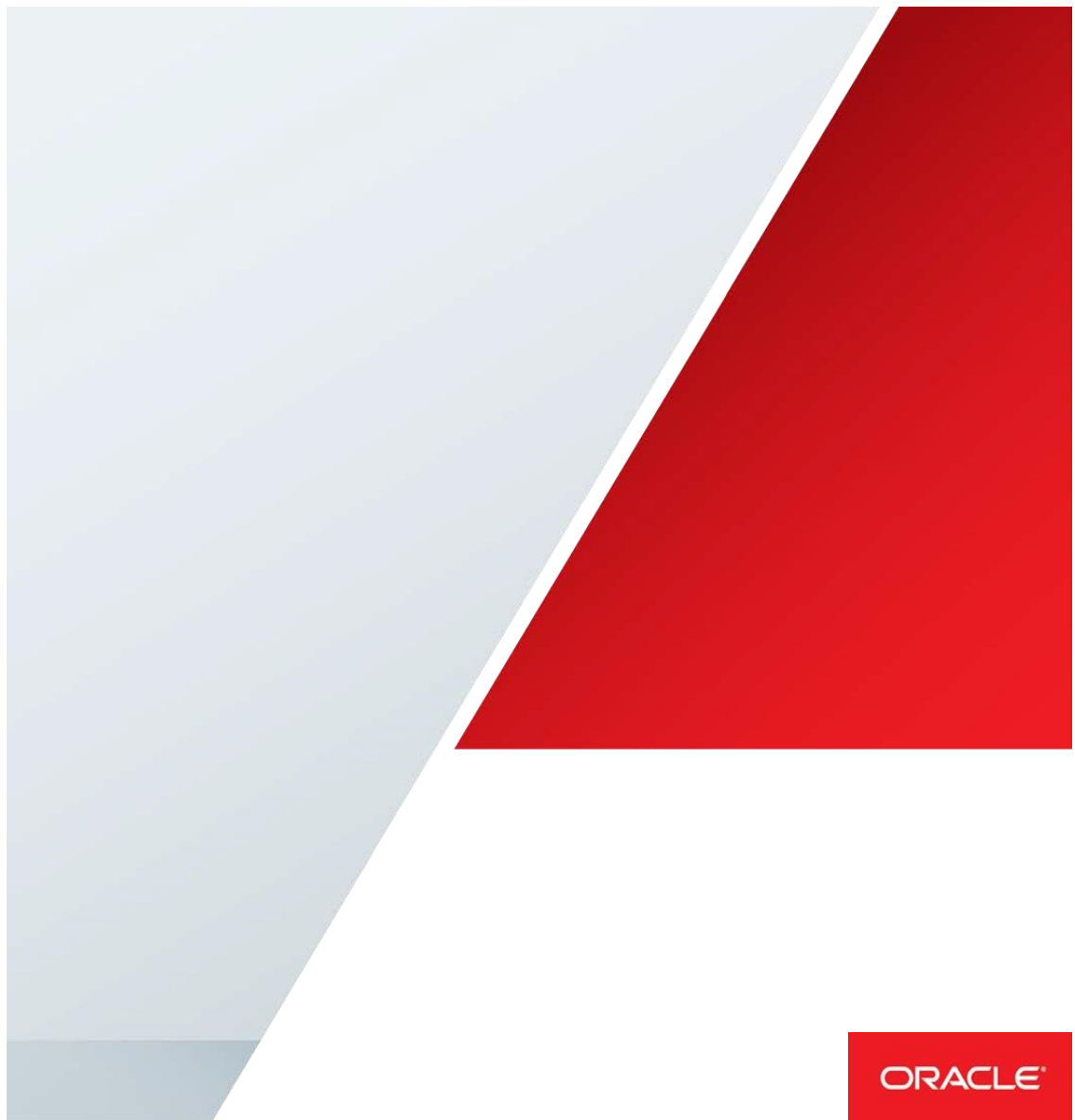




Big Data Management Systemのための 統合問合せ

ビッグ・データ・システムと
エンタープライズ・データウェアハウスの統合
Oracle ホワイト・ペーパー | 2016 年 3 月



目次

概要	3
異なるデータ・アクセス API の課題	4
統合問合せの必要性	5
統合問合せの価値	5
統合問合せと企業情報アーキテクチャ	5
統合問合せ、データ・サイエンス、およびビジネス・インテリジェンス	7
統合問合せとアプリケーション開発	7
問合せフランチाइジングと Oracle Big Data SQL	8
言語レベルのフェデレーション	8
問合せフランチाइジング：Big Data SQL のアプローチ	9
結論	11

