

详情请参见 [Write Ahead Log](#) 和 [Replication](#)。

3. 配置 PostgreSQL 允许 Proxy 拥有 replication 权限。

pg_hba.conf 示例配置：

```
host replication repl_acct 0.0.0.0/0 md5
```

详情请参见 [The pg_hba.conf File](#)。

完整流程示例

前提条件

1. 在 PostgreSQL 已准备好源端库、表、数据。

```
DROP DATABASE IF EXISTS migration_ds_0;
CREATE DATABASE migration_ds_0;

\c migration_ds_0

CREATE TABLE t_order (order_id INT NOT NULL, user_id INT NOT NULL, status
VARCHAR(45) NULL, PRIMARY KEY (order_id));

INSERT INTO t_order (order_id, user_id, status) VALUES (1,2,'ok'),(2,4,'ok'),(3,6,
'ok'),(4,1,'ok'),(5,3,'ok'),(6,5,'ok');
```

2. 在 PostgreSQL 准备目标端库。

```
DROP DATABASE IF EXISTS migration_ds_10;
CREATE DATABASE migration_ds_10;

DROP DATABASE IF EXISTS migration_ds_11;
CREATE DATABASE migration_ds_11;

DROP DATABASE IF EXISTS migration_ds_12;
CREATE DATABASE migration_ds_12;
```

操作步骤

1. 在 proxy 新建逻辑数据库并配置好资源和规则。

```
CREATE DATABASE sharding_db;

\c sharding_db

ADD RESOURCE ds_2 (
  URL="jdbc:postgresql://127.0.0.1:5432/migration_ds_10",
  USER="postgres",
  PASSWORD="root",
  PROPERTIES("minPoolSize"="1","maxPoolSize"="20","idleTimeout"="60000")
), ds_3 (
  URL="jdbc:postgresql://127.0.0.1:5432/migration_ds_11",
  USER="postgres",
  PASSWORD="root",
  PROPERTIES("minPoolSize"="1","maxPoolSize"="20","idleTimeout"="60000")
), ds_4 (
```

```

URL="jdbc:postgresql://127.0.0.1:5432/migration_ds_12",
USER="postgres",
PASSWORD="root",
PROPERTIES("minPoolSize"="1","maxPoolSize"="20","idleTimeout"="60000")
);

CREATE SHARDING TABLE RULE t_order(
RESOURCES(ds_2,ds_3,ds_4),
SHARDING_COLUMN=order_id,
TYPE(NAME="hash_mod",PROPERTIES("sharding-count"="6")),
KEY_GENERATE_STRATEGY(COLUMN=order_id,TYPE(NAME="snowflake"))
);

```

如果是迁移到异构数据库，那目前需要在 proxy 执行建表语句。

2. 在 proxy 配置源端资源。

```

ADD MIGRATION SOURCE RESOURCE ds_0 (
URL="jdbc:postgresql://127.0.0.1:5432/migration_ds_0",
USER="postgres",
PASSWORD="root",
PROPERTIES("minPoolSize"="1","maxPoolSize"="20","idleTimeout"="60000")
);

```

3. 启动数据迁移。

```
MIGRATE TABLE ds_0.t_order INTO t_order;
```

或者指定目标端逻辑库：

```
MIGRATE TABLE ds_0.t_order INTO sharding_db.t_order;
```

也可以指定源端 schema：

```
MIGRATE TABLE ds_0.public.t_order INTO sharding_db.t_order;
```

4. 查看数据迁移作业列表。

```
SHOW MIGRATION LIST;
```

示例结果：

```

+-----+-----+-----+-----+-----+
| id                                     | tables | sharding_total_count | active |
| create_time       | stop_time |
+-----+-----+-----+-----+-----+
| j015d4ee1b8a5e7f95df19babb2794395e8 | t_order | 1                     | true   |
2022-08-22 16:37:01 | NULL    |

```

5. 查看数据迁移详情。

```
SHOW MIGRATION STATUS 'j015d4ee1b8a5e7f95df19babb2794395e8';
```

item	data_source	status	active	inventory_finished_percentage	incremental_idle_seconds
0	ds_0	EXECUTE_INCREMENTAL_TASK	true	100	141

6. 执行数据一致性校验。

```
CHECK MIGRATION 'j015d4ee1b8a5e7f95df19babb2794395e8';
```

table_name	source_records_count	target_records_count	records_count_matched	records_content_matched
t_order	6	6	true	

7. 完成作业。

```
COMMIT MIGRATION 'j015d4ee1b8a5e7f95df19babb2794395e8';
```

8. 刷新元数据。

```
REFRESH TABLE METADATA;
```

更多 DistSQL 请参见 [RAL # 数据迁移](#)。

openGauss 使用手册

环境要求

支持的 openGauss 版本：2.0.1 ~ 3.0.0。

权限要求

1. 调整 WAL 配置。

postgresql.conf 示例配置：

```
wal_level = logical
max_wal_senders = 10
max_replication_slots = 10
wal_sender_timeout = 0
max_connections = 600
```

详情请参见 [Write Ahead Log](#) 和 [Replication](#)。

2. 配置 openGauss 允许 Proxy 拥有 replication 权限。

pg_hba.conf 示例配置：

```
host replication repl_acct 0.0.0.0/0 md5
```

详情请参见 [Configuring Client Access Authentication](#) 和 [Example: Logic Replication Code](#)。

完整流程示例

前提条件

1. 在 openGauss 已准备好源端库、表、数据。

```
DROP DATABASE IF EXISTS migration_ds_0;
CREATE DATABASE migration_ds_0;

\c migration_ds_0

CREATE TABLE t_order (order_id INT NOT NULL, user_id INT NOT NULL, status
VARCHAR(45) NULL, PRIMARY KEY (order_id));

INSERT INTO t_order (order_id, user_id, status) VALUES (1,2,'ok'),(2,4,'ok'),(3,6,
'ok'),(4,1,'ok'),(5,3,'ok'),(6,5,'ok');
```

2. 在 openGauss 准备目标端库。

```

DROP DATABASE IF EXISTS migration_ds_10;
CREATE DATABASE migration_ds_10;

DROP DATABASE IF EXISTS migration_ds_11;
CREATE DATABASE migration_ds_11;

DROP DATABASE IF EXISTS migration_ds_12;
CREATE DATABASE migration_ds_12;

```

操作步骤

1. 在 proxy 新建逻辑数据库并配置好资源和规则。

```

CREATE DATABASE sharding_db;

\c sharding_db

ADD RESOURCE ds_2 (
    URL="jdbc:opengauss://127.0.0.1:5432/migration_ds_10",
    USER="gaussdb",
    PASSWORD="Root@123",
    PROPERTIES("minPoolSize"="1","maxPoolSize"="20","idleTimeout"="60000")
), ds_3 (
    URL="jdbc:opengauss://127.0.0.1:5432/migration_ds_11",
    USER="gaussdb",
    PASSWORD="Root@123",
    PROPERTIES("minPoolSize"="1","maxPoolSize"="20","idleTimeout"="60000")
), ds_4 (
    URL="jdbc:opengauss://127.0.0.1:5432/migration_ds_12",
    USER="gaussdb",
    PASSWORD="Root@123",
    PROPERTIES("minPoolSize"="1","maxPoolSize"="20","idleTimeout"="60000")
);

CREATE SHARDING TABLE RULE t_order(
    RESOURCES(ds_2,ds_3,ds_4),
    SHARDING_COLUMN=order_id,
    TYPE(NAME="hash_mod",PROPERTIES("sharding-count"="6")),
    KEY_GENERATE_STRATEGY(COLUMN=order_id,TYPE(NAME="snowflake"))
);

```

如果是迁移到异构数据库，那目前需要在 proxy 执行建表语句。

2. 在 proxy 配置源端资源。

```

ADD MIGRATION SOURCE RESOURCE ds_2 (
    URL="jdbc:opengauss://127.0.0.1:5432/migration_ds_0",
    USER="gaussdb",

```

```
PASSWORD="Root@123",
PROPERTIES("minPoolSize"="1","maxPoolSize"="20","idleTimeout"="60000")
);
```

3. 启动数据迁移。

```
MIGRATE TABLE ds_0.t_order INTO t_order;
```

或者指定目标端逻辑库：

```
MIGRATE TABLE ds_0.t_order INTO sharding_db.t_order;
```

也可以指定源端 schema：

```
MIGRATE TABLE ds_0.public.t_order INTO sharding_db.t_order;
```

4. 查看数据迁移作业列表。

```
SHOW MIGRATION LIST;
```

示例结果：

```
+-----+-----+-----+-----+-----+
| id          | tables | sharding_total_count | active |
create_time   | stop_time |
+-----+-----+-----+-----+
| j015d4ee1b8a5e7f95df19babb2794395e8 | t_order | 1 | true |
2022-08-22 16:37:01 | NULL |
```

5. 查看数据迁移详情。

```
SHOW MIGRATION STATUS 'j015d4ee1b8a5e7f95df19babb2794395e8';
+-----+-----+-----+-----+-----+
| item | data_source | status | active | inventory_finished_
percentage | incremental_idle_seconds |
+-----+-----+-----+-----+
| 0 | ds_0 | EXECUTE_INCREMENTAL_TASK | true | 100
| 141 |
```

6. 执行数据一致性校验。

```
CHECK MIGRATION 'j015d4ee1b8a5e7f95df19babb2794395e8';
```

+-----+-----+-----+-----+			
+-----+-----+-----+-----+			
table_name	source_records_count	target_records_count	records_count_matched
records_content_matched			
+-----+-----+-----+-----+			
+-----+-----+-----+-----+			
t_order	6	6	true
true			
+-----+-----+-----+-----+			
+-----+-----+-----+-----+			

7. 完成作业。

```
COMMIT MIGRATION 'j015d4ee1b8a5e7f95df19babb2794395e8';
```

8. 刷新元数据。

```
REFRESH TABLE METADATA;
```

更多 DistSQL 请参见 [RAL # 数据迁移](#)。

4.2.5 可观察性

源码编译

从 Github 下载 Apache ShardingSphere 源码，对源码进行编译，操作命令如下。

```
git clone --depth 1 https://github.com/apache/shardingsphere.git
cd shardingsphere
mvn clean install -Dmaven.javadoc.skip=true -Dcheckstyle.skip=true -Drat.skip=true
-Djacoco.skip=true -DskipITs -DskipTests -Prelease
```

agent 包输出目录为 shardingsphere-agent/shardingsphere-agent-distribution/target/apache-shardingsphere-`{latest.release.version}`-shardingsphere-agent-bin.tar.gz

agent 配置

- 目录说明

创建 agent 目录，解压 agent 二进制包到 agent 目录。

```
mkdir agent
tar -zxvf apache-shardingsphere-{latest.release.version}-shardingsphere-agent-bin.
tar.gz -C agent
cd agent
tree
.
```

```

└─ apache-shardingsphere-${latest.release.version}-shardingsphere-agent-bin
    ├── LICENSE
    ├── NOTICE
    ├── conf
    │   ├── agent.yaml
    │   └─ logback.xml
    ├── plugins
    │   ├── shardingsphere-agent-logging-base-${latest.release.version}.jar
    │   ├── shardingsphere-agent-metrics-prometheus-${latest.release.version}.jar
    │   ├── shardingsphere-agent-tracing-jaeger-${latest.release.version}.jar
    │   └─ shardingsphere-agent-tracing-opentelemetry-${latest.release.version}.
jar
    ├── shardingsphere-agent-tracing-opentracing-${latest.release.version}.jar
    ├── shardingsphere-agent-tracing-zipkin-${latest.release.version}.jar
    └─ shardingsphere-agent.jar

```

- 配置说明

conf/agent.yaml 用于管理 agent 配置。内置插件包括 Jaeger、OpenTracing、Zipkin、OpenTelemetry、Logging 及 Prometheus。当需要开启插件时，只需要移除 ignoredPluginNames 中对应的插件名称即可。

```

applicationName: shardingsphere-agent
ignoredPluginNames:
  - Jaeger
  - OpenTracing
  - Zipkin
  - OpenTelemetry
  - Logging
  - Prometheus

plugins:
  Prometheus:
    host: "localhost"
    port: 9090
    props:
      JVM_INFORMATION_COLLECTOR_ENABLED : "true"
  Jaeger:
    host: "localhost"
    port: 5775
    props:
      SERVICE_NAME: "shardingsphere-agent"
      JAEGER_SAMPLER_TYPE: "const"
      JAEGER_SAMPLER_PARAM: "1"
  Zipkin:
    host: "localhost"
    port: 9411
    props:
      SERVICE_NAME: "shardingsphere-agent"

```



```
URL_VERSION: "/api/v2/spans"
SAMPLER_TYPE: "const"
SAMPLER_PARAM: "1"
OpenTracing:
  props:
    OPENTRACING_TRACER_CLASS_NAME: "org.apache.skywalking.apm.toolkit.
opentracing.SkywalkingTracer"
OpenTelemetry:
  props:
    otel.resource.attributes: "service.name=shardingsphere-agent"
    otel.traces.exporter: "zipkin"
Logging:
  props:
    LEVEL: "INFO"
```

- 参数说明;

名称	说明	取值范围	默认值
JVM_INFORMATION_COLLECTION_ENABLED	是否开启 JVM 采集器	true、false	true
SERVICE_NAME	链路跟踪的服务名称	自定义	shardingsphere-agent
JAEGER_SAMPLER_TYPE	Jaeger 采样率类型	const、probabilistic、ratelimiting、remote	const
JAEGER_SAMPLER_PARAMETER	Jaeger 采样率参数	const:0、1,probabilistic:0.0 - 1.0,ratelimiting:>0,自定义每秒采集数量, remote: 需要自定义配置远程采样率管理服务地址, JAEGER_SAMPLER_MANAGER_HOST_PORT	1 (const 类型)
SAMPLER_TYPE	Zipkin 采样率类型	const、counting、ratelimiting、boundary	const
SAMPLER_PARAMETER	Zipkin 采样率参数	const: 0、1, counting: 0.01 - 1.0, ratelimiting: >0, 自定义每秒采集数量, boundary: 0.0001 - 1.0	1 (const 类型)
otel.resource.attributes	opentelemetry 资源属性	字符串键值对 (, 分割)	service.name=shardingsphere-agent
otel.traces.exporter	Tracing exporter	zipkin、jaeger	zipkin
otel.traces.sampler	opentelemetry 采样率类型	always_on、always_off、traceidratio	always_on
otel.traces.sampler.arg	opentelemetry 采样率参数	traceidratio: 0.0 - 1.0	1.0

ShardingSphere-Proxy 中使用

- 编辑启动脚本

配置 shardingsphere-agent.jar 的绝对路径到 ShardingSphere-Proxy 的 start.sh 启动脚本中, 请注意配置自己对应的绝对路径。

```
nohup java ${JAVA_OPTS} ${JAVA_MEM_OPTS} \
-javaagent:/xxxxx/agent/shardingsphere-agent.jar \
```

```
-classpath ${CLASS_PATH} ${MAIN_CLASS} >> ${STDOUT_FILE} 2>&1 &
```

- 启动 ShardingSphere-Proxy

```
bin/start.sh
```

正常启动后，可以在 ShardingSphere-Proxy 日志中找到 plugin 的加载信息，访问 Proxy 后，可以通过配置的监控地址查看到 Metric 和 Tracing 的数据。

4.3 通用配置

本章主要介绍通用配置，包括属性配置和内置算法配置。

4.3.1 属性配置

背景信息

Apache ShardingSphere 提供属性配置的方式配置系统级配置。

参数解释

操作步骤

属性配置直接配置在 ShardingSphere-JDBC 所使用的配置文件中，格式如下：

```
props:
  sql-show: true
```

配置示例

ShardingSphere 仓库的示例中包含了多种不同场景的属性配置，请参考：<https://github.com/apache/shardingsphere/tree/master/examples/shardingsphere-jdbc-example>

4.3.2 内置算法

简介

Apache ShardingSphere 通过 SPI 方式允许开发者扩展算法；与此同时，Apache ShardingSphere 也提供了大量的内置算法以便于开发者使用。

使用方式

内置算法均通过 `type` 和 `props` 进行配置，其中 `type` 由算法定义在 SPI 中，`props` 用于传递算法的个性化参数配置。

无论使用哪种配置方式，均是将配置完毕的算法命名，并传递至相应的规则配置中。本章节根据功能区分并罗列 Apache ShardingSphere 全部的内置算法，供开发者参考。

元数据持久化仓库

背景信息

Apache ShardingSphere 为不同的运行模式提供了不同的元数据持久化方式，用户在配置运行模式的同时可以选择合适的方式来存储元数据。

参数解释

文件持久化

类型：File

适用模式：Standalone

可配置属性：

名称	数据类型	说明	默认值
path	String	元数据存储路径	.shardingsphere

ZooKeeper 持久化

类型：ZooKeeper

适用模式：Cluster

可配置属性：

名称	数据类型	说明	默认值
retryIntervalMilliseconds	int	重试间隔毫秒数	500
maxRetries	int	客户端连接最大重试次数	3
timeToLiveSeconds	int	临时数据失效的秒数	60
operationTimeoutMilliseconds	int	客户端操作超时的毫秒数	500
digest	String	登录认证密码	

Etcd 持久化

类型：Etcd

适用模式：Cluster

可配置属性：

名称	数据类型	说明	默认值
timeToLiveSeconds	long	临时数据失效的秒数	30
connectionTimeout	long	连接超时秒数	30

操作步骤

1. 在 server.yaml 中配置 Mode 运行模式
2. 配置元数据持久化仓库类型

配置示例

- 单机模式配置方式

```
mode:
  type: Standalone
  repository:
    type: File
    props:
      path: ~/user/.shardingsphere
  overwrite: false
```

- 集群模式

```
mode:
  type: Cluster
  repository:
    type: zookeeper
    props:
      namespace: governance_ds
      server-lists: localhost:2181
      retryIntervalMilliseconds: 500
      timeToLiveSeconds: 60
      maxRetries: 3
      operationTimeoutMilliseconds: 500
  overwrite: false
```

分片算法

背景信息

ShardingSphere 内置提供了多种分片算法，按照类型可以划分为自动分片算法、标准分片算法、复合分片算法和 Hint 分片算法，能够满足用户绝大多数业务场景的需要。此外，考虑到业务场景的复杂性，内置算法也提供了自定义分片算法的方式，用户可以通过编写 Java 代码来完成复杂的分片逻辑。

参数解释

自动分片算法

取模分片算法

类型：MOD

可配置属性：

属性名称	数据类型	说明
sharding-count	int	分片数量

哈希取模分片算法

类型：HASH_MOD

可配置属性：

属性名称	数据类型	说明
sharding-count	int	分片数量

基于分片容量的范围分片算法

类型：VOLUME_RANGE

可配置属性：

属性名称	数据类型	说明
range-lower	long	范围下界，超过边界的数据会报错
range-upper	long	范围上界，超过边界的数据会报错
sharding-volume	long	分片容量

基于分片边界的范围分片算法

类型: BOUNDARY_RANGE

可配置属性:

属性名称	数据类型	说明
sharding-ranges	String	分片的范围边界, 多个范围边界以逗号分隔

自动时间段分片算法

类型: AUTO_INTERVAL

可配置属性:

属性名称	数据类型 *	说明
datetime-lower	String	分片的起始时间范围, 时间戳格式: yyyy-MM-dd HH:mm:ss
datetime-upper	String	分片的结束时间范围, 时间戳格式: yyyy-MM-dd HH:mm:ss
sharding-seconds	long	单一片片所能承载的最大时间, 单位: 秒, 允许分片键的时间戳格式的秒带有时间精度, 但秒后的时间精度会被自动抹去

标准分片算法

Apache ShardingSphere 内置的标准分片算法实现类包括:

行表达式分片算法

使用 Groovy 的表达式, 提供对 SQL 语句中的 = 和 IN 的分片操作支持, 只支持单分片键。对于简单的分片算法, 可以通过简单的配置使用, 从而避免繁琐的 Java 代码开发, 如: `t_user_${u_id % 8}` 表示 `t_user` 表根据 `u_id` 模 8, 而分成 8 张表, 表名称为 `t_user_0` 到 `t_user_7`。详情请参见[行表达式](#)。

类型: INLINE

可配置属性:

时间范围分片算法

此算法主动忽视了 `datetime-pattern` 的时区信息。这意味着当 `datetime-lower`, `datetime-upper` 和传入的分片键含有时区信息时, 不会因为时区不一致而发生时区转换。当传入的分片键为 `java.time.Instant` 时存在特例处理, 其会携带上系统的时区信息后转化为 `datetime-pattern` 的字符串格式, 再进行下一步分片。

类型: INTERVAL

可配置属性:

复合分片算法

复合行表达式分片算法

详情请参见行表达式。

类型: COMPLEX_INLINE

Hint 分片算法

Hint 行表达式分片算法

详情请参见行表达式。

类型: HINT_INLINE

属性名称	数据类型	说明	默认值
algorithm-expression (?)	String	分片算法的行表达式	\${value}

自定义类分片算法

通过配置分片策略类型和算法类名, 实现自定义扩展。CLASS_BASED 允许向算法类内传入额外的自定义属性, 传入的属性可以通过属性名为 `props` 的 `java.util.Properties` 类实例取出。参考 Git 的 `org.apache.shardingsphere.example.extension.sharding.algorithm.classbased.fixture.ClassBasedStandardShardingAlgorithmFixture`。

类型: CLASS_BASED

可配置属性:

属性名称	数据类型	说明
strategy	String	分片策略类型, 支持 STANDARD、COMPLEX 或 HINT (不区分大小写)
algorithmClassName	String	分片算法全限定名

操作步骤

1. 使用数据分片时，在 `shardingAlgorithms` 属性下配置对应的数据分片算法即可；

配置示例

```
rules:
- !SHARDING
  tables:
    t_order:
      actualDataNodes: ds_${0..1}.t_order_${0..1}
      tableStrategy:
        standard:
          shardingColumn: order_id
          shardingAlgorithmName: t-order-inline
      keyGenerateStrategy:
        column: order_id
        keyGeneratorName: snowflake
    t_order_item:
      actualDataNodes: ds_${0..1}.t_order_item_${0..1}
      tableStrategy:
        standard:
          shardingColumn: order_id
          shardingAlgorithmName: t_order-item-inline
      keyGenerateStrategy:
        column: order_item_id
        keyGeneratorName: snowflake
    t_account:
      actualDataNodes: ds_${0..1}.t_account_${0..1}
      tableStrategy:
        standard:
          shardingAlgorithmName: t-account-inline
      keyGenerateStrategy:
        column: account_id
        keyGeneratorName: snowflake
  defaultShardingColumn: account_id
  bindingTables:
    - t_order,t_order_item
  broadcastTables:
    - t_address
  defaultDatabaseStrategy:
    standard:
      shardingColumn: user_id
      shardingAlgorithmName: database-inline
  defaultTableStrategy:
    none:
```

```

shardingAlgorithms:
  database-inline:
    type: INLINE
    props:
      algorithm-expression: ds_${user_id % 2}
  t-order-inline:
    type: INLINE
    props:
      algorithm-expression: t_order_${order_id % 2}
  t_order-item-inline:
    type: INLINE
    props:
      algorithm-expression: t_order_item_${order_id % 2}
  t-account-inline:
    type: INLINE
    props:
      algorithm-expression: t_account_${account_id % 2}
keyGenerators:
  snowflake:
    type: SNOWFLAKE

```

相关参考

- [核心特性：数据分片](#)
- [开发者指南：数据分片](#)

分布式序列算法

背景信息

传统数据库软件开发中，主键自动生成技术是基本需求。而各个数据库对于该需求也提供了相应的支持，比如 MySQL 的自增键，Oracle 的自增序列等。数据分片后，不同数据节点生成全局唯一主键是非常棘手的问题。同一个逻辑表内的不同实际表之间的自增键由于无法互相感知而产生重复主键。虽然可通过约束自增主键初始值和步长的方式避免碰撞，但需引入额外的运维规则，使解决方案缺乏完整性和可扩展性。

目前有许多第三方解决方案可以完美解决这个问题，如 UUID 等依靠特定算法自生成不重复键，或者通过引入主键生成服务等。为了方便用户使用、满足不同用户不同使用场景的需求，Apache ShardingSphere 不仅提供了内置的分布式主键生成器，例如 UUID、SNOWFLAKE，还抽离出分布式主键生成器的接口，方便用户自行实现自定义的自增主键生成器。

参数解释

雪花算法

类型：SNOWFLAKE

可配置属性：

注意：**worker-id** 为选配项 1. 在单机模式下支持用户自定义配置，如果用户不配置使用默认值为 0。2. 在集群模式下会由系统自动生成，相同的命名空间下不会生成重复的值。

Nanoid

类型：NANOID

可配置属性：无

UUID

类型：UUID

可配置属性：无

Cosid

类型：COSID

可配置属性：

属性名称	• 说明 数 据 类 型 *		• 默认值*
id-name	String	ID 生成器名称	'__shard__'
as-string	bool	是否生成字符串类型 ID: 将 long 类型 ID 转换成 62 进制 String 类型 (Long.MAX_VALUE 最大字符串长度 11 位), 并保证字符串 ID 有序性	'false'

CosId-Snowflake

类型: COSID_SNOWFLAKE

可配置属性:

属性名称	• 说明 数 据 类 型 *		默认值
epoch	String	雪花 ID 算法的 EPOCH	147 792960 0000
as-string	bool	是否生成字符串类型 ID: 将 long 类型 ID 转换成 62 进制 String 类型 (Long.MAX_VALUE 最大字符串长度 11 位), 并保证字符串 ID 有序性	false

操作步骤

1. 配置数据分片规则时为列配置分布式主键生成策略

配置示例

- 雪花算法

```
keyGenerators:
  snowflake:
    type: SNOWFLAKE
```

- NanoID

```
keyGenerators:
  nanoid:
    type: NANOID
```

- UUID

```
keyGenerators:
  nanoid:
    type: UUID
```

负载均衡算法

背景信息

ShardingSphere 内置提供了多种负载均衡算法，具体包括了轮询算法、随机访问算法和权重访问算法，能够满足用户绝大多数业务场景的需要。此外，考虑到业务场景的复杂性，内置算法也提供了扩展方式，用户可以基于 SPI 接口实现符合自己业务需要的负载均衡算法。

参数解释

Type	Describe	Limitations
ROUND_ROBIN	事务内, 读请求路由到 primary, 事务外, 采用轮询策略路由到 replica	
RANDOM	事务内, 读请求路由到 primary, 事务外, 采用随机策略路由到 replica	
WEIGHT	事务内, 读请求路由到 primary, 事务外, 采用权重策略路由到 replica	需配置属性, 属性名: <code>math: {replica-name }</code> , 数据类型: double, 属性名字使用读库名字, 参数填写读库对应的权重值。权重参数范围最小值 > 0, 合计 <= Double.MAX_VALUE。 TRANSACTION_RANDOM 显示/非显示开启事务, 读请求采用随机策略路由到多个 replica TRANSACTION_ROUND_ROBIN 显示/非显示开启事务, 读请求采用轮询策略路由到多个 replica TRANSACTION_WEIGHT 显示/非显示开启事务, 读请求采用权重策略路由到多个 replica 需配置属性, 属性名: <code>{replica-name}</code> , 数据类型: double, 属性名字使用读库名字, 参数填写读库对应的权重值。权重参数范围最小值 > 0, 合计 <= Double.MAX_VALUE。
FIXED_ REPLICARANDOM	显示开启事务, 读请求采用随机策略路由到一个固定 replica; 不开事务, 每次读流量使用随机策略路由到不同的 replica	
FIXED_REPLICAROUND_ROBIN	显示开启事务, 读请求采用轮询策略路由到一个固定 replica; 不开事务, 每次读流量使用轮询策略路由到不同的 replica	
FIXED_ REPLICAWEIGHT	显示开启事务, 读请求采用权重策略路由到一个固定 replica; 不开事务, 每次读流量使用权重策略路由到不同的 replica	需配置属性, 属性名: <code> \${replica-name}</code> , 数据类型: double, 属性名字使用读库名字, 参数填写读库对应的权重值。权重参数范围最小值 > 0, 合计 <= Double.MAX_VALUE。

4.3. 通用配置

操作步骤

1. 使用读写分离时，在 loadBalancers 属性下配置对应的负载均衡算法即可；

配置示例

```
rules:
- !READWRITE_SPLITTING
  dataSources:
    readwrite_ds:
      staticStrategy:
        writeDataSourceName: write_ds
        readDataSourceNames:
          - read_ds_0
          - read_ds_1
        loadBalancerName: random
  loadBalancers:
    random:
      type: RANDOM
```

相关参考

- 核心特性：读写分离
- 开发者指南：读写分离

加密算法

背景信息

加密算法是 Apache ShardingSphere 的加密功能使用的算法，ShardingSphere 内置了多种算法，可以让用户方便使用。

参数解释

MD5 加密算法

类型：MD5

可配置属性：无

AES 加密算法

类型：AES

可配置属性：

名称	数据类型	说明
aes-key-value	String	AES 使用的 KEY

RC4 加密算法

类型：RC4

可配置属性：

名称	数据类型	说明
rc4-key-value	String	RC4 使用的 KEY

SM3 加密算法

类型：SM3

可配置属性：

名称	数据类型	说明
sm3-salt	String	SM3 使用的 SALT（空或 8 Bytes）

SM4 加密算法

类型：SM4

可配置属性：

名称	数据类型	说明
sm4-key	String	SM4 使用的 KEY（16 Bytes）
sm4-mode	String	SM4 使用的 MODE（CBC 或 ECB）
sm4-iv	String	SM4 使用的 IV（MODE 为 CBC 时需指定，16 Bytes）
sm4-padding	String	SM4 使用的 PADDING（PKCS5Padding 或 PKCS7Padding，暂不支持 NoPadding）

操作步骤

1. 在加密规则中配置加密器
2. 为加密器指定加密算法类型

配置示例

```
rules:
- !ENCRYPT
  tables:
    t_user:
      columns:
        username:
          plainColumn: username_plain
          cipherColumn: username
          encryptorName: name-encryptor
      encryptors:
        name-encryptor:
          type: AES
          props:
            aes-key-value: 123456abc
```

相关参考

- 核心特性：数据加密
- 开发者指南：数据加密

影子算法

背景信息

影子库功能对执行的 SQL 语句进行影子判定。影子判定支持两种类型算法，用户可根据实际业务需求选择一种或者组合使用。

参数解释

列影子算法

列值匹配算法

类型：VALUE_MATCH

属性名称	数据类型	说明
column	String	影子列
operation	String	SQL 操作类型 (INSERT, UPDATE, DELETE, SELECT)
value	String	影子列匹配的值

列正则表达式匹配算法

类型: REGEX_MATCH

属性名称	数据类型	说明
column	String	匹配列
operation	String	SQL 操作类型 (INSERT, UPDATE, DELETE, SELECT)
regex	String	影子列匹配正则表达式

Hint 影子算法

简单 Hint 匹配影子算法

类型: SIMPLE_HINT

属性名称	数据类型	说明
foo	String	bar

配置示例

- Java API

```
public final class ShadowConfiguration {  
    // ...  
  
    private AlgorithmConfiguration createShadowAlgorithmConfiguration() {  
        Properties userIdInsertProps = new Properties();  
        userIdInsertProps.setProperty("operation", "insert");  
        userIdInsertProps.setProperty("column", "user_id");  
        userIdInsertProps.setProperty("value", "1");  
        return new AlgorithmConfiguration("VALUE_MATCH", userIdInsertProps);  
    }  
  
    // ...  
}
```

- YAML:

```
shadowAlgorithms:
  user-id-insert-algorithm:
    type: VALUE_MATCH
    props:
      column: user_id
      operation: insert
      value: 1
```

- Spring Boot Starter:

```
spring.shardingsphere.rules.shadow.shadow-algorithms.user-id-insert-algorithm.
type=VALUE_MATCH
spring.shardingsphere.rules.shadow.shadow-algorithms.user-id-insert-algorithm.
props.operation=insert
spring.shardingsphere.rules.shadow.shadow-algorithms.user-id-insert-algorithm.
props.column=user_id
spring.shardingsphere.rules.shadow.shadow-algorithms.user-id-insert-algorithm.
props.value=1
```

- Spring 命名空间:

```
<shadow:shadow-algorithm id="user-id-insert-algorithm" type="VALUE_MATCH">
  <props>
    <prop key="operation">insert</prop>
    <prop key="column">user_id</prop>
    <prop key="value">1</prop>
  </props>
</shadow:shadow-algorithm>
```

SQL 翻译

原生 SQL 翻译器

类型: NATIVE

可配置属性:

无

默认使用的 SQL 翻译器, 但目前暂未实现

使用 JooQ 的 SQL 翻译器

类型: JOOQ

可配置属性:

无

由于需要第三方的JooQ依赖, 因此 ShardingSphere 默认并未包含相关模块, 需要使用下面的 Maven 坐标引用该模块

```
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-sql-translator-jooq-provider</artifactId>
  <version>${project.version}</version>
</dependency>
```

4.4 错误码

本章列举 Apache ShardingSphere 错误码。包含 SQL 错误码和服务器错误码。

本章节所有内容均为草稿, 错误码仍可能调整。

4.4.1 SQL 错误码

SQL 错误码以标准的 SQL State, Vendor Code 和详细错误信息提供, 在 SQL 执行错误时返回给客户端。

目前内容为草稿, 错误码仍可能调整。

SQL State	Vendor Code	错误信息
01000	10000	Circuit break open, the request has been ignored
08000	10001	The URL '%s' is not recognized, please refer to the pattern '%s'
42000	10002	Can not support 3-tier structure for actual data node '%s' with JDBC '%s'
42000	10003	Unsupported SQL node conversion for SQL statement '%s'
HY004	10004	Unsupported conversion data type '%s' for value '%s'
HY004	10100	Can not register driver, reason is: %s
34000	10200	Can not get cursor name from fetch statement
42000	11000	You have an error in your SQL syntax: %s
42000	11001	configuration error
42000	11002	Resource does not exist
42000	11003	Rule does not exist
HY000	11004	File access failed, reason is: %s
42000	11200	Can not support database '%s' in SQL translation
42000	11201	Translation error, SQL is: %s
25000	11320	Switch transaction type failed, please terminate the current transaction

表 1 – 接上页

SQL State	Vendor Code	错误信息
25000	11321	JDBC does not support operations across multiple logical databases in transaction
25000	11322	Failed to create '%s' XA data source
42S02	11400	Can not get traffic execution unit
42000	12000	Unsupported command: %s
44000	13000	SQL check failed, error message: %s
HY000	14000	The table '%s' of schema '%s' is locked
HY000	14001	The table '%s' of schema '%s' lock wait timeout of %s ms exceeded
HY000	14010	Can not find '%s' file for datetime initialize
HY000	14011	Load datetime from database failed, reason: %s
HY000	15000	Work ID assigned failed, which can not exceed 1024
HY000	16000	Can not find pipeline job '%s'
HY000	16001	Failed to get DDL for table '%s'
42S02	17000	Single table '%s' does not exist
42S02	17001	Schema '%s' does not exist
0A000	17002	Can not drop schema '%s' because of contains tables
0A000	17010	DROP TABLE ...CASCADE is not supported
HY000	20000	Sharding algorithm class '%s' should be implement '%s'
44000	20001	Can not get uniformed table structure for logic table '%s', it has different meta data of actual tables are as follows: %s
42S02	20002	Can not get route result, please check your sharding rule configuration
44000	20003	Can not find table rule with logic tables '%s'
44000	20010	Sharding value can't be null in insert statement
HY004	20011	Found different types for sharding value '%s'
44000	20012	Can not update sharding value for table '%s'
0A000	20013	Can not support operation '%s' with sharding table '%s'
0A000	20014	Can not support DML operation with multiple tables '%s'
0A000	20015	The CREATE VIEW statement contains unsupported query statement
42000	20016	%s ...LIMIT can not support route to multiple data nodes
44000	20017	PREPARE statement can not support sharding tables route to same data sources
42000	20018	INSERT INTO ...SELECT can not support applying key generator with absent generate key column
44000	20019	The table inserted and the table selected must be the same or bind tables
44000	20020	Inline sharding algorithms expression '%s' and sharding column '%s' do not match
42S01	20030	Index '%s' already exists
42S02	20031	Index '%s' does not exist
44000	20032	Actual tables '%s' are in use
42S01	20033	View name has to bind to %s tables
HY000	20050	Routed target '%s' does not exist, available targets are '%s'
HY000	20051	'%s %s' can not route correctly for %s '%s'
42S02	20052	Can not find data source in sharding rule, invalid actual data node '%s'
HY004	24000	Encrypt algorithm '%s' initialize failed, reason is: %s
0A000	24001	The SQL clause '%s' is unsupported in encrypt rule
44000	24002	Altered column '%s' must use same encrypt algorithm with previous column '%s' in table '%s'

表 1 – 接上页

SQL State	Vendor Code	错误信息
42000	24003	Insert value of index '%s' can not support for encrypt
44000	24004	Can not find logic encrypt column by '%s'
HY004	25000	Shadow column '%s' of table '%s' does not support '%s' type
42000	25003	Insert value of index '%s' can not support for shadow
HY004	30000	Unknown exception: %s

4.4.2 服务器错误码

服务器发生错误时所提供的唯一错误码，打印在 Proxy 后端或 JDBC 启动日志中。

TODO

Apache ShardingSphere 可插拔架构提供了数十个基于 SPI 的扩展点。对于开发者来说，可以十分方便的对功能进行定制化扩展。

本章节将 Apache ShardingSphere 的 SPI 扩展点悉数列出。如无特殊需求，用户可以使用 Apache ShardingSphere 提供的内置实现；高级用户则可以参考各个功能模块的接口进行自定义实现。

Apache ShardingSphere 社区非常欢迎开发者将自己的实现类反馈至 [开源社区]，让更多用户从中收益。

5.1 运行模式

5.1.1 StandalonePersistRepository

全限定类名

[org.apache.shardingsphere.mode.repository.standalone.StandalonePersistRepository]

定义

单机模式配置信息持久化定义

已知实现

配置标识	详细说明	全限定类名
H2	H2-based persistence	org.apache.shardingsphere.mode.repository.standalone.h2.H2Repository

5.1.2 ClusterPersistRepository

全限定类名

org.apache.shardingsphere.mode.repository.cluster.ClusterPersistRepository

定义

集群模式配置信息持久化定义

已知实现

配置标识	详细说明	全限定类名
ZooKeeper	基于 ZooKeeper 的持久化	org.apache.shardingsphere.mode.repository.cluster.zookeeper.Curator ZookeeperRepository
etcd	基于 Etcd 的持久化	org.apache.shardingsphere.mode.repository.cluster.etcd.EtcdRepository

5.1.3 GovernanceWatcher

全限定类名

org.apache.shardingsphere.mode.manager.cluster.coordinator.registry.GovernanceWatcher

定义

治理监听器定义

已知实现

配置标识	详细说明	全限定类名 *
Types: ADDED, UPDATED, DELETED; WatchingKeys: /nodes/compute_nodes	计算节点状态变化监听器	org.apache.shardingsphere.mode.manager.cluster.coordinator.registry.status.computer.ComputeNodeStateChangeListener
Types: ADDED, DELETED; WatchingKeys: /lock/database/locks	数据库锁状态变化监听器	org.apache.shardingsphere.mode.manager.cluster.coordinator.lock.database.watcher.DatabaseLockChangeListener
Types: ADDED, DELETED; WatchingKeys: /lock/distributed/locks	分布式锁变化监听器	org.apache.shardingsphere.mode.manager.cluster.coordinator.lock.distributed.watcher.DistributedLockChangeListener
Types: UPDATED; WatchingKeys: /rules	全局规则配置变化监听器	org.apache.shardingsphere.mode.manager.cluster.registry.config.watcher.GlobalRuleChangeListener
Types: ADDED, UPDATED, DELETED; WatchingKeys: /metadata/\${databaseName}	元数据变化监听器	org.apache.shardingsphere.mode.manager.cluster.registry.metadata.watcher.MetadataChangeListener
Types: ADDED, UPDATED; WatchingKeys: /props	属性变化监听器	org.apache.shardingsphere.mode.manager.cluster.registry.config.watcher.PropertiesChangeListener
Types: ADDED, UPDATED, DELETED; WatchingKeys: /nodes/storage_nodes	存储节点状态变化监听器	org.apache.shardingsphere.mode.manager.cluster.registry.status.storage.storage.StorageNodeStateChangeListener

