

情報検索システムへのキーワード自動抽出機能の適用方法について

2N-4

石黒正典 穂吉秀嗣 田中一敏
NTT 情報通信処理研究所

1. はじめに

文献情報、新聞記事情報等を提供する情報検索サービスにおいては、原情報である“原文”からキーワードを抽出し、検索用のインデックスを作成する際に、

1. 原文を読みキーワードを抽出する、
2. 原文そのものをデータベースに登録する、
3. 登録した原文に対応するキーワード群を入力してインデックス作成をシステムに指示する、

という過程が全て人手を介して行なわれており、多大な工数を要していた。我々は、人手に頼っていた作業を可能な限り自動化し、システム全体の効率化を実現した。即ち、キーワード自動抽出用プログラムをシステムに組み込み、キーワード抽出からインデックス作成までを一貫してシステム側で自動化するようにした。本稿では、この自動化の方式について述べる。

2. 情報検索サービスにおけるインデックス

文献情報、新聞記事情報データベースを用いてキーワード検索を行なう情報検索サービスでは、インデックスは“IVL (Inverted List)”と呼ばれる構造を有する(図1)。IVLの「タグ部」の“索引語”には原文より抽出されたキーワードが登録され、「サブ部」にはRID(原文を含むレコードのID)が登録される。

3. データベース構成

文献情報の場合を例に、サービス用データベースのテーブル構成を示す(図2)。我々のシステムではデータベース管理システムとして、マルチメディアデータベース管理システムを利用している^{1),2)}。これは、マルチメディア情報をも管理することを目的に、テーブル内での長大データ(図形/画像情報や文献情報の原文等が格納対象となる)の定義を可能なようにリレーショナルモデルを拡張したものである。

また、キーワード抽出結果を格納する“抽出結果格納DB”と呼ぶ別テーブルを用意した(図3)。当該テーブルは、カラム属性が長大データである、キーワード抽出結果リスト/分析不能語リスト/同義語リストより構成され、それぞれ以下の内容が格納される。

タグ部		サブ部			
索引語	RID数	RID群			
コンピュータ	520	#001	#003	----	#715
データベース	780	#003	#005	----	#935
⋮	⋮	⋮	⋮	⋮	⋮
ワークステーション	215	#002	#004	----	#360

図1. IVLの構造

文献タイトル	著者	発表年	-----	キーワード ^{*1}				原文(長大データ) ^{*2}
				キーワード1	キーワード2	-----	キーワードn	
情報システム	A氏	1989	-----	データベース	コンピュータ	-----	情報処理	原文A
知能と知識	B氏	1988	-----	人工知能	AI	-----	知識処理	原文B
⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮
OAとデータ処理	Z氏	1983	-----	OA	帳票処理	-----	データベース	原文Z

*1: 「キーワード」; 原文に対してキーワード抽出を行なった結果として得られたキーワードが格納されるカラムであり配列構造である。この配列の個々の値がIVLのタグ部の索引語となる。

*2: 「原文」; キーワード抽出対象となる原文であり、長大データとして管理される。

図2. サービス用データベースのテーブル構成例

RID	キーワード抽出結果リスト	分析不能語リスト	同義語リスト
#003	コンピュータ, 情報処理, 人工知能, データベース, 知識処理, OA	ハイレベル, ニュロコンピュータ, ハックアップ, ハッキングモデル, ハイパーテキスト	コンピュータ, インフォメーション・ロッキング, データベース, データ管理, デレトリ
⋮	⋮	⋮	⋮

図3. 抽出結果格納DBのテーブル構成例

A Method for adapting automatic keywords extraction to I-R System

Masanori ISHIGURO, Hidetsugu AKIYOSHI, Kazutoshi TANAKA

NTT Communication and Information Processing Laboratories

