

ビッグデータ利活用を支える並列分散処理と類似検索の高速化技術

Technology for Speeding up Parallel Distributed Processing and Similarity Search to Support Utilization of Big Data

● 土屋 哲 ● 上田晴康 ● 此島真喜子

あらまし

ビッグデータという言葉が広まり始めて数年が経った。ビッグデータの利活用は、はじめは消費者系サービスで使われたが、近年は企業にも導入されてきている。従来のICTシステムでは扱えなかったデータを処理することで、ヒット商品を開発したり、事故や危険の防止に利活用したりする事例が増えてきている。企業におけるビッグデータの利活用では、消費者系サービスとは異なり、業務データを中心とし、センサーや画像など現場で新しく得られるようになった現場データ、あるいはソーシャルネットワークサービスや自治体のオープンデータなどの外部データと組み合わせて新しい知見を得ることが必要となる。その際、まず基本として必要となるのは、複数のデータ系列の結合を高速に処理することや、センサーや画像から必要な情報を検索することである。

本稿では、企業内にある大量の業務データのバッチ処理を分散して処理することで高速化するHadoop業務データ活用技術と、センサーや画像などの類似性検索を従来よりも大幅に高速化する高速近似類似検索技術を紹介する。

Abstract

Several years have passed since the phrase “big data” began to become widespread. Utilization of big data was first introduced in consumer services but it has recently been spreading among enterprises. There have been an increasing number of cases in which data that could not be handled by the conventional information and communications technology (ICT) systems are processed for use in developing good-selling products and preventing accidents or hazards. Utilization of big data in enterprises, unlike that for consumer services, involves use of business data as the basic information to be combined with on-site data from sensors or images, which have now become available on site, or with external data, such as open data from social networking services or local governments, in order to obtain new knowledge. The basic requirements in that process include high-speed processing of a combination of multiple data series and retrieval of the necessary information from sensors and images. This paper presents Hadoop business data utilization technology, which accelerates batch processing of large volumes of business data in enterprises by means of distributed processing, and high-speed approximate similarity search technology that achieves a significant speed-up of similarity search of sensors and images compared with the conventional method.

まえがき

一般に、ビッグデータを特徴づける性質は、Volume（大量）、Velocity（高頻度）、Variety（多様性）の三つであると言われている。従来のICTシステムでは取り扱えなかったデータを処理し、そこから新しい知見を得ることで、ヒット商品を開発したり、新規事業を開拓したり、事故や危険を防止する事例が近年増えてきている。

こうしたことを容易にしたのが、Hadoop⁽¹⁾に代表される大規模データ処理を実現するソフトウェアである。整然と構造化されていないサーバログなどのデータであっても、安価な計算機を多数並べて高速に処理することで、新しい知見が得られるようになってきた。

本稿では、ビッグデータの処理に関する動向を紹介し、富士通研究所が開発した企業システム向けデータ処理技術である、Hadoop業務データ活用技術と高速近似類似検索技術を紹介する。

ビッグデータと企業システム

ビッグデータという言葉が広まり始めて数年が経ち、企業での活用も広がってきている。これまでの企業システムは、データ形式が明確に規定された各種トランザクションをRDB（Relational Database）に蓄積・管理し、それに対して各種の集合演算によって、絞り込み、集計、結合を繰り返して業務に必要な計算を行うことが多かった。ビッグデータの利活用では、サーバログや電子メールのテキスト、あるいは画像や音声など、データ形式が事前に明確には規定されていないものや、項目の欠損などが多いもの、更にデータ量も数テラバイトと膨大なものを分析することで、個人が求める商品を調べて売上を伸ばしたり、倉庫や工場でのコストを削減したりといった施策を行う。データの増加率をみると、企業内トランザクションの基となる業務データも増加傾向ではあるが、ログや画像、音声などのデータ量はより急速に増加している。したがって、増え続けるビッグデータの有効な利活用が今後の企業での重要課題となっている。

データの発生源という観点からは、従来の業務用データが帳票のように企業内で閉じていたのに

対し、ビッグデータは産業・流通の現場や外部のソーシャルネットワークサービス、または政府や自治体などのオープンデータのように企業の外部から収集してくるものである。これらをまとめると「業務データ」「現場データ」「外部データ」に分類できる。

