



INNOVATE

ONLINE CONFERENCE

分会场三：大数据分析

碰撞出不一样的火花 - 如何在云中更好地使用 Apache Spark

方浩，AWS 解决方案架构师

日程

- Amazon EMR 中的 Apache Spark
- Apache Spark 与 Amazon SageMaker
- AWS Glue 中的 Apache Spark

Amazon EMR 中的 Apache Spark

© 2020, Amazon Web Services, Inc. or its affiliates. All rights reserved.

 **INNOVATE**
ONLINE CONFERENCE

AWS 中国（宁夏）区域由西云数据运营
AWS 中国（北京）区域由光环新网运营

Amazon EMR

应用清单

- Amazon EMR 是全球最大的 Spark 和 Hadoop 服务提供者之一，让客户能够以 PB 规模运行 ETL、机器学习、实时处理、数据科学和低延迟 SQL
- 请参阅 EMR 发布页面了解最新的计划。 用户也可以加载下面未列出的应用程序。

管理

- | | |
|-------------|-------------------|
| • Ganglia | 集群监控 |
| • Livy | Spark REST API 接口 |
| • Oozie | 工作流编排器 |
| • ZooKeeper | 配置和同步节点 |

机器学习

- | | |
|--------------|------|
| • Mahout | 机器学习 |
| • MXNet | 深度学习 |
| • SparkML | 机器学习 |
| • TensorFlow | 深度学习 |

数据迁移

- | | |
|---------|-----------|
| • Sqoop | 关系型数据库连接器 |
|---------|-----------|

NoSQL

- | | |
|---------|------------------|
| • HBase | Hadoop 生态非关系型数据库 |
|---------|------------------|

数据处理

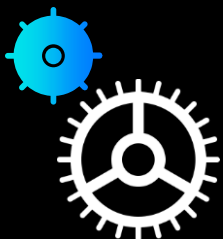
- | | |
|-------------|-------------|
| • Flink | 流处理 |
| • Hive | 数据仓库 |
| • MapReduce | 批量数据处理 |
| • Presto | 分布式 SQL 查询 |
| • Spark | 内存数据处理/机器学习 |
| • Tez | 交互式数据处理 |
| • Hudi | 更新数据湖存储 |

查询工具

- | | |
|-----------------|---------------------------|
| • EMR Notebooks | 无服务器 Notebook |
| • Hue | 可视化和查询 |
| • JupyterHub | 多用户 Jupyter Notebook（集群内） |
| • Phoenix | Hbase 查询 |
| • Pig | 脚本语言 |
| • Zeppelin | 面向数据科学家的 Notebook |



Well-Architected Framework 的五大要素



卓越运营



安全性



可靠性



性能效率



成本优化

卓越运营

© 2020, Amazon Web Services, Inc. or its affiliates. All rights reserved.

 **INNOVATE**
ONLINE CONFERENCE

AWS 中国（宁夏）区域由西云数据运营
AWS 中国（北京）区域由光环新网运营

丰富的任务提交方式

- Spark Submit
- Spark Shell
- Apache Oozie
- JupyterHub Notebook
- Apache Zeppelin Notebook
- R Studio
- Apache Livy
- Amazon EMR Step API
- AWS Step Functions
- AWS Data Pipeline;
- Apache Airflow

离线存储的 Spark History Service

Amazon EMR

Clusters

Security configurations

Block public access

VPC subnets

Events

Notebooks

Git repositories

Help

What's new

Clone Terminate AWS CLI export

Cluster: DO NOT TERMINATE Parag Test Terminated Terminated by user request

Summary Application history Monitoring Hardware Configurations Events Steps Bootstrap actions

Connections: --

Master public DNS: ec2-3-90-84-213.compute-1.amazonaws.com [SSH](#)

History service: [Spark history server UI](#) (SSH tunneling not required)

Tags: --

Summary

ID: j-3ODPTHGROJPUA

Creation date: 2019-10-25 13:30 (UTC-8)

End date: 2019-11-24 18:17 (UTC-8)

Elapsed time: 4 weeks

After last step Cluster waits completes:

Termination protection: Off

Configuration details

Release label: emr-5.27.0

Hadoop distribution: Amazon 2.8.5

Applications: Spark 2.4.4, Livy 0.6.0, Hive 2.3.5, Zeppelin 0.8.1, JupyterHub 1.0.0

Log URI: s3://aws-logs-418620260174-us-east-1/elasticmapreduce/

EMRFS consistent view: Disabled

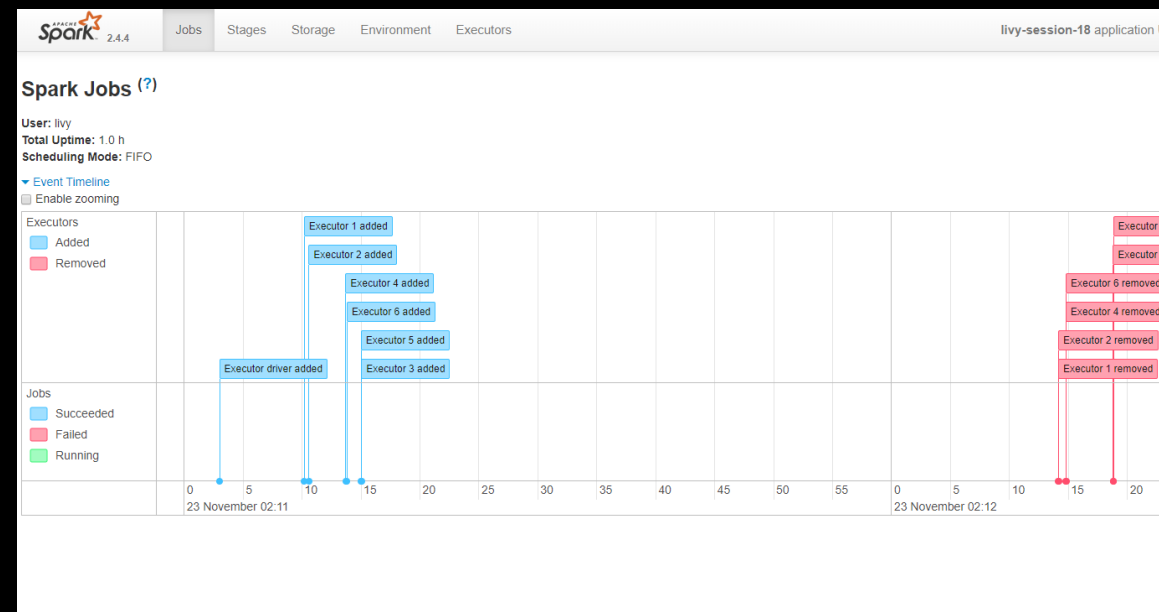
Custom AMI ID: --

Network and hardware

Availability zone: us-east-1e

Security and access

Key name: paragnc-emr-dev-new-us-east-1



安全性

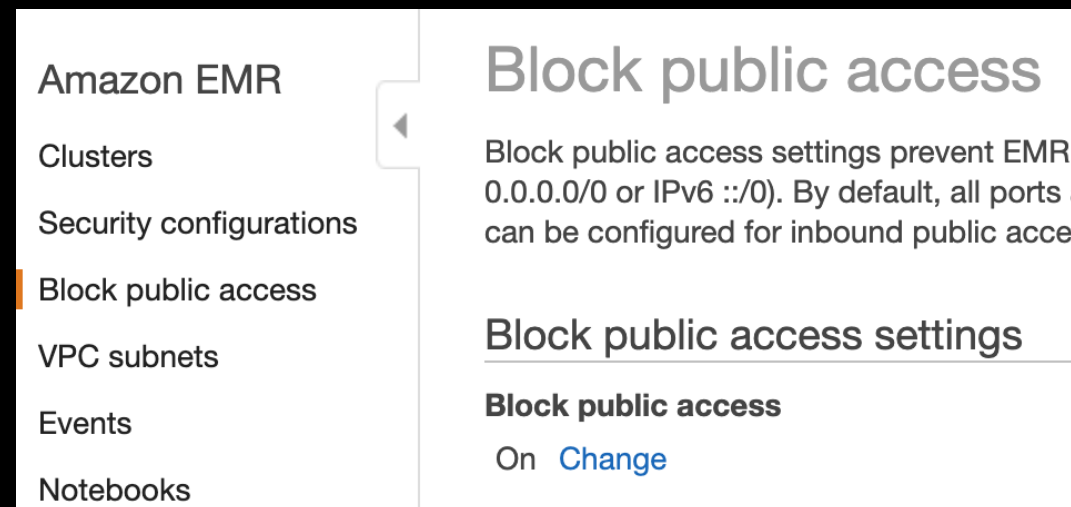
© 2020, Amazon Web Services, Inc. or its affiliates. All rights reserved.

