

論文 / 著書情報
Article / Book Information

題目(和文)	自然言語処理システムにおける文脈情報の利用に関する研究
Title(English)	
著者(和文)	住田一男
Author(English)	Kazuo Sumita
出典(和文)	学位:博士(工学), 学位授与機関:東京工業大学, 報告番号:乙第3258号, 授与年月日:1999年1月31日, 学位の種別:論文博士, 審査員:
Citation(English)	Degree:Doctor (Engineering), Conferring organization: Tokyo Institute of Technology, Report number:乙第3258号, Conferred date:1999/1/31, Degree Type:Thesis doctor, Examiner:
学位種別(和文)	博士論文
Type(English)	Doctoral Thesis

自然言語処理システムにおける 文脈情報の利用に関する研究

住 田 一 男

目次

第1章	序論	1
1.1	研究の背景と目的	1
1.2	本研究の位置づけ	3
1.3	本論文の構成	4
第2章	自然言語処理システムの定義	7
2.1	本章の概要	7
2.2	質問応答システム	7
2.3	抄録生成システム	8
2.4	文書検索システム	9
第3章	文脈情報に基づく自然な応答文の生成	11
3.1	本章の概要	11
3.2	関連研究	11
3.3	会話の自然性	12
3.3.1	入力文の理解処理	13
3.3.2	応答生成の戦略	14
3.4	質問応答システムにおける応答生成	15
3.5	実験と評価	18
3.5.1	実験システム	18
3.5.2	実験方法	21
3.5.3	実験結果	21
3.5.4	考察	23
3.6	本章のまとめ	24
第4章	解釈の情報量に基づく文解釈の曖昧性解消	25
4.1	本章の概要	25
4.2	関連研究	26
4.3	聞き手の理解モデル	27
4.4	解釈の情報量	29
4.4.1	解釈の情報量の定義	30
4.4.2	推論規則を持つ知識の場合の情報量の計算	32
4.5	ビデオの操作法に関する質問応答システム	36
4.5.1	知識処理部	37

4.5.2	文解析部	41
4.5.3	談話処理部	41
4.5.4	タスク処理部	42
4.5.5	照応解析精度に関する会話実験	44
4.6	本章のまとめ	47
第5章	文書構造に基づく抄録作成	49
5.1	本章の概要	49
5.2	関連研究	50
5.3	システムの構成	51
5.4	文書の構造化	52
5.4.1	書式解析	53
5.4.2	修辞構造解析	53
5.5	抄録生成	57
5.6	評価実験	59
5.6.1	評価実験1	59
5.6.2	評価実験2	60
5.7	本章のまとめ	62
第6章	文書構造に基づく文書検索	63
6.1	本章の概要	63
6.2	文書提示のための文書の構造化	63
6.3	検索精度向上のための文書の構造化	65
6.4	システムの構成	66
6.4.1	文書構造解析	68
6.4.2	検索結果の提示	69
6.5	全文検索との比較評価	72
6.5.1	方法	72
6.5.2	結果	72
6.5.3	考察	80
6.5.4	結果	80
6.6	意味役割に基づく検索での頻度情報の利用と評価	80
6.6.1	方法	81
6.6.2	結果	82
6.7	動的抄録提示インタフェースの評価	83
6.7.1	方法	83
6.7.2	結果	87
6.7.3	考察	87
6.8	本章のまとめ	89

第7章 結論	91
7.1 本研究のまとめと成果	91
7.2 今後の課題	93
7.2.1 質問応答システム	93
7.2.2 自動抄録システム	93
7.2.3 文書検索システム	94
7.2.4 文脈処理	94
7.3 おわりに	94
付録A 証明	97
付録B 修辞構造解析規則	101
付録C 発表論文リスト	107

目 次

2.1	質問応答システム	7
2.2	抄録生成システム	8
2.3	文書検索システム	9
3.1	量に関する公理を満足しない応答例	14
3.2	自然な応答生成の規則	16
3.3	応答例	17
3.4	規則の適用例	17
3.5	会議室予約実験システムの構成	19
3.6	規則を適用した応答文の意味表現の例	20
3.7	ユーザモデル内のデータの例	20
3.8	被験者に提示した課題の一例	22
3.9	評価実験における会話の一例	22
3.10	応答の自然性についての評価結果	23
4.1	聞き手の理解モデル	27
4.2	聞き手の知識の状態変化	30
4.3	推論規則の集合と命題の集合の関係	33
4.4	原子命題の真偽値の組合せ	35
4.5	ビデオ操作法コンサルテーションシステムの構成	37
4.6	知識表現例 (ビデオのクラススキーマ)	39
4.7	知識表現例 (「再生する」のクラススキーマ)	39
4.8	知識表現例 (操作手順)	40
4.9	条件文を扱うためのワールド機構	40
4.10	「ビデオにカセットをいれた」に対する文解析結果	41
4.11	照応解決処理の処理例	43
4.12	発話の解析用オートマトンの例	44
4.13	実験システムでの会話例	45
4.14	対話実験で得られた対話の例	46
5.1	自動抄録生成システムの構成	51
5.2	文章例	54
5.3	修辞構造	54
5.4	セグメンテーション規則の一例	56
5.5	重要文評価	58

6.1	動的な抄録提示機能を持つ文書検索システムの構成	66
6.2	文書構造解析システムの構成	68
6.3	表示インタフェースの画面構成	70
6.4	動作例	71
6.5	一つの文書の中の文の分布	73
6.6	文章例 1	74
6.7	文章例 2	75
6.8	文章例 3	75
6.9	検索された文書の分布	77
6.10	意味役割のグループ分け	77
6.11	抄録と原文を提示するインタフェース	85
6.12	原文のみ提示するインタフェース	86

表 目 次

4.1	スキーマにおける主たるスロットとファセット	38
4.2	実験のために記述した if-then 規則	45
4.3	照応解析結果	46
5.1	修辞関係の例	55
5.2	修辞関係の相対的重要性	57
5.3	自動抄録結果の評価	59
5.4	重要文捕捉率の比較	60
5.5	文書提示手段としての抄録文と原文との比較	61
6.1	意味役割と表現例	69
6.2	文の意味役割ごとの再現率と適合率	73
6.3	検索命令	76
6.4	意味役割別の適合率	76
6.5	意味役割別の再現率	76
6.6	全文検索における適合率	78
6.7	(a) の場合の適合率と再現率	78
6.8	(b) の場合の適合率と再現率	79
6.9	(c) の場合の適合率と再現率	79
6.10	意味役割に基づく検索と全文検索の精度の比較	83
6.11	各検索手法についての精度比較	84
6.12	インタフェース評価に用いた問題	84
6.13	インタフェース評価実験の条件	85
6.14	インタフェース評価実験の結果	87
6.15	抄録の質についての評価	88
B.1	修辞関係と表層表現例 (1)	101
B.2	修辞関係と表層表現例 (2)	102
B.3	修辞関係と表層表現例 (3)	103
B.4	構造選好規則表 (1)	104
B.5	構造選好規則表 (2)	105

第1章 序論

1.1 研究の背景と目的

近年の計算機やネットワークの普及は、大量の文書の流通を生み出しており、これら文書の多くの部分は、自然言語で表現されることから、自然言語で書かれた文を処理する自然言語処理の重要度が増大している。また同時に、人々が計算機とやり取りする場合や、計算機を介して人々がコミュニケーションを図る機会(電子メールなど)も増加している。このことから、計算機をより使いやすくするため、人間と人間との間のインタフェース手段として最も自然な自然言語を、人間と機械との間のインタフェース手段として利用する必要性も増している。

これまで、自然言語文を処理対象にした応用システムが、研究開発あるいは商品開発されてきている。これらの代表的なものとして、以下のようなシステムをあげることができる。

- かな漢字変換システム
- 機械翻訳システム
- 文書推敲システム
- 情報検索システム
- 自動要約システム
- 質問応答システム

上記のような応用システムを実現するため、様々な自然言語処理の要素技術が研究開発されている。自然言語の処理技術を大きく4つに分類すると以下のようなになる。

- 自然言語を解析しその構造を取り出すための解析処理
- 自然言語を異なる言語や計算機のコマンドに変換する変換処理
- 計算機内部の情報を自然な文として表現する生成処理
- 上記処理を実現するための辞書構築

これらの処理技術のうち解析処理は、通常、形態素解析、構文(意味)解析、文脈解析¹の3つの処理フェーズに分割することができる(後述するようにどのフェーズまでの処理を行うかは応用システムに依存しており、これらのすべての処理を行う必要があるわけではない)。

¹文脈(Context)は、例えば対話などにおいて、対話に参加している話者がどのような境遇や状況におかれているかといった言語外の脈絡までを意味する場合があるが、本研究では、対話や文書中での文のつながりを意味するものとし、この意味で談話(Discourse)とは特に区別しないものとする。

日本語を処理対象とした場合、形態素解析は文字列を単語に分割するとともに分割した単語の品詞を同定する処理を行う。構文解析は、単語間の修飾－非修飾関係などの依存関係を解析する処理である。形態素解析と構文解析が、基本的に一文を対象に処理する一方、文脈解析は複数の文間を見渡した処理を行うものである。

通常、形態素解析を行った結果に基づいて構文解析が行われ、そして構文解析の結果に基づいて文脈解析が行われることになる。

これらの処理には曖昧性が存在する。このため、それぞれの処理フェーズでの結果に曖昧性を持たせた形で、次の処理フェーズに情報を渡し、曖昧性の解消を行うことになる。

これら3つの階層の処理は、必ずしも順序立てて処理される必要はない。同時に実行される処理モデルも提案されているが(例えば[39])、形態素解析、構文解析、文脈解析の3つの階層が意識されていることについては変わらない。

このような処理モデルの場合、あるフェーズで生じる曖昧性は、より次段のフェーズで得られる情報を利用して解消される。

いずれの場合も、より上位のフェーズで得られる情報により曖昧性を解消する。

ここで、いくつかの応用システムの具体例について、形態素解析、構文解析、文脈解析の関係を考える。

例えば、かな漢字変換は日本語ワードプロセッサにおいて、かな文字を漢字かな混じり文に変換するための処理として研究開発された[61]。自然言語処理として実用化が最も進んでいる要素技術である。

かな漢字変換処理の基本的な処理は、単語の読み情報と漢字表記情報を対応させた辞書と単語間の接続性の情報と N 文節最長一致や文節数最少といったヒューリスティックを用いて候補を選択する。したがって、基本的な処理は、形態素解析によって実現されている。

ところが、このような情報だけでは同音意義語から正しい単語を判定することはできない。同音異義語の選択のためのユーザインタフェースとユーザが選択した同音異義語を以後の変換で優先させる処理とを導入することで実用に供することが可能となった。このユーザが選択したある同音異義語を以後の変換で優先する処理は、簡単な処理ではあるが、一種の文脈処理であるといえよう。

機械翻訳システムは、形態素解析により原文中の単語を抽出し、構文解析により単語間の依存構造を分析する。基本的な解析処理は、形態素解析と構文解析からなる。生成処理についても、解析と同様に、構文生成、形態素生成からなる。

基本的な解析処理は、形態素解析と構文解析であるが、精度の高い機械翻訳を行うための解析処理には、指示詞の照応先の検出や省略語の推定、文と文との間の接続関係の認識といった文脈処理が必要であるとされている[53]。

このように、自然言語応用システムの精度向上には、文脈解析や文脈情報の利用が必須である。しかし、実際には、形態素解析や構文解析と比較して、文脈の構造としての表現や処理方法については定式化が進んでいない。

意味論として文脈の表現を定式化する試みが行われている。Fauconnierによるメンタルスペース理論[9]、Barwiseらによる状況意味論[1]、Kampによるディスコース表示理論[31]などでは、事物への照応指示問題を中心に意味論の立場からの理論

化を図っている。しかし、これらの意味論では、言語現象の説明にとどまり、自然言語応用システムへの適用は考慮されていない。本研究では、各種自然言語応用システムにおける文脈情報の利用と、そのための文脈処理について考察するとともに、その考察に基づいたシステムを構築し、評価実験を通じ提案する方式の有効性を検証することを目的とする。

1.2 本研究の位置づけ

文脈処理で処理対象とする自然言語の形態は、対話と文書の2つの様相に分類して考える必要がある。

対話は、通常複数の人もしくは計算機が発する複数の文から構成される。そして、このような自然言語としての対話を取り扱う応用システムとしては、質問応答システム(ユーザが入力する質問に対して適切なタスクを実行し、その結果を応答する)と対話翻訳システム(機械翻訳により異なる言語での対話を可能にするため相互に翻訳を行う)とをあげることができる。

例えば、質問応答システムでは、ユーザが入力する入力文とシステムが生成する応答文から対話が構成され、この対話は一連の文脈をなすことになる。

対話中の文では、それまでユーザが入力した文やシステムが応答する文で用いられた語句などを、代名詞で参照したり、明示的には述べずに省略したりする場合が多い。ところが、質問応答システムでは、通常、ユーザから入力される文の意味を解析し、タスクとして実行するための内部コマンドに変換する必要がある。代名詞や省略の照応先を解決しなければ、その内部コマンドに変換することができない。このため、代名詞や省略の照応先を求める照応解決問題は、質問応答システムでの基本的な処理課題となる。

対話翻訳システムの場合、質問応答システムと比較して、照応の問題の重要性は低い。しかし、各入力文を正しく訳し分けるためには、必要な処理であるとされる[45]。

対話文と同様に、文書が処理対象の場合についても、個々の文中には代名詞や省略が含まれている。この解決が文を正しく解析・理解する上で必要である[49]。

また、通常、ユーザは何らかの目的を達成するために対話を行う。ユーザがどのようなプランを立てて自分の目的を達成しようとしているかを認識することが、適切な対話を実現することにつながる[94]。対話システムにおいては、ユーザのプランや目的の認識も、文脈処理の大きな課題と考えられる。

対話や文書で用いられる語の中には、複数の語義を有する語が存在している。このような多義語において、どの語義で用いられているのかを特定することは、文を正しく解析する上で必要な課題である。かな漢字変換において同音語から、正しい語を選択する問題も同種の問題と考えられる。

対話や文書は、いくつかの文が集まって内容的なまとまりを構成している。隣接している文や近傍の文がどのような関係で用いられているのか、あるいは、どの範囲で内容がまとまっているのか等を解析・認識することは、特に対話や文書全体を一括して処理するような応用システムで重要である。例えば、要約生成では、文書

全体で述べている内容を要約しなければならない。このため、文と文との間の関係や文書全体の構造を取り扱う必要がある。また、文書検索システムでは、文書を検索の単位とする。従来、文書検索システムでは、処理時間の問題から、各文書の単語の出現や頻度など単純な手掛りを利用するのみであった。しかし、文と文との間の関係や文書の構造を解析し利用することで検索精度の向上が期待できる [83]。

Groszらは、意図構造 (Intentional Structure)、注視構造 (Attentional Structure)、語彙構造 (Linguistic Structure) の3つの構造の組合せによる談話構造を提示している [16]。しかし、そこでは、具体的な応用システムは考慮されていないし、構造を抽出する処理も具体化されていない。本研究での文書構造は、上記談話構造の語彙構造に相当するといえよう。

以上をまとめると、文脈処理は以下の4つに分類することができる。

1. 照応解析
2. 語義・同音語の選択
3. プランゴール認識
4. 文書・談話構造解析

本研究では、上記の文脈処理の課題のうち、対話文については、質問応答システムにおける文脈情報を利用した照応解析問題に関して、応答生成と文解析の両面から検討を行う。

文書については、文書・談話構造解析として接続詞などの表層的な言語表現に基づく文書構造の抽出方法を検討する。そして、抽出した構造を利用した応用システムとして、抄録生成と文書検索を考えるものとする。

1.3 本論文の構成

本研究では、応用システムとして、質問応答システム、抄録生成システム、文書検索システムを題材にして文脈情報の利用について検討を行うとともに、その検討に基づきシステムを構築する。

第2章では、本論文で考察の対象とする自然言語処理システムを定義する。

第3章では、利用者の入力する自然言語文を受け付け、自然言語により応答文を返す質問応答システムにおいて、自然な会話を実現するために必要な機能を考察する。文脈情報に基づき自然な応答を生成するための応答生成規則を提案する。そして、この応答生成規則に基づく質問応答システムとして、会議室の予約状況の問い合わせと、会議室の予約を行うシステムを構築する。さらに、構築したシステムを用いて規則の有効性を確認するため評価実験を行う。

第4章では、質問応答システムにおいて入力文に含まれる曖昧性を解決する手法として、文脈情報を利用した方法を考察する。ユーザの入力する文には、その文単独では得られない情報が省略されたり、それまでの文脈を参照するという照応現象が存在する。質問応答システムでは、この照応現象が解決できなければ、ユーザの入力文を正しく処理することができない。照応解決は、照応先の抽出とその抽出した照応先の選択の2つのプロセスに分割することができる。照応先の選択時には曖

昧性が生じるため、照応解決問題は、文解釈における曖昧性解決の代表的な課題と位置づけることができる。

第4章では、入力文の解釈の情報量という概念に基づくモデル化を行い、そのモデルに基づいた入力文の解釈の曖昧性解消手法を提案する。そして、この手法に基づいた質問応答システムとして、ビデオの操作方法の問い合わせを自然言語で行うシステムを構築する。さらに、構築したシステムを用いて、指示代名詞の照応先や省略語の解決の問題に対して、提案する情報量に基づく手法の評価実験を行う。

第5章では、文書の構造に基づく抄録の生成方法について検討する。文章には、文章のまとまりやそのまとまりの間の相対的な関係が存在している。文章の理解や要約の作成のためには、これらまとまりや関係の抽出が必要であると考えられる。このような構造を取り出すための手法を提案するとともに、抽出した構造に基づく抄録生成方法について述べる。そして、提案する構造解析と抄録生成手法に基づく抄録生成システムを構築し、その評価を行う。

第6章では、文書検索での文脈情報の利用について述べる。文書の構造解析結果に基づき文書を検索する手法を提案するとともに、動的な抄録表示を検索結果の提示手段として実現する。これらの検索手法と検索結果の提示手段を有する文書検索システムを構築する。構築した検索システムを用いて、検索精度の評価と結果提示の有効性について評価を行う。

第2章 自然言語処理システムの定義

2.1 本章の概要

本章では、本研究で考察の対象とする自然言語処理システムの範囲を定義する。

第1章で述べたように、自然言語で記述された文章や文を処理対象とする自然言語処理システムには、かな漢字変換システム、機械翻訳システム、文書推敲システム、情報検索システム、自動要約システム、質問応答システムなどがある。これらのうち、本研究では、質問応答システム、抄録生成システム、文書検索システムの3つの応用システムを考察の対象とし、それらのシステムにおける文脈処理について検討するとともに、その考察に基づくシステムを構築する。

以下の節では、本研究で取り扱う質問応答システム、抄録生成システム、文書検索システムの3つのシステムを定義する。

2.2 質問応答システム

ユーザが入力する自然言語文にしたがって、ユーザの目的に沿った処理を行い、その処理の結果として得られた内容(システムの内部状態)を自然言語文で返すシステムである。質問応答システムとユーザとの関係を図2.1に示す。

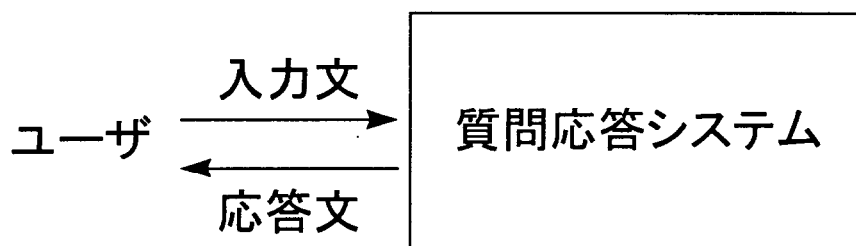


図 2.1: 質問応答システム

ELIZA は、質問応答システムの研究で最も初期の段階に開発されたシステムである [100]。あらかじめ用意されたキーワードのパターンとユーザが入力する入力文を比較し、照合可能なパターンを抽出する。パターンと対応して記憶されている応答

文のパターンに対して、入力文の一部を付加することにより、ユーザに提示する応答文を生成する。ELIZAは、自然言語による質問応答を行うシステムではあるが、単に入力文の一部を書き換えて応答を返しているにすぎず、対話を通じてユーザは何ら有益な結論を導き出すことはできない。つまりELIZAにおける対話は、無目的な対話であるといえる。

一方、例えばGUSは、ユーザの入力にしたがって飛行機の座席予約を行う質問応答システムである[2]。システムは、ユーザが目的としている旅行にふさわしい飛行便を選ぶ手助けを行うことができる。このようにユーザが何らかの目的を持っているような対話を目的指向型対話と呼ぶ。本研究では、目的指向型対話を前提とした質問応答システムを対象とする。

目的指向型対話におけるユーザの目的を達成するには、システムは何らかの問題解決を行う必要がある。対象とする問題解決の種類をタスクと呼ぶ。質問応答システムのタスクとしては、情報検索、操作指示、予約業務、コンサルテーションなどをあげることができる。

目的指向型対話では、ユーザが持つ目的を完遂するためには、通常ユーザとシステムの間での複数回のやりとりが必要となる。このため、ユーザの入力文を適切に処理しシステムの応答文を適切に生成するためには、対話の文脈を考慮した処理が必要となる。

2.3 抄録生成システム

文書を入力として抄録を生成する応用システムである。抄録生成システムとユーザとの関係を図2.2に示す。

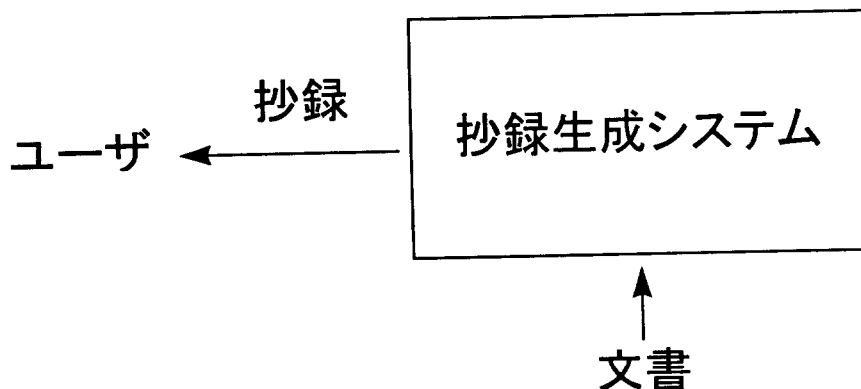


図 2.2: 抄録生成システム

ここで、抄録とはその語の定義[60]にしたがって、「原文から要点を書くぬくこと」であって、要約（「論旨などを、まとめて短く言い表すこと」）ではない。つまり、抄録生成システムにおいては、原文となる文章中から重要な箇所、つまり要旨を抽出し、それらを繋合わせることで原文の大意を表現した文章を作成することを目的とする。

一方、要約を作成するためには、その語の定義にしたがえば、原文の文章の内容理解と、その内容理解の結果に基づいた簡潔な表現への書き換えとが必要となる。通常、文章の内容に関する背景知識がなければ、文章の内容理解を行うことはできないので、要約を行うシステムにはその背景知識をあらかじめ教え込んでおかなければならない。ところが、背景知識を一般的かつ網羅的に用意することは困難であるため、実際には、例題とする文章に関する背景知識だけしか用意することができず、例題だけの解析に留まることになってしまう。

本研究では、より広範囲の文章を対象とすることを目的に、背景知識によらず原文から重要文を抜き出し、抄録文としてまとめ上げる抄録生成システムを構築するものとする。

自動的に抄録を生成する最初の試みは、Luhn による研究である [41]。このアプローチは、「一つの文書において主題と関係の深い語は、その文章中で概して繰り返して用いられる」という経験則に基づき、文章中の単語頻度の分布にしたがってキーワードを判定し、これらキーワードを多く含む文を重要文として抽出するものであり、抽出した複数の重要文を抄録とする。この手法では、同じキーワードを含む類似した文を抽出してしまうことになる。

文章は複数の文から構成されており、適切な抄録を生成するためには、文章の文脈を考慮した処理が必要といえよう。

2.4 文書検索システム

複数の文書をあらかじめ登録しておき、ユーザの検索要求にしたがって、登録された文書から検索要求に適合する文書を探し出し、その結果をユーザに提示するシステムである。図 2.3 にユーザと文書検索システムとの関係を示す。

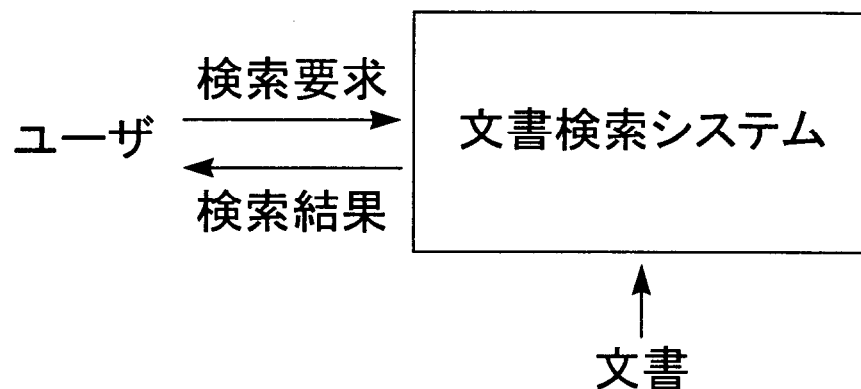


図 2.3: 文書検索システム

検索要求の表現の形態としては、検索語、検索語を要素とする論理式、自然言語文、文書などが考えられる。

検索要求を自然言語文とした場合、質問応答システムと似た動きをするようになるが、質問応答システムでは、検索された結果を自然言語に置き換えて応答文を生

成しユーザに提示するのと比較して、文書検索システムでは、検索された文書をそのままユーザに示すという点で異なる。文書検索システムでは、例えば、検索要求と適合する文書が複数存在していれば、通常、その複数の文書を検索結果としてユーザに提示することになる。

情報検索の研究は、計算機が開発された初期の段階から着手されたが、検索語あるいは検索語を要素とする論理式を入力とするブール検索が実用的なものとなった後、検索精度を向上させる研究は計算コストやロバスト性などの理由から統計的な手法が大半を占めた [83]。

ブール検索では、文書中に指定した検索語が存在するか否かだけで文書と検索要求との適合性を判定する。主題とは関連性のない箇所でも用いられた語に対しても検索語が照合する可能性があるため、検索結果には多くの検索誤りの文書が含まれることになる。

検索誤りの問題を軽減することを目的に、統計的手法では、単語の頻度情報を用い適合する文書の順位づけを行う。例えば、統計的手法の一つであるベクトル空間法 [73] では、文書は語の重みを要素としたベクトルとして表現する。語に対応する各重みには、その語がその文書中で用いられた場合には正の値が、用いられていない場合には0が設定される。検索要求も文書と同様にベクトルで表現し、ベクトル同士の内積あるいは余弦を計算することにより、文書と検索要求との適合度を算出する。ここで、ベクトルの重みは、各語の文書内の頻度と文書データベース中の分散とを用いて算出するが、これは、「一つの文書において主題と関係の深い語は、その文書中で繰り返し用いられる」という経験則と、「多くの文書で用いられている語は一般的な語であって重要度が低い」という経験則を数式として表現したものにすぎない。また、このモデルでは、各語はベクトルの要素として互いに独立なものとして取り扱われている。

文書は複数の自然言語文から構成されており、各語はそれら複数の文の文脈中で意味をなすと考えられる。このため、文書検索においても、文書中の語の頻度情報だけを利用するだけでは十分ではなく、文書の文脈を考慮した処理が必要となると考えられる。

以下の章では、特に断らない限り、自然言語処理システムとは、質問応答システム、抄録生成システム、文書検索システムを指すものとし、これらのシステムにおける文脈情報の利用について検討する。

第3章 文脈情報に基づく自然な応答文の生成

3.1 本章の概要

本章では、自然言語により利用者の入力を受け付け、自然言語により応答を返す質問応答システムにおいて、自然な会話を実現するために必要な機能について考察する。

質問応答システムが利用者と自然な会話を進めていくためには、文脈に応じて利用者の入力文を正確に理解する機能を実現するだけでは十分ではない。文脈に沿って応答文を生成していくことが重要な要素となる。本章では、この応答文生成の戦略として、文脈を考慮しその応答文に含める情報の量を制御することが重要であることを示す。考察の結果に基づき、情報の量を制御するために、具体的な8項目からなる応答生成規則を提案する。

提案する応答生成規則の有効性を検証するために、その規則に基づく質問応答システムを試作する。システムは、会議室の予約を処理タスクとしており、自然言語文による会議室の予約状況の問い合わせ、ならびに予約業務を行う。処理した結果、つまり会議室の空室情報や空時間情報、予約の結果などを、自然言語により利用者に返答するシステムである。

試作した会議室予約システムを用いて、提案した応答生成規則の評価を行う。評価実験では、14人の被験者に会議室予約の課題を与え、その課題にしたがって各被験者が質問応答システムを使用し、会議室を予約する。その際にシステムが提示する応答文の自然さ、了解度、親切さを主観的に評価する。評価実験を通じ、提案する規則を用いた応答は規則を適用せずに生成した応答と比較して、自然性が増すことが統計的に確認され、提案する規則が、質問応答システムにおいて有効である見通しが得られた。

以下の節では、関連研究について述べた後、自然な会話を行うために必要となる機能と、応答生成のための規則について考察する。そして、試作した質問応答システムの概要を述べるとともに、評価実験の内容とその結果に関する考察を行う。

3.2 関連研究

自然言語による質問応答システムを実用に供するためには、人間と機械との間の会話が、遅滞なく円滑に行われるものでなければならない、このためには、入力文の内容を理解し、問題解決を行い、適切な応答を出力する必要がある。その際シス

テムの応答は、利用者の入力に柔軟に対処するため、自然な会話を行うための何らかの戦略に基づく必要がある。

会話の戦略は、利用者とシステムとが自然な会話を行うための方法を明らかにし、使いやすいシステムを構築するために必要である。会話を自然に進めるための戦略の一つとして、どのように会話を進めて行くか制御することが考えられる。OODS[68]はプロダクションルールにより会話の制御を行おうとしているが、遷移の仕方が固定的であるので、会話の柔軟な構造には対処しきれない。また、佐川らは、会話の構造をATN(拡張遷移ネットワーク)により制御し自然な会話を行うコンサルテーションシステムについて示している[71]。大平らは、情報検索をタスクとした対話において、ATMS (Assumption based Truth Maintenance System) を用いて対話を制御する手法を提案している[62]。しかし、これらの研究では、対話制御の枠組みのみの提案にとどまり、その枠組みに基づく会話が自然なものとなるかどうかについての評価がなされていない。

会話が自然であるためには、利用者のある時点での入力に対して適切な応答が生成可能になればよいというだけではない。会話の各時点で利用者に提示されるシステムの応答が、会話全体を通じて自然なものでなければならない。Hayesらは、柔軟な会話を行えるシステムに必要な機能を、7項目にわたって列挙している[17]。これらの機能のうち幾つかは、自然な会話を実現する上で必要なものと考えられる。しかし、その十分性や、実現手段などについては明確にされていない。

データベース検索において、ユーザが指定した条件では検索結果が何も得られない場合に、その条件に近い情報を提示するなどの協調的な応答を生成するための枠組みが提案されている([34, 84]など)。データベース検索の対話は、ユーザの検索文に対してすぐに検索結果を提示するという単純なものであり、システムとユーザ間の何度かのやりとりが必要となる対話ではない。

音声による対話システムでは、入力文を処理する際に音声認識誤りが生じる可能性がある。Watanabeらは、誤りが生じる可能性のある対話に関して、対話全体での誤りや誤解を減らすために、どのような応答生成の戦略を取ればよいのかを計算機間のシミュレーションにより検証している[90]。確認のための情報をまったく提示しない戦略、常に情報を提示する戦略、要求された場合にのみ情報を提示する戦略などに関して、誤りの低減に対する効果を評価しているが、応答の自然性との関係は明らかにされていない。

本章では、質問応答システムにおける会話の自然性についての考察を行い、その一つの方法として自然な応答生成のための規則を提案する。更に、この原則についての評価を行い、その有効性を考察する。

3.3 会話の自然性

利用者が入力する自然言語文を受理し、自然言語文により応答を返す質問応答システムにおいて、自然な会話を実現するために必要となる機能について考察する。

3.3.1 入力文の理解処理

自然な会話を実現する上で、利用者が入力する自然言語文を理解する機能を、文解析、談話解析、会話戦略の面から考察する。

3.3.1.1 文解析

質問応答では、常に文法的に正しい文が入力されるわけではない。その一方で、利用者にとって長い文の入力は負担となるので、その文の文構造はそれほど複雑ではない。また、日本語の場合、文節内の構造は固定的であるが、文節間の語順等は比較的自由であり、省略の現象も多い、という特徴がある。そこで、複雑な文構造に対処する必要性よりもむしろ、単語、文節だけの断片的入力や、未知語への対応が十分考慮された処理が自然な会話を行うために必要と考えられる。

