

自己概念化モデルにおける推論過程

唐沢 博
(阪大・基礎工)

[1] まえがき

人工知能システムは、一般に図1のように表現できる。推論の際、もし知識が一方的に利用されるだけならば(図1のaの過程)、知識体系は固定的であり、この体系を越えた範囲の世界の知識については対処できない。しかしながら、対象となる世界があらかじめわかっているということは、それに適した推論方式や知識表現がシステム設計時に採用できるという点で、処理能率の向上を配慮できる。現存する Expert システム(例えば DENDRAL, MYCIN 等)の多くが、このようなタイプの構成を持つ。これらのシステムは非常に powerful であるが、対象世界を変更した場合、または一つの世界が徐々に変化していく場合、それらに応じて知識を如何に変更し、新しい知識を追加していくかが問題となる。そこで、これを克服するために、新しく獲得した知識に基き、「学習」を通じて既存の知識体系を変更していく方法が考えられる(図1のbの過程)。

知識の獲得に関しては、explicit な過程と implicit な過程との二種類の過程を定義する。

まず、explicit な過程とは、「AハBデアル。」といった宣言的な知識の獲得、及び「Aハ～シタ。」というようなある事実を表わすものである。

一方、implicit な過程とは、複数の類似事実から帰納的に得られる、一般性を伴った知識を得る過程である。これを本論では概念化と呼ぶことにする。

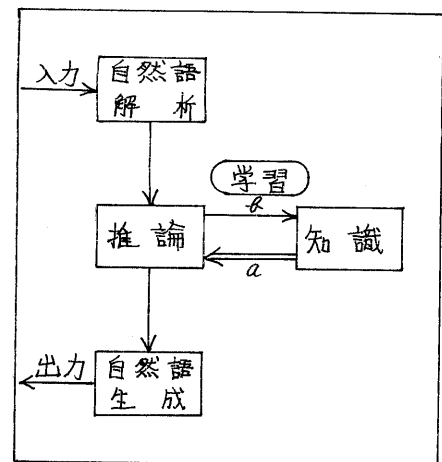


図1 人工知能システムの一般的表現

本報告は、人工知能システムの知識獲得における概念化に関する問題を取っている。特に、ある場面において一般的に見られる行動の系列に注目して、Script⁽¹⁾ (Schank, 1975) の event の系列における「学習による概念化」について論じた。その理論的背景に関しては既に報告してあるので、本報告では主として、その応用例としての人工知能システム ASPELAN-2⁽²⁾ について報告する。

[2] 自己概念化モデルについて

筆者は、知識獲得の implicit な過程における学習モデルとして、自己概念化モデル[以後、SCモデル (Self-Conceptualization Model) と表わす]を提唱し

た。それを図2に示す。

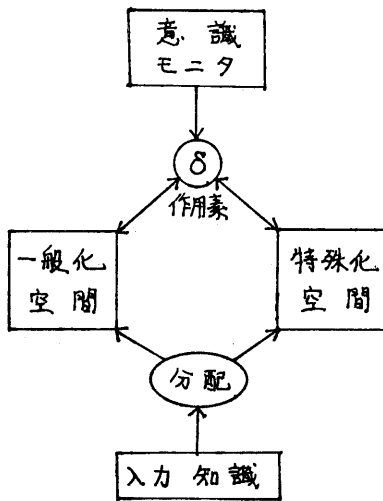


図2 SCモデルのブロック構成図

このモデルでは、implicit な知識獲得のために一般化空間を設けた。この空間上で帰納的推論が遂行される。それに対し、explicit な獲得知識は特殊化空間に収納されるものとする。先ず入力された知識は、一般化空間と特殊化空間の両者へ供給される。一般化空間では、類似した行動系列の傾向を抽出する。その結果、個々の事実についての知識は失うので、これらは特殊化空間に保存する。しかし、特殊化空間では行動系列の傾向は抽出できない。従って、両空間は相補的に作用し合って推論を進める。その仲介を果すのが作用素であり、これを管理するのが意識モニタと名付けた機能部である。

