

“Inductive Inference of Theories from Facts” の日本語訳

北上 始

ICOT 2nd Research Laboratory

〔概要〕

本報告は、1981年に、Ehud Y. Shapiro 博士がエール大学でまとめた研究報告書（事実による理論の帰納推論）に対する、ほぼ全訳（日本語訳）である。

彼は、“事実による理論の帰納推論”を、モデル推論と呼んでおり、その理論的な枠組みの定義からはじめ、帰納推論の一般的で増殖的なアルゴリズムを、与えている。

また、彼は、ロジックプログラムの虫取りという立場から、そのアルゴリズムをとらえ、そのアルゴリズムが、主に、(i) 矛盾探索(Contradiction Backtracking) アルゴリズムと、(ii) 仮説生成スキーマにより実現されると、みなすことができる、述べている。(i) は、虫の所在位置を見つけるアルゴリズムであり、(ii) は、虫の修正方法についての枠組みを与えるスキーマである。彼は、このスキーマを、精密化演算子(Refinement Operator) と呼び、そのいくつかの有益な性質を、述べている。

最後に、このモデル推論の実現手段として、ロジック的アプローチが成功した理由を、述べている。

なお、この日本語訳は、Preliminary Version であり、予告なしに修正される。

Inductive Inference of Theories
From Facts

February, 1981
Research Report 192
Ehud Y. Shapiro

Department of Computer Science
Yale University

日本語訳 Preliminary Version

初版 '83.1/20
第2版 '83.5/09
第3版 '83.7/29
第4版 '83.9/07
第5版 '84.1/11
第6版 '84.3/30

I C O T 2nd Research Laboratory
Hajime Kitakami (北上 始)

目 次

	ページ
1. はじめに	1
1. 1 いくつかのモデル推論問題	1
1. 2 科学的発見の問題について	9
1. 3 本論文の構成	10
2. 問題の定義	12
2. 1 モデル推論問題とアルゴリズム	12
2. 2 許容的な要請	14
2. 3 理論項の組み込み	16
3. モデル推論問題の複雑さ	17
3. 1 枚挙的なアルゴリズム	17
3. 2 充足的なアルゴリズムの能力限界	20
4. 決定的な実験による仮説の反ばく	22
4. 1 矛盾探索アルゴリズム	22
4. 2 科学的仮説の反ばく性	30
4. 3 応用：ロジックプログラムの虫取り	31
5. 反ばくされた仮説の精密化	40
5. 1 精密化演算子	40
5. 2 精密化グラフ	47
5. 3 最も一般的な精密化演算子	49
6. 一般的な増殖的モデル推論アルゴリズム	55
6. 1 あるモデルに関するソース集合の近似	57
6. 2 増殖的なアルゴリズム	58
6. 3 インプリメンテーション上の論点	63
7. 結論	64

概要

本論文は、モデル推論の問題とアルゴリズムについて述べられている。モデル推論問題は、科学者が直面している問題の抽象化である。科学者は、いくつかの限定されたわく組の下で、いくつかの領域を研究し、実験を行う。そして、その結果を説明できる理論を見つけようとする。これを抽象化すると、探索領域は、1 階述言語 $\mathcal{L}$ に対するいくつかの未知モデル M の領域であり、実験は、 M において、 $\mathcal{L}$ の構文に関する正当性をたしかめることである。我々の目的は、例文が全て真なる仮説の集合を見つける事にある。

本論文の主な成果は、モデル推論問題を解くための、一般的で増殖的なアルゴリズムである。モデル推論問題は、推測と反ばくからなる Popper の方式 [Popper 59, 68] に基づく。このアルゴリズムは、ある種別の複雑なクラスのモデルで、任意のモデルをある極限 [Gold 67] で同一視するように示される。また、本アルゴリズムは、その中で最も有効で、かつ、いくつかの具体的領域で首尾よくインプリメントされるのに十分に柔軟なものである。

本モデル推論アルゴリズムは、二つの調整可能なパラメータをもつ。一つは、仮説の構造をいかに複雑にするかであり、他の一つは、仮説からどの程度複雑な導出をさせるかである。両者は、同時に、このアルゴリズムの制限下で、帰納的に推論されるあるクラスのモデルを決定している。一方では、両パラメータは、このモデル推論アルゴリズムが、任意に決められた帰納関数で構成されたモデルの極限で同一視されるように定められる。他方ではそれらは、近似的にインプリメントされたアルゴリズムが、実行時間内で、少ない事実から具体的なモデルの公理化を推論するように定められる。

このモデル推論システムは、ホーン形式の理論を推論するように限定されたモデル推論アルゴリズムを、基礎にしている。本システムは、プログラミング言語 Prolog [Pereiea et al. 78] で、インプリメントされた。例として、算術式の領域で、本システムは、36個の事実から Figure 1-1 の下に書いた公理の集合を、27 sec (CPU) で導いた。本システムは、終端をもつ稠密な半順序に対して、ある公理化を発見した。それは、append, reverse そして、夏の課題論文で述べていた [Summers 76] 例のほとんどに対し首尾よく、ロジックプログラムを合成した [Kowalski 79a]。

また、boolean formulas, binary tree, inclusion, binary tree isomorphism その他についても満足のいくロジックプログラムを合成した。

この一般的アルゴリズムの一部分として、矛盾を探索するアルゴリズムが、発見された。このアルゴリズムは、ある矛盾が推測と事実に生じるときはいつでも

本アルゴリズムは、この矛盾の原因を後もどりで見つけられる。例えば、この矛盾の原因は、ある偽の仮説である。また、本アルゴリズムは、ある反例を作成する事により、仮説の虚偽性を示すことができる。その様なアルゴリズムの存在は、科学的理論 [Harding 76] の反ばくに関する哲学的議論に見られる。また、ホーン理論の制限は、ロジックプログラムの虫取りに対し、実用性を促すかもしれない。

1. はじめに

本論文は、以下の形の問題を扱う。

我々は、あるたしか領域から実際の知識を得ることができるとする。

ただし、この領域は、いくつかの未知のルールにより支配されており、これらのルールを表現するのに適切な言語が与えられていると仮定する。

我々は、いかにして、この実際の知識からこれらのルール又は理論を見つけるのだろうか？

これは、ある決められた概念の中で、いくつかの領域を研究している科学者が直面している問題である。即ち、彼らは、実験を行ったり、領域内で真なる定理を見つけようとしたり、それらの結果を説明しようとしたりしている。モデル推論問題は、この立場の抽象化である。本節では、例でモデル推論問題を説明する。

これらの例は、モデル推論問題が解けるアルゴリズムによる可能な応用を提案している。この研究の関連性を、科学的発見の問題で議論し、本論文の構成を与える。

1. 1 いくつかのモデル推論問題

Figure 1-1にモデル推論問題の例を示す。

探索領域は、整数であり、与えられた1階述領域の言語は、定数0、後続関数 X' 、そして以下の3つの述語を含む。

$X \leq Y$ …… X は Y より小か等しい

$\text{plus}(X, Y, Z)$ …… $X + Y = Z$

$\text{times}(X, Y, Z)$ …… $X \times Y = Z$

これらの関係が、具体的数値で成立するかどうか確かめてみよう。即ち、 M で真なる $0 \leq 0''', \text{plus}(0', 0', 0''), \text{times}(0'', 0'', 0''')$ としてのground atom (variable-free)を確認することができる。この立場では、モデル推論問題は、構文の有限集合を見つけることにある。この構文は、正しい算術式であり、すべて真のground atom を導いている。Figure1-1 は、その構文の一組を示している。逆の矢印“ $\leftarrow$ ”を“ $\sim$ により証明される”という記号とする。構文 $P \leftarrow Q \& R$ は“ P は Q と R の結合により算術される”と読む。

Figure 1-1: 算術式の推論

領 域: 整数

言 語: 0 — ゼロ。
 X' — X の後続要素。
 $X \leq Y$ — X は、 Y 以下である。
 $\text{plus}(X, Y, Z)$ — $X + Y = Z$ 。
 $\text{times}(X, Y, Z)$ — $X \times Y = Z$ 。

事実の例: $0 \leq 0'$ は、真である。

$\text{plus}(0, 0', 0)$ は、偽である。

$\text{times}(0'', 0'', 0''')$ は、真である。

理 論: $X \leq X$

$X \leq Y' \leftarrow X \leq Y$

$\text{plus}(X, 0, X)$

$\text{plus}(X, Y', Z') \leftarrow \text{plus}(X, Y, Z)$

$\text{times}(X, 0, 0)$

$\text{times}(X, Y', Z') \leftarrow \text{times}(X, Y, W) \ \& \ \text{plus}(W, X, Z)$

このモデル推論問題を解くために、関数と述語の性質をすべて見つけ出す必要がないことに注意しよう。特に、上記、理論は、加法の連結性と順序集合の推移性を含んでいない。もし、 T が、 M で真なる L のground atom を全ての導出する M で真なる構文であれば、 T は M に関し atomic-完全公理化 (atomic-complete axiomatization) であると呼ばれる。Figure 1-1 の構文の集合は、算術式に関し atomic-完全公理化である。それは、36の事実から27 CPU secondで、モデル推論システムにより推論された(付録I参照)。

我々は、1階述語の言語 L で2つの形の構文を区別する：それらは、実験結果の記述に関連する観測文、及びそれらの結果の説明として与えられる仮説である。科学者の探索領域は、 L に関するいくつかの未知モデル M の領域である。この領域で行なわれている実験は、 M で真なる観測文を調べることにある。事実は、その様な実験の結果の陳述である。モデル推論問題は、以下の科学的発見の問題に関連している。

いくつかの未知モデル M の実験を行う事ができる能力を与えると、 M で真なる仮説の有限集合が見つけれられる。ただし、その仮説は、すべての観測文を真にする。

仮説に使われる言語を固定するという過程は、科学的発見の主要問題の一つを回避することになる。この問題は、新しい概念スキーマの発明である。Kuhnian の用語 [Kuhn 70] では、モデル推論問題は、新しい概念スキーマを見つけるよりは、“normal science” のパズル解決の方式に似ている。新しい概念スキーマは、“凡例置き換え”過程という特性がある。

他のモデル推論問題が、Figure 1-2に示されている。この例では、探索領域は、2進ストリングの集合である。また、この1階述語の言語は、単項述語 $in(X)$ 、2つの後継関数 $0(X)$ 、 $1(X)$ 及び、null string を示す定数 $\wedge$ を、含む。この言語では、項 $(1(0(0(1(\wedge)))))$ を、ストリング 1001 で示し、 $1(0(X))$ の前置オペレータをストリングで示す。かっこは、通常除かれる。我々は、具体的なストリングが、いくつかの未知の集合の中に存在するかどうかを確かめられると仮定する。また、我々のゴールは、この集合中のすべてのストリングに対し、 $in(S)$ を証明し、他の任意の2進ストリングに対し、 $\sim in(S)$ を証明する。正しい構文の有限集合を見つける事にある。

Figure1-2 : ある形式的言語の推論

領 域: 2進ストリング

言 語: Λ - 空ストリング。
 $0(X)$ - X に0を結合する。
 $1(X)$ - X に1を結合する。
 $in(X)$ - X は、偶数の0と1の偶数の1をもつ。

事実の例: $in(0011)$ は、真である。
 $in(011)$ は、偽である。
 $in(01011)$ は、偽である。
 $in(010100)$ は、真である。

理 論: $in(\Lambda)$
 $in(00X) \leftarrow in(X)$
 $\sim in(0X) \leftarrow in(X)$
 $in(11X) \leftarrow in(X)$
 $\sim in(01X) \leftarrow in(X)$
 $\sim in(1X) \leftarrow in(X)$
 $\sim in(11X) \leftarrow in(0X) \& \sim in(X)$
 $\sim in(01X) \leftarrow in(0X) \& \sim in(X)$
 $\sim in(1X) \leftarrow in(0X) \& \sim in(X)$
 $\sim in(10X) \leftarrow in(1X) \& \sim in(0X) \& \sim in(X)$
 $\sim in(00X) \leftarrow in(1X) \& \sim in(0X) \& \sim in(X)$
 $in(00X) \leftarrow \sim in(1X) \& \sim in(0X) \& \sim in(X)$
 $\sim in(00X) \leftarrow \sim in(1X) \& \sim in(0X) \& \sim in(X)$
 $in(10X) \leftarrow \sim in(1X) \& \sim in(0X) \& \sim in(X)$

