

Oracle ホワイト・ペーパー
2014年9月

オラクル： エンタープライズ向けのビッグ・データ

ORACLE®

概要.....	2
はじめに	3
ビッグ・データの定義	3
ビッグ・データの重要性.....	4
ビッグ・データ・プラットフォームの構築.....	5
インフラストラクチャの要件.....	5
急速なテクノロジーの変化.....	6
オラクルのビッグ・データ管理システム.....	7
Oracle Big Data Appliance.....	7
Oracle Big Data SQL	8
Oracle NoSQL Database	9
Oracle Big Data Connectors	9
全データ上でのデータベース内分析.....	11
結論.....	12

概要

今日、ビッグ・データという言葉が注目を集めています。その背景にあるのはシンプルな事情です。何十年間にもわたり、企業はリレーショナル・データベースに格納されたトランザクショナル・データに基づいてビジネスの意思決定を行ってきました。しかしながら、そのような重要なデータの周囲には、あまり構造化されていない新しいデータがあり、宝の山となる可能性を持っています。つまり、Webログ、ソーシャル・メディア、電子メール、センサー、写真などから有益な情報を掘り起こせる可能性があるのです。ストレージおよびコンピューティング能力にかかるコストの低下により、これらのデータの収集が可能になりました。わずか数年前までは投げ捨てられていたようなデータです。その結果、従来とは異なるけれども非常に高い価値を持つ可能性があるデータを、従来のエンタープライズ・データとともにビジネス・インテリジェンス分析に含めようとする企業が増えています。

ビッグ・データから真のビジネス価値を抽出するには、新しいソースからさまざまなデータ・タイプを取得するだけでなく、あらゆるエンタープライズ・データのコンテキスト内でそれを簡単に分析することができる適切なツールが必要です。オラクルでは、Oracle Database、Hadoop、NoSQL データ・ストアのすべてのデータで、Oracle SQLの機能を最大限利用できる独自の製品Oracle Big Data SQLを提供しています。

はじめに

Oracle Big Data Appliance、Oracle Big Data SQL、およびOracle Big Data Connectorsを提供するオラクルは、あらゆる種類のエンタープライズ・ビッグ・データの要件に対応する、完全な統合ソリューションを提供する最初のベンダーとなりました。オラクルのビッグ・データ戦略は、現行のエンタープライズ・データ・アーキテクチャを発展させて、ビッグ・データを取り入れ、Oracle SQLの機能を最大限に活用してすべてのデータからビジネス価値を提供するという考えに基づいています。現行のエンタープライズ・アーキテクチャを発展させることで、高い信頼性、セキュリティ、パフォーマンスを備え、オラクルが新たに提供するBig Dataシステムで拡張したOracleシステムを活用できます。

ビッグ・データの定義

ビッグ・データとは、通常は以下のタイプのデータを指します。

- 従来のエンタープライズ・データ - CRMシステムからの顧客情報、トランザクショナルERPデータ、Webストアのトランザクション、総勘定元帳データなど
- 機械が生成したデータ/センサー・データ - Call Detail Record ("CDR"と呼ばれる詳細通話記録)、Webログ、スマート・メーター、製造センサー、機器ログ (Digital Exhaustと呼ばれることが多い)、取引システム・データなど
- ソーシャル・データ - 顧客フィードバック・ストリーム、Twitterなどのマイクロブログ・サイト、Facebookなどのソーシャル・メディア・プラットフォームなど

McKinsey Global Instituteの見積りによると、データ量は1年あたり40%、2009～2020年の間には44倍に増加すると見られています。これはもっとも分かりやすいパラメータですが、問題となる特性はデータ量だけではありません。実際のところ、ビッグ・データを定義する特性には、以下の4つがあります。

