

構造型 P2P ネットワークにおける キーワードを含む XPath による XML 文書検索

李 曉晨[†] 天笠 俊之^{†,††} 北川 博之^{†,††}

[†] 筑波大学システム情報工学研究科 〒305-8573 茨城県つくば市天王台 1-1-1

^{††} 筑波大学計算科学研究センター

E-mail: [†]babala0851@kde.cs.tsukuba.ac.jp, ^{††}{amagasa,kitagawa}@cs.tsukuba.ac.jp

あらまし 近年, P2P 環境における XML 文書の分散管理が注目され, 様々な研究や応用が進められている. 我々の研究グループはこれまで分散ハッシュ表を利用した XML 文書の検索手法として, XPath を利用する手法と, 転置リストを利用したキーワード検索手法を提案してきた. 本稿では, これらの手法を結合し, XML 文書の構造情報とテキスト情報の双方を考慮し, キーワードを含む XPath による XML 文書検索手法を提案する.

キーワード XML 文書, P2P 環境, 全文検索, XPath, DHT, 転置リスト

Searching XML Documents using XPath Full-Text in Structured P2P Networks

Xiaochen LI[†], Toshiyuki AMAGASA^{†,††}, and Hiroyuki KITAGAWA^{†,††}

[†] Graduate School of Systems and Information Engineering, University of Tsukuba
1-1-1 Tennoudai, Tsukuba-shi, 305-8573, Japan

^{††} Center for Computational Sciences, University of Tsukuba

E-mail: [†]babala0851@kde.cs.tsukuba.ac.jp, ^{††}{amagasa,kitagawa}@cs.tsukuba.ac.jp

Abstract Recently, Peer-to-Peer systems and XML are attracting increasing attention. There have been several attempts to deal with XML documents in P2P networks. Our research group has been proposed two methods for XML document processing in DHT-based P2P networks. One uses XPath and the other uses inverted list to realize keyword search. Here, we combine the two proposed methods to realize XPath full-text search on XML document which considers both the structural information and the text information of XML.

Key words XML, P2P network, full-text search, XPath, DHT, inverted list

1. はじめに

Peer-to-Peer(P2P) 技術は, オーバレイネットワークの基盤技術として, 近年活発に研究されている. Peer-to-Peer システムとは, ピア (Peer) とピアが対等な関係で成り立ったシステム, すなわち参加するコンピュータが対等の立場で処理を行うシステムのことである. Peer-to-Peer システムの一手法として, 分散ハッシュテーブル (DHT) がある. DHT を用いることにより, 検索対象のデータを保有するピアを, 効率的に見つけることが可能となる.

他方, XML はタグによって文書やデータにおける意味や構造を記述するためのメタ言語である. 1998 年に W3C 勧告となってから現在までの間に急速に普及し, 標準的なフォーマットとして, 様々な分野に広く応用されている. こういった XML データを P2P 環境で利用することも一般的に行われるように

なっているので, P2P 環境で XML データの格納と検索を効率的に行うための技術基盤がますます重要になっている.

P2P ネットワーク上での XML 文書の問合せには以下の三種類の方式が存在し, ユーザの専門知識の有無によって使い分けることができる.

- 構造型問合せ
- 非構造型問合せ
- 半構造型問合せ

構造型問合せでは, ユーザが XML 問合せ言語 (XQuery [7], XPath [10]) を用いて XML 文書構造を問合せ式に記述し, 処理を行う. この方式は, ユーザが XML 問合せ言語の専門知識を有し, 文書構造を事前に把握していることが前提になる. このような問合せ方式に基づき, 我々の研究グループは分散ハッシュ表に基づき XPath による XML 文書の効率的な格納・検索を可能にする手法 [4] を提案した.