もし、未知の集合が十分に簡単であれば（もっと正確にいうと、それは規則的である… [Angluin & Hoover 80]）、この言語の中に、有限集合の公理が存在する。もし、 T が M 上の構文で真であれば^{*}、 T は M のground-完全公理化 (ground-complete axiomatization) であると呼ばれている。Figure 1-2 は、偶数個の0と1をもつストリングの集合に対するground-完全公理化を示す。この特殊な公理は、Angluin とHooverのアルゴリズムによるこの集合に対し、4状態のアクセプターから作成された。モデル推論システムの古い判は、64個の事実から3分CPUでこの集合のground-完全公理化を推論した。本システムにより発見された公理は、各条件に1つのアトムを含む。その条件は、上記のいくつかの公理の2と3に対立する。本アルゴリズムのふるまいは、任意の規則集合が、ground-完全公理化をもつことを提案している。この結果は、どこかよそで [Angluin & Shapiro 81] 記述されるだろう。

同じ規則集合のatomic-完全公理化を以下に示す。それをFigure 1-2の公理化と比較せよ。Figure 1-2の公理は、1つのモデルに関し、atomic-完全公理化と、ground-完全公理化の区別を明確にできる。モデル推論システムは、28 CPUでこの公理に追いついた。同時に、35個の事実を読んだ。

```

in (Λ)
in (00X) ← in (X)
in (11X) ← in (X)
in (010X) ← in (1X)
in (011X) ← in (0X)
in (100X) ← in (1X)
in (101X) ← in (0X)

```

規則性をもった帰納的推論問題は、GoldとAngluin [Gold 78, Angluin 78]らによって研究された。それは、モデル推論アルゴリズムとシステムを開発する為のある動機づけとテストケースとして役立っている。

我々のフレームワークに自然に一致するモデル推定問題の第3番目は、自動プログラミングの領域が取り上げられる。また、それは、いくつかの例によるプログラム合成を意味する [Green et al. 74, Summers 76, Summers 77, Bierman 78]。その仕事は、入出力のふるまいから、帰納的にプログラムを推論することにある。この仕事は、モデル推論問題として言い換えることができる。また、この場合、推論されるプログラムは、ロジックプログラムである。

* 任意のground atom $P \in L$ に対し、 P が M で真であり、それ以外で $\sim P$ であるならば、 T は P を証明する。

Figure1-3 ロジックプログラムの推論

領 域：リスト

言 語： [] - nil
 [X | Y] - XとYの“cons”
 append (X, Y, Z) - XをYに加えるとZになる。
 reverse (X, Y) - Xは、Yの反転である。

事実の例：append ([], [a], [a]) は、真である。
 append ([a, b], [c, d, e], []) は、偽である。
 append ([a, b], [c, d, e], [a, b, c, d, e])
 は、真である。
 reverse ([a], [b, a]) は、偽である。
 reverse ([a, b, c], [c, b, a]) は、真である。

理 論 append ([], X, X)
 append ([A | X], Y, [A | Z]) ← append (X, Y, Z)

 reverse ([], [])
 reverse ([A | X], Y) ← reverse (X, Z) &
 append (Z, [A], Y)

ロジックプログラムは、ホーン節の集まりであるが、その形式は、 $P \leftarrow Q_1 \& Q_2 \& \dots \& Q_n$ である。 P と Q は、アトムである。この構文は、“ P は、 Q の結合により証明される”と読み、手続的には、“ゴール P を満たすためには、ゴール $Q_1, Q_2, \dots, Q_n$ を満たすことである”と解釈される。ホーン節の集まりは、手続的解釈を使ったプログラムとして実行される。通常、そのようなプログラムには、効率を考慮した、制御情報が記述されるべきである。ロジックプログラムの紹介として、[Kowalski 79a]を参照せよ。

また、ロジックプログラムの操作的でかつ表示的意味についての議論に対しては [Van Emden & Kowalski 76]を参照せよ。アルゴリズムのロジックと制御の間の区別の苦心については、[Kowalski 76b]を、参照せよ。

Prologマニュアルとしては、プログラム言語に基づくロジックの具体的インプリメントの詳細 [Pereira et al. 78]がある。

我々の設定では、もし、仮説に使われる言語が、ホーン節に、制限されれば、atomic-完全公理化は、 L の述語を計算するためのロジックプログラムになる。Figure1-3 は、そのようなプログラムの例である。この例では、言語 L は、2変数関数 $[X|Y]$ 、定数 $[]$ 、そして、二つの述語 $\text{append}(X, Y, Z)$ と $\text{reverse}(X, Y)$ を含んでいる。この言語に対するモデル M は、次のように定義される： M の要素は、 $[X|Y]$ と $[]$ から構成されるリストである；そのアトム $\text{append}(X, Y, Z)$ は、 M で真である。この場合、リスト Z はリスト X に Y を付加したリストに等しい。

例えば、 $\text{append}([a, b, c], [d, e], [a, b, c, d, e])$ ；アトム $\text{reverse}(X, Y)$ は、 M で真である。 Y は、 X を反転したリストに等しい。Figure1-3 のホーン節は、このように定義されたatomic-完全公理化である。ふり返ると、これらの節は、 append と reverse に対してのPrologプログラムである。例えば、リスト $[a, b]$ に $[c, d, e]$ を付加した結果を計算する為に、 $\text{goal} \leftarrow \text{append}([a, b], [c, d, e], X)$ は、上記プログラムをロードされたPrologインプリンタに与えられるが、 $X = [a, b, c, d, e]$ が、返される。モデル推論システムは、36の事実から11 CPUで append に対するロジックプログラムを合成し、13個の事実から6 CPU で reverse に似たProgram を合成した。ロジックプログラムのさらに進んだ例は、本論文の至る所で与えられている。また、Appendix 1に、例による、いくつかの合成の説明が見られる。

Concept-learning tasksの立場から、A1のいくつかの研究は、モデル推論の問題として言替えられる。例えば、archesとnon-archesの記述から、arch [Winston75] の概念を学習する問題は、以下のように言い換えられる：探索領域は、複数ブロックをつみ重ねるすべてのオブジェクトである。言語は、 $\text{block}(X)$ 、 $\text{column}(X)$ 、 $\text{on}(X, Y)$ 、 $\text{arch}(X, Y, Z)$ のような述語を含む。

帰納推論問題は、構文の有限集合を見つける事にある。その構文は、合成オブジェクトがarchesである事を決めることができる。archesは、archesとnon-archesの例により与えられる。そのような構文の集合は、Kowalski [Kowalski 79a] により提案された：

$$\text{arch}(X,Y,Z) \leftarrow \text{block}(X) \& \text{column}(Y) \& \text{Column}(Z) \& \text{on}(X,Y) \& \text{on}(X,Z)$$

$$\text{column}(X) \leftarrow \text{block}(X)$$

$$\text{column}(\text{stack}(X,Y)) \leftarrow \text{block}(X) \& \text{column}(Y) \& \text{on}(X,Y)$$

$$\text{on}(X,\text{stack}(Y,Z)) \leftarrow \text{on}(X,Y)$$

この構文の集合と、次式から、

$$\text{block}(a) \quad \text{on}(a,b)$$

$$\text{block}(b) \quad \text{on}(b,c)$$

$$\text{block}(c) \quad \text{on}(a,d)$$

$$\text{block}(d) \quad \text{on}(d,e)$$

$$\text{block}(e) \quad \text{on}(e,f)$$

$$\text{block}(f)$$

$\text{arch}(a, \text{stack}(b,c), \text{stack}(d, \text{stack}(e,f)))$ を証明できる。一般に、あるシーンがarchの概念のあるインスタンスになっているかどうかは、各公理から決定される。また、そのような公理を要求することにより、この概念を学べることは、妥当であるように見える。

1. 2 科学的発見の問題について

モデル推論問題に対する、任意の興味ある解は、ある予備的問題と折合いが付かねばならないが、それは、帰納の問題である。この問題は、哲学者David Humeにまでさかのぼる。帰納の問題を示すと、有限の量をもつ実際のデータは、もちろん無限の結果をもつ理論を決して確立することができないといえる。この問題の魅力的な解は、Karl R. Popper (哲学者) によって提案された [Popper 59]。Popperは、実際のデータによる一般的科学理論が、不確定であることを述べたHumeの考えを認め、科学理論の推測的立場が認識される事を提案した。科学の有効性は、理論がまちがいであれば、実際にそれは拒絶されるだろうという事実に基づいている。一般に、その理論が正しい限り、もちろん改良されない。

帰納の問題に対するPopperの方式は、科学行動に対し、いくつかの手続き的な忠告を生ずる。Popperは、実験行動が、提案された理論を拒否する力をもつ新しいデータを集める方向にある間、理論的科学行動が、簡単で容易に反ばく的な理論を作成する方向に動くべきであると提案した。ただし、この理論は、既知のデータを、説明している。Popperは、理論と実験の間、すなわち、推測と反ばくの間の相互作用は、真実へのある種の収束現象を導びいてるのかもしれないと主張した。真実らしさ [Popper 68] の概念は、競争的な科学理論がもつ “もっともらしさ” を測定したり、比較したりする事により創案された。非公式には、 T_1 が T_2 と同じくらい多くの真の観測文を導びき、 T_2 よりも多くの偽となる観測文を導びかないとき、仮説 T_1 の集合の真実らしさは、 T_2 の真実らしさより大きいとか等しいかである。それらの真実らしさを測定しする理論的数値を系統的に割り付ける事ができる。便宜上、この値は、0と1の間の範囲とする。この様な測定下では、すべて真の観測文を導びく真の理論は、真実らしさ1をもち、矛盾しかつ反復的な理論は、真実らしさ0をもつ。科学的発見の過程は、その理論の真実らしさを増加する永続的試みであるとして、見る事ができる。

真実らしさの概念は、諸理論を競争比較するのに使われる。しかしながら、競争的な科学的発見の方法論又は帰納的推論アルゴリズムのいくつかを評価・比較するために、その着目点が単一理論よりもむしろ科学的な発見プロセス又はアルゴリズムに存在するように思われる。帰納的推論アルゴリズムの能力を評価、特徴づける方式は、Gold [Gold 65] によって提案され、極限における同一視と呼ばれている。この概念の理論的根拠は、Popperのそれと良く似ている：一つのモデルについての有限の事実の数は、一般的に、すべての可能なモデルの間では、ユニークに決められない、また帰納的推論アルゴリズムおよび科学者は、いつも有限の事実を使って推測を行うので、それらのどちらもあやまちをおかざるを得ない。

ある意味では、これはPopperの命題に関する難問である。帰納的推論アルゴリズムで期待できる事は、次のようなことである。そのモデルについて有限の事実をたしかめ、有限のまちがった推測を行った後、そのようなアルゴリズムは、正しく、真なる仮説の集合を推測するだろう。そして、その仮説は、全て真なる観測構文を導くことになるだろう。そのような場合、我々は、このアルゴリズムが極限でそのモデルに同一視されると言う。その様な仮説の集合は、真実らしさが1であるが帰納的推論アルゴリズムは、一般に、その様な仮説の集合を実際に見つけたかどうかを決められないことに注意しよう。極限における同一視の概念は、帰納推論での帰納理論と複雑さの理論の研究で実りが多いことを証明した [Gold 67, Blum & Blum 75, Case & Smith 81]。

1.3 本論文の構成

第2節では、モデル推論とアルゴリズムの概念をもっと厳密に定義する。このとき、[Blum & Blum 75]の方式に従い、いくつかの複雑なモデル推論問題を考察する。第3節では、いくつかの帰納関数 h に対するいくつかの h -easyモデルのクラスを定義する。またある h について、極限で、任意の h -easyモデルを同一視できる枚挙型アルゴリズムを述べる。帰納的推論アルゴリズムが簡単で多くの要求を満足、即ち、その現在の推測では、いつも現在知られている事実を説明できるならば、そのようなアルゴリズムは、極限で h -easyモデルとだけに同一視できる。これらの結論は、任意に多くの帰納推論アルゴリズムの能力に上限を定めている。また、この結論は、この上限が、現実的である事を示している。

この一般的で複雑な理論的問題解決の方式は、任意の計算可能なモデルに適用できるがロジックの特別な性質を使っていない。この解析による動機づけを行うために、あるロジック指向のアルゴリズムと概念を開発している。このアルゴリズムと概念は、ロジックの意味と構文上の性質の比較的有利な点を取り出した比較的効率的で増殖性のある推論アルゴリズムの構成を提供できる。この増殖的なモデル推論アルゴリズムを開発するために、以下の2つの問題に注目する。