3.3.1.2 談話理解

自然な会話を行うためには、代名詞や、省略等によって前の文脈で現れた事物や、推論される実物に照応する現象が取り扱われる必要がある。質問応答では、利用者はある目的に従って会話を行うため、照応される事物は話題となっている焦点からかけ離れたものを指す場合は少ない。そこで、焦点に基づいた処理が適当であると考えられる [77]。しかしこの場合、話に矛盾が生じた場合 (つまりシステムが保持している焦点と利用者が意図している焦点との食い違いが生じる場合) についての処置を考慮しておかなければならない。また、例えば、ディスプレイ上の図形などについての質問応答や、機器の操作方法に関しての質問応答などのようにシステム外の事柄に対する質問応答を処理する質問応答システムを想定すると、談話理解は、タスクをシミュレートしたり、外部の情報をセンスする機能が必須となってくる。

3.3.1.3 会話戦略

システムは利用者の入力理解できたかどうかを、円滑な方法で利用者に知らせる戦略が必要である。Hayesらは、何らかの妨害によって入力が受け取れなかった場合や、理解できなかった場合に、明示的に応答したり、入力を繰り返したりする例を列挙している [17]。これらは、適当な応答を生成することにより、自然な会話を実現しようとするものである。この戦略は、非常に有効な手段だと考えられるが、どの現象の時に、どの戦略を使い、どの程度詳しく応答すればよいかを整理しておく必要がある。また、音声入力のように入力が必ずしも正確でない場合や、キーボード入力でのミスタイプの場合のみならず、どのようにタスクを限定しても未登録の単語が入力される場合がある。これらに対応して、システムの応答を通じ利用者に分からせる枠組が必要である。また、システムが混乱してしまった時には、自身の能力の限界を明示する必要があるが、その検知や状況に応じた内容の応答生成が必要である。更に、システム側だけが会話の主導権をもつのではなく、利用者も時と場合に応じて主導権をもてるようにしておくことも重要である [5]。

3.3.2 応答生成の戦略

一般的な会話を守らなければならない会話の原則として、Grice が会話の原則を提案している [14]。これは、量に関する公理、質に関する公理、関係に関する公理、様式に関する公理、からなり、人間と機械との間で会話を行う場合に、機械側が守らなければならないガイドラインを示す。従って、この原則にのっとって応答を生成することにより、より自然な応答が期待できる。それぞれについて考察する。

3.3.2.1 量に関する公理について

相手の会話がある前提に基づいている時に、その前提自体が間違いで、その間違いを正すような応答を行うこと、あるいは、その前提を仮定するよりもっと良い場合がある時にそれを教えることは、“協調の原則”として必要である。Kaplan は、これを“会話に必要な情報を含める”という量に関する公理に相当すると指摘している [34]。この原則を実現する単純な方法としては、応答の度にできるだけ多くの情報を、利用者に呈示するという方法が考えられる。しかし、そうすると今度は、“必要以上に情報を詳しくするな”という原則を破ってしまう場合があると考えられる。簡潔でわかりやすい文を生成するために、この原則を考慮しておく必要がある。人間が、文のなかに代名詞を用いたり、省略をしたりする現象 (照応参照現象) は、例えばこの原則に基づいているのだといえよう。例えば図 3.1 の会話例で、下線部は量に関する公理を満足していない。これらの例では、下線部を省略することで、簡潔な表現になる。応答の自然性を向上させるためには、この原則の利用が重要である。

User: 明日第 1 会議室を予約したい。

System: 明日第 1 会議室は9 時から 10 時まであいています。

U: 明日 4 時から 5 時まで会議をしたい。

S: 第 1 と第 2 会議室があいています。

U: 第 1。

S: あいています。 お名前を教えてください。

図 3.1: 量に関する公理を満足しない応答例

3.3.2.2 質に関する公理について

システムはシステム内の知識等を用いて問題解決を行い、応答を生成するものであり、その応答を生成する各時点において、偽りをいわせようとする意識などない。すなわち、応答生成の観点からは、これ以上の詳細化は必要ではない。

3.3.2.3 関係に関する公理について

会話が自然であるためには、これは考慮しておかなければならない原則であろう。McKeown は、データベースの構造を問い合わせるタスクで、焦点という概念を導入

し、関連性を保った応答を生成することを試みている [44]。一つの応答の中での関連性の保持は、パラグラフ程度の長さの応答文を生成するときに必要なが、予約業務のように 1 回の応答が短いタスクでは、一つの応答の中では特に考慮する必要はない。

しかし、利用者の要求や質問に対して、答えることが可能なことがらは即座に回答するという点は、システムを作る上で配慮しておかなければならない点である。

3.3.2.4 様式に関する公理について

あいまいな表現や、多義性を避けるという原則や、秩序立てよという原則に関しては、通常、システムを作っていくうえで、考慮しなければならない原則である。システム設計者が、とりたててあいまいな意味の単語を意図的に用いることはないだろう。しかし、結果的にあいまいな文を生成する可能性は否定できないので注意を要する。

以上をまとめると、質問応答システムにおいて自然な会話を実現するためには、文解析、談話理解、会話戦略にわたって幾つか必要な要素があるが、このうち、自然な応答生成という面から、まず Grice の会話の原則が重要であると考えられる。しかし、Grice の会話の原則は、おおまかなガイドラインを示しているだけであるので、具体的な応答生成のための詳細化を行う必要がある。

3.4 質問応答システムにおける応答生成

従来の研究では、応答を簡潔でわかりやすく自然なものとするという観点からの明確な取扱いがなされていなかった [51]。質問応答システムにおいて、自然な会話を行うためには 3.3 で考察したすべての要素が必要であるが、その中でも、自然な応答をするためには Grice の会話の原則のうちに“量に関する公理”を考慮した処理がまず必要である。

そこで本論では、自然な応答を生成するため、“量に関する公理”に基づいた図 3.2 に示す規則を提案する。これらの規則はシステム中の適当な処理過程に組み込むことができるが、以下では応答性政治に用いることを前提として説明する。なおここでは、予約業務のような比較的単純な処理タスクを想定している。

規則 (1)、(2)、(3) は、いつ何を省略するのかを決定するための規則である。これらの三つの規則は、利用者が既知の情報については繰り返さない、ということを表している。ここで“情報”とは、会議室などの予約を処理タスクとして想定した場合、“～が～から～まであいている”というような事実に関する声明のことである。“項目情報”とは、会議室名や時刻のことであり、一般的な情報検索では、検索に必要な検索条件に相当する。例えば、図 3.3 の Example 1 のように、規則を適用しないシステムは、前後の文脈に関わらず下線部を利用者に明示するが、この三つの規則を適用することにより、この部分を省略できる。すなわち、下線部のうち“第 1 会議室があいている”という情報は規則 (1) により省略され、“9 時から 10 時まで”という項目情報は、規則 (2) により省略される。これによって明示すべき情報

- (1) システムがその応答を生成する以前に生成した応答文で明示した情報は省略する。
- (2) 利用者が直前に入力した文に現れた項目情報は省略する。
- (3) (1)から(2)によって明示すべき情報がなくなった場合、その応答文は生成しない。
- (4) 利用者が直前に入力した文が質問であった場合は、(1)を満たしていてもその情報を明示する。
- (5) 利用者の入力に対して否定的な内容の応答であった場合、(3)が成立しても、その応答文は生成する。
- (6) システムが直前に要求していない項目情報を利用者が入力した場合、その項目情報をエコーバックする。
- (7) 並列的な情報がある場合は、それをまとめた形に変換する。
- (8) 会話の最後にはすべての項目情報を含めた応答文により確認を行う。

図 3.2: 自然な応答生成の規則

がなくなるので規則(3)が適用され、下線部すべてが省略され、応答として出力されない。

規則(4)は、利用者の入力に質問である時には、(1)を適用しないことを示している。

一方、(5)の規則は、(3)の適用範囲を限定している。(5)の規則は、利用者の入力に含まれる誤りを正すような応答を明示することにも含まれる。例えば図3.3のExample 2の下線部は、(1)から(3)を適用することにより省略されてしまう。しかしこの場合、利用者の要求に対して否定的な内容の応答であるので、規則(5)を設定することにより、利用者に呈示される。

(6)の規則は、図3.3のExample 3の下線部のような場合に適用される。これは、規則(2)に一見矛盾するが、システムが聞いていない項目情報が入力された場合、このような応答をすることにより、利用者の要求が正しく処理されていることを知らせ、不安を軽減することができる。

規則(7)は、図3.4のように適用される。並列的な要素をまとめて、応答文を生成する。最後に、規則(8)は、予約業務の場合、会話の最後に予約するものを確認するために必要である。図3.3のExample 4に示される形で適用される。

これらの(1)から(8)までの規則は質問応答システムにおいて、自然な応答を生成するためにいつ何を省略すればよいかを決定する規則といえる。

System: 第1会議室が8時から12時まであいています。何時から何時まで予約しますか。

User: 9時から10時まで。

System: 第1会議室が9時から10時まであいています。お名前を教えてください。

(a) 応答例 1 (Example 1)

S: 8時から12時まで第1会議室はあいていません。

:

U: 9時から10時まで第1会議室を予約したい。

S: あいていません。第2会議室があいていますが、予約しますか。

(b) 応答例 2 (Example 2)

S: 第1会議室があいています。何時から何時まで予約しますか。

U: 第3はどうですか。

S: 第3会議室ですね。あいています。

(c) 応答例 3 (Example 3)

:

U: 9時から10時まで。

S: お名前を教えてください。

U: 住田です。

S: 12月18日9:00から10:00まで第1会議室を住田様の名前で予約しました。どうもありがとうございました。

(d) 応答例 4 (Example 4)

図 3.3: 応答例

S: 第1会議室が9時から16時までと、第2会議室が9時から16時まであいています。

:

S: 第1と第2会議室が9時から16時まであいています。

図 3.4: 規則 (8) の適用例

3.5 実験と評価

3.4では、自然な応答を生成するために、“量に関する公理”についての詳細化を図った。しかしながら、この規則の有効性については必ずしも明確でない。そこで、この規則を応答生成部に実現した実験的な質問応答システムを開発し、定量的な評価実験を行う。

3.5.1 実験システム

会話実験を行うことを目的として必要最少限の機能をそなえた、文解析部、会話制御部、タスク処理部、応答生成部からなる、図3.5に示す質問応答システムを試作した[78]。

利用者はキーボードよりローマ字で入力し、システムは漢字かな混じり文をCRTに表示出力する。処理タスクは会議室予約で、利用者の入力、会議室を予約する目的のもとに必要な項目情報を埋めていく入力と、空室状況についての質問に限定している。“～以外の会議室”のような排除的な表現の取扱いや、予約の取消などは考慮していない。

文解析部は、入力中の会議室名や時間表現などをキーワード抽出した後、形態素解析などを行い格表現に変換する。この時、処理できない文字列は読みとばして処理する。単語として取扱えるものは、“明日”“会議室”のような名詞31、“あく”“予約する”のような動詞10、付属語25である。

会話制御部は、月日、時間帯、会議室名などの会議室予約に必要な項目情報を格表現から抽出し、制御フレームに埋めていく。新たな項目情報が得られるごとに、それまで埋められた項目情報とともに予約が可能か否かをタスク処理部へ問い合わせ。この時、予約が可能でない場合は、代替案の検索命令をタスク処理部へ発行する。更に、すべての項目情報が埋められた時点で予約命令をタスク処理部へ発行する。

タスク処理部の処理で明らかになった情報は、格表現の形式でそのすべてが応答生成部の入力となる。入力となる意味表現の一例を図3.6(a)に示す。規則を適用しない応答生成では、これがそのまま漢字かな混じりの表記に変換され、応答の度にシステムの全情報が利用者に呈示されることになる。規則は手続きとして記述しているが、応答生成の処理にあたってはユーザモデルを用いる。ユーザモデルには、利用者が直前に明示した項目情報(例えば図3.7(a)の形式で会話制御部が抽出する)と、システムがそれまでの会話で明示した情報(例えば図3.7(b)の形式で応答生成部が設定する)とを設定する。規則を適用する処理により、例えば図3.6(a)のような入力に対して処理を行うと、図3.6(b)に変換される。

なお、実験システムはSymbolics 3640上で試作し、約5000ステップの規模で実時間に応答を返すことができる。

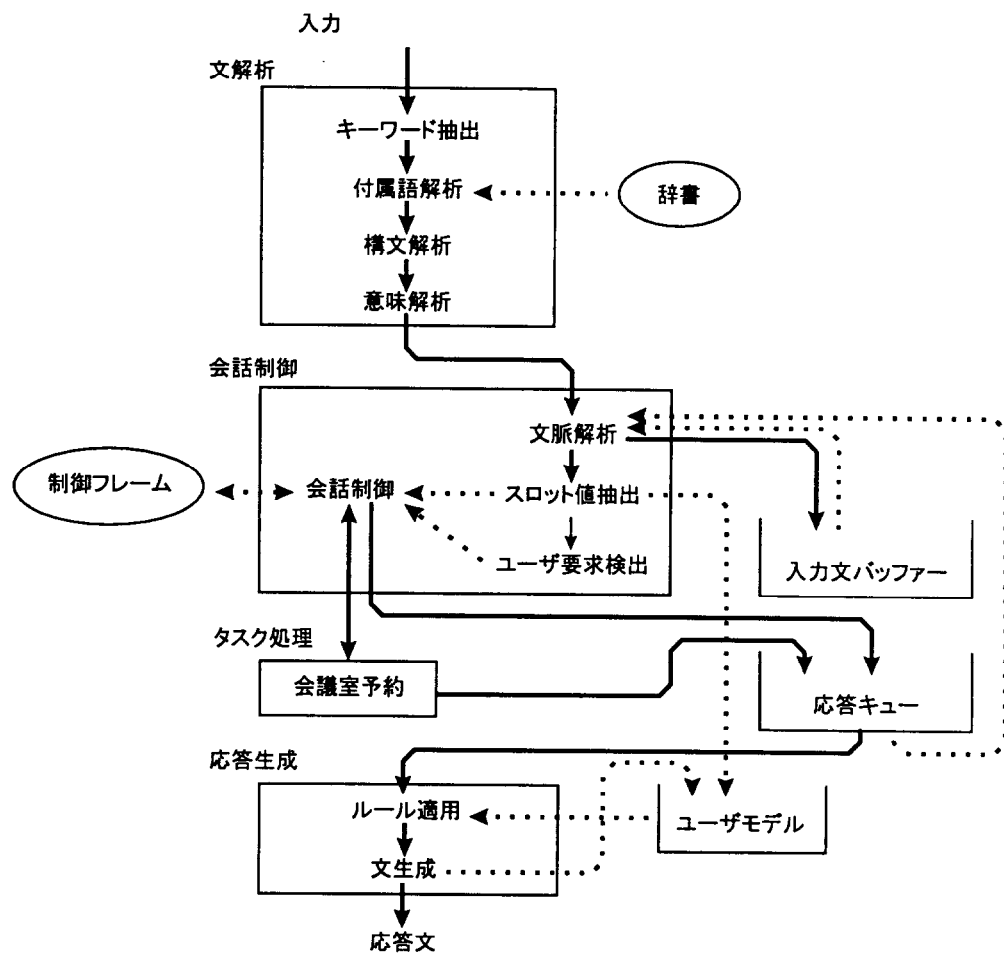


図 3.5: 会議室予約実験システムの構成

```

(((yoyaku (*gimon *teinei)))
 ((aku (*teinei))
  ((place ((1)))
   (time ((from-time 900)
          (to-time 1000))))))
 ((aku (*hitei teinei))
  ((place ((2)))
   (time ((from-time 900)
          (to-time 1000))))))

```

(a) 規則適用前

```

(((yoyaku (*gimon *teinei)))
 ((aku (*teinei))
  ((place ((1))))))
 ((aku (*hitei *teinei)))

```

(b) 規則適用後

図 3.6: 規則を適用した応答文の意味表現の例

```

((place ((1)))
 (time ((from-time 900)
        (to-time 1000))))

```

(a) ユーザの入力文に出現する情報項目

```

((1 (800 900)(1500 1700)))

```

(b) システムがすでに示した情報の記憶例

図 3.7: ユーザモデル内のデータの例

3.5.2 実験方法

この種のシステムは、利用者の入力内容によってシステムの内部状態が変化していくため、あらかじめ決められたデータセットを用いるような、定量的評価は可能でない。評価の方法としては、例えば音声合成の品質評価をする際に用いるプレファレンステストや、オピニオンテストなどが考えられる [56]。そこで今回は、応答の自然性などの項目について評価を行うため、実際に何人かの利用者にシステムを使わせ、その時のシステムが出力する応答に対して、主観的に評価させた。

会話実験の方法は、

1. 利用者に対してある会議室を予約するという課題を与える。
2. その課題に基づいて会話を行い、その一つの会話が終了するごとに、その会話におけるシステムの応答に対して、自然性を主観的に評価する (3 段階)
3. 与えられる課題は、10 個ずつの二つの課題 (時間的な提示順序にしたがって課題 1 と課題 2 と呼ぶ) からなる。課題 1 と課題 2 の内容は予約する日が異なるだけで難易差はない。またそれぞれの課題は、応答生成部規則を適用した場合か、そうでない場合かのどちらかの条件で与えられる。
4. その際気付いたことがあったら、それをコメントとしてメモする。

というものである。また、実験を行う上で次の点に留意した。

1. 規則適用、不適用の条件が、課題 1 と課題 2 のどちらかになるかは、利用者によって変え、時間的な偏りが、実験全体を通してなくなるようにする。
2. 与える課題によって被験者が入力する文の表現が決ってしまうことがないような課題の文とする。

実験では評価項目として、“自然性”、“わかりやすさ”、“親切さ”という三つの項目を設けた。この理由は、“自然性”や“わかりやすさ”などといっても、その言葉によって表される評価の尺度は、各個人によって違うものであり、得られる結果の“自然性”という評価項目～だけでは明確な結論が下せず、他の評価項目も考慮しなければならなくなるかもしれないと考えたからである。

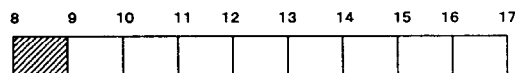
課題として与える予約項目は、図 3.8 に示すように文章の形式にはせずにその時間帯を示す帯グラフによって与えた。これは、課題がすべて文字列で与えられた場合、利用者はその文字列をそのままの字面で入力してしまう可能性があることを危惧したためである。すなわち、各被験者によって、独自の文字列で入力してもらうことにより、課題を設定したものの人為性が混入するのを避けるためである。

3.5.3 実験結果

被験者 14 人に、それぞれ課題 1 と課題 2 を与え、それらの課題に従って会議室を予約する会話を合計 20 回行わせた。評価は、“良い”(good)、“まあまあ”(fair)、“悪い”(bad) の 3 段階のいずれかに分類する。被験者は、システムの中身そのものについての知識はないが、計算機というものについてはよく知っている人たちである。

課題シート1

0. 明日第1会議室を、次の斜線部の時間帯で予約して下さい。



1. 明日の次の時間帯の会議室の空室状況を調べ、どこでもよいですから会議室を1つ、その時間帯で予約して下さい。

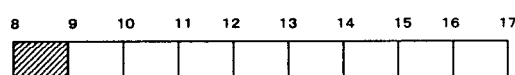


図 3.8: 被験者に提示した課題の一例

実験の会話例の一部を図 3.9 に示す。規則を適用した場合の応答に対する主観評価の結果と、規則を適用しなかった場合の応答に対する主観評価の結果をそれぞれ、図 3.10 に示す(縦軸は合計数を示しており、全合計は 138 となる)。図から明らかなように、自然性(naturalness)、わかりやすさ(intelligibility)、親切さ(kindness)、のすべての評価項目に関して、規則を適用した場合のほうが良いという結果が見られる。

会議室予約システムです。

= asuno8jikara9jimade dailkaigishitsuwoyoyakusitai

あいています。

従業員番号を教えてください。

= XXXXXXXXX

5月22日 8:00 から 9:00 まで第1会議室を住田様の名前で予約しました。

どうもありがとうございました。

図 3.9: 評価実験における会話の一例

これらの評価の差を統計的に比較するため、“良い”という評価に評価点“1”を、“まあまあ”という評価に評価点“0”を、“悪い”という評価に評価点“-1”を対応させ、その評価全体の平均を取り、統計的に平均値の有意差検定を行った。この際、評価がされていない評価は除外した。その結果、自然性に関して、規則を適用した場合の評価(平均値 0.19)とそうでない場合の評価(平均値-0.10)との間で、有意水準 1%で有意差が見られた。一方、その他の評価項目では有意差が見られなかった。

課題1と課題2は、時間的に前後になるため、これが評価に影響を及ぼすことも考えられる。課題1と課題2とについて比較した場合、有意差は見られなかった。主観的な評価尺度が時間の影響を受けずに安定したものと見なせる。

提案した会話の規則を適用することにより向上したものは“自然性”の項目だけであった。一方、“わかりやすさ”や“親切さ”という項目については、有意差が現れていない。“自然性”という評価の定義は、被験者によって個人差があると考えられるが、大局的にみて区別されていると考えられる。

以上まとめると、主観評価の結果、規則を適用した場合の応答は、そうでない場

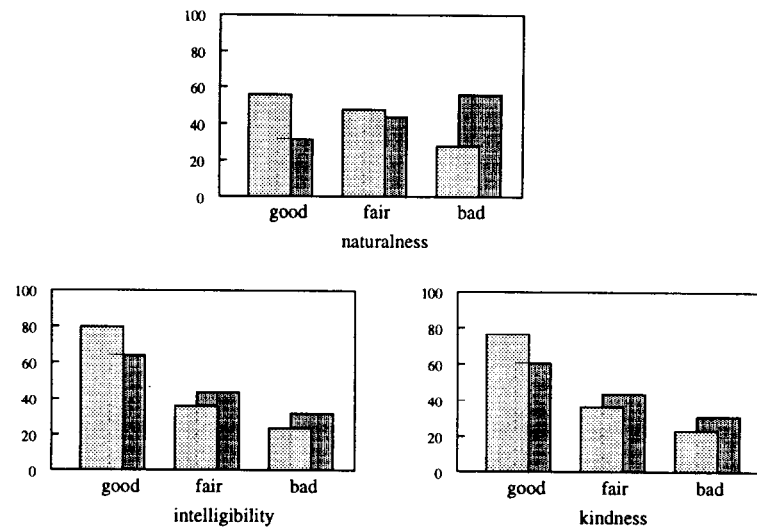


図 3.10: 応答の自然性についての評価結果 (縦軸は評価の合計数)
薄い影は規則適用の結果、濃い影は規則不適用の結果。

合の応答に比べて、“自然性”について向上する。これにより、規則が有効である見通しが得られた。

3.5.4 考察

規則を適用した場合についても、自然性が悪いと判断されたものがあった (全 138 件中 32 件)。悪いと判定された場合、その原因が何に起因しているかは、種々の要因が混在しているため明確な結論は下せない。しかし、評価実験中の利用者とシステムとの会話履歴と、コメント内容により、その主たる原因を推測すると、次のようになる。

文解析が不適當	22 件
未登録語	(13 件)
意味表現の生成が不適當	(4 件)
その他	(5 件)
会話能力が不十分	10 件

これらのすべては本論で考察している規則に関わる処理以外が原因となっている。

このうち、文解析部の処理が原因していると考えられるものには、未登録語が入力文中に存在し、その単語が文解析時に無視されたため、以後の適当な処理が行われず、応答が悪いと評価された場合が最も多い。未登録語の種類としては、タイプミス (文字の重複、文字の欠落、その他)、表記のあいまいさ (長音の欠落、長音の表記、等)、それに類する単語が実際に辞書中にある、等があった。未登録語が入力文中にある場合、その対処方法には、(1) 無視する、(2) 類似語を辞書から探し用いる、(3) 他の語で入力することを利用者に促す、(4) その語の意味を定義させる、等の対

策が考えられる。更に(1)と(2)の場合は利用者に明示するか否かの選択がある。評価実験に用いたシステムでは(1)を用い、利用者には非明示に処理を行った。この点は、未登録語の出現があまりにも頻繁である場合や、その語が重要な語ではない場合には、(3)や(4)では繁雑になってしまうので、語の登録の方法も含めて、システムを設計する必要がある。自然な会話を行うためには、入力文の解析部のロバスト性が必要である。

会話能力が不十分であるという範疇に属する会話では、あらかじめ想定した質問内容以外の質問をしている場合が多い(例えば、“その会議室は何人入れるか。”などの質問)。これについてはシステムが認める質問や要求以外の質問をした場合に、システムの能力を明示する機能を用意しておく必要がある。

3.6 本章のまとめ

質問応答システムにおける自然な会話について考察し、この際に自然な応答を生成することの重要性から Grice の会話の原則に基づく図 3.2 に示す具体的な応答生成の規則を提案した。その規則を応答生成部に実現した会議室予約をタスクとする質問応答システムを用い、14 人の被験者で会話実験を行い、提案した規則を評価した。その結果、提案した規則を用いた場合、規則を用いない場合と比較して、応答の自然性が向上することが統計的に確かめられた。本論文では、予約業務を行う比較的単純な質問応答システムを想定したが、ここで得られた結果により他の質問応答システムにも、規則が有効であることの見通しが得られた。

今後、提案した規則を処理する部分のモジュール化と、応答生成以外の要因で会話の自然性を悪化させている原因(未知語や会話能力の不十分さ)を取り除くための会話の戦略を、定式化することが必要である。また、単に不必要な項目を省略するだけでなく、代名詞や言い換えを用いたほうがより自然であると、会話実験中のコメントとして得ており、この点に関しての検討が必要である。更に、より自然な会話を実現するためには、話題の追跡などの会話の制御や高度の推論が行える知識処理機能も必要となってくるが、この主の機能の検討も課題の一つである。

本章では、自然言語を媒体にした会話の自然性について考察したが、本当に使いやすいシステムを実現するためには、視覚や聴覚など幾つかの媒体を統合した視点に立った、会話の自然性について考慮する必要がある。

第4章 解釈の情報量に基づく文解釈の曖昧性解消

4.1 本章の概要

本章では、自然言語理解システムにおいて、入力文の解釈を行う尺度として入力文解釈の情報量という概念を提案する。

機械による自然言語理解には解決しなければならない曖昧性が多く存在しており、それらは、語彙的な曖昧性、構文的な曖昧性、文脈的な曖昧性に分類される。

語彙的な曖昧性は単語の概念が何を指すかに関する曖昧性である。例えば、「Stay away from the bank」において「bank」という語は語彙的に曖昧である。というのは、この文だけでは川端を意味するのか銀行を意味するのかわからないからである。

構文的な曖昧性は文の構造の曖昧性である。「I saw a girl with a telescope」という文は「with a telescope」が修飾するのが「girl」なのか「saw」なのかが、この文だけでは決めようがない。

文脈的な曖昧性は談話上に出現する。話し手と聞き手は文脈と知識を共有しているので、話し手は代名詞や代動詞、定冠詞名詞句、省略などの照応を用いることができるのだと考えられる。したがって、照応が起因する文脈上の曖昧性は文脈や知識を用いることで減らすことができる。例えば、「Fred phoned John because he needed help」における「he」を求めるのはこのタイプの照応である。「he」は意味的にはFredとJohnのどちらを参照してもよい。したがって、「he needed help」の解釈は二つ存在することになる。これらの曖昧性を解く問題は、ある文脈における最も妥当な解釈を求めるということに相当する。人間の場合、同時にこの選択は言語的な知識と同様に言語外の知識を用いて解決していると考えられる。

本章では、言語外の知識を用いて曖昧性を解消する方法を考察する。そして、談話文脈上で文の解釈の曖昧性を解消するための尺度について提案する。具体的には、聞き手の知識を命題論理の枠組で定義し、入力文の解釈を命題と考えて、知識をもたう聞き手のモデルにしたがって情報量を定義する。与えられた入力文が構文的かつ意味的に複数の解釈候補に解析可能な場合、解釈を選択する際、最も情報量の大きい候補を優先するものとする。提案した尺度にしたがって情報量の大きい解釈を選択するという意味づけは、最小の努力で会話という行為を遂行していくということに対応するものである。

提案した情報量を用いる解釈の選択方式を、照応参照における曖昧性解消の問題に適応し評価を行う。ビデオの操作方法に関する質問応答システムを試作し、そこでの入力文解釈の処理として提案する方式を実現する。さらに、会話実験を行い、方式の有効性を検証する。

4.2 関連研究

自然言語理解において、妥当性のある解釈を求める方法が多く研究されている。Hobbsは参照、複合語、隠喩の問題を推論ルールから結論を得るためのコストを仮定して問題を解決しようとした[22]。Nishidaは仮説に基づく真理値保持システム(ATMS)を用いて自然言語システムを提案している[58]。この研究では曖昧性からくる可能性はATMSにおける仮説に対応する。そして、各仮説に対する確率は経験的に設定するものとしている。これらの2つのアプローチの精度は、ルールや仮説に設定しておくコストや確率の値に依存することは明らかである。しかしながら、これらの研究においては、コストや確率の値に関しての形式的な定義は与えられてはいない。曖昧性を解消するための客観的な指標が必要である。

曖昧性解消の尺度を定義するため話し手と聞き手との間の通信プロセスの一つとして聞き手の理解のモデルを考える。そして、この理解モデルでは、話し手と聞き手は最小の努力で最大の情報を提供しようという仮定をおき、この仮定にしたがって、話し手はメッセージを発話し、聞き手はそのメッセージを理解するものとする。この仮定の下で、文解釈において情報量といった尺度を定義することができれば、この理解プロセスにおける曖昧性を解消するための尺度として用いることができる。Griceの会話に関する量の公理は、話し手ができるだけインフォーマティブであるということを主張している[14]。先の議論は、このGriceの量の公理に対応するとも言える。機械学習ID3[69]の尺度として使用された情報量が本論でも利用しうる。

大須賀は、知識表現における情報量を many sorted logic の拡張として議論している[63]。述語の引数に対するオブジェクトの構造に関する情報量が、この研究において定義されている。この値は、入力文の解釈の情報量として適切な尺度かもしれない。といのは談話文脈における各文を理解するのは知識獲得のプロセスであると見なせるからである。しかし、そこでは、推論ルールを含む情報量を計算する手段が考慮されていない。自然言語理解のプロセスにおいては、イベント間の因果の関係は重要であるので推論ルールに関して考慮する必要がある。その上で、自然言語理解に対する文解釈の曖昧性解消のための情報量の定量的な定義が必要となる。

本章では、イベント間の因果関係を含む知識に基づく聞き手の理解のモデルと、そのモデルに基づく文解釈のための情報量を定義する。理解モデルとともに文解釈における情報量の定義のための理論的な枠組を命題論理を用いて議論する。そして、提案する情報量に関する理論を照応解決に適用する。照応解決は典型的な曖昧性解消の問題である。

自然言語処理において、照応問題の解決に関して、さまざまな視点からの研究が行われている([91, 77, 96]など)。これらの研究は、以下の3つのカテゴリに分類できる。

- 助詞や共感動詞などを用いて焦点を求め照応先を決定する試み([15, 29, 77, 30, 97]など)
- 述語や接続詞などの構文の意味的な制約を用いて照応先を決定する試み([99, 37, 55, 43]など)
- 推論規則を用いて認知的な事物を導き出したり、文脈との一貫性を照合する試

み ([50, 18, 22] など)

本研究での試みは、3番目のアプローチに分類できる。

多くの研究が、照応問題の解決には、構文や意味上の制約だけでなく、推論知識が必要なことを示している。これらの研究の一つに、“デーモン機構”と呼ぶ前向きの推論手法により、照応表現に対する参照先としての事物を生成するものがある。そして、生成される事物はテキスト上で表現されている言語要素以外のものも生成される可能性がある [50]。しかし、推論の役割は、このような事物の生成だけに限定されるものではなく、曖昧性を解消するために用いることが可能であると考えられる。Hobbsは、2つの文の間の“coherence relation”なる関係を用いて、照応の曖昧性を説明しようとしている [22]。この研究において、照応表現は2つの文の間を結びつける機能をもつものとしている。この“coherence relation”は、対比、並列、詳細化の3つの関係に分類されており、それらは推論規則を用いて認識される関係である。

Sidnerは焦点機構を用いたブートストラッピングな手法を提案している [77]。焦点は文脈を一貫性のあるものにするための参照先として機能する。推論規則は一貫性をチェックするためだけに用いられる。

推論規則を用いた研究では、推論規則は、最も妥当な参照先を決定する上で定量的な尺度を与えるためには用いられてはいない。本章で定式化する理論は、照応問題の解決において推論規則の利用について理論的な背景を与えるものである。

4.3 聞き手の理解モデル

図 4.1 は、知識を含む聞き手のモデルを示している。このモデルは二つのプロセスからなっているすなわち、文解析とプレファレンスの判定である。最初のプロセスは、新たな入力文に対して解釈の候補を作成する。そして、次のプロセスにおいて、知識ベースに基づいて最も妥当な解釈を決定する。妥当な解釈を決定したのち、この解釈が知識ベースに追加され、新たな知識として取扱われるというものである。

