

اصول ذخیره و پردازش

داده‌های حجیم (Big Data)

نویسنده:

مهندس منعم مه‌بیان

کارشناس ارشد مهندسی فناوری اطلاعات و شبکه‌های کامپیوتری

دانشگاه صنعتی امیرکبیر (پلی تکنیک تهران)



**NAGHOOS
PUBLICATION**

سرشناسه : مهدیان، محمد، ۱۳۶۴ -
 عنوان و نام پدیدآور : اصول ذخیره و پردازش داده های حجیم (Big Data) / مولف محمد مهدیان
 مشخصات نشر : تهران: ناگوس، ۱۳۹۹
 مشخصات ظاهری : ۲۱۳ص : جدول، نمودار(رنگی).
 شابک : ۹۷۸-۶۰۰-۴۷۳-۲۷۱-۰ : ۴۵۰,۰۰۰ریال
 وضعیت فهرست نویسی : فیا

موضوع : داده های کلان
 موضوع : Big Data
 موضوع : داده های کلان -- مدیریت
 موضوع : Big Data--Management
 موضوع : پایگاه های اطلاعاتی--مدیریت
 موضوع : Database Management
 رده بندی کنگره : QA۷۶/۹
 رده بندی دیویی : ۰۰۵/۷۴
 شماره کتابشناسی ملی : ۶۱۲۳۵۸۸



NAGHOOS
PUBLICATION

برای خرید Online به آدرس زیر مراجعه کنید:

www.naghoospress.ir

انتشارات ناگوس

چاپ اول

نام کتاب : اصول ذخیره و پردازش داده های حجیم (Big Data)

ناشر : انتشارات ناگوس

مولف : مهیندس محمد مهدیان

چاپ اول : ۱۳۹۹

تیراژ : ۵۰۰نسخه

ناظر چاپ : یاسمن تجملی محدث

چاپ و صحافی : حیدری

قیمت : ۴۵۰,۰۰۰ریال

شابک : ۹۷۸-۶۰۰-۴۷۳-۲۷۱-۰

ISBN : 978-600-473-271-0

کتاب به حقوق برای سرمایه گذار محفوظ است. هر گونه کپی یا قسمتی از این اثر به صورت جداگانه یا چاپ مجدد، چاپ سنت، یا کپی و کپی و انواع دیگر چاپ، توزیع است و پیگرد قانونی دارد.

انتشارات ناگوس

۱-بلوار کشاورز-خیابان ۱۶آذر-کوچه پارس-پلاک ۱۰

تلفن و فاکس : ۶۶۳۷۸۹۵۱-۶۶۳۷۸۹۵۴

فهرست مطالب

۱.....	مقدمه
۳.....	فصل اول.....
۳.....	مفاهیم محاسبات ابری و داده‌های حجیم
۴.....	۱-۱-تعریف محاسبات ابری (Cloud Computing)
۵.....	۱-۱-۱-رکمی Cloud
۷.....	۱-۱-۲-سرویس مدل Cloud
۱۱.....	۱-۱-۳-Deployment مدل‌های Cloud
۱۴.....	۲-۱-تعریف داده‌های حجیم (Big Data)
۱۵.....	۱-۲-۱-ابعاد چهارگانه Big Data
۱۶.....	۲-۲-۱-ذخیره و پردازش Big Data
۲۳.....	فصل دوم.....
۲۳.....	ذخیره سازی داده‌های حجیم
۲۴.....	۱-۲-ذخیره سازی گوگل فایل سیستم (Google File System)
۲۶.....	۱-۱-۲-کامپوننت های GFS
۳۰.....	۲-۱-۲-تعاملات کامپوننت های GFS
۳۴.....	۳-۱-۲-Fault Tolerance در GFS
۳۵.....	۲-۲-ذخیره سازی NoSQL Database (Not Only SQL)
۳۸.....	۱-۲-۲-دیتا مدل های NoSQL
۴۰.....	Consistency-۲-۲-۲
۴۱.....	۳-۲-۲-تئوری CAP
۴۲.....	۴-۲-۲-دیتابیس Dymano (استفاده در شرکت آمازون)
۴۵.....	Consistency-۱-۴-۲-۲ در Dynamo
۴۵.....	۲-۴-۲-۲-انواع API در Dynamo