1. 推測が非常に強くなったときに、いかにその推測を弱めるか？
2. 推測が非常に弱くなったときに、いかにその推測を強めるか？

仮説が M 上で偽なる観測文を証明するならば、モデル M に関して非常に強いと言う。仮説が M 上で真なる観測文を証明することがないならば、非常に弱いと言う。

矛盾探索アルゴリズム (The Contradiction Backtracing Algorithm) は、最初の問題に答えることを試みている。これは、第4節で述べる。このモデルで、ground atom が真実であることを判断する有限回の実験を行うことにより、このアルゴリズムは、推測と事実の間の矛盾の根源を逆にたどり、偽の仮説を選択することができる。

このアルゴリズムにより行なわれる種々の判断は、決定的な実験としての科学哲学で知られている。それらの重要性は、多くの方法論により認識できるが、偽の仮説の選択を保証する順序だてられたアルゴリズム的方式は、新規性をもつ。そのようなアルゴリズムの存在は、明らかに Duhem [Duhem 54] の主張と矛盾する。彼の主張は、その可能性を、簡単に否定している。この主張は、科学的理論の非反ばく性を扱っている Duhem-Quine の命題 [Harding 76] として、あとで何が知られるようになったかという事実から述べられた。矛盾探索アルゴリズムとそれを混合する一般的な帰納推論アルゴリズムの存在は、反ばく性の論争に関する哲学的論議を一新するかもしれない。矛盾探索アルゴリズムは、帰納推論の領域を越えた応用をもつ；我々は、ロジックプログラムの虫取りとしてその応用を提案している。

一度、推測の中で偽なる仮説を発見すると、とられるべく動作は、この仮説をそこからとりのぞくことである。しかしながら、これにより得られる推測は、すでに真である観測文を証明するのには弱すぎるかも知れない。この様な状況においては、第2の問合せが発生する。それを強くするために、その推測に、いかに新しい仮説を付け加えるかということを提案している。

精密化演算子(refinement operator) は、調整可能なこの推論アルゴリズムのパラメータである。さらに一般的な精密化演算子は、たとえ効率が小さいアルゴリズムになっても、より強力になる。あるクラスの仮説に対し、精密化演算子の完全性が定義され、種々のクラスの仮説を完全にするいくつかの精密化演算子が述べられる。非常に一般的な精密化演算子が定義され、任意の1階述語の言語に対して完全であることが示される。

これらの二つの問題の解決に基づいて、一般的で増殖的なモデル推論アルゴリズムが、第6節で開発されている。また、任意のh-easyモデルを、ある極限で同一視することも証明されている。また、いくつかのインプリメンテーション上の論点が論議される。本論文に述べられている陳述は、なぜこの要素を研究するのかを示している。また、付録1は、そのアルゴリズムをインプリメントしたモデル推論システムの効率が述べられている。本システムは、次の論文でもっと十分に論議されるだろう。

2. 問題の定義

本節では、もっと正確にモデル推論問題を述べ、モデル推論アルゴリズムの概念を定義する。ここでは、観測と仮説言語についての許容的な要請を紹介するが、この要請は、理論が事実によって反ばくされるべきであるというPopperの要請を現わしている。また、理論項の概念を含むモデル推論問題への拡張を議論する。

2. 1 モデル推論問題とアルゴリズム

与えられる1階述語の言語 L は、有限に多くの述語と関数記号をもつ([Robinson 65]で定義されたような)節形式である。 L の構文は、次の形をしている。

$$\{P_1, P_2, \dots, P_j\} \leftarrow \{Q_1, Q_2, \dots, Q_k\} \quad j, k \geq 0$$

ここで、 P 'sと Q 'sは、アトムである。この概念は、 P 'sが正のリテラルであり、否定 Q 'sが節の負のリテラルであるという標準節形式の概念と、等価である。この構文は、 Q 'sの結合が P の選言を導出すると解釈される。 Q 'sは構文の条件と呼ばれ、 P 'sは、その結論と呼ばれる。空の条件は、少なくとも結論の中の1アトムが真であることを示す。空の結論は、条件の中の全てのアトムが真であることを示す。 $\square$ は、空の構文であり、 L のモデルで偽を意味する。1アトムだけをもつ構文は、単一構文と呼ばれる。単一構文 $\{P\} \leftarrow$ とアトム P を区別しない。また変数の名前をつけるえると等しくなる構文も区別しない。標準的な節の概念では、構文に存在するすべての変数は、黙目的に全称的な量が定められている事に注意せよ。

L の構文において、2つの部分集合を区別する：それは、観測文 L_0 と仮説文 L_h である。 $\square \in L_0 \subset L_h \subset L_1$ であり、この両集合は、実質的に識別可能である。観測文は、 L_0 の構文であり、仮説は、 L_h の構文である。探索領域は、 L の中のあるモデル M である。また、観測文 α が与えられると、もし α が M で真ならば“true”、それ以外は“false”を返すというある仕かけがあるとする。その仕かけは、 M にとってのOracle(神託)とよばれる。そのOracleに入力を与え、その答を読む操作は、 M の実験と呼ばれる。領域 M についての事実は、形式 $\langle \alpha, V \rangle$ の組である。この形式は、 $\alpha \in L_0$ を、観測文とし、 $V \in \{\text{true}, \text{false}\}$ を M における α の真偽とする実験の結果である。モデル M 上で真なる観測文の集合は、 L_0^M で示される。ある完全な誘導手続きを仮定し、 P が T から導出される(されない)事示すために $T \vdash P$ ($T \not\vdash P$)を使う。構文 S の集合から $T \vdash S$ であるためには、すべての $P \in S$ に対し $T \vdash P$ でなければならない。それ以外は、 $T \not\vdash S$ である。

定義2.1: L_0 と L_h を L の部分集合とし、 M を L のモデルとする。ただし、 $\square \in L_0 \subset L_h \subset L$ である。構文 $T \subset L_h$ の集合が、 M で L_0 -完全公理化であるための必要十分条件は、 T が M で真であり、 $T \vdash \square$ となる事である。

ここで、さらに厳密なモデル推論アルゴリズムの問題を定義できる。

1 階述語の言語 L と、前述で定義された2つの部分集合（観測言語 L_0 と仮説 L_h ）を与え、 L に関して、ある未知モデル M に対する Oracle（神託）が与えられているとする。モデル推論問題は、 M に関し、有限の L_0 -完全公理化を見つけることにある。

モデル推定問題を解くアルゴリズムは、モデル推論アルゴリズムと呼ばれる。Gold と Blums [Gold 65, Blums & Blums 75] の定義と合わせる為に、そのようなアルゴリズムの概念をもっと正確にする。

モデル M の枚挙は、無限の並び $F_1, F_2, F_3, \dots$ である。ここで F_i は、 M についての事実である。また、 L_0 のどの構文 α も、 $i > 0$ とする事実 $F_i = \langle \alpha, V \rangle$ が存在する。モデル推論アルゴリズムは、与えられた観測言語 L_0 に対し、モデルの枚挙を読み込み、即ち、1度に1つの事実を読み、ときどき、仮説言語 L_h の構文の有限集合（アルゴリズムの推測とよばれる）を出力するアルゴリズムである。

モデル推論アルゴリズムが実際にいくつかの推論を出力し、異なる推論を二度と出力しなければ、モデル推論アルゴリズムは、モデルの枚挙に与えられた極限に収束していると言われる。モデル推論アルゴリズムが極限でモデル M に同一視されると言われる為には、 M のすべての枚挙上で、そのアルゴリズムが M に関し L_0 完全公理化となる推測に収束することである。モデル M の枚挙を与えると、 M に対して Oracle を模擬できる。また、逆もまた真である。それは、我々がこの問題の抽象化の中でこれらの概念の区別を行っていないことによる。しかしながら、以下で示すが、実験を行う能力は、推論過程で顕著な高速化を結論する。それは、具体的アルゴリズムは、通常、 M に対する Oracle が与えられていると仮定されているからである。

2. 2 許容的な要請

モデル推論アルゴリズムにより解かれるモデル推論問題の為に、許容的な要請が L_0 と L_h に行なわれなければならない。それは本質的に、 L_0 が、任意の L の理論を反ばくするのに十分な情報を含んでいるということである。これは、 L_0 が完全になる事を試行している。この許容的な要請は、Popperの方法論的要請を表わしている。彼によれば、理論は、事実により、反ばくできるといっている。

定義 2. 2: L_0 を観測言語とし、 L_h を仮説言語とする。

$\langle L_0, L_h \rangle$ が許容的であるためには、 L のすべてのモデル M とすべての $T \subseteq L_h$ に対し、 $\{\alpha \in L_0 \mid T \vdash \alpha\} = L_0^*$ で、 T が M で真になる T を導出することにある。

許容的な要請により、 L_0^* を証明する L_h で偽となる理論は、観測言語 L_0 で、虚偽性に対する証拠をもつ。これにより、 $T \vdash L_0^*$ かつ T が M で偽であれば、 M で偽となる $\alpha \in L_0$ が存在する。ここで、 $T \vdash \alpha$ である。この性質が L_0 と L_h の許容性を導出することを調べるために、この性質が成立することを前提にして $T \subseteq L_h$ が M で偽であるとしてみよう。このとき、 M で偽なる $\alpha \in L_0$ に対して $T \vdash L_0^*$ 又は $T \vdash \alpha$ のどちらかが成立する。両方とも、許容的な要請を満足するが、それらは、 L_0 と L_h の許容性を導出する。許容的な要請は、観測言語と仮説言語を一緒にする。また、それらの表現力の差は、そのような方法で縛られるべきなので、成功するすべての理論（例えば、真の観測文を全て導き、1つも偽の観測文を導かない）は、また真になるべきである。許容的な要請から次の事が言える。もし、 L_0 と L_h が許容されると、 $L_0 \subseteq L_0' \subseteq L_h' \subseteq L_h$ をみたすすべての L_0' と L_h' は、また許容される。

観測と仮説について、興味ある2組の言語がある。1つの組は、観測言語が L のground atomsであり、仮説言語が L のホーン構文となる組である。この組に対して、 M の L_0 -完全公理化は、atomic-完全公理化と呼ばれる。他の1つの組は、観測言語が L のground unit 構文であり、仮説言語が L となる組である。この組に対し、 L_0 -完全公理化は、ground-完全公理化と呼ばれる。これらの組の許容性を見る前に、組 $\langle \text{atom of } L, L \rangle$ が非許容的である事を示す。陳述的例は、十分ある。 L は、3つのアトム P, Q, R を含むとする。 P と Q が偽で、 R が真になる L のモデル M で選ぶ。このとき、 $L_0^* = (\{R\} \leftarrow)$ であり、 $T = (\{R\} \leftarrow, \{P, Q\} \leftarrow)$ が L_0 を制限した結果は、 L_0 に等しいが、 T は、 M で偽である。この例と、次の二つの定理は、ホーン節とatomic-完全公理化の親密な結合の中でいくつかの洞察力を与える。

また [Van Emden & Kowalski 76] を、この結果の議論として参照されたい。

定理 2.3: L を 1 言語として

$L_0 = \{L \text{ の ground atoms} \}$, $L_h = \{L \text{ の ホーン 構文} \}$ とする。このとき、組 $\langle L_0, L_h \rangle$ は、許容的である。

証明: $T \subseteq L_h$ は、 M で偽とし、 $T \models L$ を仮定する。 L_0 は、 T の虚偽性に対する証拠をもつとする。

T は、偽なので、 M で偽なる構文 $P = \{P\} \leftarrow \{Q_1, Q_2, \dots, Q_n\} \in T$ が存在するが、これは、 P が M で偽なる (定義 4.1 を、参照) ground instance $P' = \{P'\} \leftarrow \{Q_1', Q_2', \dots, Q_n'\}$ をもつことを導く。 $T \models L$ なので、 $T \models Q_i'$ (for $1 \leq i \leq n$) ということになる。

$P \in T$ から $T \models P'$ が導かれる。 P' は、偽の観測文である。そのために、この組は、許容される。 ■

定理 2.4: L を 1 階述語の言語とする。

$L_0 = \{L \text{ の ground unit 構文} \}$, $L_h = L$ のとき、組 $\langle L_0, L_h \rangle$ は、許容的である。

証明: 類似的証明を行う。 $T \subseteq L_h$ が、 M で偽であり、 $T \models L$ を仮定する。 T は、偽であるので、 M で偽の構文

$P = \{P_1, P_2, \dots, P_m\} \leftarrow \{Q_1, Q_2, \dots, Q_n\} \in T$ が存在する。これは、 M で偽の ground instance