しかし、実際の P2P 環境では異なる利用者がさまざまな XML スキーマを用いる可能性があり、全ての XML 文書構造を把握することが現実的ではないため、非構造型問合せのニーズが高まっている。そこで、我々は転置リストを利用した XML 文書のキーワード格納・検索手法を提案してきた [5]。

半構造型問合せでは、ユーザが検索対象の XML 文書の部分的な構造とキーワードの双方を問合せ式に記述し、処理を行う。このような問合せは、ユーザが部分的な構造しか知らない時に使用される。最近新たな XML 問合せ言語として公表された XPath Full-Text [9] では、`ftcontains` なる演算子を用いて XPath による XML 文書のキーワード検索を表現することができる。

そこで本稿では、XPath Full-Text が想定するような、パスとキーワードを組み合わせた XML キーワード検索を P2P ネットワーク上で実現するため、これまでに提案してきた二つの提案した手法を統合し、キーワードを含む XPath による XML 文書検索手法を提案する。具体的には、XML の構造情報とテキスト値の双方を検索するため、問合せ式を二つの部分（パス部分とキーワード部分）に分けて、改善した KW-DHT によって処理を行う。

本稿の構成は以下の通りである。まず第 2 章で関連研究について述べ、第 3 章で本研究に関する基本的事項を説明する。次に第 4 章で提案手法の概要を述べ、最後に第 5 章でまとめと今後の課題について述べる。

2. 関連研究

本章では、関連研究を述べる。P2P 環境における XML 文書の格納・検索には、いくつかの手法が提案されている。Reynolds 等 [2] は DHT を用いて転置リストを実現し、(XML でない) 文書の全文検索を行う手法を提案している。そこでは、AND 検索におけるコンテンツ ID の転送時にブルームフィルタを用いてトラフィックを低減する手法を用いている。本研究は、この手法を XML に拡張したものとして位置付けられる。

Abiteboul 等は、DHT に基く XML 検索手法として、KadoP [13] を提案している。KadoP にて、XML は要素あるいはテキスト中のキーワードに基いて分解され、DHT 上に格納される。検索処理としては XPath がサポートされ、XPath 問合せを構成する各要素名あるいはキーワードのリストが DHT によって検索され、それらは Holistic Twig Join によって問合せ結果へと変換される。さらに、[14] では、DHT を構成するピア間のストレージにおける負荷分散のために、DHT 上に B-tree に類似した構造を構築する手法を提案している。また、問合せ結果の候補の効率的な絞り込みのために、ブルームフィルタに基づくフィルタリング手法も提案している。

しかし、これらの手法は、キーワードを用いた XML 文書の半構造問合せを実現していない。また、転置リストにのみに実現にピア ID を使っているため、万が一ピアがオフラインの場合は、検索処理を行うことができない。

3. 基本的事項

本論に入る前に、提案手法の前提となる基本的事項について伸べる。

3.1 分散ハッシュ表 (DHT)

分散ハッシュ表 (DHT; Distributed Hash Table) は、構造型 P2P ネットワークの一種である。DHT では、システムで共通の一つのハッシュ関数を用意し、このハッシュ関数から得られる値を元に、オブジェクト (ピア、データ) をオーバーレイネットワーク上に配置する。ピアが DHT システムに参加する際、担当すべき DHT ハッシュ値の範囲が与えられる。DHT ハッシュ関数が取り得る値を、ピアで分割して担当することになる。ある検索キーで登録されたオブジェクトを検索する場合は、その検索キーの DHT ハッシュ値を計算し、このハッシュ値を担当するピアにオブジェクトの所在を尋ねる。DHT には、代表的な方式として、Chord [3]、Pastry [6] と CAN [8] などがある。以下では本研究で用いる Chord について紹介する。

Chord では、 2^N 個のオブジェクトからなる円構造のハッシュ空間 (ID サークル) を用いる。 N を ID のビット数とする。オブジェクトはハッシュ関数で ID サークル上にマップされる。ID が i のピアは、担当するデータの ID、フィンガーテーブル、前後に位置するピアの ID 情報 (*predecessor*, *successor*) を保持している。またピア i が担当するデータは、前ピア (*predecessor*) から i までのハッシュ値を持つオブジェクトである。フィンガーテーブルには、 N 個の他ピアへのリンクがあり、リンク先は $i + 2^k$ ($k = 0, 1, 2, \dots, N - 1$) 位置を担当するピアである。このため、オブジェクトの検索に要するホップ数は $O(\log_2 N)$ である。