۴۶		Dynamo در Membership	۲-۴-۲-۲
۴۶		Dynamo تشخیص خرابی در	۲-۴-۲-۲
۴۶		BigTable دیتابیس (استفاده در شرکت گوگل)	۲-۵-۲-۲
۴۸		BigTable در API انواع	۲-۵-۱-۲-۲
۴۸		BigTable معماری	۲-۵-۲-۲
۴۹		BigTable Master وظایف	۲-۵-۲-۲
۴۹		BigTable Tablet server وظایف	۲-۵-۴-۲-۲
۵۰		BigTable client وظایف	۲-۵-۵-۲-۲
۵۰		Tablet عمل خواندن و نوشتن روی	۲-۵-۶-۲-۲
۵۲		BigTable در Consistency	۲-۵-۷-۲-۲
۵۳		BigTable در (Compression) فشرده‌سازی	۲-۵-۸-۲-۲
۵۳		Cassandra دیتابیس (استفاده در شرکت فیس بوک)	۲-۶-۲-۲
۵۵		فصل سوم	
۵۵		پردازش داده‌های حجیم	
۵۷		MapReduce	۳-۱-۱
۵۷		MapReduce کلی ایده	۳-۱-۱-۱
۵۸		MapReduce کلی بخش‌های	۳-۱-۲-۱
۵۹		MapReduce در Programming Model	۳-۱-۲-۱-۱
۶۶		MapReduce در Implementation	۳-۱-۲-۲-۱
۶۷		Hadoop	۳-۱-۳
۶۸		MapReduce در Fault Tolerance	۳-۱-۴-۲
۶۹		MapReduce محدودیت‌های	۳-۱-۵-۲
۷۰		Flume	۳-۲-۲
۷۴		Spark	۳-۳-۲
۷۶		Spark کلی ایده	۳-۳-۱-۱
۷۶		MapReduce و Spark مقایسه	۳-۳-۲-۲
۷۸		RDD (Resilient Distributed DataSets)	۳-۳-۲-۲

۷۹.....	Spark در Programming model-۴-۳-۳
۸۱.....	Spark Programming model در Transformation-۱-۴-۳-۳
۸۳.....	Spark Programming model در Action-۲-۴-۳-۳
۸۴.....	Spark در SparkContext-۲-۴-۳-۳
۸۵.....	Spark در RDD ایجاد-۴-۴-۳-۳
۸۷.....	Spark در Execution-۵-۳-۳
۸۸.....	Spark در RDD ایجاد-۱-۵-۳-۳
۸۸.....	Spark در RDD ها در وابستگی-۲-۵-۳-۳
۸۹.....	Spark در Job Scheduling-۳-۵-۳-۳
۹۰.....	Spark در RDD Fault Tolerance-۴-۵-۳-۳
۹۰.....	Spark SQL (Relational Data Processing in Spark)-۶-۳-۳
۹۱.....	Spark SQL در DataFrame مدل-۱-۶-۳-۳
۹۲.....	Spark SQL در DataFrame ایجاد-۲-۶-۳-۳
۹۲.....	Spark SQL در DataFrame های Operation-۳-۶-۳-۳
۹۳.....	Spark SQL در SQL Queries اجرای-۴-۶-۳-۳
۹۴.....	Spark SQL در DataFrame به RDD تبدیل-۵-۶-۳-۳
۹۵.....	فصل چهارم
۹۵.....	پردازش داده‌های Scalable Stream
۹۵.....	سیستم‌های پردازش داده‌های (SPS) Stream-۱-۴
۹۶.....	مقایسه سیستم‌های DBMS با سیستم‌های SPS-۱-۱۰-۴
۹۷.....	SPS در دیتا مدل-۲-۱-۴
۹۸.....	SPS در پراسس مدل-۳-۱-۴
۹۸.....	SPS در Programming Model-۴-۱-۴
۹۹.....	DataFlow Composition-۱-۴-۱-۴
۱۰۱.....	DataFlow Manipulation-۲-۴-۱-۴
۱۰۶.....	SPS در Runtime System-۵-۱-۴
۱۰۸.....	SPS در Parallelization-۶-۱-۴