 **INNOVATE**
ONLINE CONFERENCE

AWS 中国（宁夏）区域由西云数据运营
AWS 中国（北京）区域由光环新网运营

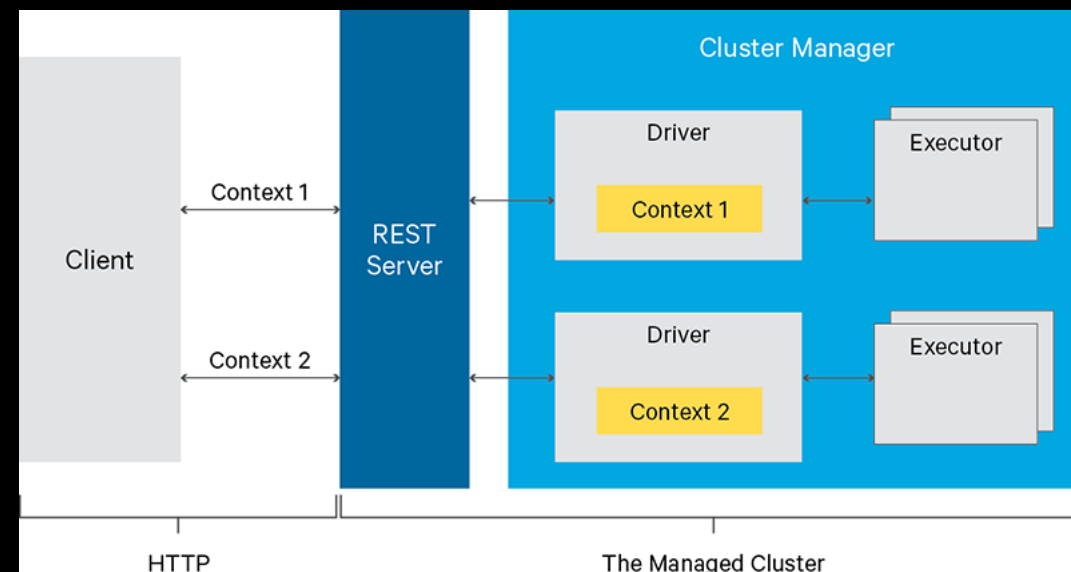
禁止公网访问

- 如果安全组中存在来自所有公网地址的规则，集群将拒绝启动；
- 我们建议您把集群放置于私有子网，如果需要访问外网的数据源，可以添加 NAT 实例；
- 如果是 AWS 中的数据源，比如 Amazon Kinesis、Amazon DynamoDB、Amazon S3，可以通过添加 VPC Endpoints 的方式启用内网访问；
- 可以按照集群管理页面的提示，添加一台 Amazon EC2 实例来提供 SSH 隧道，访问集群的各种 Web 页面



Apache Livy

- Amazon EMR 从 5.9.0 版开始集成 Livy;
- Livy 为 Spark 提供了 REST API 接口用于提交和管理任务, 用户无需 SSH 登录到主节点, 也无需部署 Spark 客户端环境;
- 编译过的 jar 包可以提前上传到 Amazon S3;
- 可以通过调用 Kerberos 来进行认证



可靠性

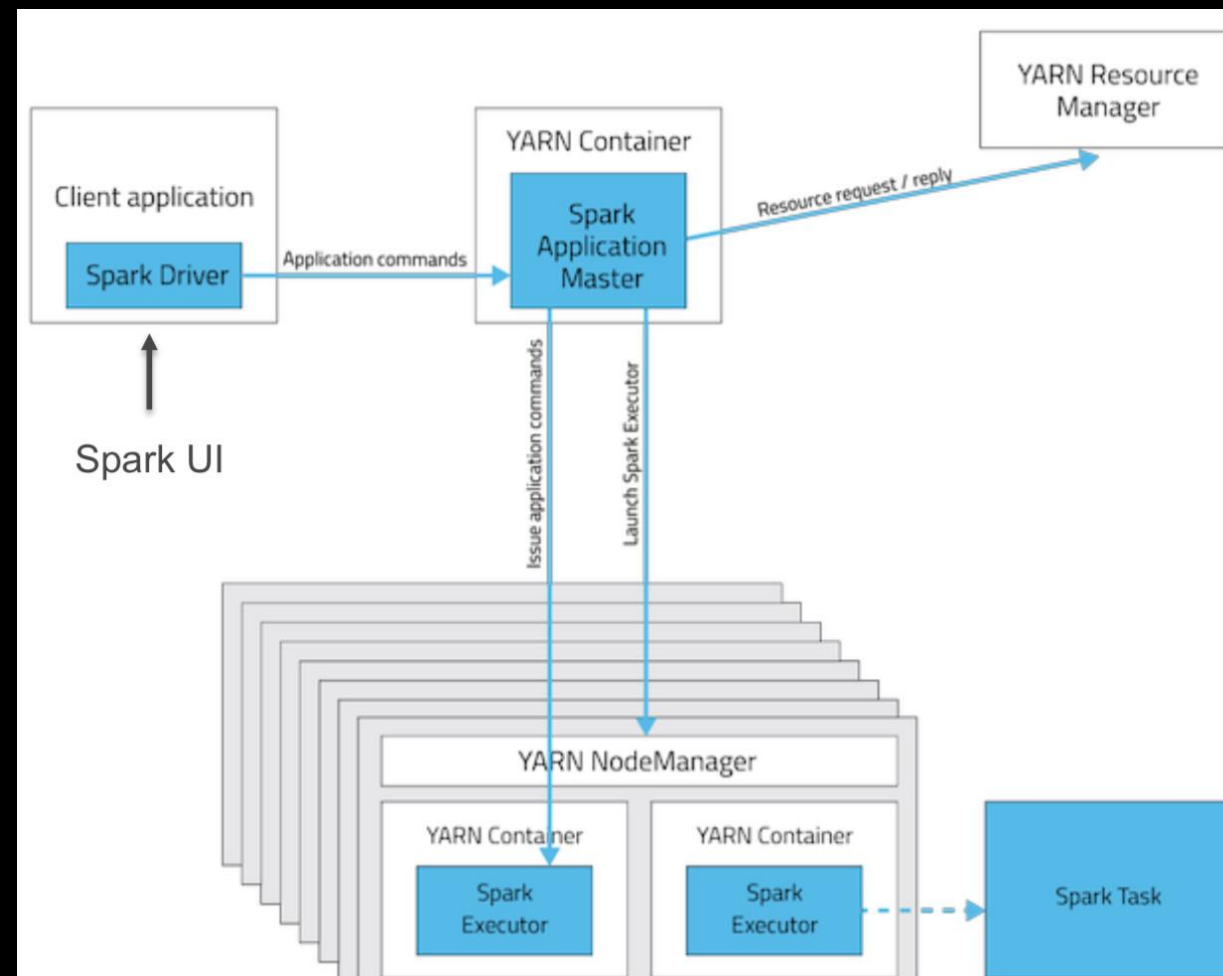
© 2020, Amazon Web Services, Inc. or its affiliates. All rights reserved.

aws **INNOVATE**
ONLINE CONFERENCE

AWS 中国（宁夏）区域由西云数据运营
AWS 中国（北京）区域由光环新网运营

Multiple Master

- 自 Amazon EMR 5.23.0 版开始支持;
- 主节点不再是潜在故障单点。如果某个主节点发生故障, Amazon EMR 将自动替换掉故障节点而不会影响集群正常运行;
- 以 YARN client 模式运行的 Spark 程序会消耗主节点资源, 多主节点也有助于缓解这一情况



性能效率

© 2020, Amazon Web Services, Inc. or its affiliates. All rights reserved.

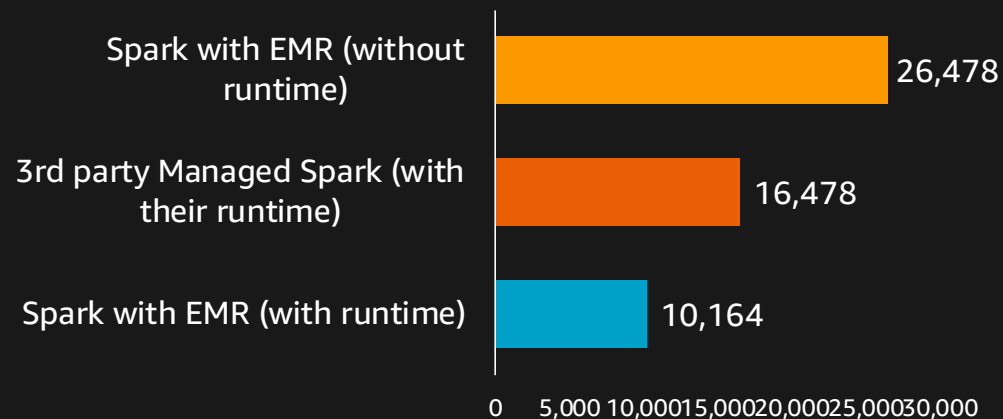
aws **INNOVATE**
ONLINE CONFERENCE

AWS 中国（宁夏）区域由西云数据运营
AWS 中国（北京）区域由光环新网运营

对 Amazon EMR 中的 Spark 进行性能优化

性能优化过后的 Apache Spark runtime, 以 10% 的成本, 取得 2.6 倍的性能提升

104 个查询的运行时总数 (秒-越低越好)



*基于 TPC-DS 3TB 性能测试, 6 节点 C4x8 xlarge 集群
(EMR 5.28, Spark 2.4)