図 1 に Chord の構造とオブジェクト検索の例を示す。ピア 12 がキー 6 の値を検索したいとする。ピア 12 のフィンガーテーブル内で最もキー 6 に近いのはピア 5 であるため、まずピア 5 にアクセスする。これを繰り返し行うことで、キー 6 の情報を保持するピアへアクセスする。この例では、キー 6 を保持しているのがピア 8 であるため、最終的にピア 12 はピア 8 にアクセスし、キー 6 に対応する値を取得する。

3.2 XPath Full-Text

XPath Full-Text は、XML 問合せ言語 XPath [10] (XML Path Language) に全文検索機能を追加した問合せ言語である。XPath Full-Text は 2008 年の 5 月に W3C 勧告として公表された、XML 文書の構造を考慮した全文検索をするための標準的な方法を提供している。例えば、従来の XPath で提供されていた述語である `contains` に対応する述語として、全文検索のマッチングを示す `ftcontains` を利用することができる。

`//book[@no = "1"/title ftcontains ("usability" ftand "testing")`

は、号数が 1 の書籍に於いて、タイトルが “usability” と “testing” に関連するかどうかを真偽値によって返す。

本研究では、この `ftcontains` に着目し、DHT 上での実現方法を検討する。

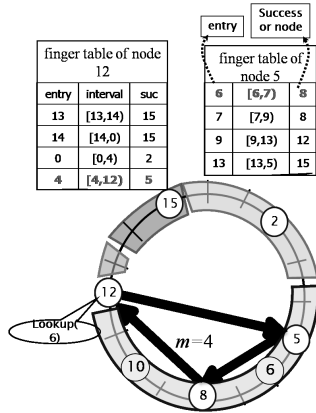


図1 Chord の ID サークル
Fig.1 An example of Chord.

4. 提案手法

ここで、キーワードを含む XPath による XML 文書検索を実現するため、我々が文献 [5] で提案してきた KW-DHT を改善した XML データの格納・検索手法を説明する。

4.1 提案手法の概要

提案手法では、DHT を利用した XML 問合せを実現するために、まず XML データを分解し DHT に格納する。XML データの構造に関する部分は [4] に従い、キーワード検索に関わる XML テキストの部分の格納方法は [5] に従う。

問合せ処理は以下のように行う：1) 問合せのパス部分が述語を含まない場合、例えば、`s/c ftcontains ("2008" ftand "XML" ftand "P2P")` のような問合せの場合は、従来手法 [5] を若干拡張することによって処理可能である。2) 問合せのパス部分が述語を含む場合、例えば、`s/c[f] ftcontains ("2008" ftand "XML" ftand "P2P")` のような問合せの場合、問合せをパス部分 `s/c[f]` とキーワード部分 `s/c ftcontains ("2008" ftand "XML" ftand "P2P")` に分けて処理を行う。基本的な戦略としては、キーワード部分は 1) と同様の処理を行う。パス部分については、[4] によって該当するノードの候補を取得し、最後に両者をマッチングすることによって最終的な解を得る。

以下、まずは DHT を利用した転置インデックスの実現方法から説明する。

4.2 パス式を導入した KW-DHT

KW-DHT とは、[5] で我々が提案した、DHT を利用した XML のための転置インデックスの構築方法である。本研究では、パス式を含むキーワード検索を扱うため、KW-DHT を拡張する必要がある。具体的には、従来はある単語に対応するエントリの値が、その単語を含む要素名とそのノードの LocationID (ピア ID + 文書 ID + XML ノード ID) だったのに対し、パス式によるフィルタリングを可能にするために、単語の所在パスを KW-DHT の転置リストに格納するようにした点である。XML

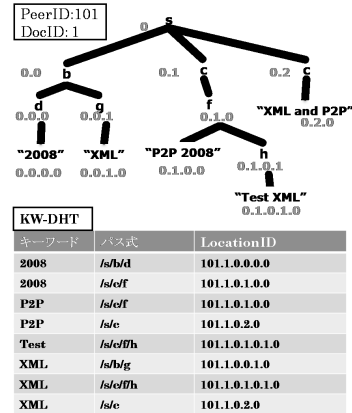


図2 XML データから KW-DHT の生成
Fig.2 An example of KW-DHT.