2-1 一般化空間と特殊化空間の相補性

一般化空間と特殊化空間とが相補的に作用し合う必要があるのは次の理由にもよる。藤崎教授（東大）の主張する⁽³⁾ように、帰納的推論は絶対的真実を必ずしも導くとは限らない。そこで、一般化空間が提供する推論のための「一般的な」知識に対して訂正機能の役割を果すのが特殊化空間である。これは、個々のケースは一般則よりも優先すべき知識バリューを持つと解釈できるからである。

次に、特殊化空間に対する一般化空間の役割について述べる。特殊化空間には既に述べたように、入力されたある事実に関する知識と宣言的事実に関する知識が詰っている。これらの知識のみで「特定事実の照会」及び「概念階層を利用した推移律に基づく演繹的推論」は可能である。しかし、それを越えた質問に対しては無力になってしまう。これをある程度の範囲までカバーしようとする機能を持つのが一般化空間である。一般化空間は、特殊化空間の演繹的推論を越えるような質問について帰納的推論を適用する役割を果すのである。これをさらに越えた範囲の質問要求については、SCモデルもついに無力になる。このレベルの応答には、ひらめきや発想の能力が必要になるであろう。これらの関係を図3に示す。また、本節で述べた相補性については図4にまとめた。

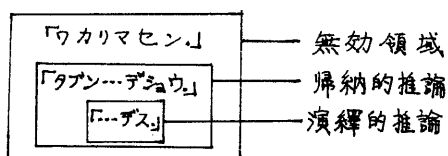


図3 推論と応答の関係

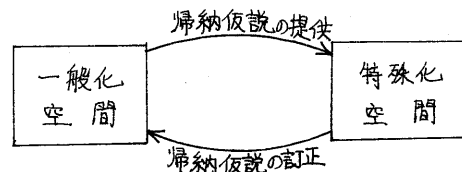


図4 一般化空間と特殊化空間の相補性

2-2 SCモデルにおける自己概念化の基盤

SCモデルの一般化空間上では、類似知識の出現度数を用いて概念化が進められる。

以下の説明では、動作を英大文字、時間的経過を矢印($\rightarrow$)で表わし、行動系列はそれらの組合せで表現する(例: $A \rightarrow B \rightarrow C$)。図5に示すように、知識の連続した入力に対して一般化空間上の知識は、個々の知識の重ね合せも行っており、その出現度数を保持する。各行動の右肩の数字は、その出現度数を示している。

入力知識	一般化空間上の知識
1) $B \rightarrow C \rightarrow D \rightarrow E$	: $B^1 \rightarrow C^1 \rightarrow D^1 \rightarrow E^1$
2) $E \rightarrow G$	: $B^1 \rightarrow C^1 \rightarrow D^1 \rightarrow E^2 \rightarrow G^1$
3) $A \rightarrow C \rightarrow D \rightarrow E \rightarrow F$	: $A^1 \rightarrow C^2 \rightarrow D^2 \rightarrow E^3 \rightarrow F^1$ $B^1 \rightarrow C^2 \rightarrow D^2 \rightarrow E^3 \rightarrow G^1$
4) $F \rightarrow H \rightarrow I$	: $A^1 \rightarrow C^2 \rightarrow D^2 \rightarrow E^3 \rightarrow F^2 \rightarrow H^1 \rightarrow I^1$ $B^1 \rightarrow C^2 \rightarrow D^2 \rightarrow E^3 \rightarrow G^1$
5) $H \rightarrow I \rightarrow J$	: $A^1 \rightarrow C^2 \rightarrow D^2 \rightarrow E^3 \rightarrow F^2 \rightarrow H^2 \rightarrow I^2 \rightarrow J^1$ $B^1 \rightarrow C^2 \rightarrow D^2 \rightarrow E^3 \rightarrow G^1$
6) $H \rightarrow I$	: $A^1 \rightarrow C^2 \rightarrow D^2 \rightarrow E^3 \rightarrow F^2 \rightarrow H^3 \rightarrow I^3 \rightarrow J^1$ $B^1 \rightarrow C^2 \rightarrow D^2 \rightarrow E^3 \rightarrow G^1$
7) A	: $A^2 \rightarrow C^2 \rightarrow D^2 \rightarrow E^3 \rightarrow F^2 \rightarrow H^3 \rightarrow I^3 \rightarrow J^1$ $B^1 \rightarrow C^2 \rightarrow D^2 \rightarrow E^3 \rightarrow G^1$