概要

今でもビッグ・データを誇大宣伝ととらえている人もいますが、ビッグ・データに関連するツールと方法論は、大規模な情報管理アーキテクチャにおいてますます重要になってきています。実際に、ビッグ・データがエンタープライズ IT に浸透するにつれて、ビッグ・データのビジネスは急成長しています。IDC の最近の予測では、ビッグ・データのテクノロジーおよびサービスの市場は 2017 年まで 27% の複合年間成長率で 324 億ドルに達すると見られ、これは IT 市場全体の成長率の約 6 倍となっています¹。変化の速いこの時勢、組織は新しいテクノロジーを採用するリスクを最小限に抑えながら、ビッグ・データのビジネスの利点を最大限に高めるよう慎重に進めていく必要があります。

幅広い Hadoop エコシステム (Apache Hadoop および Apache Spark プロジェクトを含む) のコンポーネントと NoSQL データベースのホストの両方が組織に価値を提供しています。NoSQL ストアは、もっとも堅牢で多角的なリレーショナル・データベースの何分の 1 のコストで、オンライン・アプリケーションを支えるのに適したシンプルで水平スケーラビリティ備えた短い待ち時間のデータ・アクセスを提供できます。Apache Hadoop と特にその分散ファイル・システム (HDFS) は、多くの企業が従来のストレージ式のネットワーク・デバイスにかかるコストを削減して、情報アーキテクチャのより豊かでより深い基盤を構築するのに役立っています。分析では、Apache Spark とその関連コンポーネントは多数の言語での直接の API を通して、スケーラブルで複雑なマシンの学習を可能にしています。

ビジネスにおける潜在的利益にもかかわらず、ビッグ・データ・テクノロジーは基盤の確立した企業にとって今でも難しい課題となっています。これらの課題の最大のものは、幅広いビッグ・データのエコシステムの進歩によって広がるスキルのギャップです。Gil Press は Forbes.com に次のように書いています。「スキルのあるスタッフの不足が続きます。米国だけでも 2018 年には 181,000 の詳細分析の職があり、データ管理および解析の関連スキルを必要とする多くのポジションの 5 倍となっています」²。スキルを超えて、ビッグ・データ・ソリューションを通して創出された価値を幅広い企業情報アーキテクチャに統合するには、セキュリティおよび管理ポリシーがこれらのシステムによって生産され、保持されるデータに適用されるようにする必要があります。機密情報の侵害がますます進む時代にあって、組織は分析のレベルにかかわらず、優れたデータ管理およびセキュリティを見逃すことはできません。

¹ 新しい IDC の世界的なビッグ・データのテクノロジーおよびサービス予測では市場が 2017 年に 324 億ドルまで拡大すると予想
<http://www.idc.com/getdoc.jsp?containerId=prUS24542113>

² 2015 年の 1,250 億ドルのビッグ・データ分析市場 6 つの予測、2014/12/11
<http://www.forbes.com/sites/gilpress/2014/12/11/6-predictions-for-the-125-billion-big-data-analytics-market-in-2015/>

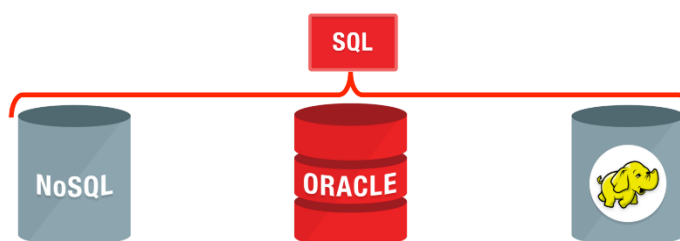
異なるデータ・アクセスAPIの課題



図1：各種データ・ストアの運用特性は企業がデータ・ストレージのパフォーマンスやコスト要件を最適化するのに役立ちますが、異なるデータ・アクセスAPIはこれらのストアで創出された価値を統合が難しいサイロに適用します。