データの各葉ノード（テキストノード）と属性ノードについて、テキスト値または属性値の中の各単語をハッシュキー、(単語所在パス, LocationID) の転置リストを値とする。単語所在パスは XPath Full-Text クエリのパス部分を処理するために利用する。LocationID は単語が出現する場所の位置情報を保持する ID であり、ピア ID、文書 ID、Dewey Number の連結で構成される。

図 2 に実際の XML データを KW-DHT に格納する例を示す。例えば、ピア ID が 101 であるノードは文書番号が 1 である文書を保持しているとする。

4.3 パス部分が単純パスである問合せ

次に、問合せ処理について伸べる。上で説明したように、問合せのパス部分が述語を含まない単純パスかそうでないかによって処理手順が異なる。まず単純な方から説明する。パス式のマッチングによるフィルタリングを行うかどうかの違い以外は、[5] と同一である。ここでは概要のみ説明する。

4.3.1 問合せ式とその答え

与えられたクエリは $q = [e_1|e_2|\dots|e_n]$ `ftcontain` k_1 `ftand` $\dots$ k_m であるとする。ここで、 e_i は要素名、 k_j はキーワード、 $|$ はロケーションステップの分離記号で、 $/$ または $//$ であるとする。

XML は木構造を有するため、キーワード検索の際には、キーワードにマッチする部分文書を決定する必要がある。与えられたキーワードにマッチする部分文書を決定するために、いくつかの概念を導入する。最小共通祖先 LCA (Lowest Common Ancestors) とは、木においていくつかのノードの共通の祖先の内、最も葉に近いノードである。LCA の中で与えられたキーワードを全て含む最小の LCA に対応する SLCA (Smallest LCA) を紹介する。

- LCA の内、与えられたノードを全てその子孫ノードに含む
- その子孫に同様に全てのノードを含むノードを持たない最小のノードである

すなわち、与えられたパス式で配下にある SLCA を検索すれば良い。

4.3.2 キーワードによる DHT の検索

まず、クエリに含まれているキーワード k_j を KW-DHT によって検索する。原理的には、各キーワードに対応するピアにおいて、そのキーワードが出現する場所の情報が得られるので、それらをまとめた上で SLCA を計算すれば良い。得られたキーワードの転置リストにおいて、クエリのパス式と単語所在パス式がマッチするかどうかを確認し、マッチした転置リストの LocationID を全て抽出する。

例えば、クエリは `s/c ftcontains ("2008" ftand "XML" ftand "P2P")` であるとする。図 2 のように、キーワード 2008 の転置リスト: $(/s/b/d, 101.1.0.0.0.0), (/s/c/f, 101.1.0.1.0.0.0)$; XML の転置リスト: $(/s/b/g, 101.1.0.0.1.0.0), (/s/c/f/h, 101.1.0.1.0.1.0.0), (/s/c, 101.1.0.2.0.0)$; P2P の転置リスト: $(/s/c/f, 101.1.0.1.0.1.0.0), (/s/c, 101.1.0.2.0.0)$ 。その中で、クエリのパス部分 `s/c` とマッチする転置リストの LocationID を抽出し、2008 の LocationID リスト $(101.1.0.1.0.0.0)$, XML の LocationID リスト $(101.1.0.1.0.1.0.0, 101.1.0.2.0.0)$ と P2P の LocationID リスト $(101.1.0.1.0.0.0, 101.1.0.2.0.0)$ が得られる。得られた LocationID を用い、P2P ネットワーク上でキーワードの SLCA を計算する。

