

日本語DB検索システムQuestにおける意味解析

加藤 恒昭, 中川 優
(日本電信電話公社 横須賀電気通信研究所)

1. はじめに

データベース(DB)技術の発展、普及に伴い、DB検索を日常的な業務としていないユーザによる検索が増加している。しかも、その検索は従来に比べ、非定型かつ複雑なものになりつつある。我々は、そのような検索のニーズに答えるため、検索言語のシンタックスやDB構成などを意識することなく、容易かつ自然にDB検索を行なえるような日本語によるDB検索システムQuest(intelligent Question answering System)の開発を進めている。

自然言語によるDB検索システムとしては、既に国内外に幾つかの提案がなされているが、その多くは、関係DB(RDB)と自然言語との対応を直接的に行っており、自然な検索という面からはまだ不十分な点が多い[1][2][3]。Questでは、この対応づけにDB世界表現モデルと名付けた表現能力の高いモデルを利用し、より自然な言い回しを理解できるようになっている。

また、Quest開発において特に重点を置いた点はシステムのportabilityである。ここで、portabilityとは以下のような条件を満たすことを言う。

- ・対象となるDBに関する情報がシステム内で局在化しており、その変更、作成が容易であること。
- ・対象とするDBMSに依存している部分が局在化していること。
- ・対象分野に固有な言い回しがある場合、その解析処理を容易に組み込めること。

Questではこれらの要求を以下の手法によって解決している。

- ・対象DBに依存する部分を基本的にはDB世界表現モデルに集約した。
 - ・中間表現として、アルファ式[4]に近い論理式を採用し、対象検索言語を意識する部分を限定した。
 - ・意味解析処理がルール集合の形をしており特殊な処理はそれに応じたルールとして追加を可能とした。
- 本資料では主にQuestの核であるDB世界表現モデルとそれを用いた意味解析手法について述べる。

2. システム構成

Questの構成を図1に示す。本システムはDEC-2060上のDEC10-Prologを用いて開発されている。本システムは大きく7つにモジュール分割できる。

・構文解析部

ローマ字ベタ書きで入力された質問文を単語辞書と文法を利用して解析し、文節の並びに展開する。このレベルでは文節間の係り受け関係を固定することは行なわない。このため、文法は形態素解析レベルの処理を行なうための汎用文法である。ただし、解析結果のalternativeを減ずるため、文節間の接続関係を調べ、文法的に適切でない解析結果が生じないように配慮している。

単語辞書は、自立語の場合には後に述べる意味世界表現モデルの要素との対応が記述されている。付属語の場合には、その文法的役割(主格を表わす格助詞等)が記述されている。

本モジュールの核部にはProlog上のBottom up parser BUP[5]を利用している。

構文解析部の出力例を図2に示す。

・DB世界表現モデル

対象とするDBの構成やその実世界との対応が記述されている。このモデルは意味解析部に利用される意味世界表現モデル(SWORD Semantic World Description model)と検索言語生成部に利用されるアクセスモデル(AH Access Model)に分割格納されている。AHは実DB中の表から意味の基本構成要素とな

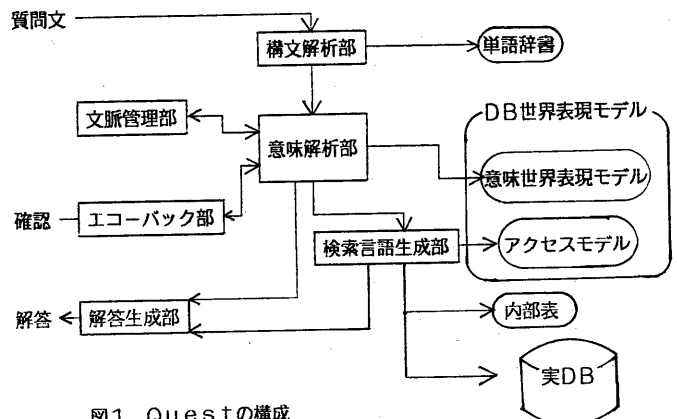
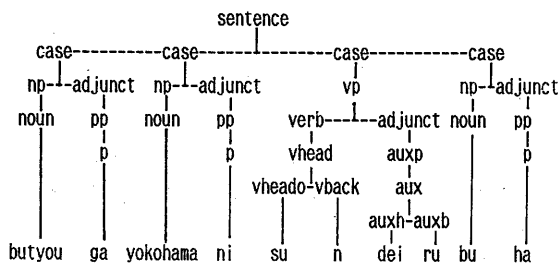