$P' = \{P_1', P_2', \dots, P_m'\} \leftarrow \{Q_1', Q_2', \dots, Q_n'\}$ をもつ。この構文は、 $Q_1', Q_2', \dots, Q_n'$ が、真で、 $P_1', P_2', \dots, P_m'$ が、偽のときだけ偽になる。 $P \in T$ なので T は、矛盾する。これにより、 T は、特殊な場合の $T \in \square$ を含めて、 L のすべての構文を導出する。 $\square$ は、偽の観測文である。 ■

2.3 理論項の組み込み

許容的な要請は、モデル推論問題の解決性に対して、必要ではあるが十分条件でない。例えば、我々の領域が算術式の標準的モデルであれば、1階述語の言語 L は、以前に述べた算術式の述語であり、 $L = L_h = L_0$ を選ぶと、組 $\langle L_0, L_h \rangle$ は、許容的である。ゲーテルの第2不完全定理は、 L の中に有限の矛盾した公理は存在しないといっているが、その公理は、算術式として真なる L のすべての構文を導出する。

算術式の場合、仮説言語 L を高めることは、このモデル推論問題を解くことにならない；しかしながら次のような事情がある。問題の非解決性は、仮説言語の表現力を弱める事になる。例えば、 L は、単一述語記号 $\text{in}(X)$ をもつFigure1-2に述べた2進ストリングの言語であるとする。また、推論される未知の集合は、0, 1の回分であるとする。ストリング S が、それ自身の逆であれば、 S は、回分である。

回分の集合は、秩序立っていないので、Angluin とHoover[Angluin & Hoover 80]は、すべての真なる回文 S により $\text{in}(S)$ を導出する有限の構文集合が存在しないと結論している。

直観的に言えば、回文の概念を理解するためには、まずストリングの反転の概念を知る必要がある。また、ある方法で、この概念を手にしたたり、その操作をシュミレートしなければ、回文と非回文の例から回文の概念を学習できない。仮説の構文を目的として使えるが、その世界で直接観測できない概念を示す項は、科学の理論がもつ哲学と呼ばれている。たとえば、理論項として append と reverse の述語を仮説言語 L_h に含めても、回文のatomic完全公理化を推論する仕事が、すぐに解けることにならない。しかし、容易に以下の公理を満たす。

$$\text{in}(X) \leftarrow \text{reverse}(X, X)$$

また、 reverse は、Figure1-3に示すように append を使って公理化され得る。

しかしながら、我々がその世界で見えるすべてのストリングが、いくつかの未知の集合の入出力にあるストリングであるとするならば、 append と reverse に関する公理を例から推論するための我々の方式は1つも存在しない。それを設定するために、我々の過程は、少なくとも append と reverse に関してもくろまれた解釈が理論的概念として推論アルゴリズムで知られているか構成されているかということである。モデル推論問題に理論的述語の概念を組み込ませる方法は、推論アルゴリズムとして知られているいくつかの固定的な解釈をもつ理論的概念を要求する事にある。この解釈は、未知モデル M で成立すると仮定する。ここで、このアイデアを公式に展開しないだろう。しかし、その関連は、科学的過程の反ばく性の論議で行なわれる(4.2節)。それは、モデル推論問題の紹介になる。

3. モデル推論問題の複雑さ

モデル推論問題に密接に関係した2つの帰納的推論問題が、文献の中に有益な注意を与えてきた：1つは、再帰的に数え上げられる集合を数え上げる為のプログラムを帰納的に推論するための問題である。この集合には、各要素に、その集合の入/出力情報が与えられている [Gold 65, Gold 67, Klette and Weihagen 80]；他は、帰納関数を計算する為のプログラムを帰納的に推論する問題である。その関数 [Gold 65, Blum & Blum 75] の値が、与えられている。Gold [Gold 65] は、再帰的に、数え上げられる集合または、すべての帰納関数に対して、これらの問題を解くための一般解が、存在しないことを示した。Blumたち [Blum & Blum 75] は、後半の問題を解くためのアルゴリズムを示しており、同一視される関数は、ある完全な測定下で十分に計算することができるとしている。また、たしかな過程の下で、彼らが示したアルゴリズムは、任意の推定アルゴリズムの能力の上限をいっているとしている。Blumたちの結果は、極限で帰納的に推論される対象の自然なクラスは、複雑なクラスであると提案している。この節では、モデル推論問題に似た結果を得る。ある帰納的関数 h に対し h -easyモデルのクラスが、定義される、また任意の h -easyモデルを極限で推論する簡単で効率的なアルゴリズムが、述べられる。もし、帰納的推論アルゴリズムが、単純な充分的要求を満たせば、即ち、現在知られている事実を、その現在の推測がいつも説明しているならば、そのアルゴリズムは、ある帰納関数 h とする、 h -easyモデルだけに同一視することができる。これは、帰納的推論アルゴリズムの能力の上限である。

枚挙型の推論アルゴリズムは、1 階述語がもつ意味と構文について、十分な優位性をもっていない。また、類似の枚挙型アルゴリズムは、ほとんどどんな他の計算モデルの中でもインプリントされ得る。ロジックの性質で優位性をとる推論過程を高速化する方法が、以下の節に示される。

3. 1 枚挙型アルゴリズム

ある許容対を $\langle L_0, L_h \rangle$ とする。 $\alpha_1, \alpha_2, \alpha_3, \dots$ を L_0 の全ての文に対して定められた効果的な枚挙であるとする。 $T_1, T_2, T_3, \dots$ を、 L の有限集合として定められた効果的な枚挙であるとする。 T が、 n 回の導出段階又はそれ以外で α を導出できる (できない) 事を示すために $T \vdash_n \alpha$ ($T \not\vdash_n \alpha$) を利用する。また、任意の有限の構文 $T \in L$ と任意の n に対し、 n 回の導出段階で T から導出される構文の集合は、計算可能で有限であると仮定される。我々は、また導出手続きは、単調である。即ち、任意の q, p, T に対して、 $T \vdash_n p \rightarrow T \cup \{q\} \vdash_n p$ であると、仮定する。このような証明手続きの例として、分解証明法 [Robinson 65] がある。

以下のさらに具体的な議論で、分解証明法は、我々の証明手続きであると仮定する。しかし、この選択は、本節のさらに抽象的な結果にとって、重要でない。アルゴリズム1は、次の簡単なアイデアをインプリメントしている：既知の事実と一致する推測を考えている間、推測が成立している。一度、推測が成立しないことがわかると、考えられる次の推測に対する枚挙をさがす。推測が事実と一致するか否かという判断は、一般に非決定的である。その時に、先験的なある境界を、選択すると、任意に与えられた事実で、この判定を行うことになるだろう。

アルゴリズム1が極限で推論するモデルの集合を特徴づけるためにいくつかの完全な論理的道具と定義を必要とする：これらのもっと完全な説明に対しては [Hachey & Young 78] を見よ。公理 T_k の各集合と階段表示関数 Φ_k を関連づける。 $\Phi_k(i) = \min \{n \mid T_k \vdash \alpha_i\}$ 。この関数は、部分帰納関数である。また読者は、階段表示関数 $\Phi_1, \Phi_2, \Phi_3, \dots$ の集合が、次のように定義される特性関数 $\phi_1, \phi_2, \dots$ の集合で、複雑さの尺度を構成していることがわかる。

$$\phi_k(i) = \begin{cases} \{T_k \vdash \alpha_i \text{ ならば、} 1. \\ \{\text{その他は不定。} \end{cases}$$

アルゴリズム1：枚挙型のモデル推論アルゴリズム。

h を、全帰納関数とする。

S_{false} を $\{\square\}$ 、 S_{true} を $\{\}$ 、そして k を 0 、とする。

repeat.

次の事実 $F_n = \langle \alpha, V \rangle$ を、読み込む。

α を、 S_{false} に加える。

while. $T_k \vdash \alpha$ となる $\alpha \in S_{\text{false}}$ が、存在するか

または、 $T_k \not\vdash \alpha_i$ となる $\alpha_i \in S_{\text{true}}$ が、存在する。 do.

k を、 $k+1$ とする。

output T_k .

forever.

L_0 の結論が、ある全帰納関数 h を法として、容易に得ることができれば、すなわち、 $T \vdash \alpha_n$ を満たすようなほとんどすべての $n > 0$ に対し、 $\Phi_k(n) \leq h(n)$ であれば、有限集合としての、構文 T は、 h -easy であるといわれる。 T_k を、 M で h -easy な L_0 -完全公理化とする $k > 0$ が、存在すれば、モデル M は、 h -easy である。もし、 M が、 h -easy であれば、アルゴリズム 1 は、極限で、 M を推論できる。すなわち、ある有限個の事実をたしかめた後、このアルゴリズムによる現在の推測は、モデル M の L_0 -完全公理化になる。また、この推測は、その後も、変化しない。いいかえると、

定理 3. 1: h を、ある全帰納関数とし、 M を L の h -easy モデルとする。このとき、アルゴリズム 1 は、極限で M と同一視する。

この定理を証明する為に、次の補題を必要とする。モデルが有限の L_0 でひどく（複雑に）ふるまう有限の L_0 -完全公理化をもつならば、異例な場合をすべて改善し、他の場合を少なくとも良くするような M の有限 L_0 完全公理化が他に存在する。

補題 3. 2: M を、ある全帰納関数 h の h -easy モデルとする。このとき、次のような k が存在する。

1. T_k は、 M で真である。
2. α_n は、 M で真なるすべての n に対して $\Phi_k(n) \leq h(n)$ である。

証明: M は、 h -easy であるので、 T_{k_1} は M の有限な L_0 -完全公理化となる k_1 であり、ほとんどすべての $n > 0$ に対し、 $\Phi_{k_1}(n) \leq h(n)$ である。この二条件を満たす k を見つけるために、すべての $i > 0$ に対し、 α_i が M で真であり $\Phi_{k_1}(i) > h(i)$ が成立する構文 α_i を、追加することによって、 T_{k_1} を修正している。このように定義された構文の有限集合は、 $L_0 \subset L_h$ なので L_h に存在する。 k_2 を構文のこの集合に関するインデックスとしよう。 T_{k_1} は、 M の L_0 完全公理化であり、 T_{k_2} は、 M で真なる構文を追加する事によって、 T_{k_1} から得られると、 T_{k_2} は、また M に関し L_0 完全公理化である。 $T_{k_1} \subset T_{k_2}$ からすべての n に対し、 $\Phi_{k_2}(n) \leq \Phi_{k_1}(n)$ であることがわかってい

る。さらに、関連する構文が T_{k_2} にあるので、 $\Phi_{k_1}(i) > h(i)$ とするすべての i について、 $\Phi_{k_2}(i) = \perp$ である。

定理3. 1の証明： k_0 を、以下のような最小の k とする。

1. T_k は、 M で真。
 2. α_n は、 M で真なる $n > 0$ に対し $\Phi_k(n) \leq h(n)$ 。
- そのような k は、補題3. 2により存在する。

$k < k_0$ とする任意の k は、少なくともこれらの条件を満足しない。

もし、 T_k が条件2を満足しなければ、 M で真なるいくつかの α_n に対して、 $T_k \not\models \Phi_k(n) \leq h(n)$ のとする $n > 0$ が存在する。 T_k は、条件2を満たすが、条件1を満足しなければ、 L_h と L_0 の許容的な要請により、 M で偽なるいくつかの α_{n_1} が存在する。ここで、 $T_k \models \alpha_{n_1}$ である。 $T_k \models \alpha_{n_1}$ から、最小の $n \geq n_1$ とする n を選択してみる。どちらの場合でも、 n 番目の事実を読むと、white-loopの条件が満足され、 k の値が増加する。ここで、有限の事実をたしかめると、アルゴリズム1は、 $k < k_0$ に対する推測 T_k のすべてを捨てるだろう、一方、このアルゴリズムがいつも T_{k_0} を推測するならば、種々の推測を決して再び出力しないだろう。その理由は、 T_{k_0} が、 M で真で、 $h(i)$ ステップで M で真なるすべての α_i を、導出するからである。そのために、有限の事実を説明した後、アルゴリズムは、 T_{k_0} を推測し、二度と、異なる推測を出力しないだろう。 ■

3. 2 充足的なアルゴリズムの能力限界

アルゴリズム1は、任意の帰納関数 h に対する、すべての h -easyモデルを、極限で、同一視できる帰納推論アルゴリズムが存在することを、示している。本節では、帰納推論アルゴリズムが充足的であるならば、ある定められた帰納関数 h に対する h -easyモデルにだけ極限で同一視が可能であることを示す。これにより、任意のこのようなアルゴリズムの能力の上限を確立する。