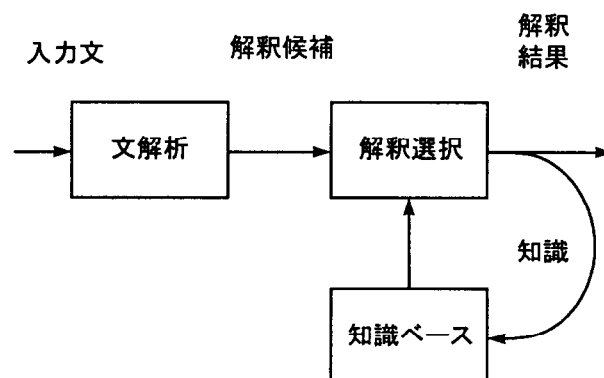


図 4.1: 聞き手の理解モデル

一般に知識ベースは世界知識と呼ぶ知識を含んでいる。世界知識というのは、目的とする分野や話し手の情報についての知識である。ここでは、目的の分野の知識(事実や因果関係)を扱うものとする。

まず、命題論理における聞き手の知識について考える。この枠組みにおいてすべての命題は原子命題とその論理的な結合によって構成されるものとする。具体的には、否定($\neg$)、論理積($\wedge$)、論理和($\vee$)と包含($\rightarrow$)などがこれら論理的結合を構成するとする。原子命題と原子命題の否定はリテラルであり、二つ以上のリテラルを結合した命題を複合命題と呼ぶこととし、ここでは、リテラルは複合命題と区別するものとする。

まず始めに、聞き手の知識が、リテラルの集合と論理的包含(つまり推論規則)の集合とからなると仮定する。リテラルの集合は談話を通じて取得した聞き手の事実についての知識に対応している。また、論理的な包含は因果関係に関する知識に相応するものである。

聞き手の知識を形式化するために、原子命題のユニバースを定義する。ここに含まれる原子命題が、聞き手の知識において用いることのできる命題と考える。

定義 4.1 すべての原子命題の集合 Ω を以下のように定義する。

$$\Omega = \{\omega_n | 1 \leq n \leq N\} \quad (4.1)$$

□

聞き手の知識におけるすべての命題は Ω に含まれる原子命題と論理的な結合を用いて構成される。つまり、聞き手の知識はリテラルの集合と論理的な包含の集合からなると言い換えることができる。そして、 Ω に属するすべての原子命題 ω_n は聞き手の知識の中の命題を構成するために用いられる。

定義 4.2 聞き手の知識 HK は D と K からなる ($HK = D \cup K$)。ここで、 D はリテラルの集合でありその要素 π_i は Ω 中のある要素 ω_n に対して $\pi_i = \omega_n$ あるいは $\pi_i = \bar{\omega}_n$ である。そして、 K は論理的な包含についての集合である。ここで K 内の要素 χ_j で用いるすべての原子命題は Ω に属する。 □

HK に属するすべての命題は一貫性があるものとし、ある解釈(すべての命題に対する真理値の割り当て)は存在するものとする。いいかえれば、 HK 内の命題は聞き手にとって真であり、その論理積も真であることになる。

ここで、一般性を失わずにすべての論理的な包含の左辺がリテラルの論理積で、右辺が一つのリテラルであるということが出来る。なぜなら以下の定理が証明できるからである。

定理 4.1 すべての D と K に対して集合 D と集合 K は論理的に等価な D_1 と K_1 に変換できる。ここで、 D_1 はリテラルの集合、 K_1 は論理的包含の集合であり、 K_1 に属する論理的包含の左辺はリテラルの論理積であり、その右辺は一つのリテラルである。 □

以下の議論では、 K に属するすべての論理的包含はその左辺がリテラルの論理積であり、右辺は一つのリテラルであると仮定する。 D と K の内容は以下のように記述できる。

$$D = \{\pi_i | 1 \leq i \leq M\} \quad (4.2)$$

$$K = \{\chi_j | 1 \leq j \leq J\} \quad (4.3)$$

次に、上で定義した知識に基づいて理解のモデルを構築する。曖昧な要素を含む文はいくつかの解釈にブレイクダウンすることができる。本章では、文の解釈は一つの命題として扱う。ここで、可能な解釈を ν として定義する。 ν が聞き手の知識に追加された時、ある ν から新たな命題が導かれ、新たなリテラルの集合 D' が構成される。

定義 4.3 ν を新たに入力される命題とする。聞き手は ν を獲得することで、そこから導き出されるリテラルの新たな集合を得ることができることになる。この集合を D' と記述する。この時、 D' は以下のように定義される。

$$D' = D \cup \{\pi'_i | (\pi'_i \in \Omega \text{ or } \overline{\pi'_i} \in \Omega) \text{ and } \nu, D, K \vdash \pi'_i\} \quad (4.4)$$

ここで、 $\nu, D, K \vdash \pi'_i$ はリテラル π'_i が ν 、 D 、 K から導き出されることを意味する。

□

事例 4.1 D と K が以下のように定義される場合、

$$\begin{aligned} D &= \{\pi_1, \pi_2, \pi_3\} = \{\omega_1, \omega_4, \overline{\omega_6}\} \\ K &= \{\chi_1, \chi_2\} = \{\omega_3 \wedge \omega_4 \rightarrow \omega_5, \omega_1 \wedge \overline{\omega_2} \rightarrow \omega_3\} \end{aligned}$$

ここで、 $\nu = \overline{\omega_2}$ は入力文に対応する新たな命題である時、新たな命題集合 D' は $\{\omega_1, \overline{\omega_2}, \omega_3, \omega_4, \overline{\omega_6}\}$ となる。

□

4.4 解釈の情報量

この節では、4.3 で述べた聞き手の理解モデルに基づいて入力文の解釈の情報量を定義する。

4.4.1 解釈の情報量の定義

入力文に対応する命題が聞き手の知識に追加された時、命題は聞き手の知識の状態を変化させる。図4.2は、ある声明(「そのボタンを押した」)を聞く前と聞いた後での、知識の状態を表している。この声明を聞く前は、聞き手は、ボタンが押されているかいないかを知らない。したがって、彼の知識の状態では2つの可能性が残されている。一方、その声明を聞いた後では、ボタンが押されていないという可能性がなくなり、もう一方の可能性だけが残される。したがって、入力文は聞き手の知識の可能性を減らすように機能している。

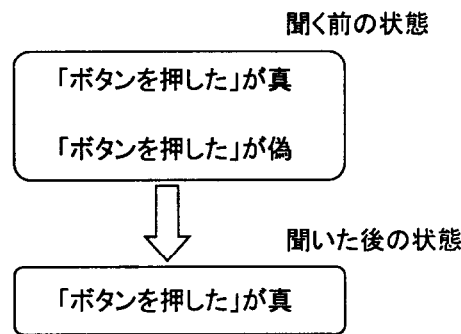


図 4.2: ある文を聞く前後の聞き手の知識の状態変化

このプロセスを形式化するために、真理値を伴うすべての原子命題の組合せを考える。可能な組合せは、 Ω 内のすべての真理値を伴う原子命題の系列によって定義する。ここで、 $\langle \cdot \rangle$ とで囲まれる式で定義する。例えば、 $\Omega = \{\omega_1, \omega_2, \omega_3\}$ なる場合において、もし ω_1 以外のすべての原子命題が真、 ω_1 が偽であるとする、命題の組合せは $\langle \overline{\omega_1}, \omega_2, \omega_3 \rangle$ と記述する。聞き手の知識の状態は以下のように定義される。

定義 4.4 聞き手の知識の状態は D と K に対して一貫性のある真理値を伴う原子命題の組合せの集合として定義できる。聞き手の知識の状態を $S_\Omega(D, K)$ と記述する。ただし、以下の議論では、情報量に関する議論において Ω は常に一定であるので、記述を簡略化するため $S_\Omega(D, K)$ において Ω を省略する。□

事例 4.2 $D = \phi$ と $K = \phi$ である場合を考える。すべての原子命題に関して肯定と否定のすべての組合せを考えるから、状態 $S(\phi, \phi)$ は以下ようになる。

$$S(\phi, \phi) = \{ \langle \omega_1, \omega_2, \omega_3, \dots, \omega_N \rangle, \langle \overline{\omega_1}, \omega_2, \omega_3, \dots, \omega_N \rangle, \dots \}$$

$$\begin{aligned} &\langle \omega_1, \overline{\omega_2}, \omega_3, \dots, \omega_N \rangle, \\ &\quad \dots \\ &\langle \overline{\omega_1}, \overline{\omega_2}, \overline{\omega_3}, \dots, \overline{\omega_N} \rangle \} \end{aligned}$$

□

ここで、 $S(D, K)$ に含まれる要素のリテラルの論理積の真理値について考える。一般に、 $S(D, K)$ における各要素は互いに排他的であり、したがって、もしある要素が真であると、他の要素は真にはなり得ない。同時に、 $S(D, K)$ 中のどの要素が真であるのかを知らないものとする。したがって、 $S(D, K)$ は情報理論における標本空間であると見なせる。もし、標本空間が互いに排他的な有限の要素によって分割され、その確率が p_i であるとする、不確定性の尺度つまりエントロピーが定義できる。

$$-\sum_{i=1}^I p_i \log p_i$$

ここで、 I は空間の要素数である。 $S(D, K)$ の状態における各要素が等確率で生じると仮定する。

仮定 4.1 $S(D, K)$ 中のすべての状態の事前確率が一定で等しいと仮定すると

$$p_i = |S(D, K)|^{-1}, \text{ for all } i$$

ここで、 $|S(D, K)|$ は $S(D, K)$ の要素数である。

□

仮定 4.1 を用いることで、 $S(D, K)$ のエントロピーを得ることができる。聞き手の知識のエントロピーは、以下のように定義できる。

定義 4.5 聞き手の知識のエントロピーを、状態 $S(D, K)$ のエントロピーとして次式で定義する。

$$\begin{aligned} E(D, K) &= - \sum_{n=1}^{|S(D, K)|} |S(D, K)|^{-1} \log_2 |S(D, K)|^{-1} \\ &= \log_2 |S(D, K)| \end{aligned} \quad (4.5)$$

新たな命題 ν が与えられた時、そして D が式 (4.4) で得られる D' で置き換えられる時、新たな命題 ν に対する情報量を以下のように定義する。

$$\begin{aligned} I(\nu, D, K) &= E(D, K) - E(D', K') \\ &= \log_2(|S(D, K)|/|S(D', K)|) \end{aligned} \quad (4.6)$$

この式によって、命題 ν の情報量は聞き手の知識において命題が付け加えられる前後での状態数によって計算可能である。

□

推論規則を含まない知識の場合について、大須賀は知識表現の定量的な考察として情報量の導出について議論している [63]。そこで定義されている情報量との対応を述べる。

事例 4.3 推論規則を含まないので、 $K = \phi$ である。 $D = \phi$ の場合、 $|S(\phi, \phi) = 2^N|$ であるので、 $E(\phi, \phi) = N$ であることは明らかである。例えば、命題集合 $D(= \phi)$ に対して追加される新たな命題を ω_1 とする。このとき、 $D' = \{\omega_1\}$ となる。 $S(\{\omega_1\}, \phi)$ は ω_1 に対して一貫性のある真理値の組合せからなるので、 $|S(\{\omega_1\}, \phi)|$ が $|S(\phi, \phi)|$ の半分であることは明らかである。したがって、 $|S(\{\omega_1\}, \phi)| = 2^{N-1}$ であり、 $I(\{\omega_1\}, \phi) = 1$ となる。 $\nu = \omega_n$ あるいは $\nu = \overline{\omega_n}$ となるような他の命題 ν を考えると、仮に $\nu \notin D$ かつ $K = \phi$ であるとする、 $A = \phi$ かつ $D' = D \cup \{\nu\}$ となる。したがって、 $|S(D', \phi)|$ は $|S(D, \phi)|$ の半分である。この結果、命題 ν に対する情報量 $I(\nu, D, \phi)$ は 1 となる。□

4.4.2 推論規則を持つ知識の場合の情報量の計算

4.4.2.1 一般の場合

推論規則を持つ場合の知識について考察する。一般に個々の推論規則はお互いに関連性を持つ。つまり、お互いに共通の命題をシェアする推論規則が存在するものと仮定できる。まず始めに、推論規則を共通な原子命題をシェアするもの同士を等価クラスとして分類する。次に、それぞれの等価クラス内の命題群に対して、真理値とともに原子命題のすべての組合せにより聞き手の知識のエントロピーを求める。

推論規則の分類を行うため、 K 内の推論規則に対する関係 R と等価関係と呼ぶ R' を導入する。

定義 4.6 K に属するすべての χ_i と χ_j に関して、 χ_i と χ_j の両方がある原子命題 ω_n を互いに含む場合に限って $\chi_i R \chi_j$ と記述する。□

次に、等価関係 R' を上記の関係 R を用いて定義する。

定義 4.7 K に属する要素に関して、以下の2つの条件を満たす場合に等価関係 R' を定義する。

1. For all i and j , if $\chi_i R \chi_j$ then $\chi_i R' \chi_j$
2. For all i, j and k , if $\chi_i R' \chi_j \wedge \chi_j R' \chi_k$ then $\chi_i R' \chi_k$

□

上記の等価関係により、集合 K は、各々共通部分を持たない集合に分割できる。

補題 4.1 等価クラス Γ_h ($\cup_h \Gamma_h = K, \Gamma_h \cap \Gamma_k = \phi$ for $h \neq k$) が存在し、 K に属する各要素は等価関係 R' を用いてこれらの Γ_h ($1 \leq h \leq H$ 、 H はクラスの全数) のどれか一つに振り分けられる。□

Γ_h は互いに直接的あるいは間接的に関係している推論規則の集合であり、異なる i と j に対しては Γ_i と Γ_j とは互いに共通部分を持たない。したがって、原子命題の組合せを考えた場合、他のクラスに属する命題規則との間の関係は考慮する必要がないことを意味している。 Γ_h に対応して、原子命題の集合 Δ_h が Ω の部分集合として定義可能である。すなわち、 Δ_h は、そこに含まれる命題規則内で用いられるすべての原子命題の集合である。図 4.3 は Γ_h と Δ_h との間の関連性を示している。ここで、 K に対して $S(D, K)$ を定義したように、 Δ_h に属する原子命題に関して真理値をとまなうすべての組合せを考える。

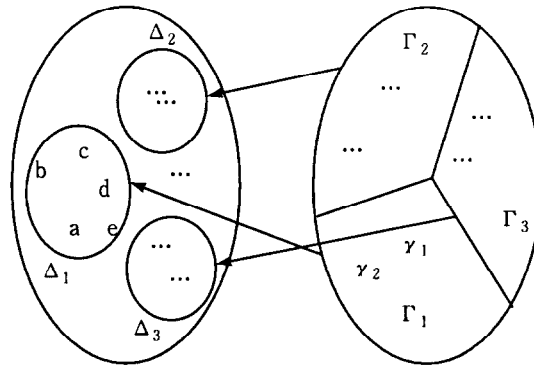


図 4.3: 推論規則の集合 Γ_h と 命題の集合 Δ_h の間の関係 (χ_1 は $a \wedge b \rightarrow c$ 、 χ_2 は $c \wedge \bar{d} \rightarrow e$)

定義 4.8 $C(D, \Gamma_h)$ は、 Δ_h に属するすべての原子命題に関する真理値をとまなうすべての組合せについての集合である。ここで、 Δ_h は Γ_h で用いられる原子命題の集合であり、 $C(D, \Gamma_h)$ は、 D と Γ_h と一貫性がある。□

これによって、以下の定理が導かれる。

定理 4.2 仮に、命題 ν が Γ_h に属する推論規則と関与する命題であるとする¹、 ν の情報量は、 Δ_h に属するすべての真理値をとまなう原子命題の組合せのみで計算可能である。

$$I(\nu, D, K) = \log_2(|C(D, \Gamma_h)|/|C(D', \Gamma_h)|),$$

(ただし ν は χ_j に関与し、 $\chi_j \in \Gamma_h$ かつ $\nu \notin D$) (4.7)

ここで、式 (4.4) で得られる D' は、 ν と D 、 K から導かれるリテラルの集合である。□

¹ ν ないしは $\bar{\nu}$ が K に属するある推論規則 χ_j の構成要素に含まれる原子命題の集合に属する場合、「命題 ν が推論規則 χ_j に関与するという」

命題 ν の情報量は、定理 4.2 にしたがって計算することができる。

$$I(\nu, D, K) = \begin{cases} \log_2(|C(D, \Gamma_h)|/|C(D', \Gamma_h)|) & (\nu \text{ は } \Gamma_h \text{ に属するある } \chi_j \text{ に関与し、かつ } \nu \notin D) \\ 1 & (\nu \text{ は } \Gamma_h \text{ に属するどの } \chi_j \text{ にも関与せず、かつ } \nu \notin D) \\ 0 & (\nu \in D) \end{cases} \quad (4.8)$$

計算過程を示すために、以下の例を考える。

事例 4.4 以下の場合を考える。

$$\begin{aligned} \Omega &= \{a, b, c, d, e, f, g, h, p, q, r\}, \\ D &= \{f\}, \\ K &= \{\chi_1, \chi_2, \chi_3, \chi_4\} \\ &= \{a \wedge b \rightarrow c, c \wedge \bar{d}, f \wedge g \rightarrow h, g \rightarrow q\}, \\ \nu &= \bar{e} \end{aligned}$$

この場合、 χ_1 と χ_2 は c を共通に含み、 χ_3 と χ_4 は g を共通に含むので、 K は Γ_1 と Γ_2 に分割される。この時、 Γ_1 と Γ_2 ならびに、 Δ_1 と Δ_2 は以下ようになる。

$$\begin{aligned} \Gamma_1 &= \{\chi_1, \chi_2\}, \\ \Gamma_2 &= \{\chi_3, \chi_4\}, \\ \Delta_1 &= \{a, b, c, d, e\}, \\ \Delta_2 &= \{f, g, h, q\} \text{ (命題 “} \chi \text{” と “} p \text{” は } \Delta_1 \text{ と } \Delta_2 \text{ のいずれにも属さない)} \end{aligned}$$

新たな命題 “ $\bar{e}$ ” の否定は、 χ_2 の右辺に等しいので、 $\bar{e}$ は Δ_1 に属する χ_2 に関与する。したがって、 Γ_1 に対応する Δ_1 に属する命題に関する真理値をとともなう組合せを考える必要がある (Δ_2 と Γ_2 を考慮する必要はない)。図 4.4 に示す行列は、これらの組合せを示しており、32 の組合せが存在する。

まず始めに、 χ_1 あるいは χ_2 と矛盾しない組合せを取り除く。これらの組合せは、図 4.4 において “ $\times$ ” と記述している。この結果、24 個の組合せが残る。個の場合、 D に属する要素に Δ_1 に含まれるものはないので、この 24 個の組合せは D と矛盾しない。

次に、命題 $\nu (= \bar{e})$ と矛盾しない組合せを取り除く。これらの組合せは、図 4.4 において、“ $\checkmark$ ” と記述している。

図 4.4 において、マークがされていない組合せは、 Δ_1 と D' と矛盾せず、 $|C(D', \Gamma_1)| = 10$ となる。この結果、“ $\bar{e}$ ” の情報量が得られ、 $I(\bar{e}, D, K) = \log_2(24/10) = 1.26$ となる。□

$\sqrt{a, b, c, d, e}$	$\sqrt{a, \bar{b}, c, d, e}$	$\sqrt{\bar{a}, b, c, d, e}$	$\sqrt{\bar{a}, \bar{b}, c, d, e}$
$a, b, c, d, \bar{e}$	$a, \bar{b}, c, d, \bar{e}$	$\bar{a}, b, c, d, \bar{e}$	$\bar{a}, \bar{b}, c, d, \bar{e}$
$\sqrt{a, b, c, \bar{d}, e}$	$\sqrt{a, \bar{b}, c, \bar{d}, e}$	$\sqrt{\bar{a}, b, c, \bar{d}, e}$	$\sqrt{\bar{a}, \bar{b}, c, \bar{d}, e}$
$\times a, b, c, \bar{d}, \bar{e}$	$\times a, \bar{b}, c, \bar{d}, \bar{e}$	$\times \bar{a}, b, c, \bar{d}, \bar{e}$	$\times \bar{a}, \bar{b}, c, \bar{d}, \bar{e}$
$\times a, b, \bar{c}, d, e$	$\sqrt{a, \bar{b}, \bar{c}, d, e}$	$\sqrt{\bar{a}, b, \bar{c}, d, e}$	$\sqrt{\bar{a}, \bar{b}, \bar{c}, d, e}$
$\times a, b, \bar{c}, d, \bar{e}$	$a, \bar{b}, \bar{c}, d, \bar{e}$	$\bar{a}, b, \bar{c}, d, \bar{e}$	$\bar{a}, \bar{b}, \bar{c}, d, \bar{e}$
$\times a, b, \bar{c}, \bar{d}, e$	$\sqrt{a, \bar{b}, \bar{c}, \bar{d}, e}$	$\sqrt{\bar{a}, b, \bar{c}, \bar{d}, e}$	$\sqrt{\bar{a}, \bar{b}, \bar{c}, \bar{d}, e}$
$\times a, b, \bar{c}, \bar{d}, \bar{e}$	$a, \bar{b}, \bar{c}, \bar{d}, \bar{e}$	$\bar{a}, b, \bar{c}, \bar{d}, \bar{e}$	$\bar{a}, \bar{b}, \bar{c}, \bar{d}, \bar{e}$

図 4.4: すべての原子命題に対する真偽値の組合せ $\{a, b, c, d, e\}$ ($\times$ のついた組合せはのどちらかと矛盾するもの。 $\sqrt$ は $\bar{e}$ と矛盾するもの)

4.4.2.2 特別な場合 (個々の推論規則が共通の命題を含まない場合)

一般の場合、入力された命題と直接的ないしは間接的に関与するリテラルについて、すべての真理値に関する組合せを考慮する必要がある。そこで、特別なケースとしてそれぞれの推論規則が共通な原子命題を持たない場合を考える。推論規則の集合 K が、次の条件を満足する時、この場合、原子命題のすべての真理値の組合せを考える必要がなくなる。

仮定 4.2 K に属するすべての χ_i と χ_j に対して (ただし、 $i \neq j$)、 χ_i と χ_j とに関与する共通の命題が存在しない。 $\square$

仮定 4.2 に基づくと、命題 ν の情報量は、 ν が関与する推論規則 χ_j に関してのみに依存することになる。そして、定理 4.2 における $|C(D, \Gamma_h)|$ を簡略化することができる。ここで、 $\Gamma_j = \{\chi_j\}$ とし、 Γ_h の添字 h を χ_j の添字 j に対応させると、以下ようになる。

定理 4.3 仮に、 ν が Γ_j に属する χ_j に関与するとすると、仮定 4.2 に基づく $|C(D, \Gamma_j)|$ が次のようになる。

$$|C(D, \Gamma_j)| = 2^{m_j - k_j} - \sigma_j(D) \quad (4.9)$$

ここで、 m_j は χ_j 内のリテラルの数であり、 k_j は χ_j 内のリテラルの内、その真理値が D において決定可能なものの数、そして、 $\sigma_j(D)$ は χ_j の否定が、 D と矛盾しない場合に 1、そうでない場合に 0 となる関数である。 $\square$

定理 4.4 仮に、 ν が χ_j に関与するとすると、仮定 3.2 に基づく ν の情報量は以下のように与えることができる。

$$I(\nu, D, K) = \begin{cases} \log_2 \{(2^{m_j - k_j} - \sigma_j(D)) / (2^{m_j - k_j - 1} - \sigma_j(D'))\} & (\nu \text{ は } \Gamma_h \text{ に属するある } \chi_j \text{ に関与し、かつ } \nu \notin D) \\ 1 & (\nu \text{ は } \Gamma_h \text{ に属するどの } \chi_j \text{ にも関与せず、かつ } \nu \notin D) \\ 0 & (\nu \in D) \end{cases} \quad (4.10)$$

$\square$

事例 4.5 次の場合を考える。

$$\begin{aligned} D &= \{a\}, \\ K &= \{a \wedge b \wedge c \rightarrow d, e \wedge f \rightarrow g\}, \\ \nu &= b \end{aligned} \quad (4.11)$$

$\chi_j \equiv a \wedge b \wedge c \rightarrow d$ であるから、 χ_j の否定は、 $\chi_j = a \wedge b \wedge c \wedge \bar{d}$ である。 $m_j = 4$ 、 $k_j = 1$ である。また、 χ_j は D と D' の両方に対して、矛盾がないので、 $\sigma_j(D)$ と $\sigma_j(D')$ の両方とも 1 となる。この結果、 $I(\nu, D, K) = \log_2(7/3) = 1.22$ $\square$

本節で、推論規則を含み知識に関して、ある命題を付け加える際の情報量が定義された。その値は、原子命題の真理値の可能な組合せの量に依存する値である。特別な場合として、仮定 4.2 で定義される命題に関係する命題については、簡単な計算によりその情報量が計算できる。

ここでの情報量に関する議論を、自然言語処理システムに入力される文に対応する解釈を命題と対応させることにより、その自然言語文の解釈の曖昧性解消のために用いる。式(4.8)は、ある入力文の解釈に対して情報量を割り当てる手段であると見なせる。

式(4.8)は、ある解釈が談話中で導入された一つ以上の事実や出来事に関与しており、しかも、その事実や出来事から一つ以上の推論が導き出された時、 $|C(D', \Gamma_h)|$ は減少するということを意味している。この結果、情報量は増大する。本論では、解釈の選択を、曖昧性の解決のため、推論規則を含む知識を使って解釈の情報量を最大化することととらえる。

式(4.8)によって与えられた情報量を自然言語処理システムの文解釈の曖昧性解消の尺度として用いた場合、すべての命題を計算対象にしなければならないため計算コストが増大する。

共通の命題を持たない推論規則を仮定した式(4.10)は一般の場合における与えられた命題についての情報量のよい近似値を与えることができるので、ここでは、式(4.10)を用いることとする。

4.5 ビデオの操作法に関する質問応答システム

前節までで議論した解釈の情報量を、ビデオの操作法に関する質問応答システムでの入力文の解釈処理に適用する。

質問応答システムは自然言語処理技術の重要な応用であり、SHRDLU[93]をはじめとして、多くのシステムが研究されている。質問システムでは、自然言語入力文の意味をシステムが持つ知識に照らし合わせてシステムのコマンドとして解釈し、そして、ユーザが要求しているタスクを実行する必要がある。ここでの質問応答システムは、フレームとルールのマルチパラダイムな知識表現を用いて、入力文の解釈と問題解決を同時に行うことを特長としている。

図 4.5 にシステム構成を示す。知識処理部、文解析部、談話処理部、タスク処理部からなる。ビデオの操作法に関する問い合わせを処理する。

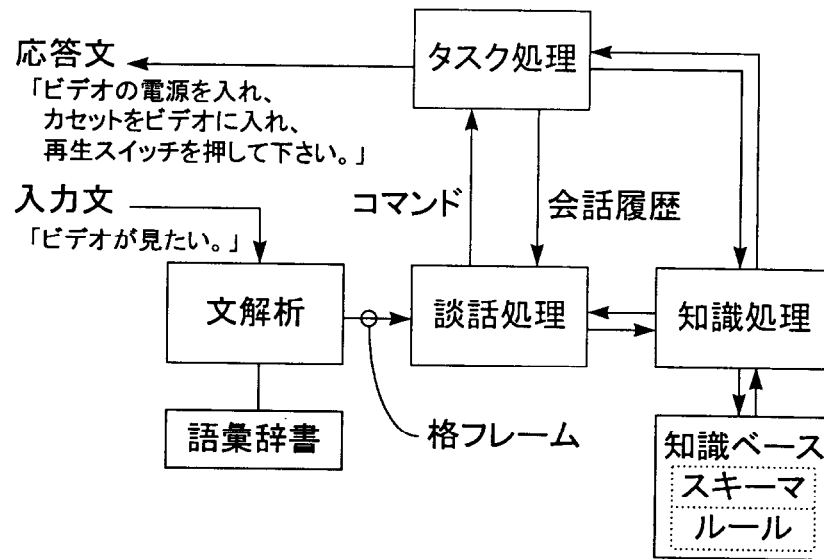


図 4.5: ビデオ操作法コンサルテーションシステムの構成

4.5.1 知識処理部

知識処理部は、対象世界の知識を表現し、他の処理部に推論機能を提供する。図 4.6 に、知識表現の一部を示す。本知識表現では、事物の動作や機能にかかる知識（深い知識）と、経験的な知識（浅い知識）とを区別して表現している。

深い知識として、ビデオ装置の部品などの構成は、フレーム形式で表現し、装置の動作はフレームに付随したルール（s-rule と呼ぶ）で表現している。これらをまとめてスキーマと呼ぶ。図 4.6 は、例としてビデオのスキーマの概要を示しており、その内容は、ビデオの構成と機能を記述している。

スキーマは、スキーマタイプ、スキーマ名、フレーム部、ルール部からなる。スキーマタイプは、そのスキーマがクラスであるか、インスタンスであるかを表す。スキーマタイプがクラスであるスキーマは、それらスキーマ間で、上位-下位関係や部分-全体関係といった関係を設定する目的に用いられる、一方、インスタンススキーマは、クラススキーマをテンプレートとして生成することができ、知識表現中でスキーマを変数的に取扱う際に用いる。

スキーマ名は、スキーマを同定するために用いる。このため、知識処理部は、インスタンススキーマの生成の際には、ユニークに名前を付与するものとする。

フレーム部は、複数のスロットから構成され、スロットは、スロット名とファセット本体からなる。さらに、ファセット本体は、複数のファセットから構成され、ファセットはファセット名とファセット値からなる。

ファセットには、スロットに課す制約を記述することができる。例えば、図 4.6 に示したビデオのクラススキーマには、superC がスロット名、value がファセット名、

表 4.1: スキーマにおける主たるスロットとファセット

スロット名	機能	例
superC/subC	概念間の階層関係	VCR superC 電器製品
has-part/part-of	部分全体関係	VCR has-part 電源スイッチ
attribute-of	属性関係	power attribute-of VCR
agent	イベントにおける動作主格	—
object	イベントにおける目的格	—
obligatory-case	イベントにおける必須格	—
power, channel, ...	自由に定義可能なスロット	—

ファセット名	機能
value	スロットに設定される値を設定する
value-class	スロットに格納可能な値のクラスを設定する
value-condition	スロットに設定される値の条件を設定する
car-num	スロットに設定可能な値の個数を設定する
enumeration	スロットに設定可能な値

denki-seihin がファセット値として設定されている。これは、「ビデオは上位のクラスが電器製品である」ということを表現している。

また、「dengen」というスロット名のスロットに、enumeration がファセット名、[on, off] がファセット値として設定されている。これは、「dengen というスロットはファセット値としてオンかオフのいずれの値しか設定することができない」という制約を表現している。

s-rule は条件部と結論部からなっている。図 4.6 に示したビデオのクラススキーマに記述されているルールは、「電源がオフの場合、電源スイッチを押すと、電源がオンになる」というビデオの動作に関する知識が記述されている。

スキーマを用いて、動作についての知識も表現可能である。図 4.7 は、「録画する」についてのクラススキーマを示している。動詞「録画する」に対して、「誰が」、「何を」、「何で」に対応する格が、各スロットに格納されている。そして、obligatory-case は、その動作を構成する上で必須の格要素を表現している。

浅い知識に対応するルールには、コンサルテーションのためのヒューリスティックな知識や、ビデオの操作手順などを記述する。図 4.8 にルールの一例を図示する。「ビデオの電源を入れ、ビデオにカセットを入れ、再生スイッチを押すと、ビデオが動作する。」という操作手順に関する知識を示している。

人の言語行為を考えると、実世界のものごとだけについて言及するだけではなく、想像上のものごとについて言及することがある。操作法のコンサルテーションのタスクでは、例えば、仮定的な行為の結果生じることがらを尋ねる場合がある。このような質問に対処するために、知識処理部にワールドと呼ぶ機構を設けた。

ワールドは、知識処理におけるワーキングメモリに相当し、新たなワールドが生

```

schema( cls,
        vcr,
        [(superC, [(value, [電器製品])])],
        (has-part, [(value, [電源スイッチ, 再生スイッチ, ...])]),
        (channeru, [(car-num, 1),
                    (value-class, integer),
                    (value-condition, (V, (V>=1, V<12)))]),
        (dengen, [(enumeratio, [on, off])]),
        [s-rule(dengen, dengon-on-1, V,
                (kr-schema(V, [(電源, [off])], true),
                event(osu, [(object, [SW])])),
                kr-schema(V, [(電源, [on])], true),
                [('$var-constraint', (
                    V # vcr;
                    SW # 電源スイッチ :- part-of(SW,V))]),
                s-rule(dengen, dengon-lamp-1, V,
                    kr-schema(V, [(電源, [on])], true),
                    kr-schema(LP, [(点灯状態, [on])], true),
                    [('$var-constraint', (
                        V # vcr;
                        LP # 電源ランプ :- part-of(LP,V)))])))]])

```

図 4.6: 知識表現例 (ビデオのクラススキーマ)

```

schema( cls, 録画する,
        [(superC, [(value, [動作])]),
        (agent, [(value-class, 人)]),
        (object, [(value-class, テレビ番組)]),
        (tool, [(value-class, vcr)]),
        (obligatory-case,
            [(value, [object, goal])])])