- 量 - 機械が生成するデータは、ほかの従来とは異なるデータよりも非常に大量に作成されます。たとえば、1台のジェット・エンジンで生成されるデータは、30分で10TBにもなる場合があります。1日あたり25,000台以上の航空機が飛行しているため、この1つのデータソースからのデータ量だけでも、1日でペタバイト単位になります。石油製油所や掘削装置などの産業用重機器やスマート・メーターでも同様のデータ量が生成されており、問題が大きくなっています。
- 速度 - ソーシャル・メディア・データ・ストリームは、機械が生成するデータほど大量ではないにしても、カスタマ・リレーションシップ・マネジメントにとって貴重な意見やリレーションシップを大量に生成します。1ツイートあたり140文字だったとしても、高速（または高頻度）のTwitterだとデータ量が多くなります（1日あたり8TB以上）。
- 種類 - 従来のデータ形式は、データ・スキーマによって比較的適切に定義されており、あまり変更されません。それとは対照的に、従来とは異なるデータ形式は目まぐるしく変化しています。新規サービスの追加、新しいセンサーのデプロイ、新規マーケティング・キャンペーンの実施などのたびに、情報を結果として取得するために新しいデータ・タイプが必要となります。
- 価値 - 異なるデータの経済的価値は、多岐にわたります。通常は、従来とは異なる大規模なデータの中に価値のある情報が隠れています。課題は、価値のあるデータを識別し、そのデータを分析できるように変換して抽出することです。

企業がビッグ・データを最大限利用するには、ITインフラストラクチャを発展させて、これらの高速かつ多彩で大量の新しいデータソースを処理し、既存のエンタープライズ・データと統合して分析する必要があります。

ビッグ・データの重要性

ビッグ・データを従来のエンタープライズ・データと組み合わせて抽出して分析すると、企業はビジネスをより徹底的に深く理解できます。それにより、生産性が向上し、競争力が増して、さらなる革新が実現します。そして、これらすべてが最終損益に大幅な影響を与えます。

たとえば、医療サービスの提供では、慢性的または長期にわたる経過の監視に費用がかかります。バイタルサインを測定して進捗を監視する在宅モニタリング装置の使用は、患者の健康状態を向上させ、診療所の訪問数と病院の受入れ数を減らすためにセンサー・データを使用できる1つの方法です。

製造業では、遠隔測定データのストリームを返すために製品にセンサーをデプロイしています。General MotorsのOnStar[®]やRenaultのR-Link[®]などの自動車業界のシステムでは、通信、セキュリティ、ナビゲーションのサービスを提供しています。さらに重要なことに、この遠隔測定データにより、使用パターン、エラー発生率や、開発コストおよび組立てコストを削減できる製品改良のための他の機会も明らかになります。

スマートフォンやその他のGPSデバイスの普及により、広告主は近隣の店、コーヒー・ショップ、レストランにいる消費者を対象に広告を展開できます。これにより、サービス・プロバイダに新しい収益が発生し、多数のビジネスに新しい顧客をターゲットとするチャンスが生まれます。

通常、小売業者は製品を買う顧客については知っています。ソーシャル・メディアやeコマース・サイトからのWebログ・ファイルを利用すると、購入しなかった人やその理由といった、小売業者が現在利用できない情報が分かります。これにより、はるかに効果的な細かな顧客のセグメント化や的を絞ったマーケティング・キャンペーンが可能になり、よりの確な需要計画によって、サプライ・チェーンの効率が向上します。

結局のところ、FacebookやLinkedInなどのソーシャル・メディア・サイトはビッグ・データなしでは成り立ちません。これらのビジネス・モデルには、Web上でのパーソナライズされたエクスペリエンスが必要であり、ユーザーやメンバーに関する使用可能なすべてのデータを取得、使用しないと提供できないためです。

ビッグ・データ・プラットフォームの構築

データウェアハウス、Webストア、ITプラットフォームなどと同様に、ビッグ・データのインフラストラクチャには固有の要件があります。ビッグ・データ・プラットフォームのコンポーネントすべてを考慮する際に重要なのは、ビッグ・データとエンタープライズ・データとを簡単に統合し、組み合わされたデータセットの深い分析を可能にすることが最終的な目標であるということです。

インフラストラクチャの要件

ビッグ・データ・インフラストラクチャの要件は、データの取得、データの体系化およびデータの分析にまで関係します。

ビッグ・データの取得

取得のフェーズは、ビッグ・データ登場前に比べてインフラストラクチャが大きく変化した点の1つです。ビッグ・データは高速度で多様なデータ・ストリームを参照するため、ビッグ・データの取得のサポートに必要なインフラストラクチャには、以下の要件があります。データの取得および短くて簡単な問合せの実行の両方において、短くて予測可能な待機時間を提供すること。多くの場合、分散環境において大量のトランザクションを処理できること。柔軟で動的なデータ構造をサポートすること。