ビッグ・データ・テクノロジーの使用の利点として考えられる、スケーラビリティ、分散処理、線形費用特性を備えた低コストのストレージは、ビジネスに対して測定可能な利点を提示することができます。しかし、これらのテクノロジーから本当に価値を創出する際の中心的な課題は、よく引き合いに出されるスキルのギャップ、つまり異なるデータ・アクセスAPIと緊密に関連しています。エンタープライズITでのデータ・アクセスはおもにSQL言語で行われます。多数のアプリケーションと分析ツールはリレーショナル・データベースと密接に関係して進歩してきたため、技術者もツールも、SQLの方言をサポートするデータベースが基盤となっています。部分的にはその宣言型の結果指向の性質により、SQLはデータ操作のための共通語となります。

Hadoop および NoSQL テクノロジーへのアクセスのための中核フレームワークを考えた場合、これらは多くのツールや技術者が使用しているSQL中心アプローチとは大きく異なります。大部分のNoSQLベンダーは



基本のCRUD操作（作成、読取り、更新、削除）をサポートしていますが、これらのメソッドの構文やセマンティックは実装間で異なることがあります。Hadoopエコシステムは、最大のデータ、とりわけSQLインタフェースを処理するためのより多様で強力な処理アプローチを急速に開発しています。しかし、最近までHadoopクラスタの主要な作業ユニットはMapReduceジョブと考えられることがありました。基礎となるMapReduceフレームワークでは、もっとも基本的なタスクを実行するだけでも、大量のJavaコードとプログラマーによる熱心な最適化が必要となります。

統合問合せの必要性

ビッグ・データ・テクノロジーの潜在的価値に基づいて本当に最適化を行うには、労働力を有効にし、方針を徹底し、異なるアクセス方法によってサイロで生じる価値のリスクを低減する能力が企業に求められます。必要な洞察を得るために多数のデータ・サイロを隈なく探して後から組み上げて明確な絵を構成するのとは反対に、値の検索は統合問合せにより大幅に時間を短縮できます。つまり、1 つの SQL 文が（方針の徹底により）すべてのデータ・ストア、リレーショナル・データベース、Hadoop クラスタ、NoSQL データベースでシームレスに実行できた場合に企業はビッグ・データから最大の利益を得られます。統合問合せにより、企業は複雑なデータの統合を調整する代わりに質問への回答に集中することができます。統合問合せにより、分析者とアプリケーションは既存のツールとスキルを利用しながらビッグ・データを最大限活用できます。統合問合せを使用することで、ビッグ・データは小さなデータと同じぐらい管理しやすくなると考えられます。

この記事では統合問合せと、特に Oracle Big Data SQL の利点について説明します。Big Data SQL は、問合せフランチाइジングという独自のアプローチにより、Oracle データベース、Hadoop クラスタ、NoSQL データ・ストアから統合問合せを作成します。Oracle Exadata Database Machine 用に最初に開発された革新的な技術に基づき、この新しい方法では、すべてのデータ・ストアでの最高のパフォーマンスと、既存のセキュリティおよび管理モデルのシンプルな拡張の両方がすべてのデータに対して実現されます。

統合問合せの価値

幅広いビッグ・データのエコシステムのコンポーネントがますます充実する企業アーキテクチャにおいて、統合問合せシステムの価値はいくら強調してもしすぎることはありません。組織の各種コンポーネントの重要な機能は異なる速度で成熟するため（たとえば、HDFS のネイティブの問合せシステムはリレーショナル・データベースよりも生産性が低く、水平スケーラビリティは一部のエンタープライズ RDBMS では桁違いの費用がかかることがある）、企業がこれらのコンポーネントを最大限に利用するよう保証することは、データ管理システム全体の最大の投資収益率を保証するための鍵となります。個々のデータ・ストアで作成された値を結合し、統合するための統合問合せシステムがなければ、各種システムのデータを統合し、調整するために膨大なリソースが必要となります。

統合問合せシステムを使用すると、ビッグ・データの課題を管理可能な程度に小さくできます。統合問合せによって実現するデータ管理システムでは、データをもっとも適切な場所に格納でき、プラットフォーム固有のアプローチによって操作と変換が行われ、既存のリレーショナル・データベースのスキーマと問合せを拡張するだけでビジネスに不可欠なプロセスに統合できます。このセクションでは、企業情報管理の3つの側面について、統合問合せの強力な効果を検討します。