```

図 4.7: 知識表現例 (「再生する」のクラススキーマ)

```

rule( operation, play-back,
  seq(event(入れる 2, 1, [(object, [PW])]),
    event(入れる 1, 1, [(object, [K]), (goal, [VCR])]),
    event(押す, 1, [(object, [SW])])),
  event(再生する, 1, [(object, [K]), (tool, [VCR])]),
  [('$var-constraint', (
    VCR # vcr;
    K # カセット;
    PW # 電源 :- attribute-of(PW, VCR);
    SW # 再生スイッチ :- part-of(SW, VCR))])

```

図 4.8: 知識表現例 (操作手順)

成された場合、その元となったワールドにおける知識の全状態が生成されたワールドにおける知識の状態に継承される。

図 4.9 に、以下の文を順に処理した場合のワールドの遷移状態を示す。

- (1) 録画スイッチを押すと、どうなるか。
- (2) その後、巻き戻しスイッチを押すと、どうなるか。
- (3) その代わりに、早送りスイッチを押すと、どうなるか。

ワールド W0 は、文(1)が入力される直前の状態であり、ワールド W1 が文(2)を処理するために生成される。文(2)の処理において、「巻き戻しスイッチを押す」という節を処理するためには、「録画スイッチを押す」という行為が行われた状態で、推論を行う必要がある。そこで、ワールド W1 から新たなワールド W2 を生成し、その上で推論を行う。また、文(3)を処理する際にも、ワールド W1 から新たなワールド W3 を生成し、その上で推論を行う。

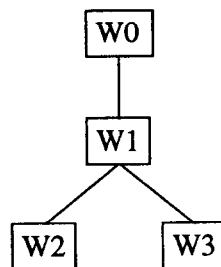


図 4.9: 条件文を扱うためのワールド機構

4.5.2 文解析部

文解析部は、ユーザが入力する仮名漢字混じりの文について形態素解析および構文解析を行い、入力述語を核とする格フレーム表現を出力する [74]。まず、入力文字列を文節に分割した後に、全文節間の係り受けの可能性を調べ、矛盾のないものを出力する。この時、係り受け関係の非交差と順方向性を制約条件として用いている。さらに入力される発話に複数の述語が含まれる時には、それらの間を「条件、目的、時間前後関係、並列、因果」などの関係に分類し出力する。

文解析の結果の例を図 4.10 に図示する。

```
{concept/入れる,
  syntax/{category/動詞型述語, source-word/入れた},
  surface-case/{
    に/{concept/ビデオ,
      syntax/{category/準用詞型述語, source-word/ビデオに}},
    を/{concept/カセット,
      syntax/{category/補語, source-word/カセットを}}},
  fixed-portion/[能動動詞, 肯定, 基本態, 過去, ...]}
```

図 4.10: 「ビデオにカセットをいれた」に対する文解析結果

4.5.3 談話処理部

処理手続きは、照応表現検出部、参照先抽出部、解釈選択部の 3 つのステップからなる。解釈選択部では、最も情報量のある解釈を選択することにより、入力文の照応表現に対する参照先の候補のうち最も妥当な候補を決定する。

文解析部から、入力される格フレームに対して、次の処理を行う。

4.5.3.1 照応表現検出

照応解決の第一のステップとして、照応表現の検出を行う。代名詞、連体詞(その、この)で修飾された名詞、格の省略などを検出する。日本語の場合、名詞単独でも文脈上の語句を参照する場合があるので、名詞単独についても照応表現とする。これらの抽出した照応表現に対応して、インスタンススキーマを生成する。

4.5.3.2 参照先抽出

各照応表現に対して、会話履歴から参照先の候補を検索する。照応表現に対応するインスタンススキーマと同じクラスか、その下位クラスか、部分-全体関係となる表現を候補として抽出する。一方、照応表現がある出来事を参照している場合もある。このため、会話履歴に格納されている格フレームについても検索対象とする。

参照候補を抽出した後、格フレームを構成する。その表現は入力文に対する解釈の候補を内部表現として保持することができる。

4.5.3.3 解釈選択

一つ以上の解釈が得られた場合、このステップの処理を行う。最も情報量の多い候補を選択する。つまり、会話履歴中の先行する文を探索し、知識ベース中の if-then 規則によって入力された解釈と結びつけられるものを探す。ここで、if-then 規則は前節の議論における推論規則 K に対応しており、会話履歴における先行文は D に対応する。式 (4.10) 情報量を最大化する解釈候補が、妥当な候補として選択される。

すべての入力格フレームに対する解釈候補から最適な候補を選んだ後、次の文を受け付ける。図 4.11 は、ある文「ビデオの再生ボタンを押したが、動かない」を入力した場合の処理例を図示している。2 番目の節から 2 つの解釈候補 ν_1 と ν_2 が得られる。そして、 ν_1 の情報量は 1.58、 ν_2 の情報量は 1.00 となる。したがって、命題 ν_1 がこの節に対する解釈候補として選ばれる。そして同時に、「動かない」に関する「何が」の部分の省略について参照先の解決が行われることになる。

4.5.4 タスク処理部

タスク処理は、実験システムの応用業務をシミュレートする処理部であり、情報検索の検索処理やエキスパートシステムの推論処理を行う部分に相当する。本論では、実験システムの対象としてビデオ機器の操作案内を想定している。ビデオ機器の操作に関わる各種問合せを処理し、案内を行う。

タスク処理部は、タスク実行処理部と応答生成処理部からなる。

タスク実行部は、談話処理部から入力される格フレーム表現に基づいて、操作方法に関わる問題解決を行う。質問を以下の 4 種類のタイプに分けて処理している。

1. 装置の状態に関する問合せ
例: ビデオの電源はオフか
2. 目的-手段の問合せ
例: 再生するにはどうすればよいか
3. 仮想的な問合せ
例: 再生スイッチを押すとどうなるか
4. 操作ミスに起因する不具合を含む問合せ
例: 電源スイッチを押しましたが、ビデオが動かない

上記 4 タイプに対して、図 4.12 に示すオートマトンを用意し、処理するようにした。入力文を受理したオートマトンに応じて、それぞれのタイプごとの問題解決用モジュールを起動する。例えば、タイプ 2 の「再生するには、どうすればよいか」という発話に対しては、図 4.12 に示すオートマトンの中央のパスを通ることになる。

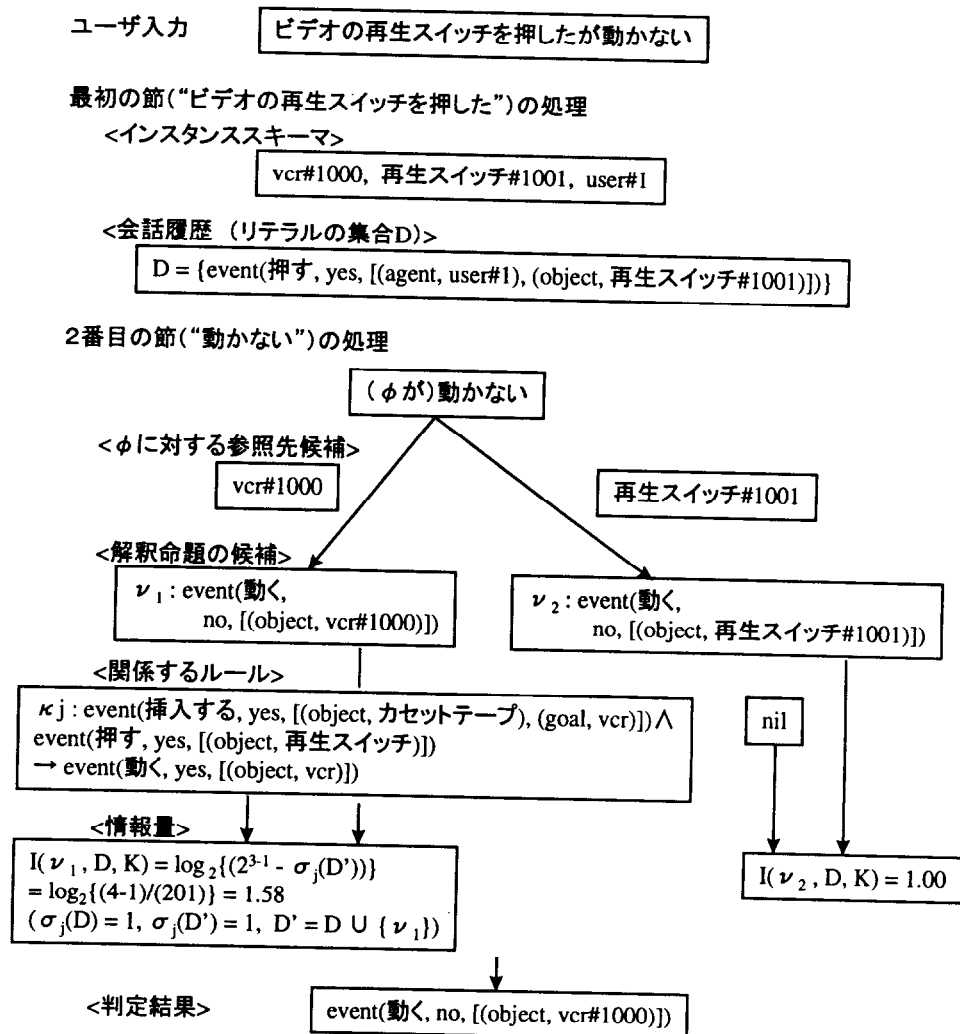


図 4.11: 照応解決処理の処理例

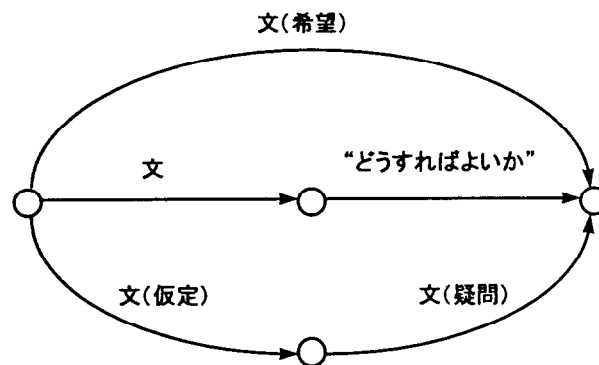


図 4.12: 発話の解析用オートマトンの例

この場合、「再生する」をゴールとして、知識処理部に対して後向き推論を行うよう命令を出し、実行されていない操作を取り出す。

タイプ1の場合、問い合わせられている対象となるスキーマのスロット値を取り出す。タイプ3の場合は、知識処理部に、ユーザが明示している事実に登録した後、前向き推論を行うよう命令を出す。そして、その結果得られた状態変化を調べる。また、タイプ4の場合は、知識処理部のルールを探索し、実行されていない操作を取り出す。

応答生成部は、タスク実行部から、応答生成のメッセージの内部表現を受け取り、自然言語に変換して表示する。自然言語だけではユーザに十分な情報を与えることができない。そこで、マニュアルの一部を提示する機能を有している。これは、会話の話題となっている操作、入力文に含まれる単語に基づいて、それらの語句を含むマニュアルを表示する。

この質問応答システムは、PSI-II²上に構築した。約1000単語の辞書、100のイベント、200のオブジェクトに対応するスキーマと、30のif-then規則からなる。

表4.2に用意したif-then規則の一例を示す。表4.2では、わかりやすくするために、自然言語で例を記述しているが、実際は、図4.8に示した形式で登録されている。実験システムによる会話例を図4.9に示す。

4.5.5 照応解析精度に関する会話実験

開発した質問応答システムを用いて会話実験を行った。この実験では、照応解決の処理手続きを通じて本論で述べた議論を評価することを目的としており、質問応答システム全体の評価は目的としていない。ここで、会話実験に以下の制約を課す。図4.14に示すような文のペアを被験者に示す。そして、その後に入力する文を被験者に考えてもらうことにより、評価文を収集した(9対の文のペアを4人の被験者に提示し、各ペアに後続しうる文を5文列挙)。収集した文のうちいくつかについては、未登録語などが原因となって、処理に失敗する。そこで、このような文を取り除き、

²第五世代コンピュータプロジェクトにおいて開発された逐次推論計算機

表 4.2: 実験のために記述した if-then 規則

内容	例
カセットの取扱い (2 ルール)	if 「カセットを VCR に入れる」 and 「イジェクトスイッチを押す」 then 「カセットが VCR から出てくる」
録画 (6 ルール)	if 「カセットを VCR に入れる」 and 「チャンネル選択スイッチを押しチャンネルを選ぶ」 and 「録画スイッチを押す」 then 「番組を録画する」
録画予約 (4 ルール)	if 「カセットを VCR に入れる」 and 「プログラムスイッチを押し予約番号を選ぶ」 and 「チャンネル選択スイッチを押しチャンネルを選ぶ」 and 「曜日スイッチを押し曜日を選ぶ」 and 「時間選択スイッチと分設定スイッチを押し開始時間を設定する」 「時間選択スイッチと分設定スイッチを押し終了時間を設定する」 「録画スピードスイッチを押し録画スピードを設定する」 「タイマースイッチを押す」 then 「録画が予約される」
再生 (12 ルール)	if 「カセットを VCR に入れる」 and 「再生ボタンを押す」 then 「カセットを再生する」
その他	電源のオン/オフ, タイマーセットなど

ユーザ: ビデオテープを再生したい。
システム: ビデオの電源を入れて、ビデオにカセットテープを入れて、再生ボタンを押して下さい。
ユーザ: ビデオにカセットテープを入れて、再生ボタンを押した。
システム: はい、わかりました。
ユーザ: 動かない。
システム: ビデオの電源はオンですか。
ユーザ: はい。
システム: 再生ランプは点灯していますか。
ユーザ: 点灯しています。
システム: VHF 出力端子はテレビに接続されていますか。
ユーザ: いいえ。
システム: VHF 出力端子をテレビに接続して下さい。

図 4.13: 実験システムでの会話例

113文を収集した。これらの文には174の文節が含まれていた。そして、204の照応表現が含まれていた。実験の結果、190の参照先について正しく見つけ出すことができた。内訳を表4.3に示す。

表 4.3: 照応解析結果

照応表現	参照タイプ	正解	合計
省略	object	86	90
	event	12	18
代名詞	object	3	4
	event	3	3
名詞	object	74	77
その(この)+名詞	object	12	12
合計	—	190	204

例えば、以下のような文が正しく処理できた。

System：ビデオカセットテープをビデオに入れましたか。

User：入れたが、出てきてしまう。

この場合、「(φが)出てきてしまう」に関するφに対応する参照先としてはビデオカセットテープが妥当である。なお、VCRと映像なども意味的な制約を満足する。

User：ビデオの再生方法を教えろ

System：ビデオカセットテープをビデオに入れてから再生ボタンを押して下さい。

(a) 被験者に与えた対話の一例

ビデオに入れる前にボタンを押すとどうなるか。

(b) 被験者が対話 (a) の後に入力される文として入力した文

図 4.14: 対話実験で得られた対話の例

処理が1つ以上の解釈を生成したのは47回であり、そのうち6回正しい候補選択を誤った。これらのエラーは、候補選択ステップで生じた誤りである。例えば、実験システムでは、最大10個の先行オブジェクトまでを対象としたため、照応表現に対して参照候補を見付けるには履歴上で距離が離れすぎていた場合、正しい候補を見付けるのに失敗している。一方、if-then規則に関係があった解釈についてはすべて正しく処理された(47回のうち32回)。全体としては、91パーセントの解釈が正しく得られた(照応表現に関する参照先としては93パーセント)。

解釈選択部を動作させない場合、得られる解釈のうち80パーセントのみが正解となる。この場合、最近傍で意味的な制約を満足する参照候補から構成した解釈に相当する。したがって、解釈選択部は、解釈の誤りの割合が半分に減らしていることになる。

4.6 本章のまとめ

自然言語理解システムにおいて曖昧性解消のための尺度として文解釈における情報量を用いる手法を述べた。この手法は、因果関係についての知識を含む聞き手の理解モデルに基づいている。妥当な解釈を求めるため、解釈の情報量を定義し、さらにこの尺度を照応の曖昧性を解消する手続きに適用した。そして、この処理を、質問応答システム上に実装し、会話実験を行った。93パーセントの参照先を正しく求めることができた。この実験によって、定義した尺度の有効性が確認できた。提案した処理を導入しない場合に比べて、誤りの約半分を減らすことができた。

本章では、聞き手の知識の形式的な定義に基づいた最初の試みであり、自然言語における曖昧性、特に照応の曖昧性を解消するための定量的な尺度を与えるものである。これらは、照応を取扱うための推論規則を用いたいくつかの研究があるが、それらは定量的な尺度を定義したものではない。本章の議論は、推論規則の利用に関する理論的な背景を与えるものである。

本章では、命題論理に基づく枠組を述べた。命題論理は、会話文に定量表現などを含まない簡単な質問応答システムでは十分な仕組みである。しかし、様相や時称などを含む自然言語文を表現するには、十分ではない。このため、この点での拡張が必要である。

第5章 文書構造に基づく抄録作成

5.1 本章の概要

近年、個人が扱わなければならない文書情報が増大しつつあり、あらかじめキーワードづけされていない文書でも効率よく検索するため、全文検索の需要が増している[57]。検索するという行為は、多くの場合それ自体が目的というわけではない。むしろ、検索した文書を読み、理解すること、内容を参考にすること、再利用することが本来の目的であると考えられる。このような点から、検索行為の効率化を支援するためには、検索結果の表示についての配慮が必要である[48]。

効率的な文書提示を可能にするための一つの解決策として、抄録の利用が考えられる。従来のいわゆる文献検索システムでは、データベースに抄録を格納しておき、検索時に利用する。この抄録の作成では、人手で行わなければならないため、多大な人的資源が必要となっている。また、抄録を作成する人によって、偏りが生じるなどの問題もあり、自動的な抄録の作成が望まれている(例えば、文献[64])。

本章では、文書構造解析に基づく抄録生成方法を提案し、試作した抄録生成システムについて述べる。文章には、文章のまとまりやそのまとまりの間の相対的な関係が存在している。文章の理解や要約の作成のためには、これらまとまりや関係の抽出が必要であると考えられる。本章では、文章のまとまりの間の関係を、修辞関係(例示、背景、換言など34種類)を用い2分木で表現した修辞構造を定義し、この構造の抽出を行う。

抄録生成は、修辞関係ごとに定義した左右のノードに対する重要度に基づいて、文の重要性を評価し、抄録文から削除すべき文を判定する。抄録文を生成する際には、修辞構造を参照し、接続表現の付け替えを行い、原文における論旨の保存を図る。

従来、この種のシステムは、例題の解析にとどまり、十分な評価が行われておらず、有効性が明確ではなかった。ここでは、システムの生成する抄録について重要文の捕捉率という観点からと、文書検索における文書提示の観点から評価を行った。その結果、重要文の捕捉率が、重要文 51%、最重要文 74% (技術論文)であることを確認した。さらに、文書提示の観点からの評価では、人が一定時間内に所望の文書を探し出せる率で、原文を読む場合に比して平均 60% の効率向上が認められ、システムの有効性が確認できた。

以下の節では、関連研究について述べた後、修辞構造の解析方法ならびに、その修辞構造に基づく抄録生成方法について述べる。そして、評価実験の内容とその結果の考察について説明する。

5.2 関連研究

従来より、自動的に抄録や要約を生成するための方式検討、実験システムの開発がなされてきている。従来研究を手法別にまとめると、以下の3つの手法に大別できる。

1. キーワード抽出に基づく手法
2. スクリプトなどの世界知識を利用する手法
3. 文の言い回し等の表層表現を手掛かりにする手法

1の代表的な手法としては、文献[41]が上げられる。この手法では、文書の特徴づける単語をキーワードと定義し、単語頻度を求めることによりキーワードを求める。キーワードを含む文を重要文として選択し、抄録を生成する。この手法には、キーワードを含む文が必ずしも重要文ではないこと、生成された文章が重要文の羅列にしかならないこと、などの問題がある。この手法の変形として、各文書が属している分野の語の頻度情報を用いてその文書のキーワードを推定し、キーワードを多く含む段落を重要段落として選び出す手法が提案されている[12]。これは、各文書に対して重要段落を抽出するものにすぎない。

一方、2の手法としては、文献[40]や[13]、[24]などが上げられる。この手法では、文章の内容に対応するスクリプト(ステレオタイプのな事象間関係)を記述することを仮定し、そのスクリプトに基づき文章中で明示されている事象の間の関係を抽出し、それに基づいて要約文を生成する。また、スクリプトのような事象間の関係以外にも、部分-全体関係のような概念間の関係知識を前提とした手法もある[98]。このような方式では、扱う文書の内容に依存した知識を必要とするため、内容を限定しない用途では、現実的とはいえない。

3番目に属する手法としては、文献[26]や[95]などがあり、今回開発した手法もこのカテゴリに属している。重要文を明示する言い回しなど、表層的な表現に基づいた処理であるため、処理対象の文書の内容や分野に依存しないという利点がある。しかし、従来の手法では、文単位の解析、あるいは隣接する文の間の解析など局所的な処理であり、全体的な文章の構造という視点からの分析が不十分である。

また、これまでの要約や抄録の研究では、数少ない例題の分析にとどまっており、開発したシステムの客観的な評価を行っておらず、その有効性が必ずしも明確にはされていないという問題がある。

開発した自動抄録システムは、文を単位とする2分木構造で文書構造を表現し、この構造を入力文書より自動抽出する。文書構造解析は、修辭的な表現から文章のまとまりを検知することにより、隣接する文間の関係だけでなく、文章のまとまりの間の関係も抽出する。この構造内の各関係は、前後の部分構造のいずれの側が重要かという情報を、関係ごとに定義することが可能である。この文書構造に基づいて、重要文の選択、接続表現の置換を行うことにより、原文との整合性が取れた抄録の生成を図っている。

以下の章では、開発した自動抄録システムの構成、検索システムの提示処理という観点から行った評価実験とその結果について述べる。

5.3 システムの構成

システムの構成を図 5.1 に図示する。システムは、文書構造解析部と抄録生成部のふたつのモジュール群に大分される。さらに、文書構造解析部は、書式解析部と修辞構造解析部とから構成される。

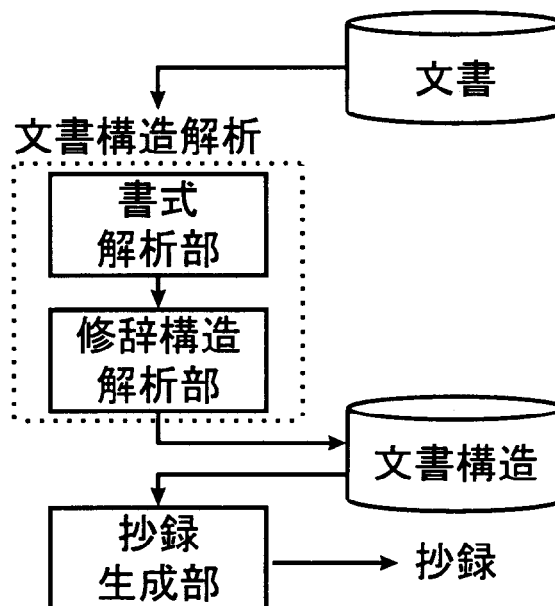


図 5.1: 自動抄録生成システムの構成

書式解析部では、文書の章や節の階層関係や、文や段落の切れ目の認識、章や節の見出しの抽出といった解析を行う。解析にあたっては、章節見出しにふられている章や節の番号、文書中の句点や空白などを手掛かりにして解析を行っている。

一方、修辞構造解析では、章や節を構成する文章本体を対象として、文と文の間の修辞的な関係（例えば、“例示”、“背景”など）を抽出するものである。書式解析結果と修辞構造解析結果とは、互いに関連づけて文書構造として取り出し、ファイルに格納する。

抄録生成は、文書構造解析で抽出された文書構造にしたがって抄録を生成する処理を行う。

開発した抄録生成システムは、検索システムにおける文書提示部での利用を前提としており、文書の提示は、実時間で行われなければならない。しかし、文書構造の解析は、現在のところ実時間処理が可能ではない。そこで、検索に先立って文書データベース中のすべての文書について、あらかじめ文書構造解析を行っておき、格納されている文書構造にしたがって抄録を生成することにより、実時間での抄録生成を実現している。

文書構造解析と抄録生成の具体的な処理は、5.4 と 5.5 において述べる。

5.4 文書の構造化

科学技術論文など比較的長い文書は、通常複数の章や節からなっている。章や節は、その文章の書き手が持っている意味的なまとまりに対応しており、検索処理や提示処理においても区別して扱える必要がある。

近年、SGMLやODAなどの標準仕様や、TEXに基づいて、タグを用い構造を定義した文書が増えつつある。このような文書では、章や節の構造は明確に定義されており、改めて解析する必要はない。しかし、今だに大半の文書は、ワープロなどで作成されたタグづけされていない文書であるのが現実である。

章や節は、その見出しとともにその本体であるいくつかの文からなる文章を含んでいる。これらの文の論理的あるいは意味的な関係を捕らえるためには、これらの文章の構造についても解析する必要がある。

また、各章や節内の文章の構造解析では、文間の関係を捕らえるだけではなく、内容上のまとまりなどの認識も必要と考えられる。

ここでは、2つの文書階層（書式構造と修辞構造）から文書の構造を抽出する。書式構造は、文書の章や節の階層構造を表現する。一方、修辞構造は、各章や節内の文や文のまとまりの間の論理的な関係を表現する。

書式構造については、章や節の単位やその階層的な関係は客観性が高く、表現上の問題はさほど生じない。また、その分析についても、章見出しや章番号、インデントーションなどを手掛かりとして分析することができる。しかし、修辞構造については、どのような言語単位を解析の対象とするのか、それらの言語単位をどのような関係で関連づけるのか、というような表現上の問題が大きい。

従来文書の構造化に対していくつかの試みがなされている（[16, 10, 11, 59] など）。

Mannらは、修辞構造の単位を節（clause）とするとともに、その間の関係を分類して文章の構造を記述することを試みている[42]。しかし、入力の記事から構造を組み立てるための具体的な処理手続きについては議論していない。

辻井は、文のタイプを事実文、判断文、要望文などに分類し、文脈自由文法でこれらの文タイプの間の関係を記述することにより、文章の構造化を試みている[86]。また、Schaらも辻井と同様に文脈自由文法による文章の構造化を試みている[75]。しかし、文章の構造を記述するには文脈自由文法は制約がきつすぎるため、実際に文書構造を記述するには無理があると考えられる。

Cohenは、接続詞などの表層的な手掛かりと命題間の論理的な関係を元に文章の構造化するモデルを提案しているが、具体的な処理手続きを計算機上で実現していない[7]。

黒橋らは、2文間の結束関係を抽出するためのヒューリスティックな規則を、主題の連鎖や接続詞などの情報に基づいて記述し、それにより文章を構造化する手法を提案している[38]。2文間の結束関係だけを手掛かりとしているので、大局的な構造を生成する手法としては無理がある。

田村らは、助詞、接続詞、指示語、時制などの複数の情報をパラメタとした重回帰分析によって、文章の構造的な切れ目を確度を推定するトップダウンな解析と、局所的な構造の解析を接続詞などを手掛かりに文間の結束性を検証するボトムアップな解析を組み合わせた文章構造解析の手法を提案している[85]。この手法では、大

局的な構造を示唆するような修辭的な表現を構造解析に利用することができない。

本研究では、文章の切れ目などをルールベースのトップダウンな処理で検出するとともに、接続詞などを分類した修辭関係の連なりから導き出される議論の流れの局所的な制約に基づき、文章の構造化を行う。修辭構造の表現の単位を文とし、抄録生成についても文単位で取捨選択を行う。文は、意味的なまとまりとして明確に区切られた単位である。したがって、節単位の処理に拡張する場合も、文を基本単位とする処理に追加する形で対応できると考える。

修辭構造の抽出は、接続詞や照応表現、文末表現など言語的な手掛かりに基づいて行うものとする。解析対象の文章の話題に関連する世界知識がなければ、正確な構造が得られない場合もある。しかし、あらゆる分野の世界知識を準備することは困難であり、検索システムを前提とした処理では、分野を限定することも現実的ではない。そこで、ここでは、世界知識には頼らない方法とした。

5.4.1 書式解析

文書自動レイアウトのため、書式構造の解析が提案されている [27]。書式解析においても、ほぼ同様の処理を行う。

具体的には、見出しを認識するために、先頭が数字や記号で始まるという情報、行の末尾に句点が存在していないという情報、などの情報を文章中より取り出している。また、ある行の先頭が空白であることを検出し、その位置が段落の開始であることを認識する。

書式構造の解析では、節ごとに見出しとその節が含む文章本体の範囲、その文章本体を構成する各段落の先頭の位置、などを出力する。そして、この解析結果は、原文の文書と対応してデータベースに格納する。また、各章の文章本体は、次の修辭構造解析においてその内部構造を解析する。

5.4.2 修辭構造解析

5.4.2.1 修辭構造の表現

修辭構造は、自然言語文の統語構造のアナログから2分木で表現している [80]。統語構造中の任意の部分木が、文法的な構成要素をなすように、修辭構造における部分木は議論上の一構成要素をなす。ある部分木は別の部分木を直接の子ノードとする木に付与される関係によって範疇化される。

修辭構造は、段落間の構造と段落内の構造の2階層からなるものとしている。すなわち、段落間の構造は段落を終端ノードとする構造を、段落内の構造は文を終端ノードとする構造をそれぞれ表現している。

図 5.2 の文章例 (下線部は修辭構造解析の手掛かりとなる表現である) に対応する修辭構造を図 5.3 に示す。図 5.3 の修辭構造において、数字は図 5.2 で示した文章中の各文の文番号を表している。また、“逆接”、“順接”などは文間の関係を表しており、修辭関係と呼ぶ。

- 1: 区間分割相関方式は、収束が保証されているという利点をもっている反面、全区間をひとわり修正するのに、トレーニング信号をN回受信しなければならない、収束に時間がかかるという欠点をもつ。
- 2: そこで、収束を速くするための便法として、区間の仕切を取り払い、トレーニング信号受信のたびに(21)式を全区間[L、M]に一斉に適用する方式を考えた。
- 3: これを、“部分相関方式”と呼ぶ。
- 4: 部分相関方式では、当該サブ区間以外のタップ利得が一時的に凍結されているという条件は成り立っていないので、収束は理論的には保証されていない。
- 5: しかし、修正係数 α を十分小さく選べば、近似的にこの条件が満足されているから、実用上は収束が期待しうる。

図 5.2: 文章例

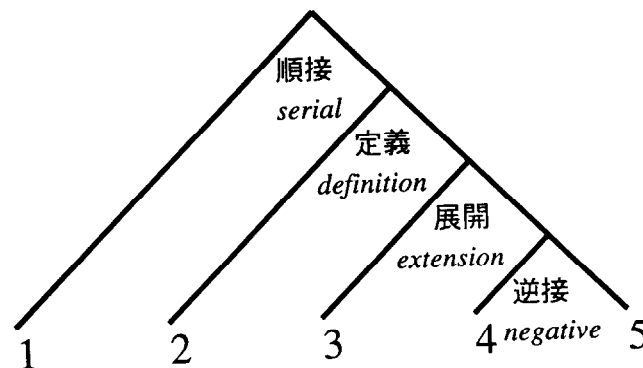


図 5.3: 修辞構造

図 5.3 に示す修辞構造は、第 5 文が第 4 文に対して逆接という関係に、第 3 文から第 5 文のまとまりが第 2 文に対して定義という関係に、そして、第 2 文から第 5 文のまとまりが第 1 文に対して順接という関係にあることを示している。

修辞関係は、接続詞や文末表現など文間の接続的な関係を明示する表現を、34 種類に分類したものである。表 5.1 に修辞関係の例と、その修辞関係に相当する表層的な表現例を示す。

表 5.1: 修辞関係の例

関係名	表現例
順接	したがって, そうすると ...
理由	なぜなら, というのも, ... だからである。
対比	一方, 他方
例示	例えば, この他, ... は一例である。
背景	従来, ... されるようになってきた。
逆接	しかし, それにもかかわらず
話題	本特集では, ここでは
定義	これを ... と呼ぶ
喚言	つまり, すなわち
展開	この, その

修辞構造解析では、この修辞関係を、文と文の間関係だけを表現する関係子としてではなく、文のまとまりの間関係表現する関係子として取り扱う。

5.4.2.2 修辞構造の解析

修辞構造の解析では、書式構造で得られた各節の文章本体に対して、以下のような手順で解析を行う。

(1) 文解析

まず始めの処理として、文章本体中の各文に対し形態素解析ならびに構文解析を行う。

(2) 修辞関係の抽出

各文の形態素解析結果あるいは構文解析結果に対して、その文に先行する文章に対する修辞関係を取り出す。そして、各文の文番号と修辞関係からなる系列を出力する。この系列を修辞関係系列と呼ぶ。例えば、図 5.2 に対しては、以下のような修辞関係系列が抽出される。