5.2 配置

5.2.1 RuleBuilder

全限定类名

`org.apache.shardingsphere.infra.rule.builder.RuleBuilder`

定义

用于将用户配置转化为规则对象的接口

已知实现

配置标识	详细说明	全限定类名 *
AuthorityRuleConfiguration	用于将权限用户配置转化为权限规则对象	org.apache.shardingsphere.authority.rule.builder.AuthorityRuleBuilder
SQLParserRuleConfiguration	用于将 SQL 解析用户配置转化为 SQL 解析规则对象	org.apache.shardingsphere.parser.rule.builder.SQLParserRuleBuilder
TransactionRuleConfiguration	用于将事务用户配置转化为事务规则对象	org.apache.shardingsphere.transaction.rule.builder.TransactionRuleBuilder
SingleTableRuleConfiguration	用于将单表用户配置转化为单表规则对象	org.apache.shardingsphere.singletable.rule.builder.SingleTableRuleBuilder
ShardingRuleConfiguration	用于将分片用户配置转化为分片规则对象	org.apache.shardingsphere.sharding.rule.builder.ShardingRuleBuilder
AlgorithmProvidedShardingRuleConfiguration	用于将基于算法的分片用户配置转化为分片规则对象	org.apache.shardingsphere.sharding.rule.builder.AlgorithmProvidedShardingRuleBuilder
ReadwriteSplittingRuleConfiguration	用于将读写分离用户配置转化为读写分离规则对象	org.apache.shardingsphere.readwrite splitting.rule.builder.ReadwriteSplittingRuleBuilder
AlgorithmReadwriteSplittingRuleConfiguration	用于将基于算法的读写分离用户配置转化为读写分离规则对象	org.apache.shardingsphere.readwrite splitting.rule.builder.AlgorithmProvidedReadwriteSplittingRuleBuilder
DatabaseDiscoveryRuleConfiguration	用于将数据库发现用户配置转化为数据库发现规则对象	org.apache.shardingsphere.database discovery.rule.builder.DatabaseDiscoveryRuleBuilder
AlgorithmDatabaseDiscoveryRuleConfiguration	用于将基于算法的数据库发现用户配置转化为数据库发现规则对象	org.apache.shardingsphere.database discovery.rule.builder.AlgorithmProvidedDatabaseDiscoveryRuleBuilder
EncryptRuleConfiguration	用于将加密用户配置转化为加密规则对象	org.apache.shardingsphere.encrypt.rule.builder.EncryptRuleBuilder
AlgorithmEncryptRuleConfiguration	用于将基于算法的加密用户配置转化为加密规则对象	org.apache.shardingsphere.encrypt.rule.builder.AlgorithmProvidedEncryptRuleBuilder
5.2. 配置		hmProvidedEncryptRuleBuilder 276
ShadowRuleConfiguration	用于将影子库用户配置转化为影子库规则对象	org.apache.shardingsphere.shadow.rule.builder.ShadowRuleBuilder

5.2.2 YamlRuleConfigurationSwapper

全限定类名

org.apache.shardingsphere.infra.yaml.config.swapper.YamlRuleConfigurationSwapper

定义

用于将 YAML 配置转化为标准用户配置

已知实现

配置标识	详细说明	全限定类名
AUTHORITY	用于将权限规则的 YAML 配置转化为权限规则标准配置	org.apache.shardingsphere.authORITY.yaml.swapper.YamlAuthORITYRuleConfigurationSwapper
SQL_PARSER	用于将 SQL 解析的 YAML 配置转化为 SQL 解析标准配置	org.apache.shardingsphere.parser.yaml.swapper.YamlSQLParserRuleConfigurationSwapper
TRANSACTION	用于将事务的 YAML 配置转化为事务标准配置	org.apache.shardingsphere.transaction.yaml.swapper.YamlTransactionRuleConfigurationSwapper
SINGLE	用于将单表的 YAML 配置转化为单表标准配置	org.apache.shardingsphere.singletable.yaml.config.swapper.YamlSingleTableRuleConfigurationSwapper
SHARDING	用于将分片的 YAML 配置转化为分片标准配置	org.apache.shardingsphere.sharding.yaml.swapper.YamlShardingRuleConfigurationSwapper
SHARDING	用于将基于算法的分片配置转化为分片标准配置	org.apache.shardingsphere.sharding.yaml.swapper.YamlShardingRuleAlgorithmProviderConfigurationSwapper
READWRITE_SPLITTING	用于将读写分离的 YAML 配置转化为读写分离标准配置	org.apache.shardingsphere.readwritesplitting.yaml.swapper.YamlReadwriteSplittingRuleConfigurationSwapper
READWRITE_SPLITTING	用于将基于算法的读写分离配置转化为读写分离标准配置	org.apache.shardingsphere.readwritesplitting.yaml.swapper.YamlReadwriteSplittingRuleAlgorithmProviderConfigurationSwapper
DB_DISCOVERY	用于将数据库发现的 YAML 配置转化为数据库发现标准配置	org.apache.shardingsphere.dbdiscovery.yaml.swapper.YamlDatabaseDiscoveryRuleConfigurationSwapper
DB_DISCOVERY	用于将基于算法的数据库发现配置转化为数据库发现标准配置	org.apache.shardingsphere.dbdiscovery.yaml.swapper.YamlDatabaseDiscoveryRuleAlgorithmProviderConfigurationSwapper
ENCRYPT	用于将加密的 YAML 配置转化为加密标准配置	org.apache.shardingsphere.encrypt.yaml.swapper.YamlEncryptRuleConfigurationSwapper
ENCRYPT	用于将基于算法的加密配置转化为加密标准配置	org.apache.shardingsphere.encrypt.yaml.swapper.YamlEncryptRuleAlgorithmProviderConfigurationSwapper
SHADOW	用于将影子库的 YAML 配置转化为影子库标准配置	org.apache.shardingsphere.shadow.yaml.swapper.YamlShadowRuleConfigurationSwapper
SHADOW	用于将基于算法的影子库配置转化为影子库标准配置	org.apache.shardingsphere.shadow.yaml.swapper.YamlShadowRuleAlgorithmProviderConfigurationSwapper
SQL_TRANSLATOR	用于将 SQL 转换的 YAML 配置转化为 SQL 转换标准配置	org.apache.shardingsphere.sqltranslator.yaml.swapper.YamlSQLTranslatorRuleConfigurationSwapper

5.2.3 ShardingSphereYamlConstruct

全限定类名

org.apache.shardingsphere.infra.yaml.engine.constructor.ShardingSphereYamlConstruct

定义

用于将定制化对象和 YAML 相互转化

已知实现

配置标识	详细说明	全限定类名 *
YamlNoneShardingStrategy-Configuration	用于将不分片策略对象和YAML相互转化	org.apache.shardingsphere.yaml.engine.constructor.NoneShardingStrategyConfigurationYamlConstruct

5.3 内核

5.3.1 SQLRouter

全限定类名

org.apache.shardingsphere.infra.route.SQLRouter

定义

用于处理路由结果

已知实现

配置标识	详细说明	全限定类名
SingleTableRule	用于处理单表路由结果	org.apache.shardingsphere.singletable.route.SingleTableSQLRouter
ShardingRule	用于处理分片路由结果	org.apache.shardingsphere.sharding.route.engine.ShardingSQLRouter
ReadwriteSplittingRule	用于处理读写分离路由结果	org.apache.shardingsphere.readwritesplitting.route.ReadwriteSplittingSQLRouter
DatabaseDiscoveryRule	用于处理数据库发现路由结果	org.apache.shardingsphere.dbdiscovery.route.DatabaseDiscoverySQLRouter
ShadowRule	用于处理影子库路由结果	org.apache.shardingsphere.shadow.route.ShadowSQLRouter

5.3.2 SQLRewriteContextDecorator

全限定类名

org.apache.shardingsphere.infra.rewrite.context.SQLRewriteContextDecorator

定义

用于处理 SQL 改写结果

已知实现

配置标识	详细说明	全限定类名
ShardingRule	用于处理分片 SQL 改写结果	org.apache.shardingsphere.sharding.rewrite.context.ShardingSQLRewriteContextDecorator
EncryptRule	用于处理加密 SQL 改写结果	org.apache.shardingsphere.encrypt.rewrite.context.EncryptSQLRewriteContextDecorator

5.3.3 SQLExecutionHook

全限定类名

org.apache.shardingsphere.infra.executor.sql.hook.SQLExecutionHook

定义

SQL 执行过程监听器

已知实现

配置标识	详细说明	全限定类名
无	基于事务的 SQL 执行过程监听器	org.apache.shardingsphere.transaction.base.transactional.sql.execution.hook.TransactionalSQLExecutionHook

5.3.4 ResultProcessEngine

全限定类名

org.apache.shardingsphere.infra.merge.engine.ResultProcessEngine

定义

用于处理结果集

已知实现

配置标识	详细说明	全限定类名
ShardingRule	用于处理分片结果集归并	org.apache.shardingsphere.sharding.merge.ShardingResultMergerEngine
EncryptRule	用于处理加密结果集改写	org.apache.shardingsphere.encrypt.merge.EncryptResultDecoratorEngine

5.4 数据源

5.4.1 DatabaseType

全限定类名

org.apache.shardingsphere.infra.database.type.DatabaseType

定义

支持的数据库类型

已知实现

配置标识	详细说明	全限定类名
SQL92	遵循 SQL92 标准的数据库类型	org.apache.shardingsphere.infra.database.type.dialect.SQL92DatabaseType
MySQL	MySQL 数据库	org.apache.shardingsphere.infra.database.type.dialect.MySQLDatabaseType
MariaDB	MariaDB 数据库	org.apache.shardingsphere.infra.database.type.dialect.MariaDBDatabaseType
PostgreSQL	PostgreSQL 数据库	org.apache.shardingsphere.infra.database.type.dialect.PostgreSQLDatabaseType
Oracle	Oracle 数据库	org.apache.shardingsphere.infra.database.type.dialect.OracleDatabaseType
SQLServer	SQLServer 数据库	org.apache.shardingsphere.infra.database.type.dialect.SQLServerDatabaseType
H2	H2 数据库	org.apache.shardingsphere.infra.database.type.dialect.H2DatabaseType
openGauss	OpenGauss 数据库	org.apache.shardingsphere.infra.database.type.dialect.OpenGaussDatabaseType

5.4.2 DialectSchemaMetadataLoader

全限定类名

org.apache.shardingsphere.infra.metadata.database.schema.loader.spi.DialectSchemaMetadataLoader

定义

使用 SQL 方言快速加载元数据

已知实现

配置标识	详细说明	全限定类名
MySQL	使用 MySQL 方言加载元数据	org.apache.shardingsphere.infra.metadata.database.schema.loader.dialect.MySQLSchemaMetadataLoader
Oracle	使用 Oracle 方言加载元数据	org.apache.shardingsphere.infra.metadata.database.schema.loader.dialect.OracleSchemaMetadataLoader
PostgreSQL	使用 PostgreSQL 方言加载元数据	org.apache.shardingsphere.infra.metadata.database.schema.loader.dialect.PostgreSQLSchemaMetadataLoader
SQLServer	使用 SQLServer 方言加载元数据	org.apache.shardingsphere.infra.metadata.database.schema.loader.dialect.SQLServerSchemaMetadataLoader
H2	使用 H2 方言加载元数据	org.apache.shardingsphere.infra.metadata.database.schema.loader.dialect.H2SchemaMetadataLoader
open-Gauss	使用 OpenGauss 方言加载元数据	org.apache.shardingsphere.infra.metadata.database.schema.loader.dialect.OpenGaussSchemaMetadataLoader

5.4.3 DataSourcePoolMetadata

全限定类名

org.apache.shardingsphere.infra.datasource.pool.metadata.DataSourcePoolMetadata

定义

数据源连接池元数据

已知实现

配置标识	详细说明	全限定类名
org.apache.commons.dbcp.BasicDataSource org.apache.tomcat.dbcp.dbcp2.BasicDataSource	DBCP 数据源连接池元数据	org.apache.shardingsphere.infra.datasource.pool.metadata.type.dbcp.DBCPDataSourcePoolMetadata
com.zaxxer.hikari.HikariDataSource	Hikari 数据源连接池元数据	org.apache.shardingsphere.infra.datasource.pool.metadata.type.hikari.HikariDataSourcePoolMetadata
com.mchange.v2.c3p0.ComboPooledDataSource	C3P0 数据源连接池元数据	org.apache.shardingsphere.infra.datasource.pool.metadata.type.c3p0.C3P0DataSourcePoolMetadata

5.4.4 DataSourcePoolActiveDetector

全限定类名

org.apache.shardingsphere.infra.datasource.pool.destroyer.detector.DataSourcePoolActiveDetector

定义

数据源连接池活跃探测器

已知实现

配置标识	详细说明	全限定类名
Default	默认数据源连接池活跃探测器	org.apache.shardingsphere.infra.datasource.pool.destroyer.detector.type.DefaultDataSourcePoolActiveDetector
com.zaxxer.hikari.HikariDataSource	Hikari 数据源连接池活跃探测器	org.apache.shardingsphere.infra.datasource.pool.destroyer.detector.type.HikariDataSourcePoolActiveDetector

5.5 SQL 解析

5.5.1 DatabaseTypedSQLParserFacade

全限定类名

org.apache.shardingsphere.sql.parser.spi.DatabaseTypedSQLParserFacade

定义

配置用于 SQL 解析的词法分析器和语法分析器入口

已知实现

配置标识	详细说明	全限定类名
MySQL	基于 MySQL 的 SQL 解析器入口	org.apache.shardingsphere.sql.parser.mysql.parser.MySQLParserFacade
PostgreSQL	基于 PostgreSQL 的 SQL 解析器入口	org.apache.shardingsphere.sql.parser.postgresql.parser.PostgreSQLParserFacade
SQLServer	基于 SQLServer 的 SQL 解析器入口	org.apache.shardingsphere.sql.parser.sqlserver.parser.SQLServerParserFacade
Oracle	基于 Oracle 的 SQL 解析器入口	org.apache.shardingsphere.sql.parser.oracle.parser.OracleParserFacade
SQL92	基于 SQL92 的 SQL 解析器入口	org.apache.shardingsphere.sql.parser.sql92.parser.SQL92ParserFacade
openGauss	基于 openGauss 的 SQL 解析器入口	org.apache.shardingsphere.sql.parser.opengauss.parser.OpenGaussParserFacade

5.5.2 SQLVisitorFacade

全限定类名

org.apache.shardingsphere.sql.parser.spi.SQLVisitorFacade

定义

SQL 语法树访问器入口

已知实现

配置标识	详细说明	全限定类名
MySQL	基于 MySQL 的 SQL 语法树访问器入口	org.apache.shardingsphere.sql.parser.mysql.visitor.statement.facade.MySQLStatementSQLVisitorFacade
PostgreSQL	基于 PostgreSQL 的 SQL 语法树访问器入口	org.apache.shardingsphere.sql.parser.postgresql.visitor.statement.facade.PostgreSQLStatementSQLVisitorFacade
SQLServer	基于 SQLServer 的 SQL 语法树访问器入口	org.apache.shardingsphere.sql.parser.sqlserver.visitor.statement.facade.SQLServerStatementSQLVisitorFacade
Oracle	基于 Oracle 的 SQL 语法树访问器入口	org.apache.shardingsphere.sql.parser.oracle.visitor.statement.facade.OracleStatementSQLVisitorFacade
SQL92	基于 SQL92 的 SQL 语法树访问器入口	org.apache.shardingsphere.sql.parser.sql92.visitor.statement.facade.SQL92StatementSQLVisitorFacade
openGauss	基于 openGauss 的 SQL 语法树访问器入口	org.apache.shardingsphere.sql.parser.opengauss.visitor.statement.facade.OpenGaussStatementSQLVisitorFacade

5.6 代理端

5.6.1 DatabaseProtocolFrontendEngine

全限定类名

org.apache.shardingsphere.proxy.frontend.spi.DatabaseProtocolFrontendEngine

定义

用于 ShardingSphere-Proxy 解析与适配访问数据库的协议

已知实现

配置标识	详细说明	全限定类名
MySQL	MySQL 协议实现	org.apache.shardingsphere.proxy.frontend.mysql.MySQLFrontendEngine
PostgreSQL	PostgreSQL 协议实现	org.apache.shardingsphere.proxy.frontend.postgresql.PostgreSQLFrontendEngine
openGauss	openGauss 协议实现	org.apache.shardingsphere.proxy.frontend.opengauss.OpenGaussFrontendEngine

5.6.2 AuthorityProvideAlgorithm

全限定类名

org.apache.shardingsphere.authority.spi.AuthorityProviderAlgorithm

定义

用户权限加载逻辑

已知实现

配置标识	详细说明	全限定类名
A L L _ P E R M I T T E D	默认授予所有权限（不鉴权）	org.apache.shardingsphere.authority.provider.simple.AllPermittedPrivilegesProvider
D A T A B A S E _ P E R M I T T E D	通过属性 u s e r - d a t a b a s e - m a p p i n g s 配置的权限	org.apache.shardingsphere.authority.provider.database.DatabasePermittedPrivilegesProvider

5.7 数据分片

5.7.1 ShardingAlgorithm

全限定类名

org.apache.shardingsphere.sharding.spi.ShardingAlgorithm

定义

分片算法

已知实现

配置标识 *	自动分片算法 *	详细说明	全限定类名
MOD	Y	基于取模的分片算法	org.apache.shardingsphere.sharding.algorithm.mod.ModShardingAlgorithm
HASH_MOD	Y	基于哈希取模的分片算法	org.apache.shardingsphere.sharding.algorithm.hash.HashModShardingAlgorithm
BOUNDARY_RANGE	Y	基于分片边界的范围分片算法	org.apache.shardingsphere.sharding.algorithm.range.BoundaryBasedRangeShardingAlgorithm
VOLUME_RANGE	Y	基于分片容量的范围分片算法	org.apache.shardingsphere.sharding.algorithm.range.VolumeBasedRangeShardingAlgorithm
AUTO_INTERVAL	Y	基于可变时间范围的分片算法	org.apache.shardingsphere.sharding.algorithm.datetime.AutoIntervalShardingAlgorithm
INTERVAL	N	基于固定时间范围的分片算法	org.apache.shardingsphere.sharding.algorithm.datetime.IntervalShardingAlgorithm
CLASS_BASED	N	基于自定义类的分片算法	org.apache.shardingsphere.sharding.algorithm.classbased.ClassBasedShardingAlgorithm
INLINE	N	基于行表达式的分片算法	org.apache.shardingsphere.sharding.algorithm.inline.InlineShardingAlgorithm

5.7.2 KeyGenerateAlgorithm

全限定类名

org.apache.shardingsphere.sharding.spi.KeyGenerateAlgorithm

定义

分布式主键生成算法

已知实现

配置标识	详细说明	全限定类名
SNOWFLAKE	基于雪花算法的分布式主键生成算法	org.apache.shardingsphere.sharding.algorithm.keygen.SnowflakeKeyGenerateAlgorithm
UUID	基于 UUID 的分布式主键生成算法	org.apache.shardingsphere.sharding.algorithm.keygen.UUIDKeyGenerateAlgorithm
NANOID	基于 NanoId 的分布式主键生成算法	org.apache.shardingsphere.sharding.algorithm.keygen.NanoIdKeyGenerateAlgorithm
COSID	基于 CosId 的分布式主键生成算法	org.apache.shardingsphere.sharding.algorithm.keygen.CosIdKeyGenerateAlgorithm
COSID_SNOWFLAKE	基于 CosId 的雪花算法分布式主键生成算法	org.apache.shardingsphere.sharding.algorithm.keygen.CosIdSnowflakeKeyGenerateAlgorithm

5.7.3 ShardingAuditAlgorithm

全限定类名

org.apache.shardingsphere.sharding.spi.ShardingAuditAlgorithm

定义

分片审计算法

已知实现

配置标识	详细说明	全限定类名
DML_SHARDING_CONDITIONS	禁止不带分片键的 DML 审计算法	org.apache.shardingsphere.sharding.algorithm.audit.DMLShardingConditionsShardingAuditAlgorithm

5.7.4 DatetimeService

全限定类名

org.apache.shardingsphere.infra.datetime.DatetimeService

定义

获取当前时间进行路由

已知实现

配置标识	详细说明	全限定类名
DatabaseDatetimeService	从数据库中获取当前时间进行路由	org.apache.shardingsphere.agent.metrics.prometheus.service.PrometheusPluginBootService
SystemDatetime	从应用系统时间中获取当前时间进行路由	org.apache.shardingsphere.datetime.system.SystemDatetimeService

5.8 读写分离

5.8.1 ReadQueryLoadBalanceAlgorithm

全限定类名

org.apache.shardingsphere.readwritesplitting.spi.ReadQueryLoadBalanceAlgorithm

定义

读库负载均衡算法

已知实现

配置标识 *	详细说明	全限定类名
ROUND_ROBIN	基于轮询的读库负载均衡算法	org.apache.shardingsphere.readwritesplitting.algorithm.loadbalance.RoundRobinReadQueryLoadBalanceAlgorithm
RANDOM	基于随机的读库负载均衡算法	org.apache.shardingsphere.readwritesplitting.algorithm.loadbalance.RandomReadQueryLoadBalanceAlgorithm
WEIGHT	基于权重的读库负载均衡算法	org.apache.shardingsphere.readwritesplitting.algorithm.loadbalance.WeightReadQueryLoadBalanceAlgorithm
TRANSACTION_RANDOM	无论是否在事务中，读请求采用随机策略路由到多个读库	org.apache.shardingsphere.readwritesplitting.algorithm.loadbalance.TransactionRandomReadQueryLoadBalanceAlgorithm
TRANSACTION_ROUND_ROBIN	无论是否在事务中，读请求采用轮询策略路由到多个读库	org.apache.shardingsphere.readwritesplitting.algorithm.loadbalance.TransactionRoundRobinReadQueryLoadBalanceAlgorithm
TRANSACTION_WEIGHT	无论是否在事务中，读请求采用权重策略路由到多个读库	org.apache.shardingsphere.readwritesplitting.algorithm.loadbalance.TransactionWeightReadQueryLoadBalanceAlgorithm
FIXED_REPLICA_RANDOM	显示开启事务，读请求采用随机策略路由到一个固定读库；不开事务，每次读流量使用指定算法路由到不同的读库	org.apache.shardingsphere.readwritesplitting.algorithm.loadbalance.FixedReplicaRandomReadQueryLoadBalanceAlgorithm
FIXED_REPLICA_ROUND_ROBIN	显示开启事务，读请求采用轮询策略路由到一个固定读库；不开事务，每次读流量使用指定算法路由到不同的读库	org.apache.shardingsphere.readwritesplitting.algorithm.loadbalance.FixedReplicaRoundRobinReadQueryLoadBalanceAlgorithm

5.8. 读写分离

5.9 高可用

5.9.1 DatabaseDiscoveryProviderAlgorithm

全限定类名

org.apache.shardingsphere.dbdiscovery.spi.DatabaseDiscoveryProviderAlgorithm

定义

数据库发现提供算法的定义

已知实现

配置标识	详细说明	全限定类名
MySQL.MGR	基于 MySQL MGR 的数据库发现算法	org.apache.shardingsphere.dbdiscovery.mysql.type.MGRMySQLDatabaseDiscoveryProviderAlgorithm
MySQL.NORMAL_REPLICATION	基于 MySQL 主从同步的数据库发现算法	org.apache.shardingsphere.dbdiscovery.mysql.type.MySQLNormalReplicationDatabaseDiscoveryProviderAlgorithm
openGauss.NORMAL_REPLICATION	基于 openGauss 主从同步的数据库发现算法	org.apache.shardingsphere.dbdiscovery.opengauss.OpenGaussNormalReplicationDatabaseDiscoveryProviderAlgorithm

5.10 分布式事务

5.10.1 ShardingSphereTransactionManager

全限定类名

org.apache.shardingsphere.transaction.spi.ShardingSphereTransactionManager

定义

分布式事务管理器

已知实现

配置标识	详细说明	全限定类名
XA	基于 XA 的分布式事务管理器	org.apache.shardingsphere.transaction.xa.XAShardingSphereTransactionManager
BASE	基于 Seata 的分布式事务管理器	org.apache.shardingsphere.transaction.base.seata.at.SeataATShardingSphereTransactionManager

5.10.2 XATransactionManagerProvider

全限定类名

org.apache.shardingsphere.transaction.xa.spi.XATransactionManagerProvider

定义

XA 分布式事务管理器

已知实现

配置标识 *	详细说明	全限定类名 *
Atomikos	基于 Atomikos 的 XA 分布式事务管理器	org.apache.shardingsphere.transaction.xa.atomikos.manager.AtomikosTransactionManagerProvider
Narayana	基于 Narayana 的 XA 分布式事务管理器	org.apache.shardingsphere.transaction.xa.narayana.manager.NarayanaXATransactionManagerProvider
Bitronix	基于 Bitronix 的 XA 分布式事务管理器	org.apache.shardingsphere.transaction.xa.bitronix.manager.BitronixXATransactionManagerProvider

5.10.3 XADataSourceDefinition

全限定类名

org.apache.shardingsphere.transaction.xa.jta.datasource.properties.XADataSourceDefinition

定义

用于非 XA 数据源转化为 XA 数据源

已知实现

配置标识	详细说明	全限定类名
MySQL	非 XA 的 MySQL 数据源自动转化为 XA 的 MySQL 数据源	org.apache.shardingsphere.transaction.xa.jta.datasource.properties.dialect.MySQLXADatasourceDefinition
MariaDB	非 XA 的 MariaDB 数据源自动转化为 XA 的 MariaDB 数据源	org.apache.shardingsphere.transaction.xa.jta.datasource.properties.dialect.MariaDBXADatasourceDefinition
PostgreSQL	非 XA 的 PostgreSQL 数据源自动转化为 XA 的 PostgreSQL 数据源	org.apache.shardingsphere.transaction.xa.jta.datasource.properties.dialect.PostgreSQLXADatasourceDefinition
Oracle	非 XA 的 Oracle 数据源自动转化为 XA 的 Oracle 数据源	org.apache.shardingsphere.transaction.xa.jta.datasource.properties.dialect.OracleXADatasourceDefinition
SQL Server	非 XA 的 SQLServer 数据源自动转化为 XA 的 SQLServer 数据源	org.apache.shardingsphere.transaction.xa.jta.datasource.properties.dialect.SQLServerXADatasourceDefinition
H2	非 XA 的 H2 数据源自动转化为 XA 的 H2 数据源	org.apache.shardingsphere.transaction.xa.jta.datasource.properties.dialect.H2XADatasourceDefinition

5.10.4 DataSourcePropertyProvider

全限定类名

org.apache.shardingsphere.transaction.xa.jta.datasource.swapper.DataSourcePropertyProvider

定义

用于获取数据源连接池的标准属性

已知实现

配置标识	详细说明	全限定类名 *
com.zaxxer.hikari.HikariDataSource	用于获取 HikariCP 连接池的标准属性	org.apache.shardingsphere.transaction.xa.jta.datasource.wrapper.impl.HikariCPPropertyProvider

5.11 SQL 检查

5.11.1 全限定类名

org.apache.shardingsphere.infra.executor.check.SQLChecker

5.11.2 定义

SQL 检查定义接口

5.11.3 已知实现

配置标识	详细说明	全限定类名
AuthorityRule.class	权限检查器	org.apache.shardingsphere.authority.checker.AuthorityChecker
ShardingRule.class	分片审计检查器	org.apache.shardingsphere.sharding.checker.audit.ShardingAuditChecker

5.12 数据加密

5.12.1 EncryptAlgorithm

全限定类名

org.apache.shardingsphere.encrypt.spi.EncryptAlgorithm

定义

数据加密算法

已知实现

配置标识	详细说明	全限定类名	
MD5	基于 MD5 的数据加密算法	org.apache.shardingsphere.encrypt ion.algorithm.MD5Encrypt	e.shardingsphere.encrypt
AES	基于 AES 的数据加密算法	org.apache.shardingsphere.encrypt ion.algorithm.AESEncrypt	e.shardingsphere.encrypt
RC4	基于 RC4 的数据加密算法	org.apache.shardingsphere.encrypt ion.algorithm.RC4Encrypt	e.shardingsphere.encrypt
SM3	基于 SM3 的数据加密算法	org.apache.shardingsphere.encrypt ion.algorithm.SM3Encrypt	e.shardingsphere.encrypt
SM4	基于 SM4 的数据加密算法	org.apache.shardingsphere.encrypt ion.algorithm.SM4Encrypt	e.shardingsphere.encrypt

5.13 影子库

5.13.1 ShadowAlgorithm

全限定类名

org.apache.shardingsphere.shadow.spi.ShadowAlgorithm

定义

影子库提供的影子算法

已知实现

已知实现类	详细说明	完全限定类名
ColumnValueMatchedShadowAlgorithm	基于字段值匹配影子算法	org.apache.shardingsphere.shadow.algorithm.shadow.column.ColumnValueMatchedShadowAlgorithm
ColumnRegexMatchedShadowAlgorithm	基于字段值正则匹配影子算法	org.apache.shardingsphere.shadow.algorithm.shadow.column.ColumnRegexMatchedShadowAlgorithm
SimpleHintShadowAlgorithm	基于 Hint 简单匹配影子算法	org.apache.shardingsphere.shadow.algorithm.shadow.hint.SimpleHintShadowAlgorithm

5.14 可观察性

5.14.1 PluginBootService

全限定类名

org.apache.shardingsphere.agent.spi.boot.PluginBootService

定义

插件启动服务定义接口

已知实现

配置标识	详细说明	全限定类名
Prometheus	Prometheus plugin 启动类	org.apache.shardingsphere.agent.metrics.prometheus.service.PrometheusPluginBootstrapService
Logging	Logging plugin 启动类	org.apache.shardingsphere.agent.plugin.logging.base.service.BaseLoggingPluginBootstrapService
Jaeger	Jaeger plugin 启动类	org.apache.shardingsphere.agent.plugin.tracing.jaeger.service.JaegerTracingPluginBootstrapService
Open-Telemetry	OpenTelemetry-Tracing plugin 启动类	org.apache.shardingsphere.agent.plugin.tracing.opentelemetry.service.OpenTelemetryTracingPluginBootstrapService
Open-Tracing	OpenTracing plugin 启动类	org.apache.shardingsphere.agent.plugin.tracing.opentracing.service.OpenTracingPluginBootstrapService
Zipkin	Zipkin plugin 启动类	org.apache.shardingsphere.agent.plugin.tracing.zipkin.service.ZipkinTracingPluginBootstrapService

5.14.2 PluginDefinitionService

全限定类名

org.apache.shardingsphere.agent.spi.definition.PluginDefinitionService

定义

探针插件定义服务接口

已知实现

配置标识	详细说明	全限定类名 *
Prometheus	Prometheus 插件定义	org.apache.shardingsphere.agent.metrics.definition.PrometheusPluginDefinitionService
Logging	Logging 插件定义	org.apache.shardingsphere.agent.plugin.logging.base.definition.BaseLoggingPluginDefinitionService
Jaeger	Jaeger 插件定义	org.apache.shardingsphere.agent.plugin.tracing.jaeger.definition.JaegerPluginDefinitionService
OpenTelemetry	OpenTelemetryTracing 插件定义	org.apache.shardingsphere.agent.plugin.tracing.opentelemetry.definition.OpenTelemetryTracingPluginDefinitionService
OpenTracing	OpenTracing 插件定义	org.apache.shardingsphere.agent.plugin.tracing.opentracing.definition.OpenTracingPluginDefinitionService
Zipkin	Zipkin 插件定义	org.apache.shardingsphere.agent.plugin.tracing.zipkin.definition.ZipkinPluginDefinitionService

Apache ShardingSphere 提供了完善的整合测试、模块测试和性能测试。

6.1 整合测试

通过真实的 Apache ShardingSphere 和数据库的连接，提供端到端的测试。

整合测试引擎以 XML 方式定义 SQL，分别为各个数据库独立运行测试用例。为了方便上手，测试引擎无需修改任何 **Java** 代码，只需修改相应的配置文件即可运行断言。测试引擎不依赖于任何第三方环境，用于测试的 ShardingSphere-Proxy 计算节点和数据库均由 Docker 镜像提供。

6.2 模块测试

将复杂的模块单独提炼成为测试引擎。

模块测试引擎同样以 XML 方式定义 SQL，分别为各个数据库独立运行测试用例，包括 SQL 解析和 SQL 改写模块。

6.3 性能测试

提供多样性的性能测试方法，包括 Sysbench、JMH、TPCC 等。

6.4 集成测试

6.4.1 设计

集成测试包括 3 个模块：测试用例、测试环境以及测试引擎。

测试用例

用于定义待测试的 SQL 以及测试结果的断言数据。每个用例定义一条 SQL，SQL 可定义多种数据库执行类型。

测试环境

用于搭建运行测试用例的数据库和 ShardingSphere-Proxy 环境。环境又具体分为环境准备方式，数据库类型和场景。

环境准备方式分为 Native 和 Docker，未来还将增加 Embed 类型的支持。

- Native 环境用于测试用例直接运行在开发者提供的测试环境中，适用于调试场景。
- Docker 环境由 Testcontainer 创建，适用于云编译环境和测试 ShardingSphere-Proxy 的场景，如：GitHub Action。
- Embed 环境由测试框架自动搭建嵌入式 MySQL，适用于 ShardingSphere-JDBC 的本地环境测试。

当前默认采用 Docker 环境，使用 Testcontainer 创建运行时环境并执行测试用例。未来将采用 Embed 环境的 ShardingSphere-JDBC + MySQL，替换 Native 执行测试用例的默认环境类型。

数据库类型目前支持 MySQL、PostgreSQL、SQLServer 和 Oracle，并且可以支持使用 ShardingSphere-JDBC 或是使用 ShardingSphere-Proxy 执行测试用例。

场景用于对 ShardingSphere 支持规则进行测试，目前支持数据分片和读写分离的相关场景，未来会不断完善场景的组合。

测试引擎

用于批量读取测试用例，并逐条执行和断言测试结果。

测试引擎通过将用例和环境进行排列组合，以达到用最少的用例测试尽可能多场景的目的。

每条 SQL 会以 数据库类型 * 接入端类型 * SQL 执行模式 * JDBC 执行模式 * 场景的组合方式生成测试报告，目前各个维度的支持情况如下：

- 数据库类型：H2、MySQL、PostgreSQL、SQLServer 和 Oracle；
- 接入端类型：ShardingSphere-JDBC 和 ShardingSphere-Proxy；
- SQL 执行模式：Statement 和 PreparedStatement；
- JDBC 执行模式：execute 和 executeQuery（查询）/ executeUpdate（更新）；
- 场景：分库、分表、读写分离和分库分表 + 读写分离。

因此，1 条 SQL 会驱动：数据库类型（5）* 接入端类型（2）* SQL 执行模式（2）* JDBC 执行模式（2）* 场景（4）= 160 个测试用例运行，以达到项目对于高质量的追求。

6.4.2 使用指南

模 块 路 径: shardingsphere-test/shardingsphere-integration-test/shardingsphere-integration-test-suite

测试用例配置

SQL 用例在 resources/cases/\${SQL-TYPE}/\${SQL-TYPE}-integration-test-cases.xml。

用例文件格式如下:

```
<integration-test-cases>
  <test-case sql="${SQL}">
    <assertion parameters="${value_1}:${type_1}, ${value_2}:${type_2}"
expected-data-file="${dataset_file_1}.xml" />
    <!-- ... more assertions -->
    <assertion parameters="${value_3}:${type_3}, ${value_4}:${type_4}"
expected-data-file="${dataset_file_2}.xml" />
  </test-case>

  <!-- ... more test cases -->
</integration-test-cases>
```

expected-data-file 的查找规则是: 1. 查找同级目录中 dataset\\${SCENARIO_NAME}\\${DATABASE_TYPE}\\${dataset_file}.xml 文件; 2. 查找同级目录中 dataset\\${SCENARIO_NAME}\\${dataset_file}.xml 文件; 3. 查找同级目录中 dataset\\${dataset_file}.xml 文件; 4. 都找不到则报错。

断言文件格式如下:

```
<dataset>
  <metadata>
    <column name="column_1" />
    <!-- ... more columns -->
    <column name="column_n" />
  </metadata>
  <row values="value_01, value_02" />
  <!-- ... more rows -->
  <row values="value_n1, value_n2" />
</dataset>
```

环境配置

`${SCENARIO-TYPE}` 表示场景名称，在测试引擎运行中用于标识唯一场景。`${DATABASE-TYPE}` 表示数据库类型。

Native 环境配置

目录: `src/test/resources/env/scenario/${SCENARIO-TYPE}`

- `scenario-env.properties`: 数据源配置;
- `rules.yaml`: 规则配置;
- `databases.xml`: 真实库名称;
- `dataset.xml`: 初始化数据;
- `init-sql\${DATABASE-TYPE}\init.sql`: 初始化数据库表结构;
- `authority.xml`: 待补充。

Docker 环境配置

目录: `src/test/resources/env/${SCENARIO-TYPE}`

- `proxy/conf/config-${SCENARIO-TYPE}.yaml`: 规则配置。

Docker 环境配置为 **ShardingSphere-Proxy** 提供了远程调试端口，可以在“`shardingsphere-test/shardingsphere-integration-test/shardingsphere-integration-test-fixture/src/test/assembly/bin/start.sh`”文件的“`JAVA_OPTS`”中找到第 2 个暴露的端口用于远程调试。

运行测试引擎

配置测试引擎运行环境

通过配置 `src/test/resources/env/engine-env.properties` 控制测试引擎。

所有的属性值都可以通过 Maven 命令行 `-D` 的方式动态注入。

```
# 运行模式，多个值可用逗号分隔。可选值: Standalone, Cluster
it.run.modes=Cluster

# 场景类型，多个值可用逗号分隔。可选值: db, tbl, dbtbl_with_replica_query, replica_query
it.scenarios=db,tbl,dbtbl_with_replica_query,replica_query

# 是否运行附加测试用例
it.run.additional.cases=false

# 配置环境类型，只支持单值。可选值: docker 或空，默认值: 空
```

```
it.cluster.env.type=${it.env}
# 待测试的接入端类型, 多个值可用逗号分隔。可选值: jdbc, proxy, 默认值: jdbc
it.cluster.adapters=jdbc

# 场景类型, 多个值可用逗号分隔。可选值: H2, MySQL, Oracle, SQLServer, PostgreSQL
it.cluster.databases=H2,MySQL,Oracle,SQLServer,PostgreSQL
```

运行调试模式

- 标准测试引擎运行 `org.apache.shardingsphere.test.integration.engine.${SQL-TYPE}.General${SQL-TYPE}IT` 以启动不同 SQL 类型的测试引擎。
- 批量测试引擎运行 `org.apache.shardingsphere.test.integration.engine.dml.BatchDMLIT`, 以启动为 DML 语句提供的测试 `addBatch()` 的批量测试引擎。
- 附加测试引擎运行 `org.apache.shardingsphere.test.integration.engine.${SQL-TYPE}.Additional${SQL-TYPE}IT` 以启动使用更多 JDBC 方法调用的测试引擎。附加测试引擎需要通过设置 `it.run.additional.cases=true` 开启。

运行 Docker 模式

```
./mvnw -B clean install -f shardingsphere-test/shardingsphere-integration-test/pom.xml -Pit.env.docker -Dit.cluster.adapters=proxy,jdbc -Dit.scenarios=${scenario_name_1,scenario_name_2,scenario_name_n} -Dit.cluster.databases=MySQL
```

运行以上命令会构建出一个用于集成测试的 Docker 镜像 `apache/shardingsphere-proxy-test:latest`。如果仅修改了测试代码, 可以复用已有的测试镜像, 无须重新构建。使用以下命令可以跳过镜像构建, 直接运行集成测试:

```
./mvnw -B clean install -f shardingsphere-test/shardingsphere-integration-test/shardingsphere-integration-test-suite/pom.xml -Pit.env.docker -Dit.cluster.adapters=proxy,jdbc -Dit.scenarios=${scenario_name_1,scenario_name_2,scenario_name_n} -Dit.cluster.databases=MySQL
```

远程 debug Docker 容器中的 Proxy 代码

IT 测试的 Proxy 镜像默认开启了 3308 端口用于远程调试容器中的实例。

使用 IDEA 等 IDE 工具可以通过如下方式连接并 debug 容器中的 Proxy 代码:

IDEA -> Run -> Edit Configurations -> Add New Configuration -> Remote JVM Debug

编辑对应的信息: - Name: 一个描述性的名字, 例如 `docker-debug`。- Host: 可以访问 docker 的 IP, 例如 `127.0.0.1`。- Port: 调试端口 3308。- use module classpath: 项目根目录 `shardingsphere`。

编辑好上面的信息后, 在 IDEA 中 Run -> Run -> `docker-debug` 即可启动 IDEA 的远程 debug。

注意事项

1. 如需测试 Oracle，请在 pom.xml 中增加 Oracle 驱动依赖；
2. 为了保证测试数据的完整性和易读性，整合测试中的分库分表采用了 10 库 10 表的方式，完全运行测试用例所需时间较长。

6.5 性能测试

提供各个压测工具的性能测试结果。

6.5.1 Sysbench ShardingSphere Proxy 空 Rules 性能测试

测试目的

对 ShardingSphere-Proxy 及 MySQL 进行性能对比 1. sysbench 直接压测 MySQL 性能 2. sysbench 压测 ShardingSphere-Proxy(底层透传 MySQL)

基于以上两组实验，得到使用 ShardingSphere-Proxy 对于 MySQL 的损耗。

测试环境搭建

服务器信息

1. DB 相关配置：推荐内存大于压测的数据量，使得数据均在内存热块中，其余可自行调整；
2. ShardingSphere-Proxy 相关配置：推荐使用高性能多核 CPU，其余可自行调整；
3. 压测涉及服务器均关闭 swap 分区。

数据库

```
[mysqld]
innodb_buffer_pool_size=${MORE_THAN_DATA_SIZE}
innodb_log_file_size=3000000000
innodb_log_files_in_group=5
innodb_flush_log_at_trx_commit=0
innodb_change_buffer_max_size=40
back_log=900
innodb_max_dirty_pages_pct=75
innodb_open_files=20480
innodb_buffer_pool_instances=8
innodb_page_cleaners=8
innodb_purge_threads=2
innodb_read_io_threads=8
innodb_write_io_threads=8
```

```
table_open_cache=102400
log_timestamps=system
thread_cache_size=16384
transaction_isolation=READ-COMMITTED
```

可参考进行适当调优，旨在放大底层 DB 性能，不让实验受制于 DB 性能瓶颈。

压测工具

可通过 [sysbench](#) 官网自行获取

ShardingSphere-Proxy

bin/start.sh

```
-Xmx16g -Xms16g -Xmn8g # 调整 JVM 相关参数
```

config.yaml

```
databaseName: sharding_db

dataSources:
  ds_0:
    url: jdbc:mysql://***.***.***.***:****/test?serverTimezone=UTC&useSSL=false # 参数可适当调整
    username: test
    password:
    connectionTimeoutMilliseconds: 30000
    idleTimeoutMilliseconds: 60000
    maxLifetimeMilliseconds: 1800000
    maxPoolSize: 200 # 最大链接池设为 ${压测并发数} 与压测并发数保持一致，屏蔽压测过程中额外的链接带来的影响
    minPoolSize: 200 # 最小链接池设为 ${压测并发数} 与压测并发数保持一致，屏蔽压测过程中初始化链接带来的影响

rules: []
```