构建在 Apache Spark 的优化版本上的 runtime

更佳性能

- 比没有 runtime 的 EMR 中的 Spark 快 2.6 倍
- 比第三方托管的 Spark (使用其 runtime) 快 1.6 倍

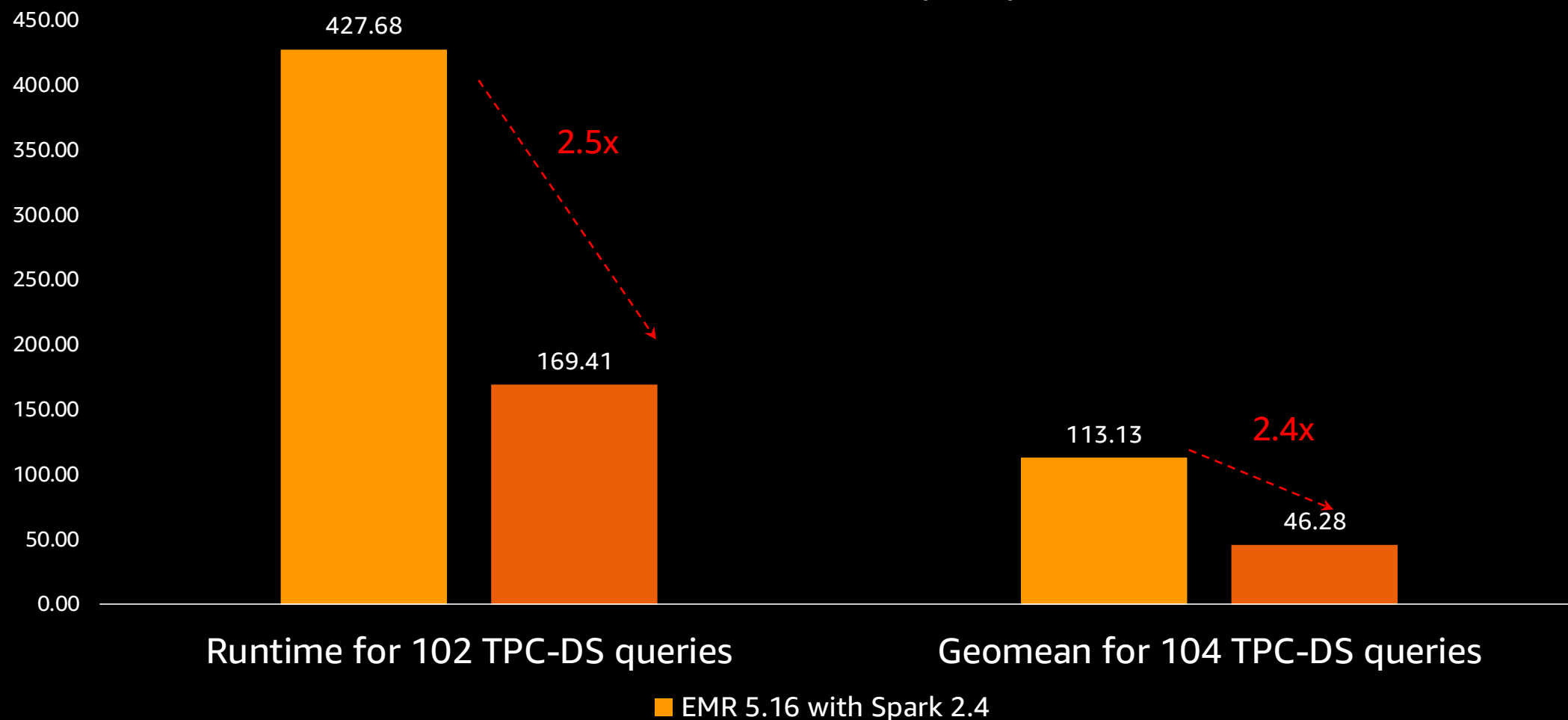
更低成本

- 是第三方托管的 Spark (使用其 runtime) 成本的 10%

100% 与 Apache Spark API 兼容

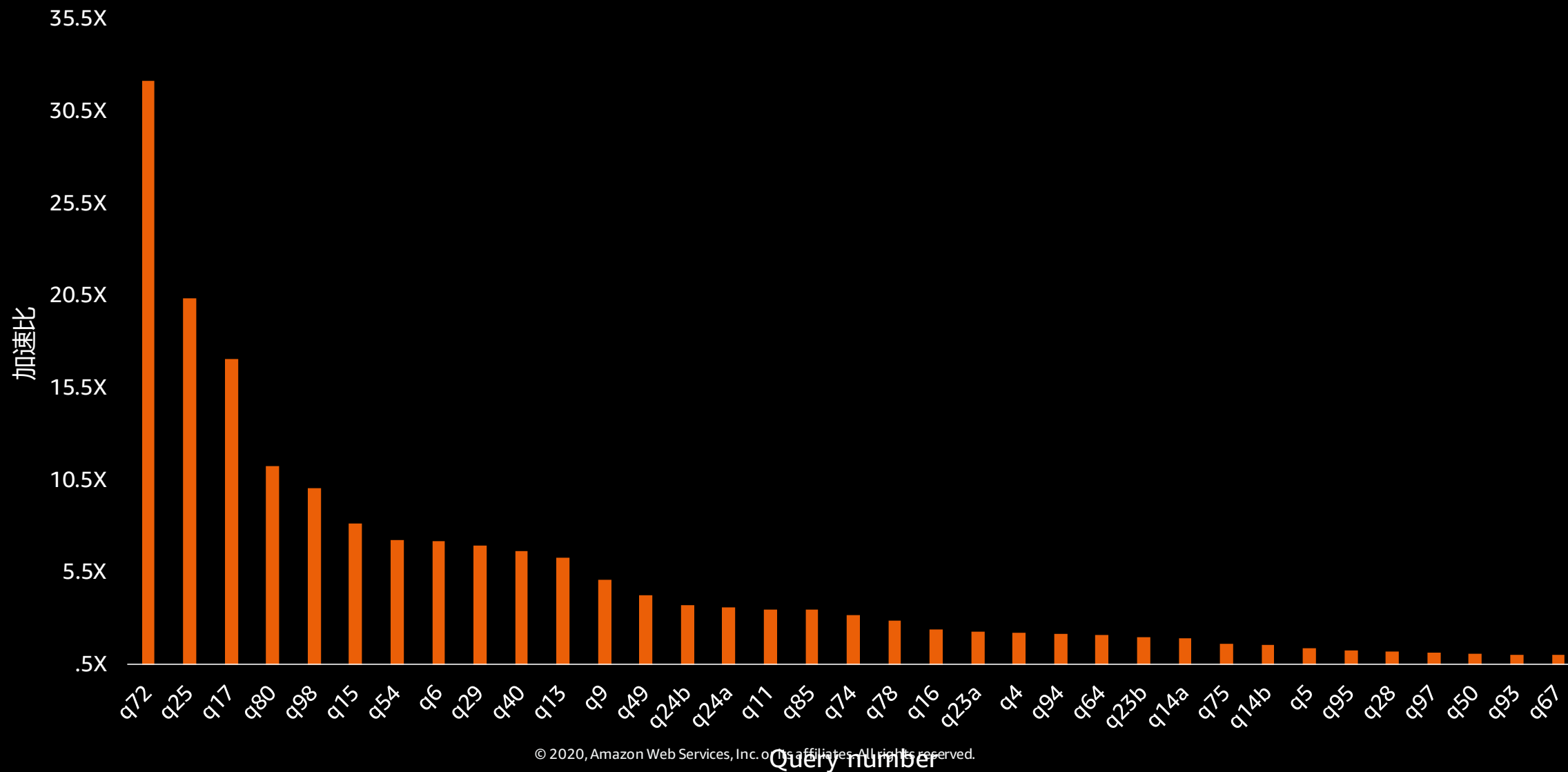
升级到最新版本有助于节省成本

自去年以来的提升 (分钟)



© 2020, Amazon Web Services, Inc. or its affiliates. All rights reserved.