(1 順接 2 定義 3 展開 4 逆接 5)

表 5.1 で例示した表現パターンと修辞関係の対応情報は、修辞関係抽出規則と呼ぶ if-then 型の規則に記述している (付録 B の表 B.1 から表 B.3 に 34 種類の修辞関

係を示す)。各条件部には、表層表現、形態素解析列までのパターンが正規表現で記述できる。そして、最長一致するパターンに対応する修辭関係が、その文の修辭関係として選ばれる。なお、関係を明示する接続表現がない場合は、“展開”という関係を割り当てている。

(3) セグメンテーション処理

(2)の修辭関係の抽出で解析される修辭関係系列内の各修辭関係は、解析対象の文ごとに求められた関係にすぎない。一方、文章中には、複数の文にわたって構造を規定するような修辭的な表現も存在する。このような複数文にわたる表現から、構造化に対する制約情報を取り出す。制約情報はセグメンテーション規則と呼ぶ規則で記述し、ルールベースで処理を行う。

図5.4に、この規則の例を示す。“Surface”には照合する文字列を正規表現で表現したパターンを、“Relation”には照合する修辭関係を、“Example”には照合する文章の一例を、“Output”には修辭関係系列に埋め込む制約を、それぞれ記述するようになっている。入力文章が、“Surface”に記述した文字列のパターンと“Relation”に記述した修辭関係との両方に照合した場合、修辭関係系列に“Output”で記述した制約が埋め込まれる。

Example	確かに ...。...。しかし...。 <i>of course</i> <i>however</i>	
Output	[1 ...] 逆接 2	
Surface	^ (確かに 本来なら) <i>of course</i> <i>properly speaking</i>	-
Relation	-	逆接 <i>negative relation</i>

図 5.4: セグメンテーション規則の一例

図示した規則では、“Surface”の項目に記述されている記号“^”が文の先頭を意味しており、“確かに”や“本来なら”という表現が文頭にある文の何文か後に、逆接の関係を持つ文が来た場合、その文の直前でまとめるという知識を記述している。なお、“Surface”や“Relation”の項目に記述されている“-”は、その部分に関しては無条件に照合することを示している。

(4) 構造候補生成

セグメンテーション処理で付与された制約を破らない2分木構造を生成する。そして、生成された構造の優先度を求め、その候補を優先して出力する。

このステップでは、隣接する修辭関係の間で成立する論旨の流れの開始/終了に着目し、優先度づけを行う。例えば、(1 例示 2 順接 3)という修辭関係系列を考えてみた場合、例を述べるという議論が文2で終了していると判断でき、((1 例示 2) 順接 3)の方が(1 例示 (2 順接 3))よりも自然性が高い。そこで、このような構造の選好性を、2つの修辭関係間のすべての組合せ(34 × 34 組)についてあらかじめ

判定し、規則として準備した。この規則を用いて構造の優先度づけを行う (付録 B の表 B.4 と表 B.5 にこの構造選好規則を示す)。

具体的な処理としては、上記例のように優先度の低い側の部分構造を修辞構造が含んでいる場合にペナルティを課す。そして、修辞構造の候補ごとにペナルティの合計を求め、その値が最も少ない候補を優先する。

段落間の構造は、段落の先頭の文の修辞関係に基づいて、段落内の構造生成と同様に解析する。段落間の構造は、各段落内の文間の構造と相似形となっており、段落間の構造では、その終端ノードが文ではなく段落ということになる。

修辞構造も、書式構造と同様に原文の文書と対応させてデータベースに格納する。

5.5 抄録生成

章や節ごとに修辞構造に基づいて抄録文を生成する [6]。抄録の生成は、あらかじめ定めた修辞関係の前後ノードの相対的な重要度に基づく処理である。修辞構造は、段落間の構造と、各段落内の文間の構造の 2 つの階層からなる構造であり、それぞれの階層における構造表現は相似形となっている。そこで、抄録生成では、すべての段落について抄録生成処理を行った後、段落間の修辞構造に基づいてまったく同じ処理を行う。

処理は、重要性評価処理、既約処理、接続表現置換処理の 3 つのステップからなる。

(1) 重要文評価処理

重要文評価処理では、修辞関係の相対的重要度に基づいて、段落を構成するすべての文についての重要度を判定する。修辞関係の相対的重要度は、3 つのカテゴリに分類している。一例を表 5.2 に図示する。例えば、*RightNucleus* に分類された順接は、右ノードが左ノードに比べ重要であることを意味している。

表 5.2: 修辞関係の相対的重要性

タイプ	修辞関係 (表現例)	重要なノード
<i>RightNucleus</i>	順接 (よって), 換言 (つまり), 逆接 (しかし), ...	右
<i>LeftNucleus</i>	例示 (例えば), 定義 (...と呼ぶ。), 補足 (ただし), ...	左
<i>BothNucleus</i>	並列 (また), 対比 (一方), 展開 (これは), ...	等価

各ノードの重要度を決定するために、修辞関係の相対的重要度に応じて左右ノードにペナルティを再帰的に与えている。例えば、*RightNucleus* に分類された関係で

は、左ノードに1ペナルティを、また、*LeftNucleus*に分類された関係では、右ノードに1ペナルティを、与える。トップノードから順次このペナルティを加えることにより、各中間ノードについてのペナルティを計算することができる。

(2) 既約処理

既約処理では、所望の文数になるまで、ペナルティの大きいノードを再帰的に刈り取っていく。そして、残った文(重要文と呼ぶ)を出現順に並べることにより、抄録文を生成する。なお、抄録文として選ばれた文数が規定の文数以下に達しない場合は、前方の文を優先して残すものとする。

図5.3の構造に対して、上記の重要文評価を行うと、図5.5のようになる。破線はペナルティの境界を表しており、例えば、第1文、第3文、第5文には1ペナルティが、第4文には2ペナルティがそれぞれ課せられている。この場合、最短の抄録は、第2文目だけとなる。また、次に長い抄録は、第1文目と第2文目の2文からなる抄録である。

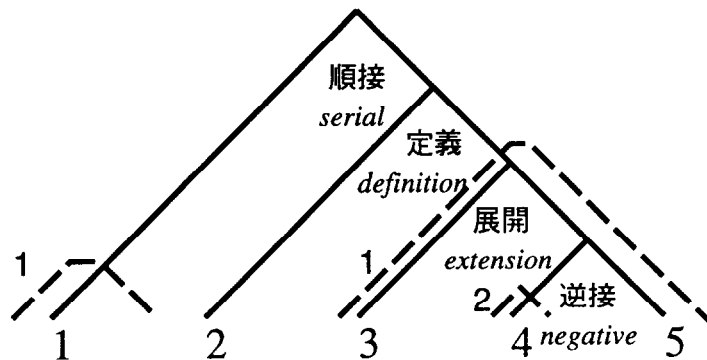


図 5.5: 重要文評価

(3) 接続表現置換処理

接続表現置換処理では、文間の修辞関係の一貫性が保たれるように、接続表現の付け換えを行う。接続表現の付け換えは、修辞構造上で、以下のような条件が成立した場合に行われる。

```

if  左部分木 L に重要文が存在
    and 右部分木 R に重要文が存在
    and R の最左端の文 C が重要文でない
then R 内の最左端の重要文 X の接続表現を
    C の接続表現で置き換える
  
```

例えば、“1。例えば、2。すなわち、3。”というような文章で、(1 例示 (2 換言 3)) という修辞構造が得られた場合を考える(ただし、1、2、3は、各文の接続表現を除いた部分)。既約処理では、1と3が重要文であると判定されるので、単に文を並べるだけでは、“1。すなわち、3。”という文章になってしまう。しかし、接続表現置換処理により、“1。例えば、3。”という文章が得られ、原文章における文間の修辞関係が保存される。

5.6 評価実験

開発した抄録生成サブシステムが生成する抄録について評価を行った。評価では、キーセンテンスの捕捉率の面と、検索システムにおける提示機能の効率化の面の2つの観点から評価を行った。評価実験の方法とその結果について述べる。

5.6.1 評価実験 1

自動抄録の精度について、人が選んだキーセンテンスの捕捉率という面から評価を行った。評価方法は、以下の通りである。

1. キーセンテンスの抽出

一人の被験者に原文章を提示しキーセンテンスを指摘させる。なお、キーセンテンスは以下の重要文と最重要文の2種類である。

- 段落内で最も重要だと思われる文を、一つ指摘させる (段落キーセンテンス)。
- 選択した段落キーセンテンスの中から、文章全体のキーセンテンスを、1/3 ~ 2/3 の比率で指摘させる (重要文)。
- 重要文の中から、文章全体を代表する文を1つ指摘させる (最重要文)。

2. 自動抄録の評価

各文章から自動生成した抄録について、キーセンテンスの捕捉状況を調べる。

上記の評価方法により、自動抄録の生成する抄録文章の評価を行った。評価の結果を表5.3に示す。なお、評価実験に用いた文章は、社説は朝日新聞社説、技術論文は東芝レビューであり、朝日新聞社説の場合は全文を、東芝レビューについては、ある節の少なくとも4文以上からなる文章を無作為に選んだ。また、評価時点での修辞関係抽出規則は約1850、セグメンテーション規則数は約150であった。表5.3において、「学習」とは修辞構造解析の解析規則を記述するために参照した文書を意味しており、「評価」と上記の記述時に参照しなかった文書を表している。表は、それぞれの場合についての捕捉率を示している。

表 5.3: 自動抄録結果の評価

(1: $\frac{\text{抄録中の文数}}{\text{原文中の文数}}$, 2: 重要文の捕捉率, 3: 最重要文の捕捉率)

文章		文章数	文数比 ¹	捕捉率 ²	捕捉率 ³
論文	学習	42	0.24	0.64	0.81
	評価	42	0.24	0.51	0.74
社説	学習	98	0.30	0.47	0.54
	評価	30	0.30	0.41	0.60

評価文書における重要文や最重要文のの補足率は、学習文書に対して若干下がるものの、それほど劣化は見られないといえる。

次に、人間がキーセンテンスを選択するという方法で抄録を作成(国文学博士課程の学生が作成、文章数10編)した場合について、作成された抄録を上記と同じ方法で評価した。人間が作成した抄録の評価と自動抄録と評価では文書数が異なるため、厳密な意味での比較にはならないが、表5.4にその比率を示す。

表 5.4: 重要文捕捉率の比較
(比率 = $\frac{\text{重要文捕捉率(自動)}}{\text{重要文捕捉率(人間)}}$)

文章	重要文			最重要文		
	自動	人間	比率	自動	人間	比率
技術論文	0.51	0.6	0.9	0.74	0.8	0.9
新聞社説	0.41	0.6	0.7	0.60	0.6	1.0

人間が作成する場合と定量的に比較すると、文章キーセンテンスに関しては7割から9割、最重要文に関しては9割以上のキーセンテンスの補足性能であるといえる。

生成された抄録を目視した結果、自動抄録システムでの誤りの原因は、以下のような内訳となった。

- 修辞構造の解析の誤りに起因するもの (約7割)
- 重要文評価での誤りに起因するもの (約1.5割)
- 意味内容的に重要であるにもかかわらず抄録文からもれてしまったもの (約1.5割)

誤りの原因の大半は、修辞構造解析であり、修辞構造の解析に関しての精度向上の必要性が認められる。

5.6.2 評価実験2

次に、検索結果の文書提示という観点から、システムの生成する抄録の有効性を検証した。評価の方法は以下の通りである。

- 問題作成者が原文の文書集合を参考に設問を作成する。
- 問題作成者とは異なる被験者に原文の文書集合あるいは抄録文の文書集合を与える。
- 被験者はすべての設問に対して、設問に対応する文書をすべて列挙する。
- 設問ごとに被験者が見つけれられた文書数と正解の文書数の比(再現率)を求める。

なお、一人の被験者に同じ内容の原文と抄録文とを提示することはできないので、別の被験者が行った結果同士で再現率を比較することになる。

実験パラメータを設定するために、次のような点を考慮した。

- 解答時間について

解答時間の設定については、問題数に対して十分な解答時間を与える方法と、少めに解答時間を与える方法とが考えられる。前者の場合、見直しの時間や考える時間などが取れる。その時間をどの程度取るかが人によって一定でないため、正解率や解答時間にその影響が及ぶ可能性がある。そこで、実験では後者の方法を取ることにした。

- 問題数について

問題作成者や被験者にあまり負担をかけないように、文書数を 20、解答時間を 5 分と決め、事前に別の文書集合と問題に対して試行実験を行った。この試行実験により、7 問程度の問題を解くのが上限であることがわかり、問題数を 7 問と設定した。

- 生成する抄録の長さについて

開発したシステムが生成する抄録は、「要旨的要約」(原文の中心思想、結論)ではなく「大意的要約」(原文の各段落の要旨を繋げた全文の縮小相似形)に相当している。大意的要約の場合、「原文の 3 分の 1 から 4 分の 1 くらいにまとめる」[72] という国語教育における目安がある。そこで、本実験では文数比を 30% という値に設定した。

以上の点を考慮し、実験のパラメータを以下のように設定した。

- 一度に与えられる原文もしくは抄録文の文書数を社説 20 文書とした (タイトルは見せない)。
- システムが生成する抄録文を原文との文数比が 30% になるように設定した。
- 設問の数を 7 問とした (設問の例を付録に示す)。
- 問題を解く時間を一定 (5 分) とした。
- 提示する文書は印刷物を用いた。
- 一人の問題作成者が用意した問題群に対して、被験者は二人とした。

問題文作成では、正解とする文章の中心的内容を、問題作成者が原文のみを読んで判断し (提示する抄録文は参照しない)、その内容を 1 行以内で表現する問題文を作成した。また、異なる文書集合については、問題作成者を別にするにより、評価結果が問題作成者の問題作成傾向に偏らないようにした。

評価結果を表 5.5 に示す。表は、与えられた設問に対する再現率について、抄録文提示の場合と原文提示の場合とで比較した表である。再現率は、全設問について正しい文書を列挙できた場合が 1 となる。したがって、再現率が高いほど文書提示として効率が良いことになる。

表 5.5: 文書提示手段としての抄録文と原文の比較 (A, B, C は被験者)

文書集合	集合 1	集合 2	集合 3	平均
抄録	0.92 (A)	0.71 (B)	0.29 (C)	0.64
原文	0.28 (C)	0.50 (A)	0.43 (B)	0.40

表 5.5 によれば、平均的に、抄録文提示の方が原文提示に比較して再現率が高い (60% (= $0.64/0.40$) 向上)。

文書集合 3 については、原文提示の方が再現率が高い。しかし、抄録文提示について評価した被験者 C は、原文提示と抄録文提示との両方に関して、他の被験者に比べ再現率が低い。このことから、文書集合 3 における再現率の逆転は、被験者間の差によるものと考えられる。

評価データの量や被験者数が少なく、客観的な結論を下すには、さらに多くの被験者で実験する必要がある。しかし、この評価結果は、抄録の提示が文書検索の提示機能として有効であることを示唆している。

5.7 本章のまとめ

文書構造解析に基づく抄録生成システムについて述べた。文書構造の解析結果を利用することにより、原文における文間の関係 (修辞関係) の保存を図っていること、構造の解析と抄録生成のフェーズとに分けることにより、実時間で抄録生成を行っていること、文書構造解析のための知識が分野に依存しないこと、などが本システムの特徴である。

従来のこの種のシステムの評価では、例題の解析にとどまり、その有効性が明確ではなかった。本論文では、システムの性能評価のため、システムの生成する抄録について重要文の捕捉率の観点からの評価 (システムの生成した抄録が、人が選んだ重要文を捕捉する割合で評価) と、文書検索時の文書提示の観点からの評価 (文書を読んで所望の文書を探す評価実験) を行った。これによれば、捕捉率は、重要文 51%、最重要文 74% (東芝レビュー) であること、また、その率は人と比較しても 9 割程度であることが明らかになった。さらに、文書提示の観点からの評価では、原文を読む場合に比して、平均 60% 程度の効率向上が認められ、システムの有効性を確認することができた。

今回のシステムは、検索システムの提示部として開発したため、評価もその応用に対応した形で行った。一方、抄録の性能自体の評価としては、文間や段落間の関係が抄録文で保存されているか否かの評価や、可読性の評価も必要である。これらの評価については今後の課題である。

開発した抄録生成システムは、文書検索システムにおける文書提示部に用いている。今後、精度の向上のため、解析規則を充実させること、現在の処理対象が技術論文、社説に限定しているので、この対象を拡大すること、実際の検索システム上で有効性を検証する必要がある。

第6章 文書構造に基づく文書検索

6.1 本章の概要

文書構造を利用した動的な抄録提示と、文書構造に基づく検索とを特長とした文書検索システムを開発した。

近年、ワークステーションの計算機パワーの増大にともない、全文文書を検索対象とした全文検索システムの実用化が進みつつある。しかし、現在実用化されている全文検索システムでは、指定した語句が文書中に存在するか否かの基準でのみ文書の適合性を判断するため、意図しない文書も検索されてしまう可能性が高い。また、検索してきた文書を表示する際、検索文書のタイトルの一覧を表示するか、あるいは始めに見つかった文書をそのまま表示するにすぎない。

大量の情報が利用可能になるにつれて、情報を様々な側面や特徴から検索し利用する技術の重要性が指摘されている ([52] など)。一方、現在の検索システムでは、出力された検索結果から必要な情報を得るための分析を行う必要があり、出力される検索結果が大量である場合にはその負担はさらに大きくなる ([48] など)。このような状況において、情報検索理論や技法の向上が求められている [23]。

効率的な検索を目的とし、文書の構造に基づく検索と抄録生成を行う全文検索システム BREVIDOC (Broadcatching System with an Essence Viewer for Retrieved Documents) を試作した [47]。効率的な文書検索のための文書提示手法について検討し、動的な抄録提示を行う検索インタフェースを提案する。また、文書の構造を解析し、その構造を利用して検索方式を提案する。上記の検索インタフェースならびに検索方式について述べるとともに、評価を行いその有効性を検証する。

6.2 文書提示のための文書の構造化

効率的な検索を行うためには、検索意図に沿った文書が的確に検索されとともに、効率的な文書提示インタフェースが必要である。本節では、上記インタフェースに必要となる項目を考察し、それに基づいた文書検索システムの構成について述べる。

検索するという行為は、多くの場合それ自体が目的というわけではない。むしろ、検索した文書を読み、理解すること、内容を参考にすること、再利用することが本来の目的であると考えられる。このような点から、検索行為の効率化を支援するためには、検索結果の表示についての配慮が必要である [57]。

例えば、複数の文書が検索された場合、文献抄録のように検索単位が十分短ければ、全文を提示しても検索者の負担にはならないかもしれない。しかし、長い文書

を検索単位とする場合、たくさんの文書が検索され、しかもその全文を読まなければならないとすると、目の疲れなど検索者の負担増をまねく。一方、文献抄録のようにあらかじめ用意された情報では、検索者がより詳しい情報を得たいと考えても、その要求を満たすことができない。

効率的な検索結果の提示を実現するためには、文書を何らかの意味で構造化することが必要である。この構造化により、必要な情報を適切な詳しさに検索者に提示できる検索インタフェースを実現できる可能性がある。ハイパーテキストは、このような観点に立った対話的な文書提示の一つのアプローチと言える。しかし、構造化を人手で行わなければならないという点が、多量の文書を処理対象とする全文検索システムに適用する上で問題となる。

従来の検索システムでは、検索対象の中心が文献抄録など比較的短い文書を検索対象とする場合が多かった。短い検索対象の場合、文書の構造を認識する必要性は小さい。しかし、今後全文文書の増大にともない、比較的長い文書を検索対象とする場合が多くなっていくと予想される[48]。

効率的な文書検索を行うための文書提示では、以下の2点が重要である。

- 一覧性

検索した文書の内容を一覧できることが必要である。

我々は、紙の上に書かれた文書を読む場合でも、いきなり頭から読むことはせず、まず始めに全体的な構成を概観した後に読み始める場合が多い。これは、始めに読む文書の全体像を知ることにより、個々の部分の文書全体に対する位置付けを把握し、理解する手助けにすると考えられる。

ハイパーテキストでは、現在読み進めている文書の断片が全体のどの部分を占めているかがわからず、ネットワーク中で読み手が迷子になってしまうことが指摘されており、全体に対する把握できる手法が検討されている(例えば、[4])。一方、紙に書かれた文書においては、一覧性を高めるために、目次のページを設けたり要約をつけたりというような工夫をしている。文書提示を効率的に行うためには、電子的なメディアであることを生かし、文書の一覧性を高めることが必要である。一覧性を高めることにより、検索作業の効率化が図れる。また、たくさんの文書から希望の文書を選択する際に経験する目の疲れなどを軽減する効果もあると考えられる。

- ユーザ主導性

利用者の指示にしたがって任意の箇所を詳しく読めることは、電子的なメディアとして必須の条件である。

利用者にとって、詳しく読みたい場所や量は、すべての文書に対して一率というわけではない。概要を知るだけで十分な文書もあるし、内容を詳しく理解しなければならない文書もある。たくさんの文書を取り扱っている場面では、読みたい箇所だけを必要最小限で読みたいというのが実情である。したがって、文書提示においても、柔軟な指定手段によって利用者が読みたい箇所や内容を提示できなければならない。

あらかじめ作成された目次や抄録は、文書の一覧性を満足しているが、ユーザ主

導性を満足することが難しい。そこで、ここでは、上記のような要件を満足する文書提示インタフェースとして、自動抄録生成を基本としたインタフェースを提案する。検索された文書から自動的に抄録を生成し、提示することにより、ユーザ主導性と一覧性を兼ね備えた文書提示インタフェースを実現する。

検索文書の提示手段として抄録生成を行う場合、十分な精度で抄録が生成できる必要がある。従来の単語頻度をベースにした抄録生成方法(例えば[41])では、十分な精度が得られない[6]。また、分野知識を用いた要約生成方法(例えば[24])では、要約を行う文書の話題に関する分野知識をあらかじめ準備しなければならない。検索の文書提示を考えた場合、この方法を取ることは現実的ではない。

ここでは、抄録を生成する手法として文書構造に基づく方法を取った[6]。この方法では、接続詞や文末表現など言語的知識だけにに基づき文章の構造化を行い、解析した構造に基づいて抄録を提示する。解析する文書の話題や内容によらない処理である。

6.3 検索精度向上のための文書の構造化

情報検索において、文書の適合性(relevancy)を識別するための様々な方法が研究開発されてきた。例えば、単語の出現頻度や段落や文の単位などを利用した単語の近接性などを利用した適合性の判定方法などがある。さらに、passage retrievalの分野などで、文書内容や文脈の構造を利用した検索方法が研究されている([3, 92]など)。文献[92]では、< purpose >, < abstract >, < start >, < summary >, < title >, < supplementary >, < misc >のタグ付けがされた TREC conferenceのテキストデータを対象とした検索実験を行い、< purpose >と< summary >の重みを2、< misc >の重みを0.5、その他の重みを1とした場合に最もよい検索結果を得たとの報告がされている。

文献[92]ではタグ付けがされたテキストデータを対象としているの対し、我々はタグ付けされていないテキストデータを対象とする点が異なる。タグ付けされていないテキストデータを対象としこれを構造解析する理由は、タグ付け、特に文書の内容に基づくタグ付けは労力を要するので、一般のテキストにタグ付けされることが期待できないことと、さらに、検索命令や文書の内容に基づく検索([76, 36]など)を行うには文書構造の解析が欠かせないからである。文書の構造化については文献[32, 33, 59]などの研究があるが、我々は文書構造化の検索への応用を図っている。

ここでは、文書構造のモデルとして、文書の種類と、その文書に含まれる典型的な情報の種類、あるいは論旨の展開に用いられる情報の種類を想定する。文書が伝達する典型的な情報の種類を、文の意味役割と呼ぶこととする。例えば、技術論文では「背景」、「目的」、「結果」、「結論」などが意味役割に相当する。また、新聞社説では、「背景」や「意見」などが考えられる。技術解説などの説明文では、「背景」や「話題」がこの意味役割に相当する。

語句がどのような意味役割で用いられているかを指定することができれば、適合する文書を絞り込みに効果があると期待される。このような意味役割を人手であらかじめ付与することは、コスト面から望ましくない。そこで、意味解析に基づく検

索では、この意味役割を自動的に抽出することとし、抽出された意味役割を検索時に利用する。

文の意味役割の抽出においては、検索への応用 (検索結果の絞り込みやランキング付け) を目的とするので、すべての文に何らかの意味役割を付与することはせず、抽出できる文のみから抽出する。また、一つの文から2つ以上の意味役割が抽出されることも許容する。この文の意味役割の抽出と、インデックスファイル作成のために文書構造解析システムを実現する。

6.4 システムの構成

6.2 ならびに 6.3 での考察に基づき、自動抄録機能を持つ文書検索システムを開発した。文書提示のための自動抄録生成と検索精度向上のための意味役割抽出は、ともに文書の構造化に基づいて処理を行う。システムは、全文検索システム [54] の全文検索機能を利用し、ワークステーション (東芝 AS4000 シリーズ) 上に実現した。図 6.1 にシステムの構成を図示する。

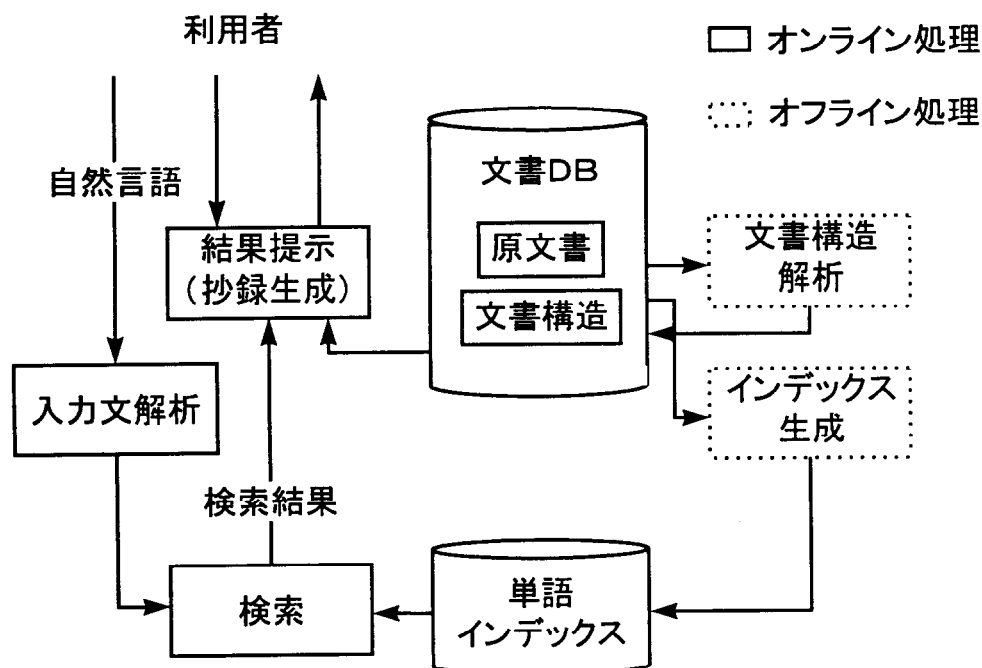


図 6.1: 動的な抄録提示機能を持つ文書検索システムの構成

システムは次の5つのモジュールから構成される。オフライン処理 (破線で記述) とオンライン処理 (実線で記述) からなる。文書DBを登録する際に、オフライン処理として、文書構造解析とインデックス生成が行われる。そして、検索時には、す

で構築されている文書DBにしたがって、利用者は、検索要求を自然言語により入力し、その結果の抄録を対話的にブラウズすることができる。

- 文書構造解析

章や節の構造、各章や節を構成する文章中の文間の修辭的な関係を抽出し、文書構造として表現する。抽出された文書構造は、原文書と対応づけて文書DB中に格納している。

- インデックス生成

高速な全文検索のために文書中に現れるすべての単語を抽出し、単語インデックスを作成する。抽出された単語は、その単語が含まれる文の“意味役割”と対応づけて単語インデックスに格納している[46]。ここで“意味役割”とは、文書の種類に対応して含まれているだろうと予測される情報を指し示すために、指標として設けたものである。例えば、科学技術論文などの場合、「話題」、「背景」、「目的」、「特徴」、「結論」、「課題」などがこれに相当する。なお、文の意味役割の抽出は、表層的な表現をベースに行っている。文書構造解析ならびにインデックス生成については、6.4.1で説明する。

- 入力文解析

利用者が入力する自然言語文を解析し、検索コマンドに変換する。検索コマンドは単語列と意味役割とのペアで表現している。例えば、「高精細なイメージ表示を目指したもの」というような入力文に対しては、次のような検索コマンドに変換されることになる。

目的:高精細/イメージ/表示

このモジュールにおいても、文書構造解析における意味役割の抽出と同様に、意味役割の抽出は表層的な表現をベースに行っている。

- 検索

単語インデックスを参照し、検索コマンドにしたがって文書を検索する[46]。この際、検索コマンドで意味役割が指定された場合、その意味役割に対応づけて単語インデックスに格納されている単語情報に基づいて文書を検索する。

- 結果提示

検索結果の文書から抄録を自動的に生成し提示する。抄録を生成する際には、文書DB中に格納されている文書構造を参照し、その構造に基づいて生成する。結果提示の内容については、6.4.2で説明する。

文書の構造解析はかなり処理時間が必要となるため、文書検索時に行うのは、現状の計算機パワーからすると現実的ではない。そこで、文書の構造解析と意味役割の抽出処理に関しては、実際の文書検索に先立って行う形を取っている。文書構造や文脈構造、文の意味役割の抽出結果は、テキストデータベース中の各文書と対応させてデータベース中にあらかじめ格納しておくようにしている。

一方、抄録生成は、処理パワーをさほど必要としないので、これらの格納情報を参照し、検索された文書ごとに検索時に実時間で処理するようにしている。これにより、詳細度の異なる抄録を動的に表示することができる。

6.4.1 文書構造解析

文書構造解析システムは、図6.2に示すように、テキストデータを入力とし、文の意味役割を抽出して意味役割ごとに単語インデックスファイルを生成する([46, 80])。日本語解析部には平川らの解析システム([19, 20])を用いており、インデックス作成部には中本らのシステム[54]を用いている。書式構造解析は、章節や段落、見出しといった構造を取り出す処理である。また、修辞構造解析は、第5で述べたように接続詞や文末の表現に基づいて文のまとまりや相対関係を取り出す処理である([65, 66]および[79])。

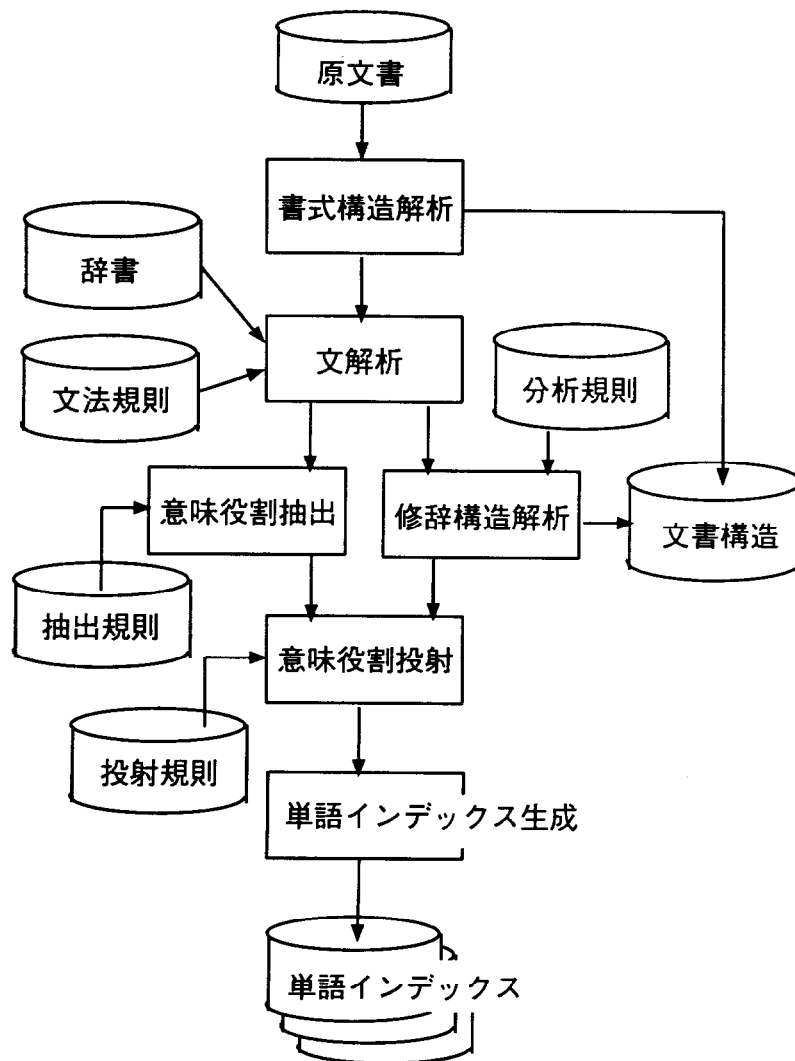


図 6.2: 文書構造解析システムの構成

文解析では、単語単位に分割する処理と各単語の品詞の判定を行う。一方、意味役割抽出は、文解析結果に意味役割抽出規則を適用し、何文目がどの意味役割であるかを取り出す処理である。抽出規則は、「近年」、「目的」のような語句や、文末の表現などを条件部とし、文解析結果との照合することで、意味役割を決定する。技術論文に関する意味役割の種類と対応する表現例を表6.1に示す。