図1 Questの構成

入力文 部長が横浜に住んでいる部は
構文木



構文情報

```
[case(np, [adj], [attr($buchou)], []),
 case(np, [att, atp], [entity(yokohama, #juusho)], []),
 case(pred, 34, [class(@sumu)], [contin]),
 case(np, [adj, rep], [class(*bu)], [])]
```

図2 構文解析の例

る表(仮想表と呼ぶ)を定義している。SWORD はそのような仮想表と実世界つまり日本語表現との対応を表現している。詳細については 3. で述べる。

・意味解析部

構文解析結果を基に SWORD を参照することにより、質問文の意味表現である中間表現を生成する。そのアルゴリズムについては 4. で述べる。

中間表現は仮想表に対する検索を表現するアルファ式に近い論理式となっている。アルファ式は関係論理における集合の内包的定義に相当する表現であるが、中間表現では、これを述語論理的に表現している。図3に中間表現の例を示す。ここからわかるように基本的にはPrologシンタックスを採用しており、大文字で始まるアトムは変項を表現し、'.' と ';' は各々and とorを表現している。また、各述語は仮想表に対応している。

中間表現では

{[変項の並び]|条件式}

に相当する集合の内包的定義を

result([変項の並び]), 条件式

で表現している。ただし、result述語が先頭に来る必要はない。条件式中に現れ、変項の並びに現れない全ての変項には存在量記号が仮定されている。また、全称量記号に代わる表現として

all([変項の並び], 述語1, 述語2)

を用意した。これは述語1を真にするすべての変項並びにおいて述語2が真である時に真となる高階述

語である。また、DB検索に固有なmax, minなどの統計関数や出力順序の指定などについては厳密さを欠くが各々max(X,Y), min(X,Y), inc(X)と述語的に表現し中間表現に含めることとした。

・文脈管理部

意味解析部から逐次、質問文や文節の意味である中間表現を受けとって、現在とひとつ前の質問の話題を記憶しており、意味解析部が代名詞の同定や省略補完のためにそれらの情報を必要とする場合に提供するモジュールである。

意味解析部の要求している情報が保持されていないときはシステム内に定義されたデフォルト値やユーザへの問い合わせによって解決を図る。

・エコーバック部

意味解析部によって得られた質問文の解釈を入力とし、それをユーザに提示し、その解釈が適切かどうかの確認を行なう部分である。ユーザが否と答えた場合は意味解析部に対してバックトラックをかけ、解釈のalternativeを捜すこととなる。

本モジュールは将来的には中間表現から疑似日本語を生成し、提示することを考えているが、現在の所、中間表現を整形したものを提示している。

・検索言語生成部

中間言語から検索言語を生成するモジュールである。Questは、検索の対象として外部のDB実表だけでなく、システム内部に内部表と呼ばれるものを持っている。内部表はPrologプログラムとして保持されており、より自然な検索のために必要となる整理された知識やルールを実DBに影響を与えることなく取り込むためのものである。

本モジュールは、まず、仮想表に対する検索表現として与えられる中間表現をAMを参照することで、実表と内部表に対する検索表現に展開する。この後、検索のスケジューリングを行ない適切な順序で、内部表と実表に対して検索コマンドを発行する。内部表への検索コマンドは中間表現に類似した形式の述

部長が横浜に住んでいる部は

```
result(A), bumei(B, A),
 buchou(B, C), sumu(C, yokohama)
```

ボスより給与の多い社員を知りたい

```
result(A), syainmei(B, A), kyuuyo(B, C)
 boss(B, D), kyuuyo(D, E), C > E
```