これまでビッグデータは、Web検索サービスプロバイダーなどで主に使われてきたが、最近は企業システムでの利活用が進み、今後は更に事例が増加すると思われる。企業が目的とする業務の改善、オムニチャネル対応などマーケティングの精度向上、新商品の開発といった用途では、Web検索サービスのような特定用途とは異なり、基本となる業務データに対して新しい二つのデータタイプ「現場データ」「外部データ」を連携して扱い、相互のつながりを発見することが必要である。

膨大なデータの結合や検索

業務データに対して、現場データや外部データとの相互のつながりを見つけるためには、膨大なデータ間の結合処理、あるいはセンサーやメディアからのデータの蓄積および高速な検索が必要である。

富士通研究所では、こうした企業システムの視点に立ったビッグデータ利活用において、重要な役割を果たす技術を研究開発している。一つは「Hadoop業務データ活用技術」である。これは、大量の業務データを分散処理し、バッチ処理を高速化する技術である。もう一つは「高速近似類似検索技術」である。これは、センサーやメディアからのデータを従来の100～1000倍以上高速に検索する技術である。

次章以降、これらを説明していく。

Hadoop業務データ活用技術

ビジネスにおけるHadoopの利用目的として、特に多いのが業務データのバッチ処理（以下、業務バッチ）の高速化である。Hadoopを使い社内にある大量の業務データを分散処理することで、バッチ処理を高速化できる。このようなデータ処理の需要の高まりを受けて、富士通研究所では、Hadoopマルチプレクサ技術⁽²⁾と処理時間の平準化技術⁽³⁾の二つの技術を開発した。これらを組み

合わせることで、他社のHadoopディストリビューションよりも、低コストかつ短時間で業務バッチの実行を終了できる。

これらの技術は、FUJITSU Software NetCOBOL Enterprise Edition V10.5⁽⁴⁾ およびInterstage Big Data Parallel Processing Server V1.1⁽⁵⁾ として提供され、また、FUJITSU Business Application Operational Data Management & Analytics^{(6),(7)} を支える技術としても利用されており、ビジネスで活用されている。⁽⁸⁾

● Hadoopマルチプレクサ技術

この技術は、複数の入出力のある業務バッチの処理時間をHadoopを用いて短縮するとともに、開発コストを下げるための技術である。複数ファイル入出力アプリケーションを修正することなくHadoop上で実行可能なため、Hadoop特有の開発スキルを持たなくても既存アプリケーションを利用したり、バッチアプリケーションを開発したりできる。

従来、業務アプリケーションをHadoop上で分散処理するには、以下の二つの方法のいずれかを取る必要があった。一つ目は、業務アプリケーションをHadoopのパラダイムに合わせて開発し直す方法である。しかし、Hadoopのパラダイムに合わせた開発のハードルは高く、開発スキルを備えた人材が少ない。そのような人材がない場合には、Hadoopの開発スキルのある企業に開発を依頼しなければならないため、コストが高くなる。二つ目は、

Hadoopの機能の一つであるHadoop Streaming⁽⁹⁾ を使う方法である。Hadoop Streamingは、標準入出力経由で必要なデータを受け渡しするアプリケーションであれば、コードを修正することなくHadoop上で並列実行できる。しかし、標準入出力をオプション指定や診断出力のために使う既存アプリケーションは実行できない。また、業務バッチでよく使用される突合せ処理を行う業務アプリケーションは、二つ以上のファイルを入出力するため、Hadoop Streamingでは実行できない。以上のことから、Hadoop上で分散処理を行う場合、適用可能なアプリケーションが大幅に限定されてしまう。

一方、Hadoopマルチプレクサ技術は、複数ファイルの入出力に対応しつつ、Hadoop Streaming同様に一般のアプリケーションをHadoop上で実行可能にする。複数ファイルの入出力を可能にするため、プロセス間通信の一種である名前付きパイプを利用し、そのパイプとの情報交換を円滑に行うために内部で並列処理している。アプリケーションが読み書きするファイルのそれぞれに対して名前付きパイプを用意し、Hadoopとのデータのやり取りをそれらのパイプ経由で行うため、大量のデータをやり取りする際でもディスク入出力を行うことなく高速に並列処理ができる。