表 6.1: 意味役割と表現例

意味役割	表現例
話題	～について解説する
目的	～を目的としたものである
特徴	～の特徴は～である
結果	～が得られた
結論	～を図った
課題	～をめざしたい
背景	～してきた
問題	～が問題となる

意味役割投射は、意味役割抽出で得られた意味役割を、修辞構造解析で抽出されている構造を用いて他の文の意味役割として複写する処理である。例えば、「以下の特徴がある。…。…。」のような列挙表現では、意味役割抽出で、1文目から「特徴」が抽出され、修辞構造解析で、列挙構造と認識される。この結果、意味役割投射により列挙構造に含まれる他の文に意味役割「特徴」が付与されることになる。

各文の単語の切り出しを行い、高速検索のための単語インデックスを、意味役割ごとに生成する。検索時には、例えば、「目的：被曝/低減」のような検索条件を指定することにより、意味役割(例の場合、「目的」)に対応する単語インデックスに基づいて、単語(例の場合、「被曝」や「低減」)を含む文書を取り出すことができる。

6.4.2 検索結果の提示

検索結果の提示では、以下の4つの詳細度で検索された結果を表示している。

- 文書群表示
検索要求との類似度にしたがって文書をグループ化して表示する。
- タイトルリスト表示
選択された文書群に属する文書についてタイトル、著者名などをリストとして表示する。
- 抄録表示
タイトル、原文アブストラクトの抄録、各節の見出しを表示する。表示されている節の見出しを指定することで、その節の抄録を新たに表示したり、詳しくしたり、簡単にすることができる。抄録の生成方法については5で述べた方法に基づいて生成する。
- 原文書表示
原文書表示では、抄録表示で詳しさを変更する操作を行うと、操作対象となった節に対応するあたりを表示する。

システムの表示インタフェースの画面構成を図 6.3 に示す。左上画面が検索要求の入力ウィンドウ、右上画面が検索結果の文書群表示ウィンドウ、左下画面がタイ

トルリスト表示ウィンドウ、中央下画面が抄録表示ウィンドウ、右下画面が原文表示ウィンドウである。

対話型検索システム

検索条件: 指定したものの 東芝レビュー 全文数: 386 文書 検索結果: 13 文書

検索結果

抄録

原文

東北電力(株)女川原子力発電所第1号機の新設

1. 東北電力(株)女川原子力発電所第1号機の新設
著者: 可児次郎, 藤田京, 滝口孝夫
内容: 東北電力女川原子力発電所第1号機は、電気
2. BWRの改良・開発
著者: 藤田京, 滝口孝夫
内容: 東北電力女川原子力発電所第1号機は、電気
3. 原子力発電所の改良・開発
著者: 藤田京, 滝口孝夫
内容: 東北電力女川原子力発電所第1号機は、電気
4. 原子力発電所の改良・開発
著者: 藤田京, 滝口孝夫
内容: 東北電力女川原子力発電所第1号機は、電気
5. 原子力発電所の改良・開発
著者: 藤田京, 滝口孝夫
内容: 東北電力女川原子力発電所第1号機は、電気

東北電力(株)女川原子力発電所第1号機の新設

可児次郎(1) 藤田京(2) 滝口孝夫(3)

東北電力(株)女川原子力発電所第1号機は、電気出力52.4
MWのBWR発電所であり、昭和59年6月1日営業運転を開始
した。

1. まえがき

2. プラント仕様と主な特徴

3. 建設工事の特徴

3.1 原子力容器圧力容器(PCV)掘削工事

3.2 (イ) 高圧容器の設置

3.3 高圧容器の掘削工事

3.4 クリーンプラントの建設

3.5 計画外プラント停止の回避

3.6 48.5か月間の建設工事の達成

4. 試運転の概要

5. あとがき

東北電力(株)女川原子力発電所第1号機の新設

Construction of Unit No.1, Onagawa Nuclear Power Station

可児次郎(1) 藤田京(2) 滝口孝夫(3)

東北電力(株)女川原子力発電所第1号機は、電気出力52.4
MWのBWR発電所であり、昭和59年6月1日営業運転を開始
した。この設備はBWR-4型で、MAK-1型炉心炉心管を採
用しているが、各種の最新技術を積極的に導入し、高度な主
による改良・開発の成果の大部分をとり入れた新設備である。当
はこれを一貫して、54年12月に着工以来、設置地を速め
14.6. 7. 8. 月の間に完成させた。東北電力の原子力
発電所が建設し、東北電力の新たな電力に貢献した。これは大
いものである。

なお、このたびの女川1号機を、当社が手がけて完成し
たBWRプラントは10基、合計出力7,274MWに達し、
さらに4,400MWを建設中である。

The unit No.1 of Onagawa Nuclear Power Station, Tokoku E
lectric Power Co., Inc. is a BWR-4 type 52.4-MW BWR plant. I
beha get the order for the design and construction as a
main contractor, and completed the plant to be put into com
ercial operation last June.

Description is made here of the highlights in its design
construction and fuel operation.

It will be of note that: (1) this unit is a BWR-4 type h
oused in a BWR-1 type primary containment vessel; (2) the
construction work has been completed in only 48 months si
nce the start in December 1979; and (3) the number of BWR u
his completed by Toshiba has reached to 10 with this unit
isolating up to 7,274 MW in output capacity.

1. まえがき

図 6.3: 表示インターフェースの画面構成

本システムにおける文書提示例を、図 6.4 に示す。

図 6.4(a) は検索結果の提示の初期画面である。図に示すように、右側のウィンドウに検索した文書の原文を、左側のウィンドウにその文書から自動的に生成した抄録を提示している。抄録の提示においては、文書のタイトル、著者名、原文のアブストラクト部分の抄録、各節の見出しを表示している。初期画面で提示される原文アブストラクト部分の抄録では、5 文が 1 文に抄録されている。

抄録提示のウィンドウにおいて、任意の節をマウスで指定して、その節をより詳しく読むことが可能である。図 6.4(b) と (c) に抄録表示ウィンドウだけを切り出している。

図 6.4(b) は、アブストラクト部分を指定して、さらに詳しく表示させた場合である。また、(c) は、4 節を指定して詳しく表示させた場合である。後者の場合、原文 13 文に対して 5 文からなる抄録を提示している。

抄 録		原 文	
詳しく () 簡単に () O K () N O () 表示項目 () 絞り込 () ファイル名: 3911-002			
東北電力(株) 女川原子力発電所第1号機の建設 可児次郎(1) 藤田京(2) 滝口幸夫(3) 東北電力(株) 女川原子力発電所第1号機は、電気出力52.4 MWeのBWR発電所であり、昭和59年6月1日営業運転を開始した。 1. まえがき 2. プラント仕様と主な特徴 3. 建設工事の特徴 3.1 原子炉格納容器(PCV)据付工事 3.2 (イ) 項使用耐震性の合理化 3.3 保守性の向上とプラント総点検 (a) RPVI 次水圧テスト時(58年2月) (b) 燃料装荷前(58年9月) (c) 営業運転開始前(59年5月) 3.4 クリーンプラントの建設 3.5 計画外プラント停止の回避 3.6 48.5か月短縮月工事の達成 4. 試験運転の概要 5. あとがき		東北電力(株)が将来の電源開発の多様化、特に脱石油電源の推進をテーマにかかげて女川地点に原子力発電所の建設を決定され、当社と500MWe級BWRの共同研究を開始されたのは昭和43年であった。わが国初めての商用軽水炉の建設が始まってまもなくのまさに原子力あけぼのの時代であった。 45年政府の電源開発促進審議会の決定を受けて、直ちに東北電力(株)の要請に基づき当社は女川原子力発電所第1号機(以下女川1号機と略す)の基本設計を開始し、同年12月に同原子炉の設置が許可された。 40年代後半から50年代初めにかけて、欧米諸国をはじめ、わが国の先行の軽水炉は幾多の故障を経験した。右の技術的な知識を超えるトラブルが発生し、プラントの稼働率は著しく低下した。また米国スリーマイル島PWR発電所の事故は、BWR発電所の建設に携わるわれわれにとっても課題に受けとめるべき重要な問題を提起した。 当社は、これらの問題を解決するためにこれまでに培った先行機的设计・建設の経験をもとに、全社を挙げて当社自主技術の確立に努めた。一方では通産省の主導により、わが国独自の軽水炉技術の定着を目ざした改良標準化計画も50年以来着々と進められた。 このような背景のもとに、女川1号機に対し当社は新技術を積極的に提案し数多くの改善工夫が盛り込まれた。最新の技術を集約した女川1号機は54年12月本格着工し、59年6月に運転した。この建設・運転に至る経緯のあらましと、当社のBWR開発の変遷を図1に示す。女川1号機は着工以来、精力的に建設が進められた結果、48.5か月の短工期で竣工することができた。その間建設工事、試験運転あるいはプラント管理においてもそれぞれ工夫が施されたので、それらの概要につきここに紹介する。 2. プラント仕様と主な特徴 女川1号機的设计・建設では当初の基本設計を固めた後、着工までに多くの年月を経過した。この期間を生かして稼働中のプラ	

(a)

抄 録		原 文	
詳しく () 簡単に () O K () N O () 表示項目 () 絞り込 () ファイル名: 3911-002			
東北電力(株) 女川原子力発電所第1号機の建設 可児次郎(1) 藤田京(2) 滝口幸夫(3) 東北電力(株) 女川原子力発電所第1号機は、電気出力52.4 MWeのBWR発電所であり、昭和59年6月1日営業運転を開始した。 1. まえがき 2. プラント仕様と主な特徴 3. 建設工事の特徴 3.1 原子炉格納容器(PCV)据付工事 3.2 (イ) 項使用耐震性の合理化 3.3 保守性の向上とプラント総点検 (a) RPVI 次水圧テスト時(58年2月) (b) 燃料装荷前(58年9月) (c) 営業運転開始前(59年5月) 3.4 クリーンプラントの建設 3.5 計画外プラント停止の回避 3.6 48.5か月短縮月工事の達成 4. 試験運転の概要 5. あとがき		東北電力(株)が将来の電源開発の多様化、特に脱石油電源の推進をテーマにかかげて女川地点に原子力発電所の建設を決定され、当社と500MWe級BWRの共同研究を開始されたのは昭和43年であった。わが国初めての商用軽水炉の建設が始まってまもなくのまさに原子力あけぼのの時代であった。 45年政府の電源開発促進審議会の決定を受けて、直ちに東北電力(株)の要請に基づき当社は女川原子力発電所第1号機(以下女川1号機と略す)の基本設計を開始し、同年12月に同原子炉の設置が許可された。 このような背景のもとに、女川1号機に対し当社は新技術を積極的に提案し数多くの改善工夫が盛り込まれた。最新の技術を集約した女川1号機は54年12月本格着工し、59年6月に運転した。この建設・運転に至る経緯のあらましと、当社のBWR開発の変遷を図1に示す。女川1号機は着工以来、精力的に建設が進められた結果、48.5か月の短工期で竣工することができた。その間建設工事、試験運転あるいはプラント管理においてもそれぞれ工夫が施されたので、それらの概要につきここに紹介する。 2. プラント仕様と主な特徴 女川1号機的设计・建設では当初の基本設計を固めた後、着工までに多くの年月を経過した。この期間を生かして稼働中のプラ	

(b)

抄 録		原 文	
詳しく () 簡単に () O K () N O () 表示項目 () 絞り込 () ファイル名: 3911-002			
東北電力(株) 女川原子力発電所第1号機の建設 可児次郎(1) 藤田京(2) 滝口幸夫(3) 東北電力(株) 女川原子力発電所第1号機は、電気出力52.4 MWeのBWR発電所であり、昭和59年6月1日営業運転を開始した。 1. まえがき 2. プラント仕様と主な特徴 3. 建設工事の特徴 3.1 原子炉格納容器(PCV)据付工事 3.2 (イ) 項使用耐震性の合理化 3.3 保守性の向上とプラント総点検 (a) RPVI 次水圧テスト時(58年2月) (b) 燃料装荷前(58年9月) (c) 営業運転開始前(59年5月) 3.4 クリーンプラントの建設 3.5 計画外プラント停止の回避 3.6 48.5か月短縮月工事の達成 4. 試験運転の概要 5. あとがき		東北電力(株)が将来の電源開発の多様化、特に脱石油電源の推進をテーマにかかげて女川地点に原子力発電所の建設を決定され、当社と500MWe級BWRの共同研究を開始されたのは昭和43年であった。わが国初めての商用軽水炉の建設が始まってまもなくのまさに原子力あけぼのの時代であった。 45年政府の電源開発促進審議会の決定を受けて、直ちに東北電力(株)の要請に基づき当社は女川原子力発電所第1号機(以下女川1号機と略す)の基本設計を開始し、同年12月に同原子炉の設置が許可された。 このような背景のもとに、女川1号機に対し当社は新技術を積極的に提案し数多くの改善工夫が盛り込まれた。最新の技術を集約した女川1号機は54年12月本格着工し、59年6月に運転した。この建設・運転に至る経緯のあらましと、当社のBWR開発の変遷を図1に示す。女川1号機は着工以来、精力的に建設が進められた結果、48.5か月の短工期で竣工することができた。その間建設工事、試験運転あるいはプラント管理においてもそれぞれ工夫が施されたので、それらの概要につきここに紹介する。 2. プラント仕様と主な特徴 女川1号機的设计・建設では当初の基本設計を固めた後、着工までに多くの年月を経過した。この期間を生かして稼働中のプラ	

(c)

図 6.4: 動作例

6.5 全文検索との比較評価

本節では、文の意味役割抽出の精度ならびに、意味役割を利用することが検索結果の絞り込みやランキング付けに有効であることを確認する。

6.5.1 方法

技術論文誌である東芝レビューを対象とし、テストデータで文の意味役割抽出規則を記述・評価し、この規則を用いて評価データの検索を行い、検索結果を主題検索[23]の観点から評価する。論文中の図や表などを除いたテキストを検索対象とする。

(1) テストデータを対象とした意味役割抽出規則の記述と評価

テストデータを対象として、文の意味役割抽出規則を記述し、評価する。意味役割として、「話題」、「背景」、「従来の問題」、「目的」、「特徴」、「課題」、「結果」、「結論」の8種類を用いる。ここで、「従来の問題」とは従来技術や方法の問題を述べた文の種類である。

(1.1) テストデータ中の文について人手で意味役割を判定し、その結果を格納する。ここでは、この結果を格納したファイルを正解ファイルとよぶ。

(1.2) 人手で判定した意味役割を抽出するための抽出規則を記述する。この抽出規則を用いて抽出した結果と正解ファイルとを照合して、文の意味役割抽出規則を修正する。

(1.3) 抽出規則を用いて抽出した結果と正解ファイルとを照合して、文の意味役割抽出規則を評価する。

(2) 評価データを用いた意味役割抽出の評価

評価データを用いて検索を行う。検索命令中の検索語の同義語や類似表現などへの展開は行わない。

(2.1) テストデータを対象として記述した抽出規則を用いて評価データ中の文の意味役割を抽出し格納する。

(2.2) 検索命令を設定し、8種類の意味役割と、意味役割なし(全文検索)で検索する。

(2.3) 検索されたすべての文書について検索命令に適合するか否か、すなわち検索命令中の検索語が文書の主題と関係があるか否かを人手で判定する。ここでは、検索命令に適合する文書を適合文書とよぶ。

(2.4) 複数の検索命令について(2.2)と(2.3)を繰り返し、集計する。

6.5.2 結果

(1) テストデータを対象とした意味役割抽出規則の記述と評価

20文書、1549文(平均77文/文書)からなるテストデータに対し、話題、背景、

従来の問題、目的、特徴、課題、結果、結論の8種類の意味役割について、290の情報抽出規則と22の意味役割投射規則を記述した。

意味役割の再現率は平均82%であり、適合率は平均72%であった。表6.2に、文の意味役割ごとの再現率(%), 適合率(%), および文書全体の文の数に対する図6.5のa、b、cの範囲の文の数の割合を示す。

表 6.2: 文の意味役割ごとの再現率と適合率

意味役割	再現率	適合率	F	G	H
話題	100	64	0.0	3.3	2.3
目的	93	67	3.8	2.4	1.5
特徴	62	82	6.8	2.7	1.3
結果	72	81	5.8	7.8	2.5
結論	95	82	1.6	5.1	3.7
課題	67	50	2.2	1.3	1.9
背景	76	83	3.6	9.6	3.5
問題	90	70	1.2	1.8	1.9
平均	82	72			

表6.2中の意味役割「問題」は「従来の問題」であることを示す。表6.2から次のことがわかる。

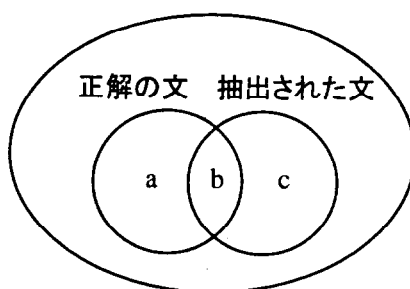


図 6.5: 一つの文書の中の文の分布

- 最も多くのファイルから抽出された意味役割は「背景」であり、すべての文書から抽出された。最も少ないファイルから抽出された意味役割は「課題」であった。
- 図6.5のa、b、cの範囲の文の数は、平均して文書全体の文の数の1.2%から9.3%であり、数文であった。

文の意味役割を抽出できなかったものには、次のようなものがあった。この説明のため、文章例を図6.6、6.7、6.8に示す。なお、説明のため、例で取り上げた文章中に数字を付与した。

- (a) 例1の第2章の第1文から第5文では背景の情報が述べられているが、第4, 5文から意味役割が抽出できなかった。これは、修辞構造解析では接続詞などの文書に依存しない言語情報だけを使った処理であるためである。「高信頼性」、「高品質」の関連性や「InGaAIP」の語の重複などを解析に利用する必要がある。

2. 結晶材料と素子構造

¹可視光領域で発振する半導体レーザを実現しようとする場合、結晶材料の選択が重要である。²現在、主流として用いられている GaAlAs 系材料では、バンドギャップが小さく $0.72\mu\text{Å}$ までは短波長化の限界とされているので、新しい材料を開発することが必要となる。

³半導体レーザ用の材料は、結晶欠陥が少なく、高信頼性をもつことが要求される。⁴高品質な基板結晶が得られる GaAs に格子整合がとれ、かつ短波長化に有利な特長をもつ材料として InGaAIP がある⁽¹⁾。

⁵InGaAIP 四元混晶は、窒化物を除いた III-V 族化合物半導体結晶の中で、もっとも大きな直接遷移形バンドギャップをもっており、短波長化にもっとも有利な結晶材料である。

..... (途中省略)

3. 素子特性

¹図3に、TOLD9200の典型的な電流-光出力特性とその温度依存性を示す。²周囲温度25℃での連続(CW:Continuous Wave)動作におけるしきい値電流は標準で76mAである。³また、スロープ効率は約0.5 mW/mAである。⁴周囲温度5℃から40℃の範囲でキックのない特性が得られている。

⁵図4に、光出力2mWでのレーザ光の遠視野像を示す。⁶ガウス分布形状をした基本横モード発振を示している。⁷ビーム広がり角の典型的な値は、接合に垂直方向で34°, 平行方向で7°である。⁸周囲温度(5~40℃)および光出力(1~3mW)の値を変えても、広がり角には変化はまったく見られない。

..... (以下省略)

図 6.6: 文章例 1

(“ $0.6\mu\text{Å}$ 帯赤色半導体レーザ TOLD9200”, 本館他, 東芝レビュー Vol.43, No.5, 1988 (説明のため上付数字を付与))

- (b) 例1の第3章の第2, 3, 4, 7, 8文はこの文書の実験結果に関することと考えられるが、意味役割を抽出できなかった。ここでは、それらに先行する第1, 5文は参照すべき図を示す文となっている。したがって、一文の中に手掛かりとなる表現を含まない場合でも、先行する文の情報を用いて意味役割を抽出するように修辞構造解析を拡張すれば対応可能である。また、図や表を示して述べる内容が文書の主題と密接な関係があるかという問題があるが、例えば当文書の第1文が文書のタイトル中の語を含んでいる。このような情報の利用も有効と考えられる。
- (c) 例2では第1, 2, 3文が背景の情報について述べているが(第2, 3文は省略)、意味役割を抽出できなかった。この文書では、第4文に「以上のような標準化経緯からわかるとおり」という表現を含んでいる。修辞構造解析でこの情報を用いれば、対応可能であると考えられる。また、この文書では、(b)の場合に反してタイトル中にある「ISDN」は本文の主題

とは関係が薄い。タイトル中の語と文書の主題との関連の解析も必要である。

2.ISDNの標準化経緯とその特長

¹ISDNは、電話網のデジタル化をねらいとして1970年代初期から検討が開始され、1980年代初期にはデジタル伝送技術とデジタル交換技術を統合して、一つのデジタル回線上でさまざまな種類・型式の情報を伝達するネットワークとしての姿が明確になった。

.....(途中省略).....

⁴以上のように標準化経緯からわかるとおり、ISDNは電話サービス網のデジタル化を基盤としているが、単に電話サービスだけでなく、データ通信や音楽・文書・ファクシミリ・ビデオテックス・ビデオなど非電話系の広範囲の通信サービスを一つのネットワークから提供することを主なねらいとしている。

.....(以下省略).....

図 6.7: 文章例 2

(“ISDNによる企業情報通信システム”, 谷野他, 東芝レビュー Vol.44, No.4, 1989)

- (d) 例3は最終章であり第1文に「完了した」という表現を含んでいる。しかし、当文書の結論ではなく、その分野での技術的な動向、つまり背景情報を述べているにすぎない。本論での表層的な表現だけの処理では、このような文を背景情報であると判断することはできない。

(2) 評価データを用いた意味役割抽出の評価

547文書からなる評価データから文の意味役割を抽出格納した結果について、表6.3に示す6種類の検索命令で検索を行った。検索命令中のすべての検索語をある意味役割の文に存在する文書が当該の意味役割での検索結果の文書となる。各検索命令での検索結果は付録に示す。各検索結果の適合率と再現率を表6.4と表6.5に示す。なお、再現率としては、本来検索対象のすべての文書中の適合文書の数に対する検索結果中の適合文書の数之比を示すべきであるが、ここでは文の意味役割を区別しない全文検索結果の中の適合文書の数との比を示している。表6.4と表6.5において、「その他」とはいずれの意味役割も抽出されなかった文にのみ検索語を含む場合である。「なし」とは文の意味役割を区別しない全文検索の場合である。表6.4中の「-」は検索された文書がなく、適合率が得られなかったことを示す。

4. あとがき

¹CT画像のS/Nや空間分解能力の向上は、その要求にしたがって相応のシステムを実現するための技術的基盤固めを完了した。²ただし、CTがより一層の飛躍をするためには透過能力、空間分解能、S/Nなど主要な画像性能を満たしたうえで、スキャン速度を向上させる必要がある。

³人体用のMRI(Magnetic Resonance Imaging)のように、1回のスキャンで数十断面のデータが得られるようなシステムが前記図7の構想では必要である。⁴20000シリーズでは1断面/スキャンで、その所要時間は約3分である。

.....(以下省略).....

図 6.8: 文章例 3

(“高品位画像産業用CTスキャナTOSCANER-20000シリーズ”, 庄司他, 東芝レビュー, Vol.44, No.5, 1989)

表 6.3: 検索命令

番号	検索命令
1	ロボット AND 点検
2	ロボット AND 画像処理
3	画像 AND 診断
4	画像 AND 伝送
5	衛星 AND 通信
6	情報 AND 検索

表 6.4: 意味役割別の適合率

検索命令	話題	目的	特徴	結果	結論	課題	背景	問題	その他	なし
1	67	100	100	100	50	—	33	—	0	22
2	100	100	100	100	100	—	50	0	0	40
3	100	—	100	100	100	100	92	100	33	57
4	50	—	—	0	67	—	33	—	4	16
5	67	100	100	67	90	—	53	—	8	39
6	100	0	0	25	50	—	75	0	2	10
平均	81	75	80	65	76	100	56	33	8	31

表 6.5: 意味役割別の再現率

検索命令	話題	目的	特徴	結果	結論	課題	背景	問題	その他	なし
1	100	100	50	100	50	0	50	0	0	100
2	50	25	25	50	50	0	25	0	0	100
3	15	0	5	40	5	5	55	15	35	100
4	17	0	0	0	33	0	50	0	17	100
5	31	15	8	31	69	0	69	0	8	100
6	17	0	0	33	17	0	50	0	17	100
平均	38	23	15	42	37	1	50	3	13	100

表 6.4 と表 6.5 から、各意味役割および「その他」の適合率と再現率について次のことがわかった。

- 「話題」、「目的」、「特徴」、「結果」、「結論」、「課題」の適合率は 65% であった。
- 文書の主題と深い関係が薄いと考えれる「背景」と「従来の問題」からも 56% と 33% の適合率であり、他の意味役割に比べて低かった。
- 文の意味役割が抽出できなかった「その他」の適合率は、意味役割の場合と比べ 8% と非常に低かった。
- 意味役割の中では「背景」が最も高く 50% であった。

一つの文が複数の意味役割をもつ場合があることと、表 6.4 の結果から、意味役割を「課題」、「話題」、「特徴」、「結論」、「目的」、「結果」の A グループと「背景」、「従来の問題」の B グループに分けて次の場合の適合率を調べた。図 6.9 に検索された文書の分布を示す。図 6.10 の中の円は、図 6.9 の中の円に対応する。

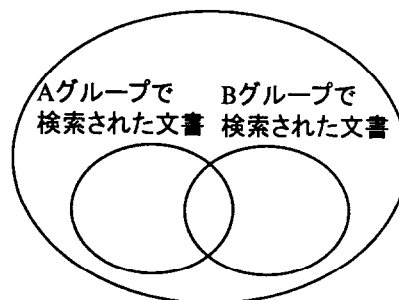


図 6.9: 検索された文書の分布

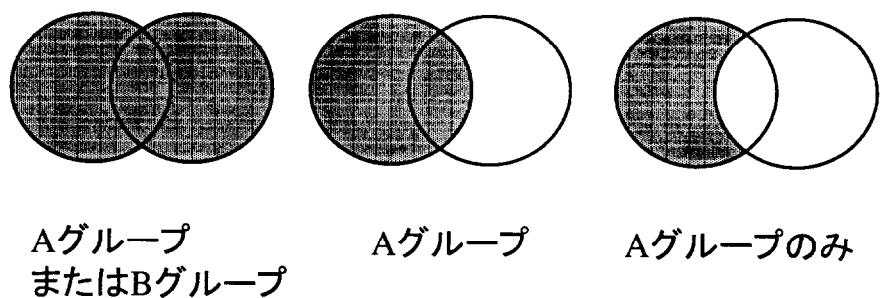


図 6.10: 意味役割のグループ分け

- (a) AグループまたはBグループのいずれかの意味役割の文で検索された場合 (図 6.10 の (a) に対応)
- (b) Aグループの意味役割の文で検索された場合 (図 6.10 の (b) に対応)
- (c) Aグループのみの意味役割の文で検索された場合 (図 6.10 の (c) に対応)

表 6.6 に全文検索の場合の検索文書、適合文書および適合率を示す。表 6.7、表 6.8、表 6.9 に、(a)、(b)、(c) の場合の検索された文書と適合文書の数、適合率と、全文検索の適合率との比、および再現率を示す。

表 6.6: 全文検索における適合率

検索命令	検索された文書数	検索された文章中の適合文書数	適合率
1	9	2	22
2	10	4	40
3	35	20	57
4	38	6	16
5	33	13	39
6	63	6	10
平均	—	—	31

表 6.7: (a) の場合の適合率と再現率

#ED: 検索された文書の数

#RD: 検索された文書の中の適合文書の数

検索命令	#ED	#RD	適合率	再現率	全文検索に対する適合率の比
1	4	2	50	100	2.3
2	6	4	67	100	1.7
3	14	13	93	65	1.7
4	13	5	38	83	2.4
5	21	12	57	92	1.5
6	15	5	33	83	3.7
平均	—	—	56	87	2.2

表 6.7, 表 6.8, 表 6.9 から次のことがわかった。

- (a), (b), (c) の順で適合率が向上し、検索文書総数比が減少した。ただし、(b) と (c) の適合率の差は大きくない
- (a), (b), (c) の順で再現率が低下した。(b) と (c) の再現率の差は大きい。

以上から、本方式による検索文書の絞り込みの有効性を確認した。

表 6.8: (b) の場合の適合率と再現率
 #ED: 検索された文書の数
 #RD: 検索された文書の中の適合文書の数

検索命令	#ED	#RD	適合率	再現率	全文検索に対する適合率の比
1	3	2	67	100	3.0
2	4	100	100	2.5	2.5
3	9	9	100	45	1.8
4	7	3	43	50	2.7
5	13	10	77	77	2.1
6	13	4	31	67	3.4
平均	—	—	70	73	2.6

表 6.9: (c) の場合の適合率と再現率
 #ED: 検索された文書の数
 #RD: 検索された文書の中の適合文書の数

検索命令	#ED	#RD	適合率	再現率	全文検索に対する適合率の比
1	1	1	100	50	4.5
2	3	3	100	75	2.5
3	2	2	100	10	1.8
4	4	2	50	33	3.1
5	4	3	75	23	2.0
6	10	2	20	33	2.2
平均	—	—	74	37	2.7

6.5.3 考察

(1) 意味役割抽出規則の記述について

情報抽出規則の記述にあたっては、意味役割の再現率を高くするような規則は不適切な文にもマッチする危険性が高く、逆に個別的な規則では抽出の網羅性を期待できない。表層的な表現を手掛かりとする本方法では、その傾向は著しい。意味役割の抽出精度の向上のためには、深層的な情報や前後の文の情報を解析し利用することが必要である。

(2) 意味役割抽出の検索における効果について

本方法では従来の全文検索方式に比べて高い適合率を得られることから、特に適合率を重視する検索の場合において、検索作業の効率化に寄与できると考えられる。

一方、本方法により得られる適合文書の数、全文検索で得られる適合文書数の平均 88% 以下であるという問題がある。また、本方式では、文からの意味役割の抽出において、その精度は形態素解析および構文・意味解析の精度に依存する。検索システムとしては、形態素解析を用いない文字レベルの検索、形態素解析を用いた単語レベルの検索などと、本方法とを組合わせて、本方式を後処理・絞り込み機能として用いることにより、システム全体では漏れがないようにすることが必要であり、これは実現可能である。

同様に、検索された文書を、検索命令に最も適合する文書から提示するランキングにおいて、検索文書の適合度の尺度の一つとして用いることが有効であると考えられる。

6.5.4 結果

文の意味役割解析に基づく全文検索の実験と評価を行った。技術論文 547 文書を対象とした 6 種類の検索命令での検索において、従来の全文検索方式で適合率が平均 31% であるのに対し、「課題」、「話題」、「特徴」、「結論」、「目的」、「結果」、「背景」、「従来の問題」の 8 種類の意味役割を利用することにより適合率が平均 56% を得た(再現率は前者に対して 87%)。このとき、検索文書の総数は前者に対し 44% であった。さらに、「背景」、「従来の問題」を除く 6 種類の意味役割では、再現率が 73% に下ルものの適合率は 70% に上昇し、検索文書総数の比は 29% に減少できた。本方式が、検索結果の絞り込みに有効であることを確認した。

6.6 意味役割に基づく検索での頻度情報の利用と評価

語句の文書中の使用頻度などの頻度情報に基づいて、語句の重要度を判定し、検索結果における文書のランキングを行う手法が存在する。ここでは、文書構造の意味解析に基づく検索手法と、頻度情報に基づく検索手法を組み合わせることにより、検索精度が向上することを示す。

6.6.1 方法

以下の2つのそれぞれの尺度を考える。

1. 文書内・文書間頻度

語の重要度を定める尺度としてよく用いられる (例えば、文献 [87])。少数の文書によく用いられる語を重要な語とする尺度である。ここでは、ある語 t のある文書 d に対する文書間頻度 idf 、文書内頻度 tf を、それぞれ式 (6.1)、(6.2) と定義する。

$$idf_t = \frac{\log \frac{df}{idf}}{\log df} \quad (6.1)$$

tdf : 語 t を含む文書数

df : 文書データベース全体の文書数

$$ntf_{td} = \frac{tf_{td}}{tf_{t_{max}d}} \quad (6.2)$$

tf_{td} : 文書 d 中の語 t の頻度

$tf_{t_{max}d}$: 文書 d 中で最も頻度の高い語 t_{max} の頻度

上記定義に基づき、文書 d 中の語 t に対する重要度評価関数 $eval_1(t, d)$ を、以下のように定義する。

$$eval_1(t, d) = ntf_{td} \cdot idf_t \quad (6.3)$$

2. 正規化頻度

式 (6.2) で定義した正規化頻度を重要度評価尺度とし、 $eval_2(t, d)$ を以下のように定義する。

$$eval_2(t, d) = ntf_{td} \quad (6.4)$$

上記評価関数 $eval_2(t, d)$ では、語の文書間の分布を考慮しないことになる。

N 個の語からなる検索語列 q で検索が行われるとして、 q に対して文書 d の適合度の評価を、以下の式で行うものとする。

$$E(q, d) = \sum_{i=0}^N eval_k(t_i, d) \quad (t_i \in q, 0 \leq i \leq N) \quad (6.5)$$