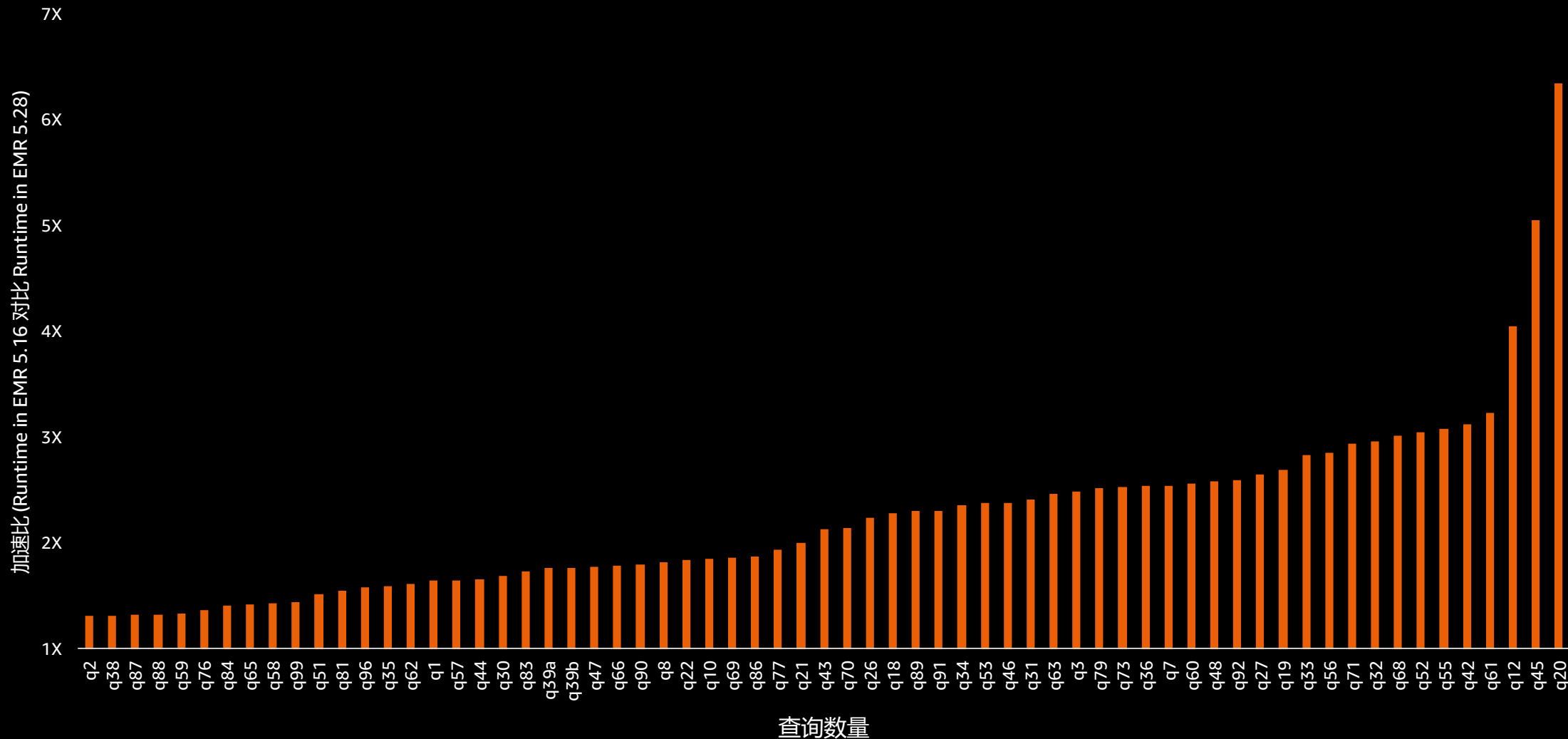
长查询平均速度提高 5 倍



© 2020, Amazon Web Services, Inc. or its affiliates. All rights reserved.

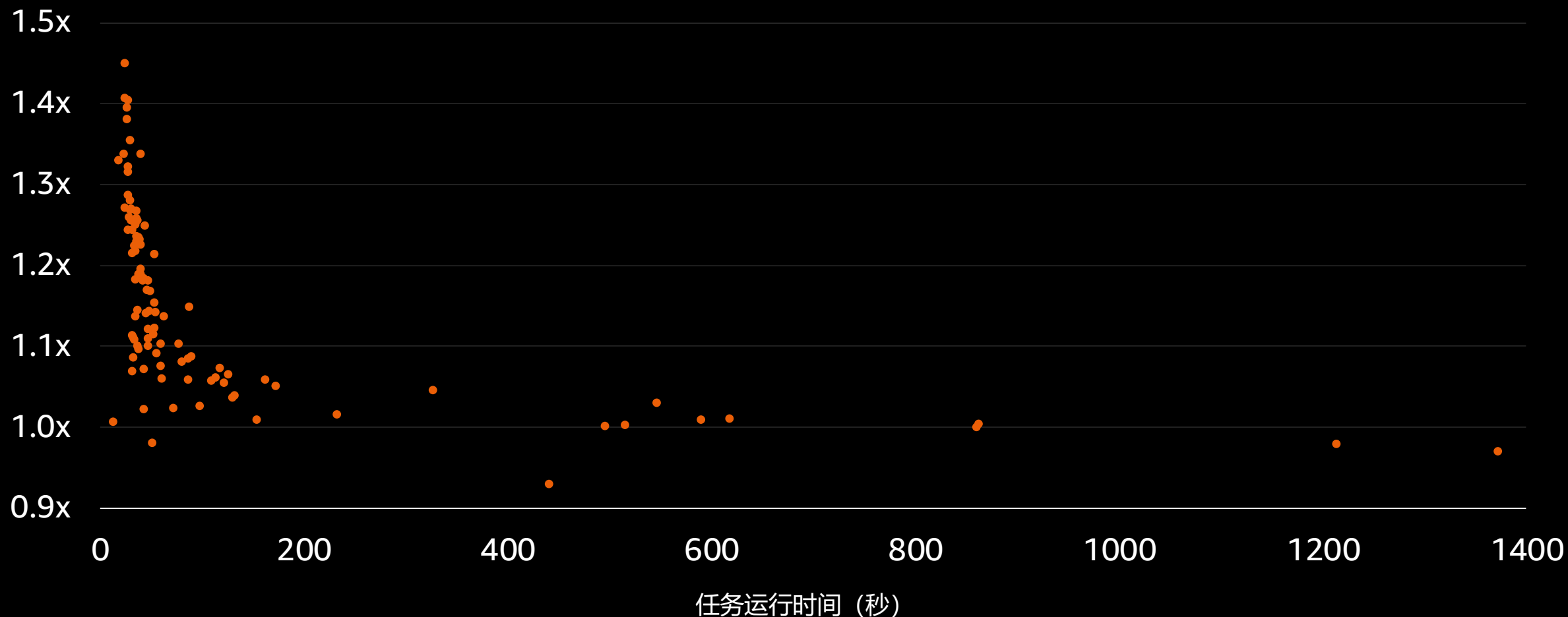
Query number

短查询的平均速度提高 2 倍



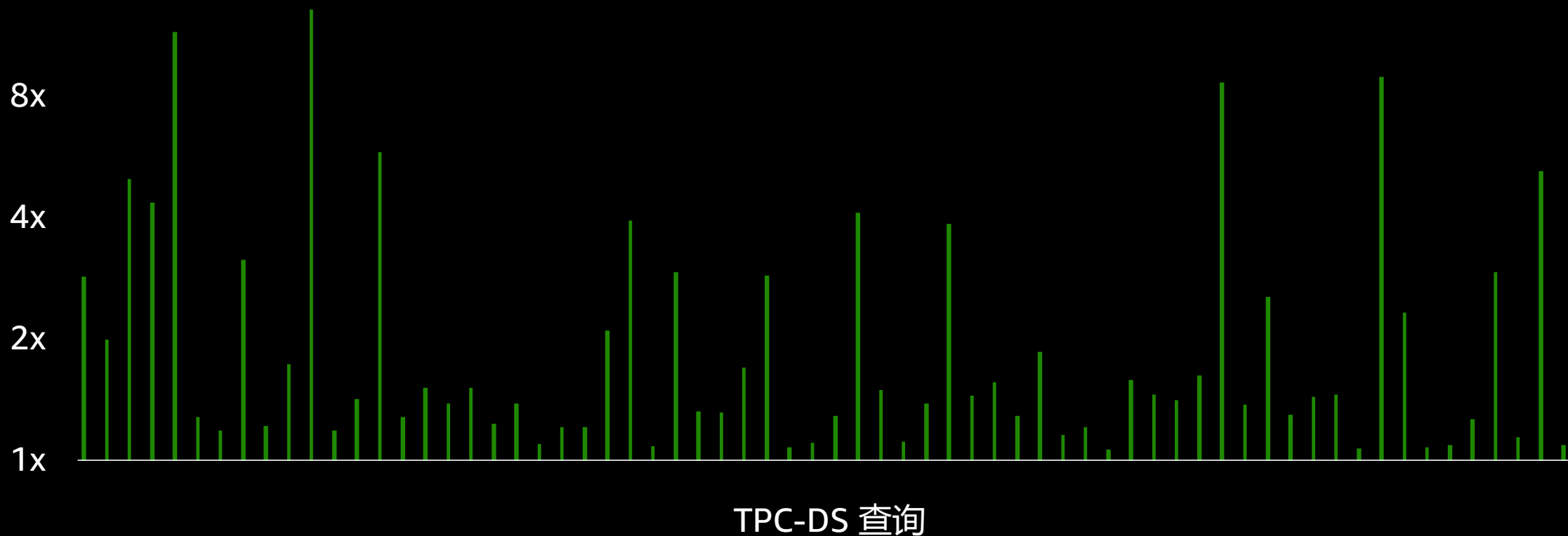
© 2020, Amazon Web Services, Inc. or its affiliates. All rights reserved.