4.3.3 LocationID による SLCA の計算方法

P2P ネットワーク上では、キーワードの転置リストが分散配置されている。与えられたキーワード集合を含む部分文書から SLCA を計算するには、LocationID をピア間でやり取りする必要がある。図 4 に例を示す。検索手順について、ピア間での通信量を最小限にするため、問合せが各検索キーワードの LocationID リストのサイズの小さいものから順に処理を行う。まず、最初のピアにおいて、KW-DHT による検索キーワードの転置リスト中の LocationID リストを抽出し、次のピアに送信する。次のピアは、送られた LocationID リストと自分が保持するキーワード集合にマッチして、LCA を探す。得られた LCA から親子もしくは祖先子孫関係にない最小部分木 (SLCA) を計算する。これを全てのキーワードについて行うことで、全キーワードを含む SLCA を計算することができる。

二つ LocationID リストのマッチについて、例えば、検索キーワードが XML と P2P であるとする。XML の LocationID リスト: $(101.1.0.0.0, 101.1.0.1.2.0.0, 101.1.0.1.1.1.0)$ と P2P の LocationID リスト: $(101.1.0.1.1.2.0, 101.1.0.1.2.1.0)$ により、三つの最大共通部分 (LCA) $(101.1.0, 101.1.0.1.2, 101.1.0.1.1)$ が得られる。この中で、`100.1.0` が他の二つ LocationID に含まれているため、最小部分木とならない。`101.1.0.1.2` と `101.1.0.1.1` が互いを含んでいないため、最小部分木となる。SLCA の定義により、この二つ LocationID はキーワード XML と P2P の SLCA であることが分かる。計算された SLCA をまた次のピアに伝送する。最終的に計算された SLCA が検索を行ったピアに返される。

しかし、複数のキーワードによる SLCA の算出を行う際、転置リストの数が多くなると LocationID を送信するためのトラ

フィックが大量に発生してしまうという問題があるため、大規模なデータに対応できない恐れがある。

4.3.4 ブルームフィルタによる通信速度の改善

上で指摘した問題を解決するため、ブルームフィルタを利用する通信速度の改善手法について述べる。まず、ブルームフィルタの概要を紹介し、続いて適用手法に XML データへどのように適用するかについて述べる。

a) ブルームフィルタ

ブルームフィルタとは、複数の要素から構成される集合の中に、任意の要素が含まれているかどうかを判定するアルゴリズムである。巨大なビット列に対し、コンテンツを表すキーワードを複数のハッシュ関数に通し、得られた値と等しい位置のビットをそれぞれ立てることでキーワードを表現する。ブルームフィルタを用いる利点は：

(1) キーワードはビットデータとして格納されているため、ノード間で情報を交換する際は少ない通信負荷で済む。

(2) コンテンツ有無の確認が XOR 演算で済むなど、処理に関する負荷が全体的に低いことが挙げられる。

例えば、二つのハッシュ関数 h_1, h_2 が定義されており、 $m=11$ ビットのビット配列を用いる。P2P ネットワーク上で、ピア A が持っているデータ集合は $\{a, b, c\}$ 、ピア B が持っているデータ集合は $\{b, c, e, f\}$ とすると、A と B のデータ集合の論理積を求める。まず、ピア A は自分が持っているデータをハッシュ関数に通して、得られたハッシュ値のブルームフィルタビット列を作成して、ピア B に自分のデータ集合の代わりに、このビット列を伝送する。ピア B は受け取ったビット列と自分が持っているデータのハッシュ値に共通で含まれているデータ $\{b, c\}$ を抽出する。抽出されたデータは A と B のデータ集合の論理積である。