図5 入力知識と一般化空間上の知識との関係

一般に、ある動作 p について、 p が出現する直前の時刻を t 、その直後の時刻を $t+1$ とし、それらの出現度数を $W_{p,t}$ 、 $W_{p,t+1}$ と表わしたとき、

$$W_{p,t+1} = W_{p,t} + 1, \quad (t \geq 1, W_{p,1} = 1) \quad (2-1)$$

で定義する⁽²⁾。これを「印象付けのプロセス」と呼ぶ。

2-3 SCモデルにおける帰納仮説の形成過程

ある質問内容が特殊化空間の手に負えなくなったことを意識モニタが検出すると(図3参照)、次に意識モニタは一般化空間上に帰納仮説を形成すべきことを指示する。SCモデルでは、この仮説をScriptのevent系列に設定した。これは、Scriptがある場面における絞切型な行動系列を示す仮説であり、行動の系列という比較的単純で扱い易い体系であることに依った。さらに、Scriptは高次仮説体系であるので、帰納的操作によってある行動系列が抽出できたとすると、それはその行動系列の背後に存在する「ある場面」を概念化したことに相当する。このことは当初のSCモデルの狙いと合致することになる。

Schank等の唱えた古典的なScriptは、図1のaの過程に従う固定的な仮説知識であった。これに対しSCモデルにおけるScriptは、学習によって逐次修正が行われるように拡張されたScriptである(図1のbの過程を付加)。これをAスクリプト(Adaptive Script)と呼ぶ。Aスクリプトには「均一なAスクリプト」と「入れ子構造のAスクリプト」が定義される。

均一なAスクリプトは、注目動作 w について、

$$\frac{\min\{|W_i - W_p|\}}{W_p} \leq \varepsilon, \quad (0 < \varepsilon \ll 1) \quad (2-2)$$

で意味ネットワークとして表現された行動系列の探索が打ち切られるような範囲の意味ネットワーク断片である。

入れ子構造のAスクリプトは、注目動作 w について、

$$W_i \geq (1 - \varepsilon) W_p \quad (2-3)$$

で探索が打ち切られる範囲の意味ネットワーク断片である。

均一なAスクリプトと入れ子構造のAスクリプトの持つ意味は、(2-1)式と合せ考えればそれぞれ「同程度に印象付けられた行動の系列」及び「ある程度以上に印象付けられた行動の系列」としてとらえることができる。

さらに重要なことは、入れ子構造のAスクリプトを用いて行動系列の「要約」が可能であるという点である。要約は次のように定義される。

ある動作 w についての入れ子構造のAスクリプト E を構成する任意の動作 f_i について、その出現度数を W_{f_i} とすれば、

$$W_{f_i} \geq \theta, \quad [\max(W_{f_i}) \geq \theta > (1 - \varepsilon) W_p] \quad (2-4)$$

であるような f_i が形成するAスクリプトは、 E の要約である。このとき、

$$\eta = \frac{\theta}{\max(W_{f_i})}, \quad \left[1 \geq \eta > \frac{(1 - \varepsilon) W_p}{\max(W_{f_i})} \right] \quad (2-5)$$

なる η を要約率と呼ぶ。 η が1に近い程、要約性は高くなり、 $\eta=1$ の場合は E の中で最も重要な動作が抽出される。つまり要約とは、特に印象付けられた動作(の系列)を抽出する過程である。