测试阶段

环境准备

```
sysbench oltp_read_write --mysql-host=${DB_IP} --mysql-port=${DB_PORT} --mysql-user=${USER} --mysql-password=${PASSWD} --mysql-db=test --tables=10 --table-size=1000000 --report-interval=10 --time=100 --threads=200 cleanup
sysbench oltp_read_write --mysql-host=${DB_IP} --mysql-port=${DB_PORT} --mysql-user=${USER} --mysql-password=${PASSWD} --mysql-db=test --tables=10 --table-size=1000000 --report-interval=10 --time=100 --threads=200 prepare
```

压测命令

```
sysbench oltp_read_write --mysql-host=${DB/PROXY_IP} --mysql-port=${DB/PROXY_PORT} --mysql-user=${USER} --mysql-password=${PASSWD} --mysql-db=test --tables=10 --table-size=1000000 --report-interval=10 --time=100 --threads=200 run
```

压测报告分析

```
sysbench 1.0.20 (using bundled LuaJIT 2.1.0-beta2)
Running the test with following options:
Number of threads: 200
Report intermediate results every 10 second(s)
Initializing random number generator from current time
Initializing worker threads...
Threads started!
# 每 10 秒钟报告一次测试结果, tps、每秒读、每秒写、95% 以上的响应时长统计
[ 10s ] thds: 200 tps: 11161.70 qps: 223453.06 (r/w/o: 156451.76/44658.51/22342.80)
lat (ms,95%): 27.17 err/s: 0.00 reconn/s: 0.00
...
[ 120s ] thds: 200 tps: 11731.00 qps: 234638.36 (r/w/o: 164251.67/46924.69/23462.00)
lat (ms,95%): 24.38 err/s: 0.00 reconn/s: 0.00
SQL statistics:
  queries performed:
    read:                19560590                # 读总数
    write:               5588740                 # 写总数
    other:               27943700                # 其他操作总数 (COMMIT 等)
    total:               27943700                # 全部总数
    transactions:       1397185 (11638.59 per sec.) # 总事务数 (每秒事务数)
    queries:            27943700 (232771.76 per sec.) # 执行语句总数 (每秒执行语句次数)
    ignored errors:      0 (0.00 per sec.)          # 忽略错误数 (每秒忽略错误数)
```

reconnects:	0	(0.00 per sec.)	# 重连次数 (每秒重连次数)
General statistics:			
total time:	120.0463s		# 总共耗时
total number of events:	1397185		# 总共发生多少事务数
Latency (ms):			
min:	5.37		# 最小延时
avg:	17.13		# 平均延时
max:	109.75		# 最大延时
95th percentile:	24.83		# 超过 95% 平均耗时
sum:	23999546.19		
Threads fairness:			
events (avg/stddev):	6985.9250/34.74		# 平均每线程完成 6985.9250 次 event, 标准差为 34.74
execution time (avg/stddev):	119.9977/0.01		# 每个线程平均耗时 119.9977 秒, 标准差为 0.01

压测过程中值得关注的点

1. ShardingSphere-Proxy 所在服务器 CPU 利用率, 充分利用 CPU 为佳;
2. DB 所在服务器磁盘 IO, 物理读越低越好;
3. 压测中涉及服务器的网络 IO。

6.5.2 BenchmarkSQL ShardingSphere Proxy 分片性能测试

测试目的

使用 BenchmarkSQL 工具测试 ShardingSphere Proxy 的分片性能。

测试方法

ShardingSphere Proxy 支持通过 [BenchmarkSQL 5.0](#) 进行 TPC-C 测试。除本文说明的内容外, BenchmarkSQL 操作步骤按照原文档 `HOW-TO-RUN.txt` 即可。

测试工具微调

与单机数据库压测不同，分布式数据库解决方案难免在功能支持上有所取舍。使用 BenchmarkSQL 压测 ShardingSphere Proxy 建议进行如下调整。

移除外键与 extraHistID

修改 BenchmarkSQL 目录下 run/runDatabaseBuild.sh，文件第 17 行。

修改前：

```
AFTER_LOAD="indexCreates foreignKeys extraHistID buildFinish"
```

修改后：

```
AFTER_LOAD="indexCreates buildFinish"
```

压测环境或参数建议

注意：本节中提到的任何参数都不是绝对值，都需要根据实际测试结果进行调整或取舍。

建议使用 Java 17 运行 ShardingSphere

编译 ShardingSphere 可以使用 Java 8。

使用 Java 17 可以在默认情况下尽量提升 ShardingSphere 的性能。

ShardingSphere 数据分片建议

对 BenchmarkSQL 的数据分片，可以考虑以各个表中的 warehouse id 作为分片键。

其中一个表 bmsql_item 没有 warehouse id，数据量固定 10 万行：- 可以取 i_id 作为分片键。但可能会导致同一个 Proxy 连接同时持有多个不同数据源的连接。- 或考虑不做分片，存在单个数据源内。可能会导致某一数据源压力较大。- 或对 i_id 进行范围分片，例如 1-50000 分布在数据源 0、50001-100000 分布在数据源 1。

BenchmarkSQL 中有如下 SQL 涉及多表：

```
SELECT c_discount, c_last, c_credit, w_tax
FROM bmsql_customer
      JOIN bmsql_warehouse ON (w_id = c_w_id)
WHERE c_w_id = ? AND c_d_id = ? AND c_id = ?
```

```
SELECT o_id, o_entry_d, o_carrier_id
FROM bmsql_oorder
WHERE o_w_id = ? AND o_d_id = ? AND o_c_id = ?
      AND o_id = (
```

```
SELECT max(o_id)
FROM bmsql_oorder
WHERE o_w_id = ? AND o_d_id = ? AND o_c_id = ?
)
```

如果以 warehouse id 作为分片键，以上 SQL 涉及的表可以配置为 bindingTable：

```
rules:
- !SHARDING
  bindingTables:
    - bmsql_warehouse, bmsql_customer
    - bmsql_stock, bmsql_district, bmsql_order_line
```

以 warehouse id 为分片键的数据分片配置可以参考本文附录。

PostgreSQL JDBC URL 参数建议

对 BenchmarkSQL 所使用的配置文件中的 JDBC URL 进行调整，即参数名 conn 的值：- 增加参数 defaultRowFetchSize=50 可能减少多行结果集的 fetch 次数，需要根据实际测试结果适当增大或减小。- 增加参数 rewriteBatchedInserts=true 可能减少批量插入的耗时，例如准备数据或 New Order 业务的批量插入，需要根据实际测试结果决定是否启用。

props.pg 文件节选，建议修改的位置为第 3 行 conn 的参数值：

```
db=postgres
driver=org.postgresql.Driver
conn=jdbc:postgresql://localhost:5432/postgres?defaultRowFetchSize=50&
rewriteBatchedInserts=true
user=benchmarksql
password=PWbmsql
```

ShardingSphere Proxy server.yaml 参数建议

proxy-backend-query-fetch-size 参数值默认值为 -1，修改为 50 左右可以尽量减少多行结果集的 fetch 次数。proxy-frontend-executor-size 参数默认值为 CPU * 2，可以根据实际测试结果减少至 CPU * 0.5 左右；如果涉及 NUMA，可以根据实际测试结果设置为单个 CPU 的物理核数。

server.yaml 文件节选：

```
props:
  proxy-backend-query-fetch-size: 50
  # proxy-frontend-executor-size: 32 # 4 路 32C aarch64
  # proxy-frontend-executor-size: 12 # 2 路 12C24T x86
```

附录

BenchmarkSQL 数据分片参考配置

Pool size 请根据实际压测情况适当调整。

```
databaseName: bmsql_sharding
dataSources:
  ds_0:
    url: jdbc:postgresql://db0.ip:5432/bmsql
    username: postgres
    password: postgres
    connectionTimeoutMilliseconds: 3000
    idleTimeoutMilliseconds: 60000
    maxLifetimeMilliseconds: 1800000
    maxPoolSize: 1000
    minPoolSize: 1000
  ds_1:
    url: jdbc:postgresql://db1.ip:5432/bmsql
    username: postgres
    password: postgres
    connectionTimeoutMilliseconds: 3000
    idleTimeoutMilliseconds: 60000
    maxLifetimeMilliseconds: 1800000
    maxPoolSize: 1000
    minPoolSize: 1000
  ds_2:
    url: jdbc:postgresql://db2.ip:5432/bmsql
    username: postgres
    password: postgres
    connectionTimeoutMilliseconds: 3000
    idleTimeoutMilliseconds: 60000
    maxLifetimeMilliseconds: 1800000
    maxPoolSize: 1000
    minPoolSize: 1000
  ds_3:
    url: jdbc:postgresql://db3.ip:5432/bmsql
    username: postgres
    password: postgres
    connectionTimeoutMilliseconds: 3000
    idleTimeoutMilliseconds: 60000
    maxLifetimeMilliseconds: 1800000
    maxPoolSize: 1000
    minPoolSize: 1000
rules:
  - !SHARDING
    bindingTables:
      - bmsql_warehouse, bmsql_customer
```

```

- bmsql_stock, bmsql_district, bmsql_order_line
defaultDatabaseStrategy:
  none:
defaultTableStrategy:
  none:
keyGenerators:
  snowflake:
    type: SNOWFLAKE
tables:
  bmsql_config:
    actualDataNodes: ds_0.bmsql_config

  bmsql_warehouse:
    actualDataNodes: ds_${0..3}.bmsql_warehouse
    databaseStrategy:
      standard:
        shardingColumn: w_id
        shardingAlgorithmName: mod_4

  bmsql_district:
    actualDataNodes: ds_${0..3}.bmsql_district
    databaseStrategy:
      standard:
        shardingColumn: d_w_id
        shardingAlgorithmName: mod_4

  bmsql_customer:
    actualDataNodes: ds_${0..3}.bmsql_customer
    databaseStrategy:
      standard:
        shardingColumn: c_w_id
        shardingAlgorithmName: mod_4

  bmsql_item:
    actualDataNodes: ds_${0..3}.bmsql_item
    databaseStrategy:
      standard:
        shardingColumn: i_id
        shardingAlgorithmName: mod_4

  bmsql_history:
    actualDataNodes: ds_${0..3}.bmsql_history
    databaseStrategy:
      standard:
        shardingColumn: h_w_id
        shardingAlgorithmName: mod_4

  bmsql_oorder:

```

```

actualDataNodes: ds_${0..3}.bmsql_oorder
databaseStrategy:
  standard:
    shardingColumn: o_w_id
    shardingAlgorithmName: mod_4

bmsql_stock:
actualDataNodes: ds_${0..3}.bmsql_stock
databaseStrategy:
  standard:
    shardingColumn: s_w_id
    shardingAlgorithmName: mod_4

bmsql_new_order:
actualDataNodes: ds_${0..3}.bmsql_new_order
databaseStrategy:
  standard:
    shardingColumn: no_w_id
    shardingAlgorithmName: mod_4

bmsql_order_line:
actualDataNodes: ds_${0..3}.bmsql_order_line
databaseStrategy:
  standard:
    shardingColumn: ol_w_id
    shardingAlgorithmName: mod_4

shardingAlgorithms:
  mod_4:
    type: MOD
    props:
      sharding-count: 4

```

BenchmarkSQL 5.0 PostgreSQL 语句列表

Create tables

```

create table bmsql_config (
  cfg_name    varchar(30) primary key,
  cfg_value   varchar(50)
);

create table bmsql_warehouse (
  w_id        integer not null,
  w_ytd       decimal(12,2),
  w_tax       decimal(4,4),

```

```
w_name      varchar(10),
w_street_1  varchar(20),
w_street_2  varchar(20),
w_city      varchar(20),
w_state     char(2),
w_zip       char(9)
);

create table bmsql_district (
  d_w_id      integer      not null,
  d_id        integer      not null,
  d_ytd       decimal(12,2),
  d_tax       decimal(4,4),
  d_next_o_id integer,
  d_name      varchar(10),
  d_street_1  varchar(20),
  d_street_2  varchar(20),
  d_city      varchar(20),
  d_state     char(2),
  d_zip       char(9)
);

create table bmsql_customer (
  c_w_id      integer      not null,
  c_d_id      integer      not null,
  c_id        integer      not null,
  c_discount  decimal(4,4),
  c_credit    char(2),
  c_last      varchar(16),
  c_first     varchar(16),
  c_credit_lim decimal(12,2),
  c_balance   decimal(12,2),
  c_ytd_payment decimal(12,2),
  c_payment_cnt integer,
  c_delivery_cnt integer,
  c_street_1  varchar(20),
  c_street_2  varchar(20),
  c_city      varchar(20),
  c_state     char(2),
  c_zip       char(9),
  c_phone     char(16),
  c_since     timestamp,
  c_middle    char(2),
  c_data      varchar(500)
);

create sequence bmsql_hist_id_seq;
```



```
create table bmsql_history (
  hist_id integer,
  h_c_id integer,
  h_c_d_id integer,
  h_c_w_id integer,
  h_d_id integer,
  h_w_id integer,
  h_date timestamp,
  h_amount decimal(6,2),
  h_data varchar(24)
);

create table bmsql_new_order (
  no_w_id integer not null,
  no_d_id integer not null,
  no_o_id integer not null
);

create table bmsql_oorder (
  o_w_id integer not null,
  o_d_id integer not null,
  o_id integer not null,
  o_c_id integer,
  o_carrier_id integer,
  o_ol_cnt integer,
  o_all_local integer,
  o_entry_d timestamp
);

create table bmsql_order_line (
  ol_w_id integer not null,
  ol_d_id integer not null,
  ol_o_id integer not null,
  ol_number integer not null,
  ol_i_id integer not null,
  ol_delivery_d timestamp,
  ol_amount decimal(6,2),
  ol_supply_w_id integer,
  ol_quantity integer,
  ol_dist_info char(24)
);

create table bmsql_item (
  i_id integer not null,
  i_name varchar(24),
  i_price decimal(5,2),
  i_data varchar(50),
  i_im_id integer
```

```
);

create table bmsql_stock (
  s_w_id      integer      not null,
  s_i_id      integer      not null,
  s_quantity  integer,
  s_ytd       integer,
  s_order_cnt integer,
  s_remote_cnt integer,
  s_data      varchar(50),
  s_dist_01   char(24),
  s_dist_02   char(24),
  s_dist_03   char(24),
  s_dist_04   char(24),
  s_dist_05   char(24),
  s_dist_06   char(24),
  s_dist_07   char(24),
  s_dist_08   char(24),
  s_dist_09   char(24),
  s_dist_10   char(24)
);
```

Create indexes

```
alter table bmsql_warehouse add constraint bmsql_warehouse_pkey
  primary key (w_id);

alter table bmsql_district add constraint bmsql_district_pkey
  primary key (d_w_id, d_id);

alter table bmsql_customer add constraint bmsql_customer_pkey
  primary key (c_w_id, c_d_id, c_id);

create index bmsql_customer_idx1
  on bmsql_customer (c_w_id, c_d_id, c_last, c_first);

alter table bmsql_oorder add constraint bmsql_oorder_pkey
  primary key (o_w_id, o_d_id, o_id);

create unique index bmsql_oorder_idx1
  on bmsql_oorder (o_w_id, o_d_id, o_carrier_id, o_id);

alter table bmsql_new_order add constraint bmsql_new_order_pkey
  primary key (no_w_id, no_d_id, no_o_id);

alter table bmsql_order_line add constraint bmsql_order_line_pkey
```

```

primary key (ol_w_id, ol_d_id, ol_o_id, ol_number);

alter table bmsql_stock add constraint bmsql_stock_pkey
primary key (s_w_id, s_i_id);

alter table bmsql_item add constraint bmsql_item_pkey
primary key (i_id);

```

New Order 业务

stmtNewOrderSelectWhseCust

```

UPDATE bmsql_district
  SET d_next_o_id = d_next_o_id + 1
  WHERE d_w_id = ? AND d_id = ?

```

stmtNewOrderSelectDist

```

SELECT d_tax, d_next_o_id
  FROM bmsql_district
  WHERE d_w_id = ? AND d_id = ?
  FOR UPDATE

```

stmtNewOrderUpdateDist

```

UPDATE bmsql_district
  SET d_next_o_id = d_next_o_id + 1
  WHERE d_w_id = ? AND d_id = ?

```

stmtNewOrderInsertOrder

```

INSERT INTO bmsql_oorder (
  o_id, o_d_id, o_w_id, o_c_id, o_entry_d,
  o_ol_cnt, o_all_local)
VALUES (?, ?, ?, ?, ?, ?, ?, ?)

```

stmtNewOrderInsertNewOrder

```

INSERT INTO bmsql_new_order (
  no_o_id, no_d_id, no_w_id)
VALUES (?, ?, ?)

```

stmtNewOrderSelectStock

```

SELECT s_quantity, s_data,
  s_dist_01, s_dist_02, s_dist_03, s_dist_04,
  s_dist_05, s_dist_06, s_dist_07, s_dist_08,
  s_dist_09, s_dist_10

```

```
FROM bmsql_stock
WHERE s_w_id = ? AND s_i_id = ?
FOR UPDATE
```

stmtNewOrderSelectItem

```
SELECT i_price, i_name, i_data
FROM bmsql_item
WHERE i_id = ?
```

stmtNewOrderUpdateStock

```
UPDATE bmsql_stock
SET s_quantity = ?, s_ytd = s_ytd + ?,
    s_order_cnt = s_order_cnt + 1,
    s_remote_cnt = s_remote_cnt + ?
WHERE s_w_id = ? AND s_i_id = ?
```

stmtNewOrderInsertOrderLine

```
INSERT INTO bmsql_order_line (
    ol_o_id, ol_d_id, ol_w_id, ol_number,
    ol_i_id, ol_supply_w_id, ol_quantity,
    ol_amount, ol_dist_info)
VALUES (?, ?, ?, ?, ?, ?, ?, ?, ?)
```

Payment 业务

stmtPaymentSelectWarehouse

```
SELECT w_name, w_street_1, w_street_2, w_city,
       w_state, w_zip
FROM bmsql_warehouse
WHERE w_id = ?
```

stmtPaymentSelectDistrict

```
SELECT d_name, d_street_1, d_street_2, d_city,
       d_state, d_zip
FROM bmsql_district
WHERE d_w_id = ? AND d_id = ?
```

stmtPaymentSelectCustomerListByLast

```
SELECT c_id
FROM bmsql_customer
WHERE c_w_id = ? AND c_d_id = ? AND c_last = ?
ORDER BY c_first
```

stmtPaymentSelectCustomer

```
SELECT c_first, c_middle, c_last, c_street_1, c_street_2,
       c_city, c_state, c_zip, c_phone, c_since, c_credit,
       c_credit_lim, c_discount, c_balance
FROM   bmsql_customer
WHERE  c_w_id = ? AND c_d_id = ? AND c_id = ?
FOR UPDATE
```

stmtPaymentSelectCustomerData

```
SELECT c_data
FROM   bmsql_customer
WHERE  c_w_id = ? AND c_d_id = ? AND c_id = ?
```

stmtPaymentUpdateWarehouse

```
UPDATE bmsql_warehouse
SET    w_ytd = w_ytd + ?
WHERE  w_id = ?
```

stmtPaymentUpdateDistrict

```
UPDATE bmsql_district
SET    d_ytd = d_ytd + ?
WHERE  d_w_id = ? AND d_id = ?
```

stmtPaymentUpdateCustomer

```
UPDATE bmsql_customer
SET    c_balance = c_balance - ?,
       c_ytd_payment = c_ytd_payment + ?,
       c_payment_cnt = c_payment_cnt + 1
WHERE  c_w_id = ? AND c_d_id = ? AND c_id = ?
```

stmtPaymentUpdateCustomerWithData

```
UPDATE bmsql_customer
SET    c_balance = c_balance - ?,
       c_ytd_payment = c_ytd_payment + ?,
       c_payment_cnt = c_payment_cnt + 1,
       c_data = ?
WHERE  c_w_id = ? AND c_d_id = ? AND c_id = ?
```

stmtPaymentInsertHistory

```
INSERT INTO bmsql_history (
    h_c_id, h_c_d_id, h_c_w_id, h_d_id, h_w_id,
    h_date, h_amount, h_data)
VALUES (?, ?, ?, ?, ?, ?, ?, ?)
```

Order Status 业务

stmtOrderStatusSelectCustomerListByLast

```
SELECT c_id
  FROM bmsql_customer
 WHERE c_w_id = ? AND c_d_id = ? AND c_last = ?
 ORDER BY c_first
```

stmtOrderStatusSelectCustomer

```
SELECT c_first, c_middle, c_last, c_balance
  FROM bmsql_customer
 WHERE c_w_id = ? AND c_d_id = ? AND c_id = ?
```

stmtOrderStatusSelectLastOrder

```
SELECT o_id, o_entry_d, o_carrier_id
  FROM bmsql_oorder
 WHERE o_w_id = ? AND o_d_id = ? AND o_c_id = ?
    AND o_id = (
      SELECT max(o_id)
        FROM bmsql_oorder
       WHERE o_w_id = ? AND o_d_id = ? AND o_c_id = ?
    )
```

stmtOrderStatusSelectOrderLine

```
SELECT ol_i_id, ol_supply_w_id, ol_quantity,
       ol_amount, ol_delivery_d
  FROM bmsql_order_line
 WHERE ol_w_id = ? AND ol_d_id = ? AND ol_o_id = ?
 ORDER BY ol_w_id, ol_d_id, ol_o_id, ol_number
```

Stock level 业务

stmtStockLevelSelectLow

```
SELECT count(*) AS low_stock FROM (
  SELECT s_w_id, s_i_id, s_quantity
    FROM bmsql_stock
   WHERE s_w_id = ? AND s_quantity < ? AND s_i_id IN (
      SELECT ol_i_id
        FROM bmsql_district
       JOIN bmsql_order_line ON ol_w_id = d_w_id
          AND ol_d_id = d_id
          AND ol_o_id >= d_next_o_id - 20
          AND ol_o_id < d_next_o_id
    )
)
```

```

        WHERE d_w_id = ? AND d_id = ?
    )
) AS L

```

Delivery BG 业务

stmtDeliveryBGSelectOldestNewOrder

```

SELECT no_o_id
FROM bmsql_new_order
WHERE no_w_id = ? AND no_d_id = ?
ORDER BY no_o_id ASC

```

stmtDeliveryBGDeleteOldestNewOrder

```

DELETE FROM bmsql_new_order
WHERE no_w_id = ? AND no_d_id = ? AND no_o_id = ?

```

stmtDeliveryBGSelectOrder

```

SELECT o_c_id
FROM bmsql_oorder
WHERE o_w_id = ? AND o_d_id = ? AND o_id = ?

```

stmtDeliveryBGUpdateOrder

```

UPDATE bmsql_oorder
SET o_carrier_id = ?
WHERE o_w_id = ? AND o_d_id = ? AND o_id = ?

```

stmtDeliveryBGSelectSumOLAmount

```

SELECT sum(ol_amount) AS sum_ol_amount
FROM bmsql_order_line
WHERE ol_w_id = ? AND ol_d_id = ? AND ol_o_id = ?

```

stmtDeliveryBGUpdateOrderLine

```

UPDATE bmsql_order_line
SET ol_delivery_d = ?
WHERE ol_w_id = ? AND ol_d_id = ? AND ol_o_id = ?

```

stmtDeliveryBGUpdateCustomer

```

UPDATE bmsql_customer
SET c_balance = c_balance + ?,
    c_delivery_cnt = c_delivery_cnt + 1
WHERE c_w_id = ? AND c_d_id = ? AND c_id = ?

```

6.6 模块测试

提供复杂模块的测试引擎。

6.6.1 SQL 解析测试

数据准备

SQL 解析无需真实的测试环境，开发者只需定义好待测试的 SQL，以及解析后的断言数据即可：

SQL 数据

在集成测试的部分提到过 `sql-case-id`，其对应的 SQL，可以在不同模块共享。开发者只需要在 `shardingsphere-sql-parser/shardingsphere-sql-parser-test/src/main/resources/sql/supported/${SQL-TYPE}/*.xml` 添加待测试的 SQL 即可。

断言数据

断言的解析数据保存在 `shardingsphere-sql-parser/shardingsphere-sql-parser-test/src/main/resources/case/${SQL-TYPE}/*.xml` 在 xml 文件中，可以针对表名，token，SQL 条件等进行断言，例如如下的配置：

```
<parser-result-sets>
  <parser-result sql-case-id="insert_with_multiple_values">
    <tables>
      <table name="t_order" />
    </tables>
    <tokens>
      <table-token start-index="12" table-name="t_order" length="7" />
    </tokens>
    <sharding-conditions>
      <and-condition>
        <condition column-name="order_id" table-name="t_order" operator=
"EQUAL">
          <value literal="1" type="int" />
        </condition>
        <condition column-name="user_id" table-name="t_order" operator=
"EQUAL">
          <value literal="1" type="int" />
        </condition>
      </and-condition>
      <and-condition>
        <condition column-name="order_id" table-name="t_order" operator=
"EQUAL">
          <value literal="2" type="int" />
        </condition>
      </and-condition>
    </sharding-conditions>
  </parser-result>
</parser-result-sets>
```



```

        </condition>
        <condition column-name="user_id" table-name="t_order" operator=
"EQUAL">
            <value literal="2" type="int" />
        </condition>
    </and-condition>
</sharding-conditions>
</parser-result>
</parser-result-sets>

```

设置好上面两类数据，开发者就可以通过 `shardingsphere-sql-parser/shardingsphere-sql-parser-test` 下对应的测试引擎启动 SQL 解析的测试了。

6.6.2 SQL 改写测试

目标

面向逻辑库与逻辑表书写的 SQL，并不能够直接在真实的数据库中执行，SQL 改写用于将逻辑 SQL 改写为在真实数据库中可以正确执行的 SQL。它包括正确性改写和优化改写两部分，所以 SQL 改写的测试都是基于这些改写方向进行校验的。

测试

SQL 改写测试用例位于 `sharding-core/sharding-core-rewrite` 下的 `test` 中。SQL 改写的测试主要依赖如下几个部分：

- 测试引擎
- 环境配置
- 验证数据

测试引擎是 SQL 改写测试的入口，跟其他引擎一样，通过 Junit 的 `Parameterized` 逐条读取 `test/resources` 目录中测试类型下对应的 `xml` 文件，然后按读取顺序一一进行验证。

环境配置存放在 `test/resources/yaml` 路径中测试类型下对应的 `yaml` 中。配置了 `dataSources`, `shardingRule`, `encryptRule` 等信息，例子如下：

```

dataSources:
  db: !!com.zaxxer.hikari.HikariDataSource
    driverClassName: org.h2.Driver
    jdbcUrl: jdbc:h2:mem:db;DB_CLOSE_DELAY=-1;DATABASE_TO_UPPER=false;MODE=MYSQL
    username: sa
    password:

## sharding 规则
rules:
- !SHARDING

```

```

tables:
  t_account:
    actualDataNodes: db.t_account_${0..1}
    tableStrategy:
      standard:
        shardingColumn: account_id
        shardingAlgorithmName: account_table_inline
    keyGenerateStrategy:
      column: account_id
      keyGeneratorName: snowflake
  t_account_detail:
    actualDataNodes: db.t_account_detail_${0..1}
    tableStrategy:
      standard:
        shardingColumn: order_id
        shardingAlgorithmName: account_detail_table_inline
bindingTables:
  - t_account, t_account_detail
shardingAlgorithms:
  account_table_inline:
    type: INLINE
    props:
      algorithm-expression: t_account_${account_id % 2}
  account_detail_table_inline:
    type: INLINE
    props:
      algorithm-expression: t_account_detail_${account_id % 2}
keyGenerators:
  snowflake:
    type: SNOWFLAKE

```

验证数据存放在 test\resources 路径中测试类型下对应的 xml 文件中。验证数据中, yaml-rule 指定了环境以及 rule 的配置文件, input 指定了待测试的 SQL 以及参数, output 指定了期待的 SQL 以及参数。其中 db-type 决定了 SQL 解析的类型, 默认为 SQL92, 例如:

```

<rewrite-assertions yaml-rule="yaml/sharding/sharding-rule.yaml">
  <!-- 替换数据库类型需要在这里更改 db-type -->
  <rewrite-assertion id="create_index_for_mysql" db-type="MySQL">
    <input sql="CREATE INDEX index_name ON t_account ('status')" />
    <output sql="CREATE INDEX index_name ON t_account_0 ('status')" />
    <output sql="CREATE INDEX index_name ON t_account_1 ('status')" />
  </rewrite-assertion>
</rewrite-assertions>

```

只需在 xml 文件中编写测试数据, 配置好相应的 yaml 配置文件, 就可以在不更改任何 Java 代码的情况下校验对应的 SQL 了。

6.7 Scaling 集成测试

6.7.1 测试目的

验证数据迁移以及依赖模块功能的正确性。

6.7.2 测试环境

目前支持 Native 和 Docker 两种环境。

1. Native 环境直接运行在开发者提供的测试环境中，需要用户自己启动 ShardingSphere-Proxy 和对应的数据库实例，适于调试场景；
2. Docker 环境由 Maven 运行，适用于云编译环境和测试 ShardingSphere-Proxy 的场景，如：GitHub Action。

目前支持的数据库类型：MySQL、PostgreSQL、openGauss。

6.7.3 使用指南

模 块 路 径 shardingsphere-test/shardingsphere-integration-test/
shardingsphere-integration-test-scaling。

环境配置

`${DOCKER-IMAGE}` 表示 docker 镜像名称，如 `mysql:8`。`${DATABASE-TYPE}` 表示数据库类型。目录：`src/test/resources/env-it-env.properties`: 集成测试启动参数。`-${DATABASE-TYPE}/server.yaml`: 数据库对应的 ShardingSphere-Proxy 配置文件。`-${DATABASE-TYPE}/initdb.sql`: 数据库初始化 SQL。`-${DATABASE-TYPE}/*.cnf,*.conf`: 以 `cnf` 或者 `conf` 结尾的文件，是数据库的配置文件，用于 Docker 挂载。`- common/command.xml`: 测试中用到的 DistSQL。`- scenario/`: 存放测试场景中的 SQL。

测试用例

目前所有的测试用例，都直接继承自 `BaseExtraSQLITCase`，间接继承了 `BaseITCase`。`-BaseITCase`: 提供了通用方法给子类。`-BaseExtraSQLITCase`: 提供了建表、CRUD 语句执行方法

用例示例: `MySQLGeneralScalingIT`。覆盖的功能点如下: `- 库级别迁移（所有表）` - `表级别迁移（任意多个表）` - `迁移数据一致性校验` - `数据迁移过程中支持停写` - `数据迁移过程中支持重启` - `数据迁移支持整型主键` - `数据迁移支持字符串主键` - `使用非管理员账号进行数据迁移`

运行测试用例

`it-env.properties` 所有属性值都可以通过 Maven 命令行 `-D` 的方式传入，优先级高于配置文件。

Native 环境启动

使用者在本地提前启动 ShardingSphere-Proxy 以及依赖的配置中心（如 ZooKeeper）和数据库。要求 ShardingSphere-Proxy 的端口是 3307。以 MySQL 为例，`it-env.properties` 可以配置如下：

```
scaling.it.env.type=NATIVE
scaling.it.native.database=mysql
scaling.it.native.mysql.username=root
scaling.it.native.mysql.password=root
scaling.it.native.mysql.port=3306
```

找到对应的用例，在 IDE 下使用 Junit 的方式启动即可。

Docker 环境启动

第一步：打包镜像

```
./mvnw -B clean install -am -pl shardingsphere-test/shardingsphere-integration-test/shardingsphere-integration-test-scaling -Dit.env.docker -DskipTests
```

运行以上命令会构建出一个用于集成测试的 Docker 镜像 `apache/shardingsphere-proxy-test:latest`，该镜像设置了远程调试的端口，默认是 3308。如果仅修改了测试代码，可以复用已有的测试镜像，无须重新构建。

Docker 模式下，如果需要对 Docker 镜像启动参数进行调整，可以对修改 `ShardingSphereProxyDocker-Container` 文件中的相关配置。

ShardingSphere-Proxy 输出的日志带有 `:Scaling-Proxy` 前缀。

使用 Maven 的方式运行用例。以 MySQL 为例：

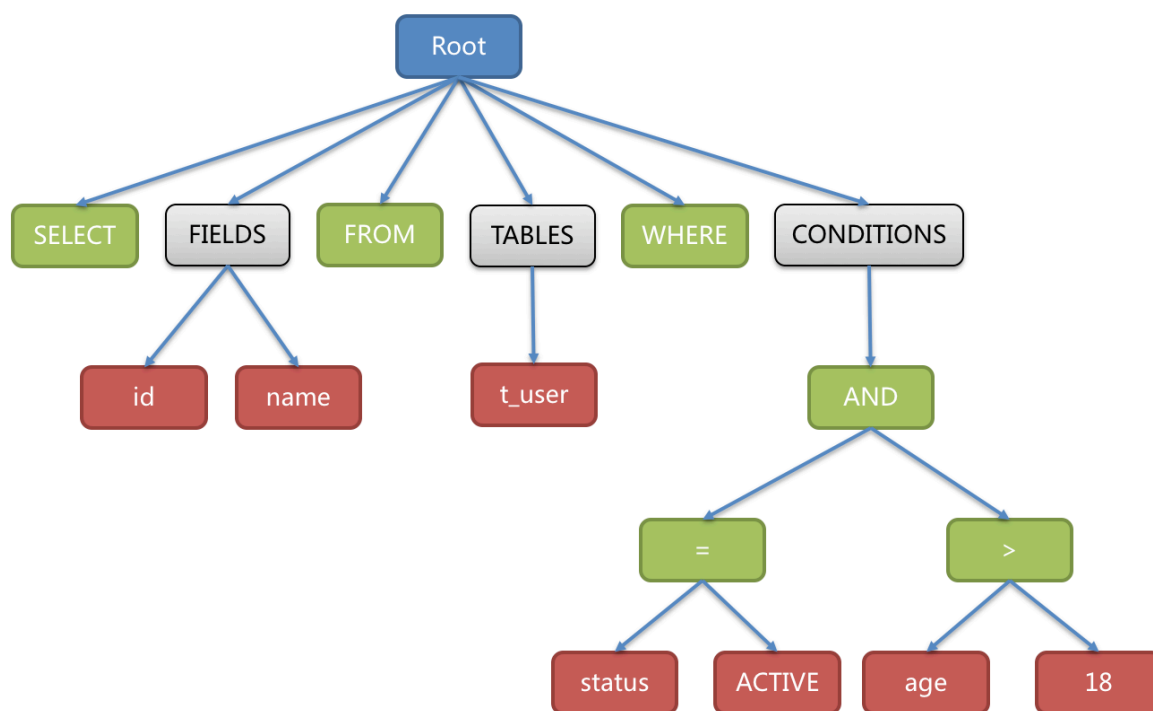
```
./mvnw -nsu -B install -f shardingsphere-test/shardingsphere-integration-test/shardingsphere-integration-test-scaling/pom.xml -Dscaling.it.env.type=DOCKER -Dscaling.it.docker.mysql.version=${image-name}
```

也可以使用 IDE 的方式运行用例。`it-env.properties` 可以配置如下：

```
scaling.it.env.type=DOCKER
scaling.it.docker.mysql.version=mysql:5.7
```

本章包含了 Apache ShardingSphere 的技术实现细节，供开发者和用户参考。

7.1 数据兼容性



• SQL 兼容

SQL 是使用者与数据库交流的标准语言。SQL 解析引擎负责将 SQL 字符串解析为抽象语法树，供 Apache ShardingSphere 理解并实现其增量功能。

ShardingSphere 目前支持 MySQL, PostgreSQL, SQLServer, Oracle, openGauss 以及符合 SQL92 规范的 SQL 方言。由于 SQL 语法的复杂性，目前仍然存在少量不支持的 SQL。

• 数据库协议兼容

Apache ShardingSphere 目前根据不同的数据协议，实现了 MySQL 和 PostgreSQL 协议。

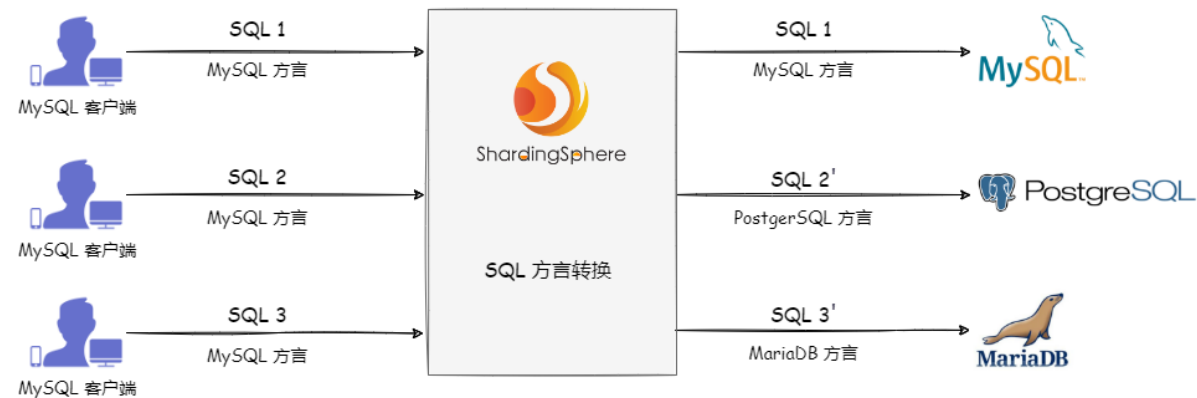
- 特性支持

Apache ShardingSphere 为数据库提供了分布式协作的能力，同时将一部分数据库特性抽象到了上层，进行统一管理，以降低用户的使用难度。

因此，对于统一提供的特性，原生的 SQL 将不再下发到数据库，并提示该操作不被支持，用户可使用 ShardingSphere 提供的的方式进行代替。

7.2 数据库网关

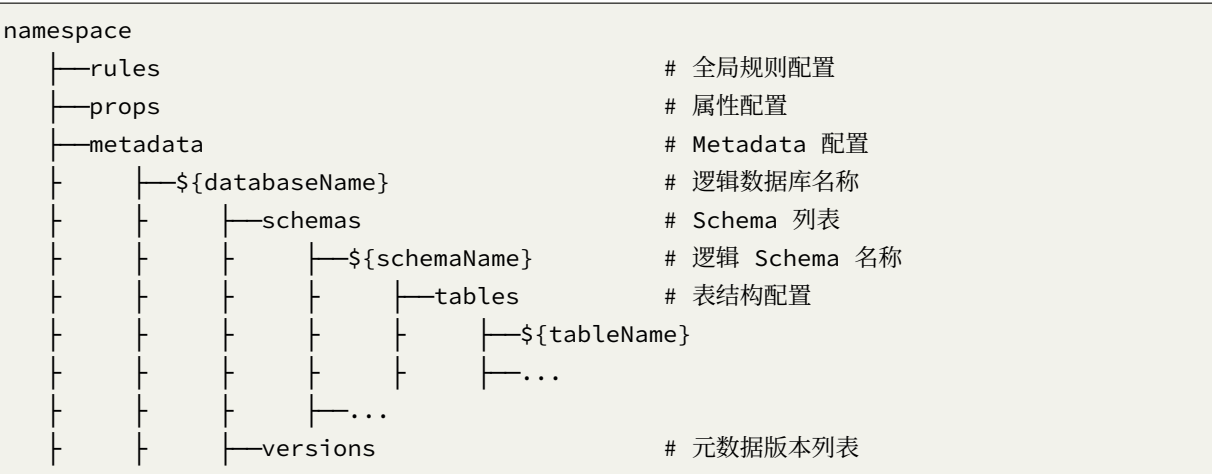
Apache ShardingSphere 提供了 SQL 方言翻译的能力，能否实现数据库方言之间的自动转换。例如，用户可以使用 MySQL 客户端连接 ShardingSphere 并发送基于 MySQL 方言的 SQL，ShardingSphere 能自动识别用户协议与存储节点类型，自动完成 SQL 方言转换，访问 PostgreSQL 等异构存储节点。



7.3 管控

7.3.1 注册中心数据结构

在定义的命名空间下，rules、props 和 metadata 节点以 YAML 格式存储配置，可通过修改节点来实现对于配置的动态管理。nodes 存储数据库访问对象运行节点，用于区分不同数据库访问实例。





/rules

全局规则配置，可包括访问 ShardingSphere-Proxy 用户名和密码的权限配置。

```

- !AUTHORITY
users:
  - root@%:root
  - sharding@127.0.0.1:sharding
provider:
  type: ALL_PERMITTED
    
```

/props

属性配置，详情请参见[配置手册](#)。

```
kernel-executor-size: 20
sql-show: true
```

/metadata/databaseName/versions/{versionNumber}/dataSources

多个数据库连接池的集合，不同数据库连接池属性自适应（例如：DBCP，C3P0，Druid，HikariCP）。

```
ds_0:
  initializationFailTimeout: 1
  validationTimeout: 5000
  maxLifetime: 1800000
  leakDetectionThreshold: 0
  minimumIdle: 1
  password: root
  idleTimeout: 60000
  jdbcUrl: jdbc:mysql://127.0.0.1:3306/ds_0?serverTimezone=UTC&useSSL=false
  dataSourceClassName: com.zaxxer.hikari.HikariDataSource
  maximumPoolSize: 50
  connectionTimeout: 30000
  username: root
  poolName: HikariPool-1
ds_1:
  initializationFailTimeout: 1
  validationTimeout: 5000
  maxLifetime: 1800000
  leakDetectionThreshold: 0
  minimumIdle: 1
  password: root
  idleTimeout: 60000
  jdbcUrl: jdbc:mysql://127.0.0.1:3306/ds_1?serverTimezone=UTC&useSSL=false
  dataSourceClassName: com.zaxxer.hikari.HikariDataSource
  maximumPoolSize: 50
  connectionTimeout: 30000
  username: root
  poolName: HikariPool-2
```