ビッグ・データの取得や保存には、NoSQLデータベースがよく使用されます。NoSQLデータベースは動的データ構造に最適で、非常にスケーラブルです。通常、NoSQLデータベースに保存されるデータは多種多様です。このシステムの用途が、固定スキーマに分類、解析せずに、あらゆるデータを単に取得することだからです。

たとえばNoSQLデータベースは、ソーシャル・メディア・データの収集と保存によく使用されます。顧客が実際に使用するアプリケーションがしばしば変更される一方で、基礎となるストレージ・ストラクチャは単純なままです。エンティティ間のリレーションシップによってスキーマを設計するかわりに、これらの単純な構造には、多くの場合、データ・ポイントを識別する主要キーと、関連データ（顧客のIDやプロフィールなど）を保持するコンテンツ・コンテナのみが含まれます。この単純で動的な構造により、ストレージ・レイヤーにおけるコストのかかる再編成（顧客プロフィールへの新規フィールドの追加など）なしに変更を実施できます。

ビッグ・データの体系化

従来のデータウェアハウジングの用語では、データの体系化はデータ統合と呼ばれています。ビッグ・データの量が非常に大きいため、データを初期の宛先の場所で体系化し、大量のデータを移動させないことによって時間と費用を節約する傾向があります。ビッグ・データの体系化に必要なインフラストラクチャの要件は以下のとおりです。データを元のストレージ・ロケーションで処理および操作できること。大規模なデータ処理手順に対処できる非常に高いスループット（多くの場合バッチ処理）をサポートすること。非構造化データから構造化データまで、多様なデータ形式に対処できること。

Hadoopは、データを元のデータ・ストレージ・クラスタに維持しながら大量のデータを体系化して処理できる新しいテクノロジーです。たとえばHadoop Distributed File System (HDFS) は、Webログの長期間のストレージ・システムです。これらのWebログは、クラスタ上でMapReduceプログラムを実行し、集計された結果を同じクラスタ上で生成することにより、ブラウジング動作（セッション）情報に変換されます。これらの集計結果が、リレーショナルDBMSシステムにロードされます。

ビッグ・データの分析

データの移動は常に体系化段階で行われるとは限らないため、分析も分散環境で行われる場合があります。その場合、データによっては、最初に格納されていた場所にとどまり、データウェアハウスから透過的にアクセスされることになります。ビッグ・データの分析に必要とされるインフラストラクチャの要件は以下のとおりです。統計分析やデータ・マイニングなどのより深い分析を、多様なシステムに保存されたさまざまなデータ・タイプでサポートできること。きわめて大きなデータ量へ拡張できること。動作の変更による応答時間を高速化できること。そして、分析モデルに基づく決定を自動化できること。もっとも重要なのは、インフラストラクチャによって、ビッグ・データと従来のエンタープライズ・データとを組み合わせた上で分析を統合できる必要があるということです。新しい発見は、新しいデータを分析することだけで得られるのではなく、新しいデータを古いコンテキストの中で分析し、古い問題に新しい視点を与えることによって得られるものです。

たとえば、自動販売機の在庫データを自動販売機が置いてある場所のイベント・カレンダーと組み合わせることで分析することにより、その自動販売機にとって最適な製品構成と補充スケジュールが決まります。

急速なテクノロジーの変化

上記のITインフラストラクチャ要件に対応するため、多くの新規テクノロジーが生まれています。最新の調査によれば、ビッグ・データの取得に使用されるオープンソースのkey-valueデータベースは120を超えています。Hadoopは、ビッグ・データを体系化する主要システムとして登場し、リレーショナル・データベースはデータウェアハウスとしての地位を維持しつつ、ビッグ・データを分析するために構造化が進んでいないデータセットに範囲を広げています。

このような新しいシステムは当初、データ・アクセスのための専用APIを使用する傾向にあり、テクノロジー・ランドスケープを分散させる原因となりました。言語としてのSQLは、NoSQLという名前が示しているように、あまり使用されなくなりました。

このような傾向は最近になって完全に逆転し、ビッグ・データやNoSQL領域において最も重要なことは、このようなkey-valueストアやNoSQLストア（Hadoopなど）でSQLが使用できることです。