[3] SCモデルの動作過程

本論で述べているSCモデルは、一般的なQAシステムとしてインプリメントし、これをASPELAN-2と命名した。

ASPELAN-2は、

- 1) 知識の付加
- 2) 所有知識の照会
- 3) 非所有知識の照会
- 4) 入力文の要約命令

を形式的な日本語(カタカナ文)で受理し、応答する。

3-1 知識の付加

ASPELAN-2が受け付ける知識は、一連の行動系列を表わす単文の連鎖（重文）及び「AハBデアル」といった宣言的単文である。

この重文についてASPELAN-2は、文を単文単位に分割し、一つの単文についてeventを中心としたフレームを形成し、各フレームは時間的な因果リンクを通じて連結させる。そして、これらのフレームは特殊化空間に格納する。さらに動作に関する情報のみ（2-1）式及び図5に示すルールに従って一般化空間上に蓄える。図6はその様子を示している。

また、主語に関して、重文の最初の単文に主語が含まれていれば、二番目以後の単文中には同一主語である限り省略してもよい機構になっている。重文の途中で新しい主語が出現した場合は、そこから以後の各単文の主語はいずれもその新しい主語と解釈される。

宣言文は、二つの事物ノードを階層リンクで結び付ける役目を果すフレームに変換される。このフレームを通じてシステムは、概念階層の処理を行なう（図7参照）。

3-2 所有知識の照会

「…デスカ?」という形式の文章は質問文として処理し、システムは応答を返す。ASPELAN-2は質問文を受けとるとそれを解析し、以下に示すような種類に分類する。

- 1) 事実の真偽を問うもの
- 2) 疑問詞を含むもの
- 3) 直前行為・後続行為に関するもの
- 4) 目的・方法に関するもの

この内、1), 2), 3) について所有知識を探索する。その過程は、特殊化空間内に格納されて

*タロウハ ミセハ ハイッテ、タバモノ チュウモンシテ、ユウショクオ タベテ、ミセオ デタ。

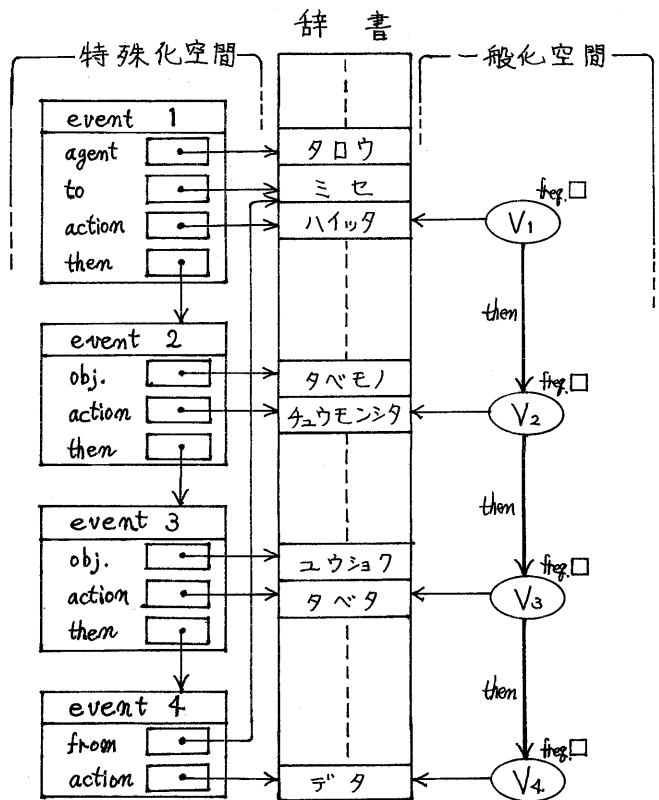


図6 ASPELAN-2における知識表現

*タロウハ ニンゲンデアル。

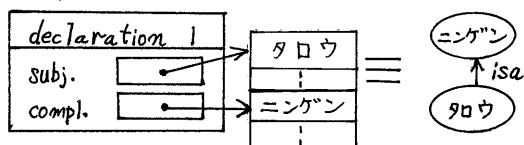


図7 概念階層を表わす宣言的知識

いるフレーム群を探索し該当フレームを呼び出すことに相当する。ASPELAN-2は次のような手法により、該当フレームを連想的に呼び出す。