- 企業情報アーキテクチャ
- データ・サイエンスとビジネス・インテリジェンス
- アプリケーション開発

統合問合せと企業情報アーキテクチャ

企業情報アーキテクチャは、大まかに言うと中核ビジネス・プロセスを IT システムに変換することと考えられます。実際、Ross などが規定したように、企業情報アーキテクチャはビジネスの執行の

ための基礎となっています³。デジタル化がますます進む時代において、ビッグ・データ・エコシステムの部分を含む優れた情報アーキテクチャが最大の価値を実現するために重要となります。

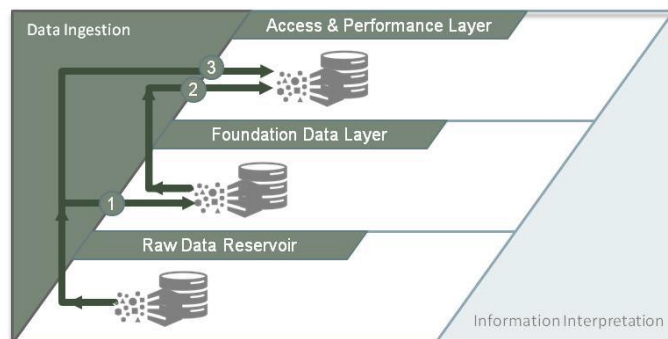


図2：Oracle Information Management Reference Architectureの物理コンポーネント

図 2 の Oracle Information Management Reference Architecture について考えてみます。多くの参照アーキテクチャと同様、ビッグ・データ・コンポーネントに適した箇所が増えていきます。アーキテクチャの基礎レイヤーには、多くの最適なレポート・システムやレポート・チェーンへの情報提供が可能な、ビジネス・プロセスに依存しない標準的なモデルが配置されます。長年、基礎レイヤーは企業のデータウェアハウスとして実装されており、このようなマッピングが続いていると考えられます。しかし、規制の要件から洞察のための検索まで、多くの要素によって生データの大きなリザーバの要件が高まっています。この生データ・リザーバは運用の起源を忠実に反映するデータの不変のプールとなります。生データ・リザーバは、普遍性と真の潜在サイズを考えると、Apache Hadoop、特に HDFS の実装に最適であると言えます。実際に、「データ・リザーバ」という用語は Hadoop と関連して使われることが多くなっています⁴。HDFS 上に構築された生データ・リザーバを追加する上での課題は、特にデータをリザーバから基礎レイヤーに昇格させるときに発生します。基礎レイヤーが RDBMS に実装されている場合、データについて以下を実行する必要があります。

1. リザーバからスキャンする
2. 適切なデータ型に変換する
3. システムの品質保証契約に影響を及ぼさずに RDBMS にロードする

統合問合せシステムがない場合、この一連の作業は明確なバッチ・プロセスのセットとして実装されるのが普通です。スキャンは MapReduce、Spark、SQL-on-Hadoop 問合せエンジンを通して行うことができます。型変換とロードはスケジュールされたバッチ・ロードによって行われます。ここで、データの品質が検証される必要があり、新しい変換サイクル全体が行われることもあります。統合問合せシステムを使用すると、データベースから HDFS のデータに直接問合せすることによってこれらの作業が行われます。つまり、リザーバのデータは、既存の品質管理作業に含まれ、進行中の変換にも含まれ、必要に応じて臨機応変にアクセスできることになります。

統合問合せとは

統合問合せでは、1 つの宣言型の文で多数の異種データ・ストアのデータにアクセスし、分析することができます。たとえば、リレーショナル・データベース、NoSQL データ・ストア、大規模なパラレル・ファイル・システム（HDFS など）などが対象となります。統合問合せはデータ・ストア間でシームレスに機能し、多くの形式でのデータの保存が可能でなければなりません。

3 Jeanne W. Ross, Peter Weill, David Robinson 著『Enterprise Architecture as Strategy: Creating a Foundation for Business Execution』Harvard Business Press, 2006 年

4 『Are You Pouring Pollution into the Data Lake?』Merv Adrian, Gartner, 2014 年 10 月 10 日