オラクルのビッグ・データ管理システム

オラクルは、あらゆる種類のエンタープライズ・ビッグ・データの要件に対処する、完全な統合データ管理ソリューションを提供する最初のベンダーです。オラクルのビッグ・データ管理システムは、エンジニアド・システムで現行のエンタープライズ情報アーキテクチャを発展させることにより、ビッグ・データを取り入れ、Oracle SQLの機能を使用してこれらのシステムのすべてのデータに対応するという考えに基づいています。HadoopやOracle NoSQLなどのテクノロジーをOracleデータウェアハウスと同時に実行することで、現行のOracleシステムのセキュリティ・ポリシーに従いつつ、ビジネス価値を統合および拡大し、ビッグ・データの要件に対応します。

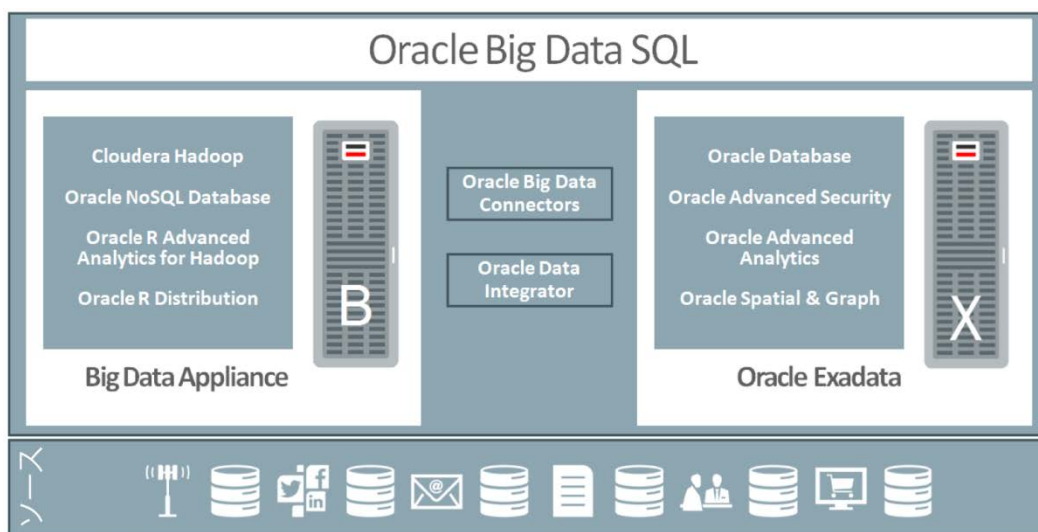


図1 オラクルのビッグ・データ管理システム

Oracle Big Data Appliance

Oracle Big Data Applianceは、最適化されたハードウェアと包括的なビッグ・データ・ソフトウェア・スタックを組み合わせたエンジニアド・システムで、ビッグ・データを取得、体系化するための完全なデプロイしやすいソリューションを提供します。

Oracle Big Data Applianceの拡張可能なラック構成により、開発と本番のワークロード問題を解決し、企業全体のデータ・ニーズの増加に合わせて拡張することができます。

Oracle Big Data Applianceでは、オープン・ソース・ソフトウェアと、エンタープライズ・ビッグ・データの要件に対応するためにオラクルが開発したソフトウェアとが組み合わされています。

Oracle Big Data Applianceソフトウェアには、以下が含まれます。

- Cloudera Managerを含むCloudera Enterprise (Data Hub Edition)
- Oracle Big Data SQL¹
- Enterprise Manager用Oracle Big Data Applianceプラグイン
- オラクルによる統計パッケージRのディストリビューション
- Oracle NoSQL Database Community Edition²
- Oracle Enterprise Linuxオペレーティング・システムとOracle Java VM

Oracle Big Data SQL

Oracle Big Data SQLは、ビッグ・データ・ソースへのアクセスと統合を簡略化するオラクルの画期的なアプローチで、Hadoop、NoSQLデータ・ストア、またはOracle Databaseのすべてのデータを単一のSQL文で問合せできる機能を提供します。また、拡張外部表としてHadoopや他のソースを提供します。この機能は、Oracle Database 12.1.0.2以降で利用可能です。これらの表は、データ・アクセスの外部セマンティクス(並列度、位置、スキーマ)をOracle内部に透過的にマッピングするよう設計されています。