業務アプリケーションを例に、Hadoopマルチプレクサ技術を用いて並列実行する仕組みを示したのが図-1である。

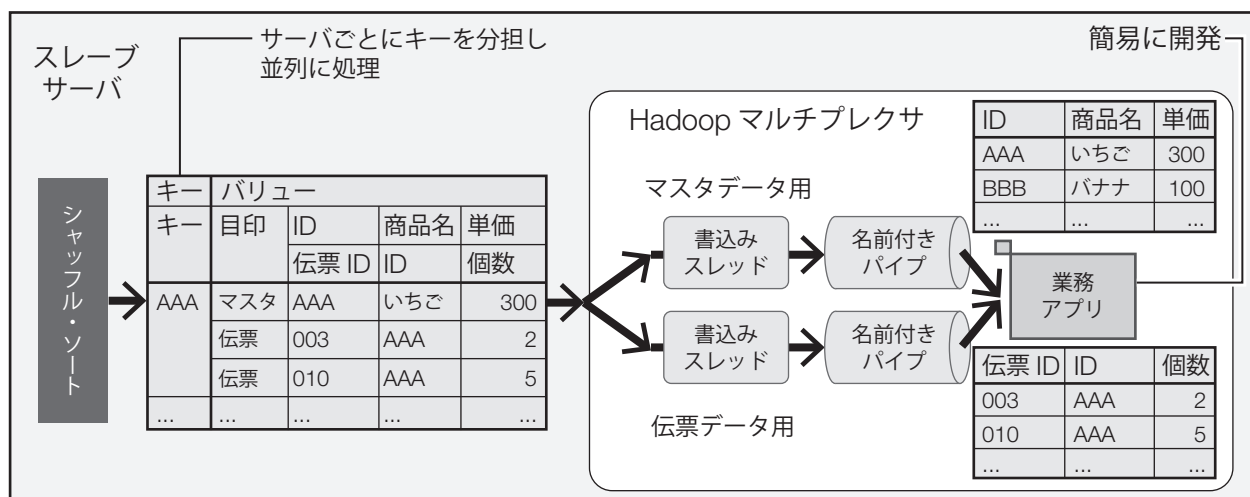


図-1 Hadoopマルチプレクサ技術

● 処理時間の平準化技術

この技術は、各サーバにおけるHadoopの処理時間を短縮する技術である。Hadoopは、主キーの数を平準化することで各サーバの処理時間を平準化できる。しかし、一部の主キーに属するレコード数がほかの主キーのものと比べて極端に多い場合、各サーバの処理時間がデータ量に比例してばらつくため、全体のバッチ処理時間が遅くなることがある^{(1),(10)}。特に、Hadoopマルチプレクサ技術の主な適用事例である突合せ処理では、一般に主キーの種類が少ないため出現数の偏りが大きくなりやすい。例えば、金融業で店舗ごとに集計処理をするような場合は、店舗によって口座数が大きく異なるためにこのような処理の偏りが起きる。

これを解決するために開発したのが処理時間の平準化技術である。入力データのレコードに含まれる主キーを事前に調べ、各主キーに属するレコード数をカウントする。このカウントされた数に基づき、バッチ処理実行時にHadoopのチューニングパラメーターの一つである並列度を自動的に決定する。その上で、各サーバでの実行順序とデータ配置を最適化することで分散処理時間を均等化し、バッチ処理全体の処理時間を短縮する。

なお、一部の企業からHadoopを使った分散処理の処理時間の平準化技術が公開されているが、富士通研究所の平準化技術は実行時に生じる想定外の実行時間のばらつきに対しても有効である。通常、主キーの出現数を基に各サーバが処理するデータ数が均等になるように処理をスケジューリングする。この方式では、実行中に生じる想定外の処理遅延が発生すると、処理時間の平準化が崩れてしまう。そこで、あえて小数のデータグループを複数作成し、このデータグループを処理するプロセスのスケジューリングをバッチの最後に行うことにした。これによって、想定外の処理遅延が発生したとしても、遅延が発生していない別のサーバにプロセスをスケジューリングできる。その結果、全体として処理時間を平準化できた。