工作启动——积极的工作节点分配

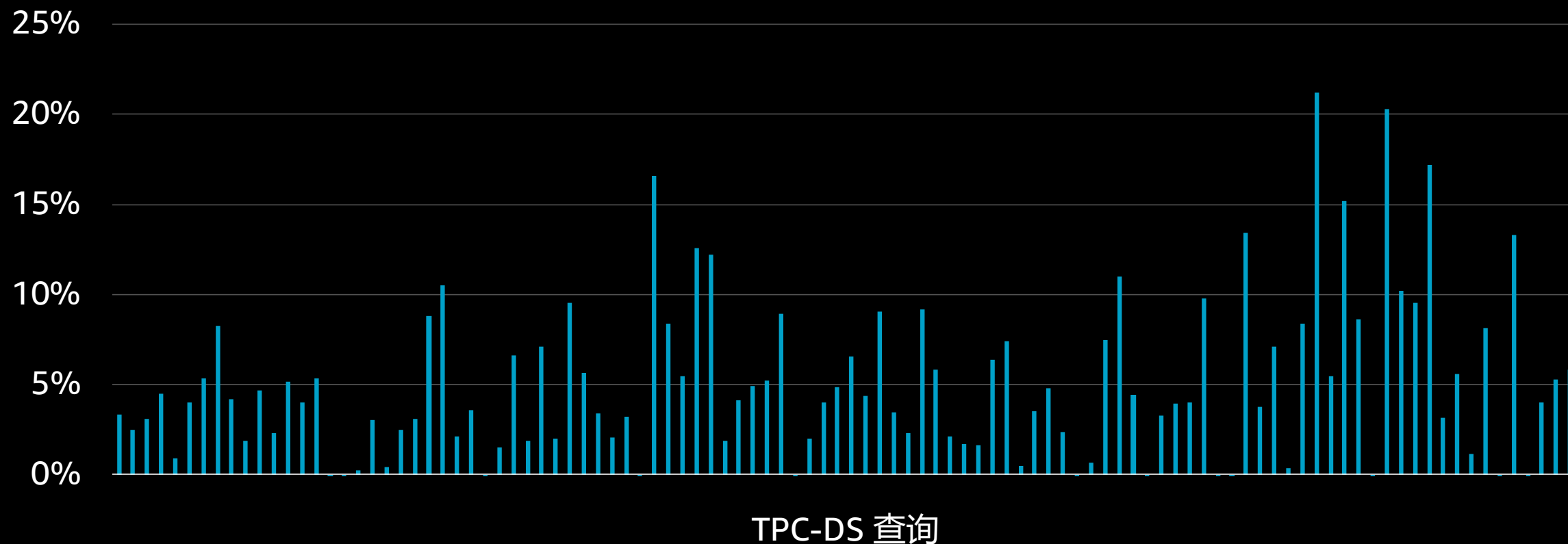


© 2020, Amazon Web Services, Inc. or its affiliates. All rights reserved.

规划/优化——动态分区修剪



查询执行——数据预取



成本优化

© 2020, Amazon Web Services, Inc. or its affiliates. All rights reserved.