۱۱۱.....	SPS در Fault Tolerant-۷-۱-۴
۱۱۴.....	SPS در Optimization-۸-۱-۴
۱۱۷.....	سیستم Storm (استفاده در شرکت Twitter)
۱۱۷.....	Storm در دیتامدل-۱-۲-۴
۱۱۸.....	Storm در Parallelization-۲-۲-۴
۱۱۹.....	Storm در Programming Model-۳-۲-۴
۱۲۱.....	معماری Storm-۴-۲-۴
۱۲۲.....	Storm Deployment-۵-۲-۴
۱۲۳.....	Storm در Fault Tolerance-۶-۲-۴
۱۲۳.....	Storm در Reliable Processing-۷-۲-۴
۱۲۶.....	سیستم SEEP (استفاده در دانشگاه لندن)
۱۲۶.....	SEEP در ایده کلی-۱-۳-۴
۱۲۶.....	SEEP در State انواع-۱-۱-۳-۴
۱۲۸.....	SEEP در Operator State مدیریت-۲-۱-۳-۴
۱۲۹.....	SEEP در Scale out-۲-۳-۴
۱۳۰.....	SEEP در Fault Tolerant-۳-۳-۴
۱۳۱.....	Spark Streaming-۴-۴
۱۳۱.....	Spark Streaming در ایده کلی-۱-۴-۴
۱۳۲.....	Spark Streaming در Programing model-۲-۴-۴
۱۳۳.....	Spark Streaming در DStream دیتامدل-۱-۲-۴-۴
۱۳۳.....	Spark Streaming در StreamingContext-۲-۲-۴-۴
۱۳۴.....	Spark Streaming در Source of Streaming-۳-۲-۴-۴
۱۳۵.....	Spark Streaming در DStream Transformation-۴-۲-۴-۴
۱۳۹.....	Spark Streaming و DataFrame-۵-۲-۴-۴
۱۴۰.....	Spark Streaming در Implementation-۳-۴-۴
۱۴۲.....	Structured Streaming-۴-۴-۴
۱۴۵.....	Flink در Stream-۵-۴
۱۴۵.....	Stream Processing یا Batch Processing مقایسه‌ی-۱-۵-۴

۱۴۶.....	Flink و Spark مقایسه‌ی مفهومی ۲-۵-۴
۱۴۸.....	Flink در Programming Model ۳-۵-۴
۱۴۸.....	Flink در Windowing Semantics ۱-۳-۵-۴
۱۴۸.....	Flink در API ۲-۳-۵-۴
۱۵۰.....	Flink در DataStream sources ۳-۳-۵-۴
۱۵۲.....	Flink در Implementation ۴-۵-۴
۱۵۳.....	Flink در Fault Tolerance ۵-۵-۴
۱۵۷.....	فصل پنجم
۱۵۷.....	پردازش Large Scale Graph
۱۶۰.....	مدل Pregel (استفاده در شرکت Facebook) ۱-۵
۱۶۰.....	ایده کلی Pregel ۱-۱-۵
۱۶۱.....	Pregel در Programming Model ۲-۱-۵
۱۶۱.....	Pregel در Execution Model ۳-۱-۵
۱۶۵.....	Pregel در Implementation ۴-۱-۵
۱۶۵.....	Pregel در Fault Tolerance ۱-۴-۱-۵
۱۶۶.....	Pregel محدودیت ۲-۴-۱-۵
۱۶۶.....	GraphLab (Furi) ۲-۵
۱۶۶.....	GraphLab در Programming model ۱-۲-۵
۱۶۷.....	GraphLab در Execution model ۲-۲-۵
۱۶۸.....	GraphLab در Consistency ۳-۲-۵
۱۶۹.....	GraphLab در Graph Partitioning ۴-۲-۵
۱۷۱.....	GraphLab در Fault Tolerance ۵-۲-۵
۱۷۱.....	(GraphLab2) PowerGraph ۳-۵
۱۷۱.....	PowerGraph در Programming model ۱-۳-۵
۱۷۲.....	PowerGraph در Execution model ۲-۳-۵
۱۷۴.....	PowerGraph در Graph partitioning ۳-۳-۵
۱۷۶.....	GraphX ۴-۵