● 性能評価

まず、既存の業務アプリケーションを使用した場合と、Hadoopマルチプレクサ技術を使用して既存の業務アプリケーションを並列処理した場合の性能を評価した。従来のバッチ方式とHadoopマル

チプレクサの処理の各多重度（1・2・4）における、処理時間と入力データサイズの関係を図-2に示す。実行した業務アプリケーションは、二つのファイルを突き合わせるものである。

データサイズが1～4 Gバイトのとき、従来のバッチ形式と比較し性能が劣化するが、もしくは同等程度の性能であった。しかし、8 Gバイトより大きな入力データでは多重度の増加に合わせて性能向上が見られた。このように、既存バッチ処理をHadoopに適用すると、一定のデータサイズ以上の場合に非常に効果的である。また、データサイズが大きい場合は、従来のバッチ形式と同じ1多重の場合でも大きく性能向上している。これは、業務アプリケーションの実行前にHadoopによってデータの分割とソートが並列に実行されたことが寄与している。

次に、Hadoopマルチプレクサ技術のみを使用してバッチ処理を実行した場合と、Hadoopマルチプレクサ技術と処理時間の平準化技術を組み合わせて実行した場合の処理時間の比較評価結果を示す。実行するアプリケーションは同じく二つのファイルを突き合わせるものである。処理時間の平準化技術は、入力データに含まれる主キーの出現数の偏りによって異なるため、入力データを複数用意し評価した。用意したデータは、Hadoopマルチプレクサを使って定期的にバッチ処理を行っている金融業におけるデータ2種類と、社内で夜間バッチによって定期的に処理される経理データ2種類であ

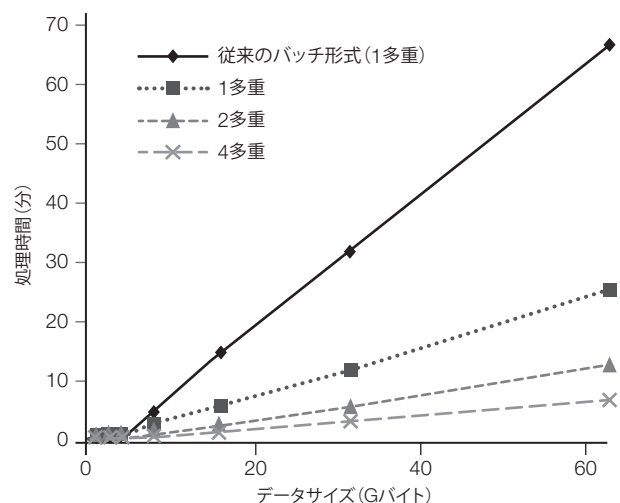


図-2 データ量と処理時間の関係

る。金融業のデータのの一つにおいて、最大で30%処理時間が短縮した(図-3)。また、経理業務のように、処理時間が100秒程度と短いバッチ処理に対しても、22.3%の処理時間が短縮できることを確認した。Hadoopマルチプレクサ技術のみを使用して実行した際に、サーバ間の処理時間の差が大きいケースほど実行時間の短縮効果が大きい。

● Hadoop業務データ活用技術の今後

本章で紹介した技術の利用分野は、単に業務データを用いるデータ処理の高速化にとどまらない。今後ますます重要度が高まる外部データや現場データを組み合わせたビッグデータを分析処理する際に、開発工数を抑えながら高速に処理するための必須の技術である。

現場データの高速近似類似検索技術

本章以降では、富士通研究所が独自開発したセンサーやメディアからの現場の類似検索を、従来の100～1000倍高速に実行する高速近似類似検索技術と、その適用事例について説明する。

類似検索とは、探したいデータの値に類似しているデータをデータベースから検索することである。典型的な対象は、現場に設置された無数のセンサーから収集される膨大なデータである。これらは多数の属性から成り、連続値をとる。例えば温度や湿度、動体の位置情報、画像、音声などが

挙げられ、レコード数は数億件以上、属性数も数百となる場合がある。したがって、IoT時代のデータの利活用において高速な検索が求められる。

従来の類似検索との比較を図-4に示す。データは、取得したままの値を用いるか、検索目的に合った特徴を生データから抽出し、特徴量ベクトルとして用いる。いずれも整数や浮動小数点などの連続値となり、非類似度は二つのデータ間の距離で表わされる。従来は、特徴量ベクトルをそのまま使用したが、この類似検索技術ではビットベクトルというビット列に変換して近似し、検索処理そのものを高速化する。その理由を以下に述べる。