このマッピングにより、アクセスおよびネイティブ・プロセッシング全体が可能な限り最適化されます。Oracle Big Data SQLによって以下のことが実現できます。

- 世界で最も機能豊富なSQL言語を使用して、すべてのデータで問合せを表現
- 既存のインタフェースを使用して、ビッグ・データを即座にレポートまたはアプリケーションに組み入れる
- 既存のOracleセキュリティおよびアクセス制御ポリシーをHadoopに保存されているデータへ拡張

ビッグ・データは膨大になる可能性がある一方で、特定の問合せに関連するデータ量が桁違いに少ないことはよくあることです。これによって、問合せのパフォーマンスが大幅に最適化される機会が得られます。Oracle Exadata Storage Servers SoftwareのSmart Scan for Hadoopの次の機能は、Oracle Big Data SQLのパフォーマンスを最大化します。

- データのローカル・スキャン：ストレージ側でデータの読み取りと処理をします。
- 述語の評価と射影：関連データのみHadoopから転送されます。
- 複合解析：JSONやXMLなどのデータはソース側のローカルで処理されます。
- Bloomフィルタ：最適化された結合がHadoopのBloomフィルタに変換されます。

¹ Oracle Big Data SQLは、Oracle Big Data Appliance上でのみ利用でき、別ライセンス・コンポーネントです。

² Oracle NoSQL Database Enterprise Editionは、Oracle Big Data Appliance用の別ライセンス・コンポーネントとして入手できます。

Oracle NoSQL Database

Oracle NoSQL Databaseは、Oracle Berkeley DBに基づく、非常にスケーラブルな分散key-valueデータベースです。分散Berkeley DB上にインテリジェント・ドライバを追加することで、汎用のエンタープライズ・クラスのkey-valueストアを提供します。このインテリジェント・ドライバは、基礎になるストレージ・トポロジを追跡し、データをシャーディングして、最小の待機時間でデータを置く場所を示します。Oracle NoSQL Databaseは競合他社のソリューションとは異なり、インストール、構成、管理が簡単で、幅広いワークロードをサポートしており、エンタープライズ・クラスのOracleサポートにより、エンタープライズ・クラスの信頼性を実現します。

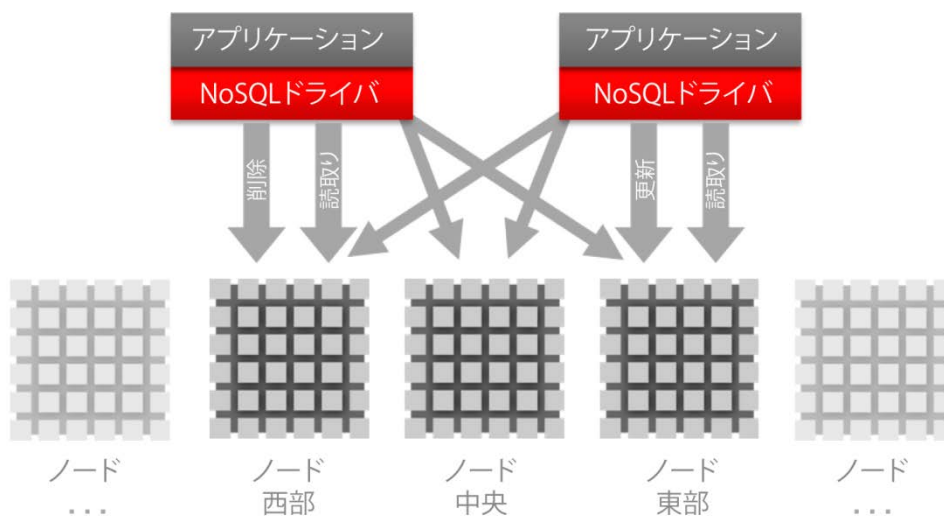


図2 NoSQL Databaseのアーキテクチャ

Oracle NoSQL Databaseのおもなユースケースは、短い待機時間でのデータ取得と、一般にはキー検索によるそのデータの高速の問合せです。Oracle NoSQL Databaseには、使いやすいJava APIおよび管理フレームワークが含まれています。この製品は、オープン・ソース・コミュニティ・エディションか、有償のエンタープライズ・エディション（大規模な分散データセンター用）として入手できます。オープンソース・コミュニティ・エディションは、Big Data Appliance統合ソフトウェアの一部としてインストールされています。