式 (6.5) において、 k が 1 の場合、文書内・文書間頻度による評価、2 の場合、正規化頻度による評価に相当する。

具体的な検索は、以下の通りである。

1. 検索語列 q に含まれるすべての語を含む文書を、全文検索によって検索する。
2. 検索された文書の適合度を式 (6.5) により求め、降順にソートする。

検索手法の精度を評価する場合、検索要求に適合した文書 (以下、適合文書と呼ぶ) の判定条件にその優劣が大きく依存する。判定条件を明確にしておく必要がある。意味役割に基づく検索では、意味役割を指定して検索することができるが、他の手法 (例えば語の頻度情報に基づく手法) では、このようなことは不可能である。比較のため、何らかの点で共通な視点で評価する必要がある。

本節では、検索で指定した語 (以下、検索語と呼ぶ) が主題となっていることで評価するものとする。意味役割に基づく検索では、中心的な内容に関わる意味役割としては、「話題」、「目的」、「特徴」、「結論」などがあげられる。そこで、意味役割に基づく検索では、検索語が「話題」、「目的」、「特徴」、「結論」のいずれかの意味役割の文に含まれる文書を取り出し、それらの文書を式 (6.5) を用いて適合度を求め、ソートする。

一方、比較のため、意味役割による制限を行わずに全文を対象にした検索結果の式 (6.5) によるランキングについても行った。なお実験において、同義語や類似表現などへの展開は考えないものとした。

上記の検索方式について、上記の適合文書にしたがい、6つの検索語ペア (例えば「CCD/イメージセンサ」) を主題とする文書を検索し、適合率、再現率を評価した (東芝レビュー 386 文書)。

6.6.2 結果

意味役割に基づいて検索した結果と、全文検索の検索結果との比較を、表 6.10 に示す (語の頻度情報によるランキングは行っていない)。

全文検索では、検索語を含む文書すべてが取り出されるため、再現率は 100% となるが、適合率は 36% (6つの検索についての平均) にすぎない。一方、意味役割に基づき、文書の適合性を判定した場合、検索もれが生じ、再現率が低下 (平均 86%) するものの、適合率は平均 66% に向上することがわかる。

語の頻度に基づく検索では、式 (6.5) に示す評価値にしたがって、検索した文書をソートし、ある閾値より高い評価値となる文書を適合文書とする。したがって、閾値を適宜設定することにより、適合率、再現率を変化させることが可能である。

ここでは、意味役割に基づく検索と比較のため、再現率がほぼ同じになる閾値を最適な閾値とし、それぞれの検索手法についての適合率を求めた。この結果を表 6.11 に示す。

なお、表 6.11 の全文検索についての適合率、再現率は、以下のように求めた。

1. 全文検索結果の文書のファイル名をアルファベット順にソートする。
2. 1 の文書を始めから順に一定の比率 G で取り出したものを検索結果とする。
3. 6つの検索条件についての再現率の平均が意味役割に基づく検索と同じになるよう G を調整する。

表 6.11 より、検索手法の適合率は、意味役割と正規化頻度の組合せた場合が最も高く、以下、意味役割、正規化頻度、文書内・文書間頻度、全文検索の順であることがわかる。この結果は、意味役割の情報と統計的な情報を組み合わせることによ

表 6.10: 意味役割に基づく検索と全文検索の精度の比較

番号	C	意味役割				全文検索			
		A	B	P	R	A	B	P	R
No.1	5	5	7	71	100	5	10	50	100
No.2	3	3	5	60	100	3	13	23	100
No.3	3	3	4	75	100	3	13	23	100
No.4	12	6	7	86	50	12	26	46	100
No.5	3	2	12	17	67	3	20	15	100
No.6	6	6	7	86	100	6	10	60	100
平均	-	-	-	66	86	-	-	36	100

A: 検索文書中の適合文書数

B: 検索文書数

C: 全文書中の適合文書数

P: 適合率 (A/B) (%)

R: 再現率 (A/C) (%)

り、検索の精度を向上させることができることを示している。

6.7 動的抄録提示インタフェースの評価

システムの動的抄録提示について評価を行う。動的抄録提示インタフェースが文書の主題に関する要不要 (relevance) の判定に与える影響を調べることを目的とする。抄録の提示により要不要の判定時間が短縮される一方で、判定を行った時点での抄録が重要な文を含まなかったり、重要でない文を含んでいたりするために誤判定が起こる可能性があるので、判定時間と共に判定の正確さをも測定することにより本インタフェースの有効性を検証する。

6.7.1 方法

動的抄録と原文とを並べて提示する本インタフェースの効果を、原文のみを提示するインタフェースの場合と比較する。このために、本インタフェース (図 6.11) に加えて、抄録画面が表示されない以外は全く同等の機能をもつ実験用インタフェースを用意した (図 6.12)。以後、前者を (a)、後者を (b) と呼ぶ。いずれの場合も、被験者は以下の手順で実験を行った。

- 表 6.12 のいずれかの与えられた問題を読む。
- 表 6.12 の検索語を検索語入力画面に入力し [82]、全文検索の結果を表示させる。
- 提示された文書群中の各文書が問題に該当するかどうかを、(a) 抄録と原文、あるいは (b) 原文のみを見ることにより判定する。

表 6.11: 各検索手法についての精度比較

番号	意味役割*1		$tf \cdot idf$		正規化頻度*2		*1 + *2		全文検索	
	P	R	P	R	P	R	P	R	P	R
No.1	71	100	62	100	57	80	83	100	62	80
No.2	60	100	67	67	67	67	60	100	30	100
No.3	75	100	30	100	27	100	75	100	30	100
No.4	86	50	60	75	65	92	83	41	36	58
No.5	17	67	43	100	43	100	20	67	20	100
No.6	86	100	80	67	100	67	100	100	71	83
平均	66	86	57	85	60	84	70	85	41	87

P : 適合率 (%)

R : 再現率 (%)

- 該当するかどうかの判定に要した時間 (文書群中の各文書の判定に要した時間の総和) と、
判定再現率 = (被験者が該当するとした正解文書数) / (正解文書数)

表 6.12: インタフェース評価に用いた問題

検索対象：朝日新聞社説			
文書群	正解文書数/文書数	検索語	問題「Xについて主に述べたもの」
A	5/10	雇用	X=雇用のあるべき姿
B	4/16	円高	X=円高対策
C	7/11	核	X=核保有問題
D	7/17	選挙区制	X=選挙区制改革
E	5/21	憲法	X=憲法擁護を唱える政党や政治家
F	2/22	水	X=水質汚濁
G	5/17	戦争	X=戦争責任
H	2/16	欧州	X=欧州統合問題

判定適合率 = (被験者が該当するとした正解文書数) / (被験者が該当するとした文書数)

を記録する。ここで、正解文書とは問題を用意した筆者自身が問題に該当すると判定したものである。

例えば、表 6.12 の文書群 A は、「雇用」という検索語を含む文書 10 件からなる。検索システムにこの検索語を入力するとこれらの文書が検索されるので、この中から「雇用のあるべき姿について主に述べたもの」を被験者に選ばせる。

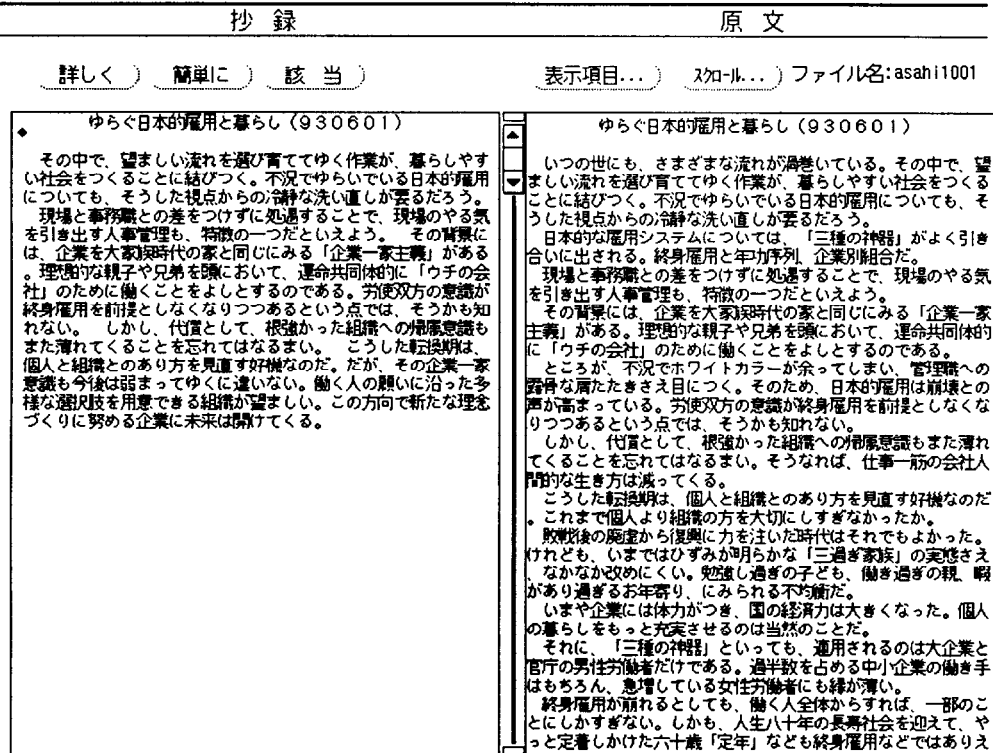


図 6.11: (a) 抄録と原文を提示するインタフェース

表 6.13: 実験の条件

	(a) 抄録+原文	(b) 原文のみ
文書群 A,B,C,D	被験者 (1)(2)(3)(4)	被験者 (5)(6)(7)(8)
文書群 E,F,G,H	被験者 (5)(6)(7)(8)	被験者 (1)(2)(3)(4)

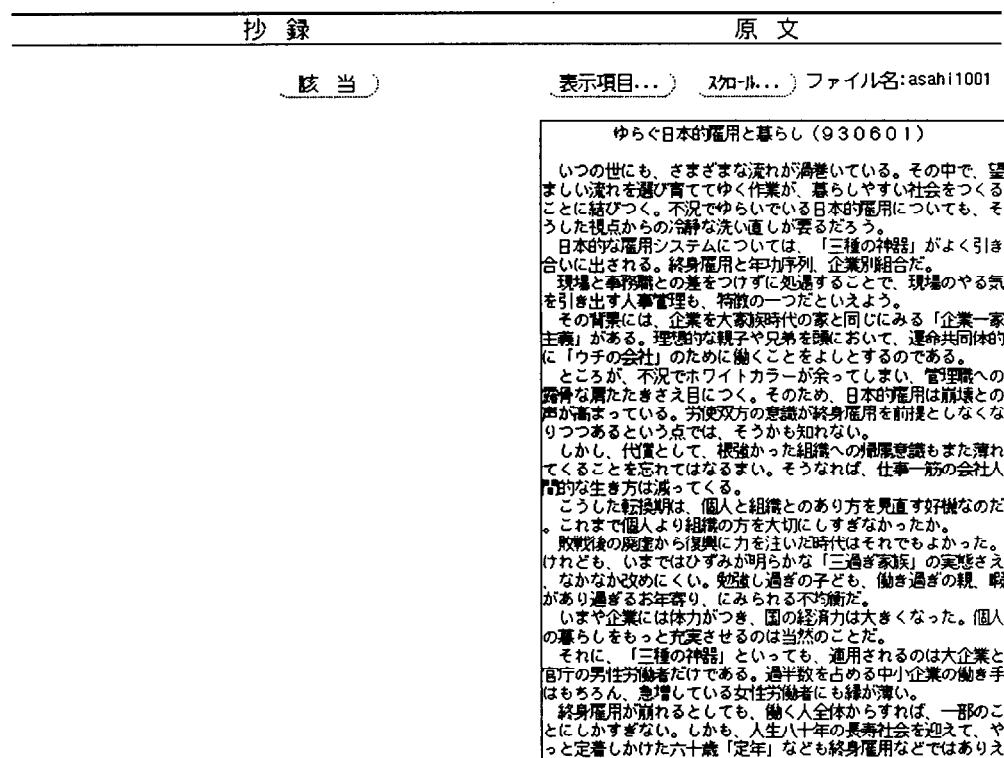


図 6.12: (b) 原文のみ提示するインタフェース

表 6.14: 実験結果

4 回の平均→	判定時間 (秒)			判定再現率 (%)			判定適合率 (%)		
文書群 ↓	(a)	(b)	(a)/(b)	(a)	(b)	(a)/(b)	(a)	(b)	(a)/(b)
A	331	525	63.0	80	80	100	100	96	104
B	506	663	76.3	56	88	64	79	95	83
C	275	358	76.8	71	71	100	100	100	100
D	441	610	72.3	64	71	90	92	94	98
E	636	656	97.0	25	50	50	45	62	73
F	290	351	82.6	100	100	100	75	71	106
G	432	435	99.3	80	75	107	89	100	89
H	383	407	94.1	100	100	100	92	92	100
	平均 82.7			平均 89			平均 94		

被験者は 8 人とし、表 6.13 のように、各被験者に (a) と (b) で同じ問題を与えないように配慮した。この表から、例えば被験者 (1) は文書群 A,B,C,D に対して抄録付きの実験を行い、文書群 E,F,G,H に対して抄録なしの実験を行ったことがわかる。また、(a) の場合、動的抄録の圧縮率の初期値は 30% とし、被験者に「詳しく」「簡単に」のボタンを押して抄録の長さを調節することを許した。

6.7.2 結果

表 6.14 に実験結果を示す。文書群 A,C,F,H については、動的抄録提示インタフェースにより、判定再現率及び適合率を損なうことなく判定時間が短縮されるという結果が得られている。一方、文書群 B,D,E,G については、判定時間は短縮されているものの判定再現率や適合率が若干犠牲にされている。

6.7.3 考察

文書群によって実験結果に差が出た原因としては、抄録の質の差、問題の難易度の差、個人差などが考えられる。今回は、個人差はないものとみなし、抄録の質の差及び問題の難易度の差の 2 つについて考察する。

6.7.3.1 抄録の質の差

文書群 A,C,F,H と文書群 B,D,E,G の実験結果に差をもたらした要因として第一に考えられるのは、抄録の質のばらつきである。すなわち、文書群 A,C,F,H の抄録の質に比べて文書群 B,D,E,G の抄録の質が劣っていたために、判定再現率や適合率に悪影響が出たという仮説を立てることができる。そこで、抄録の質と今回の結果の相関の有無を調べるために、今回実験に用いた文書の抄録の質を文書集合毎に評価

した。

以下に、抄録の質の評価の手順を示す。

- 画面にちょうど収まるくらいに抄録を「詳しく」し、この時の文数 N を数える。
- 正解作成者に原文を与え、重要文を N 文選ばせる。
- 上記の抄録文と重要文を比較し、(抄録文でありかつ重要文である文の数)/ N を算出する。

このようにして評価した抄録の質の、文書群毎の平均値を表 6.15 に示す。

表 6.15: 抄録の質についての評価

文書群	検索語	抄録の質平均 (%)
A	雇用	62.3
B	円高	50.1
C	核	53.9
D	選挙区制	52.7
E	憲法	57.1
F	水	52.5
G	戦争	49.9
H	欧州	49.4
		平均 53.5

なお、ここでの抄録の質は高いとは言い難いが、人手で作成した抄録の質を同じ方法により評価しても平均 6 割程度であることが報告されているので [67]、ある程度実用にたえるものであると考えられる。

文書群 A,C,F,H の抄録の質は平均 54.5%、文書群 B,D,E,G は 52.5% で、有意差は見られない。これより、抄録の質のばらつきが判定再現率及び適合率に直接影響を与えているわけではないことがわかる。

6.7.3.2 問題の難易度の差

文書群 A,C,F,H と文書群 B,D,E,G の判定再現率及び適合率に差が出たのは、問題の難易度のばらつきのせいである可能性がある。

判定再現率及び適合率を損なわずに判定時間を短縮することができた文書群のうち、C の適合率及び F,H の再現率は抄録なしの場合で既に 100% になっているので、これらは選定の比較的容易な問題であると考えられる。例えば、F の文書群には「水をさす」という表現を含む文書が含まれているが、このような文書を「水質汚濁について主に述べたもの」ではないとして排除することはユーザーにとって容易である。これに対し、例えば再現率の低下が著しかった E を見てみると、問題文が「憲法擁護を唱える政党や政治家について主に述べたもの」と他よりも複雑であり、選定の難易度が高い。このため、抄録なしの場合でも再現率は 50% という低い値をとっ

ている。Fのような問題とは違い、字面を眺めるだけではなく内容を把握することが要求される問題であると言える。このように考えると、字面を眺めるような浅いレベルで要不要の判定ができる平易な問題では判定再現率及び適合率に影響が出ないが、判定にある程度深い内容理解を必要とするような問題では抄録の質の影響が判定再現率及または適合率の低下という形で現れる傾向がある。

6.8 本章のまとめ

動的な自動抄録機能を持つ文書検索システムを開発した。検索結果の抄録を提示し、提示された抄録の任意の箇所を詳しく表示するという文書検索のための新しい文書提示インタフェースである。

また、文の意味役割に基づく検索方式をこの検索システム上に実現した。上記検索手法について、従来の検索手法(語の頻度などの統計情報に基づく検索)との比較評価を行った。評価実験は、複数の語を指定し、それらの語が主題となっている文書を検索するという観点で行い、従来方式に比して適合率が向上できることを確認した。

また、検索結果の文書を効率的にブラウズするための動的抄録提示インタフェースの評価実験を行った。文書の要不要の判定時間は、原文のみを提示した場合に比べ80%程度に短縮されることがわかった。この代償として、平均では判定再現率及び適合率が原文のみを提示した場合の90%程度に低下してしまう。しかし、内容の深い理解を必要としない比較的容易な文書の選定問題においては判定の質を保持することができた。

第7章 結論

7.1 本研究のまとめと成果

本研究では自然言語処理システムにおける文脈情報の利用について検討を行った。具体的には、自然言語処理における以下の3つの処理局面に関して検討した。

1. 応答生成
2. 文解析
3. 文書構造解析

応答文生成に関しては、質問応答システムにおいて自然な応答文を生成する際に必要となる要件を考察し、文脈情報に基づく応答生成の規則を提案した。

質問応答システムにおける応答生成に関して、従来、生成の戦略や方針という形式での整理は試みられていたが、具体的な規則という形式では定義されていなかった。本研究で提案した応答生成規則は、質問応答システムにおける応答生成時に、省略すべき項目を決定するための規則であり、この規則に基づくことにより適切な量でシステムの応答を生成することができる。この規則を考察するにあたり、会議室の予約を自然言語で行うタスクを例題として用いた。しかし、提案した規則は会議室予約に依存する内容を含んでいないため、その他の予約業務での質問応答にも適用可能であると考えられる。

文解析に関しては、質問応答システムにおいて、ユーザが入力する文を解析する際に生じる曖昧性解消のための尺度として、文解釈における情報量を定義し、その尺度を用いた文解釈の曖昧性を解消する手法を提案した。この手法は、因果関係についての知識を含む聞き手の理解モデルに基づいている。

本研究では、文解釈の過程で生じる曖昧性として、照応解析の問題を取り扱った。照応解析についての従来手法は、ヒューリスティックな方法が大半であるが、提案した手法は、聞き手の理解モデルから導き出した定量的な手法である。文の解釈時に曖昧性が生じる要因としては、文中に用いられる多義語等の存在もあげられる。複数の文解釈候補から選択する方式として、提案した手法は、これらの問題にも適用可能であると考えられる。

文書構造解析に関しては、接続詞を手掛かりにした文間の関係の解析と、内容のまとまりの抽出について検討した。そして、修辞関係を抽出するための辞書と、修辞関係に基づく構造選好規則等を整備し、それら知識辞書に基づいた文書構造解析処理を試作した。

本研究で提案した文書構造解析の手法は、文書の記述内容や分野に依存しない処理であり、文書を単位として取り扱う様々な自然言語応用システムに適用が可能で

あると考えられる。本研究では、この文書構造解析の具体的な応用として、抄録生成と文書検索への利用を検討した。

以上、応答生成、文解析、文書構造解析という自然言語処理における3つの処理局面に関する文脈処理についての検討を行った。さらに、上記の検討に基づいて、以下の応用システムを試作した。

1. 会議室予約システム

ユーザが自然言語で入力する会議室の予約状況の問い合わせや予約命令に対して、予約状況などの情報を自然言語で応答するシステム。

2. ビデオ操作法案内システム

ユーザが自然言語で入力するビデオの操作方法の問い合わせに対して、操作方法を自然言語で応答するシステム。

3. 抄録生成システム

文書を入力として、抄録文章を生成するシステム。論説文、技術解説文などの文書を処理対象とする。

4. 文書検索システム

検索結果の動的抄録提示と、文書構造に基づく語句の重要度判定を特長とする文書検索システム。

応答生成生成規則は、会議室予約システムを用いて評価実験を行った。被験者に会議室予約の課題を与え、実際にシステムを用いて会議室を予約する会話実験を行った。システムが生成する応答文の自然性についての、被験者による主観的な評価の結果、提案した規則を用いて応答文を生成することにより、自然性が向上することを確認し、規則の有効性が確かめられた。

文解釈の情報量に基づく曖昧性解消の処理については、ビデオ操作法案内システムを用いて評価実験を行った。被験者に、システムとユーザと間の対話対を被験者に示し、その後に入力する文を被験者に考えてもらうことにより、評価文を収集した。収集した評価文について文解析を行い、照応先の選択精度を評価した。この実験の結果、提案した処理を用いることにより、誤りの半分程度を減らすことができることが確認でき、提案する手法の有効性が確かめられた。

文書構造解析に関しては、文書構造に基づく抄録生成について評価実験を行った。抄録生成結果の評価では、人があらかじめ選んだ重要文をシステムが生成した抄録が捕捉しているかどうかについての評価と、所望の文書を探す際に原文と抄録のいずれを読むのが効率的であるかについての評価を行った。この結果、提案する文書構造解析に基づく抄録生成方式の有効性を確認することができた。

また、文書構造解析については、文書検索での適用し評価を行った。評価実験では、複数の語を検索語として、それらの語が主題となっている文書を検索するという観点で行い、従来方式(全文検索)に比して適合率が向上できることを確認した。また、統計的な手法との組合せにより、より精度の向上が図れることを確かめた。

上記の直接的な研究上の成果に加え、本研究で述べた方式を実運用システムに適用し、その実用性を確認している。

本研究で研究を行った文書構造解析処理をベースに、情報フィルタリングシステムを開発した。このシステムでは、文書構造としてタイトルや段落などの書式情報と文の意味役割を文書より抽出する。抽出される構造に基づいた検索条件記述にしたがって入力文書の類似度を算出し、入力文書を分類する。

新聞記事を対象とした情報フィルタリングの有料サービスを開始した。現在、このサービスへ加入するユーザ数は着実に伸びており、今後の発展が期待される。

さらに、WWW ページを対象とした情報フィルタリングサービスをインターネット上での新規サービスとして開始した。このサービスについても、今後アクセス数の増加などが期待される。

7.2 今後の課題

本節では、自然言語応用システムならびに、文脈処理の今後の課題を述べる。

7.2.1 質問応答システム

本研究で述べた質問応答システムのうち、ビデオ操作案内システムでは、利用者から与えられた問題を解決するために、問題解決を行うモジュールとその問題解決のための知識の作成が必要となった。

会議室予約システムのようなデータベースアクセスを中心とする質問応答システムと比較して、問題解決型の質問応答システムでは、この問題解決モジュールの開発と知識作成にコストがかかるため、システムを大規模化なものにする上で問題があった。

上記のような問題を解決する一つ的手段として、加工しない文書データを対象にすることが考えられる (ビデオ操作案内システムでは、入力文の解析やユーザのゴールの認識に失敗した時に、関連するマニュアルを表示することにより、この機能を実現しようとした)。つまり、蓄積されている文書データを提示することを中心に、ユーザが持つ問題を解決していくことになる。

問題解決のための文書データは、比較的低コストで収集することができる。このような方法を取ることで、規模の拡大が望める。この場合、文書中の情報を有効に利用するための情報抽出処理や、対話戦略を含めた知識処理との融合が課題となる。

7.2.2 自動抄録システム

本研究では、論説文や技術解説文を対象としてシステムならびに、文書構造解析のための知識辞書を構築した。論説文などの論旨の流れが意識される文書についての方式の有効性はある程度確かめることができたが、それ以外の文書 (例えば、新聞記事、随筆、物語など) に対する提案した方式の適用可能性は明らかではない。

例えば、新聞記事では、重要な内容は、記事の先頭に記述することが一般的であ

る。他の文書を対象にした文書構造解析と、抄録生成のための分析と、その分析に基づく知識辞書の構築が必要である。

7.2.3 文書検索システム

本研究では、文書構造に基づく検索手法について検討を行った。統計的な手法と組み合わせることにより、精度が向上する見通しが得られた。今後、さらに、統計的な手法との融合を進めるとともに、大規模なテキストデータを用いたより定量的な評価が必要である。

7.2.4 文脈処理

本研究では辞書や規則といった知識をあらかじめ人手で記述するアプローチにより、自然言語処理における文脈の解析や文脈情報の利用を検討した。知識記述によるアプローチは、処理に必要な知識が無駄なく作成できるという面があるが、一方で、大規模に作成しようとした場合に、一貫性や作成コストの面で問題が生じる。現在、対話コーパスなど文脈処理の研究のための土台が、整いつつある。従来の知識記述によるアプローチと、コーパスから得られる統計情報に基づくアプローチとの融合が必要である。

7.3 おわりに

自然言語処理システムにおける文脈情報の利用について研究を行った。自らの言語理解過程を内省すれば、対話や文書中の一つ一つの文は、その文が語られている文脈の中で理解していることは明らかである。したがって、自然言語を対象としたシステムの高精度化のためには、文脈情報の利用が必須である。文脈情報の利用を含め、形態素解析、文解析などの高精度化を今後とも進めていく必要がある。

一方、実際の商品やサービスとして使用可能な応用システムを実現するにあっては、100%の精度が得られない不完全な自然言語処理をベースにシステムを構築していかざるを得ない。これをカバーするためのユーザインタフェースの工夫が必要である。

自然言語は、人と人との間で、意図や意味などの情報を伝える手段として、最も基本的なコミュニケーション手段である。自然言語処理技術の高度化と同時に不完全な処理をカバーするユーザインタフェースにより、このコミュニケーションの一助となることができればと熱望する。

謝辞

本研究ならびに本論文執筆にあたり、終始ご指導、ご鞭撻を賜りました東京工業大学 大学院総合理工学研究科の小林隆夫教授に心から感謝致します。

東京工業大学 精密工学研究所の佐藤誠教授ならびに、東京工業大学 大学院総合理工学研究科の新田克己教授、山田誠二助教授、内海彰講師には、本論文に対し、適切にご指導とご助言をいただきました。心から感謝致します。

本論文の内容は、筆者が株式会社東芝 総合研究所 情報システム研究所ならびに研究開発センター 情報・通信システム研究所において行った研究に基づいており、研究の機会を与えていただいた、情報システム研究所の森健一 元所長、河田勉 元所長、情報・通信システム研究所の南正名 元所長、杉山文夫 所長、天野真家 元主幹に感謝致します。また、日頃ご指導をいただきます情報・通信システム研究所 ヒューマンインタフェース技術センターの麻田治男 元センター長、竹林洋一 元センター長、平川秀樹 センター長に感謝いたします。また、情報・通信システム研究所の浮田輝彦 元主任研究員には、筆者が自然言語処理の研究を開始した際に直接ご指導を賜りました。感謝申し上げます。

本研究を進めるにあたっては、情報システム研究所ならびに情報・通信システム研究所の上司ならびに同僚の方々のご指導とご協力を賜りました。特に、情報システム研究所の佐野洋 元研究員と情報・通信システム研究所の木下聡 研究主務とは、知識ベースに基づく質問応答システムについての研究開発を共同して行いました。また、情報・通信システム研究所の知野哲朗 元研究員、小野顕司 研究員、上原龍也 研究主務とは抄録生成システムについての研究開発を共同して行いました。情報・通信システム研究所の三池誠司 元研究主務、小野顕司 研究員、伊藤悦雄 元研究員、酒井哲也 研究員とは、検索システムに関する研究開発を共同して行いました。皆さまに心から感謝致します。

最後に、日ごろ研究開発をささえてくれる妻ならびに家族の皆に感謝致します。

付 録 A 証明

定理 4.1

すべての D と K に対して集合 D と集合 K は論理的に等価な D_1 と K_1 に変換できる。ここで、 D_1 はリテラルの集合、 K_1 は論理的包含の集合であり、 K_1 に属する論理的包含の左辺はリテラルの論理積であり、その右辺は一つのリテラルである。

(証明) 定理の証明のため、 $HK (= D \cup K)$ における命題の結言が真であることを用いる。すべての複合命題は結言標準形に変形することが可能であるので、 K に属する任意の論理的包含は、リテラルの選言形を論理積で結合した形式に書き換えることが可能である。ここで、これらの結言の各要素に関して、リテラル部分については D_1 に、選言部分については K_1 に分割することを考える。このような分割の後では、 K_1 に属する要素はすべてリテラルを論理和で結合した形になるので、これらを論理的に等しく、左辺がリテラルの結言で右辺がリテラルである論理的包含に変形することが可能である。また、 D_1 の要素はすべてリテラルとなる。 $\square$

補題 4.1

等価クラス Γ_h ($\cup_h \Gamma_h = K, \Gamma_h \cap \Gamma_k = \phi$ for $h \neq k$) が存在し、 K に属する各要素は等価関係 R' を用いてこれらの Γ_h ($1 \leq h \leq H$, H はクラスの全数) のどれか一つに振り分けられる。

(証明) 定義 4.6 と定義 4.7 より、 R' は反射的、対称的かつ遷移的であるので、 K 上で関係 R' が等価関係であることは明らかである。したがって、 K は等価関係 R' によって等価クラスに分割される。 $\square$

定理 4.2

仮に、命題 ν が Γ_h に属する推論規則と関与する命題であるとする、 ν の情報量は、 Δ_h に属するすべての真理値をともなう原子命題の組合せのみで計算可能である。

$$I(\nu, D, K) = \log_2(|C(D, \Gamma_h)|/|C(D', \Gamma_h)|),$$

(ただし ν は χ_j に関与し、 $\chi_j \in \Gamma_h$ かつ $\nu \notin D$)

ここで、式 (4.4) で得られる D' は、 ν と D 、 K から導かれるリテラルの集合である。

(証明) 定義より、 $i \neq j$ の時、命題の集合 Δ_i と Δ_j とは共通な命題は含まない。今、どの Δ_i にも属さない命題の組合せの数を C_0 とする (ただし、 $C_0 = 2^{N_0-L}$ であり、 N_0 はどの Δ_i にも属さない原子命題の数、 D ですでに真偽値が既定された原子命題の数である)。この時、 $S(D, K)$ の要素数は、命題の組合せ $C(D, \Gamma_h)$ の要素数の積で得られる。

$$|S(D, K)| = C_0 \times \prod_{h=1}^H |C(D, \Gamma_h)| \quad (\text{A.1})$$

ν が D に付け加えられる前後では、 ν が関与する Γ_h について $C(D, \Gamma_h)$ だけが値が変化する。したがって、式 (A.1) を式 (4.6) に代入することにより、式 (4.7) を得る。 $\square$

定理 4.3

仮に、 ν が Γ_j に属する χ_j に関与するとすると、仮定 4.2 に基づく $|C(D, \Gamma_j)|$ が次のようになる。

$$|C(D, \Gamma_j)| = 2^{m_j - k_j} - \sigma_j(D) \quad (4.9)$$

ここで、 m_j は χ_j 内のリテラルの数であり、 k_j は χ_j 内のリテラルの内、その真値が D において決定可能なものの数、そして、 $\sigma_j(D)$ は χ_j の否定が、 D と矛盾しない場合に 1、そうでない場合に 0 となる関数である。

(証明) Ψ_j と Ψ'_j を推論規則 χ_j に対応するリテラルの集合として以下のように仮定する。

$$\Psi_j = \{\overline{\rho_{j1}}, \overline{\rho_{j2}}, \dots, \overline{\rho_{j(k_j-1)}}, \rho_{jk_j}\} \quad (A.2)$$

$$\Psi_j = \{\rho_{j1}, \rho_{j2}, \dots, \rho_{j(k_j-1)}, \overline{\rho_{jk_j}}\} \quad (A.3)$$