統合問合せ、データ・サイエンス、およびビジネス・インテリジェンス

ビッグ・データ・システムによって新しい洞察を得るための大きな可能性が生まれていますが、価値を創出するにはこれらの洞察を視覚化し、組織の最高レベルで理解できるようにする必要があります。広範な粒状データからアクションのためのストーリーと推奨事項を引き出す必要性からデータ・サイエンスの分野が生まれています。しかし、データ・サイエンティストの作業は純粋な分析よりも、膨大な「用務員の作業」によって埋もれてしまうことがよくあります⁵。洞察を得るための予備段階としてデータを整備し、修正するために膨大な手作業が必要となるためです。これとは対照的に、統合問合せシステムを使用すると、基本の問合せシステムにデータの処理と結合を任せながら、データ・サイエンティストは分析に集中することができます。統計分析ツールは SAS Advanced Analytics や Stata などの業界標準のものから、Python、R などのよく使用される言語のための最新のパッケージまで幅広く、これらで統合問合せを使用してサイエンティストの選択したプラットフォームで異種データをすばやく分析することができます。

同様に、洞察が最初のアクションにつながったら、アクションの結果とビジネスへの影響をレポートするためにビジネス・インテリジェンス (BI) システムが必要となります。会社の Web サイトを利用した販売の向上と、ユーザーが新製品のブラウズに費やす時間の向上の両方を目標としたデータ・サイエンスの研究を考えてみます。ユーザーのトランザクションが RDBMS システムに格納され、ユーザーのクリック・データが Hadoop に格納される場合、両方の動作の変化を同じダッシュボード内でレポートするには、Hadoop から統合問合せシステムまたは BI システムをサポートする RDBMS にバッチ・ロードする必要があります。ますますリアルタイム性が求められる世界において、すべてのシステムで必要なレポートを提供できるのは統合問合せのみです。

統合問合せとアプリケーション開発

統合問合せはアプリケーション開発の領域に 2 つの利点をもたらします。1 つめは直接的なものです。データ処理アプリケーションの大部分は、オブジェクト・リレーション・マッパー (ORM) を通してリレーショナル・データベースとインタフェースに支えられています。既存のアプリケーションでビッグ・データを利用するために再調整する問題により、アプリケーション・データ・アクセス・レイヤーの完全な書換えが必要になるおそれがあります。このため、NoSQL データベース専用のデータ・アクセス・レイヤーを追加したり、Hadoop クラスタからアプリケーションのバックアップ・データベースにバッチ・ロードしたりすることが必要となる場合があります。統合問合せによりこれらの問題が解消します。SQL を使用してリレーショナル・データベースのデータにアクセスする ORM は、既存のモデルにオブジェクトの関係を追加するだけで NoSQL および Hadoop ストアにアクセスできます。前述のように、基本のシステムによってアプリケーションの問合せが処理されますが、統合問合せでは関連する新しいソースのデータを対象とすることができます。

ビッグ・データによるアプリケーションの強化では統合問合せが有効ですが、新規アプリケーションのレポートは大きな利点です。短い待機時間の NoSQL ストアをバックアップ・データベースとして利用するモバイルおよび Web アプリケーションの数を考えてみます。これらのシステムから既存の BI システムへのデータの統合は非常に難しい場合があります。ここでも、ソース・システムからデータウェアハウスへのバッチ・ロードが必要になることがあります。統合問合せシステムでは、新規アプリケーションを NoSQL プラットフォームで開発することができますが、既存のデータウェアハウスペースのレポート・プロセスと一緒に容易にレポートできます。

5 『For Big Data Scientists, 'Janitor-Work' Is Key Hurdle to Insights.』 New York Times. 2014 年 08 月 17 日

問合せフランチャイジングとOracle Big Data SQL

既存のデータウェアハウスと進歩するビッグ・データ・ツールのセットから最大の利益を引き出すとする企業にとっての統合問合せの価値は明白です。既存のスキルとプロセスを有効に利用してすべてのデータを管理する能力はコストを大幅に節約し、革新を直接促進することができます。しかし、統合問合せの価値を本当に実現するには、その実装が特に重要となります。十分に設計された実装により、SLA が満たされ、リソースが適切に管理され、管理している一連のシステムが RDBMS システムか Hadoop システムか NoSQL システムかに関係なく、1 つのものとして応答することを保証できます。十分に考慮されていない実装では、システム全体のパフォーマンスと重要な機能に大きなばらつきが生じます。