Oracle Big Data Connectors

Oracle Big Data Applianceによって新しい種類のデータを簡単に取得、体系化できる一方、Oracle Big Data Connectorsによってビッグ・データ環境をOracle ExadataやOracle Databaseと緊密に統合できるため、すべてのデータを一緒に高いパフォーマンスで分析できます。Oracle Big Data Connectorsは、次の4つのコンポーネントで構成されます。

Oracle Loader for Hadoop

Oracle Loader for Hadoopにより、ユーザーはHadoop MapReduce処理を使用して、Oracle Database 11gへの効率的なローディングを通して、分析のために最適化されたデータセットを作成できます。他のHadoopローダーと異なり、Oracle Loader for HadoopはOracleの内部フォーマットを生成して

データを高速でロードし、データベース・システム・リソースの使用をより少量に抑えます。Oracle Loader for Hadoopは、MapReduce処理のマップ - パーティション - リデュースのステップの最後に追加されます。この最後の手順では、Hadoopクラスタ内のCPUを使用して、オラクルの内部データベースの形式にデータをフォーマットします。このため、Oracle DatabaseプラットフォームのCPU利用率が下がり、データの取得速度が上がります。いったんロードされたデータは永続的にデータベースに保存され、SQLまたはビジネス・インテリジェンス・ツールを利用する一般のデータベース・ユーザーからきわめて高速でデータにアクセスできるようになります。

Oracle SQL Connector for Hadoop Distributed File System

Oracle SQL Connector for Hadoop Distributed File System (HDFS) は、Oracle DatabaseからHDFS上のデータに直接アクセスするための高速コネクタです。Oracle SQL Connector for HDFSにより、ユーザーはアプリケーションの必要性に応じて、HDFSへいつでも柔軟に問合せができます。

つまり、これにより、Oracle Databaseで外部表の作成が可能になり、HDFSに格納されたデータへのSQLによる直接アクセスが可能になります。HDFSに保存されているデータはSQL経由で問い合わせられ、Oracle Databaseに保存されているデータと結合したり、Oracle Databaseにロードしたりできます。HDFS上のデータへのアクセスを最適化し、自動ロードバランシングによって高速に移動、パラレル化できます。HDFSへは、デリミタ付きファイルまたはOracle Loader for Hadoopによって作成されるOracle Data Pumpファイルを保存できます。

Oracle Data Integrator Application Adapter for Hadoop

Oracle Data Integrator Application Adapter for Hadoopは、Oracle Data Integratorの使いやすいインタフェースを通して、HadoopとOracle Databaseとのデータ統合を容易にします。データベースでデータがアクセス可能になると、エンドユーザーはSQLとOracle BI Enterprise Editionを使用してデータにアクセスできます。

Hadoopソリューションをすでに使用しており、Oracle Big Data Applianceなどの統合製品が不要な企業は、Big Data Connectorsをスタンドアロンのソフトウェア・ソリューションとして使用して、HDFSからのデータを統合できます。

Oracle R Connector for Hadoop

Oracle R Connector for Hadoopは、HadoopやHDFSに保存されたデータに透過的にアクセスするためのRパッケージです。

R Connector for Hadoopにより、オープンソースの統計環境Rのユーザーは、HDFSに格納されたデータを分析し、MapReduce処理を利用する大量のデータに対してRモデルをスケーラブルに実行できます。この際、Rのユーザーは、別のAPIや言語を学ぶ必要はありません。ユーザーは3,500を超えるオープンソースのRパッケージを使用してHDFSに格納されたデータを分析できますし、管理者は本番環境でR MapReduceモデルをスケジューリングするためにRを学ぶ必要はありません。

R Connector for Hadoopは、Oracle DatabaseのOracle Advanced Analytics Optionと一緒に使用できます。Oracle Advanced Analytics Optionを使用すると、RユーザーはSQLやデータベースの概念を学ぶ必要なく、R演算をデータベース内で直接実行することによってデータベースのデータを透過的に処理できます。