表-1は、類似検索全体の技術マップである。木構造を持つB-Treeやkd-treeなどが従来の主な検索手段である。kd-tree⁽¹¹⁾は、空間を探索する際の半径を調整することで、検索するレコード数を削減する。この方法では、属性数が多くなると単位空間あたりのデータの個数が極端に少なくなるため、空間を探索する際の半径を調整し、検索するレコード数を削減することが困難となる。一般的に、10数個の属性になるとレコード数を削減できなくなるため、全件検索と同様になることが知られている。特にセンサーデータや音声など、時系列性のあるデータに対して、長い時間幅で横通しして特徴量を作るため、属性数は数百と大きくなることが多い。

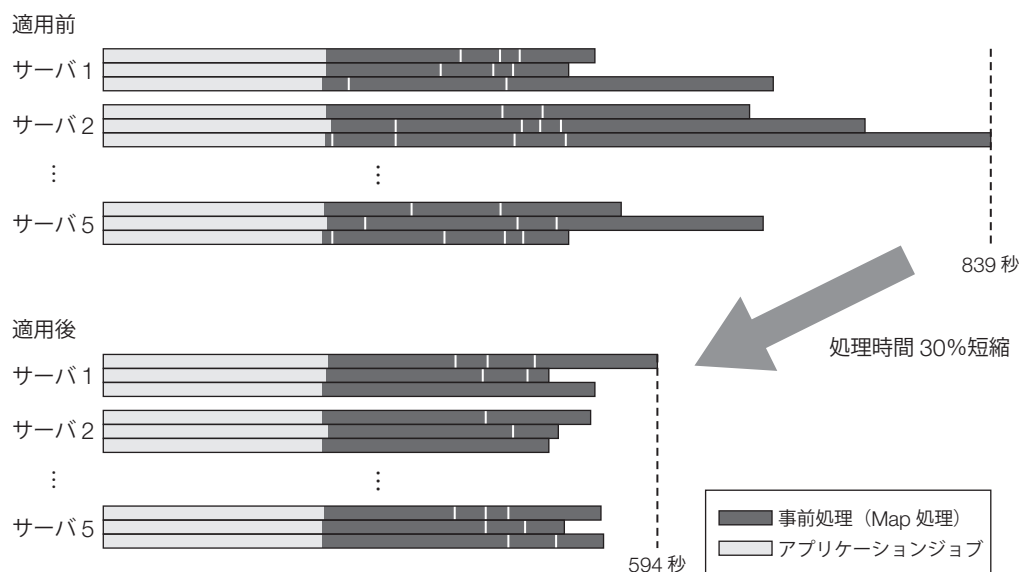


図-3 処理時間の平準化技術による高速化の効果

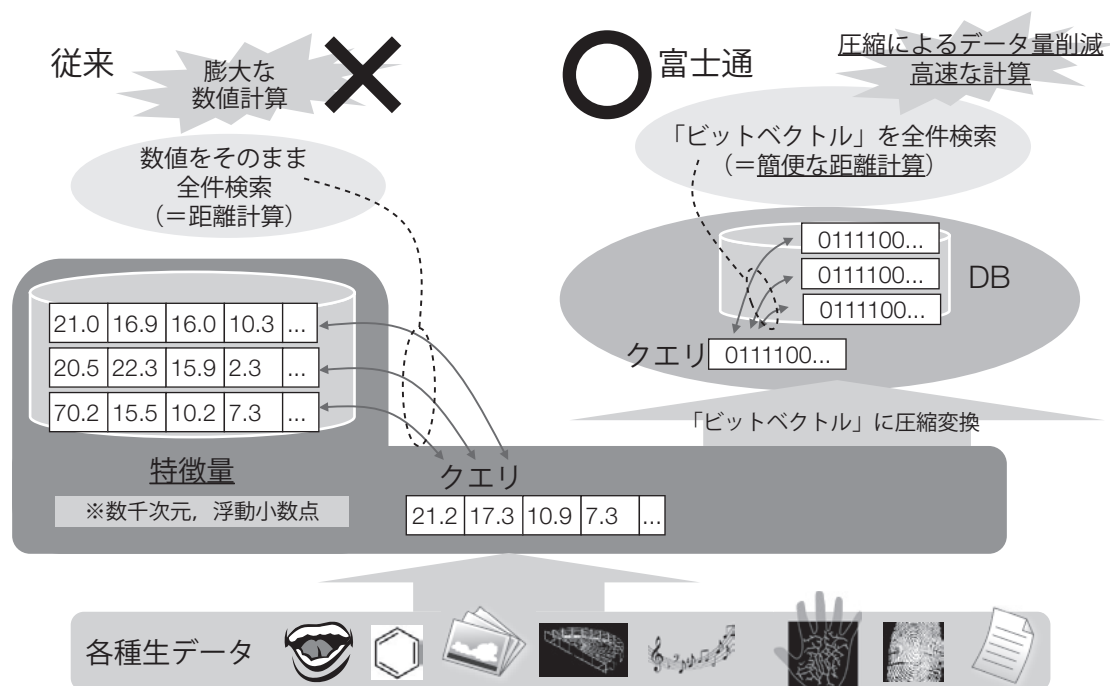


図-4 従来の類似検索との差

表-1 類似検索とその周辺技術の全体像

	近似なし	近似あり（丸めを行うため必ずしも近傍とは限らない）
木構造	B-Tree/kd-Tree	B-Tree/kd-Tree探索処理を途中で打ち切る
ハッシュ	D-index	富士通研究所独自技術

次章以降、全件検索することを前提に、検索処理そのものを高速化する技術について説明する。

富士通研究所の高速近似類似検索技術