ここで、 $\chi_j \equiv \rho_{j1} \wedge \rho_{j2} \wedge \dots \wedge \rho_{j(k_j-1)} \rightarrow \rho_{jk_j}$ である。 $\Psi_j \cup \Psi'_j$ は χ_j 内の原子命題に関するすべてのリテラルを含んでいるので、 $D \cap (\Psi_j \cup \Psi'_j)$ は D に属するリテラルのうち、 χ_j に関与するすべてのリテラルを含んでいることになるしたがって、 $|D \cap (\Psi_j \cup \Psi'_j)| (= k_j)$ は、 D に属し真偽値が既定である χ_j 内のリテラルの数となる。 χ_j 内の真偽値を伴った原子命題の組合せ総数は、 2^{m_j} であるから、 D と矛盾しない組合せ総数は、 $2^{m_j - k_j}$ となる。これらの組合せのうち、組合せ $\langle \rho_{j1}, \rho_{j2}, \dots, \rho_{j(k_j-1)}, \overline{\rho_{jk_j}} \rangle$ のみが χ_j と矛盾するので、 $D \cap (\Psi_j \cup \Psi'_j)$ が $\langle \rho_{j1}, \rho_{j2}, \dots, \rho_{j(k_j-1)}, \overline{\rho_{jk_j}} \rangle$ の可能性を含む場合、つまり $D \cap (\Psi_j \cup \Psi'_j) \subset \Psi'_j$ の場合、組合せ数から 1 減じる必要がある。したがって、 χ_j の否定が D と矛盾しない場合、 $|C(D, \Gamma_j)| = 2^{m_j - k_j} - 1$ であり、その他の場合、 $|C(D, \Gamma_j)| = 2^{m_j - k_j}$ となる。 $\square$

定理 4.4

仮に、 ν が χ_j に関与するとすると、仮定 4.2 に基づく ν の情報量は以下のように与えることができる。

$$I(\nu, D, K) = \begin{cases} \log_2\{(2^{m_j-k_j} - \sigma_j(D))/(2^{m_j-k_j-1} - \sigma_j(D'))\} \\ \quad (\nu \text{ は } \Gamma_h \text{ に属するある } \chi_j \text{ に関与し、かつ } \nu \notin D) \\ 1 \quad (\nu \text{ は } \Gamma_h \text{ に属するどの } \chi_j \text{ にも関与せず、かつ } \nu \notin D) \\ 0 \quad (\nu \in D) \end{cases} \quad (4.10)$$

(証明) 式 (4.9) を 式 (4.8) に代入することにより、式 (4.10) を得る。

付 録 B 修辞構造解析規則

表 B.1: 修辞関係と表層表現例 (1)

表層表現例において、“～” は任意の文字列を表し、“N～” は文頭、“～N～” は文中、“～N” は文末を意味する。

記号	関係名	表層表現例
bi	背景	～以前～, ～一般に～, 思えば～, ～過去～, かつて～, かの～, ～元来～, ～この時に～, ～このところ～, ～さきごろ～
ib	背景予告	この特集では～, ～近年～, ～今日～, ～今回の～, これから～, ～最近～, ～今のところ～, ～いまや～, ～我が国において～, 当面は～
de	定義	ここで～, ～相当する, ～というのは, ～, ～のだ～, ～ということだ, ～ので, ～, ～と呼ぶ, ～と定義する, ～と称する, いわゆる～
≡	重複	～なのである, が, それに当たる, ～と言ってよい, ～とも言えるだろう, ～同様である, ～とも言えよう, ～ものである, ～それが～目的である
<=	前重	～のみである, ～のみ, ～その～によれば～, ～例を～きりがない, 同じ～, ～たとえ～は～であれ, その方が～, ～するものであった
=>	後重	それと同じように～, その～というのは, ～本当は～, むしろ～, 代わりに～, それほど～, それによると～, そうでないなら～
～	補足	確かに～, この他～, ただし～, ちなみに～, なお～, むろん～, もちろん～, もっとも～, 本来～, 少なくとも～, ～言うまでもない
←	理由	～一因だろう, その理由は～, なぜなら～, というのも～, これというのも～, 一つには～, というのは～, ～によるものである, ～おかげである, この原因～
hl	ハイライト	まして～, ことに～, 特に～, とりわけ～, ～むしろ～, この典型～, まさに～

表 B.2: 修辞関係と表層表現例 (2)

記号	関係名	表層表現例
&	添加	この上～, さらに～, さらに～, しかも～, その上～, その代わり～, それに～, そればかりか～, ～でもあろう, <名詞>も
ap	展開	この～, ここで～, こうした～, その～, それは～, そんな～
a2	展開 (複数)	これら～, このような～, そのうちで～, それらの～, どちらも～, ～これらのうち～
≥	時間展開	以来～, やがて～, その後～, これを境に～, ～いずれ～
▽	例示	その中でも～, 例えば～, そうかと思うと～, ～などである, その一つが～, この代表～
→	順接	おかげで～, こうして～, この結果～, このため～, これによって～, これにより～, したがって～, そのため～, だから～, ～から言えることは～
結	結論	結論は～, ～達成した, ～について述べた, ～わかった, ～明らかになった
意	意見	残念なのは～, 不思議なのは～, ～大切だ, ～重要だ, ～注目される, ～必要といえる, 今後注目～, ～欲しいと思う, ～べきだ, ～望む
×	逆接	しかし～, 逆に～, けれども～, ～にもかかわらず～, これでは～, これに反して～, しかるに～, それなのに～, ところが～, とはいえ～
+	並列	あるいは～, かとおもうと～, これはまた～, それから～, なおかつ～, また～, 同じことが～, ～は同時にまた～, これに加え～
—	対比	この半面～, 一方～, これに対して～, その一方～, 他方～, それに対して～, それに比べ～, それでは～, また他方では～, それより～
=	換言	いいかえると～, すなわち～, 実は～, つまり～, それどころか～, ひいては～, いわば～, 重ねて～, あまつさえ～, 換言すれば～
rq	修辞疑問	一体～だろうか, ～べきか, ～か

表 B.3: 修辞関係と表層表現例 (3)

記号	関係名	表層表現例
※	参照	～参照されたい, ～節で～述べた, ～前述したように～, ～先に述べたように～
/	転換	さて～, では～, ところで～, それはさておき～, それはともかく～, それでは～, ときに～
□	概括	こうして～, こうすると～, かくして～, 結局～, 一言で～, そうなると～, という訳で～, 以上～, いずれにせよ～, いずれにせよ～
1s	第 1	第 1 に～, 最初に～, 一つは～, はじめに～, まず～, ここでは最初に～, ここではまず～, ここでははじめに～, (a)
2n	第 2	第 2 に～, 二つ目～, 2 番目に～, もう一つは～, ついでに～, (b)～
3r	第 3	第 3 に～, 三つ目～, 3 番目に～, (c)
4t	第 4	第 4 に～, 四つ目～, 4 番目に～, (d)
5t	第 5	第 5 に～, 五つ目～, 5 番目に～, (e)
○	次	ついで～, 続いて～, 次に～
●	最後	終わりに～, 最後に～
話	話題	～概説する, ～解説を試みる, ～ここで～述べる, ～紹介したい, この章では～, 本論では～, この特集では～, 本提案は～, ～述べたい, ～以下に～示す
ε	展開 (イプシロン)	(接続的な表現を持たない)

表 B.4: 構造選好規則表 (1)

p(pop) : ...rel-1 S) rel-2 ... の構造を優先する

n(no-pop) : ...rel-1 (S rel-2 ... の構造を優先する

5(hi-kouzouka):1 rel-1 2) rel-1 3 rel-2 4 ... のような場合にrel-1 が n に変更される

	bi	ib	de	Ξ	$\leq$	$\Rightarrow$	$\sim$	$\leftarrow$	hl	&	ap	a2	$\geq$	∇	$\rightarrow$	結	意
bi	-	p	n	n	-	-	n	n	n	n	n	p	-	n	p	p	p
ib	-	p	n	n	-	-	n	n	n	n	n	p	-	n	p	p	p
de	-	p	n	n	-	-	n	n	n	n	n	p	-	n	p	p	p
Ξ	p	p	n	n	-	-	p	n	n	n	p	p	-	n	p	p	p
$\leq$	p	p	n	n	-	-	p	n	n	n	p	p	-	n	p	p	p
$\Rightarrow$	p	p	n	n	-	-	p	n	n	n	n	p	-	n	p	p	p
$\sim$	-	p	n	n	-	-	n	n	n	n	n	p	-	n	p	p	p
$\leftarrow$	-	p	n	n	-	-	-	n	n	n	n	-	-	n	p	p	p
hl	p	p	n	n	-	-	p	p	n	n	n	p	-	n	p	p	p
&	p	p	n	n	-	-	p	p	n	n	n	p	-	n	p	p	p
ap	-	p	n	n	-	-	-	n	n	n	5	p	p	n	p	p	p
a2	-	p	n	n	-	-	-	n	n	n	n	p	-	n	p	p	p
$\geq$	-	-	-	-	-	-	-	-	-	-	-	p	5	-	-	p	-
∇	-	p	n	n	-	-	p	n	n	n	n	p	-	n	p	p	p
$\rightarrow$	p	p	n	n	-	-	n	n	n	n	n	p	-	n	p	p	p
結	n	n	n	n	n	n	n	n	n	n	n	n	n	n	n	n	n
意	p	p	n	n	-	-	n	n	n	n	n	p	-	n	p	p	p
$\times$	-	p	n	n	-	-	-	n	n	n	n	-	-	n	p	p	p
+	-	p	n	n	-	-	n	n	n	n	n	-	-	n	p	p	p
-	-	p	n	n	-	-	n	n	n	n	n	p	-	n	p	p	p
=	p	p	n	n	-	-	p	n	n	n	p	p	-	n	p	p	p
rq	n	p	n	n	-	-	n	n	n	n	n	p	-	n	n	p	n
※	n	p	n	n	-	-	n	n	n	n	n	p	-	n	n	p	n
/	n	p	n	n	-	-	n	n	n	n	n	p	-	n	n	p	n
□	n	n	n	n	-	-	n	n	n	n	n	p	-	n	n	p	n
1s	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
2n	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
3r	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
4t	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
5t	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
○	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
●	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
話	n	-	n	n	-	-	n	n	n	n	n	n	-	n	n	p	n
ϵ	-	p	-	-	-	-	n	-	-	-	n	p	-	n	p	p	p

表 B.5: 構造選好規則表 (2)

	$\times$	$+$	$-$	$=$	rq	$\ast$	$/$	$\square$	1s	2n	3r	4t	5t	$\bigcirc$	$\bullet$	話	ϵ
bi	n	-	-	n	n	p	p	p	-	-	-	-	-	p	p	p	n
ib	p	-	-	n	n	p	p	p	-	-	-	-	-	p	p	p	n
de	-	-	p	n	n	p	p	p	-	-	-	-	-	p	p	p	n
Ξ	p	p	p	n	n	p	p	p	-	-	-	-	-	p	p	p	n
$\leq$	p	p	p	n	n	p	p	p	-	-	-	-	-	p	p	p	n
$\Rightarrow$	p	p	p	n	n	p	p	p	-	-	-	-	-	p	p	p	n
$\sim$	p	-	p	n	n	p	p	p	-	-	-	-	-	p	p	p	n
$\leftarrow$	p	p	p	n	n	p	p	p	-	-	-	-	-	p	p	p	n
hl	p	p	p	p	n	p	p	p	-	-	-	-	-	p	p	p	n
&	p	p	p	p	n	p	p	p	-	-	-	-	-	p	p	p	n
ap	p	p	-	n	n	p	p	p	-	-	-	-	-	p	p	p	n
a2	-	p	-	n	n	p	p	p	-	-	-	-	-	p	p	p	n
$\geq$	-	-	-	-	-	-	-	-	-	-	-	-	-	p	p	-	-
∇	p	p	p	n	n	p	p	p	-	-	-	-	-	p	p	p	n
$\rightarrow$	p	p	p	n	n	p	p	p	-	-	-	-	-	p	p	p	n
結意	n	n	n	n	n	n	n	n	-	-	-	-	-	p	p	p	n
$\times$	p	p	p	n	n	p	p	p	-	-	-	-	-	n	n	n	n
$+$	p	n	p	n	n	p	p	p	-	-	-	-	-	p	p	p	n
$-$	p	n	p	n	n	p	p	p	-	-	-	-	-	p	p	p	n
$=$	p	p	p	n	n	p	p	p	-	-	-	-	-	p	p	p	p
rq	n	n	n	n	n	p	p	p	-	-	-	-	-	p	p	p	p
$\ast$	n	n	n	n	n	p	p	p	-	-	-	-	-	p	p	p	n
$/$	n	n	n	n	n	p	p	p	-	-	-	-	-	p	p	p	n
$\square$	n	n	n	n	n	n	n	p	-	-	-	-	-	p	p	p	p
1s	-	-	-	-	-	-	-	-	5	5	5	5	5	5	5	p	n
2n	-	-	-	-	-	-	-	-	5	5	5	5	5	5	5	p	n
3r	-	-	-	-	-	-	-	-	5	5	5	5	5	5	5	p	n
4t	-	-	-	-	-	-	-	-	5	5	5	5	5	5	5	p	n
5t	-	-	-	-	-	-	-	-	5	5	5	5	5	5	5	p	n
$\bigcirc$	-	-	-	-	-	-	-	-	5	5	5	5	5	5	5	p	n
$\bullet$	-	-	-	-	-	-	-	-	5	5	5	5	5	5	5	p	n
話	n	n	n	n	n	n	p	p	-	-	-	-	-	p	p	p	n
ϵ	p	p	p	n	n	-	p	p	-	-	-	-	-	p	p	p	n

付 録 C 発表論文リスト

第3章に関連する論文

住田一男, 浮田輝彦: “質問応答システムにおける会話の自然性に関する考察”, 電子情報通信学会論文誌, **J70-D**, 11, pp.2287-2293 (1987).

第4章に関連する論文

Kazuo SUMITA, Teruhiko UKITA and Shin-ya AMANO: “Disambiguation in Natural Language Interpretation Based on Amount of Information,” *IEICE Trans.*, **E74**, 6, pp.1735-1746 (1991).

Teruhiko UKITA, Satoshi KINOSHITA, Kazuo SUMITA, Hiroshi SANO and Shin-ya AMANO: “Understanding Conversational Sentences Using Mult-Paradigm World Knowledge,” *IEICE Trans. Inf. & Syst.*, **E75-D**, 3, pp.352-362 (1992).

第5章に関連する論文

住田一男, 知野哲朗, 小野顕司, 三池誠司: “文書構造解析に基づく自動抄録生成と検索提示機能としての評価”, 電子情報通信学会論文誌, **J78-D-II**, 3, pp.511-519 (1995).

Kazuo SUMITA, Kenji ONO, Tetsuro CHINO, Teruhiko UKITA and Shin-ya AMANO: “A Discourse Structure Analyzer for Japanese Text,” *Proc. Int. Conf. FGCS 1992*, pp.1133-1140 (1992).

Kazuo SUMITA, Kenji ONO and Seiji MIIKE: “Document Structure Extraction for Interactive Document Retrieval Systems,” *Proc. ACM SIGDOC'93*, pp.301-310 (1993).

Kenji ONO, Kazuo SUMITA and Seiji MIIKE : “Abstract Generation Based on Rhetorical Structure Extraction,” *Proc. COLING '94*, pp.343-348 (1994).

第6章に関連する論文

Kazuo SUMITA, Kenji ONO and Seiji MIIKE: "Document Structure Extraction for Interactive Full-Text Retrieval," *Natural Language Processing Pacific Rim Symposium '93*, pp.144-152 (1993).

Seiji MIIKE, Etsuo ITOH, Kenji ONO and Kazuo SUMITA : "A Full-Text Retrieval System with a Dynamic Abstract Generation Function," *Proc. ACM SIGIR '94*, pp.152-161 (1994).

酒井哲也, 三池誠司, 住田一男 : "文書検索システムの動的抄録提示インタフェースの評価," 情報処理研究会資料, **HI** 57-7, pp49-54 (1994)

三池誠司, 住田一男 : "文の意味解析に基づく全文検索", 情報処理研究会資料, **FI** 34-3, pp.17-24 (1994)

住田一男, 三井誠司 : "文の意味解析に基づく全文検索", 第8回人工知能学会全国大会, **24-2**, pp.667-670 (1994)

その他の発表論文

今井聖, 住田一男, 古市千枝子, "音声合成のためのメル対数スペクトル近似 (MLSA) フィルタ", 電子通信学会論文誌, **J66-A**, 2, pp.122-129 (1983)

G.Jones, T.Sakai, M.Kajiura and K.Sumita : "Experiment in Japanese Text Retrieval and Routing using the NEAT System," *Proc. ACM SIGIR '98*, pp.197-205 (1998)

参考文献

- [1] J.Barwise and J.Perry : *Situations and Attitudes*, MIT Press (1883).
- [2] D.G.Bobrow, R.M.Kaplan, M.Kay, D.A.Norman, H.Thompson and T.Winograd : “GUS, A Frame-Driven Dialog System,” *Artificial Intelligence*, **8**, pp.155-173 (1977).
- [3] J.P.Callan : “Passage-Level Evidence in Document Retrieval,” *Proc. ACM SIGIR’94*, pp.302-310 (1994).
- [4] F.R.Campagnoni and K.Ehrlich : “Information Retrieval Using a Hypertext-Based Help System,” *ACM Trans. on Information Systems*, **7**, 3, pp.271-291 (1989).
- [5] J.R.Carbonell : “AI in CAI,” *IEEE Trans. Man-Machine Systems*, **11**, pp.190-202 (1970).
- [6] 知野哲郎, 浮田輝彦, 小野顕司, 住田一男 : “日本語論説文自動抄録システムの試作と評価”, 第46回情処全大, 第3分冊, pp.189-190 (1993).
- [7] R.Cohen : “Analyzing the Structure of Argumentative Discourse,” *Computational Linguistics*, **13**, 1-2, pp.11-24 (1987).
- [8] J. de Kleer : “An Assumption-based TMS,” *Artificial Intelligence*, **28**, pp.127-162 (1986).
- [9] G.Fauconnier : *Mental Spaces : Aspects of Meaning Construction in Natural Language*, MIT Press (1985).
- [10] 福本淳一, 安原宏 : “日本語文章の構造解析”, 情処研資, **NL 85-11** (1991).
- [11] 福本淳一, 安原宏 : “文の連接関係解析に基づく文章構造解析”, 情処研資, **NL 88-2** (1992).
- [12] 福本文代, 福本淳一, 鈴木良弥 : “文脈依存の度合を考慮した重要パラグラフの抽出”, 自然言語処理, **4**, 2, pp.89-109 (1997).
- [13] D.Fum : “Evaluating Importance: A Step Towards Text Summarization,” *Proc. IJCAI’85*, pp.840-844 (1985).
- [14] H.P.Grice : “Logic and Conversation,” *Syntax and Semantics*, **3**, Speech Acts, pp.41-58 (1975).
- [15] B.J.Grosz : “The Representation and Use of Focus in a System for Understanding Dialogs,” *Proc. IJCAI ’77*, pp.67-76 (1977).

- [16] B.J.Grosz and C.L.Sidner : “The Structures of Discourse Structure,” *Technical Report CSLI*, CSLI-85-39 (1985).
- [17] P.Hayes and D.R.Reddy : “An anatomy of graceful interaction in spoken and written man-machine communication,” *Technical Report CMU*, CMU-CS-79-144, Computer Science Department, CMU (1979).
- [18] 平井誠, 北橋忠宏 : “省略語と事象間関係から見た日本語文の意味解析と文脈の枠組みについて”, 信学技報, **NLC** 86-8 (1986).
- [19] 平川秀樹, 天野真家 : “構文／意味優先規則による日本語解析”, 人工知能学会第3回全国大会 8-4 (1989).
- [20] 平川秀樹, 天野真家 : “日本語解析における最適解探索”, 情処研資, **NL** 74-2 (1989).
- [21] J.R.Hobbs : “Coherence and Coreference,” *Cognitive Science*, **3**, pp.67-90 (1979).
- [22] J.R.Hobbs, M.Stickel, P.Martin and D.Edwards : “Interpretation as Abduction,” *Proc. ACL'88*, pp.95-103 (1988).
- [23] 細野公男 : “情報検索理論・技法の問題点とその解決の方向”, 情処研資, **FI** 24-5 (1991).
- [24] 稲垣博人 : “事象解析による要約情報の抽出”, 情処研資, **NL** 84-3 (1991).
- [25] P.Ingwensen : *Information Retrieval Interaction*, Taylor Graham, London (1992).
- [26] 石橋弘義, 重永信一, 新納浩幸, 福重貴雄, 安川秀樹 : “英文要約システム『DIET』”, 第38回情処全大, pp.293-294 (1989).
- [27] M.Doi, M.Fukui and I.Iwai : “Research on Model Based Document Processing System DARWIN,” *Proc. INTERACT'87*, pp.1101-1106 (1987).
- [28] I.Iwai, M.Doi, K.Yamaguchi, M.Fukui and Y.Takebayashi : “A Document Layout System Using Automatic Document Architecture Extraction,” *Proc of CHI'89*, pp.369-374 (1989).
- [29] A.K.Joshi and S.Weinstein : “Control of Inference of Some Aspects of Discourse Structure — Centering,” *Proc. IJCAI '81*, pp.385-387 (1981).
- [30] M.Kameyama : “A Property-sharing Constraint in Centering,” *Proc. ACL'86*, pp.200-206 (1986).
- [31] H.Kamp : “A Theory of Truth and Semantic Representation,” *Formal Methods in the Study of Language*, pp.277-322, Mathematisch Centrum (1981).
- [32] 神門典子 : “構成要素カテゴリを用いた原著論文の内部構造分析”, 情処研資, **FI** 25-7 (1992).
- [33] 神門典子 : “文章特性が異なる複数領域の原著論文の構造”, 情処研資, **FI** 33-4 (1994).

- [34] J.Kaplan : “Cooperative Responses from Portable Natural Language Database Query System,” *Computational Models of Discourse*, pp.167–208, MIT Press (1983).
- [35] S.Kinosita, H.Sano, T.Ukita, K.Sumita and S.Amano : “Knowledge Representation and Reasoning for Discourse Understanding,” *Logic Programming '88*, pp.238–251 (1989).
- [36] 岸本行生, 須之内美幸, 塚田康博, 千葉滋, 石川徹也: “テキストの構造化に基づく検索システム”, 情処論, **35**, 5, pp.908–916 (1994).
- [37] 工藤育男, 友清睦子: “日本語の述部の特性を用いた省略の補完機構について”, 信学論, **J76-D-II**, pp.624–635 (1993).
- [38] 黒橋禎夫, 長尾真: “表層表現中の情報に基づく文章構造の自動抽出”, 自然言語処理, **1**, 1, pp.3–20 (1994).
- [39] 小松英二, 福本淳一, 木下哲男: “マルチエージェントによる頑健な自然言語処理の協調方式”, 情処研資, **NL** 109–8 (1995).
- [40] W.Lehnert: “Narrative Text Summarization,” *Proc. AAAI*, pp.337–339 (1980).
- [41] H.P.Luhn: “The Automatic Creation of Literature Abstracts,” *IBM Journal*, pp.159–165 (Apr. 1958).
- [42] W.C.Mann and S.A.Thompson: “Rhetorical Structure Theory: A Framework for the Analysis of Texts,” *USC/Information Science Institute Research Report*, RR-87-190 (1987).
- [43] 松尾衛, 森辰則, 中川裕志: “条件表現による日本語マニュアル文のゼロ主語同定”, 情処論, **38**, 4, pp.737–475 (1997).
- [44] K.R.McKeown: “Discourse Strategies for Generating Natural-Language Text,” *Artificial Intelligence*, **27**, pp.1–41 (1985).
- [45] S.Miike, K.Hasebe, H.Somers and S.Amano: “Experiences with an On-line Translating Dialogue System,” *Proc. COLING'88*, pp.155–162 (1988).
- [46] 三池誠司, 小野顕司, 住田一男: “文書の構造解析に基づく文書情報検索”, 情処研資, **FI** 31-6 (1993).
- [47] S.Miike, E.Itoh, K.Ono and K.Sumita: “A Full-Text Retrieval System with a Dynamic Abstract Generation Function,” *Proc. ACM SIGIR '94*, pp.152–161 (1994).
- [48] 三輪真木子: “データベースサーチャーの視点”, 情報処理, **33**, 10, pp.1162–1170 (1992).
- [49] 村田真樹, 長尾真: “用例や表層表現を用いた日本語文書中の指示詞・代名詞・ゼロ代名詞の指示対象の推定”, 自然言語処理, **4**, 1, pp.87–109 (1997).
- [50] 長尾真, 辻井潤一, 田中一敏: “意味および文脈情報を用いた日本語文の解析 — 文脈を考慮した処理”, 情処論, **17**, 1, pp.19–28 (1976).

- [51] 長尾真：言語工学, pp.222-283, 昭晃堂 (1983).
- [52] 長尾真：“情報社会の生態学”, 情処研資, CH 11-6 (1991).
- [53] 長尾真：“新しい機械翻訳のための自然言語処理”, 人工知能学会, 11, 4, 500-506 (1996).
- [54] 中本幸夫, 田野崎康雄, 岩井勇：“日本語解析を用いたフルテキストサーチの実験”, 第46回情処全大, 4B-4, pp.3-125-126 (1993).
- [55] 中岩浩巳, 池原悟：“語用論的・意味論的制約を用いた日本語ゼロ代名詞の文内照応解析”, 自然言語処理, 3, 4, pp.49-63 (1996).
- [56] 中田和男：音声, 音響工学講座7, p.178, コロナ社 (1977).
- [57] 根岸正光：“フルテキスト・データベースの応用動向”, 情報処理, 33, 4, pp.413-420 (1992).
- [58] T.Nishida, X.Liu, S.Doshita and A.Yamada：“Maintaining Consistency and Plausibility in Integrated Natural Language Understanding,” *Proc. COLING '88*, pp.482-487 (1988).
- [59] 西村健士, 島津秀雄：“特定表現の重点的解析による科学技術論文構造化手法”, 情処研資, FI 29-5 (1993).
- [60] 西尾実, 岩淵悦太郎, 水谷静夫編：岩波国語辞典 第三版, p.533, p.1130 (1979).
- [61] 野上宏康, 齋藤裕美：“仮名漢字変換技術”, 人工知能学会, 11, 6, pp.845-851 (1996).
- [62] 大平栄二, 阿部正博, 小松昭男, 市川熹：“情報検索における柔軟な対話制御方式”, 信学論, J76-D-II, 3, pp.586-595 (1993).
- [63] 大須賀節雄：“知識の表現に関する一考察 — 情報理論的観点から —”, 情処論, 25, 4, pp.685-694 (1984).
- [64] オンラインデータベースディレクトリ '91, p.21, 東洋経済新報社 (1991).
- [65] 小野顕司, 浮田輝彦, 天野真家：“文脈構造の分析”, 情処研資, NL 70-2 (1989).
- [66] 小野顕司, 住田一男, 知野哲郎：“日本語論説文の自動抄録のための文脈構造解析”, 第46回情処全大, 7B-10, pp.3-187-188 (1993).
- [67] K.Ono, K.Sumita and S.Miike：“Abstract Generation Based on Rhetorical Structure Extraction,” *Proc. COLING '94*, pp.343-348 (1994).
- [68] 大澤一郎, 米澤明憲：“オブジェクト指向方式による対話理解システム”, 情処研資, NL 44-7 (1984).
- [69] J.R.Quinlan：“Learning Efficient Classification Procedures and their Application to Chess and Games,” *Machine Learning, an Artificial Intelligence Approach*, eds. R.S.Michalski, et al., pp.463-482, Tioga (1983).
- [70] F.M.Reza：“An Introduction to Information Theory,” McGraw-Hill (1961).

- [71] 佐川雄二, 杉江昇, 杉原厚吉: “コンサルテーション・システムにおける対話の階層構造モデルについて”, 情処研資, **NL** 59-7 (1987).
- [72] 佐久間あゆみ: 文書構造と要約文の諸相, p.9, くろしお出版 (1989).
- [73] G.Salton: *Automatic Text Processing - The Transformation Analysis, and Retrieval of Information by Computer*, Addison-Wesley Publishing Company Inc. (1989).
- [74] 佐野洋, 杉村領一, 赤坂宏二, 久保幸弘: “語構成に基づく形態素解析”, 情処研資, **NL** 66-3 (1988).
- [75] R.Scha and L.Polanyi: “An Augmented Context Free Grammar for Discourse,” *Proc. COLING'88*, pp.573-577 (1988).
- [76] 新開正史, 田淵仁浩, 村岡洋一: “主題の意味構造に基づく論文検索法の提案”, 第46回情処全大, 6G-6, pp.4-207-208 (1993).
- [77] C.L.Sider: “Focusing in Comprehension of Definite Anaphora,” *Computational Models of Discourse*, eds. M.Brady, et al., pp.267-330, MIT Press (1983).
- [78] 住田一男, 浮田輝彦: “質問応答システムにおける応答の自然性に関する考察”, 信学技報, **NLC** 86-16 (1986).
- [79] K.Sumita, T.Chino, K.Ono, T.Ukita and S.Amano: “A Discourse Structure Analyzer for Japanese Text,” *Proc. Int. Conf. Fifth Generation Computer Systems (FGCS)'92*, pp.1133-1140 (1992).
- [80] K.Sumita, K.Ono and S.Miike: “Document Structure Extraction for Interactive Document Retrieval Systems,” *Proc. ACM SIGDOC'93*, pp.301-310 (1993).
- [81] K.Sumita, K.Ono and S.Miike: “Document Structure Extraction for Interactive Full-Text Retrieval,” *Proc. Natural Language Processing Pacific Rim Symposium (NLPRS)'93*, pp.144-152 (1993).
- [82] 住田一男, 伊藤悦雄, 三池誠司, 武田公人: “対話的抄録生成機能を持つ文書検索システム”, 情処研資, **HI** 52-3 (1994).
- [83] 住田一男, 三池誠司: “知的情報検索の動向”, 人工知能学会, **11**, 1, pp.10-16 (1996).
- [84] 高野敦子, 平井誠, 北橋忠宏: “自然言語によるデータベース検索における協調的応答生成”, 信学技報, **NLC** 95-4 (1995).
- [85] 田村直良, 和田啓二: “セグメントの分割と統合による文書の構造解析”, 自然言語処理, **5**, 1, pp.59-78 (1998).
- [86] 辻井潤一: “論説文における文脈構造”, 日本学術振興会 文字言語・音声言語の知能的処理第152委員会第7回研究会資料 7-1 (1988).
- [87] H.Turtle and R.Croft: “Evaluation of an Inference Network-Based Retrieval Model,” *ACM Trans. Information Systems*, **9**, 3, pp.187-222 (1991).

- [88] T.Ukita, K.Sumita, S.Kinoshita, H.Sano and S.Amano : "Preference Judgment in Comprehending Conversational Sentences Using Multi-Paradigm World Knowledge," *Proc. Int. Conf. Fifth Generation Computer Systems (FGCS)*'88, pp.1133-1140 (1988).
- [89] 浮田輝彦, 木下聡 : "自然言語理解のための知識表現と推論 — 知識表現", 情報処理, **30**, 10, pp.1224-1231 (1989).
- [90] T.Watanabe, M.Araki and S.Doshita : "Evaluating Dialogue Strategies under Communication Errors Using Computer-to-Computer Simulation," *IEICE Trans. Inf. & Syst.*, **E81-D**, 9, pp.1025-1033 (1998).
- [91] B.L.Webbeer : "Syntax beyond the Sentence : Anaphora," *Theoretical Issues in Reading Comprehension*, eds. R.J.Spiro, et al., Lawrence Erlbaum Associates (1980).
- [92] R.Wilkinson : "Effective Retrieval of Structured Documents," *Proc. ACM SIGIR '94*, pp.311-317 (1994).
- [93] T.Winograd : *Understanding Natural Language*, NY., Academic Press (1972).
- [94] 山本幹雄, 中川聖一 : "対話における談話構造とユーザ・モデルの関係", 信学技報, **NLC 89-4** (1989).
- [95] 山本和英, 増山繁, 内藤昭三 : "文章内構造を複合的に利用した論説文要約システム GREEN", 情処研資, **NL 99-3** (1994).
- [96] 山村毅, 大西昇, 杉江昇 : "日本語指示詞の前方照応現象の分類", 信学論, **J75-D-II**, 2, pp.371-378 (1992).
- [97] 山村毅, 大西昇, 杉江昇 : "観察の概念に基づく照応解決法", 信学論, **J77-D-II**, 1, pp.162-169 (1994).
- [98] 安原宏, 小松英二, 日比孝, 加藤安彦 : "要約支援システム COGITO", 情報処理, **30**, 10, pp.1258-1267 (1989).
- [99] 吉本啓 : "談話処理における日本語ゼロ代名詞の扱いについて", 情処研資, **NL 56-4** (1986).
- [100] J.Weizenbaum : "ELIZA-A Computer Program for the Study of Natural Language Communication between Man and Machine," *ACM Comm.*, **9**, 1, pp.36-45 (1966).