`/metadata/databaseName/versions/{versionNumber}/rules`

规则配置，可包括数据分片、读写分离、数据加密、影子库压测等配置。

```
- !SHARDING
  xxx

- !READWRITE_SPLITTING
  xxx

- !ENCRYPT
  xxx
```

`/metadata/databaseName/schemas/{schemaName}/tables`

表结构配置，每个表使用单独节点存储，暂不支持动态修改。

```
name: t_order                                # 表名
columns:                                     # 列
  id:                                         # 列名
    caseSensitive: false
    dataType: 0
    generated: false
    name: id
    primaryKey: true
  order_id:
    caseSensitive: false
    dataType: 0
    generated: false
    name: order_id
    primaryKey: false
indexes:                                     # 索引
  t_user_order_id_index:                     # 索引名
    name: t_user_order_id_index
```

`/nodes/compute_nodes`

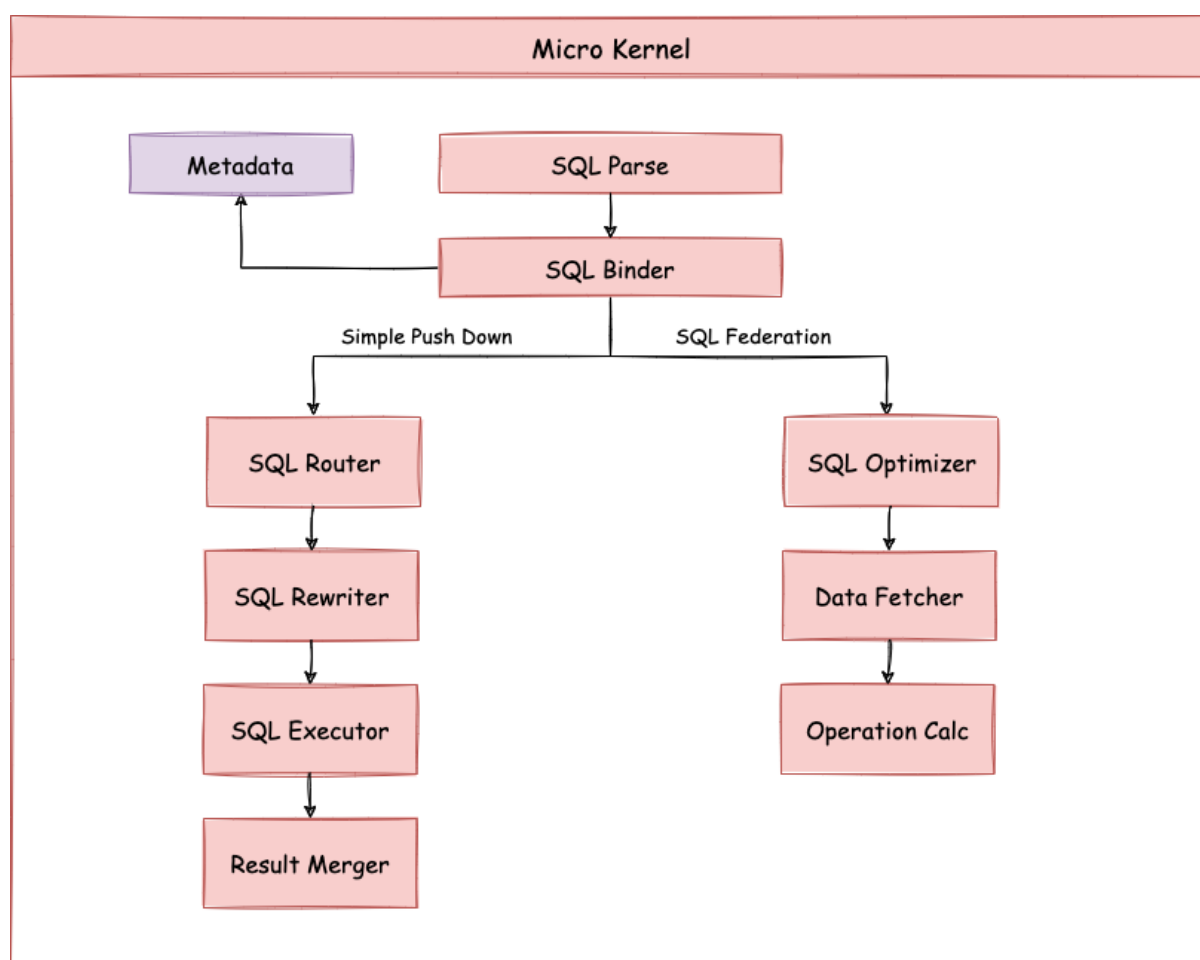
数据库访问对象运行实例信息，子节点是当前运行实例的标识。运行实例标识使用 UUID 生成，每次启动重新生成。运行实例标识均为临时节点，当实例上线时注册，下线时自动清理。注册中心监控这些节点的变化来治理运行中实例对数据库的访问等。

/nodes/storage_nodes

可以治理读写分离从库，可动态添加删除以及禁用。

7.4 数据分片

ShardingSphere 数据分片的原理如下图所示，按照是否需要进行查询优化，可以分为 Simple Push Down 下推流程和 SQL Federation 执行引擎流程。Simple Push Down 下推流程由 SQL 解析 => SQL 绑定 => SQL 路由 => SQL 改写 => SQL 执行 => 结果归并组成，主要用于处理标准分片场景下的 SQL 执行。SQL Federation 执行引擎流程由 SQL 解析 => SQL 绑定 => 逻辑优化 => 物理优化 => 数据拉取 => 算子执行组成，SQL Federation 执行引擎内部进行逻辑优化和物理优化，在优化执行阶段依赖 Standard 内核流程，对优化后的逻辑 SQL 进行路由、改写、执行和归并。



7.4.1 SQL 解析

分为词法解析和语法解析。先通过词法解析器将 SQL 拆分为一个个不可再分的单词。再使用语法解析器对 SQL 进行理解，并最终提炼出解析上下文。解析上下文包括表、选择项、排序项、分组项、聚合函数、分页信息、查询条件以及可能需要修改的占位符的标记。

7.4.2 SQL 路由

根据解析上下文匹配用户配置的分片策略，并生成路由路径。目前支持分片路由和广播路由。

7.4.3 SQL 改写

将 SQL 改写为在真实数据库中 can 正确执行的语句。SQL 改写分为正确性改写和优化改写。

7.4.4 SQL 执行

通过多线程执行器异步执行。

7.4.5 结果归并

将多个执行结果集归并以便于通过统一的 JDBC 接口输出。结果归并包括流式归并、内存归并和使用装饰者模式的追加归并这几种方式。

7.4.6 查询优化

由 Federation 执行引擎（开发中）提供支持，对关联查询、子查询等复杂查询进行优化，同时支持跨多个数据库实例的分布式查询，内部使用关系代数优化查询计划，通过最优计划查询出结果。

7.4.7 解析引擎

相对于其他编程语言，SQL 是比较简单的。不过，它依然是一门完善的编程语言，因此对 SQL 的语法进行解析，与解析其他编程语言（如：Java 语言、C 语言、Go 语言等）并无本质区别。

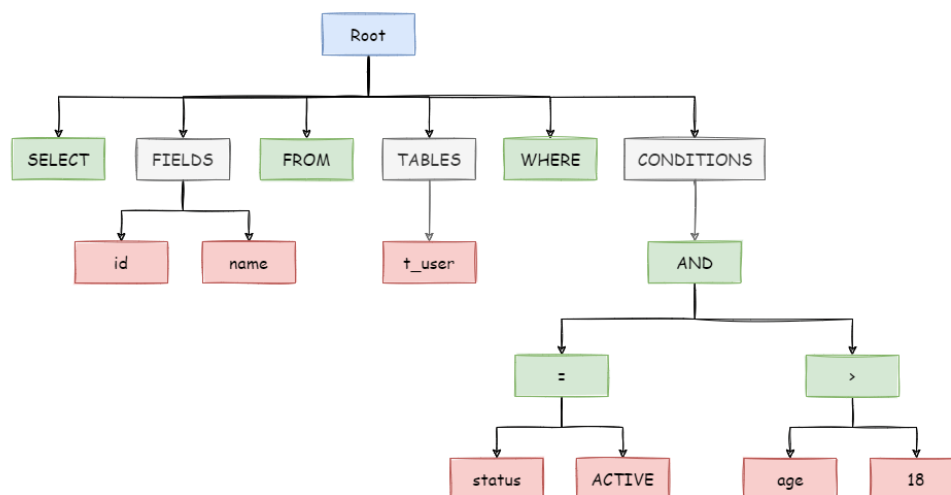
抽象语法树

解析过程分为词法解析和语法解析。词法解析器用于将 SQL 拆解为不可再分的原子符号，称为 Token。并根据不同数据库方言所提供的字典，将其归类为关键字，表达式，字面量和操作符。再使用语法解析器将词法解析器的输出转换为抽象语法树。

例如，以下 SQL：

```
SELECT id, name FROM t_user WHERE status = 'ACTIVE' AND age > 18
```

解析之后的为抽象语法树见下图。



为了便于理解，抽象语法树中的关键字的 Token 用绿色表示，变量的 Token 用红色表示，灰色表示需要进一步拆分。

最后，通过 visitor 对抽象语法树遍历构造域模型，通过域模型（SQLStatement）去提炼分片所需的上下文，并标记有可能需要改写的位置。供分片使用的解析上下文包含查询选择项（Select Items）、表信息（Table）、分片条件（Sharding Condition）、自增主键信息（Auto increment Primary Key）、排序信息（Order By）、分组信息（Group By）以及分页信息（Limit、Rownum、Top）。SQL 的一次解析过程是不可逆的，一个个 Token 按 SQL 原本的顺序依次进行解析，性能很高。考虑到各种数据库 SQL 方言的异同，在解析模块提供了各类数据库的 SQL 方言字典。

SQL 解析引擎

历史

SQL 解析作为分库分表类产品的核心，其性能和兼容性是最重要的衡量指标。ShardingSphere 的 SQL 解析器经历了 3 代产品的更新迭代。

第一代 SQL 解析器为了追求性能与快速实现，在 1.4.x 之前的版本使用 Druid 作为 SQL 解析器。经实际测试，它的性能远超其它解析器。

第二代 SQL 解析器从 1.5.x 版本开始，ShardingSphere 采用完全自研的 SQL 解析引擎。由于目的不同，ShardingSphere 并不需要将 SQL 转为一颗完全的抽象语法树，也无需通过访问器模式进行二次遍历。它采用对 SQL 半理解的方式，仅提炼数据分片需要关注的上下文，因此 SQL 解析的性能和兼容性得到了进一步的提高。

第三代 SQL 解析器从 3.0.x 版本开始，尝试使用 ANTLR 作为 SQL 解析引擎的生成器，并采用 Visit 的方式从 AST 中获取 SQL Statement。从 5.0.x 版本开始，解析引擎的架构已完成重构调整，同时通过将第一次解析得到的 AST 放入缓存，方便下次直接获取相同 SQL 的解析结果，来提高解析效率。因此我们建议用户采用 PreparedStatement 这种 SQL 预编译的方式来提升性能。

功能点

- 提供独立的 SQL 解析功能
- 可以非常方便的对语法规则进行扩充和修改（使用了 ANTLR）
- 支持多种方言的 SQL 解析

数据库	支持状态
MySQL	支持，完善
PostgreSQL	支持，完善
SQLServer	支持
Oracle	支持
SQL92	支持
openGauss	支持

API 使用

- 引入 Maven 依赖

```
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-sql-parser-engine</artifactId>
  <version>${project.version}</version>
</dependency>
<!-- 根据需要引入指定方言的解析模块（以 MySQL 为例），可以添加所有支持的方言，也可以只添加使用到的 -->
<dependency>
  <groupId>org.apache.shardingsphere</groupId>
  <artifactId>shardingsphere-sql-parser-mysql</artifactId>
  <version>${project.version}</version>
</dependency>
```

- 获取语法树

```
CacheOption cacheOption = new CacheOption(128, 1024L);
SQLParserEngine parserEngine = new SQLParserEngine(sql, cacheOption);
ParseASTNode parseASTNode = parserEngine.parse(sql, useCache);
```

- 获取 SQLStatement

```
CacheOption cacheOption = new CacheOption(128, 1024L);
SQLParserEngine parserEngine = new SQLParserEngine(sql, cacheOption);
ParseASTNode parseASTNode = parserEngine.parse(sql, useCache);
SQLVisitorEngine sqlVisitorEngine = new SQLVisitorEngine(sql, "STATEMENT",
useCache, new Properties());
SQLStatement sqlStatement = sqlVisitorEngine.visit(parseASTNode);
```

- SQL 格式化

```
ParseASTNode parseASTNode = parserEngine.parse(sql, useCache);
SQLVisitorEngine sqlVisitorEngine = new SQLVisitorEngine(sql, "FORMAT", useCache,
new Properties());
String result = sqlVisitorEngine.visit(parseASTNode);
```

例子:

原 SQL	格式化 SQL
select a+1 as b, name n from table1 join table2 where id=1 and name= 'lu' ;	SELECT a + 1 AS b, name n FROM table1 JOIN table2 WHERE id = 1 and name = 'lu' ;
select id, name, age, sex, ss, yy from table1 where id=1;	SELECT id , name , age , sex , ss , yy FROM table1 WHERE id = 1;
select id, name, age, count(*) as n, (select id, name, age, sex from table2 where id=2) as sid, yyyy from table1 where id=1;	SELECT id , name , age , COUNT(*) AS n, (SELECT id , name , age , sex FROM table2 WHERE id = 2) AS sid, yyyy FROM table1 WHERE id = 1;
select id, name, age, sex, ss, yy from table1 where id=1 and name=1 and a=1 and b=2 and c=4 and d=3;	SELECT id , name , age , sex , ss , yy FROM table1 WHERE id = 1 and name = 1 and a = 1 and b = 2 and c = 4 and d = 3;
ALTER TABLE t_order ADD column4 DATE, ADD column5 DATETIME, engine ss max_rows 10,min_rows 2, ADD column6 TIMESTAMP, ADD column7 TIME;	ALTER TABLE t_order ADD column4 DATE, ADD column5 DATETIME, ENGINE ss MAX_ROWS 10, MIN_ROWS 2, ADD column6 TIMESTAMP, ADD column7 TIME
CREATE TABLE IF NOT EXISTS runoob_tbl(runoob_id INT UNSIGNED AUTO_INCREMENT,runoob_title VARCHAR(100) NOT NULL,runoob_author VARCHAR(40) NOT NULL,runoob_test NATIONAL CHAR(40),submission_date DATE,PRIMARY KEY (runoob_id))ENGINE=InnoDB DEFAULT CHARSET=utf8;	CREATE TABLE IF NOT EXISTS runoob_tbl (runoob_id INT UNSIGNED AUTO_INCREMENT, runoob_title VARCHAR(100) NOT NULL, runoob_author VARCHAR(40) NOT NULL , runoob_test NATIONAL CHAR(40), submission_date DATE, PRIMARY KEY (runoob_id)) ENGINE = InnoDB DEFAULT CHARSET = utf8;
INSERT INTO t_order_item(order_id, user_id, status, creation_date) values (1, 1, 'insert', '2017-08-08'),(2, 2, 'insert', '2017-08-08') ON DUPLICATE KEY UPDATE status = 'init' ;	INSERT INTO t_order_item (order_id , user_id , status , creation_date) VALUES (1, 1, 'insert', '2017-08-08'), (2, 2, 'insert', '2017-08-08') ON DUPLICATE KEY UPDATE status = 'init' ;
INSERT INTO t_order SET order_id = 1, user_id = 1, status = convert(to_base64(aes_encrypt(1, 'key')) USING utf8) ON DUPLICATE KEY UPDATE status = VALUES(status);	INSERT INTO t_order SET order_id = 1, user_id = 1, status = CONVERT(to _base64(aes_encrypt(1 , 'key')) USING utf8) ON DUPLICATE KEY UPDATE status = VALUES(status);
INSERT INTO t_order (order_id, user_id, status) SELECT order_id, user_id, status FROM t_order WHERE order_id = 1;	INSERT INTO t_order (order_id , user_id , status) SELECT order_id , user_id , status FROM t_order WHERE order_id = 1;

7.4.8 路由引擎

根据解析上下文匹配数据库和表的分片策略，并生成路由路径。对于携带分片键的 SQL，根据分片键的不同可以划分为单片路由（分片键的操作符是等号）、多片路由（分片键的操作符是 IN）和范围路由（分片键的操作符是 BETWEEN）。不携带分片键的 SQL 则采用广播路由。

分片策略通常可以采用由数据库内置或由用户方配置。数据库内置的方案较为简单，内置的分片策略大致可分为尾数取模、哈希、范围、标签、时间等。由用户方配置的分片策略则更加灵活，可以根据使用方需求定制复合分片策略。如果配合数据自动迁移来使用，可以做到无需用户关注分片策略，自动由数据库中间层分片和平衡数据即可，进而做到使分布式数据库具有的弹性伸缩的能力。在 ShardingSphere 的线路规划中，弹性伸缩将于 4.x 开启。

分片路由

用于根据分片键进行路由的场景，又细分为直接路由、标准路由和笛卡尔积路由这 3 种类型。

直接路由

满足直接路由的条件相对苛刻，它需要通过 Hint（使用 HintAPI 直接指定路由至库表）方式分片，并且是只分库不分表的前提下，则可以避免 SQL 解析和之后的结果归并。因此它的兼容性最好，可以执行包括子查询、自定义函数等复杂情况的任意 SQL。直接路由还可以用于分片键不在 SQL 中的场景。例如，设置用于数据库分片的键为 3，

```
hintManager.setDatabaseShardingValue(3);
```

假如路由算法为 $\text{value} \% 2$ ，当一个逻辑库 `t_order` 对应 2 个真实库 `t_order_0` 和 `t_order_1` 时，路由后 SQL 将在 `t_order_1` 上执行。下方是使用 API 的代码样例：

```
String sql = "SELECT * FROM t_order";
try (
    HintManager hintManager = HintManager.getInstance();
    Connection conn = dataSource.getConnection();
    PreparedStatement pstmt = conn.prepareStatement(sql)) {
    hintManager.setDatabaseShardingValue(3);
    try (ResultSet rs = pstmt.executeQuery()) {
        while (rs.next()) {
            //...
        }
    }
}
```


标准路由

标准路由是 ShardingSphere 最为推荐使用的分片方式，它的适用范围是不包含关联查询或仅包含绑定表之间关联查询的 SQL。当分片运算符是等于号时，路由结果将落入单库（表），当分片运算符是 BETWEEN 或 IN 时，则路由结果不一定落入唯一的库（表），因此一条逻辑 SQL 最终可能被拆分为多条用于执行的真实 SQL。举例说明，如果按照 order_id 的奇数和偶数进行数据分片，一个单表查询的 SQL 如下：

```
SELECT * FROM t_order WHERE order_id IN (1, 2);
```

那么路由的结果应为：

```
SELECT * FROM t_order_0 WHERE order_id IN (1, 2);
SELECT * FROM t_order_1 WHERE order_id IN (1, 2);
```

绑定表的关联查询与单表查询复杂度和性能相当。举例说明，如果一个包含绑定表的关联查询的 SQL 如下：

```
SELECT * FROM t_order o JOIN t_order_item i ON o.order_id=i.order_id WHERE order_id IN (1, 2);
```

那么路由的结果应为：

```
SELECT * FROM t_order_0 o JOIN t_order_item_0 i ON o.order_id=i.order_id WHERE order_id IN (1, 2);
SELECT * FROM t_order_1 o JOIN t_order_item_1 i ON o.order_id=i.order_id WHERE order_id IN (1, 2);
```

可以看到，SQL 拆分的数目与单表是一致的。

笛卡尔路由

笛卡尔路由是最复杂的情况，它无法根据绑定表的关系定位分片规则，因此非绑定表之间的关联查询需要拆解为笛卡尔积组合执行。如果上个示例中的 SQL 并未配置绑定表关系，那么路由的结果应为：

```
SELECT * FROM t_order_0 o JOIN t_order_item_0 i ON o.order_id=i.order_id WHERE order_id IN (1, 2);
SELECT * FROM t_order_0 o JOIN t_order_item_1 i ON o.order_id=i.order_id WHERE order_id IN (1, 2);
SELECT * FROM t_order_1 o JOIN t_order_item_0 i ON o.order_id=i.order_id WHERE order_id IN (1, 2);
SELECT * FROM t_order_1 o JOIN t_order_item_1 i ON o.order_id=i.order_id WHERE order_id IN (1, 2);
```

笛卡尔路由查询性能较低，需谨慎使用。

广播路由

对于不携带分片键的 SQL，则采取广播路由的方式。根据 SQL 类型又可以划分为全库表路由、全库路由、全实例路由、单播路由和阻断路由这 5 种类型。

全库表路由

全库表路由用于处理对数据库中与其逻辑表相关的所有真实表的操作，主要包括不带分片键的 DQL 和 DML，以及 DDL 等。例如：

```
SELECT * FROM t_order WHERE good_prority IN (1, 10);
```

则会遍历所有数据库中的所有表，逐一匹配逻辑表和真实表名，能够匹配得上则执行。路由后成为

```
SELECT * FROM t_order_0 WHERE good_prority IN (1, 10);  
SELECT * FROM t_order_1 WHERE good_prority IN (1, 10);  
SELECT * FROM t_order_2 WHERE good_prority IN (1, 10);  
SELECT * FROM t_order_3 WHERE good_prority IN (1, 10);
```

全库路由

全库路由用于处理对数据库的操作，包括用于库设置的 SET 类型的数据库管理命令，以及 TCL 这样的事务控制语句。在这种情况下，会根据逻辑库的名字遍历所有符合名字匹配的真实库，并在真实库中执行该命令，例如：

```
SET autocommit=0;
```

在 t_order 中执行，t_order 有 2 个真实库。则实际会在 t_order_0 和 t_order_1 上都执行这个命令。

全实例路由

全实例路由用于 DCL 操作，授权语句针对的是数据库的实例。无论一个实例中包含多少个 Schema，每个数据库的实例只执行一次。例如：

```
CREATE USER customer@127.0.0.1 identified BY '123';
```

这个命令将在所有的真实数据库实例中执行，以确保 customer 用户可以访问每一个实例。

单播路由

单播路由用于获取某一真实表信息的场景，它仅需要从任意库中的任意真实表中获取数据即可。例如：

```
DESCRIBE t_order;
```

t_order 的两个真实表 t_order_0, t_order_1 的描述结构相同，所以这个命令在任意真实表上选择执行一次。

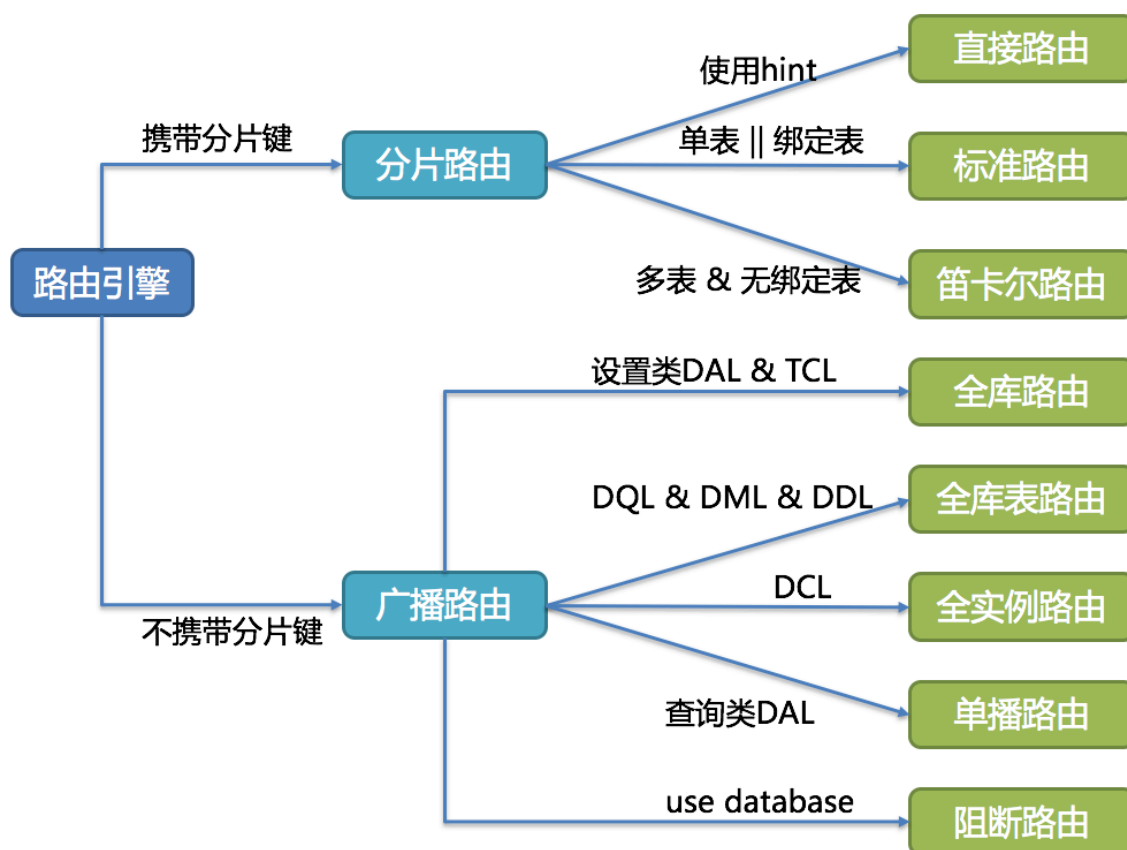
阻断路由

阻断路由用于屏蔽 SQL 对数据库的操作，例如：

```
USE order_db;
```

这个命令不会在真实数据库中执行，因为 ShardingSphere 采用的是逻辑 Schema 的方式，无需将切换数据库 Schema 的命令发送至数据库中。

路由引擎的整体结构划分如下图。



7.4.9 改写引擎

工程师面向逻辑库与逻辑表书写的 SQL，并不能够直接在真实的数据库中执行，SQL 改写用于将逻辑 SQL 改写为在真实数据库中正确执行的 SQL。它包括正确性改写和优化改写两部分。

正确性改写

在包含分表的场景中，需要将分表配置中的逻辑表名称改写为路由之后所获取的真实表名称。仅分库则不需要表名称的改写。除此之外，还包括补列和分页信息修正等内容。

标识符改写

需要改写的标识符包括表名称、索引名称以及 Schema 名称。

表名称改写是指将找到逻辑表在原始 SQL 中的位置，并将其改写为真实表的过程。表名称改写是一个典型的需要对 SQL 进行解析的场景。从一个最简单的例子开始，若逻辑 SQL 为：

```
SELECT order_id FROM t_order WHERE order_id=1;
```

假设该 SQL 配置分片键 order_id，并且 order_id=1 的情况，将路由至分片表 1。那么改写之后的 SQL 应该为：

```
SELECT order_id FROM t_order_1 WHERE order_id=1;
```

在这种最简单的 SQL 场景中，是否将 SQL 解析为抽象语法树似乎无关紧要，只要通过字符串查找和替换就可以达到 SQL 改写的效果。但是下面的场景，就无法仅仅通过字符串的查找替换来正确的改写 SQL 了：

```
SELECT order_id FROM t_order WHERE order_id=1 AND remarks=' t_order xxx';
```

正确改写的 SQL 应该是：

```
SELECT order_id FROM t_order_1 WHERE order_id=1 AND remarks=' t_order xxx';
```

而非：

```
SELECT order_id FROM t_order_1 WHERE order_id=1 AND remarks=' t_order_1 xxx';
```

由于表名之外可能含有表名称的类似字符，因此不能通过简单的字符串替换的方式去改写 SQL。

下面再来看一个更加复杂的 SQL 改写场景：

```
SELECT t_order.order_id FROM t_order WHERE t_order.order_id=1 AND remarks=' t_order xxx';
```

上面的 SQL 将表名作为字段的标识符，因此在 SQL 改写时需要一并修改：

```
SELECT t_order_1.order_id FROM t_order_1 WHERE t_order_1.order_id=1 AND remarks=' t_order xxx';
```

而如果 SQL 中定义了表的别名，则无需连同别名一起修改，即使别名与表名相同亦是如此。例如：

```
SELECT t_order.order_id FROM t_order AS t_order WHERE t_order.order_id=1 AND
remarks=' t_order xxx';
```

SQL 改写则仅需要改写表名称就可以了：

```
SELECT t_order.order_id FROM t_order_1 AS t_order WHERE t_order.order_id=1 AND
remarks=' t_order xxx';
```

索引名称是另一个有可能改写的标识符。在某些数据库中（如 MySQL、SQLServer），索引是以表为维度创建的，在不同的表中的索引是可以重名的；而在另外的一些数据库中（如 PostgreSQL、Oracle），索引是以数据库为维度创建的，即使是作用在不同表上的索引，它们也要求其名称的唯一性。

在 ShardingSphere 中，管理 Schema 的方式与管理表如出一辙，它采用逻辑 Schema 去管理一组数据源。因此，ShardingSphere 需要将用户在 SQL 中书写的逻辑 Schema 替换为真实的数据库 Schema。

ShardingSphere 目前还不支持在 DQL 和 DML 语句中使用 Schema。它目前仅支持在数据库管理语句中使用 Schema，例如：

```
SHOW COLUMNS FROM t_order FROM order_ds;
```

Schema 的改写指的是将逻辑 Schema 采用单播路由的方式，改写为随机查找到的一个正确的真实 Schema。

补列

需要在查询语句中补列通常由两种情况导致。第一种情况是 ShardingSphere 需要在结果归并时获取相应数据，但该数据并未能通过查询的 SQL 返回。这种情况主要是针对 GROUP BY 和 ORDER BY。结果归并时，需要根据 GROUP BY 和 ORDER BY 的字段项进行分组和排序，但如果原始 SQL 的选择项中若并未包含分组项或排序项，则需要对原始 SQL 进行改写。先看一下原始 SQL 中带有结果归并所需信息的场景：

```
SELECT order_id, user_id FROM t_order ORDER BY user_id;
```

由于使用 user_id 进行排序，在结果归并中需要能够获取到 user_id 的数据，而上面的 SQL 是能够获取到 user_id 数据的，因此无需补列。

如果选择项中不包含结果归并时所需的列，则需要进行补列，如以下 SQL：

```
SELECT order_id FROM t_order ORDER BY user_id;
```

由于原始 SQL 中并不包含需要在结果归并中需要获取的 user_id，因此需要对 SQL 进行补列改写。补列之后的 SQL 是：

```
SELECT order_id, user_id AS ORDER_BY_DERIVED_0 FROM t_order ORDER BY user_id;
```

值得一提的是，补列只会补充缺失的列，不会全部补充，而且，在 SELECT 语句中包含 * 的 SQL，也会根据表的元数据信息选择性补列。下面是一个较为复杂的 SQL 补列场景：

```
SELECT o.* FROM t_order o, t_order_item i WHERE o.order_id=i.order_id ORDER BY
user_id, order_item_id;
```

我们假设只有 t_order_item 表中包含 order_item_id 列，那么根据表的元数据信息可知，在结果归并时，排序项中的 user_id 是存在于 t_order 表中的，无需补列；order_item_id 并不在 t_order 中，因此需要补列。补列之后的 SQL 是：

```
SELECT o.*, order_item_id AS ORDER_BY_DERIVED_0 FROM t_order o, t_order_item i
WHERE o.order_id=i.order_id ORDER BY user_id, order_item_id;
```

补列的另一种情况是使用 AVG 聚合函数。在分布式的场景中，使用 $avg1 + avg2 + avg3 / 3$ 计算平均值并不正确，需要改写为 $(sum1 + sum2 + sum3) / (count1 + count2 + count3)$ 。这就需要将包含 AVG 的 SQL 改写为 SUM 和 COUNT，并在结果归并时重新计算平均值。例如以下 SQL：

```
SELECT AVG(price) FROM t_order WHERE user_id=1;
```

需要改写为：

```
SELECT COUNT(price) AS AVG_DERIVED_COUNT_0, SUM(price) AS AVG_DERIVED_SUM_0 FROM t_
order WHERE user_id=1;
```

然后才能够通过结果归并正确的计算平均值。

最后一种补列是在执行 INSERT 的 SQL 语句时，如果使用数据库自增主键，是无需写入主键字段的。但数据库的自增主键是无法满足分布式场景下的主键唯一的，因此 ShardingSphere 提供了分布式自增主键的生成策略，并且可以通过补列，让使用方无需改动现有代码，即可将分布式自增主键透明的替换数据库现有的自增主键。分布式自增主键的生成策略将在下文中详述，这里只阐述与 SQL 改写相关的内容。举例说明，假设表 t_order 的主键是 order_id，原始的 SQL 为：

```
INSERT INTO t_order (`field1`, `field2`) VALUES (10, 1);
```

可以看到，上述 SQL 中并未包含自增主键，是需要数据库自行填充的。ShardingSphere 配置自增主键后，SQL 将改写为：

```
INSERT INTO t_order (`field1`, `field2`, order_id) VALUES (10, 1, xxxxx);
```

改写后的 SQL 将在 INSERT FIELD 和 INSERT VALUE 的最后部分增加主键列名称以及自动生成的自增主键值。上述 SQL 中的 xxxxx 表示自动生成的自增主键值。

如果 INSERT 的 SQL 中并未包含表的列名称，ShardingSphere 也可以根据判断参数个数以及表元信息中的列数量对比，并自动生成自增主键。例如，原始的 SQL 为：

```
INSERT INTO t_order VALUES (10, 1);
```

改写的 SQL 将只在主键所在的列顺序处增加自增主键即可：

```
INSERT INTO t_order VALUES (xxxxx, 10, 1);
```

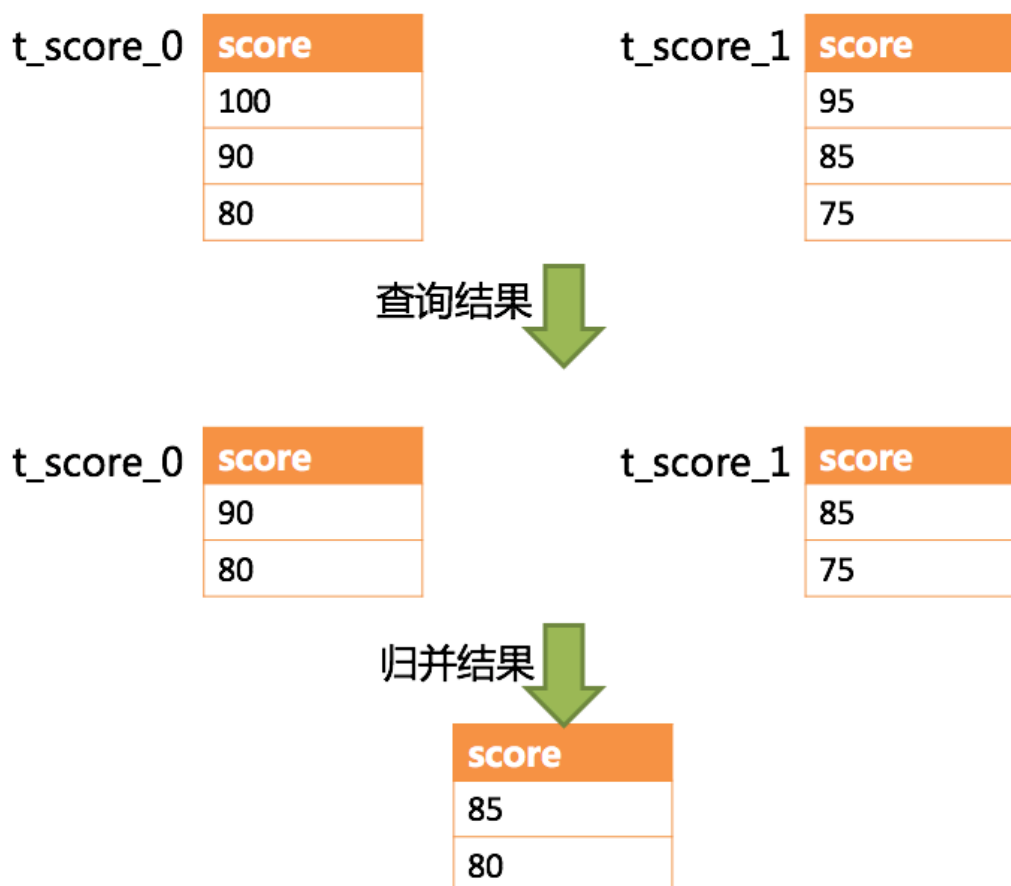
自增主键补列时，如果使用占位符的方式书写 SQL，则只需要改写参数列表即可，无需改写 SQL 本身。

分页修正

从多个数据库获取分页数据与单数据库的场景是不同的。假设每 10 条数据为一页，取第 2 页数据。在分片环境下获取 `LIMIT 10, 10`，归并之后再根据排序条件取出前 10 条数据是不正确的。举例说明，若 SQL 为：

```
SELECT score FROM t_score ORDER BY score DESC LIMIT 1, 2;
```

下图展示了不进行 SQL 的改写的分页执行结果。



通过图中所示，想要取得两个表中共同的按照分数排序的第 2 条和第 3 条数据，应该是 95 和 90。由于执行的 SQL 只能从每个表中获取第 2 条和第 3 条数据，即从 `t_score_0` 表中获取的是 90 和 80；从 `t_score_1` 表中获取的是 85 和 75。因此进行结果归并时，只能从获取的 90，80，85 和 75 之中进行归并，那么结果归并无论怎么实现，都不可能获得正确的结果。

正确的做法是将分页条件改写为 `LIMIT 0, 3`，取出所有前两页数据，再结合排序条件计算出正确的数据。下图展示了进行 SQL 改写之后的分页执行结果。

SELECT score FROM t_score ORDER BY score DESC LIMIT 0, 3

查询结果



t_score_0	score
	100
	90
	80

t_score_1	score
	95
	85
	75

归并结果



score
95
90

越获取偏移量位置靠后数据，使用 LIMIT 分页方式的效率就越低。有很多方法可以避免使用 LIMIT 进行分页。比如构建行记录数量与行偏移量的二级索引，或使用上次分页数据结尾 ID 作为下次查询条件的分页方式等。

分页信息修正时，如果使用占位符的方式书写 SQL，则只需要改写参数列表即可，无需改写 SQL 本身。

批量拆分

在使用批量插入的 SQL 时，如果插入的数据是跨分片的，那么需要对 SQL 进行改写来防止将多余的数据写入到数据库中。插入操作与查询操作的不同之处在于，查询语句中即使用了不存在于当前分片的分片键，也不会对数据产生影响；而插入操作则必须将多余的分片键删除。举例说明，如下 SQL：

```
INSERT INTO t_order (order_id, xxx) VALUES (1, 'xxx'), (2, 'xxx'), (3, 'xxx');
```

假设数据库仍然是按照 order_id 的奇偶值分为两片的，仅将这条 SQL 中的表名进行修改，然后发送至数据库完成 SQL 的执行，则两个分片都会写入相同的记录。虽然只有符合分片查询条件的数据才能够被查询语句取出，但存在冗余数据的实现方案并不合理。因此需要将 SQL 改写为：

```
INSERT INTO t_order_0 (order_id, xxx) VALUES (2, 'xxx');
INSERT INTO t_order_1 (order_id, xxx) VALUES (1, 'xxx'), (3, 'xxx');
```

使用 IN 的查询与批量插入的情况相似，不过 IN 操作并不会导致数据查询结果错误。通过对 IN 查询的改写，可以进一步的提升查询性能。如以下 SQL：

```
SELECT * FROM t_order WHERE order_id IN (1, 2, 3);
```

改写为：


```
SELECT * FROM t_order_0 WHERE order_id IN (2);  
SELECT * FROM t_order_1 WHERE order_id IN (1, 3);
```

可以进一步的提升查询性能。ShardingSphere 暂时还未实现此改写策略，目前的改写结果是：

```
SELECT * FROM t_order_0 WHERE order_id IN (1, 2, 3);  
SELECT * FROM t_order_1 WHERE order_id IN (1, 2, 3);
```