富士通研究所では、検索処理そのものを高速化する方法として、Locality-Sensitive Hashing (LSH) をベースとした技術を考案した（表-1の右下部分に相当）。

既存のアルゴリズム^{(12),(13)}の中で、検索処理を高速化できるアイデアを基に、実用に耐え得る検索精度となるように、以下に述べるような独自の技術を開発した。

アイデアのベースは、浮動小数点からなる特徴量ベクトルを、短い長さのビットベクトルに変換し、元の特徴量のデータ配置を大まかに反映したインデックスとして用いることである。ビットベクトルへの変換方法は、まず空間を平面で分割し、データが平面のどちら側にあるかで0または1を割

り当てる。分割する平面ごとに0または1を割り当ててビットを決定し、実際の位置を近似する。二つのデータ間の距離は、同じビット位置の値が相違する個数、すなわちハミング距離で表される。分割を行う平面は複数であり、平面数が多いほどデータの座標の精度は高くなるが、ビット数は大きくなる。

筆者らの経験上、多くても数千ビットで特徴量が十分表現できることが多い。そして、ハミング距離計算は通常距離計算より容易である。また、ビットを決定する際の計算は、基本的に行列による線型な計算を行い、結果の符号で0または1を決定するため高速である。これらのことから、筆者らのアプローチは大量のデータのビットベクトル生成と全件検索の双方に最適である。