 **INNOVATE**
ONLINE CONFERENCE

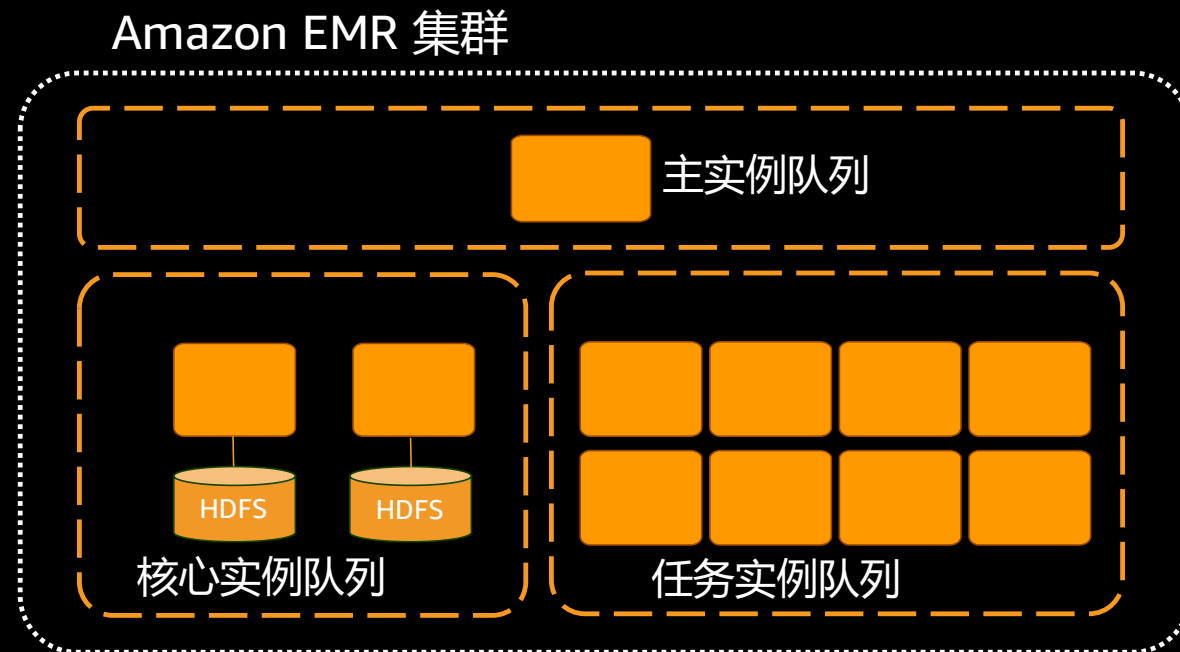
AWS 中国（宁夏）区域由西云数据运营
AWS 中国（北京）区域由光环新网运营

Amazon EMR 节点类型

主节点：管理群集节点。主节点跟踪任务的状态并监视群集的运行状况

核心节点：用于运行任务并在群集上的 Hadoop 分布式文件系统 (HDFS) 中存储数据的节点

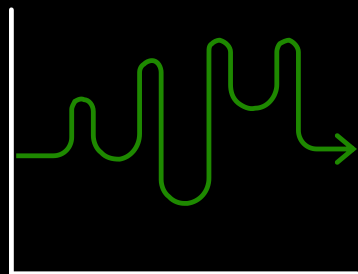
任务节点：仅运行任务，不在 HDFS 中存储数据的节点。任务节点是可选的



Amazon EC2 购买选择

按需

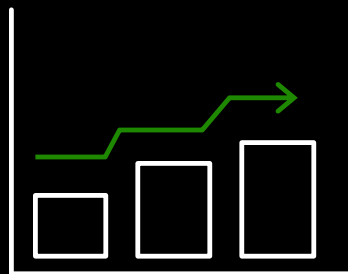
以秒为单位为计算容量付费
无需长期承诺



突发的工作负荷

预留

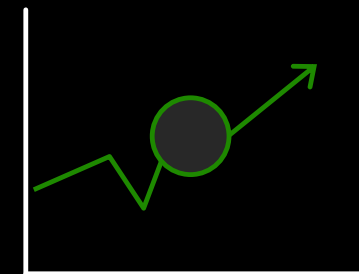
作出 1 年或 3 年的承诺
在按需价格的基础上获得显著折扣



承诺和稳定使用

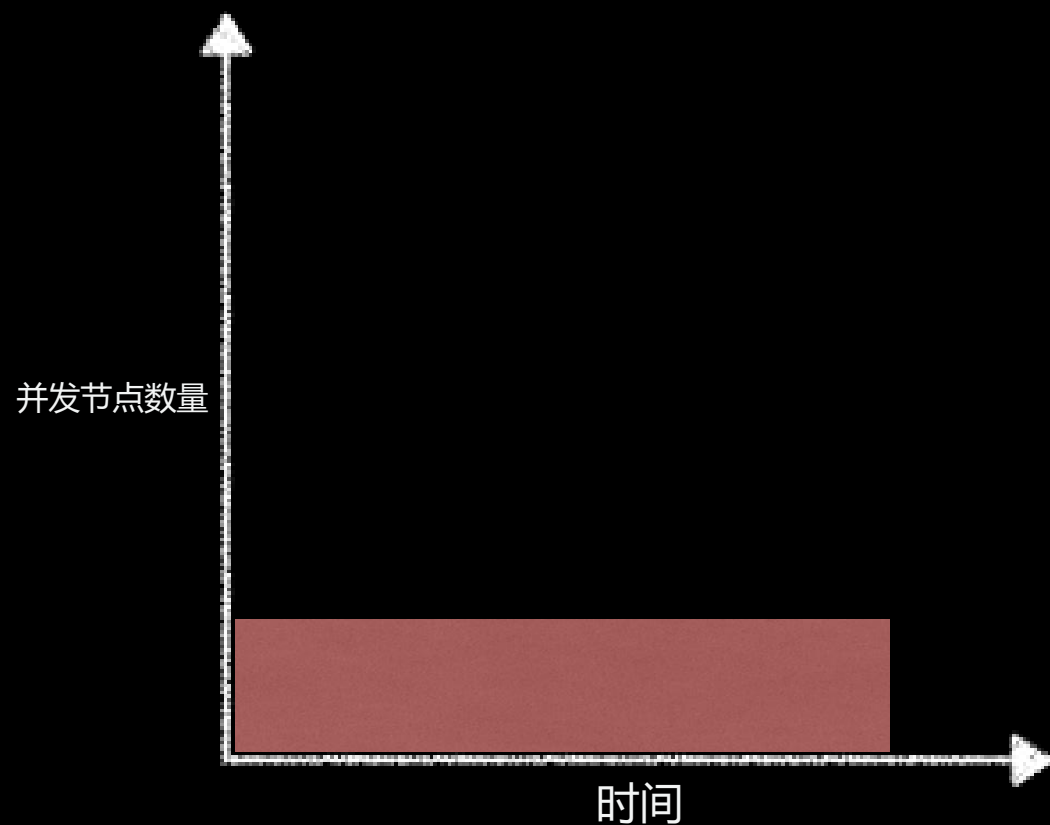
竞价

闲置 EC2 容量
相比按需价格，折扣高达 90%

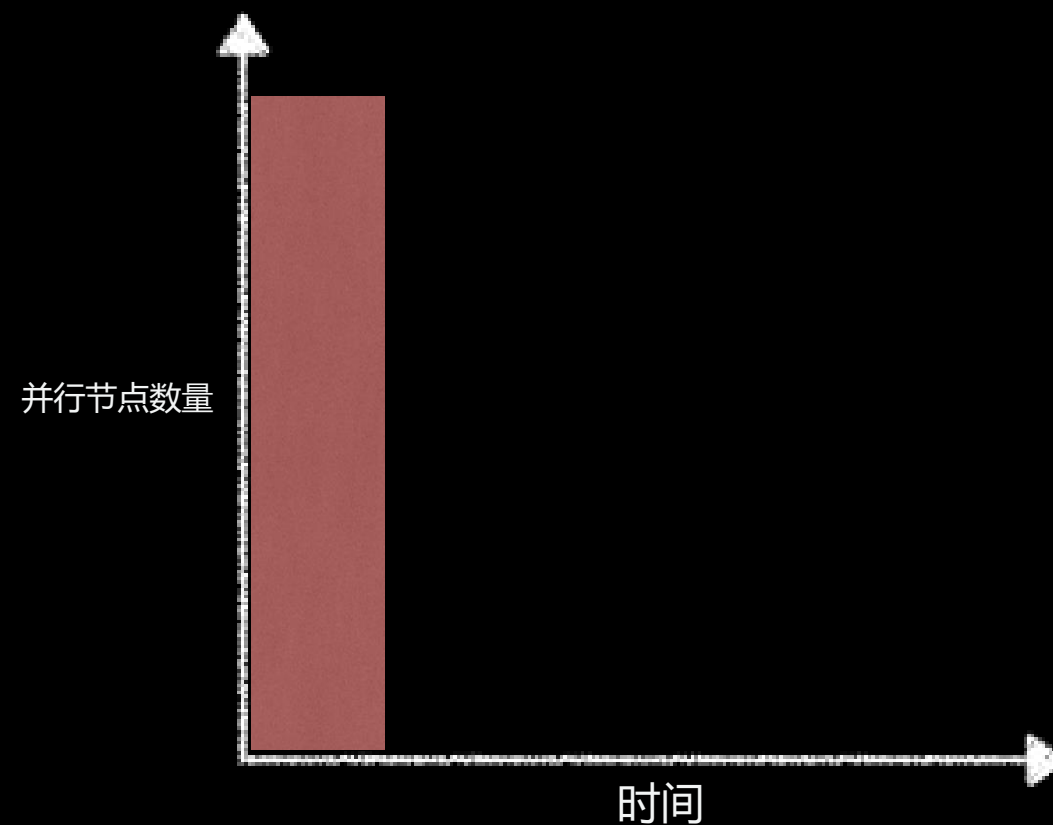


容错、灵活、无状态的工作负载

并行处理

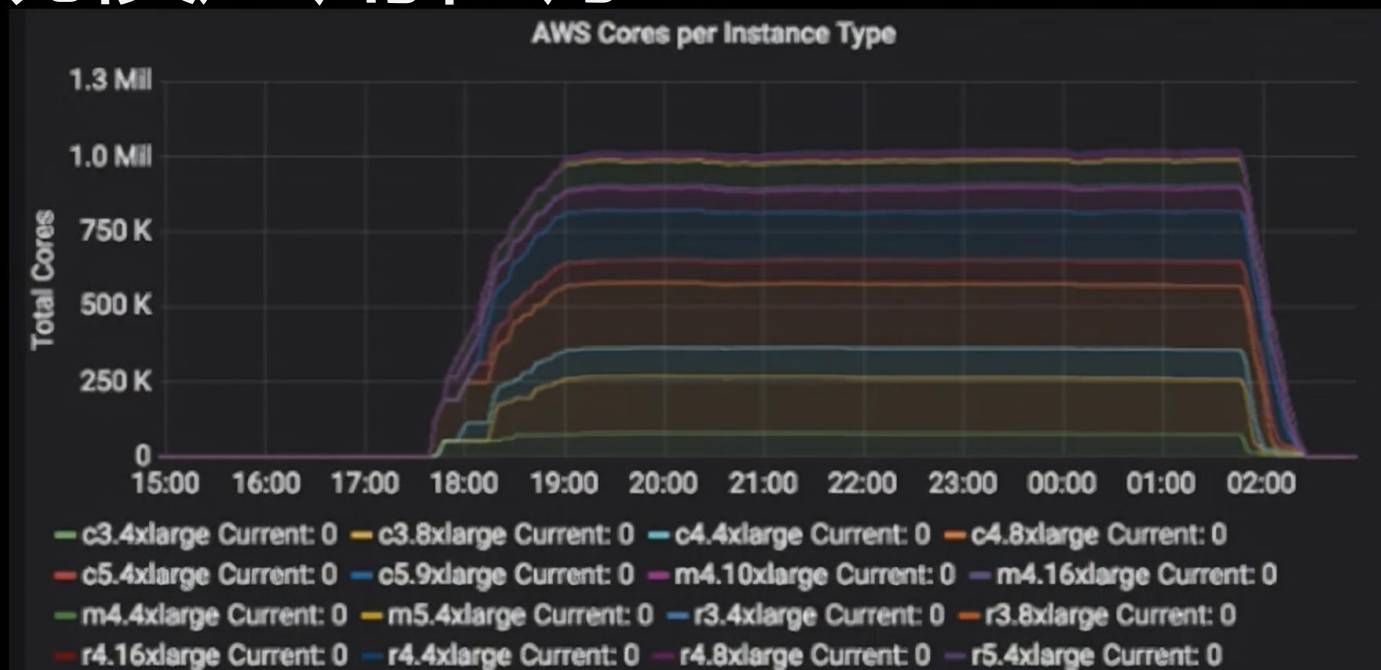


任务运行时间：10小时



任务运行时间：1小时

运行超大规模应用程序



Western Digital®

超过 230 万个模拟作业运行在基于 Amazon EC2 竞价实例构建的，由 **100 万个 vCPU** 组成的单个 HPC 集群
计算时间：从 20 天 → 8 小时

EC2 竞价池——实例类型灵活选择

C4	1a	1b	1c	On Demand
8XL	\$0.50	\$0.27	\$0.29	\$1.76
4XL	\$0.21	\$0.30	\$0.16	\$0.88
2XL	\$0.08	\$0.07	\$0.08	\$0.44
XL	\$0.04	\$0.05	\$0.04	\$0.22
L	\$0.01	\$0.01	\$0.04	\$0.11

每个实例类型

每个实例规格

每个可用区

每个区域

都是单独的竞价池

R5

M4

R4

I3

C4

M5d

D2

C5

R5d

竞价型实例顾问



Region: US East (N. Virginia)

OS: Linux/UNIX

Instance type filter:

vCPU (min): 1

Memory GiB (min): 0

☐ Instance types supported by EMR

Instance Type	vCPU	Memory GiB	Savings over On-Demand*	Frequency of interruption ▼
i3.xlarge	4	30.5	70%	<5% □□□□□
m3.large	2	7.5	78%	<5% □□□□□
r4.16xlarge	64	488	76%	<5% □□□□□
m4.4xlarge	16	64	70%	<5% □□□□□
i3.8xlarge	32	244	70%	<5% □□□□□

<https://aws.amazon.com/ec2/spot/instance-advisor/>

© 2020, Amazon Web Services, Inc. or its affiliates. All rights reserved.

Apache Spark 与 Amazon Sagemaker

© 2020, Amazon Web Services, Inc. or its affiliates. All rights reserved.