虽然 SQL 的执行结果是正确的，但并未达到最优的查询效率。

优化改写

优化改写的目的是在不影响查询正确性的情况下，对性能进行提升的有效手段。它分为单节点优化和流式归并优化。

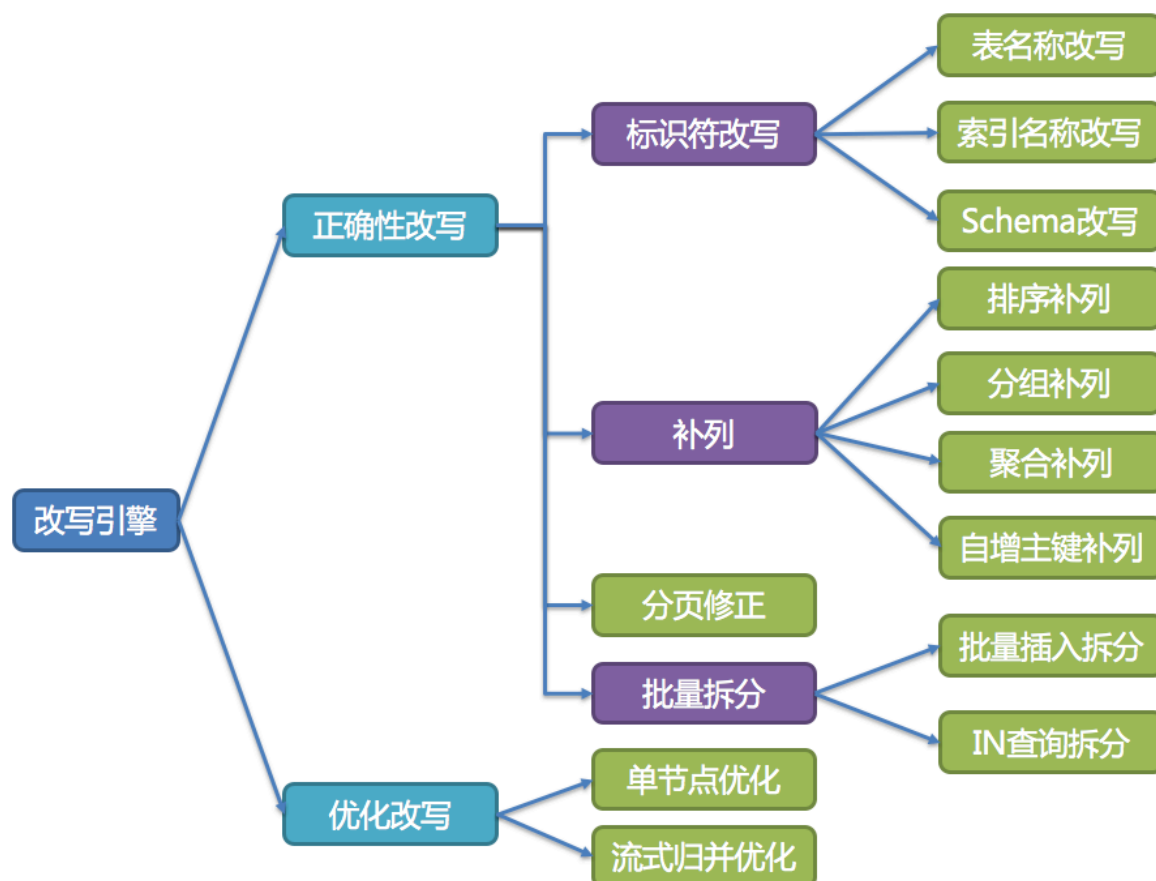
单节点优化

路由至单节点的 SQL，则无需优化改写。当获得一次查询的路由结果后，如果是路由至唯一的数据节点，则无需涉及到结果归并。因此补列和分页信息等改写都没有必要进行。尤其是分页信息的改写，无需将数据从第 1 条开始取，大量的降低了对数据库的压力，并且节省了网络带宽的无谓消耗。

流式归并优化

它仅为包含 GROUP BY 的 SQL 增加 ORDER BY 以及和分组项相同的排序项和排序顺序，用于将内存归并转化为流式归并。在结果归并的部分中，将对流式归并和内存归并进行详细说明。

改写引擎的整体结构划分如下图所示。



7.4.10 执行引擎

ShardingSphere 采用一套自动化的执行引擎，负责将路由和改写完成之后的真实 SQL 安全且高效发送到底层数据源执行。它不是简单地将 SQL 通过 JDBC 直接发送至数据源执行；也并非直接将执行请求放入线程池去并发执行。它更关注平衡数据源连接创建以及内存占用所产生的消耗，以及最大限度地合理利用并发等问题。执行引擎的目标是自动化的平衡资源控制与执行效率。

连接模式

从资源控制的角度看，业务方访问数据库的连接数量应当有所限制。它能够有效地防止某一业务操作过多的占用资源，从而将数据库连接的资源耗尽，以致于影响其他业务的正常访问。特别是在一个数据库实例中存在较多分表的情况下，一条不包含分片键的逻辑 SQL 将产生落在同库不同表的大量真实 SQL，如果每条真实 SQL 都占用一个独立的连接，那么一次查询无疑将会占用过多的资源。

从执行效率的角度看，为每个分片查询维持一个独立的数据库连接，可以更加有效的利用多线程来提升执行效率。为每个数据库连接开启独立的线程，可以将 I/O 所产生的消耗并行处理。为每个分片维持一个独立的数据库连接，还能够避免过早的将查询结果数据加载至内存。独立的数据库连接，能够持有查询结果集游标位置的引用，在需要获取相应数据时移动游标即可。

以结果集游标下移进行结果归并的方式，称之为流式归并，它无需将结果数据全数加载至内存，可以有效的节省内存资源，进而减少垃圾回收的频次。当无法保证每个分片查询持有一个独立数据库连接时，则

需要在复用该数据库连接获取下一张分表的查询结果集之前，将当前的查询结果集全数加载至内存。因此，即使可以采用流式归并，在此场景下也将退化为内存归并。

一方面是对数据库连接资源的控制保护，一方面是采用更优的归并模式达到对中间件内存资源的节省，如何处理好两者之间的关系，是 ShardingSphere 执行引擎需要解决的问题。具体来说，如果一条 SQL 在经过 ShardingSphere 的分片后，需要操作某数据库实例下的 200 张表。那么，是选择创建 200 个连接并行执行，还是选择创建一个连接串行执行呢？效率与资源控制又应该如何抉择呢？

针对上述场景，ShardingSphere 提供了一种解决思路。它提出了连接模式（Connection Mode）的概念，将其划分为内存限制模式（MEMORY_STRICTLY）和连接限制模式（CONNECTION_STRICTLY）这两种类型。

内存限制模式

使用此模式的前提是，ShardingSphere 对一次操作所耗费的数据库连接数量不做限制。如果实际执行的 SQL 需要对某数据库实例中的 200 张表做操作，则对每张表创建一个新的数据库连接，并通过多线程的方式并发处理，以达成执行效率最大化。并且在 SQL 满足条件情况下，优先选择流式归并，以防止出现内存溢出或避免频繁垃圾回收情况。

连接限制模式

使用此模式的前提是，ShardingSphere 严格控制对一次操作所耗费的数据库连接数量。如果实际执行的 SQL 需要对某数据库实例中的 200 张表做操作，那么只会创建唯一的数据库连接，并对其 200 张表串行处理。如果一次操作中的分片散落在不同的数据库，仍然采用多线程处理对不同库的操作，但每个库的每次操作仍然只创建一个唯一的数据库连接。这样即可以防止对一次请求对数据库连接占用过多所带来的问题。该模式始终选择内存归并。

内存限制模式适用于 OLAP 操作，可以通过放宽对数据库连接的限制提升系统吞吐量；连接限制模式适用于 OLTP 操作，OLTP 通常带有分片键，会路由到单一的分片，因此严格控制数据库连接，以保证在线系统数据库资源能够被更多的应用所使用，是明智的选择。

自动化执行引擎

ShardingSphere 最初将使用何种模式的决定权交由用户配置，让开发者依据自己业务的实际场景需求选择使用内存限制模式或连接限制模式。

这种解决方案将两难的选择的决定权交由用户，使得用户必须要了解这两种模式的利弊，并依据业务场景需求进行选择。这无疑增加了用户对 ShardingSphere 的学习和使用的成本，并非最优方案。

这种一分为二的处理方案，将两种模式的切换交由静态的初始化配置，是缺乏灵活应对能力的。在实际的使用场景中，面对不同 SQL 以及占位符参数，每次的路由结果是不同的。这就意味着某些操作可能需要使用内存归并，而某些操作则可能选择流式归并更优，具体采用哪种方式不应该由用户在 ShardingSphere 启动之前配置好，而是应该根据 SQL 和占位符参数的场景，来动态的决定连接模式。

为了降低用户的使用成本以及连接模式动态化这两个问题，ShardingSphere 提炼出自动化执行引擎的思路，在其内部消化了连接模式概念。用户无需了解所谓的内存限制模式和连接限制模式是什么，而是交由执行引擎根据当前场景自动选择最优的执行方案。

自动化执行引擎将连接模式的选择粒度细化至每一次 SQL 的操作。针对每次 SQL 请求，自动化执行引擎都将根据其路由结果，进行实时的演算和权衡，并自主地采用恰当的连接模式执行，以达到资源控制和效率的最优平衡。针对自动化的执行引擎，用户只需配置 `maxConnectionSizePerQuery` 即可，该参数表示一次查询时每个数据库所允许使用的最大连接数。

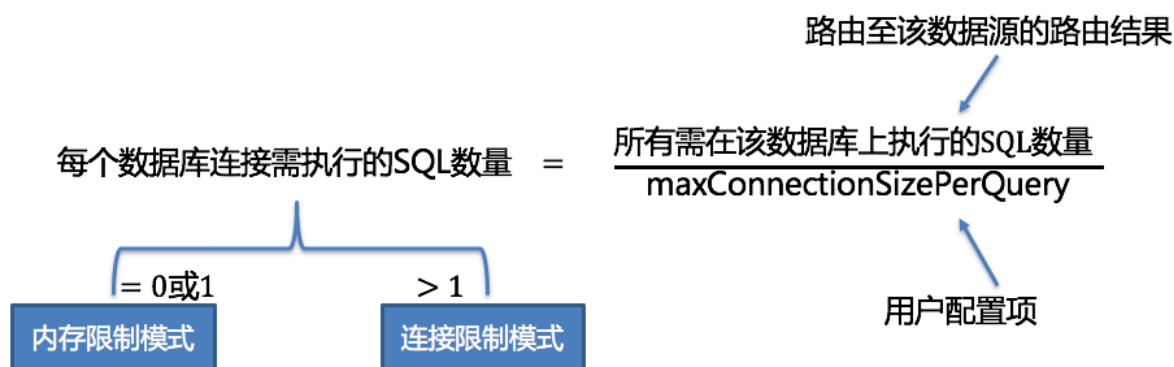
执行引擎分为准备和执行两个阶段。

准备阶段

顾名思义，此阶段用于准备执行的数据。它分为结果集分组和执行单元创建两个步骤。

结果集分组是实现内化连接模式概念的关键。执行引擎根据 `maxConnectionSizePerQuery` 配置项，结合当前路由结果，选择恰当的连接模式。具体步骤如下：

1. 将 SQL 的路由结果按照数据源的名称进行分组。
2. 通过下图的公式，可以获得每个数据库实例在 `maxConnectionSizePerQuery` 的允许范围内，每个连接需要执行的 SQL 路由结果组，并计算出本次请求的最优连接模式。

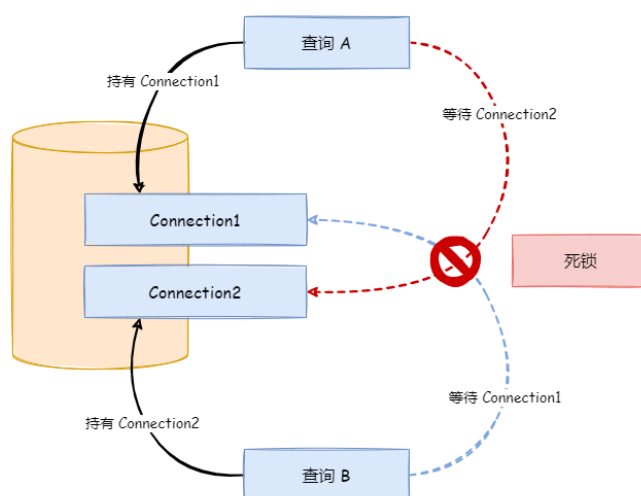


在 `maxConnectionSizePerQuery` 允许的范围内，当一个连接需要执行的请求数量大于 1 时，意味着当前的数据库连接无法持有相应的数据结果集，则必须采用内存归并；反之，当一个连接需要执行的请求数量等于 1 时，意味着当前的数据库连接可以持有相应的数据结果集，则可以采用流式归并。

每一次的连接模式的选择，是针对每一个物理数据库的。也就是说，在同一次查询中，如果路由至一个以上的数据库，每个数据库的连接模式不一定一样，它们可能是混合存在的形态。

通过上一步骤获得的路由分组结果创建执行的单元。当数据源使用数据库连接池等控制数据库连接数量的技术时，在获取数据库连接时，如果不妥善处理并发，则有一定几率发生死锁。在多个请求相互等待对方释放数据库连接资源时，将会产生饥饿等待，造成交叉的死锁问题。

举例说明，假设一次查询需要在某一数据源上获取两个数据库连接，并路由至同一个数据库的两个分表查询。则有可能出现查询 A 已获取到该数据源的 1 个数据库连接，并等待获取另一个数据库连接；而查询 B 也已经在该数据源上获取到的一个数据库连接，并同样等待另一个数据库连接的获取。如果数据库连接池的允许最大连接数是 2，那么这 2 个查询请求将永久的等待下去。下图描绘了死锁的情况。



ShardingSphere 为了避免死锁的出现，在获取数据库连接时进行了同步处理。它在创建执行单元时，以原子性的方式一次性获取本次 SQL 请求所需的全部数据库连接，杜绝了每次查询请求获取到部分资源的可能。由于对数据库的操作非常频繁，每次获取数据库连接时时都进行锁定，会降低 ShardingSphere 的并发。因此，ShardingSphere 在这里进行了 2 点优化：

1. 避免锁定一次性只需要获取 1 个数据库连接的操作。因为每次仅需要获取 1 个连接，则不会发生两个请求相互等待的场景，无需锁定。对于大部分 OLTP 的操作，都是使用分片键路由至唯一的数据节点，这会使得系统变为完全无锁的状态，进一步提升了并发效率。除了路由至单分片的情况，读写分离也在此范畴之内。
2. 仅针对内存限制模式时才进行资源锁定。在使用连接限制模式时，所有的查询结果集将在装载至内存之后释放掉数据库连接资源，因此不会产生死锁等待的问题。

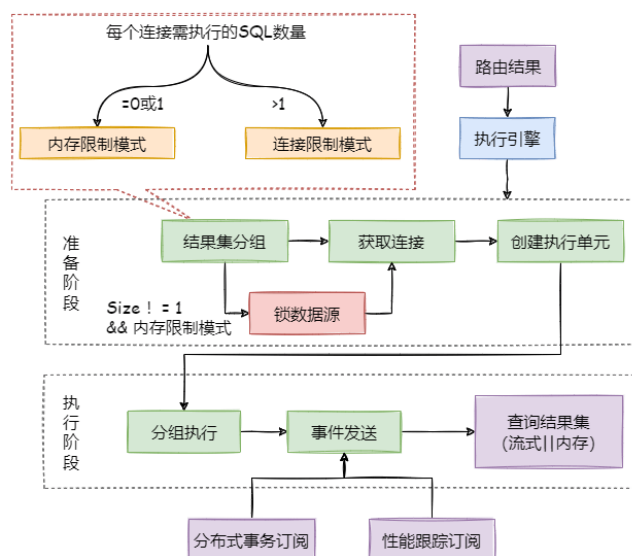
执行阶段

该阶段用于真正的执行 SQL，它分为分组执行和归并结果集生成两个步骤。

分组执行将准备执行阶段生成的执行单元分组下发至底层并发执行引擎，并针对执行过程中的每个关键步骤发送事件。如：执行开始事件、执行成功事件以及执行失败事件。执行引擎仅关注事件的发送，它并不关心事件的订阅者。ShardingSphere 的其他模块，如：分布式事务、调用链路追踪等，会订阅感兴趣的事件，并进行相应的处理。

ShardingSphere 通过在执行准备阶段的获取的连接模式，生成内存归并结果集或流式归并结果集，并将其传递至结果归并引擎，以进行下一步的工作。

执行引擎的整体结构划分如下图所示。



7.4.11 归并引擎

将从各个数据节点获取的多数据结果集，组合成为一个结果集并正确的返回至请求客户端，称为结果归并。

ShardingSphere 支持的结果归并从功能上分为遍历、排序、分组、分页和聚合 5 种类型，它们是组合而非互斥的关系。从结构划分，可分为流式归并、内存归并和装饰者归并。流式归并和内存归并是互斥的，装饰者归并可以在流式归并和内存归并之上做进一步的处理。

由于从数据库中返回的结果集是逐条返回的，并不需要将所有的数据一次性加载至内存中，因此，在进行结果归并时，沿用数据库返回结果集的方式进行归并，能够极大减少内存的消耗，是归并方式的优先选择。

流式归并是指每一次从结果集中获取到的数据，都能够通过逐条获取的方式返回正确的单条数据，它与数据库原生的返回结果集的方式最为契合。遍历、排序以及流式分组都属于流式归并的一种。

内存归并则是需要将结果集的所有数据都遍历并存储在内存中，再通过统一的分组、排序以及聚合等计算之后，再将其封装成为逐条访问的数据结果集返回。

装饰者归并是对所有的结果集归并进行统一的功能增强，目前装饰者归并有分页归并和聚合归并这 2 种类型。

遍历归并

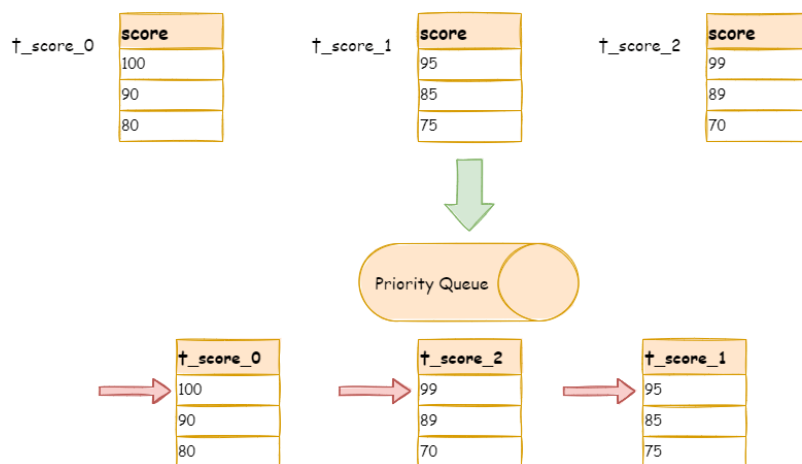
它是最为简单的归并方式。只需将多个数据结果集合并为一个单向链表即可。在遍历完成链表中当前数据结果集之后，将链表元素后移一位，继续遍历下一个数据结果集即可。

排序归并

由于在 SQL 中存在 ORDER BY 语句，因此每个数据结果集自身是有序的，因此只需要将数据结果集当前游标指向的数据值进行排序即可。这相当于对多个有序的数组进行排序，归并排序是最适合此场景的排序算法。

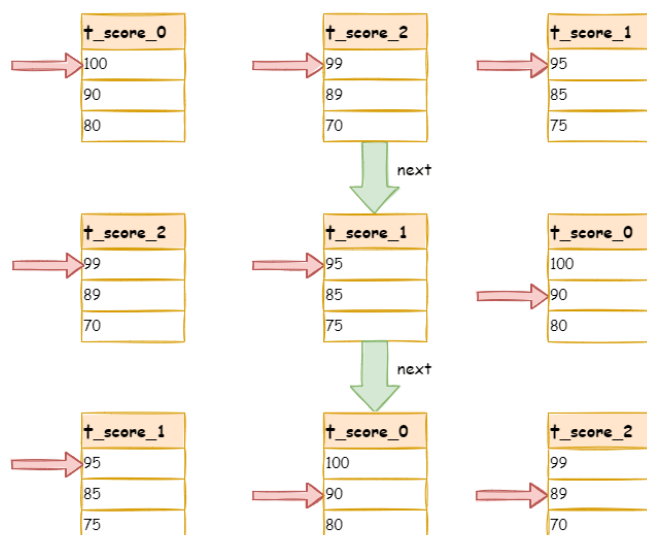
ShardingSphere 在对排序的查询进行归并时，将每个结果集的当前数据值进行比较（通过实现 Java 的 Comparable 接口完成），并将其放入优先级队列。每次获取下一条数据时，只需将队列顶端结果集的游标下移，并根据新游标重新进入优先级排序队列找到自己的位置即可。

通过一个例子来说明 ShardingSphere 的排序归并，下图是一个通过分数进行排序的示例图。图中展示了 3 张表返回的数据结果集，每个数据结果集已经根据分数排序完毕，但是 3 个数据结果集之间是无序的。将 3 个数据结果集的当前游标指向的数据值进行排序，并放入优先级队列，t_score_0 的第一个数据值最大，t_score_2 的第一个数据值次之，t_score_1 的第一个数据值最小，因此优先级队列根据 t_score_0, t_score_2 和 t_score_1 的方式排序队列。



下图则展现了进行 next 调用的时候，排序归并是如何进行的。通过图中我们可以看到，当进行第一次 next 调用时，排在队列首位的 t_score_0 将会被弹出队列，并且将当前游标指向的数据值（也就是 100）返回至查询客户端，并且将游标下移一位之后，重新放入优先级队列。而优先级队列也会根据 t_score_0 的当前数据结果集指向游标的数据值（这里是 90）进行排序，根据当前数值，t_score_0 排列在队列的最后一位。之前队列中排名第二的 t_score_2 的数据结果集则自动排在了队列首位。

在进行第二次 `next` 时，只需要将目前排列在队列首位的 `t_score_2` 弹出队列，并且将其数据结果集游标指向的值返回至客户端，并下移游标，继续加入队列排队，以此类推。当一个结果集中已经没有数据了，则无需再次加入队列。



可以看到，对于每个数据结果集中的数据有序，而多数数据结果集整体无序的情况下，ShardingSphere 无需将所有数据都加载至内存即可排序。它使用的是流式归并的方式，每次 `next` 仅获取唯一正确的一条数据，极大的节省了内存的消耗。

从另一个角度来说，ShardingSphere 的排序归并，是在维护数据结果集的纵轴和横轴这两个维度的有序性。纵轴是指每个数据结果集本身，它是天然有序的，它通过包含 `ORDER BY` 的 SQL 所获取。横轴是指每个数据结果集当前游标所指向的值，它需要通过优先级队列来维护其正确顺序。每一次数据结果集当前游标的下移，都需要将该数据结果集重新放入优先级队列排序，而只有排列在队列首位的数据结果集才可能发生游标下移的操作。

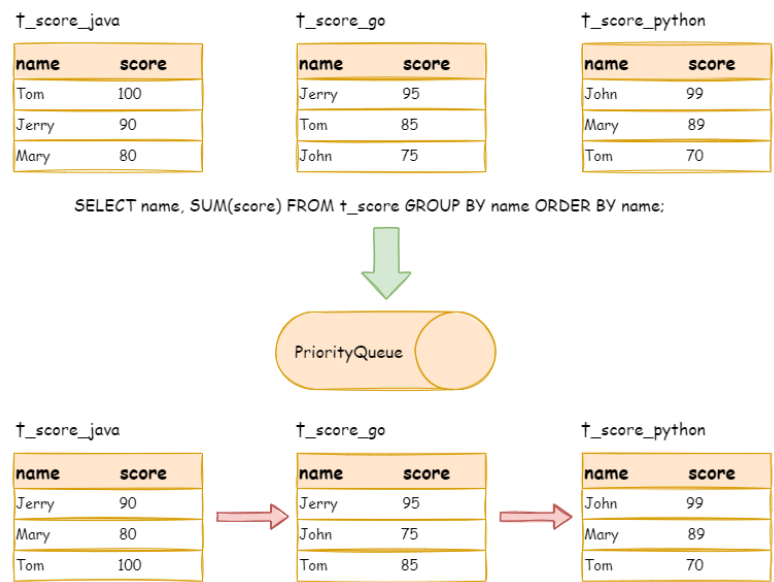
分组归并

分组归并的情况最为复杂，它分为流式分组归并和内存分组归并。流式分组归并要求 SQL 的排序项与分组项的字段以及排序类型（ASC 或 DESC）必须保持一致，否则只能通过内存归并才能保证数据的正确性。

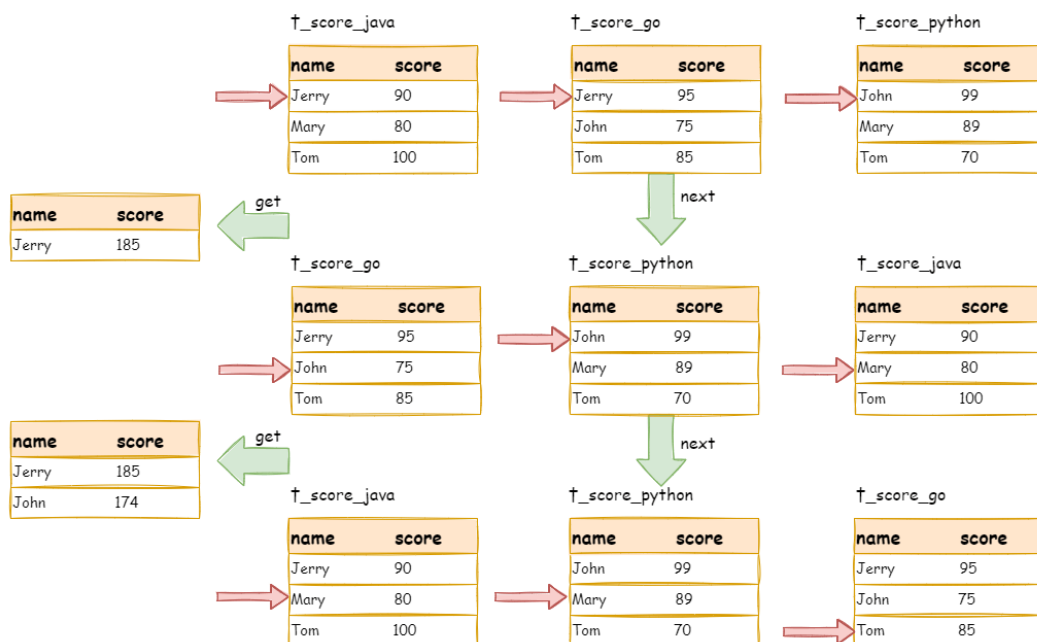
举例说明，假设根据科目分片，表结构中包含考生的姓名（为了简单起见，不考虑重名的情况）和分数。通过 SQL 获取每位考生的总分，可通过如下 SQL：

```
SELECT name, SUM(score) FROM t_score GROUP BY name ORDER BY name;
```


在分组项与排序项完全一致的情况下，取得的数据是连续的，分组所需的数据全数存在于各个数据结果集的当前游标所指向的数据值，因此可以采用流式归并。如下图所示。



进行归并时，逻辑与排序归并类似。下图展现了进行 `next` 调用的时候，流式分组归并是如何进行的。



通过图中我们可以看到，当进行第一次 `next` 调用时，排在队列首位的 `t_score_java` 将会被弹出队列，并且将分组值同为“Jerry”的其他结果集中的数据一同弹出队列。在获取了所有的姓名为“Jerry”的同学的分数之后，进行累加操作，那么，在第一次 `next` 调用结束后，取出的结果集是“Jerry”的分数总和。与此同时，所有的数据结果集中的游标都将下移至数据值“Jerry”的下一个不同的数据值，并且根据数据结果集当前游标指向的值进行重排序。因此，包含名字顺着第二位的“John”的相关数据结果集则排在队列的前列。

流式分组归并与排序归并的区别仅仅在于两点：

1. 它会一次性的将多个数据结果集中的分组项相同的数据全数取出。
2. 它需要根据聚合函数的类型进行聚合计算。

对于分组项与排序项不一致的情况，由于需要获取分组的相关的数据值并非连续的，因此无法使用流式归并，需要将所有的结果集数据加载至内存中进行分组和聚合。例如，若通过以下 SQL 获取每位考生的总分并按照分数从高至低排序：

```
SELECT name, SUM(score) FROM t_score GROUP BY name ORDER BY score DESC;
```

那么各个数据结果集中取出的数据与排序归并那张图的上半部分的表结构的原始数据一致，是无法进行流式归并的。

当 SQL 中只包含分组语句时，根据不同数据库的实现，其排序的顺序不一定与分组顺序一致。但由于排序语句的缺失，则表示此 SQL 并不在意排序顺序。因此，ShardingSphere 通过 SQL 优化的改写，自动增加与分组项一致的排序项，使其能够从消耗内存的内存分组归并方式转化为流式分组归并方案。

聚合归并

无论是流式分组归并还是内存分组归并，对聚合函数的处理都是一致的。除了分组的 SQL 之外，不进行分组的 SQL 也可以使用聚合函数。因此，聚合归并是在之前介绍的归并类的之上追加的归并能力，即装饰者模式。聚合函数可以归类为比较、累加和求平均值这 3 种类型。

比较类型的聚合函数是指 MAX 和 MIN。它们需要对每一个同组的结果集数据进行比较，并且直接返回其最大或最小值即可。

累加类型的聚合函数是指 SUM 和 COUNT。它们需要将每一个同组的结果集数据进行累加。

求平均值的聚合函数只有 AVG。它必须通过 SQL 改写的 SUM 和 COUNT 进行计算，相关内容已在 SQL 改写的内容中涵盖，不再赘述。

分页归并

上文所述的所有归并类型都可能进行分页。分页也是追加在其他归并类型之上的装饰器，ShardingSphere 通过装饰者模式来增加对数据结果集进行分页的能力。分页归并负责将无需获取的数据过滤掉。

ShardingSphere 的分页功能比较容易让使用者误解，用户通常认为分页归并会占用大量内存。在分布式的场景中，将 LIMIT 100000000, 10 改写为 LIMIT 0, 10000010，才能保证其数据的正确性。用户非常容易产生 ShardingSphere 会将大量无意义的的数据加载至内存中，造成内存溢出风险的错觉。其实，通过流式归并的原理可知，会将数据全部加载到内存中的只有内存分组归并这一种情况。而通常来说，进行 OLAP 的分组 SQL，不会产生大量的结果数据，它更多的用于大量的计算，以及少量结果产出的场景。除了内存分组归并这种情况之外，其他情况都通过流式归并获取数据结果集，因此 ShardingSphere 会通过结果集的 next 方法将无需取出的数据全部跳过，并不会将其存入内存。

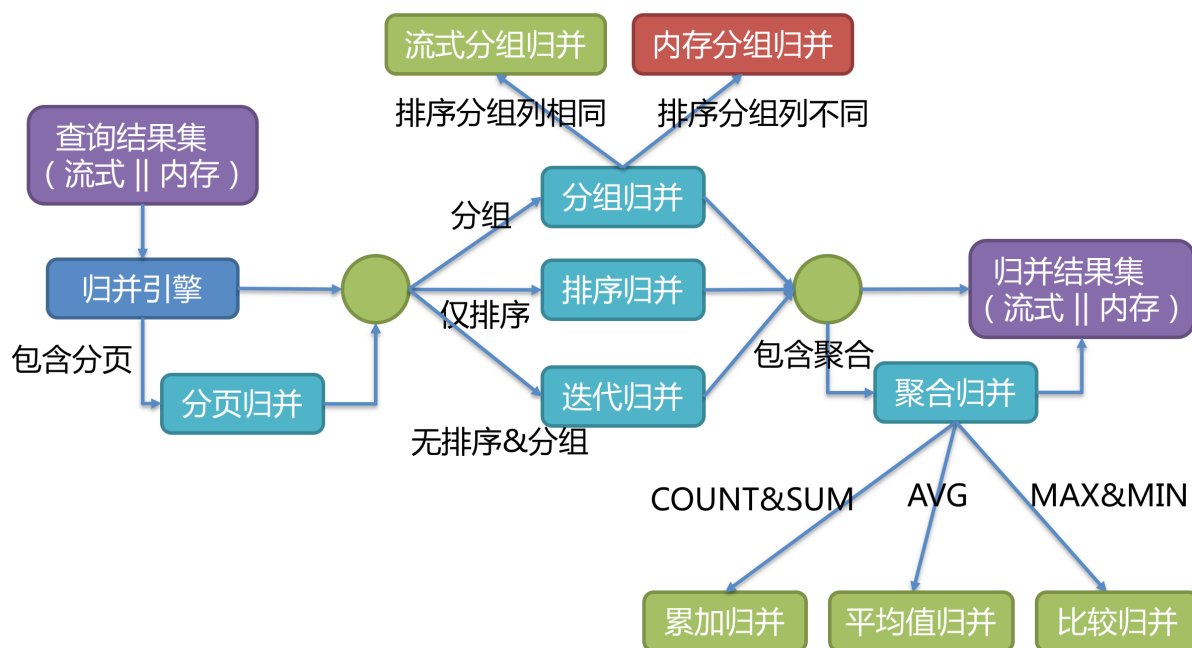
但同时需要注意的是，由于排序的需要，大量的数据仍然需要传输到 ShardingSphere 的内存空间。因此，采用 LIMIT 这种方式分页，并非最佳实践。由于 LIMIT 并不能通过索引查询数据，因此如果可以保证 ID 的连续性，通过 ID 进行分页是比较好的解决方案，例如：

```
SELECT * FROM t_order WHERE id > 100000 AND id <= 100010 ORDER BY id;
```

或通过记录上次查询结果的最后一条记录的 ID 进行下一页的查询，例如：

```
SELECT * FROM t_order WHERE id > 100000000 LIMIT 10;
```

归并引擎的整体结构划分如下图。



7.5 分布式事务

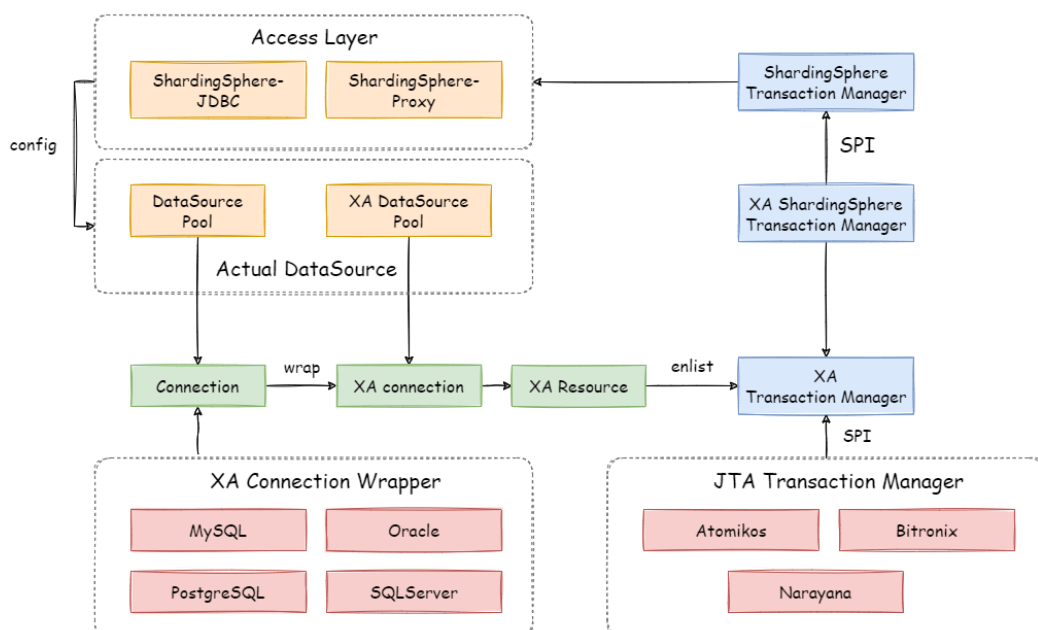
7.5.1 导览

本小节主要介绍 Apache ShardingSphere 分布式事务的实现原理

- 基于 XA 协议的两阶段事务
- 基于 Seata 的柔性事务

7.5.2 XA 事务

XAShardingSphereTransactionManager 为 Apache ShardingSphere 的分布式事务的 XA 实现类。它主要负责对多数据源进行管理和适配，并且将相应事务的开启、提交和回滚操作委托给具体的 XA 事务管理器。



开启全局事务

收到接入端的 `set autoCommit=0` 时, `XAShardingSphereTransactionManager` 将调用具体的 `XA` 事务管理器开启 `XA` 全局事务, 以 `XID` 的形式进行标记。

执行真实分片 SQL

`XAShardingSphereTransactionManager` 将数据库连接所对应的 `XAResource` 注册到当前 `XA` 事务中之后, 事务管理器会在此阶段发送 `XAResource.start` 命令至数据库。数据库在收到 `XAResource.end` 命令之前的所有 `SQL` 操作, 会被标记为 `XA` 事务。

例如:

```

XAResource1.start      ## Enlist 阶段执行
statement.execute("sql1");  ## 模拟执行一个分片 SQL1
statement.execute("sql2");  ## 模拟执行一个分片 SQL2
XAResource1.end        ## 提交阶段执行
  
```

示例中的 `sql1` 和 `sql2` 将会被标记为 `XA` 事务。

提交或回滚事务

XAShardingSphereTransactionManager 在接收到接入端的提交命令后，会委托实际的 XA 事务管理进行提交动作，事务管理器将收集到的当前线程中所有注册的 XAResource，并发送 XAResource.end 指令，用以标记此 XA 事务边界。接着会依次发送 prepare 指令，收集所有参与 XAResource 投票。若所有 XAResource 的反馈结果均为正确，则调用 commit 指令进行最终提交；若有任意 XAResource 的反馈结果不正确，则调用 rollback 指令进行回滚。在事务管理器发出提交指令后，任何 XAResource 产生的异常都会通过恢复日志进行重试，以保证提交阶段的操作原子性，和数据强一致性。

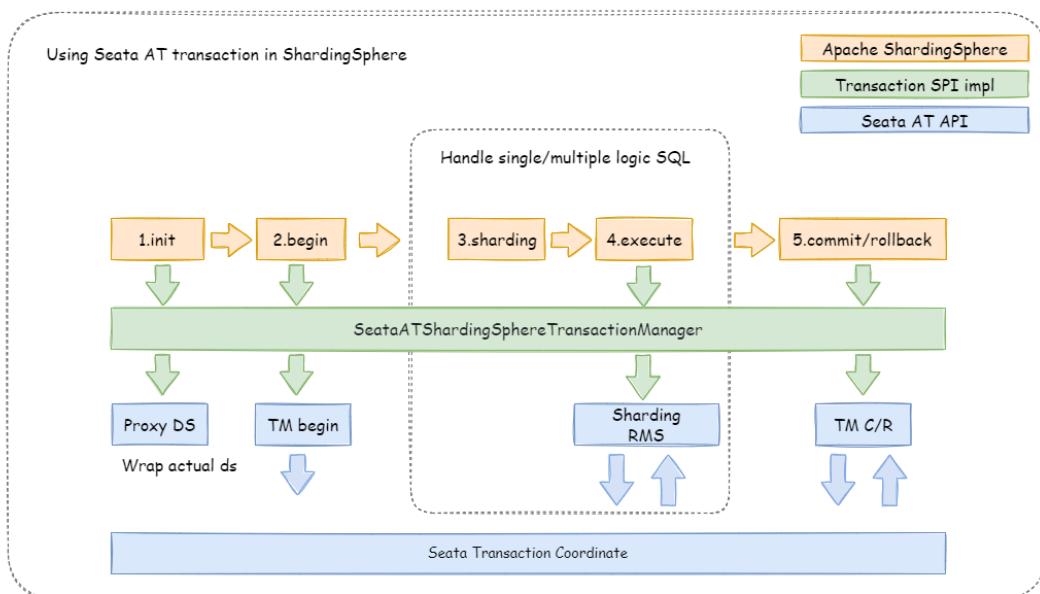
例如：

```
XAResource1.prepare      ## ack: yes
XAResource2.prepare      ## ack: yes
XAResource1.commit
XAResource2.commit

XAResource1.prepare      ## ack: yes
XAResource2.prepare      ## ack: no
XAResource1.rollback
XAResource2.rollback
```

7.5.3 Seata 柔性事务

整合 Seata AT 事务时，需要将 TM，RM 和 TC 的模型融入 Apache ShardingSphere 的分布式事务生态中。在数据库资源上，Seata 通过对接 DataSource 接口，让 JDBC 操作可以同 TC 进行远程通信。同样，Apache ShardingSphere 也是面向 DataSource 接口，对用户配置的数据源进行聚合。因此，将 DataSource 封装为基于 Seata 的 DataSource 后，就可以将 Seata AT 事务融入到 Apache ShardingSphere 的分片生态中。



引擎初始化

包含 Seata 柔性事务的应用启动时，用户配置的数据源会根据 `seata.conf` 的配置，适配为 Seata 事务所需的 `DataSourceProxy`，并且注册至 RM 中。

开启全局事务

TM 控制全局事务的边界，TM 通过向 TC 发送 `Begin` 指令，获取全局事务 ID，所有分支事务通过此全局事务 ID，参与到全局事务中；全局事务 ID 的上下文存放在当前线程变量中。

执行真实分片 SQL

处于 Seata 全局事务中的分片 SQL 通过 RM 生成 undo 快照，并且发送 `participate` 指令至 TC，加入到全局事务中。由于 Apache ShardingSphere 的分片物理 SQL 采取多线程方式执行，因此整合 Seata AT 事务时，需要在主线程和子线程间进行全局事务 ID 的上下文传递。