図3 中間表現の例

語で実現している。実DBへの検索言語としてはSQLを利用しており、その生成に関しては検索効率を考慮して、表結合と副照会の適切な使い分けを行っている。

・解答生成部

Questは従来のDB検索システムでは扱っていなかったyes-no型の疑問文をもその対象としているため、検索結果から質問文の真偽を判定する処理が必要となる。この処理と検索結果の出力を行なうのが本モジュールである。

検索結果の出力は表イメージで行なわれ、yes-no型の場合は判定の根拠となる値が、wh型の場合は、ユーザが陽に求めた検索項目とそれを一意識別するための情報が出力される。

3. DB世界表現モデル

DB世界表現モデルはDBもしくはそれが対象としている世界の意味を記述するためのものである。

DB技術の分野でも、DBの自然な設計を目的としてこのようなデータモデルの提案が幾つかなされている。Questで利用しているDB世界表現モデルはこれらの中で、SDM [6]、DAPLEX [7]、E-Rモデル[8]を基礎として、日本語の表現や省略がモデル構造に対応するよう拡張整理したものである。

3.1 DB世界表現モデルの要素

DB世界表現モデルの基本構成要素は、entity, class, attribute である。

entityは、世界内に存在する「ものごと」を表現している。例えば、「加藤という名前の社員」という個人、「加藤が営業部に属す」という状態、更に、「加藤」という名前そのものなどがentityとして扱われる。これらentityはその性質によって3種類に分類される。

- ・もの(object) 人や建物などの具体的な物
- ・こと(event) 所属や存在などの出来事や状態
- ・名前(name) 「横浜」、「加藤」などの識別子

classはこれらentityの集合であり、共通の特徴を持ったentityが集められている。classはそこに属すentityによって以下のように分類される。

- ・oc(Object Class) ものを要素とするclass
社員、部など
- ・ec(Event Class) ことを要素とするclass
属す、住むなど
- ・nc(Name class) 名前を要素とするclass
地名、社員名など

attributeはclassの要素であるentityの特徴を

表わすためのもので、あるentityから同じもしくは異なるclassのentityもしくはその集合への写像である。例えば、部長attributeは部classのentityを特徴づけるもので、その値として社員classのentityをとる。

attributeが定義されているclassを定義classその値となるentityを要素とするclassを値classと呼ぶ。部長attributeの場合、定義classは部、値classは社員である。

attributeには、classにおけるその役割によって特別な名前を与えられているものがある。

・pa(Primary Attributes)

このattributeの値によってclassのentityが一意識別できるようなattributeもしくはその集合

・da(Default Attribute)

日本語によって名指しできるのは、ncのentityのみであるが、一般にはocのentityはそこに定義されているdaの値によって名指しされる。例えば、「加藤」は単なる名前であるが、これによって、名前attributeの値が「加藤」であるoc社員のentityを表現している。この名前attributeをdaと呼ぶ。daの定義classは常にocで、その値classはncである。

・fa(Functional Attribute)

ecに定義されるattributeの中で、そのecを定義classとする他のattributeの値classが定義classで、値classがそのattributeの値classであるようなattributeとしてふるまうものを指す。

以上がDB世界表現の基本構成要素である。これらに加えて基本構成要素間の関連や基本要素から導出される要素が定義される。そのようなものとして、subclass, grouping, 導出attributeがある。

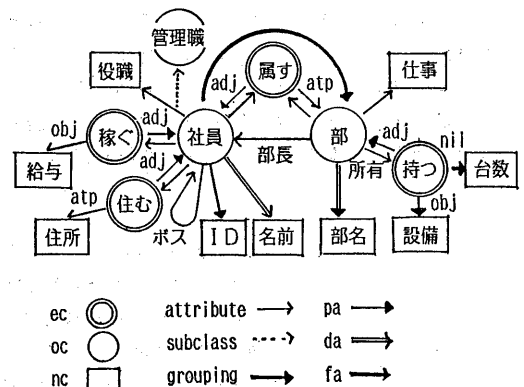


図4 SHORD ダイアグラム

subclassはあるocのentityに条件を課してそれを満すentityのみからなる部分集合である。subclassはその基となったclassのattributeを全て受けつぐ。例えば、oc社員からその役職attributeの値が課長以上という条件で、そのsubclass管理職を作ることができる。

grouping関連は他のclassのentityによってあるclassのentityがグループ分けされているような時、換言すれば、あるclassのentityが他のclassのentityの集合としての側面を持っているときに定義される。例えば、社員は彼が属す部によってグループ分けされるからoc部とoc社員の間にはgrouping関連が定義される。