任意のフレームAの各slot内容を $a_1, a_2, \dots, a_m$ として $A(a_1, a_2, \dots, a_m)$ という記法を用いると、質問文を表わすフレーム $Q(g_1, g_2, \dots, g_n)$ について、

$$Q \subseteq A \quad (3-1)$$

であるようなフレームAを探す。そのために $g_i (i=1 \dots n)$ を同一slot内容として含むフレームをすべて探し出し(リンクを辿るだけでよい)、そのフレームの集合を A_i としたとき、

$$P_{i+1} = A_i \cap A_{i+1} \neq \phi \quad (3-2)$$

である限り、 i を1から $n-1$ まで増加する。 $i=n-1$ に達しない内に $P_{i+1} = \phi$ となった場合、該当フレームが存在しないと判断する。

一方、 $i=n-1$ かつ(3-2)式を満たすならば、 P_n の要素が該当フレーム(群)である。

該当フレームの存否と質問文との関係は、図8のようなフローに従う。

また、2)型の質問文に現われる疑問詞は表1のものに制限されている。

3)型、4)型の質問文は、次のkey wordによって検出される。

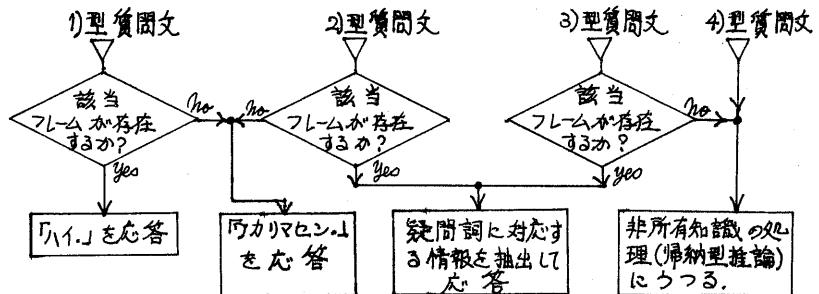


図8 種別質問文の処理プロセス

* 3)型の質問文

a) --- (シテ), どうシタデスカ?

〔後続行為〕

b) --- (シタ)マエ, どうシタデスカ?

〔直前行為〕

* 4)型の質問文

c) --- ナンノタメニ --- デスカ?

〔目的〕

d) --- ドノヨウニシテ --- デスカ?

〔方法〕

このような質問文が特別扱いされるのは、Aスクリプトを用いた帰納型推論が有効に働く対象だからである。

疑問詞	意味	対応する助詞
ダレ(か)	行為者	…ハ
イツ	時	…ニ
ドコ(デ)	地点	…デ
ダレ(ト)	共同者	…ト
ナニ(オ)	目的物	…オ
ドコ(カラ)	初期地点	…カラ
ドコ(ヘ)	目標地点	…ヘ
ドウ(シタ)	行為	(動詞)

表1 システムが扱う疑問詞と意味

3-3 非所有知識の照会

ASPELAN-2システムは、3)型質問の内容に合致する知識を探索して発見できない場合、及び4)型質問を受け付けた場合は、次に帰納推論の過程に入り、「ワカリマセン」という応答を極力回避するように機能する(2-1節の図3,図4参照)。

システムは、推論過程でAスクリプトを形成させて、それをもとに応答内容を生成する。換言すれば、『一般的にはこう考えられるであろう』というような応答を行うのである。従って、その応答文の形式は、「タブン…デショウ」という表現をとる。