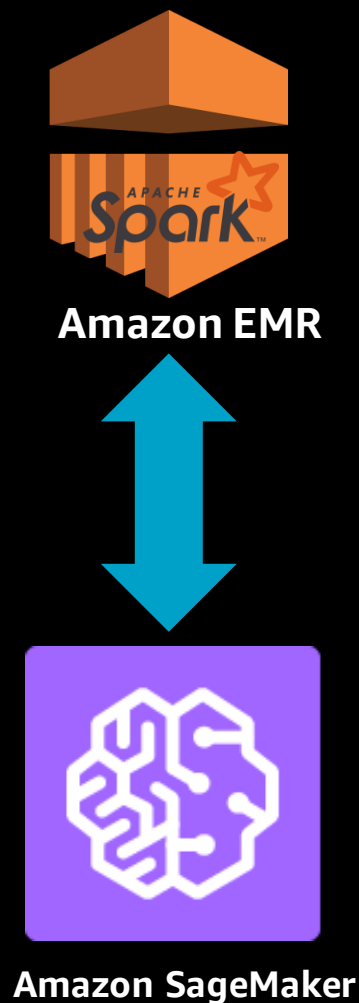
AWS 中国（宁夏）区域由西云数据运营
AWS 中国（北京）区域由光环新网运营

Amazon SageMaker



- 适用于机器学习/深度学习的完全托管服务
- 快速轻松地构建、训练和部署任何规模的模型
- 通过 Web 控制台或以编程方式使用的编排平台
- 内置云规模算法、深度学习框架，也支持自带算法
- 根据需要进行合适的计算资源，包括 GPU
- Notebook 可以从各种 Amazon Elastic Compute Cloud (Amazon EC2) 实例类型中进行选择
- 训练服务使用短暂计算资源生成模型，然后自动终止计算资源
- 模型托管可以使用 Auto Scaling 来动态调整大小，以响应客户端应用程序的推理需求

Apache Spark 与 Amazon SageMaker 结合的优势



健壮的机器学习和深度学习模型需要大量正确预处理的训练数据

一些有限的 ETL 可以在 Amazon SageMaker Notebook 中完成，但主要用于开发

Amazon EMR 中的 Spark 提供自定义 ETL 所需的大规模并行处理

Amazon SageMaker 提供按需、可扩展的分布式计算学习/深度学习训练和推理

Amazon SageMaker-Spark 库支持 Spark Amazon SageMaker 双向集成

Amazon SageMaker-Spark SDK

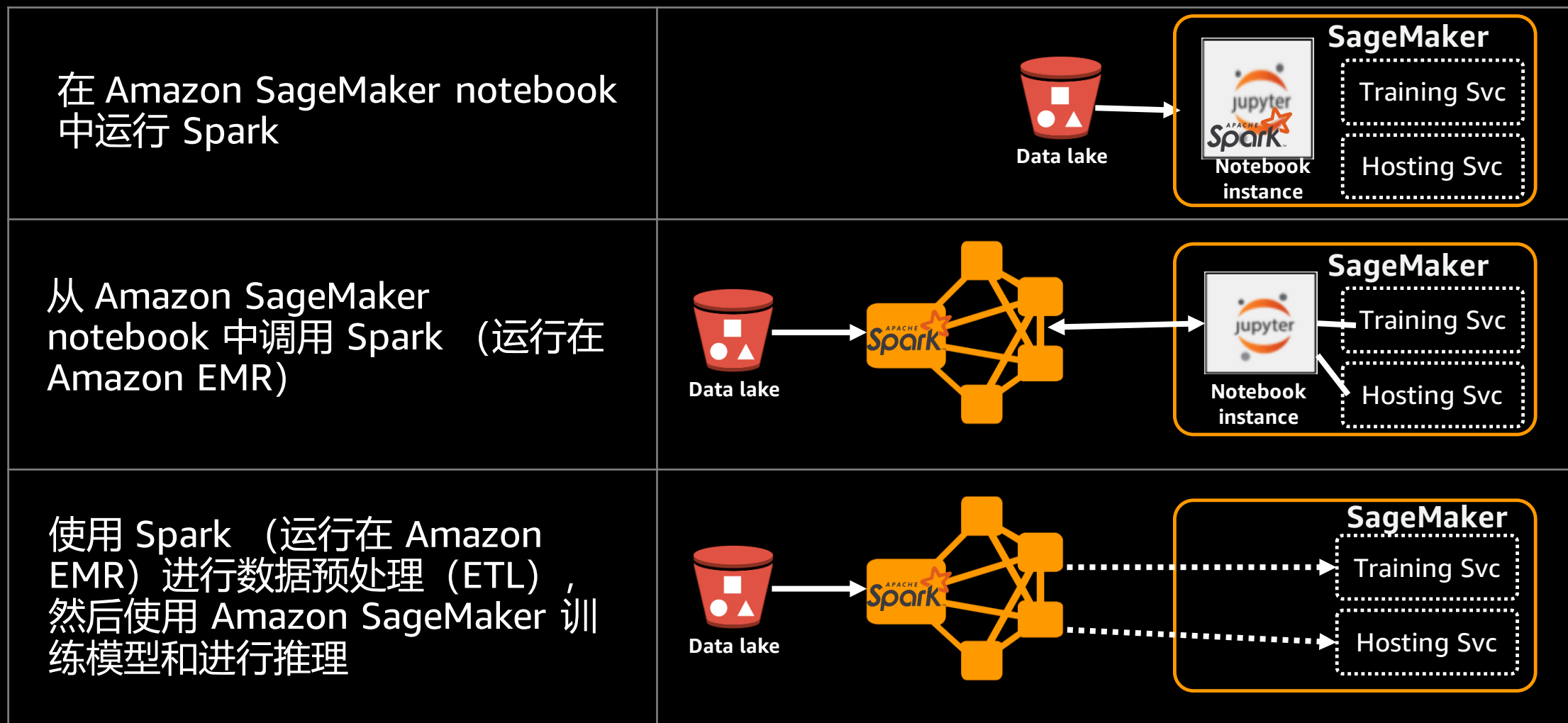
- 开放源代码——GitHub
- Scala 和 Python SDKs
- Spark 和 Amazon SageMaker 之间的通信
 - 序列化与反序列化
 - RecordIO protobuf 格式
- 开箱即用的类
 - 模型训练: `SageMakerEstimator`
 - 推理: `SageMakerModel`

Amazon SageMaker-Spark SDK

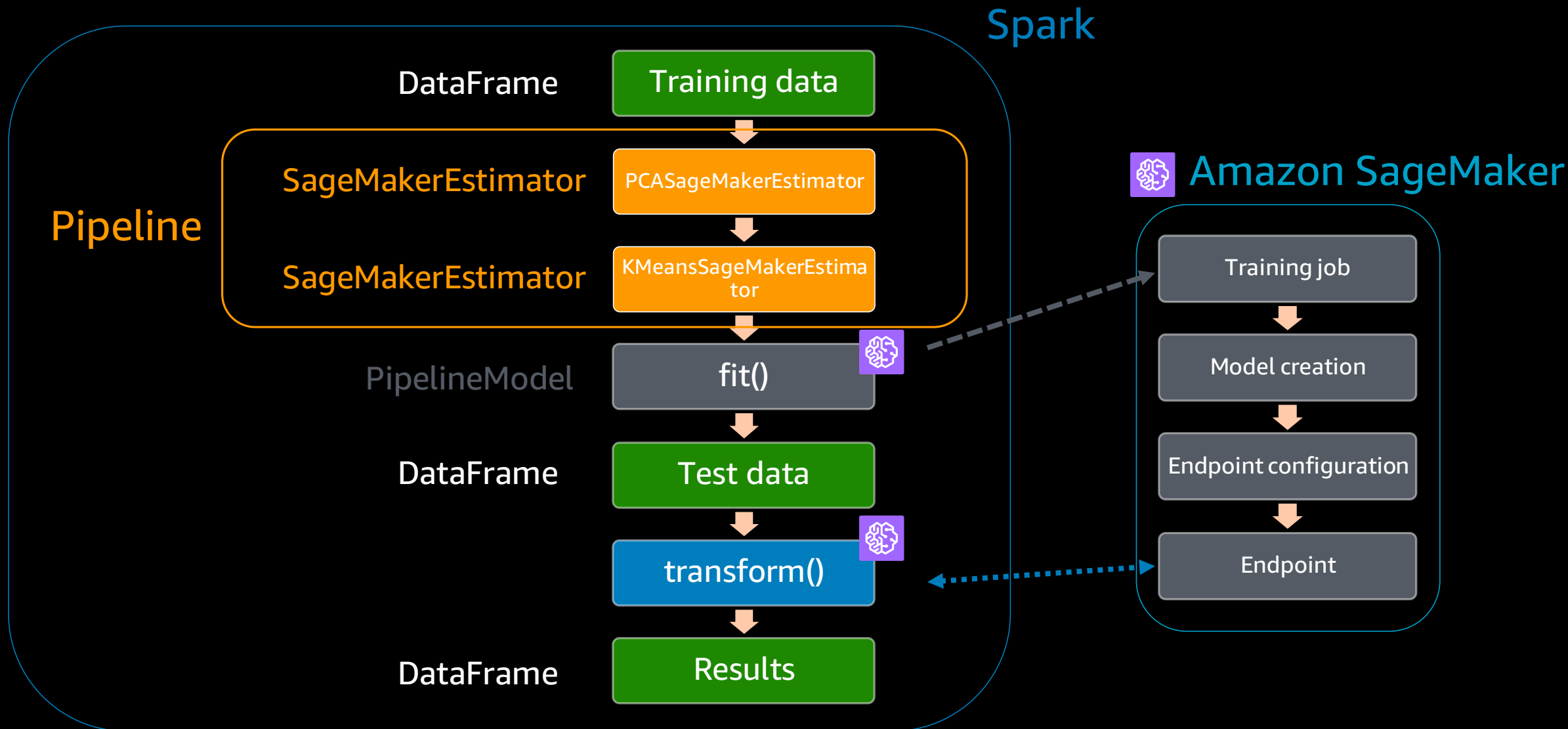
算法 (Amazon SageMaker Estimators):