導出attributeは複数のattributeを合成することによって得られるattributeである。例えば、ボスというattributeは社員が属す部の部長という形で定義できる。

3.2 DB世界表現モデルの表現

DB世界表現モデルは意味解析に利用される述語の定義を行なっていると考えることができる。この述語は先にも述べたように仮想表と呼ばれる。DB世界表現モデルの要素は述語と以下のように対応する。

- ・ecはそこに定義されているattributeに対応する引数を持った述語となる。
- ・値class, 定義classが共にecでないattributeは導出attributeか否かにかかわらず、2引数の述語となる。

これらの引数の定義域は全てそのclassによって決定される。

DB世界表現モデルはこのように定義された仮想表と日本語表現とを対応づけるSWORDと実表および内部表との対応をとるAHから構成される。

SWORDに記述される情報は以上で述べたモデル構成と各要素の名前である。この名前が辞書に記述されている意味情報となる。ただし、SWORDのレベルでは、導出attributeは一般のattributeと区別されない。また、ecは日本語における動詞に対応し、

```
trans(boss(A,B), (#syain(A,-,C,-,-,-), #bu(C,B,-)))
```

-は無名変項

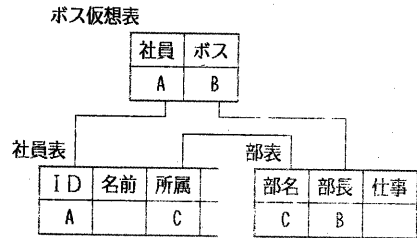


図6 アクセスモデル内部表現

そこに定義されたattributeはその格に対応するため格標識が付けられる。このダイアグラム表現を社員情報を例に図4に示す。

Quest内部においてはSWORDをこのダイアグラムのリンク単位にfact形式で表現している。その例を図5に示す。

仮想表は実表の写像や結合によって得ることができる一種のビューであるから、その定義をAHに記述している。AHは仮想表単位にfact形式で表現されており、これを参照することで、仮想表に対する検索表現を実表に対する検索表現に変換できる。AHの例を図6に示す。

4. 意味世界表現モデルと意味解析

質問文に現れる自立語にはその意味としてSWORDの要素が以下のように対応している。

- ・用言 ecのclass名
- ・名詞 oc,ec,ncのclass名
attribute名
ncのentity

質問文中の名詞句にはSWORDのentityが対応づけられ、文には、真偽値が対応する。ただし、質問文中の文は埋め込み文となって、名詞句に係ることが多く、これによって生成された名詞句にはentityが

```
attr(type(A,*bu), type(B,*syain), $buchou, buchou(A,B), naux, no-omit, single, nil).
attr(type([A,B], @zokusu), type(A,*bu), nil, zokusu(A,B), atp, omit, single, pa).
attr(type([A,B], @zokusu), type(B,*syain), nil, zokusu(A,B), adj, omit, single, nil).
subclass(type(A, &kanriskyoku), *syain, (yakusyoku(A,B), (B=kachou;B=buchou))).
grouping(type(A,*bu), type(B,*syain), zokusu(A,B), [$sigoto]).
```

図5 SWORD 内部表現

対応する。名詞句を修飾する文は、その中にentityを表現する自由変項を含んでおり、修飾される名詞句の意味となるentityはその自由変項と同じclassに属している。

これら、質問文の部分的意味を表現するために以下に示す途中表現を利用する。

$\text{type}(X, \text{CLASS})$: 条件式

これは、条件式を満たすようなclass CLASSのentity Xを意味する。条件式は必ずXを含んでおり、そこに現れるX以外の変項のあるものには存在量記号が仮定されている。図7に日本語とSWORDと中間表現との基本的な対応を示す。