統合問合せの原則の多くは 1990 年代のフェデレーテッド・データベース・システムやデータ仮想化ソリューションに類似しています。一般的に、これらのシステムはデータ・.shipping（つまり、計算作業のソースへのデータの移動）または言語レベルのフェデレーションに大きく依存します。データ・shipping の場合、ビッグ・データ・システムへの適用は相反するものと考えられます。ビッグ・データ・システム、特に Hadoop の動機となる原則の多くは、管理対象のデータの量が、計算に移送するには多すぎることが前提となります。言語レベルのフェデレーションは妥当性が高いように見え、各システムは問合せ全体の一部をローカルで実行することを求められます。しかし、これから検証するように、言語レベルのフェデレーションには大きな欠点があります。

Oracle Big Data SQL では統合問合せの実装のために別のアプローチを取っています。Big Data SQL では、問合せフランチャイジングというアプローチを取ることで、パフォーマンスが最大になり、言語レベルのフェデレーションの落とし穴が避けられるように Oracle Database、Hadoop、NoSQL データ・ストア全体で統合問合せが提供されます。このセクションでは、言語レベルのフェデレーションの欠点と、問合せフランチャイジングによって Big Data Management System の統合問合せのための最良の解決策が提供される方法について検証します。

言語レベルのフェデレーション

異種データ・ストアに格納されたデータにアクセスする 1 つの問合せを実行する場合、言語レベルのフェデレーションは問合せの統合に対するもっとも単純なアプローチと考えられます。特に、各データ・ストアで SQL アクセスがサポートされている場合、理論上は SQL 副問合せを各システムにディスパッチし、問合せコーディネータ全体でディスパッチ結果をマージすることができます。

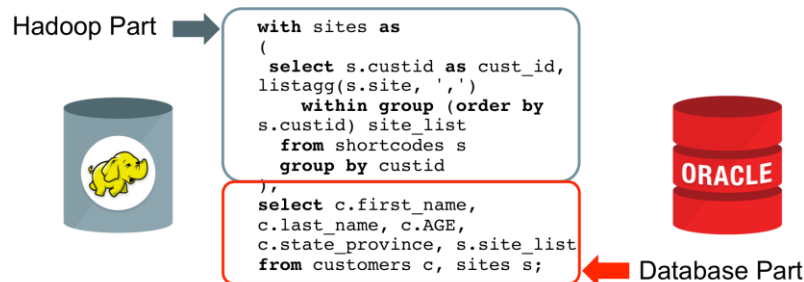


図3：言語ベースのフェデレーションによるSQL処理エンジン間での問合せ文の分割

図 3 の例を見てみましょう。RDBMS と Hadoop の間の統合問合せを提供する言語フェデレーテッド・アプローチでは、1 つの副問合せを Apache Hive に送信し、別の副問合せを Oracle Database に送信することがあります。

これは論理的にはシステムに求められる機能ですが、元の問合せを各システムの副問合せに自動的



に書き直したものを呼び出すだけでは重大な問題が発生します。言語レベルのフェデレーションを可能にするにはいくつかの条件を満たす必要があります。Apache Hive の例に戻ると、これらの一部が明確に理解できます。

まず、2 つのシステムが SQL の同じ方言をサポートしているかを確認します。Hive は一部のウィンドウ機能をサポートしていますが、within group は HiveQL では表示されません。同様に、listagg 関数は HiveQL で表示されません。そのため、Apache Hive による言語フェデレーションは、HiveSQL によって実装された SQL のサブセットに制限されるか、または問合せオーケストレータで結果を模倣することができるコードを動的に生成する必要があります。実装されていないウィンドウ関数の場合、同じデータセットの多数回のスキャンと大量の一時領域が必要となる可能性があります。

2 番目に、システム間でどのようにリソースが管理されるかを確認します。Oracle Database で動作するコードのセクションは高度なリソース管理の対象となり、そのシステムの機能に従ってスケジューリングおよびパラレル化されます。Hive システムに送信された単純副問合せでも、YARN リソース・マネージャから実行時リソースを要求し、ジョブを実行する必要があります。リソース管理ソリューションも間違いではありませんが、予測可能なレスポンス・タイムを提示するため、統合問合せのリソース管理はグローバルで行わなければなりません。グローバル・リソース管理がなければ、言語フェデレーテッド問合せは問合せ対象のシステムのいずれかの予期しない負荷によってブロックされることがあります。