- K-Means
- Linear Learner Regressor
- Linear Learner Binary Classifier
- PCA
- XGBoost
- Factorization Machine
- Latent Dirichlet Allocation (LDA)
- SageMakerEstimator → 自有算法可运行于 Amazon SageMaker

Apache Spark 和 Amazon SageMaker 整合的场景



Spark ML Pipeline 集成 Amazon SageMaker



AWS Glue 中的 Apache Spark

© 2020, Amazon Web Services, Inc. or its affiliates. All rights reserved.

aws **INNOVATE**
ONLINE CONFERENCE

AWS 中国（宁夏）区域由西云数据运营
AWS 中国（北京）区域由光环新网运营

AWS Glue 组件



数据目录

探索

- 自动爬取
- 兼容 Apache Hive 源数据
- 与 Amazon Web Services (AWS) 分析服务集成



无服务引擎

开发

- Apache Spark
- Python shell
- 交互式和批处理任务



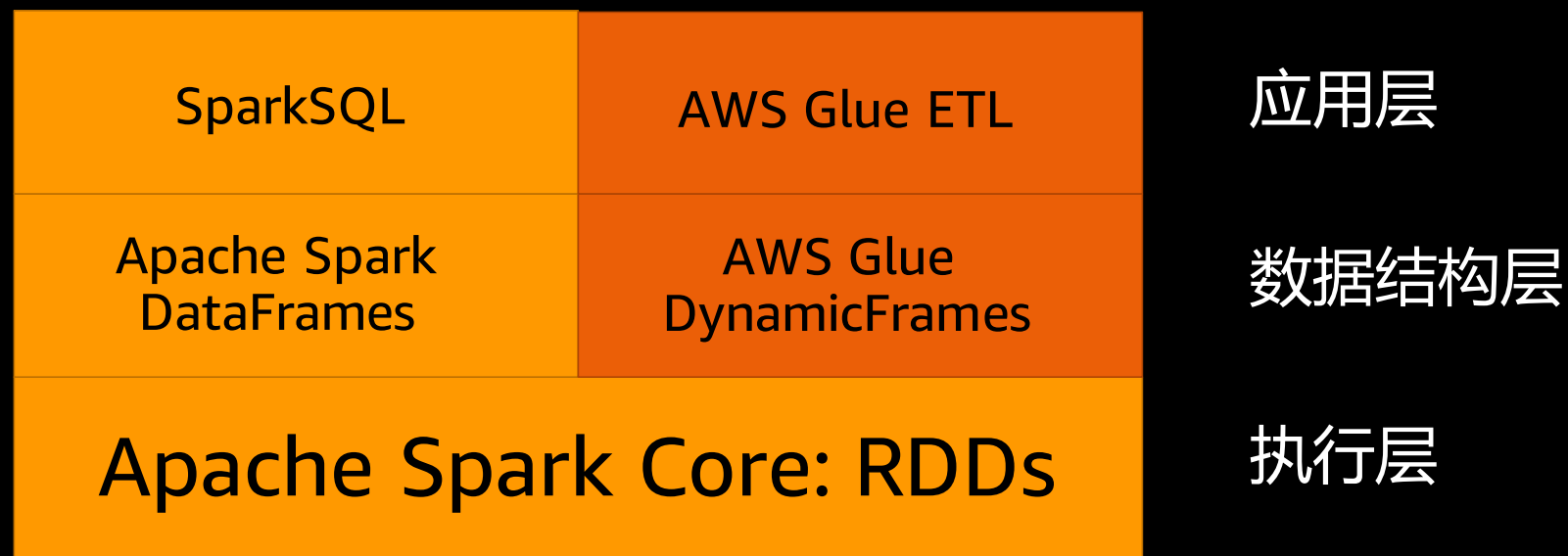
编排器

部署

- 灵活的调度
- 监控和警报
- 外部集成

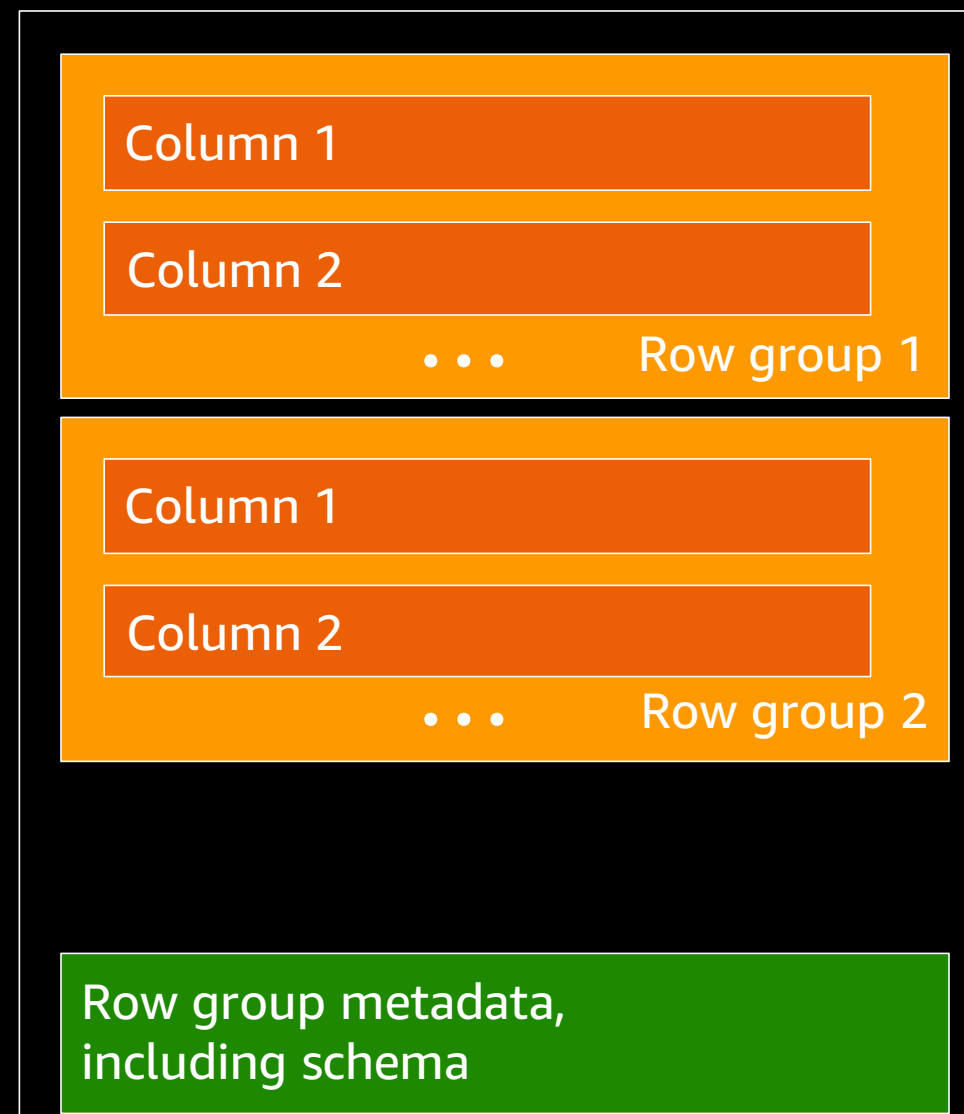
幕后细节

无服务器的 Apache Spark 和重要扩展！



Glue Parquet Writer

- 我们创建了一个定制的 Parquet writer, 提供了 schema 灵活性
- 标准 Parquet writer:
设置 schema -> 写入行组
- Glue Parquet writer:
 1. 开始写入列, 按需增加字段;
 2. 关闭开始的行组并写入 schema
- 额外的架构更改会触发新文件创建



感谢参加 AWS INNOVATE 2020 在线技术大会

我们希望您在这里找到感兴趣的内容！

也请帮助我们完成**投票打分**和**反馈问卷**。

欲获取关于 AWS 的更多信息和技术内容，可以通过以下方式找到我们：

-  微信订阅号：AWS 云计算 ([awschina](#))
-  微信服务号：AWS Builder 俱乐部 ([amazonaws](#))
-  新浪微博：<https://www.weibo.com/amazonaws/>
-  抖音：亚马逊云计算 (抖音号：266052872)
-  视频中心：<http://aws.amazon.bokecc.com/>
-  博客：<https://aws.amazon.com/cn/blogs/china/>
-  更多线上活动：<https://aws.amazon.com/cn/about-aws/events/webinar/>



AWS 中国账户注册



AWS 全球账户注册