b) ブルームフィルタによる転置リストのマッチング

ピア間での通信速度の改善についてブルームフィルタを XML データへどのように適用するかについて述べる。基本的には、二つ問題点がある。

(1) SLCA を計算するため、ブルームフィルタのような情報を登録すれば良いか。

(2) ブルームフィルタでは、ハッシュ関数を用いるため、ハッシュ関数による誤検出 (False Positive) がある

c) ブルームフィルタの構成

まず、問題点 1 についての解決手法を説明する。XML データは木構造を有するため、ノード間の先祖・子孫関係を考慮して SLCA を算出しなければならない。このため、各 LocationID だけではなく、それらの全ての先祖ノードをブルームフィルタに格納する。例えば、LocationID が `101.2.0.1.1.2` とすると、`0.1.1.2` による親を抽出して、ブルームフィルタに集合 $(101.2.0, 101.2.0.1, 101.2.0.1.1, 101.2.0.1.1.2)$ を入れる。ブルームフィルタを用いる SLCA を計算する際、ピア A は自分が持っている検索キーワードの LocationID リストのハッシュ値を計算する。続いては、受け取ったブルームフィルタビット列によって、計算されたハッシュ値についての該当部分がすべて 1 かどうかをチェックして、ビット列に共通で含まれている

LocationID を得る。得られた LocationID から互いに含まれていない LocationID を抽出する。SLCA の定義により、この LocationID はピア A が担当しているキーワードと受け取ったブルームフィルタが担当しているキーワードの SLCA である。最後に、ピア A が計算した SLCA をまた次のピアに伝送する。

d) 誤検出の解消

続いて、ブルームフィルタの誤検出の解消について述べる。基本的な考え方は、以下の式が示す通りである。

$$A \cap (B \cup F(A)) = A \cup B$$

ここで、 A と B は集合、 $F(A)$ は誤りを含む集合 A の部分集合である。仮に $F(A)$ に偽陽性の要素が含まれていたとしても、再度 A との論理積を取ってやれば、正しい答えが得られる。

実際のキーワードを検索する場合は、まず、ブルームフィルタのみで SLCA を計算し、(誤りを含む)SLCA を求める。いったん求めた SLCA を今度は全てのピアに再度送信し、各ピアで論理積を計算することで、誤検出した LocationID を排除する。

詳しく例で説明する。ピア A からピア Z までブルームフィルタを転送しながら SLCA の計算を行う。末端まで行き着いたら、次はピア Z から逆に、今度はブルームフィルタではなく候補解を直前のピアに手渡し。受け取ったピアは、自分が持っているデータと受け取った LocationID リストの論理積を計算する。これによって、偽陽性のエントリが取り除かれ、最終的に正しい SLCA を得ることができる。SLCA を算出した後は一般的に LocationID の数が前より少なくなるため、例え候補解を伝送しても、LocationID のサイズがかなり小さいことが期待できる。よって、全体の検索効率を上げることができる。

4.4 述語を含むパス部分の問合せ処理

述語を含むパス部分の場合では、KW-DHT のみで処理できない。そこで、本研究では、[4] の手法を用い、述語を含むパス部分の問合せ処理を行う。

4.4.1 C-DHT と S-DHT による XML データの格納

まず、[4] を説明する。基本的には、XML データからテキスト情報を保持するノードと葉ノードに分解し、その要素名とテキスト値の組合せを、DHT に格納する。これを C-DHT と呼ぶ。具体的には、XML データのノードのうち、(1) テキストノードだけを含む要素ノード、(2) 属性ノード、(3) テキストノードを含まない葉ノードなどを DHT に格納する。この際、要素名をハッシュキー、(ピア ID、文書 ID、パス式、DeweyOrder、テキスト) の組を値とする。

図 3 に XML データを C-DHT に格納した例を示す。例えば、c ノードはテキスト値 “XML and P2P” を持ち、それについてのパス式や順序ラベルを取得することができる。さらに、任意の部分 XML データを検索、再構築するための手がかりとして、P2P ネットワーク上の全 XML データから Strong DataGuide [11] を構築し、別途 DHT に格納する。これを S-DHT と呼ぶ。具体的には、DG の各要素名をハッシュキー、(パス式、子ノード集合) の組を値として格納する。子ノード集合には、子ノードとして現れるノードを列挙するのに加え