全データ上でのデータベース内分析

Oracle Exadataと組み合わせて使用するOracle Big Data Applianceのデータは、Oracle Big Data SQLにより、完全にSQLに対応しています。この独自機能により、エンドユーザーは、両方のデータ・ストアのデータにSQL分析機能をフルに使用できるため、独自のビジネス価値を創出できます。

また、次の使いやすいツールを使用すれば、データの分析時にデータベース内で高度な分析を行うことができます。

- Oracle R Enterprise – 広く使用されているProject R統計環境のOracleバージョンを使用すると、統計担当者は、エンドユーザーのエクスペリエンスを変更することなしにRを非常に大きなデータセット上で使用できます。Rは、特定の空港の飛行機の遅延予測や、臨床試験の分析と結果の提出などに使用されます。
- インデータベース・データ・マイニング - 複雑なモデルを作成し、それらを非常に大量のデータ上へ配置して予測分析を行うための機能。エンドユーザーは、モデルのビルド方法を知らなくても、BIツールでこれらの予測モデルを使用できます。たとえば、リグレーション・モデルを使用することにより、購買行動や人口統計データに基づいて顧客の年齢を予測できます。
- インデータベース・テキスト・マイニング - Oracle Text とOracle Data Miningを組み合わせ、マイクロ・ブログ、CRMシステム・コメント・フィールドおよびレビュー・サイトからテキストをマイニングする機能。テキスト・マイニングの例は、コメントに基づく感情分析です。感情分析では、顧客がある特定の企業、製品、アクティビティなどにどのような感情を抱いているかを明らかにしようとします。
- インデータベース・グラフ分析 - 各種のデータ・ポイントとデータセットのグラフを作成して、これらの間の関係性を分析する機能。たとえば、グラフ分析では、顧客の親しい友人の価値を判断するリレーションシップのネットワークを作成します。顧客離れを見ると、ある顧客が離れた場合に、その周辺にいる関係性を重視している顧客まで失う可能性があります。
- インデータベース・Spatial – データに空間の次元を追加し、地図にデータを表示する機能。この機能により、エンドユーザーは地理上の関係およびトレンドを非常に効率的に認識できます。たとえば、空間データにより、人々のネットワークおよび地理的近接性を視覚化できます。きわめて近い場所にいる顧客は、すぐに互いの購買行動に影響を与えることが可能で、空間の視覚化がない場合、機会が簡単に失われる可能性があります。

Oracle Databaseの分析コンポーネントには、それぞれ価値があります。これらのコンポーネントを組み合わせることにより、ビジネスにさらに大きな価値が創造されます。SQLやBIツールを利用してこれらの分析結果をエンドユーザーに公開することで、Oracle Databaseの分析の潜在能力を一部しか利用していない他社より優位に立つことができます。

Oracle Big Data ApplianceとOracle ExadataはInfiniBand経由で接続されるため、バッチや問合せのワークロードで、データを高速に転送できます。Oracle Big Data SQLは、Smart Scan on Hadoopデータを提供することによって、問合せのパフォーマンスを大幅に高めます。

結論

新しいさまざまなデジタル・データ・ストリームを分析することにより、経済的価値の新しいソースを明らかにし、顧客の行動について新しい見解をもたらして、市場動向を早めに特定できます。しかし、この新しいデータの流入がIT部門の課題となっています。ビッグ・データから真のビジネス価値を抽出するには、異なるソースからさまざまなデータ・タイプを取得して体系化し、あらゆるエンタープライズ・データのコンテキスト内でそれを簡単に分析するための適切なツールが必要です。Oracle Big Data ApplianceおよびOracle Big Data ConnectorsをOracle Exadataとともに使用することによって、企業は構造化データおよび非構造化データを含むあらゆるエンタープライズ・データを取得、体系化および分析し、最適な決定を行うことができます。



オラクル：エンタープライズ向けのビッグ・データ

2014年9月

著者：Jean-Pierre Dijcks

Oracle Corporation
World Headquarters
500 Oracle Parkway
Redwood Shores, CA 94065
U.S.A.

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
oracle.com



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

Copyright © 2014, 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、ここに記載される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。オラクル社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。本文書はオラクル社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は米国Oracle Corporationおよびその子会社、関連会社の登録商標です。Cloudera、Cloudera CDH、およびCloudera Managerは、Cloudera, Inc.の登録商標および未登録商標です。その他の名称はそれぞれの会社の商標です。

0109