意味解析で行なわれる処理は、単語の意味から始めて、名詞句の意味、文の意味と合成してゆき、最終的に質問文全体の意味を得ることに他ならない。

質問文の各部分の意味を上記のようにentityを表現する変項と考えると、その処理は変項どうしの関係づけを行なうことに相当する。

Questでは、このような処理を係り受け関係の決定ととらえている。各文節はそのタイプや付属語に応じて、各々係りの役割と受けの役割を有している。これは表1のようにまとめられる。

Questの意味解析の基本的処理はSWORDを参照することで、係り受けのチェックを行ないながら、各単語の役割に応じた処理を行なうことである。

5. 意味解析部の構成と動作

5.1 構成

意味解析部は図8に示すような構成を持っている。本構成は一種のプロダクションシステムと見ることもできるし、turing機械と考えてもよい。

入力となる構文解析結果は文節単位のセルを持つ入力テープであり、これをp-infoと呼ぶ。w-stack

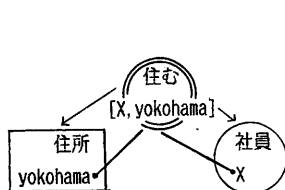
表1 係りの役割とその受けの分類

係り	役割	受け
名詞句+格助詞	格要素	用言(文)
名詞句+連体修飾助詞	名詞句の連体修飾	名詞句
名詞句+並列助詞	名詞句の並列	名詞句
連体文	文の連体修飾	名詞句
連用文	文の並列	連体文

は係りの情報が入れられるスタックである。具体的には、先に述べた途中表現の $\text{type}(X, \text{CLASS})$ と係りの役割が入れられている。s-infoは解析結果である中間表現が格納されるが、解析の途中では、先に述べた途中表現の条件式が基本的には","(and)で結合されて入れられている。u-infoには解析途中で得られた自由変項が $\text{type}(X, \text{CLASS})$ の形式で格納されている。これは埋め込み文の解析に利用され、その内容は文の解析が終了した時点でw-stackに移される。o-infoは意味解析部が後処理として必要とする処理が何であるかを記憶している部分であり、compositionalityが保たれない幾つかの処理で利用される。SWORD探索部は、動作部の要求に応じて、SWORDを探索し、必要とする情報を与える。

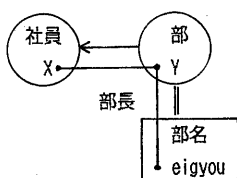
エンジン部の基本的な動作は、次に示すサイクルの反復である。

- ・p-infoとw-stackの状態から適当なルールを選ぶ。
 - ・ルール適用が可能であるかをSWORD探索部に問い合わせ、可能であれば必要な情報を得る。実際には、Prologのunificationを利用しているので、この2つの処理は同時に実現される。
 - ・ルールのaction部に従って、p-info, s-info, w-stack, u-info, o-infoなどの変更を行なう。
- 動作部の停止条件はp-infoを最後まで読み取り、



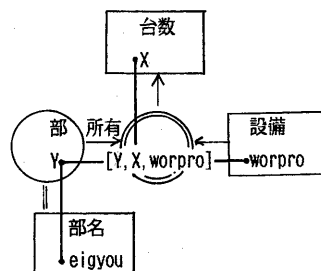
横浜に住む社員

$\text{type}(X, \text{*syain}): \text{sumu}(X, \text{yokohama})$



営業部の部長

$\text{type}(X, \text{*syain}): \text{bumei}(Y, \text{eigyou}), \text{buchou}(Y, X)$



営業部のワープロ所有台数

$\text{type}(X, \text{*daisuu}): \text{bumei}(Y, \text{eigyou}), \text{motu}(Y, X, \text{worpro})$