定義3. 3： $F_i = \langle \alpha_i, \text{true} \rangle$ 、 $1 \leq i \leq n$ のすべての事実に対し、 T が α_i を導出すれば、事実 $F_1, F_2, \dots, F_n$ に関し、構文 T は、充足的であると言われる。

記述されているアルゴリズム1は、充足的である。この性質は、推論アルゴリズムに関し、妥当な要求であるかのように思われる。それは、Blums の論文で、それを満たないいくつかの強力な推論アルゴリズムが、述べられている事を知っている事である。次の定理は、充足的な帰納アルゴリズムを要求すれば、 h -easyモデルがそれ自身である事を示している。

定理 3.4: I を、言語 L に対する充足的な帰納アルゴリズムであるとす
 る。このとき、 I で、一様な全帰納関数 h が存在する。すなわち、 I が、 L の
 モデル M に同一視するならば、 M は、 h -easy である。

証明：下記の補題 3.5 により、 I で一様な充足アルゴリズム I' が存在する。
 I' は、 I と同じくらい強力であり、すべての事実をいくつか読むだろう。

$S(n)$ を形成 $\sigma = (\langle \alpha_1, V_2 \rangle, \langle \alpha_2, V_n \rangle, \dots, \langle \alpha_n, V_n \rangle)$
 をもつ矛盾のない有限列とする。ただし、空節は、含まない。すべての $\sigma \in S$
 (n) に対し、この順序で、 σ の事実を読んだ後と、 $n+1$ 番目の事実を読み込
 む前に $I[\sigma]$ を、最後の推測アルゴリズム I' の出力のインデックスであると
 定義する。 I' はいつかすべての事実を読むので、 $I[\sigma]$ は、妥当な定義であ
 り、計算可能である。

I' は、充足的であるので、すべての $\sigma \in S(n)$ と、 σ のすべての事実
 $F_i = \langle \alpha_i, \text{true} \rangle$ に対し、いくつかの有限ステップで $k = I[\sigma] \rightarrow T_k \vdash$
 σ_i を導く。また、そのために、 $\Phi_k(n)$ が定義され、 $V_i = \text{true}$ とするすべ
 ての i に対し有限になる ($0 \leq i \leq n$)。我々は、次のような複雑な境界 h を選
 択する。

$h(n) = \max \{ \Phi_k(n) \mid k = I[\sigma], \sigma \in S(n), V_n = \text{true} \}$
 $S(N)$ の並びには、空節を含まないので、すべての n に対する $\sigma \in S(n)$ が、
 存在する。これにより、 h は、全帰納的である。

I' が極限で推論できるすべてのモデルは、 h -easy である事を示す。 M を、その
 モデルとし、その列 $\langle \alpha_1, V_1 \rangle, \langle \alpha_2, V_2 \rangle, \dots$ に I' を、適用してみ
 よう。 I' が M で同一視できるならば、十分に大きな n に対し n 個の事実を読ん
 だ後、 M のある L_0 完全公理化 T_i に収束するだろう。 h の定義により、十分に
 大きな n に対し、 $\Phi_i(n) \leq h(n)$ であり、 M は、 h -easy になる。■

補題 3.5: I を、十分に帰納的な推論アルゴリズムとする。そのとき、
 I で、一様かつ充足的な帰納推論アルゴリズムが存在する。即ち、 I' は、 I
 と同じくらい強力であり、すべての事実をいくつかは読む。

証明：この補題は、Minicozzi のユニオン理論から結論されるが、
 [Blum & Blum 75] にある定理 3 の証明と、変わらない。■

4. 決定的な実験(Crucial Experiments)による仮説の反ばく

アルゴリズム1の実行不可能性に関する主な原因は、その大局的性格にある。ある一組の構文が、モデルの公理化でないことがわかれば、単にそれをすて、次の発見的推測を見つけるまで、有限なすべての構文の組を調べる。本節では、増殖的推論アルゴリズムを開発するために、解かなければならない最初の問題を提供する。即ち、現在の推測が非常に強いこと（例えば、偽の観測結果を導く）が分ったとき、何をすべきかということである。この場合、少なくとも仮説の一つが偽であると結論できる。あるアルゴリズムは、偽の観測文から推測の中で偽なる仮説を、検出できるように開発されている。このアルゴリズムは、矛盾探索アルゴリズムと呼ばれており、仮説が偽になった原因をさかのぼって推測と事実の矛盾をたどる事ができる。このアルゴリズムと、科学理論の反ばく性に対する哲学的問題の関連性が議論される。また、ロジックプログラムをデバックする為のアプリケーションが示される。

4. 1 矛盾探索アルゴリズム

決定的な実験は、競争的な科学理論を選択する可能性をもつ実験である。良い結果をもつ決定的な実験は、その予測の1つと反対の例を与える競争理論の少なくとも1つを除去する事ができる。科学の進歩における決定的な実験の重要な役割が、科学哲学をもつ多くの学校によって認識されている。しかしながら、その極限的能力として、強い議論がある。哲学者Pierre Duhemが言うには“実験が我々におしえてくれる唯一のものとしては、使われている命題には、すくなくとも1つの誤りがあるということである；しかし、このエラーが存在する場所は、我々に教えられていない”[Duhem 54]。Popperでさえ科学の進歩に寄与する反ばくの主な仕事を次のように言っている。“我々は、しばしば、理論システムの、ある大きな単位と、ときどきおそらく全システムにだけテストができるということとを、認めなければならない。また、それは、この場合、その要素が虚偽性に対し応答的になるべきであるというけわしい推量(a sheer guesswork)である”[Popper 68]。

矛盾探索アルゴリズムは、この制限をある意味で克服することができる。一般に、決定的な実験は、収集された仮説だけを反ばくできるが、矛盾探索アルゴリズムは、連続的に決定的な実験方法を、提案している。この方法は、唯一の偽の仮説を選択する事を保証している。この方法は、エラーが見い出される場所を正確に伝えており、任意の推量を含まない。この連続的方法は、動的である。即ち、実施されるべく実験の順序は、アブリアリに決められない；むしろ、すべての新しい実験の結果は、その後任を、提案する。

最後に、偽の仮説を明白に指摘する。さらに、これらの実験の正しい結果は、反ばく的仮説に反例を与えている。即ち、この仮説構文は、仮説から理論的に導き出され、偽であることが実験的に決定された構文である。

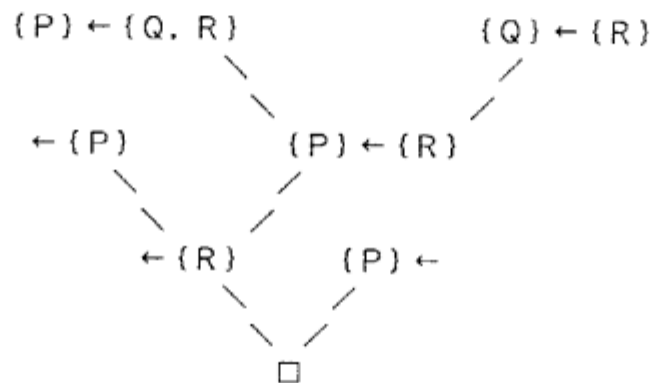
矛盾探索アルゴリズムは、矛盾が、いくつかの仮説と事実の間に導出されると、いつでも適用される。その入力は、仮説と真の観測構文 S から空節 “ $\square$ ” をもつ順序づけられた導出木である。

このアルゴリズムには、 L のすべての ground atom の真理を告げる M での Oracle が与えられている。そのアルゴリズムは、 M で有限回の実験を行っているが、この実験は、その導出木の深さにより制限され、 M で偽の仮説 $p \in S$ を出力する。

最初に、このアルゴリズムが Figure4-1 で示される命題計算例で何をやるかを説明する。この例では、

$S = \{ (P) \leftarrow \{Q, R\}, \{Q\} \leftarrow \{R\}, \leftarrow \{P\}, \{R\} \leftarrow \}$ である。

Figure4-1 : 命題論理における矛盾のバックトレース。



このアルゴリズムは、ルート $\square$ からスタートしている。また、導出されるアトムを反復的に確かめている。もし、アトムが M で真であれば左の補木をえらび、そうでなければ右の補木を選ぶ。これは、リーフにたどりつくまで行なわれる。リーフにある仮説は、 M で偽なことが保証されている。示されている例では、仮説 $(P) \leftarrow \{Q, R\}$ が、偽なる事を仮定している。これは、 R と Q が M で真で P が偽なる事を意味する。

このアルゴリズムは、RがMで真なるかどうかを、最初にためし、我々の仮定によりそれが真なので左側の補木を選んでいく。次に、Pがそのモデルで偽になっていることをたしかめ右側の補木を選んでいく。最後に、Qが真であることをたしかめ、左側のブランチを選択したところ、リーフになっていることをたしかめている。Mで偽となっている仮説 $\{P\} \leftarrow \{Q, R\}$ を出力している。

一つの実験に対する答は、次の実験を決定することに注意せよ。また、実験の回数は、まさに到達したリーフの深さであり、これらの実験を結合した結果は、そのリーフでの仮説の反例を構成する。もし、いくつかのリーフが、(仮説よりも)むしろ真の観測構文をもてば、必要とされる実験の回数は、非常に少なくなるかもしれない。

計画的な場合、アトムをすべて確認することによって偽の仮説を、簡単に検出できる。その確認は、Mでそれらの真理を許しているように見える。このだまされやすい手続きは、指数的に必要とされる実験の数を増大するかもしれないが、しかし、原理的には、同じ仕事を行うことができる。述語計算の場合に、この仕事は、より複雑になる。全称的仮説は、無限領域の上に、効果的に確認され得ないが、もしそれが偽ならば、それは偽のground instanceをもつ。述語計算の場合に、このだまされやすい手続きは、その少しばかりの領域を増加して、仮説をシステムチェックに例証することである。また、仮説の反例、すなわち、それに関する偽の例証が発見されるまで、作成されるground instance (すべて L_0 に存在すると仮定) を、すべてたしかめる。許可されている仮説の少なくとも1つが偽であると、そのアルゴリズムは、停止することが、保証されている。また、そのような仮説は、有限領域でそれを例証したとき、偽の例証をもっている。しかしながら、このアプローチは、明らかに実行不可能である。

1階述語の言語における矛盾探索アルゴリズムは、それに対する反例をシステムチェックに構成して偽の仮説を検出するのと、同じアイデアに基づいている。しかしながら、この反例が構成される方法はさらに細く含まれている。このアルゴリズムを記述する前に、導出に関するいくつかの用語を、さらに正確にさせておく必要がある。これらの問題については、ロビンソン [Robinson 65] にならう。置換は、形式 V/t の対に観測する有限集合 (空も可能) である。ここで、 V は、変数で t はタームである。これらの変数は、どれでも同じでない。任意の置換 $\theta = \{V_1/t_1, V_2/t_2, \dots, V_n/t_n\}$ と表現 s に対し、表現 $s\theta$ は、 s で変数 V_i の各でき事をターム t_i で置き換えた結果である。もし、置換 θ に対して、 $t = s\theta$ であれば、 t は、 s の例証と呼ばれている。 S が表現の集合であれば、 $S\theta = \{s\theta \mid s \in S\}$ である。もし、二つの構文が、おたがいに例証であれば、それらは変数に名をつけかえることに等しい。また、その他に指摘されなければ、それらを区別しない。

θ を、ある置換とする。もし、 $S \theta$ がある別物であれば、アトム S の有限集合（空ではない）は、 θ によりユニファイできる。

もし、 $\theta_1 = \{t_1 / V_1, t_2 / V_2, \dots, t_n / V_n\}$ と θ_2 を、2つの置換とすると、 $\theta_1 \circ \theta_2 = \theta_1' \cup \theta_2'$ が、成立する。ここで、 θ_2' は、 θ_2 のすべての要素の集合である。 θ_2 の変数は、 $V_1, \dots, V_n$ の中になく、 $\theta_1' = \{t_i / \theta_2 \mid 1 \leq i \leq n, t_i / \theta_2 \neq V_i\}$ の中にある。任意のストリング s と2つの置換 θ_1 と θ_2 に対し、等号 $(s \theta_1) \theta_2 = s(\theta_1 \circ \theta_2)$ が成立することが示される。もし、 θ が S のユニファイアであれば、置換 θ は、 S に関し最も一般的なユニファイアである。また、 S に対する任意の他のユニファイア θ_1 に対し、 $\theta = \theta_1 \circ \theta_2$ である様な置換 θ_2 が存在する。もしそれが存在するとき、Robinson は、任意の集合に関する最も一般的なユニファイアを計算するアルゴリズムを、述べている。