従来のアプローチでは、インデックスを作るために多くても32ビット程度のビットベクトルを目標としたものが大半であった。しかし筆者らは、

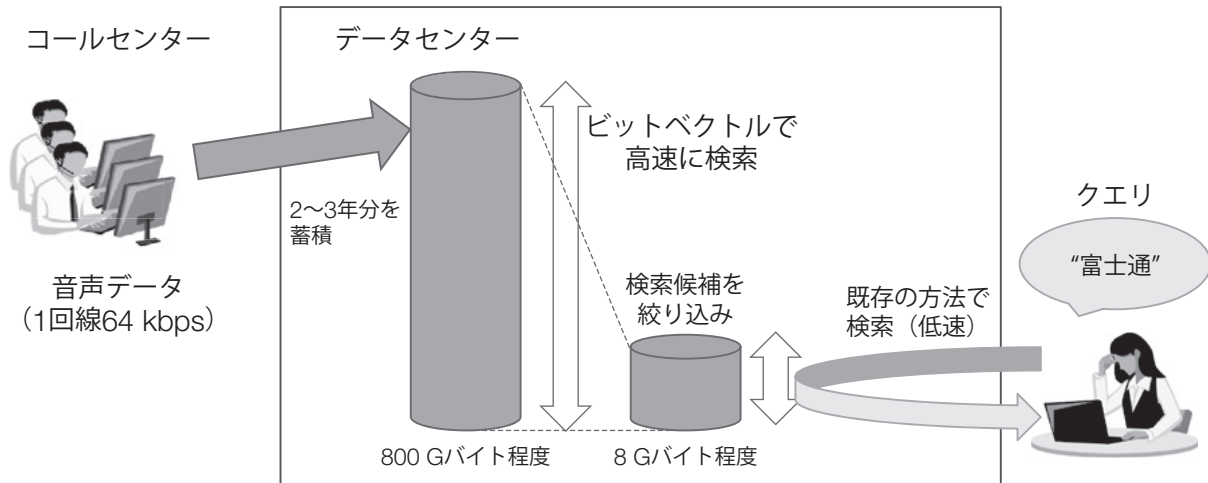


図-5 高速音声検索システムイメージ

これを数百次元以上からなる巨大な特徴量の代替とするため、ビットベクトルを数百ビット以上とした上で、各レコードの特徴量ベクトル間の類似性を極力保ち、検索精度を向上させることを目指し検討を行った。

その一つが、特徴量空間内でのデータの分布を、統計的推定的一种であるマルコフ連鎖モンテカルロ法で推定した結果に基づいて平面を配置し分割する方法であり、⁽¹⁴⁾ もう一つが曲面を用いて分割する方法である。⁽¹⁵⁾ いずれも、従来手法と比較してより高い精度が得られ、演算量の面でも実用的な方法となった。次章では、この技術の適用事例について述べる。

高速近似類似検索の適用事例

特徴的な適用事例を二つ挙げる。

一つは、類似していると推測される箇所を大まかに検索した後、厳密な類似性を判定する事例である。具体例として高速音声検索を挙げる。従来は、長期間蓄積された音声データの中から音声認識に基づいて検索し、指定した単語を発話した時刻を検出していた。発話をテキストに変換してから検索する場合と比較して、精度が高いことが特長である。ただし、処理時間が長く、発話時間の1/10程度の処理時間を要してしまう。そこで、高速な類似検索を用いて、発話している可能性が高い時間帯を検索し、その時間帯のみ音声認識に基づい

て検索することで、高い精度を保ったまま処理時間を短くできる。これにより、数年間分の音声データが数時間程度で処理できる。この事例におけるシステムイメージを図-5に示す。コールセンターからデータセンターに蓄積された全ての音声データの中から、類似発話箇所を高速近似類似検索して候補を検出した後、厳密で処理時間のかかる音声認識を行う。音声検索については前記のアイデアに基づき、音声検索システムを試作し、事前に予測した処理速度が達成できることを確認している。この技術は音声だけでなく、ほかのセンサーデータにも同様に適用できる。同じ構成で、100万人規模の生体認証デモシステムも試作した。⁽¹⁶⁾

もう一つは、ビットベクトルを使った類似検索のみを行う事例である。従来の検索において、特徴量の類似度だけで検索している事例に適用し、高速化できる。例えば、Deep Learning⁽¹⁷⁾などで自動学習した特徴量をそのまま検索に使用する。従来の特徴量による検索と比較し、100～1000倍ほど高速化できる。レントゲンなどの医療画像によるスクリーニング作業の補助や、自然言語における類語の検索、書籍やデザインのリコメンドのための類似構図画像検索、類似形状の部品検索への適用が考えられる。

む す び

ビッグデータの処理技術は、Web検索サービス

プロバイダーの消費者向けサービスで発展してきた。こうしたサービスでは、Web検索や電子メールなど特定のサービスをいかに多数の消費者に滞りなく提供するかに集中し、整合性などは割りきって大容量・低コストを実現している。

これに対し、企業でのビッグデータ活用は、収集したデータからいかに効率良く自社に有用な知見が得られるかが中心となる。富士通研究所では、多数のデータ系列の相関を見出すこと、あるいはセンサーや画像などのメディアデータから得られる新しいタイプのデータを高速に検索することが有用な知見を得られる基本要素になると考え、これらの基盤技術を開発してきた。