提交或回滚事务

提交 Seata 事务时，TM 会向 TC 发送全局事务的提交或回滚指令，TC 根据全局事务 ID 协调所有分支事务进行提交或回滚。

7.6 数据迁移

7.6.1 原理说明

目前的数据迁移解决方案为：使用一个全新的数据库集群作为迁移目标库。

这种实现方式有以下优点：

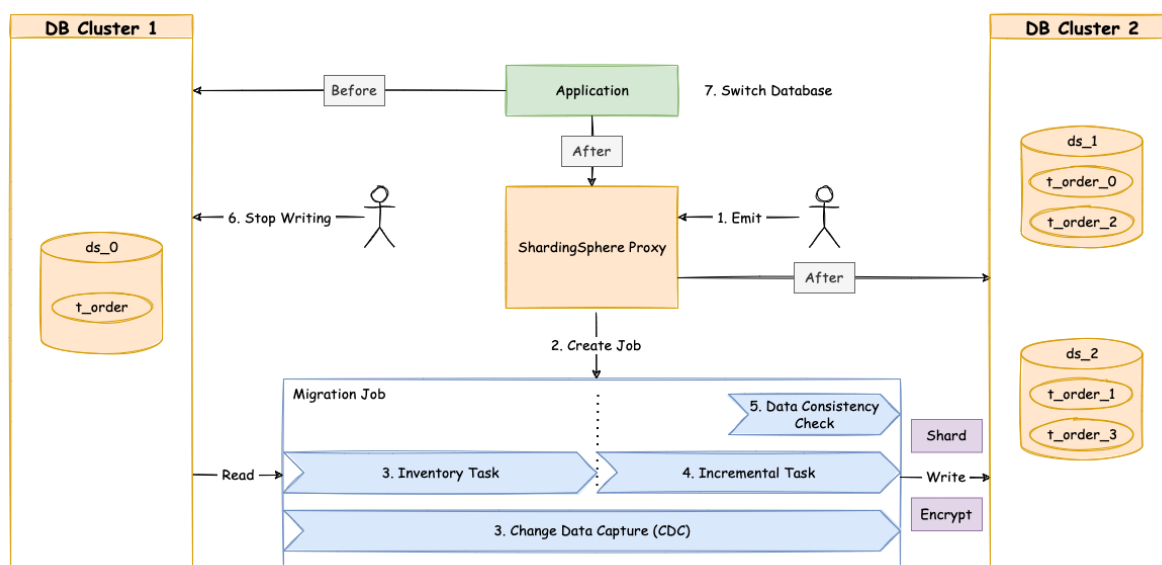
1. 迁移过程中，原始数据没有任何影响；
2. 迁移失败无风险；
3. 不受分片策略限制。

同时也存在一定的缺点：

1. 在一定时间内存在冗余服务器；
2. 所有数据都需要移动。

一次数据迁移包括以下几个主要阶段：

1. 准备阶段；
2. 存量数据迁移阶段；
3. 增量数据同步阶段；
4. 流量切换阶段。



7.6.2 执行阶段说明

准备阶段

在准备阶段，数据迁移模块会进行数据源连通性及权限的校验，同时进行存量数据的统计、日志位点的记录，最后根据数据量和用户设置的并行度，对任务进行分片。

存量数据迁移阶段

执行在准备阶段拆分好的存量数据迁移任务，存量迁移阶段采用 JDBC 查询的方式，直接从源端读取数据，基于配置的分片等规则写入到目标端。

增量数据同步阶段

由于存量数据迁移耗费的时间受到数据量和并行度等因素影响，此时需要对这段时间内业务新增的数据进行同步。不同的数据库使用的技术细节不同，但总体上均为基于复制协议或 WAL 日志实现的变更数据捕获功能。

- MySQL：订阅并解析 binlog；
- PostgreSQL：采用官方逻辑复制 `test_decoding`。

这些捕获的增量数据，同样会由数据迁移模块写入到新数据节点中。当增量数据基本同步完成时（由于业务系统未停止，增量数据是不断的），则进入流量切换阶段。

流量切换阶段

在此阶段，可能存在一定时间的业务只读窗口期，通过设置数据库只读、控制源头写流量等方式，让源端数据节点中的数据短暂静态，确保增量同步完全完成。

这个只读窗口期时长取决于用户是否需要对数据进行一致性校验以及数据量。确认完成后，数据迁移完成。然后用户可以把读流量或者写流量切换到 Apache ShardingSphere。

7.6.3 相关参考

数据迁移的配置

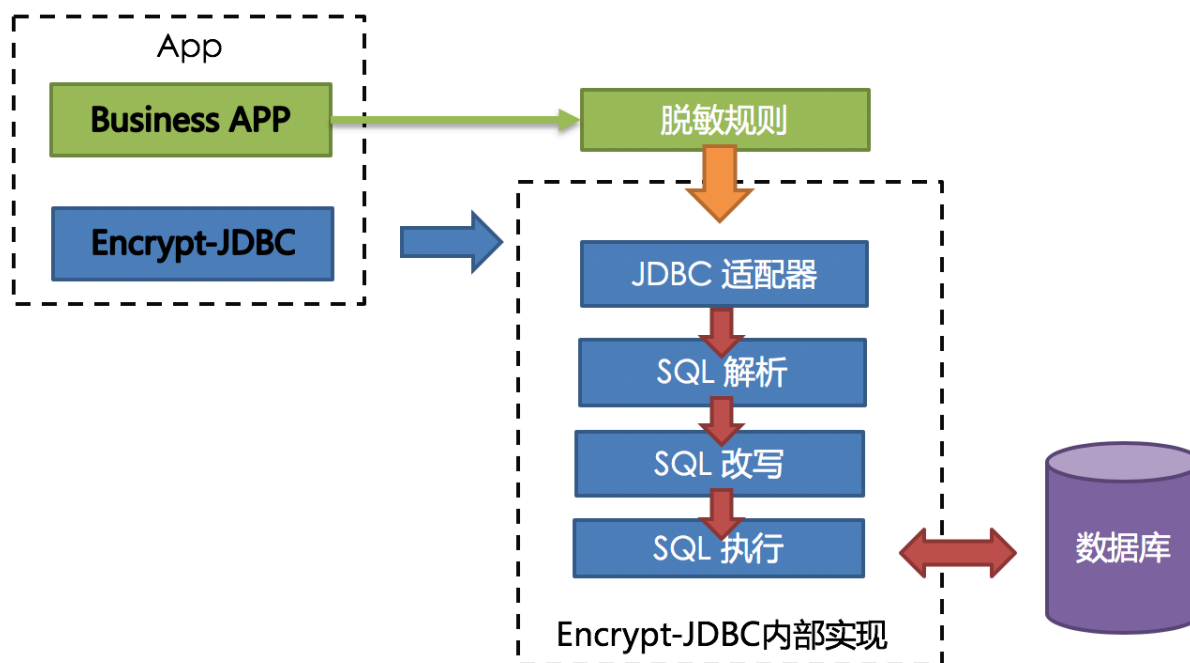
7.7 数据加密

7.7.1 处理流程详解

Apache ShardingSphere 通过对用户输入的 SQL 进行解析，并依据用户提供的加密规则对 SQL 进行改写，从而实现对原文数据进行加密，并将原文数据（可选）及密文数据同时存储到底层数据库。在用户查询数据时，它仅从数据库中取出密文数据，并对其解密，最终将解密后的原始数据返回给用户。Apache ShardingSphere 自动化 & 透明化了数据加密过程，让用户无需关注数据加密的实现细节，像使用普通数

据那样使用加密数据。此外，无论是已在线业务进行加密改造，还是新上线业务使用加密功能，Apache ShardingSphere 都可以提供一套相对完善的解决方案。

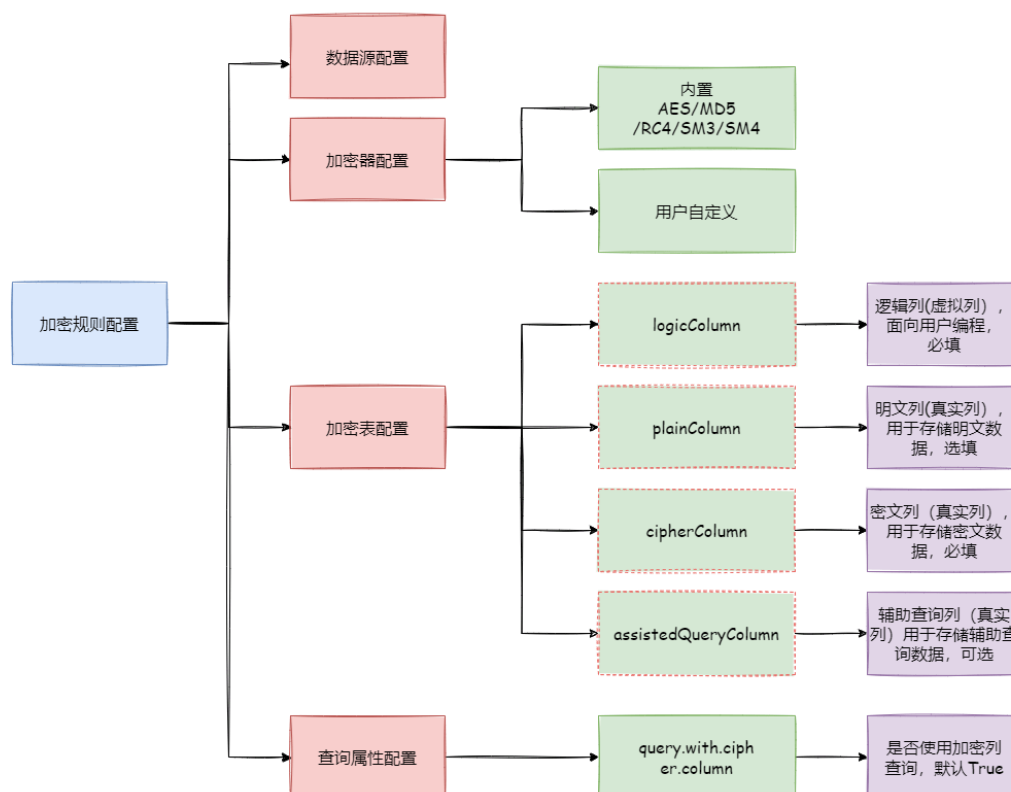
整体架构



加密模块将用户发起的 SQL 进行拦截，并通过 SQL 语法解析器进行解析、理解 SQL 行为，再依据用户传入的加密规则，找出需要加密的字段和所使用的加解密算法对目标字段进行加解密处理后，再与底层数据库进行交互。Apache ShardingSphere 会将用户请求的明文进行加密后存储到底层数据库；并在用户查询时，将密文从数据库中取出进行解密后返回给终端用户。通过屏蔽对数据的加密处理，使用户无需感知解析 SQL、数据加密、数据解密的处理过程，就像在使用普通数据一样使用加密数据。

加密规则

在详解整套流程之前，我们需要先了解下加密规则与配置，这是认识整套流程的基础。加密配置主要分为四部分：数据源配置，加密算法配置，加密表配置以及查询属性配置，其详情如下图所示：



数据源配置：指数据源配置。

加密算法配置：指使用什么加密算法进行加解密。目前 ShardingSphere 内置了五种加解密算法：AES, MD5, RC4, SM3 和 SM4。用户还可以通过实现 ShardingSphere 提供的接口，自行实现一套加解密算法。

加密表配置：用于告诉 ShardingSphere 数据表里哪个列用于存储密文数据（cipherColumn）、使用什么算法加解密（encryptorName）、哪个列用于存储辅助查询数据（assistedQueryColumn）、使用什么算法加解密（assistedQueryEncryptorName）、哪个列用于存储明文数据（plainColumn）以及用户想使用哪个列进行 SQL 编写（logicColumn）。

如何理解 用户想使用哪个列进行 SQL 编写（logicColumn）？

我们可以从加密模块存在的意义来理解。加密模块最终目的是希望屏蔽底层对数据的加密处理，也就是说我们不希望用户知道数据是如何被加解密的、如何将明文数据存储到 plainColumn，将密文数据存储到 cipherColumn，将辅助查询数据存储到 assistedQueryColumn。换句话说，我们不需要用户知道 plainColumn、cipherColumn 和 assistedQueryColumn 的存在和使用。所以，我们需要给用户提供一个概念意义上的列，这个列可以脱离底层数据库的真实列，它可以是数据库表里的一个真实列，也可以不是，从而使得用户可以随意改变底层数据库的 plainColumn、cipherColumn 和 assistedQueryColumn 的列名。或者删除 plainColumn，选择永远不再存储明文，只存储密文。只要用户的 SQL 面向这个逻辑列进行编写，并在加密规则里给出 logicColumn 和 plainColumn、cipherColumn、assistedQueryColumn 之间正确的映射关系即可。

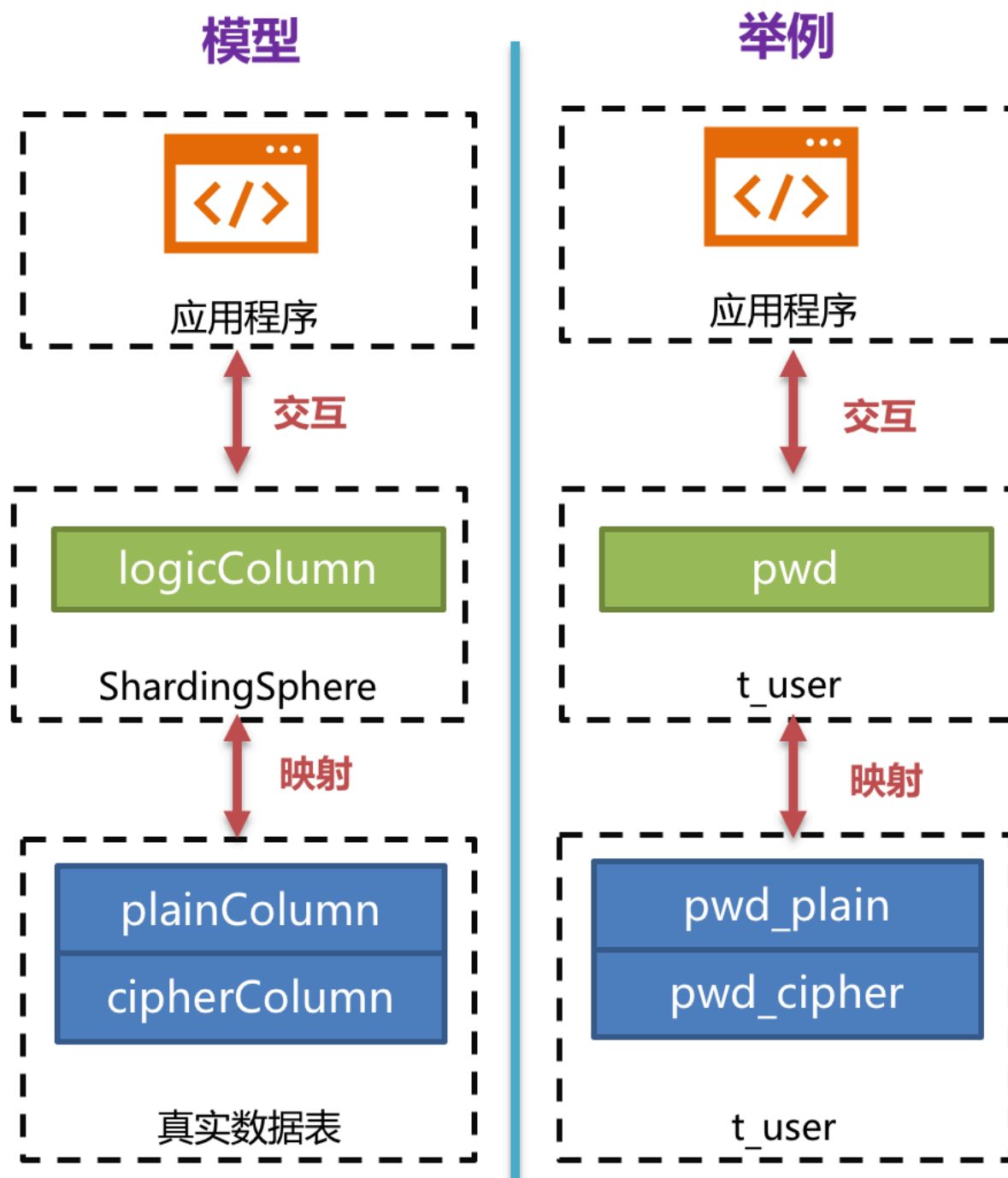
为什么要这么做呢？答案在文章后面，即为了让已上线的业务能无缝、透明、安全地进行数据加密迁移。

查询属性的配置：当底层数据库表里同时存储了明文数据、密文数据后，该属性开关用于决定是直接查

询数据库表里的明文数据进行返回，还是查询密文数据通过 Apache ShardingSphere 解密后返回。该属性开关支持表级别和整个规则级别配置，表级别优先级最高。

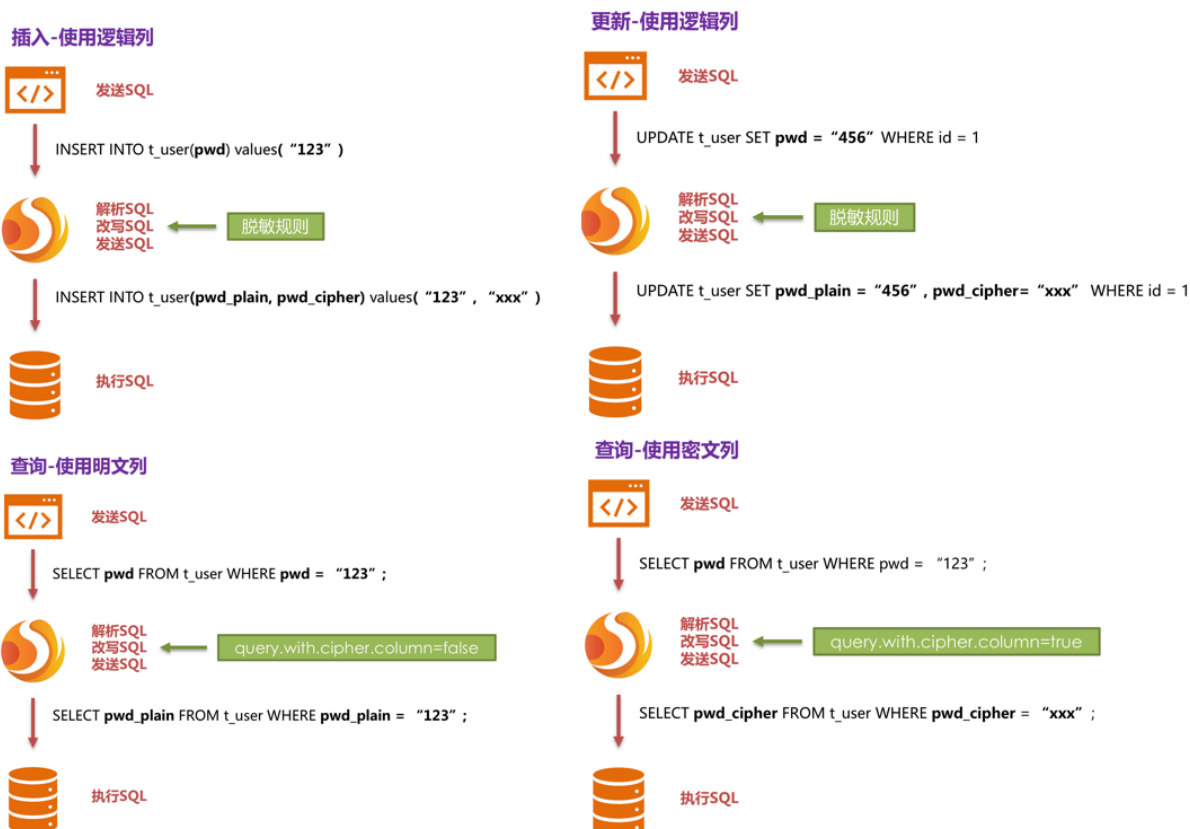
加密处理过程

举例说明，假如数据库里有一张表叫做 `t_user`，这张表里实际有两个字段 `pwd_plain`，用于存放明文数据、`pwd_cipher`，用于存放密文数据、`pwd_assisted_query`，用于存放辅助查询数据，同时定义 `logicColumn` 为 `pwd`。那么，用户在编写 SQL 时应该面向 `logicColumn` 进行编写，即 `INSERT INTO t_user SET pwd = '123'`。Apache ShardingSphere 接收到该 SQL，通过用户提供的加密配置，发现 `pwd` 是 `logicColumn`，于是便对逻辑列及其对应的明文数据进行加密处理。**Apache ShardingSphere** 将面向用户的逻辑列与面向底层数据库的明文列和密文列进行了列名以及数据的加密映射转换。如下图所示：



即依据用户提供的加密规则，将用户 SQL 与底层数据表结构割裂开来，使得用户的 SQL 编写不再依赖于真实的数据库表结构。而用户与底层数据库之间的衔接、映射、转换交由 Apache ShardingSphere 进行处理。

下方图片展示了使用加密模块进行增删改查时，其中的处理流程和转换逻辑，如下图所示。



7.7.2 解决方案详解

在了解了 Apache ShardingSphere 加密处理流程后，即可将加密配置、加密处理流程与实际场景进行结合。所有的设计开发都是为了解决业务场景遇到的痛点。那么面对之前提到的业务场景需求，又应该如何使用 Apache ShardingSphere 这把利器来满足业务需求呢？

新上线业务

业务场景分析：新上线业务由于一切从零开始，不存在历史数据清洗问题，所以相对简单。

解决方案说明：选择合适的加密算法，如 AES 后，只需配置逻辑列（面向用户编写 SQL）和密文列（数据表存密文数据）即可，逻辑列和密文列可以相同也可以不同。建议配置如下（YAML 格式展示）：

```

-!ENCRYPT
encryptors:
  aes_encryptor:
    type: AES
    props:
      aes-key-value: 123456abc
tables:
  t_user:
    columns:
      pwd:
        cipherColumn: pwd_cipher

```

```

encryptorName: aes_encryptor
assistedQueryColumn: pwd_assisted_query
assistedQueryEncryptorName: pwd_assisted_query_cipher
queryWithCipherColumn: true

```

使用这套配置，Apache ShardingSphere 只需将 logicColumn 和 cipherColumn，assistedQueryColumn 进行转换，底层数据表不存储明文，只存储了密文，这也是安全审计部分的要求所在。如果用户希望将明文、密文一同存储到数据库，只需添加 plainColumn 配置即可。整体处理流程如下图所示：

新上线业务



已上线业务改造

业务场景分析：由于业务已经在线上运行，数据库里必然存有大量明文历史数据。现在的问题是如何让历史数据得以加密清洗、如何让增量数据得以加密处理、如何让业务在新旧两套数据系统之间进行无缝、透明化迁移。

解决方案说明：在提供解决方案之前，我们先来头脑风暴一下：首先，既然是旧业务需要进行加密改造，那一定存储了非常重要且敏感的信息。这些信息含金量高且业务相对基础重要。不应该采用停止业务禁止新数据写入，再找个加密算法把历史数据全部加密清洗，再把之前重构的代码部署上线，使其能把存量和增量数据进行在线加密解密。

那么另一种相对安全的做法是：重新搭建一套和生产环境一模一样的预发环境，然后通过相关迁移洗数工具把生产环境的存量原文数据加密后存储到预发环境，而新增数据则通过例如 MySQL 主从复制及业务方自行开发的工具加密后存储到预发环境的数据库里，再把重构后可以进行加解密的代码部署到预发环境。这样生产环境是一套以明文为核心的查询修改的环境；预发环境是一套以密文为核心加解密查询修改的环境。在对比一段时间无误后，可以夜间操作将生产流量切到预发环境中。此方案相对安全可靠，只是时间、人力、资金、成本较高，主要包括：预发环境搭建、生产代码整改、相关辅助工具开发等。

业务开发人员最希望的做法是：减少资金费用的承担、最好不要修改业务代码、能够安全平滑迁移系统。于是，ShardingSphere 的加密功能模块便应运而生。可分为 3 步进行：

1. 系统迁移前

假设系统需要对 t_user 的 pwd 字段进行加密处理，业务方使用 Apache ShardingSphere 来代替标准化的 JDBC 接口，此举基本不需要额外改造（我们还提供了 Spring Boot Starter，Spring 命名空间，YAML 等接入方式，满足不同业务方需求）。另外，提供一套加密配置规则，如下所示：

```

-!ENCRYPT
encryptors:
  aes_encryptor:
    type: AES
    props:
      aes-key-value: 123456abc
tables:
  t_user:
    columns:
      pwd:
        plainColumn: pwd
        cipherColumn: pwd_cipher
        encryptorName: aes_encryptor
        assistedQueryColumn: pwd_assisted_query
        assistedQueryEncryptorName: pwd_assisted_query_cipher
        queryWithCipherColumn: false

```

依据上述加密规则可知, 首先需要在数据库表 `t_user` 里新增一个字段叫做 `pwd_cipher`, 即 `cipherColumn`, 用于存放密文数据, 同时我们把 `plainColumn` 设置为 `pwd`, 用于存放明文数据, 而把 `logicColumn` 也设置为 `pwd`。由于之前的代码 SQL 就是使用 `pwd` 进行编写, 即面向逻辑列进行 SQL 编写, 所以业务代码无需改动。通过 Apache ShardingSphere, 针对新增的数据, 会把明文写到 `pwd` 列, 并同时把明文进行加密存储到 `pwd_cipher` 列。此时, 由于 `queryWithCipherColumn` 设置为 `false`, 对业务应用来说, 依旧使用 `pwd` 这一明文列进行查询存储, 却在底层数据库表 `pwd_cipher` 上额外存储了新增数据的密文数据, 其处理流程如下图所示:

已上线业务改造-迁移前



新增数据在插入时, 就通过 Apache ShardingSphere 加密为密文数据, 并被存储到了 `cipherColumn`。而现在就需要处理历史明文存量数据。由于 Apache ShardingSphere 目前并未提供相关迁移洗数工具, 此时需要业务方自行将“`pwd`”中的明文数据进行加密处理存储到“`pwd_cipher`”。

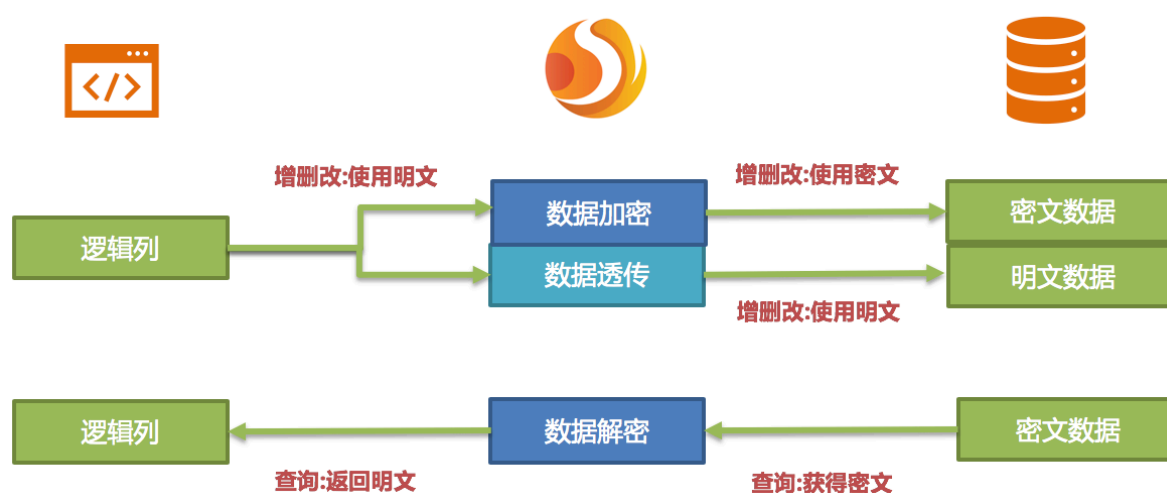
2. 系统迁移中

新增的数据已被 Apache ShardingSphere 将密文存储到密文列, 明文存储到明文列; 历史数据被业务方自行加密清洗后, 将密文也存储到密文列。也就是说现在的数据库里即存放着明文也存放着密文, 只是

由于配置项中的 `queryWithCipherColumn = false`，所以密文一直没有被使用过。现在我们为了让系统能切到密文数据进行查询，需要将加密配置中的 `queryWithCipherColumn` 设置为 `true`。在重启系统后，我们发现系统业务一切正常，但是 Apache ShardingSphere 已经开始从数据库里取出密文列的数据，解密后返回给用户；而对于用户的增删改需求，则依旧会把原文数据存储到明文列，加密后密文数据存储到密文列。

虽然现在业务系统通过将密文列的数据取出，解密后返回；但是，在存储的时候仍旧会存一份原文数据到明文列，这是为什么呢？答案是：为了能够进行系统回滚。因为只要密文和明文永远同时存在，我们就可以通过开关项配置自由将业务查询切换到 `cipherColumn` 或 `plainColumn`。也就是说，如果将系统切到密文列进行查询时，发现系统报错，需要回滚。那么只需将 `queryWithCipherColumn = false`，Apache ShardingSphere 将会还原，即又重新开始使用 `plainColumn` 进行查询。处理流程如下图所示：

已上线业务改造-迁移中



3. 系统迁移后

由于安全审计部门要求，业务系统一般不可能让数据库的明文列和密文列永久同步保留，我们需要在系统稳定后将明文列数据删除。即我们需要在系统迁移后将 `plainColumn`，即 `pwd` 进行删除。那问题来了，现在业务代码都是面向 `pwd` 进行编写 SQL 的，把底层数据表中的存放明文的 `pwd` 删除了，换用 `pwd_cipher` 进行解密得到原文数据，那岂不是意味着业务方需要整改所有 SQL，从而不使用即将要被删除的 `pwd` 列？还记得我们 Apache ShardingSphere 的核心意义所在吗？

这也正是 Apache ShardingSphere 核心意义所在，即依据用户提供的加密规则，将用户 SQL 与底层数据库表结构割裂开来，使得用户的 SQL 编写不再依赖于真实的数据库表结构。而用户与底层数据库之间的衔接、映射、转换交由 Apache ShardingSphere 进行处理。

是的，因为有 `logicColumn` 存在，用户的编写 SQL 都面向这个虚拟列，Apache ShardingSphere 就可以把这个逻辑列和底层数据表中的密文列进行映射转换。于是迁移后的加密配置即为：

```

-!ENCRYPT
encryptors:
  aes_encryptor:
    type: AES
    props:
      aes-key-value: 123456abc
tables:

```

```
t_user:
  columns:
    pwd: # pwd 与 pwd_cipher 的转换映射
    cipherColumn: pwd_cipher
    encryptorName: aes_encryptor
    assistedQueryColumn: pwd_assisted_query
    assistedQueryEncryptorName: pwd_assisted_query_cipher
    queryWithCipherColumn: true
```

其处理流程如下：

已上线业务改造-迁移后



4. 系统迁移完成

安全审计部门再要求，业务系统需要定期或某些紧急安全事件触发修改密钥，我们需要再次进行迁移洗数，即使用旧密钥解密后再使用新密钥加密。既要又要还要的问题来了，明文列数据已删除，数据库表中数据量千万级，迁移洗数需要一定时间，迁移洗数过程中密文列在变化，系统还需正确提供服务。怎么办？答案是：辅助查询列 因为辅助查询列一般使用不可逆的 MD5 和 SM3 等算法，基于辅助列进行查询，即使在迁移洗数过程中，系统也是可以提供正确服务。

至此，已在线业务加密整改解决方案全部叙述完毕。我们提供了 Java、YAML、Spring Boot Starter、Spring 命名空间多种方式供用户选择接入，力求满足业务不同的接入需求。该解决方案目前已在京东数科不断落地上线，提供对内基础服务支撑。

7.7.3 中间件加密服务优势

1. 自动化 & 透明化数据加密过程，用户无需关注加密中间实现细节。
2. 提供多种内置、第三方（AKS）的加密算法，用户仅需简单配置即可使用。
3. 提供加密算法 API 接口，用户可实现接口，从而使用自定义加密算法进行数据加密。
4. 支持切换不同的加密算法。
5. 针对已上线业务，可实现明文数据与密文数据同步存储，并通过配置决定使用明文列还是密文列进行查询。可实现在不改变业务查询 SQL 前提下，已上线系统对加密前后数据进行安全、透明化迁移。

7.7.4 加密算法解析

Apache ShardingSphere 提供了加密算法用于数据加密，即 `EncryptAlgorithm`。

一方面，Apache ShardingSphere 为用户提供了内置的加解密实现类，用户只需进行配置即可使用；另一方面，为了满足用户不同场景的需求，我们还开放了相关加解密接口，用户可依据这两种类型的接口提供具体实现类。再进行简单配置，即可让 Apache ShardingSphere 调用用户自定义的加解密方案进行数据加密。

EncryptAlgorithm

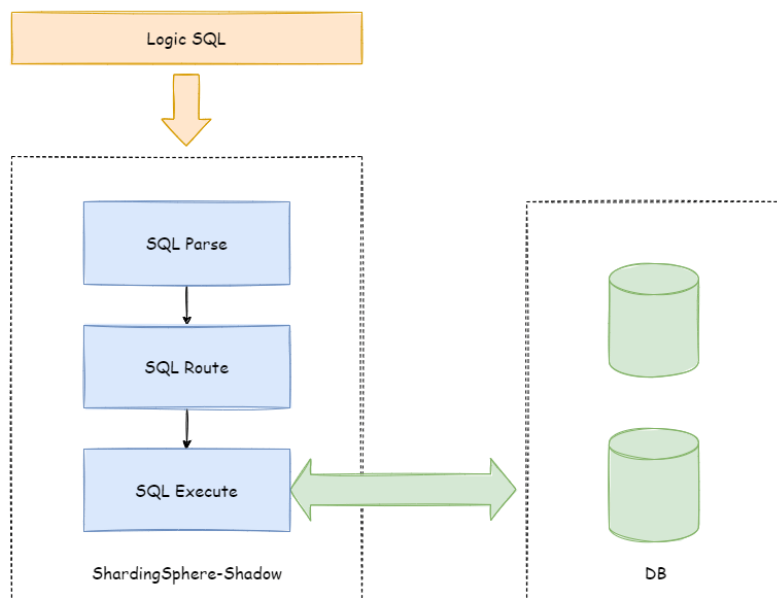
该解决方案通过提供 `encrypt()`，`decrypt()` 两种方法对需要加密的数据进行加解密。在用户进行 INSERT，DELETE，UPDATE 时，ShardingSphere 会按照用户配置，对 SQL 进行解析、改写、路由，并调用 `encrypt()` 将数据加密后存储到数据库，而在 SELECT 时，则调用 `decrypt()` 方法将从数据库中取出的加密数据进行逆向解密，最终将原始数据返回给用户。

当前，Apache ShardingSphere 针对这种类型的加密解决方案提供了五种具体实现类，分别是 MD5（不可逆），AES（可逆），RC4（可逆），SM3（不可逆），SM4（可逆），用户只需配置即可使用这五种内置的方案。

7.8 影子库

7.8.1 原理介绍

Apache ShardingSphere 通过解析 SQL，对传入的 SQL 进行影子判定，根据配置文件中用户设置的影子规则，路由到生产库或者影子库。



以 INSERT 语句为例，在写入数据时，Apache ShardingSphere 会对 SQL 进行解析，再根据配置文件中的规则，构造一条路由链。在当前版本的功能中，影子功能处于路由链中的最后一个执行单元，即，如果有其他需要路由的规则存在，如分片，Apache ShardingSphere 会首先根据分片规则，路由到某一个数据库，再执行影子路由判定流程，判定执行 SQL 满足影子规则的配置，数据路由到与之对应的影子库，生产数据则维持不变。

DML 语句

支持两种算法。影子判定会首先判断执行 SQL 相关表与配置的影子表是否有交集。如果有交集，依次判定交集部分影子表关联的影子算法，有任何一个判定成功，SQL 语句路由到影子库。影子表没有交集或者影子算法判定不成功，SQL 语句路由到生产库。

DDL 语句

仅支持注解影子算法。在压测场景下，DDL 语句一般不需要测试。主要在初始化或者修改影子库中影子表时使用。影子判定会首先判断执行 SQL 是否包含注解。如果包含注解，影子规则中配置的 HINT 影子算法依次判定。有任何一个判定成功，SQL 语句路由到影子库。执行 SQL 不包含注解或者 HINT 影子算法判定不成功，SQL 语句路由到生产库。

7.8.2 相关参考

JAVA API：影子库配置

YAML 配置：影子库配置

Spring Boot Starter：影子库配置

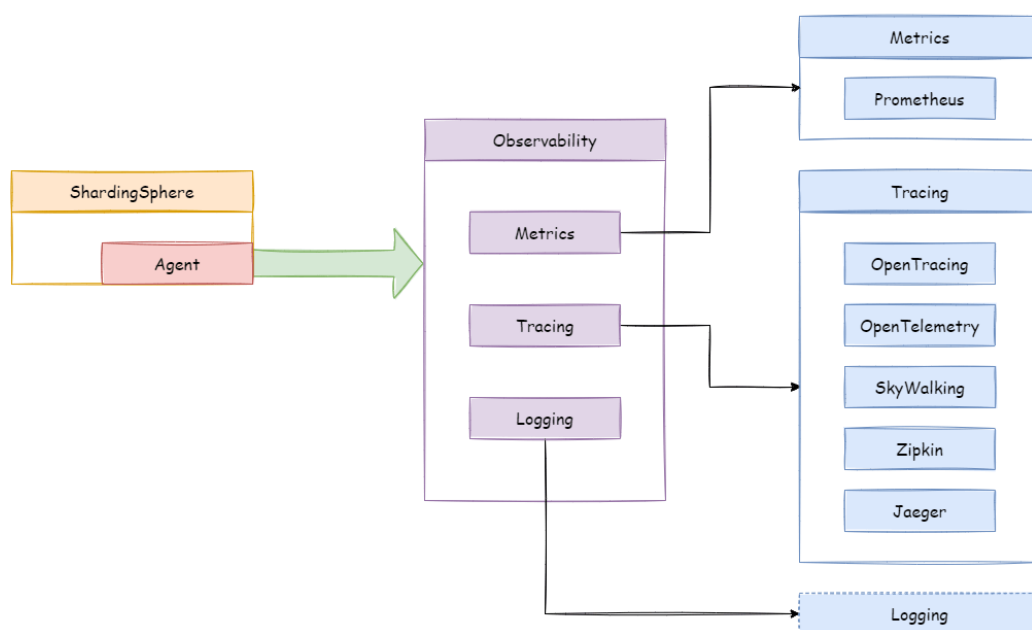
Spring 命名空间：影子库配置

7.9 可观察性

7.9.1 原理说明

ShardingSphere-Agent 模块为 ShardingSphere 提供了可观察性的框架，它是基于 Java Agent 技术实现的。

Metrics、Tracing 和 Logging 等功能均通过插件的方式集成在 Agent 中，如图：



- Metrics 插件用于收集和展示整个集群的统计指标。Apache ShardingSphere 默认提供了对 Prometheus 的支持。
- Tracing 插件用于获取 SQL 解析与 SQL 执行的链路跟踪信息。Apache ShardingSphere 默认提供了对 Jaeger、OpenTelemetry、OpenTracing (SkyWalking) 和 Zipkin 的支持，也支持用户通过插件化的方式开发自定义的 Tracing 组件。

- 默认的 Logging 插件展示了如何在 ShardingSphere 中记录额外的日志，实际应用中需要用户根据自己的需求进行探索。

7.10 DistSQL

本章节将介绍 DistSQL 的详细语法。

7.10.1 语法

本章节将对 DistSQL 的语法进行详细说明，并以实际的例子介绍 DistSQL 的使用。

RDL 语法

RDL (Resource & Rule Definition Language) 为 Apache ShardingSphere 的资源 and 规则定义语言。

资源定义

本章节将对资源定义的语法进行详细说明。

ADD RESOURCE

描述

ADD RESOURCE 语法用于为当前所选逻辑库 (DATABASE) 添加资源。

语法

```
AddResource ::=
  'ADD' 'RESOURCE' resourceDefinition (',' resourceDefinition)*

resourceDefinition ::=
  resourceName '(' ( 'HOST' '=' hostName ',' 'PORT' '=' port ',' 'DB' '=' dbName |
  'URL' '=' url ) ',' 'USER' '=' user (',' 'PASSWORD' '=' password )? (','
  proerties)? ')'

resourceName ::=
  identifier

hostname ::=
  string

port ::=
  int
```

```
dbName ::=
    string

url ::=
    string

user ::=
    string

password ::=
    string

properties ::=
    PROPERTIES '(' property ( ',' property )* ')'

property ::=
    key '=' value

key ::=
    string

value ::=
    string
```

特别说明

- 添加资源前请确认已经在 Proxy 中创建逻辑数据库，并执行 use 命令成功选择一个逻辑数据库；
- 确认添加的资源是可以正常连接的，否则将不能添加成功；
- resourceName 区分大小写；
- resourceName 在当前逻辑库中需要唯一；
- resourceName 命名只允许使用字母、数字以及 _，且必须以字母开头；
- poolProperty 用于自定义连接池参数，key 必须和连接池参数名一致，value 支持 int 和 String 类型；
- 当 password 包含特殊字符时，建议使用 string 形式；例如 password@123 的 string 形式为 "password@123"。

示例

- 使用标准模式添加资源

```
ADD RESOURCE ds_0 (  
  HOST=127.0.0.1,  
  PORT=3306,  
  DB=db_0,  
  USER=root,  
  PASSWORD=root  
);
```

- 使用标准模式添加资源并设置连接池参数

```
ADD RESOURCE ds_1 (  
  HOST=127.0.0.1,  
  PORT=3306,  
  DB=db_1,  
  USER=root,  
  PASSWORD=root,  
  PROPERTIES("maximumPoolSize"=10)  
);
```

- 使用 URL 模式添加资源并设置连接池参数

```
ADD RESOURCE ds_2 (  
  URL="jdbc:mysql://127.0.0.1:3306/db_2?serverTimezone=UTC&useSSL=false",  
  USER=root,  
  PASSWORD=root,  
  PROPERTIES("maximumPoolSize"=10,"idleTimeout"="30000")  
);
```

保留字

ADD、RESOURCE、HOST、PORT、DB、USER、PASSWORD、PROPERTIES、URL

相关链接

- [保留字](#)

ALTER RESOURCE

描述

ALTER RESOURCE 语法用于修改当前所选逻辑库（DATABASE）的资源。

语法

```
AlterResource ::=
    'ALTER' 'RESOURCE' resourceDefinition (',' resourceDefinition)*

resourceDefinition ::=
    resourceName '(' ( 'HOST' '=' hostName ',' 'PORT' '=' port ',' 'DB' '=' dbName |
    'URL' '=' url ) ',' 'USER' '=' user (',' 'PASSWORD' '=' password )? (','
    proerties)? ')'

resourceName ::=
    identifier

hostname ::=
    string

port ::=
    int

dbName ::=
    string

url ::=
    string

user ::=
    string

password ::=
    string

proerties ::=
    PROPERTIES '(' property ( ',' property )* ')'

property ::=
    key '=' value
```

```
key ::=
    string

value ::=
    string
```