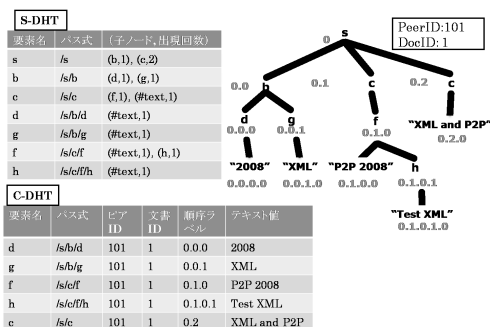


図 3 XML データから C-DHT, S-DHT を生成する

Fig. 3 Storing XML in C-DHT and S-DHT

て、ノードの出現回数も記録しておく。図 3 に実際の DG を C-DHT と S-DHT に格納した例を示す。例えば、c ノードは子要素とするテキストノードと f を持ち、それぞれ一つずつ存在することが分かる。

4.4.2 C-DHT と S-DHT による XPath 検索

C-DHT と S-DHT による検索処理の概要を述べる。まず、 $/, //$ だけを含む単純パス式の場合を説明する。与えられた検索式は $q = |e_1|e_2|\dots|e_n$ であるとする。まず、 q の末端要素名 e_n を手がかりに、テキスト値を直接含む情報を C-DHT に探索する。続いて e_n を根とする部分 XML データを辿るために、S-DHT の情報を手がかりにパス式を拡張し、再度 C-DHT を参照する。具体的には、S-DHT から e_n の子要素である要素名を取得し、それをキーとして再度 C-DHT を参照することで、テキストを得る。これを部分 XML データの再構築が完了するまで繰り返す。

C-DHT と S-DHT を用い、述語を含むパス式の検索処理を説明する。ここでは簡単な例を用いて処理の概要を示す。例として、 $q = /a/b[c]$ なる問合せを考える。まず q から全てのパス式を抽出する (この例の場合は $/a/b$ と $/a/b/c$ になる)。得られたそれぞれのパス式に対して、上記で説明した C-DHT と S-DHT を用いる単純パスの問合せ処理を行い、解の候補集合を得る (例では、それぞれ R_1, R_2 とする)。次に、 R_1 と R_2 が XML データ上で関連をもっているかどうか (先祖子孫関係にあるかどうか) を構造結合 (Structural Join) [12] によって調べる。具体的には、DeweyOrder の性質から、(パス式、ピア ID、文書 ID、DeweyOrder) について、ピア ID、文書 ID が等しく、DeweyOrder がプレフィックスを共有しているかどうかを調べる。構造結合の結果残ったデータが検索結果となる。

4.4.3 C-DHT と S-DHT による述語を含むキーワード検索

C-DHT と S-DHT によって述語を含む XPath を処理することができる。そこで、問合せのパス部分が述語を含む場合、例えば、 $s/c[f]$ ftcontains ("2008" ftand "XML" ftand "P2P") のような問合せの場合、問合せをパス部分 $s/c[f]$ とキーワード部分 s/c ftcontains ("2008" ftand "XML" ftand "P2P") に分けて処理を行なうことにする。キーワード部分は上で述べ

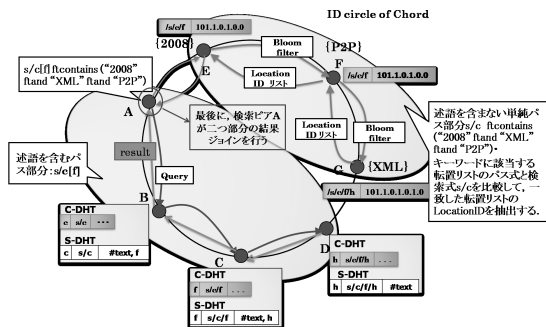


図4 C-DHT と S-DHT による述語を含むキーワード検索処理の例
Fig. 4 Keyword search including predicate by using C-DHT and S-DHT