今後は、IoT時代の様々な機器から収集された現場データと、企業内の業務データとを組み合わせることで知見を得ることが更に重要になると考えられる。今後も引き続き、より大量かつ多数のデータ源を組み合わせることで高速にデータを処理する技術の研究を進めていきたい。

参考文献

- (1) Apache : Hadoop.
<http://hadoop.apache.org/>
- (2) 上田晴康ほか：NetCOBOLのHadoop連携機能の開発と実践事例. デジタルプレクティス, 情報処理学会, Vol.5, No.2, p.120-129 (2014).
- (3) 黒松信行ほか：データ処理平準化による業務アプリ並列実行の高速化. 情報処理学会研究報告, Vol.2014-HPC-145, No.3, p.1-5 (2014).
- (4) 富士通：オープンプラットフォームCOBOL開発環境 FUJITSU Software NetCOBOL.
<http://software.fujitsu.com/jp/cobol/>
- (5) 富士通：並列分散処理ソフトウェア FUJITSU Software Interstage Big Data Parallel Processing Server.
<http://interstage.fujitsu.com/jp/bigdatapps/>
- (6) 柴田 徹ほか：現場の意志決定プロセスを革新するビッグデータ分析ソリューション. *FUJITSU*, Vol.66, No.1, p.32-40 (2015).
<http://img.jp.fujitsu.com/downloads/jp/jmag/vol66-1/paper05.pdf>
- (7) 倉知陽一ほか：消費行動の現場を捉えるマーケティング～食品業のイノベーション事例～. *FUJITSU*, Vol.66, No.4, p.89-97 (2015).
<http://img.jp.fujitsu.com/downloads/jp/jmag/vol66-4/paper14.pdf>
- (8) 富士通：ビッグデータ活用を支えるソフトウェア.
<http://software.fujitsu.com/jp/middleware/bigdata/>
- (9) Apache : Hadoop Streaming.
<http://hadoop.apache.org/docs/r1.2.1/streaming.html>
- (10) Kwon. Y et al. : SkewTune : Mitigating Skew in MapReduce Applications. *SIGMOD '12 Proceedings of the 2012 ACM SIGMOD International Conference on Management of Data*, New York, USA, ACM, p.25-36 (2012).
- (11) S. Arya et al. : An optimal algorithm for approximate nearest neighbor searching fixed dimensions. *Journal of the ACM*, Vol.45 (6), p.891-923 (1998).
- (12) P. Indyk et al. : Approximate nearest neighbors : towards removing the curse of dimensionality. *STOC '98, Proceedings of the thirtieth annual ACM symposium on Theory of computing*, In New York, USA, ACM, p.604-613 (1998).
- (13) M. S. Charikar : Similarity estimation techniques from rounding algorithms. *STOC '02, Proceedings of the thirty-fourth annual ACM symposium on Theory of computing*, In New York, USA, ACM, p.380-388 (2002).
- (14) Y. Noma et al. : Markov Chain Monte Carlo for Arrangement of Hyperplanes in Locality-Sensitive Hashing. *情報処理学会誌*, Vol.55, No.1, p.44-55 (2014).
- (15) Noma. Y et al. : Eclipse Hashing : Alexandrov Compactification and Hashing with Hyperspheres for Fast Similarity Search. arXiv:1406.3882.
- (16) 富士通研究所：世界初！手のひら静脈と指紋を用いた100万人規模の認証技術を開発.
<http://pr.fujitsu.com/jp/news/2011/06/1.html>
- (17) M. A. Ranzato : Deep Learning Tutorial. *ICML* 2013.

著者紹介



土屋 哲 (つちや さとし)

知識情報処理研究所
ビッグデータ処理プロジェクト 所属
現在、ビッグデータ処理の研究開発に
従事。



此島真喜子 (このしま まきこ)

コンピュータシステム研究所
データシステムプロジェクト 所属
現在、ビッグデータ処理の研究開発に
従事。



上田晴康 (うえだ はるやす)

知識情報処理研究所
ビッグデータ処理プロジェクト 所属
現在、ビッグデータ処理の研究開発に
従事。