最後に、問合せにセキュリティおよび管理ポリシーが適用される方法を確認します。言語フェデレーテッド問合せでは、1 つのシステムですべてのシステムのすべてのポリシーを解析して適用する必要があるか、またはアクセス制御をベスト・エフォート方式で実装します。これらのベスト・エフォートでは、ユーザーの資格証明が環境全体で整合性を持たない可能性があるため、最小限のアクセス制御しか行われない可能性があります。Hadoop システムのユーザーは Kerberos の認証を得ている可能性があります。同じ Kerberos のチケットは共有されていて、フェデレーテッド問合せに参加している RDBMS に格納されているのでしょうか。ここでも、言語レベルでの単純なフェデレーション問合せは、問合せ対象のシステムのセットで示される最低限の機能までしか持ちません。

問合せフランチャイジング：Big Data SQLのアプローチ

Oracle Big Data SQL は、問合せフランチャイジングという技法を用いた、Hadoop、NoSQL、Oracle Database 環境全体の統合問合せを提供する言語レベル・フェデレーションの代替を提示します。これはシステムのセットの最大の機能を備えています。すべてのシステムでサポートされる最小限のサブセットとは反対に、すべてのデータで Oracle SQL の全体を利用できます。すべての非常に適切なアクセス制御、管理およびセキュリティ・ポリシーは、データ保護に対する緩いベスト・エフォートのアプローチではなく、すべてのデータに適用できます。この多機能の統合問合せシステムは、データの格納位置で Smart Scan 機能を実行する専用のエージェントを使用して、最大のパフォーマンスと予測可能性により実行されます。

問合せフランチャイジングとは

問合せフランチャイジング・システムは、運用の忠実度を失うことなく、異種システムの自己相似コンピュート・エージェントに問合せ処理をディスパッチします。簡単に言うと、各サブシステムの軽量の専用エージェントが問合せ実行に参加し、すべてのデータが同じ演算子を使用してスキャンされ、処理されることが保証されます。

Big Data SQL は統合問合せインタフェースを Oracle Database の外部表機能の拡張として提示します。これにより、開発者、ツール、および Hadoop ソースと NoSQL ソースへのアプリケーション・アクセスが既存のデータベース動作に基づいて 1 対 1 となることが保証されます。これらの拡張は、

Hadoop の InputFormat インタフェースや Hive の SerDe インタフェースなど、広く実装されたオープンソース API を十分に活用できるように設計されています。そのため、ビッグ・データのソースはユーザーにとっては通常のデータベース表に見えますが、ベースのビッグ・データ・システム上にアクセスし、Schema-on-Read セマンティックにより問合せ時に Oracle Database によって自動的に処理されます。

問合せフランチャイズは、格納されているシステム上でデータに対して大規模パラレル・プロセスを行うため、これらの外部表インタフェースを大量に供給します。言語のフェデレーションを行って個別のジョブをディスパッチするのではなく、軽量のエージェントを使用してデータレベルで作業の分散を行います。Smart Scan 機能は大部分のデータの行単位の処理を、そのデータを実際に保持しているサーバーに任せます。これはつまり、Big Data SQL エージェントが動作している Exadata Storage Cell または Hadoop DataNode です。このローカル処理により、システム間で転送されるデータの量が大幅に削減（問合せの実行も高速化）され、各システムで同じ演算子を使用可能であることが保証されます。さらに、これらの専用のエージェントにより、データが統一形式に変換され、ここで必要なすべてのセキュリティ・ポリシーを適用できます。

図 4 に示すように、Oracle Big Data Appliance の Hadoop クラスタと、Oracle Exadata で動作する Oracle Database にまたがる問合せの実行を検証することで、ポイントが明確になります。問合せのコンパイル（図の①）で Oracle 内のデータの位置がデータベースによって検出され、HDFS NameNode に対して Hadoop のデータ位置（または InputSplits）の問合せが行われます。これらの位置からグローバル問合せ計画が作成され、