「…(シテ),ドウシタデスカ?」というような、後続行為を問う3)型質問に対しては、ASPELAN-2は一般化空間内で(2-2)式に従うような順向探索を行ない、選ばれた一つの動作を所有するようなeventを特殊化空間内より引き出し、現在の質問に整合するような内容にそれを修正して応答とする(次頁、例1参照)。

「…(シタ)マエ,ドウシタデスカ?」という直前行為についての質問の場合は、一般化空間上で逆行探索を行なう以外は上述したプロセスと同様である(次頁、例2参照)。

「…ナンノタメニ、 α シタデスカ?」といった目的を問う4)型質問については次のようなプロセスを踏む。先ず一般化空間内をノード α から順向探索していき、 α についての均一なAスクリプトの端まで到達する。次に、入れ子になっているAスクリプトがそのすぐ前方に存在するかどうかが調べる。存在していれば、その入れ子のAスクリプトに入った最初の動作ノードが目的を示しているものとする。入れ子のAスクリプトが存在しなければ、均一なAスクリプトの前方端の動作を目的とみなす(次頁、例3参照)。

「…ドノヨウニシテ、 α シタデスカ?」は、方法を問う4)型質問であるが、これも α についての均一なAスクリプトを生成して、応答のベースをつくる。すなわち、動作 α に到った行動の系列をそのAスクリプトから引き出し、これを応答とする(次頁、例4参照)。

3-4 入力文の要約命令

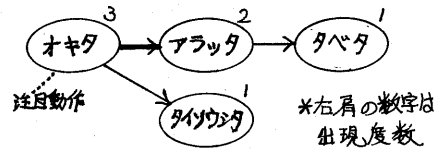
「ヨウヤクシナサイ。」の命令文の後に行動系列を表わす重文が組合されたものが入力されると、システムは入れ子構造のAスクリプトを形成し、要約率 η に従って行動系列の要約を行い、要約文を出力する。これは、ASPELAN-2の最も重要な機能である。この過程は(2-4)式及び(2-5)式に従うものであり、(2-4)式で選択される動作ノードの集合を R とすると、入力文中の全動作ノードの集合を I としたとき、要約文中の動作ノードの集合 Q は、

$$Q = R \cap I$$

(3-3)

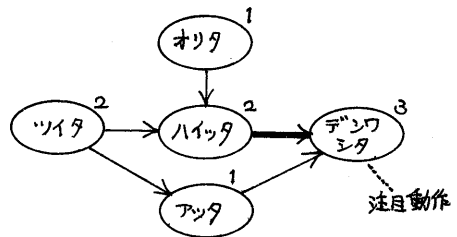
で表わされる。Aスクリプトを用いた要約は、入力知識からimplicitな知識獲得を行うことによって常にその時点で最も一般的であろうと判断されるような要約が生成される点で、古典的なScriptを用いた要約と大きな相異がある。その動作例を例5(次頁)に示す。

*タロウハ 7シニ オキテ、カオオ アラツタ。
 *シロウハ 8シニ オキテ、カオオ アラツテ、チヨウシヨクオ タヘタ。
 *ハナコハ 7シニ オキテ、タイソウシタ。
 *サフシロウハ オキテ、トウシタテスカ?
 "タフン カオオ アラツタテシヨウ。"



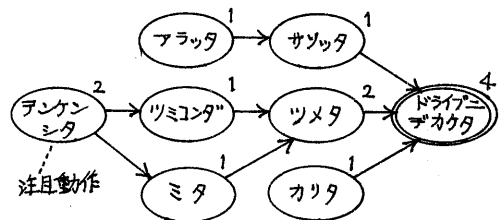
例1 後続行為を問う質問文

*JOHNハ ニューヨークへ ツイテ、ホテルへ ハイツテ、デモンシタ。
 *TOMハ ロンドンへ ツイテ、ユウシメント アツテ、デモンシタ。
 *FREDハ クルマカラ オリテ、ホテルへ ハイツテ、デモンシタ。
 *FRANKハ デモンシタエ、トウシタテスカ?
 "タフン ホテルへ ハイツタテシヨウ。"