补充说明

- 修改资源前请确认已经在 Proxy 中创建逻辑数据库，并执行 use 命令成功选择一个逻辑数据库；
- ALTER RESOURCE 不允许改变该资源关联的真实数据源；
- ALTER RESOURCE 会发生连接池的切换，这个操作可能对进行中的业务造成影响，请谨慎使用；
- 确认添加的资源是可以正常连接的，否则将不能添加成功；
- resourceName 区分大小写；
- resourceName 在当前逻辑库中需要唯一；
- resourceName 命名只允许使用字母、数字以及 _，且必须以字母开头；
- poolProperty 用于自定义连接池参数，key 必须和连接池参数名一致，value 支持 int 和 String 类型；
- 当 password 包含特殊字符时，建议使用 string 形式；例如 password@123 的 string 形式为 "password@123"。

示例

- 使用标准模式修改资源

```
ALTER RESOURCE ds_0 (
    HOST=127.0.0.1,
    PORT=3306,
    DB=db_0,
    USER=root,
    PASSWORD=root
);
```

- 使用标准模式修改资源并设置连接池参数

```
ALTER RESOURCE ds_1 (
    HOST=127.0.0.1,
    PORT=3306,
    DB=db_1,
    USER=root,
    PASSWORD=root,
```

```
PROPERTIES("maximumPoolSize"=10)
);
```

- 使用 URL 模式修改资源并设置连接池参数

```
ALTER RESOURCE ds_2 (
    URL="jdbc:mysql://127.0.0.1:3306/db_2?serverTimezone=UTC&useSSL=false",
    USER=root,
    PASSWORD=root,
    PROPERTIES("maximumPoolSize"=10,"idleTimeout"="30000")
);
```

保留字

ALTER、RESOURCE、HOST、PORT、DB、USER、PASSWORD、PROPERTIES、URL

相关链接

- [保留字](#)

DROP RESOURCE

描述

DROP RESOURCE 语法用于从当前逻辑库中移除资源。

语法

```
DropResource ::=
    'DROP' 'RESOURCE' ( 'IF' 'EXISTS' )? resourceName ( ',' resourceName )* (
    'IGNORE' 'SINGLE' 'TABLES' )?

resourceName ::=
    identifier
```

补充说明

- DROP RESOURCE 只会移除 Proxy 中的资源，不会删除与资源对应的真实数据源；
- 无法移除已经被规则使用的资源。移除被规则使用的资源时会提示 Resources are still in used；
- 将要移除的资源中仅包含 SINGLE TABLE RULE，且用户确认可以忽略该限制时，可添加 IGNORE SINGLE TABLES 关键字移除资源。

示例

- 移除资源

```
DROP RESOURCE ds_0;
```

- 移除多个资源

```
DROP RESOURCE ds_1, ds_2;
```

- 忽略单表移除资源

```
DROP RESOURCE ds_3 IGNORE SINGLE TABLES;
```

- 如果资源存在则移除

```
DROP RESOURCE IF EXISTS ds_4;
```

保留字

DROP、RESOURCE、IF、EXISTS、IGNORE、SINGLE、TABLES

相关链接

- [保留字](#)

规则定义

本章节将对规则定义的语法进行详细说明。

数据库发现

本章节将对数据库发现特性的语法进行详细说明。

CREATE DB_DISCOVERY RULE

描述

CREATE DB_DISCOVERY RULE 语法用于创建数据库发现规则

语法定义

```
CreateDatabaseDiscoveryRule ::=
    'CREATE' 'DB_DISCOVERY' 'RULE' ( databaseDiscoveryDefinition |
databaseDiscoveryConstruction ) ( ',' ( databaseDiscoveryDefinition |
databaseDiscoveryConstruction ) ) *

databaseDiscoveryDefinition ::=
    ruleName '(' 'RESOURCES' '(' resourceName ( ',' resourceName ) * ')' ',' 'TYPE'
'(' 'NAME' '=' typeName ( ',' 'PROPERTIES' 'key' '=' 'value' ( ',' 'key' '=' 'value'
' ) * ) ? ',' 'HEARTBEAT' '(' 'key' '=' 'value' ( ',' 'key' '=' 'value' ) * ')' ')'

databaseDiscoveryConstruction ::=
    ruleName '(' 'RESOURCES' '(' resourceName ( ',' resourceName ) * ')' ',' 'TYPE'
'=' discoveryTypeName ',' 'HEARTBEAT' '=' discoveryHeartbeatName ')'

ruleName ::=
    identifier

resourceName ::=
    identifier

typeName ::=
    identifier

discoveryHeartbeatName ::=
    identifier
```

补充说明

- `discoveryType` 指定数据库发现服务类型，ShardingSphere 内置支持 `MySQL.MGR`；
- 重复的 `ruleName` 将无法被创建；
- 正在被使用的 `discoveryType` 和 `discoveryHeartbeat` 无法被删除；
- 带有 - 的命名在改动时需要使用 " "。

示例

创建 `discoveryRule` 时同时创建 `discoveryType` 和 `discoveryHeartbeat`

```
CREATE DB_DISCOVERY RULE db_discovery_group_0 (
  RESOURCES(ds_0, ds_1, ds_2),
  TYPE(NAME='MySQL.MGR',PROPERTIES('group-name'='92504d5b-6dec')),
  HEARTBEAT(PROPERTIES('keep-alive-cron'='0/5 * * * * ?'))
);
```

使用已有的 `discoveryType` 和 `discoveryHeartbeat` 创建 `discoveryRule`

```
CREATE DB_DISCOVERY RULE db_discovery_group_1 (
  RESOURCES(ds_0, ds_1, ds_2),
  TYPE=db_discovery_group_1_mgr,
  HEARTBEAT=db_discovery_group_1_heartbeat
);
```

保留字

CREATE、DB_DISCOVERY、RULE、RESOURCES、TYPE、NAME、PROPERTIES、HEARTBEAT

相关链接

- [保留字](#)

数据加密

本章节将对数据加密特性的语法进行详细说明。

CREATE ENCRYPT RULE

描述

CREATE ENCRYPT RULE 语法用于创建数据加密规则

语法定义

```

CreateEncryptRule ::=
    'CREATE' 'ENCRYPT' 'RULE' encryptDefinition ( ',' encryptDefinition )*

encryptDefinition ::=
    tableName '(' 'COLUMNS' '(' columnDefinition ( ',' columnDefinition )* ')' ','
    'QUERY_WITH_CIPHER_COLUMN' '=' queryWithCipherColumn ')'

columnDefinition ::=
    'NAME' '=' columnName ',' ( 'PLAIN' '=' plainColumnName )? 'CIPHER' '='
    cipherColumnName ',' 'TYPE' '(' 'NAME' '=' encryptAlgorithmType ( ',' 'PROPERTIES'
    '(' 'key' '=' 'value' ( ',' 'key' '=' 'value' )* ')' )? ')'

tableName ::=
    identifier

queryWithCipherColumn ::=
    identifier

columnName ::=
    identifier

plainColumnName ::=
    identifier

cipherColumnName ::=
    identifier

encryptAlgorithmType ::=
    identifier

```

补充说明

- PLAIN 指定明文数据列，CIPHER 指定密文数据列；
- encryptAlgorithmType 指定加密算法类型，请参考加密算法；
- 重复的 tableName 将无法被创建；
- queryWithCipherColumn 支持大写或小写的 true 或 false。

示例

创建数据加密规则

```
CREATE ENCRYPT RULE t_encrypt (
  COLUMNS(
    (NAME=user_id,PLAIN=user_plain,CIPHER=user_cipher,TYPE(NAME='AES',PROPERTIES('aes-
key-value'='123456abc'))),
    (NAME=order_id, CIPHER =order_cipher,TYPE(NAME='MD5'))
  ),QUERY_WITH_CIPHER_COLUMN=true),
t_encrypt_2 (
  COLUMNS(
    (NAME=user_id,PLAIN=user_plain,CIPHER=user_cipher,TYPE(NAME='AES',PROPERTIES('aes-
key-value'='123456abc'))),
    (NAME=order_id, CIPHER=order_cipher,TYPE(NAME='MD5'))
  ), QUERY_WITH_CIPHER_COLUMN=FALSE);
```

保留字

CREATE、ENCRYPT、RULE、COLUMNS、NAME、CIPHER、PLAIN、QUERY_WITH_CIPHER_COLUMN、TYPE、TRUE、FALSE

相关链接

- [保留字](#)

读写分离

本章节将对读写分离特性的语法进行详细说明。

CREATE READWRITE_SPLITTING RULE

描述

CREATE DEFAULT SINGLE TABLE RULE 语法用于创建读写分离规则

语法定义

```

CreateReadwriteSplittingRule ::=
    'CREATE' 'READWRITE_SPLITTING' 'RULE' readwriteSplittingDefinition ( ','
readwriteSplittingDefinition )*

readwriteSplittingDefinition ::=
    ruleName '(' ( staticReadwriteSplittingDefinition |
dynamicReadwriteSplittingDefinition ) ( ',' loadBanlancerDefinition )? ')'

staticReadwriteSplittingDefinition ::=
    'WRITE_RESOURCE' '=' writeResourceName ',' 'READ_RESOURCES' '(' ruleName ( ','
ruleName)* ')'

dynamicReadwriteSplittingDefinition ::=
    'AUTO_AWARE_RESOURCE' '=' resourceName ( ',' 'WRITE_DATA_SOURCE_QUERY_ENABLED'
'=' ( 'TRUE' | 'FALSE' ) )?

loadBanlancerDefinition ::=
    'TYPE' '(' 'NAME' '=' loadBanlancerType ( ',' 'PROPERTIES' '(' 'key' '=' 'value'
' ( ',' 'key' '=' 'value' )* ')' )? ')'

ruleName ::=
    identifier

writeResourceName ::=
    identifier

resourceName ::=
    identifier

```

补充说明

- 支持创建静态读写分离规则和动态读写分离规则；
- 动态读写分离规则依赖于数据库发现规则；
- loadBanlancerType 指定负载均衡算法类型，请参考负载均衡算法；
- 重复的 ruleName 将无法被创建。

示例

创建静态读写分离规则

```
CREATE READWRITE_SPLITTING RULE ms_group_0 (
    WRITE_RESOURCE=write_ds,
    READ_RESOURCES(read_ds_0,read_ds_1),
    TYPE(NAME="random")
);
```

创建动态读写分离规则

```
CREATE READWRITE_SPLITTING RULE ms_group_1 (
    AUTO_AWARE_RESOURCE=group_0,
    WRITE_DATA_SOURCE_QUERY_ENABLED=false,
    TYPE(NAME="random",PROPERTIES("read_weight"="2:1"))
);
```

保留字

CREATE、READWRITE_SPLITTING、RULE、WRITE_RESOURCE、READ_RESOURCES、
AUTO_AWARE_RESOURCE、WRITE_DATA_SOURCE_QUERY_ENABLED、TYPE、NAME、PROPERTIES、
TRUE、FALSE

相关链接

- [保留字](#)

影子库压测

本章节将对影子库压测特性的语法进行详细说明。

CREATE SHADOW RULE

描述

CREATE SHADOW RULE 语法用于创建影子库压测规则

语法定义

```

CreateShadowRule ::=
    'CREATE' 'SHADOW' 'RULE' shadowDefinition ( ',' shadowDefinition )*

shadowDefinition ::=
    ruleName '(' resourceMapping shadowTableRule ( ',' shadowTableRule )* ')'

resourceMapping ::=
    'SOURCE' '=' resourceName ',' 'SHADOW' '=' resourceName

shadowTableRule ::=
    tableName '(' shadowAlgorithm ( ',' shadowAlgorithm )* ')'

shadowAlgorithm ::=
    ( algorithmName ',' )? 'TYPE' '(' 'NAME' '=' shadowAlgorithmType ','
    'PROPERTIES' '(' 'key' '=' 'value' ( ',' 'key' '=' 'value' ) ')'

ruleName ::=
    identifier

resourceName ::=
    identifier

tableName ::=
    identifier

algorithmName ::=
    identifier

shadowAlgorithmType ::=
    identifier

```

补充说明

- 重复的 ruleName 无法被创建；
- resourceMapping 指定源数据库和影子库的映射关系，需使用 RDL 管理的 resource，请参考 [数据源资源](#)；
- shadowAlgorithm 可同时作用于多个 shadowTableRule；
- algorithmName 未指定时会根据 ruleName、tableName 和 shadowAlgorithmType 自动生成；
- shadowAlgorithmType 目前支持 VALUE_MATCH、REGEX_MATCH 和 SIMPLE_HINT。

示例

创建影子库压测规则

```
CREATE SHADOW RULE shadow_rule(
  SOURCE=demo_ds,
  SHADOW=demo_ds_shadow,
  t_order((simple_hint_algorithm, TYPE(NAME="SIMPLE_HINT", PROPERTIES("shadow"=
"true", "foo"="bar"))),(TYPE(NAME="REGEX_MATCH", PROPERTIES("operation"="insert",
"column"="user_id", "regex"='[1]')))),
  t_order_item((TYPE(NAME="VALUE_MATCH", PROPERTIES("operation"="insert","column"=
"user_id", "value"='1'))))
);
```

保留字

CREATE、SHADOW、RULE、SOURCE、SHADOW、TYPE、NAME、PROPERTIES

相关链接

- [保留字](#)

分片

本章节将对分片特性的语法进行详细说明。

CREATE SHARDING TABLE RULE

描述

CREATE SHARDING TABLE RULE 语法用于为当前所选逻辑库添加分片规则

语法定义

```
CreateShardingTableRule ::=
  'CREATE' 'SHARDING' 'TABLE' 'RULE' ( tableDefinition | autoTableDefinition ) ( ',',
  ' ( tableDefinition | autoTableDefinition ) ) *

tableDefinition ::=
  tableName '(' 'DATANODES' '(' dataNode ( ',' dataNode ) * ')' ( ',' 'DATABASE_
STRATEGY' '(' strategyDefinition ')' ) ? ( ',' 'TABLE_STRATEGY' '('
strategyDefinition ')' ) ? ( ',' 'KEY_GENERATE_STRATEGY' '('
keyGenerateStrategyDefinition ')' ) ? ')' )
```

```

autoTableDefinition ::=
    tableName '(' 'RESOURCES' '(' resourceName ( ',' resourceName )* ')' ','
'SHARDING_COLUMN' '=' columnName ',' algorithmDefinition ( ',' 'KEY_GENERATE_
STRATEGY' '(' keyGenerateStrategyDefinition ')' )? ')'

strategyDefinition ::=
    'TYPE' '=' strategyType ',' ( 'SHARDING_COLUMN' | 'SHARDING_COLUMNS' ) '='
columnName ',' algorithmDefinition

keyGenerateStrategyDefinition ::=
    'KEY_GENERATE_STRATEGY' '(' 'COLUMN' '=' columnName ',' ( 'KEY_GENERATOR' '='
algorithmName | algorithmDefinition ) ')'

algorithmDefinition ::=
    ( 'SHARDING_ALGORITHM' '=' algorithmName | 'TYPE' '(' 'NAME' '=' algorithmType (
', ' 'PROPERTIES' '(' propertyDefinition ')' )? ')' )

propertyDefinition ::=
    ( key '=' value ) ( ',' key '=' value ) *

tableName ::=
    identifier

resourceName ::=
    identifier

columnName ::=
    identifier

algorithmName ::=
    identifier

```

补充说明

- `tableDefinition` 为标准分片规则定义；`autoTableDefinition` 为自动分片规则定义。标准分片规则和自动分片规则可参考[数据分片](#)；
- 当使用标准分片时：
 - DATANODES 只能使用已经添加到当前逻辑库的资源，且只能使用 `INLINE` 表达式指定需要的资源；
 - `DATABASE_STRATEGY`、`TABLE_STRATEGY` 表示分库和分表策略，均为可选项，未配置时使用默认策略；
 - `strategyDefinition` 中属性 `TYPE` 用于指定分片算法的类型，目前仅支持 `STANDARD`、`COMPLEX`。使用 `COMPLEX` 时需要用 `SHARDING_COLUMNS` 指定多个分片键。

- 当使用自动分片时：
 - RESOURCES 只能使用已经添加到当前逻辑库的资源，可通过枚举或 `INLINE` 表达式指定需要的资源；
 - 只能使用自动分片算法，可参考[自动分片算法](#)。
- `algorithmType` 为分片算法类型，分片算法类型请参考[分片算法](#)；
- 自动生成的算法命名规则为 `tableName_strategyType_algorithmType`；
- 自动生成的主键策略命名规则为 `tableName_‘strategyType`；
- `KEY_GENERATE_STRATEGY` 用于指定主键生成策略，为可选项，关于主键生成策略可参考[分布式主键](#)。

示例

1. 标准分片规则

- 指定分片算法创建标准分片规则

```
-- 创建分片算法
CREATE SHARDING ALGORITHM database_inline (
    TYPE(NAME="inline", PROPERTIES("algorithm-expression"="t_order_${user_id % 2}
"))
), table_inline (
    TYPE(NAME="inline", PROPERTIES("algorithm-expression"="t_order_${order_id % 2}
"))
);

-- 指定分片算法创建分片规则
CREATE SHARDING TABLE RULE t_order (
    DATANODES("ds_${0..1}.t_order_${0..1}"),
    DATABASE_STRATEGY(TYPE="standard", SHARDING_COLUMN=user_id, SHARDING_
ALGORITHM=database_inline),
    TABLE_STRATEGY(TYPE="standard", SHARDING_COLUMN=order_id, SHARDING_
ALGORITHM=table_inline)
);
```

- 在默认分库策略下，通过指定分片算法创建标准分片规则

```
-- 创建分片算法
CREATE SHARDING ALGORITHM database_inline (
    TYPE(NAME="inline", PROPERTIES("algorithm-expression"="t_order_${user_id % 2}
"))
), table_inline (
    TYPE(NAME="inline", PROPERTIES("algorithm-expression"="t_order_${order_id % 2}
"))
);
```

```
);

-- 创建默认分库策略
CREATE DEFAULT SHARDING DATABASE STRATEGY (
    TYPE="standard", SHARDING_COLUMN=order_id, SHARDING_ALGORITHM=database_inline
);

-- 指定分片算法创建分片规则
CREATE SHARDING TABLE RULE t_order (
    DATANODES("ds_${0..1}.t_order_${0..1}"),
    TABLE_STRATEGY(TYPE="standard", SHARDING_COLUMN=order_id, SHARDING_
ALGORITHM=table_inline)
);
```

- 在默认分库分表策略下，通过指定分片算法创建标准分片规则

```
-- 创建分片算法
CREATE SHARDING ALGORITHM database_inline (
    TYPE(NAME="inline", PROPERTIES("algorithm-expression"="t_order_${user_id % 2}
"))
), table_inline (
    TYPE(NAME="inline", PROPERTIES("algorithm-expression"="t_order_${order_id % 2}
"))
);

-- 创建默认分库策略
CREATE DEFAULT SHARDING DATABASE STRATEGY (
    TYPE="standard", SHARDING_COLUMN=order_id, SHARDING_ALGORITHM=database_inline
);

-- 创建默认分表策略
CREATE DEFAULT SHARDING TABLE STRATEGY (
    TYPE="standard", SHARDING_COLUMN=order_id, SHARDING_ALGORITHM=table_inline
);

-- 创建分片规则
CREATE SHARDING TABLE RULE t_order (
    DATANODES("ds_${0..1}.t_order_${0..1}")
);
```

- 创建标准分片规则的同时创建分片算法

```
CREATE SHARDING TABLE RULE t_order (
    DATANODES("ds_${0..1}.t_order_${0..1}"),
    DATABASE_STRATEGY(TYPE="standard", SHARDING_COLUMN=user_id, SHARDING_
ALGORITHM(TYPE(NAME="inline", PROPERTIES("algorithm-expression"="ds_${user_id % 2}
")))),
    TABLE_STRATEGY(TYPE="standard", SHARDING_COLUMN=user_id, SHARDING_
ALGORITHM(TYPE(NAME="inline", PROPERTIES("algorithm-expression"="ds_${order_id % 2}
"))))
);
```

2. 自动分片规则

- 创建自动分片规则

```
CREATE SHARDING TABLE RULE t_order (
    RESOURCES(ds_0, ds_1),
    SHARDING_COLUMN=order_id, TYPE(NAME="MOD", PROPERTIES("sharding-count"="4"))
);
```

保留字

CREATE、SHARDING、TABLE、RULE、DATANODES、DATABASE_STRATEGY、TABLE_STRATEGY、KEY_GENERATE_STRATEGY、RESOURCES、SHARDING_COLUMN、TYPE、SHARDING_COLUMN、KEY_GENERATOR、SHARDING_ALGORITHM、COLUMN、NAME、PROPERTIES

相关链接

- [保留字](#)
- [CREATE SHARDING ALGORITHM](#)
- [CREATE DEFAULT_SHARDING STRATEGY](#)

ALTER SHARDING TABLE RULE

描述

ALTER SHARDING TABLE RULE 语法用于修改当前所选逻辑库的分片规则

语法定义

```

AlterShardingTableRule ::=
    'ALTER' 'SHARDING' 'TABLE' 'RULE' ( tableDefinition | autoTableDefinition ) ( ','
( tableDefinition | autoTableDefinition ) )*

tableDefinition ::=
    tableName '(' 'DATANODES' '(' dataNode ( ',' dataNode )* ')' ( ',' 'DATABASE_
STRATEGY' '(' strategyDefinition ')' )? ( ',' 'TABLE_STRATEGY' '('
strategyDefinition ')' )? ( ',' 'KEY_GENERATE_STRATEGY' '('
keyGenerateStrategyDefinition ')' )? ')'

autoTableDefinition ::=
    tableName '(' 'RESOURCES' '(' resourceName ( ',' resourceName )* ')' ','
'SHARDING_COLUMN' '=' columnName ',' algorithmDefinition ( ',' 'KEY_GENERATE_
STRATEGY' '(' keyGenerateStrategyDefinition ')' )? ')'

strategyDefinition ::=
    'TYPE' '=' strategyType ',' ( 'SHARDING_COLUMN' | 'SHARDING_COLUMNS' ) '='
columnName ',' algorithmDefinition

keyGenerateStrategyDefinition ::=
    'KEY_GENERATE_STRATEGY' '(' 'COLUMN' '=' columnName ',' ( 'KEY_GENERATOR' '='
algorithmName | algorithmDefinition ) ')'

algorithmDefinition ::=
    ( 'SHARDING_ALGORITHM' '=' algorithmName | 'TYPE' '(' 'NAME' '=' algorithmType (
',' 'PROPERTIES' '(' propertyDefinition ')' )? ) ')'

propertyDefinition ::=
    ( key '=' value ) ( ',' key '=' value ) *

tableName ::=
    identifier

resourceName ::=
    identifier

columnName ::=
    identifier

algorithmName ::=
    identifier

```

补充说明

- `tableDefinition` 为标准分片规则定义；
- `autoTableDefinition` 为自动分片规则定义。标准分片规则和自动分片规则可参考[数据分片](#)；
- 当使用标准分片时：
 - DATANODES 只能使用已经添加到当前逻辑库的资源，且只能使用 `INLINE` 表达式指定需要的资源；
 - DATABASE_STRATEGY、TABLE_STRATEGY 表示分库和分表策略，均为可选项，未配置时使用默认策略；
 - strategyDefinition 中属性 TYPE 用于指定分片算法的类型，目前仅支持 STANDARD、COMPLEX。使用 COMPLEX 时需要用 SHARDING_COLUMNS 指定多个分片键。
- 当使用自动分片时：
 - RESOURCES 只能使用已经添加到当前逻辑库的资源，可通过枚举或 `INLINE` 表达式指定需要的资源；
 - 只能使用自动分片算法，可参考[自动分片算法](#)。
- algorithmType 为分片算法类型，分片算法类型请参考[分片算法](#)；
- 自动生成的算法命名规则为 `tableName_strategyType_algorithmType`；
- 自动生成的主键策略命名规则为 `tableName_strategyType`；
- KEY_GENERATE_STRATEGY 用于指定主键生成策略，为可选项，关于主键生成策略可参考[分布式主键](#)。

示例

1. 标准分片规则

- 修改分片算法并修改标准分片规则为指定分片算法

```
-- 修改分片算法
ALTER SHARDING ALGORITHM database_inline (
    TYPE(NAME="inline", PROPERTIES("algorithm-expression"="t_order_${user_id % 4}"))
), table_inline (
    TYPE(NAME="inline", PROPERTIES("algorithm-expression"="t_order_${order_id % 4}"))
);

-- 修改分片规则为指定分片算法
ALTER SHARDING TABLE RULE t_order (
    DATANODES("resource_${0..3}.t_order_item${0..3}"),
    DATABASE_STRATEGY(TYPE="standard", SHARDING_COLUMN=user_id, SHARDING_
```

```
ALGORITHM=database_inline),
    TABLE_STRATEGY(TYPE="standard", SHARDING_COLUMN=order_id, SHARDING_
ALGORITHM=table_inline)
);
```

- 修改默认分库策略并修改标准分片规则为指定分片算法

```
-- 修改分片算法
ALTER SHARDING ALGORITHM database_inline (
    TYPE(NAME="inline", PROPERTIES("algorithm-expression"="t_order_${user_id % 4}
"))
), table_inline (
    TYPE(NAME="inline", PROPERTIES("algorithm-expression"="t_order_${order_id % 4}
"))
);

-- 修改默认分库策略
ALTER DEFAULT SHARDING DATABASE STRATEGY (
    TYPE="standard", SHARDING_COLUMN=order_id, SHARDING_ALGORITHM=database_inline
);

-- 修改分片规则为指定分片算法
ALTER SHARDING TABLE RULE t_order (
    DATANODES("resource_${0..3}.t_order_item${0..3}"),
    TABLE_STRATEGY(TYPE="standard", SHARDING_COLUMN=order_id, SHARDING_
ALGORITHM=table_inline)
);
```

- 修改默认分库分表策略并修改标准分片规则为指定分片算法

```
-- 修改分片算法
ALTER SHARDING ALGORITHM database_inline (
    TYPE(NAME="inline", PROPERTIES("algorithm-expression"="t_order_${user_id % 4}
"))
), table_inline (
    TYPE(NAME="inline", PROPERTIES("algorithm-expression"="t_order_${order_id % 4}
"))
);

-- 修改默认分库策略
ALTER DEFAULT SHARDING DATABASE STRATEGY (
    TYPE="standard", SHARDING_COLUMN=order_id, SHARDING_ALGORITHM=database_inline
);

-- 修改默认分表策略
```

```
ALTER DEFAULT SHARDING TABLE STRATEGY (
    TYPE="standard", SHARDING_COLUMN=order_id, SHARDING_ALGORITHM=table_inline
);

-- 修改分片规则
ALTER SHARDING TABLE RULE t_order (
    DATANODES("resource_${0..3}.t_order_item${0..3}")
);
```

- 修改标准分片规则的同时创建分片算法

```
ALTER SHARDING TABLE RULE t_order (
    DATANODES("resource_${0..3}.t_order_item${0..3}"),
    DATABASE_STRATEGY(TYPE="standard", SHARDING_COLUMN=user_id, SHARDING_
ALGORITHM(TYPE(NAME="inline", PROPERTIES("algorithm-expression"="ds_${user_id % 2}
")))),
    TABLE_STRATEGY(TYPE="standard", SHARDING_COLUMN=user_id, SHARDING_
ALGORITHM(TYPE(NAME="inline", PROPERTIES("algorithm-expression"="ds_${order_id % 2}
")))))
);
```

2. 自动分片规则

- 修改自动分片规则

```
ALTER SHARDING TABLE RULE t_order (
    RESOURCES(ds_0, ds_1),
    SHARDING_COLUMN=order_id, TYPE(NAME="MOD", PROPERTIES("sharding-count"="4"))
);
```

保留字

ALTER、SHARDING、TABLE、RULE、DATANODES、DATABASE_STRATEGY、TABLE_STRATEGY、KEY_GENERATE_STRATEGY、RESOURCES、SHARDING_COLUMN、TYPE、SHARDING_COLUMN、KEY_GENERATOR、SHARDING_ALGORITHM、COLUMN、NAME、PROPERTIES

相关链接

- [保留字](#)
- [ALTER SHARDING ALGORITHM](#)
- [ALTER DEFAULT_SHARDING STRATEGY](#)

CREATE SHARDING ALGORITHM

描述

CREATE SHARDING ALGORITHM 语法用于为当前所选的逻辑库添加分片算法

语法定义

```
CreateShardingAlgorithm ::=
    'CREATE' 'SHARDING' 'ALGORITHM' shardingAlgorithmName '(' algorithmDefinition ')'

algorithmDefinition ::=
    'TYPE' '(' 'NAME' '=' algorithmType ( ',' 'PROPERTIES' '(' propertyDefinition
    ')' )? ')'

propertyDefinition ::=
    ( key '=' value ) ( ',' key '=' value ) *

shardingAlgorithmName ::=
    identifier

algorithmType ::=
    identifier
```

补充说明

- `algorithmType` 为分片算法类型，详细的分片算法类型信息请参考[分片算法](#)。

示例

1. 创建分片算法

```
-- 创建类型为 INLINE 的分片算法
CREATE SHARDING ALGORITHM inline_algorithm (
    TYPE(NAME="inline", PROPERTIES("algorithm-expression"="t_order_${user_id % 2}"))
);

-- 创建类型为 AUTO_INTERVAL 的分片算法
CREATE SHARDING ALGORITHM interval_algorithm (
    TYPE(NAME="auto_interval", PROPERTIES("datetime-lower"="2022-01-01 00:00:00",
"datetime-upper"="2022-01-03 00:00:00", "sharding-seconds"="86400"))
);
```

保留字

CREATE、SHARDING、ALGORITHM、TYPE、NAME、PROPERTIES

相关链接

- [保留字](#)

CREATE DEFAULT SHARDING STRATEGY

描述

CREATE DEFAULT SHARDING STRATEGY 语法用于创建默认的分片策略

语法定义

```
CreateDefaultShardingStrategy ::=
    'CREATE' 'DEFAULT' 'SHARDING' ('DATABASE' | 'TABLE') 'STRATEGY' '('
    shardingStrategy ')'

shardingStrategy ::=
    'TYPE' '=' strategyType ',' ( 'SHARDING_COLUMN' '=' columnName | 'SHARDING_
    COLUMNS' '=' columnNames ) ',' ( 'SHARDING_ALGORITHM' '=' algorithmName |
```

```

algorithmDefinition )

algorithmDefinition ::=
    'TYPE' '(' 'NAME' '=' algorithmType ( ',' 'PROPERTIES' '(' propertyDefinition
    ')' )? ')'

columnNames ::=
    columnName ( ',' columnName)+

columnName ::=
    identifier

algorithmName ::=
    identifier

algorithmType ::=
    identifier

```

补充说明

- 当使用复合分片算法时，需要通过 SHARDING_COLUMNS 指定多个分片键；
- algorithmType 为分片算法类型，详细的分片算法类型信息请参考[分片算法](#)。

示例

1. 通过已有的分片算法创建默认分库策略

```

-- 创建分片算法
CREATE SHARDING ALGORITHM database_inline (
    TYPE(NAME="inline", PROPERTIES("algorithm-expression"="t_order_${order_id % 2}
"))
);

-- 创建默认分库策略
CREATE DEFAULT SHARDING DATABASE STRATEGY (
    TYPE="standard", SHARDING_COLUMN=user_id, SHARDING_ALGORITHM=database_inline
);

```

2. 同时创建分片算法和默认分表策略

```
-- 创建默认分表策略
CREATE DEFAULT SHARDING TABLE STRATEGY (
    TYPE="standard", SHARDING_COLUMN=user_id, SHARDING_ALGORITHM(TYPE(NAME=inline,
    PROPERTIES("algorithm-expression"="t_order_${user_id % 2}")))
);
```

保留字

CREATE、DEFAULT、SHARDING、DATABASE、TABLE、STRATEGY、TYPE、SHARDING_COLUMN、SHARDING_COLUMNS、SHARDING_ALGORITHM、NAME、PROPERTIES

相关链接

- [保留字](#)
- [CREATE SHARDING ALGORITHM](#)

CREATE SHARDING BINDING TABLE RULE

描述

CREATE SHARDING BINDING TABLE RULE 语法用于为具有分片规则的表（分片表）添加绑定关系并创建绑定规则

语法定义

```
CreateBindingTableRule ::=
    'CREATE' 'SHARDING' 'BINDING' 'TABLE' 'RULES' bindingRelationshipDefinition (' ,
    ' bindingRelationshipDefinition ) *

bindingRelationshipDefinition ::=
    '(' tableName (' , ' tableName ) * ')'

tableName ::=
    identifier
```


补充说明

- 只有分片表才能创建绑定关系；
- 一个分片表只能具有一个绑定关系；
- 添加绑定关系的分片表需要使用相同的资源，并且分片节点个数相同。例如 `ds_${0..1}.t_order_${0..1}` 与 `ds_${0..1}.t_order_item_${0..1}`；
- 添加绑定关系的分片表需要对分片键使用相同的分片算法。例如 `t_order_${order_id % 2}` 与 `t_order_item_${order_item_id % 2}`；
- 只能存在一个绑定规则，但可包含多个绑定关系，因此无法重复执行 `CREATE SHARDING BINDING TABLE RULE`。当绑定规则已经存在但还需要添加绑定关系时，需要使用 `ALTER SHARDING BINDING TABLE RULE` 来修改绑定规则。

示例

1. 创建绑定关系

```
-- 创建绑定关系之前需要先创建分片表 t_order,t_order_item
CREATE SHARDING BINDING TABLE RULES (t_order,t_order_item);
```

2. 创建多个绑定关系

```
-- 创建绑定关系之前需要先创建分片表 t_order,t_order_item,t_product,t_product_item
CREATE SHARDING BINDING TABLE RULES (t_order,t_order_item),(t_product,t_product_item);
```

保留字

CREATE、SHARDING、BINDING、TABLE、RULES

相关链接

- [保留字](#)
- [CREATE SHARDING TABLE RULE](#)

CREATE SHARDING BROADCAST TABLE RULE

描述

CREATE SHARDING BROADCAST TABLE RULE 语法用于为需要广播的表（广播表）创建广播规则

语法定义

```
CreateDefaultShardingStrategy ::=
    'CREATE' 'DEFAULT' 'SHARDING' ('DATABASE' | 'TABLE') 'STRATEGY' '('
shardingStrategy ')'

shardingStrategy ::=
    'TYPE' '=' strategyType ',' ( 'SHARDING_COLUMN' '=' columnName | 'SHARDING_
COLUMNS' '=' columnNames ) ',' ( 'SHARDING_ALGORITHM' '=' algorithmName |
algorithmDefinition )

algorithmDefinition ::=
    'TYPE' '(' 'NAME' '=' algorithmType ( ',' 'PROPERTIES' '(' propertyDefinition
')' )? ')'

columnNames ::=
    columnName ( ',' columnName)+

columnName ::=
    identifier

algorithmName ::=
    identifier

algorithmType ::=
    identifier
```

补充说明

- tableName 可使用已经存在的表或者将要创建的表；
- 只能存在一个广播规则,但可包含多个广播表,因此无法重复执行 CREATE SHARDING BROADCAST TABLE RULE。当广播规则已经存在但还需要添加广播表时,需要使用 ALTER BROADCAST TABLE RULE 来修改广播规则。

示例

创建广播规则

```
-- 将 t_province, t_city 添加到广播规则中
CREATE SHARDING BROADCAST TABLE RULES (t_province, t_city);
```

保留字

CREATE、SHARDING、BROADCAST、TABLE、RULES

相关链接

- [保留字](#)

CREATE SHARDING KEY GENERATOR

描述

CREATE SHARDING KEY GENERATOR 语法用于为当前所选的逻辑库添加分布式主键生成器

语法定义

```
CreateShardingAlgorithm ::=
  'CREATE' 'SHARDING' 'KEY' 'GENERATOR' keyGeneratorName '(' algorithmDefinition ')'
  ' '

algorithmDefinition ::=
  'TYPE' '(' 'NAME' '=' algorithmType ( ',' 'PROPERTIES' '(' propertyDefinition
  ')' )? ')'

propertyDefinition ::=
  ( key '=' value ) ( ',' key '=' value ) *

keyGeneratorName ::=
  identifier

algorithmType ::=
  identifier
```

补充说明

- `algorithmType` 为分布式主键生成算法类型, 详细的分布式主键生成算法类型信息请参考[分布式序列算法类型](#)。

示例

创建分布式主键生成器

```
CREATE SHARDING KEY GENERATOR snowflake_key_generator (  
    TYPE(NAME="SNOWFLAKE", PROPERTIES("max-vibration-offset"="3"))  
);
```

保留字

CREATE、SHARDING、KEY、GENERATOR、TYPE、NAME、PROPERTIES

相关链接

- [保留字](#)

单表

本章节将对单表特性的语法进行详细说明。

CREATE DEFAULT SINGLE TABLE RULE

描述

CREATE DEFAULT SINGLE TABLE RULE 语法用于创建默认的单表规则

语法定义

```
CreateDefaultSingleTableRule ::=  
    'CREATE' 'DEFAULT' 'SINGLE' 'TABLE' 'RULE' singleTableDefinition  
  
singleTableDefinition ::=  
    'RESOURCE' '=' resourceName  
  
resourceName ::=  
    identifier
```

补充说明

- RESOURCE 需使用 RDL 管理的数据源资源。

示例

创建默认单表规则

```
CREATE DEFAULT SINGLE TABLE RULE RESOURCE = ds_0;
```

保留字

CREATE、SHARDING、SINGLE、TABLE、RULE、RESOURCE

相关链接

- [保留字](#)

RQL 语法

RQL (Resource & Rule Query Language) 为 Apache ShardingSphere 的资源 and 规则查询语言。

资源查询

本章节将对资源查询的语法进行详细说明。

SHOW RESOURCE

描述

SHOW RESOURCE 语法用于查询指定逻辑库已经添加的资源。

语法

```
ShowResource ::=  
    'SHOW' 'DATABASE' 'RESOURCES' ('FROM' databaseName)?  
  
databaseName ::=  
    identifier
```



```

-----+
-----+
-----+
| ds_0 | MySQL | 127.0.0.1 | 3306 | db_0 | 30000 | 60000
|      |      |      |      |      | 1800000 | 50 | 1 |
false | {"dataSourceProperties":{"cacheServerConfiguration":"true",
"elideSetAutoCommits":"true","useServerPrepStmts":"true","cachePrepStmts":"true",
"rewriteBatchedStatements":"true","cacheResultSetMetadata":"false",
"useLocalSessionState":"true","maintainTimeStats":"false","prepStmtCacheSize":
"200000","tinyIntTrisBit":"false","prepStmtCacheSqlLimit":"2048",
"zeroDateTimeBehavior":"round","netTimeoutForStreamingResults":"0"},
"healthCheckProperties":{},"initializationFailTimeout":1,"validationTimeout":5000,
"leakDetectionThreshold":0,"registerMbeans":false,"allowPoolSuspension":false,
"autoCommit":true,"isolateInternalQueries":false} |
| ds_1 | MySQL | 127.0.0.1 | 3306 | db_1 | 30000 | 60000
|      |      |      |      |      | 1800000 | 50 | 1 |
false | {"dataSourceProperties":{"cacheServerConfiguration":"true",
"elideSetAutoCommits":"true","useServerPrepStmts":"true","cachePrepStmts":"true",
"rewriteBatchedStatements":"true","cacheResultSetMetadata":"false",
"useLocalSessionState":"true","maintainTimeStats":"false","prepStmtCacheSize":
"200000","tinyIntTrisBit":"false","prepStmtCacheSqlLimit":"2048",
"zeroDateTimeBehavior":"round","netTimeoutForStreamingResults":"0"},
"healthCheckProperties":{},"initializationFailTimeout":1,"validationTimeout":5000,
"leakDetectionThreshold":0,"registerMbeans":false,"allowPoolSuspension":false,
"autoCommit":true,"isolateInternalQueries":false} |
+-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
2 rows in set (0.26 sec)

```

- 查询当前逻辑库的资源

```
SHOW DATABASE RESOURCES;
```

```

+-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+

```

```

-----+
-----+
-----+
| name | type | host      | port | db   | connection_timeout_milliseconds | idle_
timeout_milliseconds | max_lifetime_milliseconds | max_pool_size | min_pool_size |
read_only | other_attributes

```

```

-----+
| ds_0 | MySQL | 127.0.0.1 | 3306 | db_0 | 30000 | 60000
| 1800000 | 50 | 1 |
false | {"dataSourceProperties":{"cacheServerConfiguration":"true",
"elideSetAutoCommits":"true","useServerPrepStmts":"true","cachePrepStmts":"true",
"rewriteBatchedStatements":"true","cacheResultSetMetadata":"false",
"useLocalSessionState":"true","maintainTimeStats":"false","prepStmtCacheSize":
"200000","tinyInt1isBit":"false","prepStmtCacheSqlLimit":"2048",
"zeroDateTimeBehavior":"round","netTimeoutForStreamingResults":"0"},
"healthCheckProperties":{"initializationFailTimeout":1,"validationTimeout":5000,
"leakDetectionThreshold":0,"registerMbeans":false,"allowPoolSuspension":false,
"autoCommit":true,"isolateInternalQueries":false} |
| ds_1 | MySQL | 127.0.0.1 | 3306 | db_1 | 30000 | 60000
| 1800000 | 50 | 1 |
false | {"dataSourceProperties":{"cacheServerConfiguration":"true",
"elideSetAutoCommits":"true","useServerPrepStmts":"true","cachePrepStmts":"true",
"rewriteBatchedStatements":"true","cacheResultSetMetadata":"false",
"useLocalSessionState":"true","maintainTimeStats":"false","prepStmtCacheSize":
"200000","tinyInt1isBit":"false","prepStmtCacheSqlLimit":"2048",
"zeroDateTimeBehavior":"round","netTimeoutForStreamingResults":"0"},
"healthCheckProperties":{"initializationFailTimeout":1,"validationTimeout":5000,
"leakDetectionThreshold":0,"registerMbeans":false,"allowPoolSuspension":false,
"autoCommit":true,"isolateInternalQueries":false} |