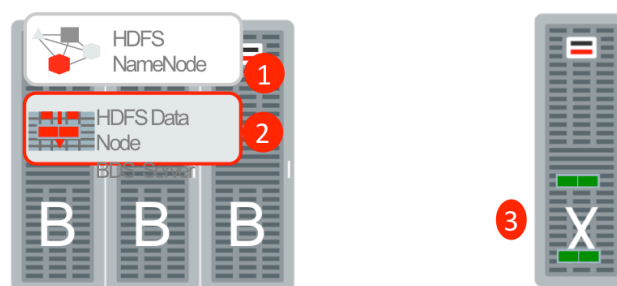


図4：Big Data SQLはSmart Scanテクノロジーを使用して統合問合せを実行します。問合せではHadoopシステムのデータローカルエージェントが実行し、結合は問合せデータベースによって処理されます。

Big Data Applianceの Big Data SQL エージェントと Exadataの Storage Cell に作業がディスパッチされます（図の②）。データに対してローカルである同じ Smart Scan 機能が Exadata 内の Oracle 形式のデータと Hadoop 内に格納された Schema-on-Read データの両方に適用されます。Smart Scan では、関連するデータだけが問合せデータベースに返されるように、多数の操作が適用されます。たとえば、

- WHERE 句の評価
- 列射影
- 結合のパフォーマンスを高めるための Bloom フィルタ
- JSON の解析、データ・マイニング・モデルの評価

Smart Scan の処理後に、データ・ストリームはどちらも Oracle 形式になり、問合せデータベースのコンピュータ・サーバーに返されます。ここで、これらを結合、共用体、PL/SQL ファンクション、高度な分析の対象とすることができます。また、Oracle Database 内のセキュリティまたは管理ポリシーをデータ・ストリームに適用できます。結果的に生成される問合せは、両方のシステムのデータに対して高速で、セキュアで、予測可能なアクセスが行えます。



結論

ビッグ・データのエコシステムの急速な進歩に伴い、Hadoop ストアと NoSQL ストアも急速に幅広い情報管理資産の一部になっています。これらの各システムはビジネスに真の価値を提供することができますが、異なるデータ・アクセス API によって洞察がサイロ化されないよう注意が必要です。多くの組織は統合問合せシステムを通して最大の価値を見出します。統合問合せシステムを使用すると、1 つの問合せで複数のストア（Hadoop、NoSQL、リレーショナル・データベース）のデータにアクセスすることができます。

同じ Smart Scan サービスを提供する軽量のエージェントに問合せ実行をフランチャイズすることにより、Big Data SQL は統合問合せの優れたアプローチを可能にしています。ユーザーとツールはすべてのデータに対して Oracle SQL の性能を存分に活用できます。管理者は既存のセキュリティおよび管理ポリシーの適用範囲を Hadoop または NoSQL データベースに格納されたデータまで容易に拡張できます。企業は成長を続けるビッグ・データのエコシステムの力をシームレスに利用することにより、価値を異種データ・ストアの背後に隠すことなく、迅速に価値を実現できます。



Oracle Corporation, World Headquarters

500 Oracle Parkway
Redwood Shores, CA 94065, USA

海外からのお問い合わせ窓口

電話：+1.650.506.7000
ファクシミリ：+1.650.506.7200

CONNECT WITH US



blogs.oracle.com/oracle



facebook.com/oracle



twitter.com/oracle



oracle.com

Hardware and Software, Engineered to Work Together

Copyright © 2016, Oracle and/or its affiliates. All rights reserved. 本文書は情報提供のみを目的として提供されており、記載内容は予告なく変更されることがあります。本文書は一切間違いがないことを保証するものではなく、さらに、口述による明示または法律による黙示を問わず、特定の目的に対する商品性もしくは適合性についての黙示的な保証を含み、いかなる他の保証や条件も提供するものではありません。オラクルは本文書に関するいかなる法的責任も明確に否認し、本文書によって直接的または間接的に確立される契約義務はないものとします。本文書はオラクルの書面による許可を前もって得ることなく、いかなる目的のためにも、電子または印刷を含むいかなる形式や手段によっても再作成または送信することはできません。

Oracle および Java は Oracle およびその子会社、関連会社の登録商標です。その他の名称はそれぞれの会社の商標です。

Intel および Intel Xeon は Intel Corporation の商標または登録商標です。すべての SPARC 商標はライセンスに基づいて使用される SPARC International, Inc. の商標または登録商標です。AMD, Opteron, the AMD



Oracle is committed to developing practices and products that help protect the environment