図7 日本語、中間表現、SWORDの対応

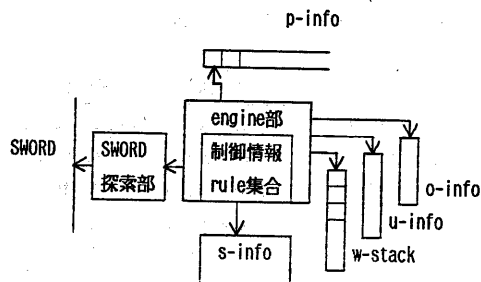


図8 意味解析部構成

かつ、w-stack が空であることであり、これ以前に適用できるルールがなくなった場合にはバックトラックを生じる。また、処理の高速化のために、ルール適用後の状態が既に生じた状態と同じであるかを check し、そうであれば強制的に fail して無駄な探索を減らしている。

ルールは大きく以下の3種類に分類される。

・受けのルール

w-stack の要素と p-info の要素の間に係り受け関係が成立すれば、p-info をその係りを受けた後の意味に変換し、w-stack の要素を pop する。更に、その関係を表現する仮想表を s-info に加える。

・係りのルール

p-info の先頭要素を取り除き、その意味と係りの役割を w-stack に push する。必要があれば、u-info を更新する。

・変型ルール

p-info の表層的な構造の変換を行なう。例えば、「最も給与が多い」を「給与が最も多い」に変換する。このようなルールを設けることで、上記の係り受けルールの数を減らすことができる。

これらのルールは、head に各作業域のルール適用前後の状態がバアで記述されており、body に SHWORD 探索部の呼び出しと必要となる情報の抽出手続きが記述されている。図9に、ルールの一例を示す。実際には、clause index を利用するためや、デバックのための付加的な引数が加えられている。

5.2 動作例

簡単な例によって、意味解析部の動作を説明する。質問文が、「部長が横浜に住んでいる部は」であったとする。初期状態として、p-info には図2に示した構文解析結果が入り、それ以外は全て空である。

「部長が」文節

部長 attribute の値 class *syain とその entity を表現する変項 B、係りの役割が主格の格要素であることが、w-stack に push される。仮想表 buchou(A, B) が s-info に入れられる。ここで、A は部長 attribute の定義 class *bu の entity である。この A は自由変項であるから、class 情報と共に u-info に格納される。

「横浜に」文節

同様に処理が行なわれる。

「住んでいる」文節

用言の受け処理として sumu(X, Y) の X, Y が各々 B, yokohama に unify される。この unify と同時に、w-stack の内容が pop され、s-info には、sumu(X, Y) が加えられる。prolog の unification を利用しているため、s-info を直接操作することなく、変項の関係づけがなされる。また、w-stack により係り受けの非交差が自然に実現できている。

係りの処理として、u-info にある type(A, *bu) を w-stack に push する。

「部は」文節

受けの処理として、文節の意味を entity(A, *bu) に変換し、w-stack の先頭を pop する。

係りのルールによって、case([adj], type(A, *bu)) が w-stack に push される。

終了ルールとして、w-stack の先頭要素の entity を result の引数と unify する処理が起動されるが、この entity が nc の entity でない場合には、その da の値 class の entity が unify され、仮想表 bumei(A, C) が s-info に加えられる。最終的に得られる中間表現は図3に示したものとなる。

以上の様に意味解析の基本的動作は受けのルールによって各文節の意味である entity とそれを制約する条件を明らかにし、係りのルールでその entity の係りの役割を明らかにするというサイクルの反復で

```
rule([case(np, P-list, [attr(Name)|other], [])|P-info], P-info,
      W-stack, [case(P-list, type(E2, C2))|W-stack],
      U-info, New-U-info,
      O-info, New-O-info,
      S-info, (Term, S-info))
:-not-overlap(P-list, [or, and, para, naux]),
check-sword(attribute, [attr(Name)|other],
              , O-info, New-O-info, -, type(E2, C2), Term),
free-var(E2, (Term, S-info), U-info, New-U-info).
```

図9 ルールの例

ある。

5.3 幾つかの個別処理

(1) 省略を含んだ言い回しへの対処

質問文において、2つの要素に係り受けが可能であることはSWORDにおいて、その2つの要素の間にリンクがあることと等価である。その最も基本的な形式は「aのb」という表現において、aがentityを意味しており、bがそのentityを要素とするclassを定義classとするattributeを意味しているとき、係り受けが可能で、その意味はbの値classのentityになるというものである。例えば、「加藤が属す部の部長」という表現では「加藤が属す部」がaにあたり、「部長」がbにあたる。