Paterson と Wegman [Paterson & Wegman 76] は、linear time で操作するアルゴリズムを述べている。

定理 4. 1 $A \leftarrow B, C \leftarrow D$ を、 L の2つの構文とする。ただし、 $R_1 \subset B, R_2 \subset C$ である。もし、集合 R_1 と R_2 が、最も一般的なユニファイア θ をもてば、 $\{P\} = R_1 \theta = R_2 \theta$ であり、このとき $[A \cup (C - R_2) \leftarrow (B - R_1) \cup D] \theta$ は、 $A \leftarrow B$ と $C \leftarrow D$ の、ある導出形である。また、 P は、上で導出されたアトムである。導出形を形成する過程は、左成分に $A \leftarrow B$ 、右成分に $C \leftarrow D$ をもつ導出と呼ばれる。

矛盾探索アルゴリズムは、ある反例の断片を作成している。また、そのような反例を見つけるために、関係する実験だけを行っている。これは、そのアルゴリズムをさらに複雑にしている。この述語計算の場合、導出されるアトムは、ground にならないので、 M についての Oracle を使って、その真実を直接たしかめることが必ずしもできるとは限らない。この問題の解は、Oracle にそれを与える前に、ある ground atom に対する P を例証することにある。 P をいかに例証するかという選択は、任意であるが、しかし一度それが行なわれると、リーフにたどりついた仮説に対する反例を作るために、同じ置換によるすべてのすすんだ実験が行なわれてるべきである。

アルゴリズム2：決定的な実験による矛盾のバックトレース。

入力：モデルMのあるOracleと、次の性質をもつ構文のOrdered 2進木。

1. ルートは、空節である。
2. すべてのリーフは、仮説又は、観測文のどちらかがMで真であり、
2つのリーフは、1つの変数を共有しない。
3. リーフに存在しないすべてのノードは、その子孫に関する2進導出形である。ここで、導出に関し左の子孫は、左の成分であり、右の子孫は、右の成分である。

出力：導出木のリーフに発生し、Mで偽なる一つの仮説。

アルゴリズム：kを0、 N_0 をその木のルートとし、 θ_0 を {} とする。

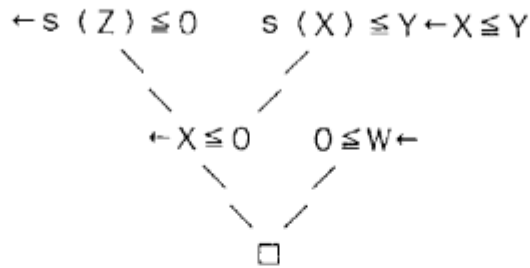
```

while  $N_k$  がリーフにならない do.
    Pを、最も一般的なユニファイア $\theta$ をもつ $N_k$ で導出されるアトムとする。
     $P\theta_k$ をground atom  $P_k$ にインスタンス化して置換 $\theta'$ 、を選択する。
     $\theta_{k+1}$ を、 $\theta \circ \theta_k \circ \theta'$ とする。
     $P_k$ が、Mで真であるか否かを調べる。
    case  $P_k$ が、
        真のとき： $N_{k+1}$ を $N_k$ の左子孫とする。
        偽のとき： $N_{k+1}$ を $N_k$ の右子孫とする。
    kを、k+1とする。
output  $N_k$ .

```

以下は、モデル推論アルゴリズムによる矛盾探索の利用例である。このアルゴリズムは、自然数で関数 $\leq$ の有限近似を推論する事を行なっている間、この矛盾探索が行なわれる。ここで自然数は、定数0と後続関数 $s(X)$ で作られる。このアルゴリズムは、仮説 $0 \leq X \leftarrow$ と $\leftarrow s(X) \leq 0$ をすでに推測しており、事実 $\langle s(0) \leq s(s(0)), \text{true} \rangle$ に出くわしていると仮定する。仮説 $s(X) \leq X \leftarrow X = Y$ を提案すると、仮説 $0 \leq X \leftarrow$ をもつ構文 $s(0) \leq s(s(0))$ が導出される。しかしながら、この新しい仮説を追加後、Figure4-2の導出が、また構成される。包含されている仮説が偽になることを検証するために、矛盾探索アルゴリズムを適用してみよう。

Figure4-2 : 述語論理での矛盾のバックトレース.



ルートに導出されるアトムは、 $0 \leq 0$ である。Oracleは、 $0 \leq 0$ に対して、'true'を返す。そのため左側の木が選ばれる。

即ち、 N_1 が、 $\leftarrow X \leq 0$ に設定され、 θ に $\{X/0\}$ が設定される。

N_1 で導出されるアトムは、 $s(X) \leq 0$ であり、 $(s(X) \leq 0) \theta_1$ は $s(0) \leq 0$ である。このOracleは、この確認で'false'を返す。右側の木は、 $\theta_2 = \{Y/0\} \cdot \theta_1 = \{X/0, Y/0\}$ で選択され、リーフ $N_2 = s(X) \leq Y \leftarrow X \leq Y$ は、偽になることが、検出される。この実験の結果により構成される反例は、 $s(0) \leq 0 \leftarrow 0 \leq 0$ であるが、これは、また $N_2 \theta_2$ である。仮説 $s(X) \leq X \leftarrow X \leq Y$ のすべての具象例が、偽でないことに注意しよう。例えば、置換 $\{X/0, Y/s(0)\}$ は、真なる具象例 $s(0) \leq s(0) \leftarrow 0 \leq s(0)$ を生ずる。

このアルゴリズムの正しさを示す前に、実験の結果が、仮説に対する反証を、いかに構成できるかをさらに正確に定義しよう。

$A' \theta \subset A$ かつ $B' \theta \subset B$ あるいは要するに、 $(A' \leftarrow B') \theta \subset (A \leftarrow B)$ (5.3節を見よ)となるような置換 θ が、存在すれば、構文 $A' \leftarrow B'$ は、構文 $A \leftarrow B$ を包含する、と言う。ground atom $P_1, P_2, \dots, P_k$ が確められると仮定してみよう。BをMで真なるアトム $\{P_1, P_2, \dots, P_k\}$ の部分集合とし、AをMで偽なる部分集合とする。このとき、ground sentence $A \leftarrow B$ は、Mで偽であり、それを含むすべての構文に類似している。そのような場合、 $A \leftarrow B$ は、 $P_1, P_2, \dots, P_k$ を確認した成果から構成された θ による ρ の反例として参照される。

補第4.2: MをLのモデルとし、TをLの構文から空節をもつOrdered 2進導出木とする。もし、アルゴリズム2をTに適用し、構文 N_k が、ground atom $P_1, \dots, P_k$ を確めた後にたどりつくならば、 N_k は、Mで偽であり、これらの確認の成果は、 θ_k による N_k の反例を構成する。

証明: この証明は、kの帰納と、多くの確認により存在する。もし、 $k=0$ ならば、 $N_0 = \square$ である。これば、Mで偽であり $\theta_0 = \{\}$ である。 $\square \theta_0$ の実験なしに構成される反例は、また空節である。

ground atom $P_1, \dots, P_k$ を確認するアルゴリズム2は、ノード N_k にたどりつくとし、この確認の成果は、置換 θ_k により構文 $N_k = (A' \leftarrow B')$ に対する反例 $A \leftarrow B$ を構成すると、帰納的に仮定しよう。Pを N_k で導出されたアトムとしよう。 N_k の左側の子孫は、 $A_1 \leftarrow B_1 \cup R_1$ であり、 N_k の右側の子孫は、 $A_2 \cup R_2 \leftarrow B_2$ である。集合 R_1 と R_2 は、 $R_1 \theta = R_2 \theta = \{P\}$ かつ $A' \leftarrow B' = (A_1 \cup A_2 \leftarrow B_1 \cup B_2) \theta$ とする、最も一般的なユニファイヤ θ をもつ。

$$\begin{array}{ccc} A_1 \leftarrow B_1 \cup R_1 & & A_2 \cup R_2 \leftarrow B_2 \\ & \searrow \quad \swarrow & \\ N_k = A' \leftarrow B' & = & (A_1 \cup A_2 \leftarrow B_1 \cup B_2) \theta \end{array}$$

$P \theta_k$ は、ground atom P_{k+1} に対する、置換 θ' により例証され、 θ_{k+1} は、 $\theta \circ \theta_k \circ \theta'$ になり、 P_{k+1} は、Mで置換されると、仮定する。次のような、2つの場合がある:

場合1: P_{k+1} は、 M で真である。このとき、アルゴリズム2は、 $A_1 \leftarrow B_1 \cup R_1$ である N_k の左側の子孫を N_{k+1} に設定する。帰納的仮定により、 $A \leftarrow B$ は M で偽である。 P_{k+1} は、 M で真なので、構文 $A \leftarrow B \cup \{P_{k+1}\}$ は、また M で偽である。 $(A_1 \leftarrow B_1) \theta \in N_k$ なので、帰納 $(A_1 \leftarrow B_1) \theta \theta_k \theta' = (A_1 \leftarrow B_1) \theta_{k+1} \in N_k \theta_k \theta'$ が成立する。また $N_k \theta_k \theta'$ 最後に ground になるので $N_k \theta_k$ に等しい。 $N_k \theta_k \in (A \leftarrow B)$ であるという事実から、 $(A_1 \leftarrow B_1) \theta_{k+1} \in (A \leftarrow B)$ である。また、 $R_1 \theta = P$ で $P \theta_k \theta' = P_{k+1}$ であるので、 $R_1 \theta_{k+1} = P_{k+1}$ である。そのため、 $N_{k+1} \theta_{k+1} = (A_1 \leftarrow B_1 \cup R_1) \theta_{k+1} \in (A \leftarrow B \cup \{P_{k+1}\})$ が成立する。

一方、 N_{k+1} は、 θ_{k+1} により $A \leftarrow B \cup \{P_{k+1}\}$ を含む。そのため、 N_{k+1} は、 M で偽であり、 $P_1, \dots, P_{k+1}$ を確めた結果は、置換 θ_{k+1} による N_{k+1} の反例を、構成する。

場合2: P_{k+1} が M で偽である。対照的な議論が、適用される。 ■

次の定理は、補題4.2から、すぐわかる系である。

定理4.3: M を L のモデルとし、 T を L の構文による空節をもつ Ordered 2進導出木とする。もし、 T の深さが n とすると、 T に適用される。アルゴリズム2は、 M で n 回以上の実験をもはや行わず、 M で偽の仮説 p を出力する。 ■

4. 2 科学的仮説の反ばく性

反ばく理論で、偽の仮説を検出する事は、むしろ困難な仕事であるかの様に見える。それが、哲学者Pierre Duhemにより、不可能である事が主張された。つまり、“物理学者は実験的確認による孤立した仮説を決して指示しないが、仮説の完全なかたまりだけを指示する” [Duhem 54]。私は、Duhem の言葉をレビューし、矛盾探索アルゴリズムが動作するときの仮定で、それを関係づけた。本質的に、Duhem の言葉は、唯一に偽にする実験の可能性とは逆に、2つの要素：論理的要素と認識論的要素をもつ。Duhem の言葉の論理的要素は、次のように示される。“実験が、我々におしえてくれる唯一のものは、すくなくとも1つのエラーが存在することである；しかし、このエラーが生ずる場所は、我々に知らせていない、ということである。物理学者は、このエラーが、反ばくを望んでいるその信条に正しく含まれていると、定義するかも知れない。しかし、彼は、他の信条に、それが存在すると確信しているだろうか？ もし、彼がいれば、彼は、自分が利用した他の信条の正当性をすべて絶対的に受け入れる。また、彼の結果の有効性は、彼の信賴の有効性と同じくらい大きい。”

このDuhem の主張は、決定的な実験の真理である。しかしながら、矛盾探索アルゴリズムで示されたように、その事実と理論の間に存在する矛盾が発見されると、アルゴリズム的に次にやるべき実験を定めることができる。それは、その結果が、反ばく理論で偽の仮説を選択することによる。さらにDuhem は、次のようなことを述べている。矛盾が仮説と事実の間に生ずると、ある仮説は、その導出で利用された他のすべての仮説の正当性を正しく受け入れることによってだけ偽である。この矛盾探索アルゴリズムは、次の論証により、この主張を、反ばくしている。ground atom(Variable-free)に関してだけ実際の判断の正当性を許容することにより、仮説の虚偽性を、結論づけることができる。このground atom は、具体的な対象と、でき事を参照する構文である。