۱۷۷.....	GraphX ایده کلی ۱-۴-۵
۱۷۹.....	GraphX در Programming Model ۲-۴-۵
۱۸۲.....	GraphX در Implementation ۳-۴-۵
۱۸۳.....	X-Stream-۵-۵
۱۸۴.....	X-Stream ایده کلی ۱-۵-۵
۱۸۶.....	X-Stream در Programming model ۲-۵-۵
۱۹۳.....	X-Stream در Streaming Partitioning ۳-۵-۵
۱۹۷.....	فصل ششم
۱۹۷.....	مدیریت Resource در Mesos
۱۹۸.....	Mesos اهداف ۱-۶
۱۹۹.....	Mesos در Computation مدل ۲-۶
۱۹۹.....	Mesos عوامل موثر در طراحی ۳-۶
۲۰۱.....	Mesos Framework ها در ۴-۶
۲۰۳.....	Mesos تخصیص منابع در ۵-۶
۲۰۴.....	Single Resource تخصیص ۱-۵-۶
۲۰۶.....	Multi-Resource تخصیص ۲-۵-۶
۲۰۶.....	Asset fairness-۱-۲-۵-۶
۲۰۸.....	(DRF) Dominant Resource Fairness ۲-۲-۵-۶
۲۱۱.....	مراجع

مقدمه

امروزه منابع زیادی در حال تولید داده هستند و دیوایس های خیلی زیادی به اینترنت متصل اند که توسط این منابع و دیوایس ها، حجم دیتای زیادی با سرعت بالا در حال تولید شدن است. مثلا اطلاعات شرکت مخابرات و شبکه های اجتماعی و بازار بورس، اطلاعات ژنتیکی و... در Google روزانه ۱۵۰۰۰ پنتا بایت داده در حال ذخیره و ۱۰۰ پنتا بایت پردازش اطلاعات داریم^[25] این مقادیر آن قدر بالا هست که در یک کامپیوتر عادی نمی توانیم نه این حجم ذخیره سازی و نه این حجم پردازشی را داشته باشیم و نیاز به مجموعه ای از کامپیوترها داریم که با هم کار کنند تا این حجم داده را ذخیره و پردازش کنند.

مفهوم داده های جدید (Big Data) زمانی مطرح شد که داده هایمان را نتوانستیم ذخیره و پردازش کنیم و با عنایت به اینکه داده های حجیم (Big Data) فقط دیتا هستند که تنها حجم و یک سری مشخصاتشان تغییر کرده است و با توجه به محدودیت هایی که از لحاظ Resource داریم نمی توانیم با استفاده از روش های سنتی، دیتا را ذخیره و تحلیل کنیم پس به یک راه و یک بستر جدید برای ذخیره و تحلیل این دیتا احتیاج داریم. طبیعتا این راه را در بعضی شرکت ها مثل Google که داده های بیشتری داشتند بیشتر احساس شد.

در رابطه با داده های حجیم سه سوال اصلی پیش روی ما قرار دارد و در این کتاب سعی می کنیم آن ها را بررسی کنیم که این سوالات به شرح زیر می باشد:

۱- چگونه داده های حجیم را ذخیره (Store) کنیم؟ ذخیره سازی داده ها اولین مسئله ای است که باید آن را بررسی کنیم که در این رابطه Distributed File System، NoSQL database ها را معرفی و بررسی می کنیم.

۲- چگونه داده های حجیم را پردازش (Process) کنیم؟ حال بسته به اینکه داده های ما از چه جنسی باشند Platform های مختلفی را معرفی می کنیم. مثلا اگر داده ها ما ثابت غیر قابل تغییر باشند و اگر همه داده ها را یکجا داشته باشیم از Batch Processing Platform استفاده می کنیم و اگر داده های Stream (جاری) داشته باشیم از Streaming Processing Platform استفاده می کنیم و اگر داده ها ساختاریافته باشند می توانیم از Platform هایی مثل Spark SQL یا ... استفاده کنیم و اگر داده ها از جنس گراف باشند از Graph Platform استفاده می کنیم.