例2 直前行為を問う質問文

*Aハ クルマオ デンゲンシテ、ニモツオ ツミコンテ、カソリンオ ツメテ、トライフニテカゲタ。
 *Bハ クルマオ デンゲンシテ、チスオ ミテ、カソリンオ ツメテ、トライフニテカゲタ。
 *Cハ クルマオ アラツテ、ユウシメント サソツテ、トライフニテカゲタ。
 *Dハ クルマオ カリテ、トライフニテカゲタ。
 *タロウハ ナンノタメニ クルマオ デンゲンシタテスカ?
 "タフン トライフニテカゲタテシヨウ。"



例3 目的を問う質問文

例2 の 3個の事実文

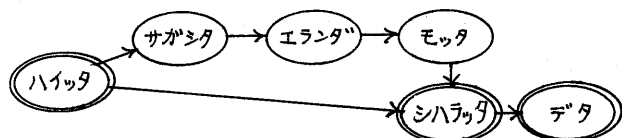
*TERRYハ トノヨウニシテ デモンシタテスカ?
 "タフン ニューヨークへ ツイテ、ホテルへ ハイツテ、デモンシタテシヨウ。"

図は 例2 の図を参照

例4 方法を問う質問文

*タロウハ ミセへ ハイツテ、タキケンオ シハラツテ、ミセオ テタ。
 *ヨウヤクシナサ1. シロウハ ミセへ ハイツテ、シナモノオ サカシテ、ホシモノオ エランデ、ソレオ モツテ、タキケンオ シハラツテ、ミセオ テタ。
 "タロウハ ミセへ ハイツテ、タキケンオ シハラツテ、ミセオ テタ。"

例5 要約の命令文



[4] まとめ

本報告は、人工知能システムの学習における implicit な知識獲得に関する手法を、その評価用システム ASPELAN-2 への具体的応用を核に論じたものである。このシステムの大きな特徴は、明示的に与えた知識をもとにして演繹的推論で応答できるような範囲を越えた質問に対しても、現在までに入力された多くの類似知識をもとに帰納的推論を働かせて応答範囲を拡大した点にある。但し、それは行動の系列という特殊な型の知識に限定してある。これにより、「解る」と「解らない」という二極の断定的応答の中間を埋めることがある程度可能となった。

本論の方式では、意味まで考慮した学習ではないために、内部に生成される概念は時として価値のないものが発生する可能性がある。この解決が今後の課題である。意味を考慮することにより、所有知識に矛盾するような概念の形成が防止されることが期待できる。

もう一つの問題点は、A スクリプトは動作系列のパターンではなくある動作から別なある動作への遷移を処理できるにすぎないという点である。この点、A スクリプトは Script の概念を完全に包括しているとはいえない。従って、パターンを考慮した A スクリプトの問題が残されている。

最後に、この研究をまとめるにあたり御指導頂いた山梨大学田中正次教授、御助言頂いた電気通信大学高澤嘉光教授、山梨大学重永実教授・森英雄助教授・御援助頂いた松原康夫助手に厚く感謝いたします。

〈参考文献〉

- (1) Schank, R.C.: SAM - A Story Understander; Research Report 43, Yale Univ. (1975).
- (2) 唐沢博: 意識作用を重視した自己概念化能力を持つ知識構造; 電子通信学会技術報告 PRL79-73, pp.1~6 (1980).
- (3) 藤崎博也: 言語的思考の過程の定式化; 人工知能に関する総括的研究, 昭和53年度文部省科学研究費補助金・総合研究(B)研究報告書 pp243~252
- (4) 唐沢博: 上向き仮説体系に基づく自己概念化モデル; 情報処理学会第21回全国大会 (1980) [予定].