Duhem の言葉の認識論的要素は、具体的な対象と、でき事に関する実際の判断が、次のいくつかの論理的仮定を含んでいることにある：“実験を行うか、1つの報告を与える物理学者は、理論の完全なかたまりが正確であることを、絶対的に認識している” 理論項の概念を組み合わせるモデル推論問題の拡張は、2. 3節で論議した。矛盾探索アルゴリズムは、うまくこの拡張を、利用することができる。非公式に、矛盾探索アルゴリズムが、このモデルでground atom を確かめる方法は、次のようになる：もし、ためされるアトムの述語が可観測であれば、このモデルで直接、確かめられる。さもなければ、アルゴリズムが理論的述語について、いくつかの組み込み知識をもち、この知識をもってground atom が真か偽かのどちらかを定める事ができると仮定する。どちらの場合も、矛盾探索アルゴリズムは、このモデルで、偽の仮説を導く。さらに、公式的に、この仮定は、L

のすべてのground atom $P = p(t_1, t_2, \dots, t_n)$ に対し、 $P \in L_0$ であるか、又は、 p が、ある固定的解釈をもつ理論述語であるか、のどちらかである。この固定的解釈は、その帰納的推論アルゴリズムで、知られている。

Duhem の言葉の議論をまとめると矛盾探索アルゴリズムは、彼の言葉の論理的要素を反ばくするかのように見える。

1. 決定的な実験は、収集されている仮説だけを反ばくすることができる。
2. ある仮説を反ばくするために、他の仮説が真理であることを受容しなければならない。

ここで定義されたように、たとえ、科学理論で使われている項を、理論項と観測項に分解できないとしても、Duhem の言葉の認識論的要素は、今でも成立している。

4. 3 応用：ロジックプログラムの虫取り

矛盾探索のアイデアは、ロジックプログラムの虫取りに容易に利用できる。もし、ロジックプログラムが、不正な入出力関係を計算するか、与えられた出力で不正の入力を返すならば、少なくとも次のような節がプログラム中に存在する。その節は、述語の予期された解釈の下で偽である。例として、リストの後続を計算する次のプログラムから、

```
subsequence([], []).
subsequence(L1, [X|L2]) ← subsequence(L1, L2)
subsequence([X|L1], [Y|L2]) ← subsequence(L1, L2)
```

我々は、 $\text{subsequence}([a, b, c], [a, a, c, s])$ を、計算できる。述語 $\text{subsequence}(X, Y)$ の予期された解釈の下で（即ち、 X は、 Y の後続である）プログラムの第3目の節は偽であり、次の例は、反例である。

$\text{subsequence}([b], [a]) : -\text{subsequence}([], [])$ 。

この節の‘虫’は、任意の他のリスト L_2 の後続として、任意のリスト L_1 を許すプログラムを発生する。即ち、もはや L_1 は、 L_2 より長くない。

矛盾探索アルゴリズムを $\text{subsequence}([a], [b])$ の計算（証明）に利用すると、最初に、 $\text{subsequence}([], [])$ を確認すると、真である事がわかる。次に $\text{subsequence}([a], [b])$ を確認すると、偽であることがわかる。そして、 subsequence の第3番目の節に対する上記の反例を与える。

読者は、subsequence に対する正しいロジックプログラムを見つけることになる。解は、後で与えられている。

以下では、矛盾探索アルゴリズムをインプリメントしているPrologプログラムを記述する。そこでは、仮説は、ホーン節であり、観測言語は、すべてアトムであり、Oracleはユーザである。また、ユーザが、虫のあるquicksort プログラムの虫取りをしようとしているときの注釈付きリストを示す。

矛盾探索アルゴリズムにより利用されている導出の書き換え収集法は、Green Green 69] により提案されたのが、はじまりである。彼のゴールは、証明の中で使われた仮説と、置換の収集からこの節を証明する事により、ロジカルな構文として形成された質問に対する応答を構成する事にあった。Prologでは、次の点を除くと、非常に一致したアイデアを、用いている。除くべき点は、証明される構文が、 $\leftarrow A$ の形をしており、仮説がホーン説になるように制限されている、という点である。

Prologインプリメンタは、導出過程としてのすべての置き換えを、自動的に維持しているので、その中で、矛盾探索アルゴリズムをインプリメントするのが極度に簡単になっている。

Figure4-3 は、そのようなインプリメンテーションの全リストである。同じようなインプリメンテーションは、本モデル推論システムに取り入れられている。

Figure4-3: 矛盾探索のPrologプログラム

```
backtrace((P,Q),CE) ←!, backtrace(P,CE), backtrace(Q,CE).
backtrace(P,CE) ←clause(P,Q), backtrace(Q,CE), resolve((P←Q),CE).

resolve((P←Q),CE)←var(CE),!, ask(P,V), (V=false, CE=(P ←Q);V=true).
resolve((P←Q),CE).

ask(P,V)←recorded(fact, (P,V)),!.
ask(P,V)←repeat, write(P), put(63), nl, read(V), (V=true;V=false),
record(fact, (P,V)),!.
```

意図されたモデルMで、偽となるアトムに対して、最初の手続きは、アトムAtomの証明の組み立てを試み、それが真ならば、Atomの導出で使用された為の仮説CE（反例に対応する）のインスタンスを置換していることになる。読者は、Prologに関する詳細についてPrologマニュアルを参照されたい[Pereira et al. 78]。特に、構文 $\{P\} \leftarrow \{Q_1, Q_2, \dots, Q_n\}$ に関する内部表現は、 $P \leftarrow (Q_1, (Q_2, (\dots, Q_n) \dots))$ であり、 $\{P\} \leftarrow$ は、 $P \leftarrow \text{true}$ である事に注意せよ。

backtrace の手続きは、単純なPrologインプリンタであるが、このインタプリタは、ある節の導出がある毎に増大する。

backtrace(Atom, CE)の呼び出しは、まず、アトムAtomの証明を最初に組み立てている。これは、ある証明が存在し、証明の再帰的な組み立てから 'popping' backしている途中、節 $P \leftarrow Q$ を証明するこめに $\text{resolve}(P \leftarrow Q, CE)$ を呼び出す時にかさる。導出に関する基礎的なロジックは、すべて巻きこんである。この決定的なアイデアは、次のようなものである。

Pが偽になる $\text{resolve}(P \leftarrow Q, CE)$ を最初に呼び出すと、CEに節 $P \leftarrow Q$ がユニファイされる。このアイデアの正当性は、次のように帰納的に論じられる：もし、 $Q = \text{true}$ でPが偽であれば、節 $P \leftarrow Q$ は、偽であり、backtrace は、変数CEに偽節を返す。backtrace が証明の再帰的組み立てからもどってきたときに作られる分解の為の事前呼びだしで、Pの上に分解されるアトムがテストされ、真なることが見つけられると、帰納的に仮定しよう。

このとき、条件Qは、真である。そのため、もし、Pが現在テストされ、偽であれば節 $P \leftarrow Q$ は、偽になる。その結果、変数CEは、 $P \leftarrow Q$ をユニファイし、backtrace は、偽節のインスタンスを返す。補第4.2の証明での、みごとな点は、証明の組み立てと反例の組み立ての中で使われた置換を収集するための必要性によって生じたが、この点は、Prologのユニフィケーション手続きと、共有変数の機能によって、自動的に配慮されている。

以下は、次に示す quicksort プログラムをデバックするために、ユーザが、Figure 4-3 のプログラムを実行した実行トレースである。この quicksort プログラムは、最初、いくつかの虫をもっている。それらの虫は、つまらない虫であったり、むずかしい虫であったりする。

```

qsort([X | L], L0) ← partition(L, X, L1, L2).
                        qsort(L1, L3), qsort(L2, L4), append(L3, L4, L0).
qsort([], []).

partition([X | L], Y, L1, [X | L2]) ← partition(L, Y, L1, L2).
partition([X | L], Y, [X | L1], L2) ← X ≤ Y, partition(L, Y, L1, L2).
partition([], X, [], []).

append([X | L1], L2, L3) ← append(L1, L2, L3).
append([], L, L).

```

プログラムの虫取をするために、その意味する解釈を知る必要がある。qsort(X, Y) は、リスト Y がリスト X をソートしたものであれば、真である。partition(L, X, L1, L2) は、L1 が、X 以下の要素 L を含み、L2 が、X より大きい要素 L を含めば、真である。以下の実行トレースで、'|' で、はじまる行は、ユーザにより入力されており、残りの行は、システムの出力である。quicksort プログラムを実行してみよう。

```

| ?-qsort([2,1,2], X).

X=[]

```

これは、正しくない結果である。そこで、backtrace を呼び出してみる。この backtrace は、上記、qsort, partition, append に対する 'true' または 'false' で答える真偽の質問をする。

```

| ?-backtrace(qsort([2,1,2],X,CE).
partition([],2,[],[]))?
| :true.
partition([2],2,[],[2]).
| :false.

X=[],
CE=(partition([2],2,[],[2])←partition([],2,[],[]))

```

backtrace は、反例として、partition の最初の節を、選出した。この節は、その帰結部が偽のあいだ、その条件部が真であるので、偽である。X>Y の判定が考慮されていないという明白な誤りがある。この判定をつけ加えて、この節を、修正すると、

```

partition([X | L],Y,L1,[X | L2])←X>Y,partition(L,Y,L1,L2).

```

そして、その新しいプログラムを、実行してみよう。

```

| ?-qsort([2,1,2],X).

X=[]

```

それは、今までどおり偽の答を返すので、ふたたび、backtrace を呼び出してみよう。

```

| ?-backtrace(qsort([2,1,2],X), CE).
partition([2],2,[2],[]).
| :true.
partition([1,2],2,[1,2],[])?
| :true.
partition([],1,[],[])?
| :true.
partition([2],1,[],[2])?
| :true.
qsort([],[])?
| :true.
append([],[],[])?
| :true.
qsort([2],[])?
| :false.

X=[],
CE=(qsort([2],[]) $\leftarrow$ partition([],2,[],[]),
    qsort([],[]),qsort([],[]),append([],[],[]))

```

今、CEは、反例として、qsort それ自身を、制限している。qsort([2],[]) に対するすべてのサブゴールが、ただしい答を返したことは、明白である。我々は、そのpartition が、最終結果で実施されるその要素を含めることを、単に、わすれたのである。すなわち、appendの呼び出しを修正すれば、それが可能になる。

```

qsort([X | L],L0) $\leftarrow$ partition(L,X,L1,L2),qsort(L1,L3),
    qsort(L2,L4),append([X | L3],L4,L0).

```

qsort([],[])は、いつでも、そのアルゴリズムによりテストされるが、最初のテスト後、その事実 <qsort([],[]),true> がデータベース内に記録されるので、ユーザは、一度だけ、その質問になやまされる事に、注意しよう。虫取りが進行するので、データベースは、虫取りが実施されている手続きについて、多くの事実をもつようになり、その過程のなかでのユーザの干渉は、だんだん少なくなる。このアルゴリズムの増殖的ふるまいは、比較的、大きなシステムの虫取りをするのに、非常に有用である。ふたたび、qsort を実行してみよう。

```
| ?-qsort([2,1,2],X).
```

```
X=[]
```

```
yes
```

```
| ?-backtrace(qsort([2,1,2],X),CE).
```

```
append([2],[],[2])?
```

```
| :false.
```

```
X=[],
```

```
CE=(append([2],[],[2])←append([],[2],[2]))
```

appendの中で明白な虫が、発見されている。プログラムは、全体として、デバックされ、そのメインの手続きをデバックするまえに、そのサブの手続きをデバックを終了する必要がない。appendの最初の節を、訂正してみよう。

```
append([X|L1],L2,[X|L3])←append(L1,L2,L3).
```

再度、qsort を、実行してみよう。

```
| ?-qsort([2,1,2],X).
```

```
X=[2,1,2]
```

```
yes
```

```
| ?-backtrace(qsort([2,1,2],X),CE).
```

```
append([2],[],[2])?
```

```
| :true.
```

```
qsort([2],[2])?
```

```
| :true.
```

```
append([],[2],[2])?
```

```
| :true.
```

```
append([1],[2],[1,2])?
```

```
| :true.
```

```

qsort([1,2],[1,2])?
| :true.
append([1,2],[],[1,2])?
| :true.
append([2,1,2],[],[2,1,2])?
| :true.
qsort([2,1,2],[2,1,2])?
| :false.

X=[2,1,2],
CE=(qsort([2,1,2],[2,1,2]) ←partition([1,2],2,[1,2],[ ]),
    qsort([1,2],[1,2]),qsort([],[ ]),
    append([2,1,2],[],[2,1,2]))

```

ここで、何が誤りであるかを、みつけることは、少しむずかしいが、丁度、訂正したqsort の節をみると、その誤りの訂正が明白になる。そのpartition を実行するときのその要素は、その最後の結果に追加されるが、まちがった立場では、その要素は、それよりも大きい方のリストではなく、小さい方のリストのヘッドにある。これは、次のように、容易に訂正できる。

```

qsort([X | L],L0)←partition(L,X,L1,L2),qsort(L1,L3),
    qsort(L2,L4),append(L3,[X | L4],L0).