```



```
-----+
2 rows in set (0.26 sec)
```

SHOW UNUSED RESOURCE

描述

SHOW UNUSED RESOURCE 语法用于查询指定逻辑库中还未被规则引用的资源。

语法

```
ShowUnusedResource ::=
    'SHOW' 'UNUSED' 'DATABASE'? 'RESOURCES' ('FROM' dbName)?

dbName ::=
    identifier
```

特别说明

- 未指定 dbName 时，默认是当前使用的 DATABASE；如未使用 DATABASE 则会提示 No database selected。

返回值说明

列	说明
name	数据源名称
type	数据源类型
host	数据源地址
port	数据源端口
db	数据库名称
attribute	数据源参数

示例

- 查询指定逻辑库的资源

```
SHOW UNUSED DATABASE RESOURCES FROM sharding_db;
```



```
-----+
1 rows in set (0.26 sec)
```

- 查询当前逻辑库的资源

```
SHOW UNUSED DATABASE RESOURCES;
```

```
-----+
| name | type | host      | port | db   | connection_timeout_milliseconds | idle_
timeout_milliseconds | max_lifetime_milliseconds | max_pool_size | min_pool_size |
read_only | other_attributes
```

```
-----+
| ds_0 | MySQL | 127.0.0.1 | 3306 | db_0 | 30000 | 60000
| 1800000 | 50 | 1 |
false | {"dataSourceProperties":{"cacheServerConfiguration":"true",
"elideSetAutoCommits":"true","useServerPrepStmts":"true","cachePrepStmts":"true",
"rewriteBatchedStatements":"true","cacheResultSetMetadata":"false",
"useLocalSessionState":"true","maintainTimeStats":"false","prepStmtCacheSize":
"200000","tinyInt1isBit":"false","prepStmtCacheSqlLimit":"2048",
"zeroDateTimeBehavior":"round","netTimeoutForStreamingResults":"0"},
"healthCheckProperties":{},"initializationFailTimeout":1,"validationTimeout":5000,
```

```
"leakDetectionThreshold":0,"registerMbeans":false,"allowPoolSuspension":false,
"autoCommit":true,"isolateInternalQueries":false} |
+-----+-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
-----+-----+-----+-----+-----+-----+-----+-----+
1 rows in set (0.26 sec)
```

SHOW RULES USED RESOURCE

描述

SHOW RULES USED RESOURCE 语法用于查询指定逻辑库中使用指定资源的规则。

语法

```
showRulesUsedResource ::=
    'SHOW' 'RULES' 'USED' 'RESOURCES' resourceName ('FROM' databaseName)?

resourceName ::=
    IDENTIFIER | STRING

databaseName ::=
    IDENTIFIER
```

特别说明

- 未指定 `databaseName` 时, 默认是当前使用的 DATABASE; 如未使用 DATABASE 则会提示 No database selected。

返回值说明

列	说明
type	特性
name	数据源名称

示例

- 查询指定逻辑库中使用指定资源的规则

```
SHOW RULES USED RESOURCE ds_0 FROM sharding_db;
```

```
+-----+-----+
| type   | name       |
+-----+-----+
| sharding | t_order    |
| sharding | t_order_item |
+-----+-----+
2 rows in set (0.00 sec)
```

- 查询当前逻辑库中使用指定资源的规则

```
SHOW RULES USED RESOURCE ds_0;
```

```
+-----+-----+
| type   | name       |
+-----+-----+
| sharding | t_order    |
| sharding | t_order_item |
+-----+-----+
2 rows in set (0.00 sec)
```

规则查询

本章节将对规则查询的语法进行详细说明。

分片

本章节将对分片特性的语法进行详细说明。

SHOW SHARDING TABLE RULE

描述

SHOW SHARDING TABLE RULE 语法用于查询指定逻辑库中的分片规则。

语法

```
ShowShardingTableRule ::=
    'SHOW' 'SHARDING' 'TABLE' ('RULE' tableName | 'RULES') ('FROM' databaseName)?

tableName ::=
    identifier

databaseName ::=
    identifier
```

补充说明

- 未指定 databaseName 时, 默认是当前使用的 DATABASE。如果也未使用 DATABASE 则会提示 No database selected。

返回值说明

列	说明
table	逻辑表名
actual_data_nodes	实际的数据节点
actual_data_sources	实际的数据源（通过 RDL 创建的规则时显示）
database_strategy_type	数据库分片策略类型
database_sharding_column	数据库分片键
database_sharding_algorithm_type	数据库分片算法类型
database_sharding_algorithm_props	数据库分片算法参数
table_strategy_type	表分片策略类型
table_sharding_column	表分片键
table_sharding_algorithm_type	表分片算法类型
table_sharding_algorithm_props	表分片算法参数
key_generate_column	分布式主键生成列
key_generator_type	分布式主键生成器类型
key_generator_props	分布式主键生成器参数

示例 - 查询指定逻辑库的分片规则

```
SHOW SHARDING TABLE RULES FROM sharding_db;
```

```
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
| table          | actual_data_nodes | actual_data_sources | database_strategy_type |
database_sharding_column | database_sharding_algorithm_type | database_sharding_
algorithm_props | table_strategy_type | table_sharding_column | table_sharding_
algorithm_type | table_sharding_algorithm_props | key_generate_column | key_
generator_type | key_generator_props |
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
+-----+-----+-----+-----+
| t_order        |                  | ds_0,ds_1           |
|               |                  |                     |
|      mod      |                  | order_id             | mod
| sharding-count=4 |                  |                     |
|               |                  |                     |
| t_order_item   |                  | ds_0,ds_1           |
|               |                  |                     |
|      mod      |                  | order_id             | mod
| sharding-count=4 |                  |                     |
|               |                  |                     |
+-----+-----+-----+-----+
```

```
+-----+-----+-----+
+-----+-----+-----+
+-----+-----+-----+
+-----+-----+-----+
2 rows in set (0.12 sec)
```

- 查询当前逻辑库的分片规则

```
SHOW SHARDING TABLE RULES;
```

```
+-----+-----+-----+-----+-----+
+-----+-----+-----+-----+-----+
+-----+-----+-----+-----+-----+
+-----+-----+-----+-----+-----+
+-----+-----+-----+-----+-----+
| table          | actual_data_nodes | actual_data_sources | database_strategy_type |
database_sharding_column | database_sharding_algorithm_type | database_sharding_
algorithm_props | table_strategy_type | table_sharding_column | table_sharding_
algorithm_type | table_sharding_algorithm_props | key_generate_column | key_
generator_type | key_generator_props |
+-----+-----+-----+-----+-----+
+-----+-----+-----+-----+-----+
+-----+-----+-----+-----+-----+
+-----+-----+-----+-----+-----+
+-----+-----+-----+-----+-----+
| t_order        |                  | ds_0,ds_1          |                        |
|                |                  | order_id           |                        |
|                | mod             |                    | mod                   |
|                | sharding-count=4 |                    |                    |
|                |                  |                    |                    |
| t_order_item   |                  | ds_0,ds_1          |                        |
|                |                  | order_id           |                        |
|                | mod             |                    | mod                   |
|                | sharding-count=4 |                    |                    |
|                |                  |                    |                    |
+-----+-----+-----+-----+-----+
+-----+-----+-----+-----+-----+
+-----+-----+-----+-----+-----+
+-----+-----+-----+-----+-----+
+-----+-----+-----+-----+-----+
2 rows in set (0.12 sec)
```

- 查询指定逻辑表的分片规则

```
SHOW SHARDING TABLE RULE t_order;
```

```
+-----+-----+-----+-----+-----+
+-----+-----+-----+-----+-----+
```



```

-----+-----+-----+-----+
-----+-----+-----+-----+
--++-----+
| table          | actual_data_nodes | actual_data_sources | database_strategy_type |
database_sharding_column | database_sharding_algorithm_type | database_sharding_
algorithm_props | table_strategy_type | table_sharding_column | table_sharding_
algorithm_type | table_sharding_algorithm_props | key_generate_column | key_
generator_type | key_generator_props |
+-----+-----+-----+-----+
+-----+-----+-----+-----+
-----+-----+-----+-----+
-----+-----+-----+-----+
--++-----+
| t_order        |                  | ds_0,ds_1          |                  |
|               |                  |                   |                   |
|               | mod              | order_id            | mod              |
| sharding-count=4 |                  |                     |                  |
|               |                  |                     |                  |
+-----+-----+-----+-----+
+-----+-----+-----+-----+
-----+-----+-----+-----+
--++-----+
1 rows in set (0.12 sec)

```

RAL 语法

RAL (Resource & Rule Administration Language) 为 Apache ShardingSphere 的管理语言，负责强制路由、事务类型切换、弹性伸缩、分片执行计划查询等增量功能的操作。

保留字

资源定义

ADD、ALTER、DROP、RESOURCE、IF、EXISTS、HOST、PORT、DB、USER、PASSWORD、PROPERTIES、URL、IGNORE、SINGLE、TABLES

规则定义

分片

CREATE、ALTER、DEFAULT、SHARDING、BROADCAST、BINDING、DATABASE、TABLE、STRATEGY、RULE、RULES、ALGORITHM、DATANODES、DATABASE_STRATEGY、TABLE_STRATEGY、KEY_GENERATE_STRATEGY、RESOURCES、SHARDING_COLUMN、KEY、GENERATOR、TYPE、SHARDING_COLUMNS、KEY_GENERATOR、SHARDING_ALGORITHM、COLUMN、NAME、PROPERTIES

单表

CREATE、SHARDING、SINGLE、TABLE、RULE、RESOURCE

读写分离

CREATE、READWRITE_SPLITTING、RULE、WRITE_RESOURCE、READ_RESOURCES、AUTO_AWARE_RESOURCE、WRITE_DATA_SOURCE_QUERY_ENABLED、TYPE、NAME、PROPERTIES、TRUE、FALSE

数据加密

CREATE、ENCRYPT、RULE、COLUMNS、NAME、CIPHER、PLAIN、QUERY_WITH_CIPHER_COLUMN、TYPE、TRUE、FALSE

数据库发现

CREATE、DB_DISCOVERY、RULE、RESOURCES、TYPE、NAME、PROPERTIES、HEARTBEAT

影子压测

CREATE、SHADOW、RULE、SOURCE、SHADOW、TYPE、NAME、PROPERTIES

补充说明

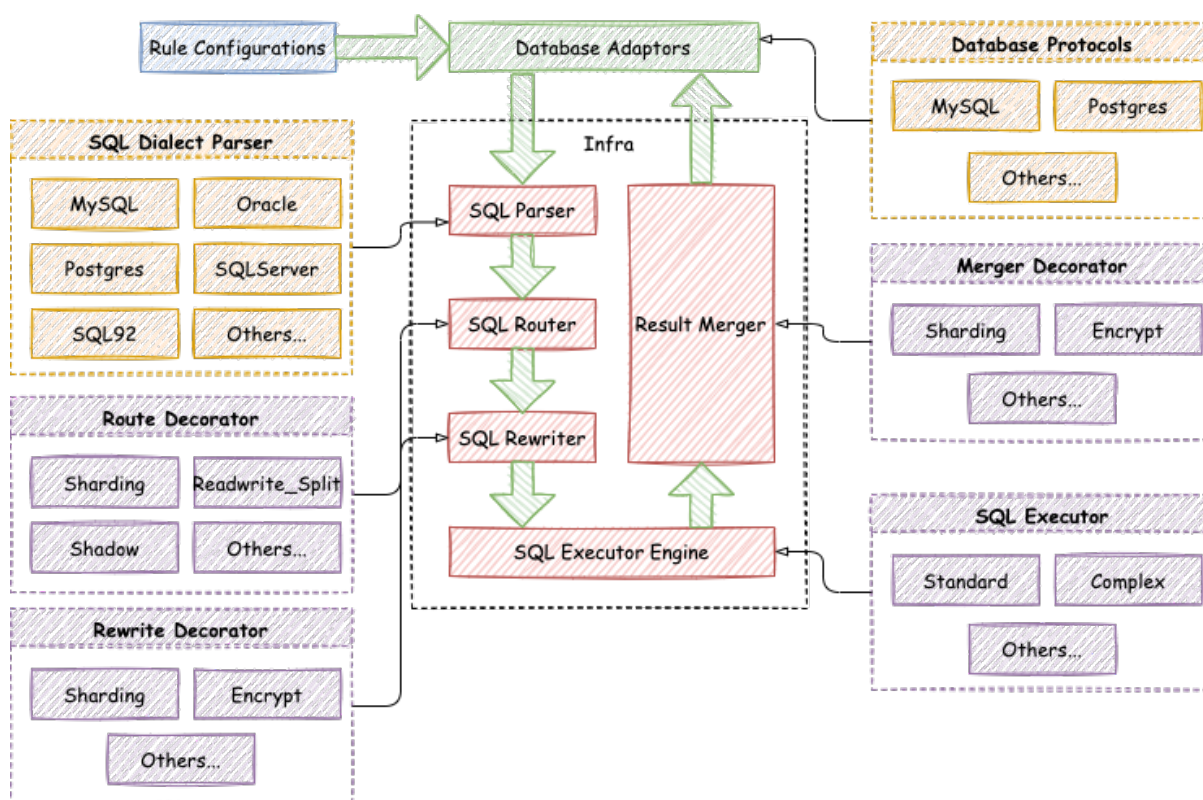
- 上述保留字大小写不敏感

7.11 基础架构

让开发者能够像使用积木一样定制属于自己的独特系统，是 Apache ShardingSphere 可插拔架构的设计目标。

可插拔架构对程序架构设计的要求非常高，需要将各个模块相互独立，互不感知，并且通过一个可插拔内核，以叠加的方式将各种功能组合使用。设计一套将功能开发完全隔离的架构体系，既可以最大限度的将开源社区的活力激发出来，也能够保障项目的质量。

Apache ShardingSphere 5.x 版本开始致力于可插拔架构，项目的功能组件能够灵活的以可插拔的方式进行扩展。目前，数据分片、读写分离、数据库高可用、数据加密、影子库压测等功能，以及对 MySQL、PostgreSQL、SQLServer、Oracle 等 SQL 与协议的支持，均通过插件的方式织入项目。Apache ShardingSphere 目前已提供数十个 SPI（Service Provider Interface）作为系统的扩展点，而且仍在不断增加中。



8.1 JDBC

8.1.1 JDBC 为什么配置了某个数据连接池的 **spring-boot-starter**（比如 **druid**）和 **shardingsphere-jdbc-spring-boot-starter** 时，系统启动会报错？

回答：

1. 因为数据连接池的 starter（比如 druid）可能会先加载并且其创建一个默认数据源，这将会使得 ShardingSphere-JDBC 创建数据源时发生冲突。
2. 解决办法为，去掉数据连接池的 starter 即可，ShardingSphere-JDBC 自己会创建数据连接池。

8.1.2 JDBC 使用 **Spring** 命名空间时找不到 **xsd**？

回答：

Spring 命名空间使用规范并未强制要求将 **xsd** 文件部署至公网地址，但考虑到部分用户的需求，我们也将相关 **xsd** 文件部署至 ShardingSphere 官网。实际上 **shardingsphere-jdbc-spring-namespace** 的 jar 包中 **META-INF:raw-latex:spring.schemas** 配置了 **xsd** 文件的位置：**META-INF:raw-latex:namespace:raw-latex:'sharding'**.xsd 和 **META-INF:raw-latex:namespace:raw-latex:'replica'**-query.xsd，只需确保 jar 包中该文件存在即可。

8.1.3 JDBC 引入 **shardingsphere-transaction-xa-core** 后，如何避免 **spring-boot** 自动加载默认的 **JtaTransactionManager**？

回答：

1. 需要在 **spring-boot** 的引导类中添加 **@SpringBootApplication(exclude = JtaAutoConfiguration.class)**。### **JDBC** Oracle 表名、字段名配置大小写在加载 metadata 元数据时结果不正确？回答：需要注意，Oracle 表名和字段名，默认元数据都是大写，除非建表语句中带双引号，如 **CREATE TABLE "TableName"("Id" number)** 元数据为双引号中内容，可参考以下 SQL 查看元数据的具体情况：

```
SELECT OWNER, TABLE_NAME, COLUMN_NAME, DATA_TYPE FROM ALL_TAB_COLUMNS WHERE TABLE_
NAME IN ('TableName')
```

ShardingSphere 使用 `OracleTableMetaDataLoader` 对 Oracle 元数据进行加载，配置时需确保表名、字段名的大小写配置与数据库中的一致。ShardingSphere 查询元数据关键 SQL:

```
private String getTableMetaDataSQL(final Collection<String> tables, final
DatabaseMetaData metaData) throws SQLException {
    StringBuilder stringBuilder = new StringBuilder(28);
    if (versionContainsIdentityColumn(metaData)) {
        stringBuilder.append(", IDENTITY_COLUMN");
    }
    if (versionContainsCollation(metaData)) {
        stringBuilder.append(", COLLATION");
    }
    String collation = stringBuilder.toString();
    return tables.isEmpty() ? String.format(TABLE_META_DATA_SQL, collation)
        : String.format(TABLE_META_DATA_SQL_IN_TABLES, collation, tables.
stream().map(each -> String.format("%s'", each)).collect(Collectors.joining(",
"))));
}
```

8.2 Proxy

8.2.1 Proxy Windows 环境下，运行 ShardingSphere-Proxy，找不到或无法加载主类 `org.apache.shardingsphere.proxy.Bootstrap`，如何解决？

回答：

某些解压缩工具在解压 ShardingSphere-Proxy 二进制包时可能将文件名截断，导致找不到某些类。解决方案：打开 `cmd.exe` 并执行下面的命令：

```
tar zxvf apache-shardingsphere-${RELEASE.VERSION}-shardingsphere-proxy-bin.tar.gz
```

8.2.2 Proxy 在使用 ShardingSphere-Proxy 的时候，如何动态在添加新的逻辑库？

回答：

使用 ShardingSphere-Proxy 时，可以通过 `DistSQL` 动态的创建或移除逻辑库，语法如下：

```
CREATE DATABASE [IF NOT EXISTS] databaseName;
DROP DATABASE [IF EXISTS] databaseName;
```

例：

```
CREATE DATABASE sharding_db;  
DROP DATABASE sharding_db;
```

8.2.3 Proxy 在使用 ShardingSphere-Proxy 时，怎么使用合适的工具连接到 ShardingSphere-Proxy?

回答：

1. ShardingSphere-Proxy 可以看做是一个 database server，所以首选支持 SQL 命令连接和操作。
2. 如果使用其他第三方数据库工具，可能由于不同工具的特定实现导致出现异常。
3. 目前已测试的第三方数据库工具如下：
 - Navicat：11.1.13、15.0.20。
 - DataGrip：2020.1、2021.1（使用 IDEA/DataGrip 时打开 introspect using JDBC metadata 选项）。
 - WorkBench：8.0.25。

8.2.4 Proxy 使用 Navicat 等第三方数据库工具连接 ShardingSphere-Proxy 时，如果 ShardingSphere-Proxy 没有创建 Database 或者没有添加 Resource，连接失败？

回答：

1. 第三方数据库工具在连接 ShardingSphere-Proxy 时会发送一些 SQL 查询元数据，当 ShardingSphere-Proxy 没有创建 database 或者没有添加 resource 时，ShardingSphere-Proxy 无法执行 SQL。
2. 推荐先创建 database 和 resource 之后再使用第三方数据库工具连接。
3. 有关 resource 的详情请参考。[相关介绍](#)

8.3 分片

8.3.1 分片 Cloud not resolve placeholder …in string value …异常的解决方法？

回答：

行表达式标识符可以使用 `${...}` 或 `$->{...}`，但前者与 Spring 本身的属性文件占位符冲突，因此在 Spring 环境中使用行表达式标识符建议使用 `$->{...}`。

8.3.2 分片 inline 表达式返回结果为何出现浮点数?

回答:

Java 的整数相除结果是整数, 但是对于 inline 表达式中的 Groovy 语法则不同, 整数相除结果是浮点数。想获得除法整数结果需要将 A/B 改为 A.intdiv(B)。

8.3.3 分片如果只有部分数据库分库分表, 是否需要将不分库分表的表也配置在分片规则中?

回答:

不需要, ShardingSphere 会自动识别。

8.3.4 分片指定了泛型为 Long 的 SingleKeyTableShardingAlgorithm, 遇到 ClassCastException: Integer can not cast to Long?

回答:

必须确保数据库表中该字段和分片算法该字段类型一致, 如: 数据库中该字段类型为 int(11), 泛型所对应的分片类型应为 Integer, 如果需要配置为 Long 类型, 请确保数据库中该字段类型为 bigint。

8.3.5 [分片、PROXY] 实现 StandardShardingAlgorithm 自定义算法时, 指定了 Comparable 的具体类型为 Long, 且数据库表中字段类型为 bigint, 出现 ClassCastException: Integer can not cast to Long 异常。

回答:

实现 doSharding 方法时, 不建议指定方法声明中 Comparable 具体的类型, 而是在 doSharding 方法实现中对类型进行转换, 可以参考 ModShardingAlgorithm#doSharding 方法

8.3.6 分片 ShardingSphere 提供的默认分布式自增主键策略为什么不连续的, 且尾数大多为偶数?

回答:

ShardingSphere 采用 snowflake 算法作为默认的分布式自增主键策略, 用于保证分布式的情况下可以无中心化的生成不重复的自增序列。因此自增主键可以保证递增, 但无法保证连续。而 snowflake 算法的最后 4 位是在同一毫秒内的访问递增值。因此, 如果毫秒内并发度不高, 最后 4 位为零的几率则很大。因此并发度不高的应用生成偶数主键的几率会更高。在 3.1.0 版本中, 尾数大多为偶数的问题已彻底解决, 参见: <https://github.com/apache/shardingsphere/issues/1617>

8.3.7 分片如何在 inline 分表策略时，允许执行范围查询操作（BETWEEN AND、>、<、>=、<=）？

回答：

1. 需要使用 4.1.0 或更高版本。
2. 调整以下配置项（需要注意的是，此时所有的范围查询将会使用广播的方式查询每一个分表）：
 - 4.x 版本：allow.range.query.with.inline.sharding 设置为 true 即可（默认为 false）。
 - 5.x 版本：在 InlineShardingStrategy 中将 allow-range-query-with-inline-sharding 设置为 true 即可（默认为 false）。

8.3.8 分片为什么我实现了 KeyGenerateAlgorithm 接口，也配置了 Type，但是自定义的分布式主键依然不生效？

回答：

Service Provider Interface (SPI) 是一种为了被第三方实现或扩展的 API，除了实现接口外，还需要在 META-INF/services 中创建对应文件来指定 SPI 的实现类，JVM 才会加载这些服务。具体的 SPI 使用方式，请大家自行搜索。与分布式主键 KeyGenerateAlgorithm 接口相同，其他 ShardingSphere 的扩展功能也需要用相同的方式注入才能生效。

8.3.9 分片 ShardingSphere 除了支持自带的分布式自增主键之外，还能否支持原生的自增主键？

回答：

是的，可以支持。但原生自增主键有使用限制，即不能将原生自增主键同时作为分片键使用。由于 ShardingSphere 并不知晓数据库的表结构，而原生自增主键是不包含在原始 SQL 中内的，因此 ShardingSphere 无法将该字段解析为分片字段。如自增主键非分片键，则无需关注，可正常返回；若自增主键同时作为分片键使用，ShardingSphere 无法解析其分片值，导致 SQL 路由至多张表，从而影响应用的正确性。而原生自增主键返回的前提条件是 INSERT SQL 必须最终路由至一张表，因此，面对返回多表的 INSERT SQL，自增主键则会返回零。

8.4 数据加密

8.4.1 数据加密 JPA 和数据加密无法一起使用，如何解决？

回答：

由于数据加密的 DDL 尚未开发完成，因此对于自动生成 DDL 语句的 JPA 与数据加密一起使用时，会导致 JPA 的实体类（Entity）无法同时满足 DDL 和 DML 的情况。解决方案如下：1. 以需要加密的逻辑列名编写 JPA 的实体类（Entity）。2. 关闭 JPA 的 auto-ddl，如 auto-ddl=none。3. 手动建表，建表时应使用数据加密配置的 cipherColumn,plainColumn 和 assistedQueryColumn 代替逻辑列。

8.5 DistSQL

8.5.1 DistSQL 使用 DistSQL 添加数据源时，如何设置自定义的 JDBC 连接参数或连接池属性？

回答：

1. 如需自定义 JDBC 参数，请使用 `urlSource` 的方式定义 `dataSource`。
2. ShardingSphere 预置了必要的连接池参数，如 `maxPoolSize`、`idleTimeout` 等。如需增加或覆盖参数配置，请在 `dataSource` 中通过 `PROPERTIES` 指定。
3. 以上规则请参考 [相关介绍](#)。

8.5.2 DistSQL 使用 DistSQL 删除资源时，出现 `Resource [xxx] is still used by [SingleTableRule]`。

回答：

1. 被规则引用的资源将无法被删除。
2. 若资源只被 `single table rule` 引用，且用户确认可以忽略该限制，则可以添加可选参数 `ignore single tables` 进行强制删除。

8.5.3 DistSQL 使用 DistSQL 添加资源时，出现 `Failed to get driver instance for jdbcURL=xxx`。

回答：

ShardingSphere-Proxy 在部署过程中没有添加 `jdbc` 驱动，需要将 `jdbc` 驱动放入 ShardingSphere-Proxy 解压后的 `ext-lib` 目录，例如：`mysql-connector`。

8.6 其他

8.6.1 其他如果 SQL 在 ShardingSphere 中执行不正确，该如何调试？

回答：

在 ShardingSphere-Proxy 以及 ShardingSphere-JDBC 1.5.0 版本之后提供了 `sql.show` 的配置，可以将解析上下文和改写后的 SQL 以及最终路由至的数据源的细节信息全部打印至 `info` 日志。`sql.show` 配置默认关闭，如果需要请通过配置开启。> 注意：5.x 版本以后，`sql.show` 参数调整为 `sql-show`。

8.6.2 其他阅读源码时为什么会出现编译错误? IDEA 不索引生成的代码?

回答:

ShardingSphere 使用 lombok 实现极简代码。关于更多使用和安装细节, 请参考 [lombok 官网](https://lombok.org)。org.apache.shardingsphere.sql.parser.autogen 包下的代码由 ANTLR 生成, 可以执行以下命令快速生成:

```
./mvnw -Dcheckstyle.skip=true -Drat.skip=true -Dmaven.javadoc.skip=true -Djacoco.skip=true -DskipITs -DskipTests install -T1C
```

生成的代码例如 org.apache.shardingsphere.sql.parser.autogen.PostgreSQLStatementParser 等 Java 文件由于较大, 默认配置的 IDEA 可能不会索引该文件。可以调整 IDEA 的属性: idea.max.intellisense.filesize=10000。

8.6.3 其他使用 SQLSever 和 PostgreSQL 时, 聚合列不加别名会抛异常?

回答:

SQLServer 和 PostgreSQL 获取不加别名的聚合列会改名。例如, 如下 SQL:

```
SELECT SUM(num), SUM(num2) FROM tablexxx;
```

SQLServer 获取到的列为空字符串和 (2), PostgreSQL 获取到的列为空 sum 和 sum(2)。这将导致 ShardingSphere 在结果归并时无法找到相应的列而出错。正确的 SQL 写法应为:

```
SELECT SUM(num) AS sum_num, SUM(num2) AS sum_num2 FROM tablexxx;
```

8.6.4 其他 Oracle 数据库使用 Timestamp 类型的 Order By 语句抛出异常提示 “Order by value must implements Comparable” ?

回答:

针对上面问题解决方式有两种: 1. 配置启动 JVM 参数 “-oracle.jdbc.J2EE13Compliant=true” 2. 通过代码在项目初始化时设置 System.getProperties().setProperty(“oracle.jdbc.J2EE13Compliant”, “true”); 原因如下: org.apache.shardingsphere.sharding.merge.dql.orderby.OrderByValue#getOrderValues() 方法如下:

```
private List<Comparable<?>> getOrderValues() throws SQLException {
    List<Comparable<?>> result = new ArrayList<>(orderByItems.size());
    for (OrderItem each : orderByItems) {
        Object value = resultSet.getObject(each.getIndex());
        Preconditions.checkState(null == value || value instanceof Comparable,
            "Order by value must implements Comparable");
        result.add((Comparable<?>) value);
    }
    return result;
}
```

使用了 `resultSet.getObject(int index)` 方法, 针对 `TimeStamp oracle` 会根据 `oracle.jdbc.J2EE13Compliant` 属性判断返回 `java.sql.TimeStamp` 还是自定义 `oralce.sql.TIMESTAMP` 详见 `ojdbc` 源码 `oracle.jdbc.driver.TimestampAccessor#getObject(int var1)` 方法:

```
Object getObject(int var1) throws SQLException {
    Object var2 = null;
    if(this.rowSpaceIndicator == null) {
        DatabaseError.throwSQLException(21);
    }
    if(this.rowSpaceIndicator[this.indicatorIndex + var1] != -1) {
        if(this.externalType != 0) {
            switch(this.externalType) {
                case 93:
                    return this.getTimestamp(var1);
                default:
                    DatabaseError.throwSQLException(4);
                    return null;
            }
        }
        if(this.statement.connection.j2ee13Compliant) {
            var2 = this.getTimestamp(var1);
        } else {
            var2 = this.getTIMESTAMP(var1);
        }
    }
    return var2;
}
```

8.6.5 其他 Windows 环境下, 通过 Git 克隆 ShardingSphere 源码时为什么提示文件名过长, 如何解决?

回答:

为保证源码的可读性, ShardingSphere 编码规范要求类、方法和变量的命名要做到顾名思义, 避免使用缩写, 因此可能导致部分源码文件命名较长。由于 Windows 版本的 Git 是使用 `msys` 编译的, 它使用了旧版本的 Windows Api, 限制文件名不能超过 260 个字符。解决方案如下: 打开 `cmd.exe` (你需要将 `git` 添加到环境变量中) 并执行下面的命令, 可以让 `git` 支持长文件名:

```
git config --global core.longpaths true
```

如果是 Windows 10, 还需要通过注册表或组策略, 解除操作系统的文件名长度限制 (需要重启) : > 在注册表编辑器中创建 `HKLM\SYSTEM\CurrentControlSet\Control\FileSystem LongPathsEnabled`, 类型为 `REG_DWORD`, 并设置为 1。> 或者从系统菜单点击设置图标, 输入“编辑组策略”, 然后在打开的窗口依次进入“计算机管理”>“管理模板”>“系统”>“文件系统”, 在右侧双击“启用 win32 长路径”。参考资料: <https://docs.microsoft.com/zh-cn/windows/desktop/FileIO/naming-a-file>

<https://ourcodeworld.com/articles/read/109/how-to-solve-filename-too-long-error-in-git-powershell-and-github-application-for-windows>

8.6.6 其他 **Type is required** 异常的解决方法?

回答:

ShardingSphere 中很多功能实现类的加载方式是通过 **SPI** 注入的方式完成的, 如分布式主键, 注册中心等; 这些功能通过配置中 **type** 类型来寻找对应的 SPI 实现, 因此必须在配置文件中指定类型。

8.6.7 其他服务启动时如何加快 **metadata** 加载速度?

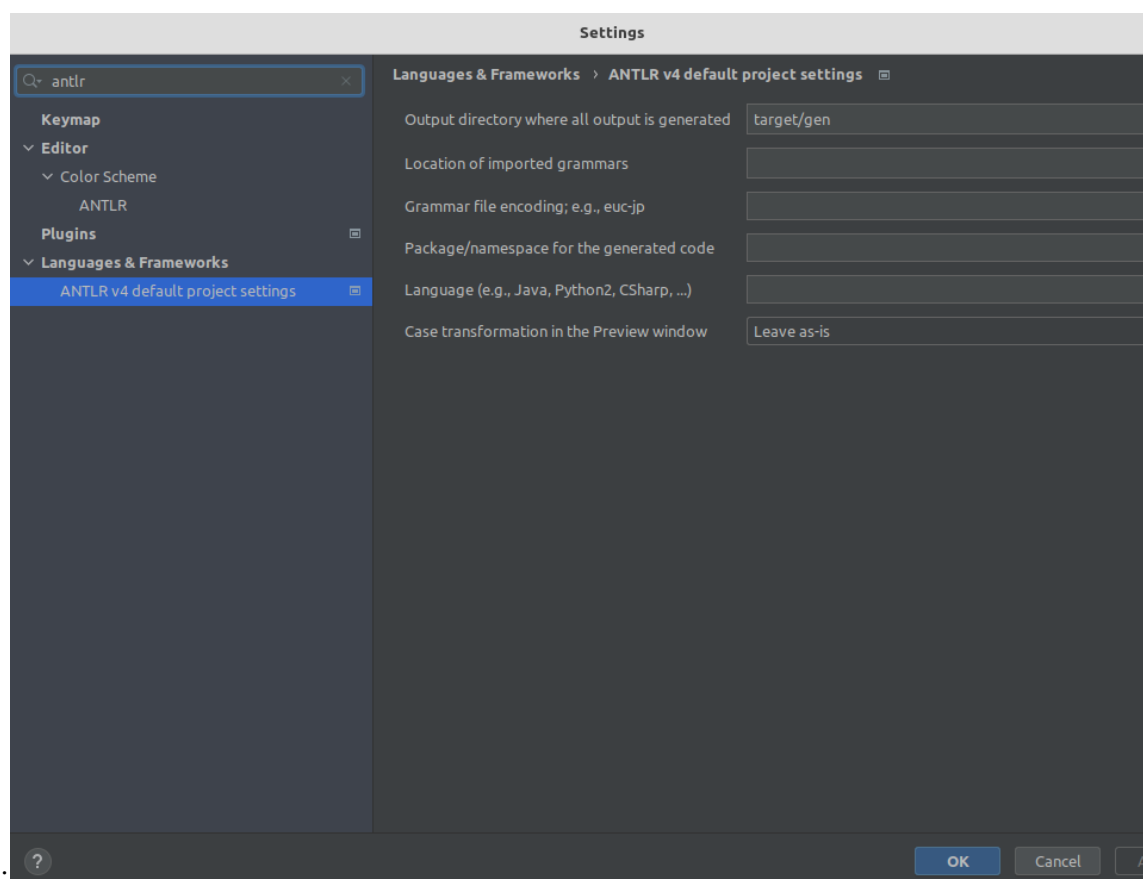
回答:

1. 升级到 4.0.1 以上的版本, 以提高 **metadata** 的加载速度。
2. 参照你采用的连接池, 将:
 - 配置项 `max.connections.size.per.query` (默认值为 1) 调高 (版本 $\geq 3.0.0.M3$ 且低于 5.0.0)。
 - 配置项 `max-connections-size-per-query` (默认值为 1) 调高 (版本 $\geq 5.0.0$)。

8.6.8 其他 **ANTLR** 插件在 **src** 同级目录下生成代码, 容易误提交, 如何避免?

回答:

进入 Settings -> Languages & Frameworks -> ANTLR v4 default project settings 配置生成代码的输出目录



为 target/gen,如图:

8.6.9 其他使用 Proxool 时分库结果不正确?

回答:

使用 Proxool 配置多个数据源时, 应该为每个数据源设置 alias, 因为 Proxool 在获取连接时会判断连接池中是否包含已存在的 alias, 不配置 alias 会造成每次都只从一个数据源中获取连接。以下是 Proxool 源码中 ProxoolDataSource 类 getConnection 方法的关键代码:

```
if(!ConnectionPoolManager.getInstance().isPoolExists(this.alias)) {
    this.registerPool();
}
```

更多关于 alias 使用方法请参考 [Proxool 官网](#)。PS: sourceforge 网站需要翻墙访问。

8.6.10 其他使用 Spring Boot 2.x 集成 ShardingSphere 时, 配置文件中的属性设置不生效?

回答:

需要特别注意, Spring Boot 2.x 环境下配置文件的属性名称约束为仅允许小写字母、数字和短横线, 即 [a-z][0-9] 和 -。原因如下: Spring Boot 2.x 环境下, ShardingSphere 通过 Binder 来绑定配置文件, 属性名称不规范(如: 驼峰或下划线等)会导致属性设置不生效从而校验属性值时抛出 NullPointerException 异常。参考以下错误示例: 下划线示例: database_inline

```
spring.shardingsphere.rules.sharding.sharding-algorithms.database_inline.  
type=INLINE  
spring.shardingsphere.rules.sharding.sharding-algorithms.database_inline.props.  
algorithm-expression=ds-${user_id % 2}
```

Caused by: org.springframework.beans.factory.BeanCreationException: Error creating bean with name 'database_inline': Initialization of bean failed; nested exception is java.lang.NullPointerException: Inline sharding algorithm expression cannot be null.

...

Caused by: java.lang.NullPointerException: Inline sharding algorithm expression cannot be null.

```
    at com.google.common.base.Preconditions.checkNotNull(Preconditions.  
java:897)  
    at org.apache.shardingsphere.sharding.algorithm.sharding.inline.  
InlineShardingAlgorithm.getAlgorithmExpression(InlineShardingAlgorithm.java:58)  
    at org.apache.shardingsphere.sharding.algorithm.sharding.inline.  
InlineShardingAlgorithm.init(InlineShardingAlgorithm.java:52)  
    at org.apache.shardingsphere.spring.boot.registry.  
AbstractAlgorithmProvidedBeanRegistry.  
postProcessAfterInitialization(AbstractAlgorithmProvidedBeanRegistry.java:98)  
    at org.springframework.beans.factory.support.  
AbstractAutowireCapableBeanFactory.  
applyBeanPostProcessorsAfterInitialization(AbstractAutowireCapableBeanFactory.  
java:431)  
    ...
```

驼峰示例: databaseInline

```
spring.shardingsphere.rules.sharding.sharding-algorithms.databaseInline.type=INLINE  
spring.shardingsphere.rules.sharding.sharding-algorithms.databaseInline.props.  
algorithm-expression=ds-${user_id % 2}
```

Caused by: org.springframework.beans.factory.BeanCreationException: Error creating bean with name 'databaseInline': Initialization of bean failed; nested exception is java.lang.NullPointerException: Inline sharding algorithm expression cannot be null.

...

Caused by: java.lang.NullPointerException: Inline sharding algorithm expression cannot be null.

```
    at com.google.common.base.Preconditions.checkNotNull(Preconditions.  
java:897)  
    at org.apache.shardingsphere.sharding.algorithm.sharding.inline.  
InlineShardingAlgorithm.getAlgorithmExpression(InlineShardingAlgorithm.java:58)  
    at org.apache.shardingsphere.sharding.algorithm.sharding.inline.  
InlineShardingAlgorithm.init(InlineShardingAlgorithm.java:52)  
    at org.apache.shardingsphere.spring.boot.registry.  
AbstractAlgorithmProvidedBeanRegistry.
```

```
postProcessAfterInitialization(AbstractAlgorithmProvidedBeanRegistry.java:98)
    at org.springframework.beans.factory.support.
AbstractAutowireCapableBeanFactory.
applyBeanPostProcessorsAfterInitialization(AbstractAutowireCapableBeanFactory.
java:431)
    ...
```

从异常堆栈中分析可知：AbstractAlgorithmProvidedBeanRegistry.registerBean 方法调用 PropertyUtil.containsPropertyPrefix(environment, prefix) 方法判断指定前缀 prefix 的配置是否存在, 而 PropertyUtil.containsPropertyPrefix(environment, prefix) 方法, 在 Spring Boot 2.x 环境下使用了 Binder, 不规范的属性名称（如：驼峰或下划线等）会导致属性设置不生效。

9.1 最新版本

Apache ShardingSphere 的发布版包括源码包及其对应的二进制包。由于下载内容分布在镜像服务器上，所以下载后应该进行 GPG 或 SHA-512 校验，以此来保证内容没有被篡改。

9.1.1 Apache ShardingSphere - 版本: 5.2.0 (发布日期: Sept 5th, 2022)

- 源码: [SRC](#) ([ASC](#), [SHA512](#))
- ShardingSphere-JDBC 二进制包: [TAR](#) ([ASC](#), [SHA512](#))
- ShardingSphere-Proxy 二进制包: [TAR](#) ([ASC](#), [SHA512](#))
- ShardingSphere-Agent 二进制包: [TAR](#) ([ASC](#), [SHA512](#))

9.2 全部版本

全部版本请到 [Archive repository](#) 查看。全部孵化器版本请到 [Archive incubator repository](#) 查看。

9.3 校验版本

PGP 签名文件

使用 PGP 或 SHA 签名验证下载文件的完整性至关重要。可以使用 GPG 或 PGP 验证 PGP 签名。请下载 KEYS 以及发布的 asc 签名文件。建议从主发布目录而不是镜像中获取这些文件。

```
gpg -i KEYS
```

或者


```
pgpk -a KEYS
```

或者

```
pgp -ka KEYS
```

要验证二进制文件或源代码，您可以从主发布目录下载相关的 asc 文件，并按照以下指南进行操作。

```
gpg --verify apache-shardingsphere-*****.asc apache-shardingsphere-*****
```

或者

```
pgpv apache-shardingsphere-*****.asc
```

或者

```
pgp apache-shardingsphere-*****.asc
```