たのと同様の処理を行えばよい。パス部分については、C-DHT と S-DHT により平行して処理を行ない、両者の結果が得られた後、両者の間で実際に XML データ上で関連があるかどうかを構造結合によって判定する。図4に実際の P2P ネットワーク上での C-DHT と S-DHT による述語を含むキーワード検索の例を示す。図4のように検索ピア A がパス部分とキーワード部分の二つの結果を持ち、構造結合を行う。

5. まとめと今後の課題

本稿では、XPath を利用する XML 文書検索手法と転置リストを利用したキーワード検索手法を統合し、キーワードを含む XPath による XML 文書検索手法を提案した。

今後の課題として、まず提案した手法の実装が挙げられる。また、XPath Full-Text の ftcontains と fland 以外の述語を含む問合せを対応していく予定である。

謝辞 本研究の一部は科学研究費補助金若手研究（＃19700083）、科学技術振興機構 CREST「自律連合型基盤システムの構築」による。

文 献

- [1] Y. Xu and Y. Papakonstantinou. Efficient keyword search for smallest LCAs in XML databases. In Proceedings of the 2005 ACM SIGMOD international conference on Management of data, pages 527-538, 2005.
- [2] Reynolds, P., Vahdat, A. Efficient peer-to-peer keyword searching. Technical Report 2002, Duke University, CS Department, Feb. 2002.
- [3] Stoica, I., Morris, R., Karger, D., Kaashoek, M.F., Balakrishnan, H. Chord: A scalable peer-to-peer lookup service for Internet applications. In Proc. ACM SIGCOMM' 01, San Diego, CA, Aug. 2001.
- [4] T. Amagasa, C. Wu, and H. Kitagawa. Retrieving arbitrary XML fragments from structured peer-to-peer networks. In Joint Conference of the 9th Asia-Pacific Web Conference and the 8th International Conference on Web-Age Information Management (APWeb/WAIM 2007), pages 317-328, 2007.
- [5] Xiaochen Li, Toshiyuki Amagasa, and Hiroyuki Kitagawa, "Searching XML Documents by Keywords in Structured P2P Networks", 3rd International Workshop on XML Data Management Tools and Techniques (XANTEC 2008), Turin, Italy, Sept 1 - 5, 2008. (to appear)
- [6] A. Rowstron and P. Druschel. Pastry: Scalable, decentralized object location and routing for large-scale peer-to-peer systems. In In the 18th IFIP/ACM International Conference on Distributed Systems Platforms(Middleware), pages 329- 350, 2001.
- [7] W3C. XQuery 1.0: An XML Query Language. <http://www.w3.org> January 2007. Recommendation.
- [8] S. Ratnasamy, P. Francis, M. Handley, R. Karp, and S. Shenker. A scalable content-addressable network. In ACM SIGCOMM 2001, pages 161-172, 2001.
- [9] W3C. XQuery and XPath Full-Text Language Version 1.0. <http://www.w3.org/TR/xpath-full-text-10>, May 2008. Recommendation.
- [10] W3C. XML Path Language (XPath) Version 1.0. <http://www.w3.org/TR/xpath.html>, November 1999. Recommendation.
- [11] R. Goldman and J. Widom, DataGuides:Enabling Query Formulation and Optimization in Semistructured Databases. Proc. V LDB, pp.436-445(1997).
- [12] S. Al-Khalifa, H. Jagadish, N. Koudas, J. Patel, D. Srivastava and Y. Wu, Structural Joins: A Primitive for Efficient XML Query Pattern Matching. Proc. ICDE, 2002.
- [13] S. Abiteboul, I. Manolescu, and N. Preda, "Constructing and querying peer-to-peer warehouses of XML resources." in SWDB, 2004.
- [14] Abiteboul, S. Manolescu, I. Polyzotis, N. Preda, N. Chong Sun. XML processing in DHT networks. Data Engineering, 2008. ICDE 2008. IEEE 24th International Conference. pp.606-615, 7-12 April 2008