質問文に現れるのは、このような厳密な言い回しのみではないが、それらもSWORDとの対応によって理解することができる。以下に例を挙げる。

- ・「営業部の部長」はこの形式にあてはまらない。なぜなら「営業部」は部名classのentityであって部classのentityではないからである。この場合は、daとの関係で、「営業部」の意味を名前attributeの値が「営業部」である部classのentityと解釈する。

- ・「加藤の部」や「加藤の住所」という表現は、上の省略補完を行なっても形式にあてはまらない。なぜなら「部」や「住所」はclass名であってattributeを意味していないからである。この場合、そのclass名で、ecを介した導出attributeを意味していると考えられる。

- ・「営業部のコピー機の台数」の場合は、faの名前によって省略がなされていると考える。これによってec名を補い、「営業部のコピー機の所有台数」の省略と理解する。

- ・「加藤の仕事」の場合は、上記の省略補完を行なっても、「加藤」は社員classのentityであり、「仕事」は部classを定義classとするattributeである。このような言い回しではgrouping関連の定義が利用される。この例では、部と社員とがgrouping関連で結ばれているから、その解釈は「加藤が属す部の仕事」となる。grouping関連に伴うもうひとつの省略は統計関数と共に生じる。「営業部の平均給与」は「営業部に属す社員の平均給与」として解釈される。

- ・subclassにおけるattributeの継承も省略と同じように扱える。「営業部に属す管理職」は管理職classが社員classのsubclassであるから意味を持つ。この解釈はsubclass定義の条件を加え、「営業部に属して、役職attributeの値が課長以上の

社員」となる。subclass定義の条件が付加されるのは、この例のようにsubclass名が陽に示されたときとsubclassにのみ定義されたattributeが指定された場合である。これらがひとつのentityに重複して現れた場合にsubclass定義の条件が重複して中間表現に現れないように、subclassのentityが現れたという情報をo-info中に保存しておき、解析の後処理として条件を中間表現に加えている。

これら省略補完の処理は、そのほとんどがSWORD探索部によって行なわれ、ルールがそれを意識する必要はない。

(2) 並列の処理

質問文に現れる並列は以下の2つに分類される。

- ・検索項目の並列 「加藤の住所と給与」
- ・条件要素の並列 「営業部と販売部の社員」

検索項目の並列の時に問題となるのは、上の例で言えば、加藤が給与にも係ることの判定である。これを実現するため、「住所と」文節に対する係りルールでは、w-stackに

```
para(entity, type(X, #juusyo), [type(Y, *syain)])
```

という形式で、この文節に係っているentityを保存しておく。これらの情報はs-infoを探索することで得る。この後、「給与」文節で並列を受ける時点で、自由変項があれば、このw-stack中の変項とunifyする。このような処理によって得られる名詞句の意味は`type([X,Z], [#juusho, #kyuuyo])`と表現し、result述語以外では受けられないようにする。

条件項目の並列での問題は、主にそれがand並列であるかor並列であるかの判定である。Questでは、これを以下のように判定する。

- ・埋め込み文の格でないものはor並列
- ・埋め込み文の格であるものの内、その文の用言の意味となるecに定義されているattributeがpa, faのみであればand並列
- ・埋め込み文の格であるものの内、その文の用言の意味となるecにpa, fa以外のattributeが定義されていればor並列

これによって、

「営業部と販売部の社員」 or 並列

「ワープロとコピーを持っている部」 and 並列

「横浜と横須賀に住んでいる社員」 or 並列

となる。この判定はandに判定されることが多少多いが、or並列は「ワープロかコピー」というように陽に指定できるため問題はないと考えている。

この判定を実現するために、並列を含んだ格は`case([p], type(para([X,Y]), CLASS), [OP, X=A, Y=B])`の形式でw-stackにpushされる。用言の受け処理で、

このような要素が現れたら、ecを参照して、OPを",
"もしくは";"にunifyし、項形式に変換してs-infoに加える。更にこのecを表現する仮想表のコピーを作り、それぞれの引数をX,Yとunifyする。例えば、「ワープロとコピーのある部は」の中間表現は、
`result(A),bumei(B,A),motu(B,C,X),
 motu(B,C,Y),(X=worpro,Y=copy)`
 となる。

(3) 否定

「ワープロを持っていない部」などの否定を含む文については、以下のような中間表現を生成することを考えている。

```
result(A),bumei(B,A),  

all(C,motu(B,D,C),not(C=worpro))
```

しかし、この形式では検索言語生成部の負荷が大きいため、これを意味解析後処理で、

```
result(A),bumei(B,A)  

all(E,motu(E,D,worpro),not(B=E))
```

に変換する。このような否定が現れたという情報はo-infoに格納して、後処理の起動を行なう。

(4) 最上級

「営業部で、給与が最も多い社員は」という最上級を含んだ言い回しに対しては次のような中間表現が対応する。

```
result(A),syainmei(B,A),zokusu(B,C),  

bumei(C,eigyau),kyuuyo(B,D),max(E,D),  

zokusu(F,G),bumei(G,eigyau),kyuuyo(F,E)
```

このような言いまわしの解析が困難である理由は、何の中で「最も」であるのかが、「最も」を解析しているタイミングでは明らかになっていない点にある。現在、「最も多い」を解析する時点では、s-infoとo-infoにmax(E,D)を加え、意味解析の後処理で、o-infoにmax(E,D)があった場合は、s-info中のDに関連する中間表現を抽出し、これをEの条件に変換して、s-infoに加えることを考えている。

(5) 一意化情報の付加

DB検索においては、質問文中に陽に示された意味だけを忠実に抽出するだけでは不充分であることがある。例えば、「加藤と中川の住所」という質問文は、住所の検索を要求しているが、それだけでは出力されるのは住所だけであり、どちらがどちらの住所であるかがユーザにわからない。このような不親切な解答になることを防ぐため、Questでは意味解析後処理として、一意化情報の付加を行なっている。これは中間表現がor並列や自由変項を含んでいて、解が複数になる場合には、それらの変項をresult述語の引数に加えるものである。

6. おわりに

現在、我々はここに述べた機能の実現と文脈処理部の設計を進めている。文脈処理部では、さしあたり、entityを意味する代名詞や省略を扱う予定である。基本的には、文節単位の意味が確定した時点で、その意味である変項とs-infoの内容をclass名に対応させてstackしておき、代名詞がどのclassのentityかが明らかになった時点でその情報を引くという方式をとる予定である。また、前文の検索結果の参照などのために前文で何が話題の中心であったかも保存しておく必要があると考えている。

Questの適用領域は、ここで述べた事務用のDBや案内DBであると思われるので、この2つのDBに対してシステムを作成し、オビニオンテストなどによって有用性を確認する予定である。

謝辞

日頃御指導いただく寺島信義知識ベース研究室室長に感謝します。

参考文献

- [1] 藤崎他 データベース照会システム「ヤチマタ」と名詞句データ模型
情報論文誌vol.20 no.1 (1979)
- [2] 泉田他 対象世界のモデルを利用したデータベース検索システム
情報データベースシステム研究会43-2 (1984)
- [3] Martin.P et al: Transpaortability and Generality in a Natural-Language Interface System
IJCAI pp573-581
- [4] 植村: データベースシステムの基礎
オーム社 (1979)
- [5] 松本他 Prologに埋め込まれたBottom-up parser:BUP
情報自然言語研究会34-6 (1982)
- [6] Hammer.M: Database Description with SDH : A Semantic Database Model
ACM Trans. Database System vol.6 no.3 (1981)
- [7] Shipman.D.W: The Functional Data Model and The Data Language DAPLEX
ACM Trans. Database System vol.6 no.1 (1981)
- [8] Chen.P.P : The Entity-Relationship Model - Toword a Uified View of Data
ACM Trans. Database System vol.1 no.1 (1976)