```

再度、qsort を実行してみよう。

```

| ?-qsort([2,1,2],X).

X=[1,2,2]

yes
| ?-qsort([5,2,45,3,24,543,2,3,1, ],X).

X=[1,2,2,3,3,5,24,45,543]

yes

```

それは、正しく動作している。虫取り後、その最終的なquicksort プログラムは、次のようになる。

```
qsort([X | L], L0) ← partition(L, X, L1, L2),  
                    qsort(L1, L3), qsort(L2, L4), append(L3, [X | L4], L0).  
qsort([], []).  
  
partition([X | L], Y, L1, [X | L2]) ← X > Y, partition(L, Y, L1, L2).  
partition([X | L], Y, [X | L1], L2) ← X ≤ Y, partition(L, Y, L1, L2).  
partition([], X, [], []).  
  
append([X | L1], L2, [X | L3]) ← append(L1, L2, L3).  
append([], L, L).
```

実際の矛盾探索にもとづくいたデバッガーをつくるために、少なくとも、さらに二つの道具が、必要である。第一番目は、ユーザが入力した誤りを訂正する方式をもつべき道具、である。第二番目は、矛盾探索アルゴリズムが、あるプログラムの呼び出しに成功する場合にだけ適用され得る道具、である。しかし、そのプログラムは、誤った答をもっている。もし、その呼び出しが、失敗したならば、他のアプローチが必要である。次の節では、それが失敗した場合、いかにして、ロジックプログラムに節集合をつけ加えるかについて、一つのアプローチを提案するが、新しい節をつけ加えるよりも、既存の節の修正を基本とする他のアプローチは、このようなデバッキングの仕事にとって、より实际的であるかもしれない。

5. 反ばくされた仮説の精密化

1. 2節で述べた2つの問題を思い出してみよう。矛盾探索アルゴリズムは、非常に強い推測の中から、偽の仮説を検出することにより最初の1つを解いている。そのような仮説を推測から取り除くと、非常に弱い推測になるかもしれない。この節では、その推測をいかに強くするかという問題を述べる。この仕事のために、 L_h の構文上に精密化演算子 (refinement operator) を考案する。直観的に、精密化演算子は、反ばく仮説に対して論理的に比較的弱い発見の置き換えをする。次の方法で、増殖的な帰納推論アルゴリズムを利用することができる：ある仮説に対する反例が発見されたときはいつでも、その仮説の除去により、反ばくされた推測を修正する。もし、結論された推測が、弱すぎるならば、他の仮説を付加し、その推測を強くする。ここで、他の仮説は、この精密化と、以前に反ばくされた仮説から選択される。

以下では、精密化の概念と、精密化演算子を明らかにする。精密化演算子の完全性は、与えられた仮説言語に対して定義されている。また、いくつかの例が、種々の精密化演算子として与えられるが、この演算子は、仮説言語の種々のクラスに対して完全である。

L_h の構文で精密化演算子により導かれる精密化グラフの特性が述べられる。任意の1階述語言語 L に対して、完全である最も一般的な精密化演算子についても述べられる。

5.1 精密化演算子

ある構造的な複雑度を表わす尺度を、仮定する。この尺度は、 L の構文から自然数へ対応づける関数であり、すべての $n > 0$ に対し、尺度 n の構文は、有限であるという特性をもつ。

構文 S の任意の集合と任意の $n > 0$ に対し、集合 $\{p \in S \mid \text{size}(p) \leq n\}$ を $S(n)$ と定義する。

定義5.1: L を、1階述語とし、 p と q を L の構文とする。このとき、 p が q を導出し $\text{size}(p) \leq \text{size}(q)$ ならば、 q は p の精密化であるという。

定義5.2: 精密化演算子 ρ は、 L の構文から、それらの精密化に関する部分集合への写像である。したがって、任意の $p \in L$ と任意の $n > 0$ に対し集合 $\rho(p)(n)$ 、すなわち、 $\text{size} \leq n$ の構文として制限される集合 $\rho(p)$ は、計算可能である。

L 上の精密化演算子は、 L 上の半順序 $\leq$ を導出するが、この半順序は、最少要素として空節をもつ： $p_{i+1} \leftarrow \rho(p_i)$ を満たす有限の節

$p = p_0, p_1, \dots, p_n = q$ は、

有限の全 ρ チェーン(finite total ρ -chain)と呼ばれる；もし、 p から q への有限の全 ρ チェーンが存在するならば $p \leq q$ とする；もし、 $p \leq q$ かつ $q \leq p$ が成立するならば、 $p \sim q$ とする。

関係 $\leq$ は、構文の集合を一般化する。すべての $q \in S$ に対して、 $p \leq q$ なる構文 $p \in T$ が存在すれば、構文中の二つの集合 $T, S \subset L$ に対し、 $T \leq S$ とする。

任意の構文 P に対し、集合 $\{q \in L \mid p \leq q\}$ は、 $\rho(p)$ で示される。集合 $\rho(\square)$ は、 ρ^* で示される。

定義5.3: $S \subset L$ を空節 $\square$ を含む構文の集合とする。このとき、 $\rho^* = S$ ならば、 L 上の精密化演算子 ρ は、 S に対し完全であるといわれる。

精密化演算子に関する本来の用法にもどってみよう：仮説が反ばくされるとき、推論アルゴリズムは、その精密化の中で、仮説の置換を探索する。そのため、そのようなアルゴリズムをもつ精密化演算子は、仮説言語 L_h に対して完全になる事が必要とされる。

以下は、モデル推論システムにより利用される具体的な精密化演算子の例であるが、この精密化演算子は、仮説言語の種々のクラスに対し、完全である。

[Reynolds 70]で述べられているように、関係“ $\rightarrow$ ”を定義し、 $\rho_1(p) = \{q \mid p \rightarrow q\}$ により定義される演算子 ρ_1 を考えてみよう。

$p \in L$ を、ある構文とする。このとき、以下の1つが成立するならば、 $p \rightarrow q$ が成立する。

1. $p = \square$ かつ $q = a(X_1, X_2, \dots, X_n)$ である。
ただし、 L の n 引数の述語記号 a に対し、 $n \geq 0$ であり、 $X_1, X_2, \dots, X_n$ は、 n 個の独立変数である。
2. あるアトム P に対し $p = P$ かつ $q = P(U/V)$ である。ただし、 U と V は、 P の中に存在する独立変数である。
3. あるアトム P に対し $p = P$ であり $n \geq 0$ となる、 n 引数の関数記号 F に対し、 $q = P(V/F(X_1, X_2, \dots, X_n))$ である。
ここで、 V は、 P に存在する変数であり、 $X_1, X_2, \dots, X_n$ は、 P に存在する独立変数ではない。

以下では、 L 上の具体的な構造的複雑さに対する尺度である構文のサイズについて、Reynold の定義を採用する：構文 P のサイズ、即ち、 $\text{size}(p)$ は p の記号オカレンス（句読点シンボルを除く）の数であるが、これは、 p に存在する独立変数の数を除く。 ρ_1 が、この尺度に関する精密化演算子であることを示す。

定理 5.4: ρ_1 は、 L 上の精密化演算子であるとき、 L のアトムに対して完全である。

証明: p と q の 2 つの構文について、もし、 p が q インスタンスであれば $p \vdash q$ である。構文のサイズに関する定義より、もし、 $q \in \rho_1(p)$ であれば、 $\text{size}(q) = \text{size}(p) + 1$ である。ここで、 ρ_1 は、構文からそれらの精密化への写像である。 ρ_1 の定義から、次のことが明らかである。任意の $p \in L$ に対し、 $\rho_1(p)$ の集合は、計算可能であり、有限である。ここで、 $\rho_1(p)(n)$ の集合は、任意の $n \leq \infty$ に対し計算可能である。そのため、 ρ_1 は、ある精密化演算子である。

[Reynolds 70] の定理 4 は、任意のアトム $q \in L$ に対し、 $\square$ から q への有限の全 ρ_1 チェインが存在することを示している。これは、 L のアトムに対し精密化演算子の完全性を定めている。■

アトムだけを含む公理系は、相対的につまらない。しかしながら、 ρ_1 のわずかな一般化は、ある精密化演算子を結論づける。この演算子は、前述で与えたいくつかの例を公理化するのに十分余裕がある。上で与えた例には、 $\leq$, plus, 整数リストに対する append があつた。

定義 5.5: ρ と ρ' を L 上の二つの精密化演算子とする。 $\rho * \subset \rho' *$ であれば、 ρ' は、 ρ よりも一般的であるという。

ρ_1 に関し、次の一般化を考えてみよう。 $p \in L$ を、ある構文とする。このとき、次の一つが成立すれば $q \in \rho_2(p)$ である。

1. $q \leftarrow \rho_1(p)$
2. いくつかの n 引数の述語記号 a と項 $t_1, t_2, \dots, t_n$ 及び $q = a(t_1, t_2, \dots, t_n) \leftarrow a(X_1, X_2, \dots, X_n)$ に対して $p = a(t_1, t_2, \dots, t_n)$ である。ここで、 $X_1, X_2, \dots, X_n$ は、独立変数であり、 X_i は t_j の中に存在する ($1 \leq j \leq n$)。

Reynoldsは、形式 $\{P\} \leftarrow \{Q\}$ の構文を、変換と呼んだ。我々は、文脈自由変換で、条件2を満たす変換を呼びだしている。

いくつか意味のある述語は、アトムと文脈自由変換に対するatomic-完全公理化をもっており、モデル推論システムでは、それらを推論するための精密化演算子 ρ_2 が使用されている。

これらの述語は、次のようなものを含む：順序関係、0と後続関数 X' により作成される整数のたし算としては、

```

X ≤ X
X ≤ Y' ← X ≤ Y
plus( 0, X, X)
plus( X', Y, Z') ← plus( X, Y, Z)

```

空節リスト $[]$ とリスト構造 $[X | Y]$ により作成されるリストの前置、後置、後続関係としては、

```

prefix( [], X)
prefix( [A | X], [A | Y]) ← prefix( X, Y)

suffix(X, X)
suffix(X, [A | Y]) ← suffix(X, Y)

subsequence( [], X)
subsequence( [A | X], [A | Y]) ← subsequence(X, Y)
subsequence(X, [A | Y]) ← subsequence(X, Y)

```

リストでの結合関係については、

```

append( [], X, X)
append( [A | X], Y, [A | Z]) ← append( X, Y, Z)

conc( [], X, X)
conc( [A | X], Y, [A | Z]) ← conc( X, Y, Z)

```

$t(X, Y)$ と、いくつかの定数により作成される 2 進木に対する subtree 関係については、 $\text{subtree}(X, X)$

```

subtree(X, X)
subtree(X, t(Y, Z)) ← subtree(X, Y)
subtree(X, t(Y, Z)) ← subtree(X, Z)

```

定理 5.6: ρ_2 は、 L に対する精密化演算子であり、 L のアトムと文脈自由変換に対して、完全である。

証明: $p = \{P\} \leftarrow$ と $q = \{P\} \leftarrow \{Q\}$ とする、任意の二つの構文 p と q に対し、 $p \vdash q$ と、 $\text{size}(p) < \text{size}(q)$ が成立する。よって、 ρ_2 は、構文からそれらの精密化への写像である。任意の $p \in L$ に対し、集合 $\rho_2(\{p\})$ は、有限である。したがって、 $\rho_2(\{p\})(n)$ は、任意の $n \geq 0$ に対して計算可能である。そのため、 ρ_2 は精密化演算子である。

前の証明は、 ρ_1 が L のアトムに対して完全であることを示している。任意の文脈自由変換 $\{P\} \leftarrow \{Q\}$ に対して、 $\square$ から $\{P\} \leftarrow$ への有限な全 ρ_1 -チェーンが存在するが、このチェーンは、また ρ_2 -チェーンでもある。 ρ_2 の定義により $(\{P\} \leftarrow \{Q\}) \in \rho_2(\{P\} \leftarrow)$ が成立する。

ゆえに、 $\square$ から $\{P\} \leftarrow \{Q\}$ への有限な全 ρ_2 -チェーンが存在する。このチェーンは、 L のアトムと文脈自由変換に対して、 ρ_2 の完全性を定めている。 ■

文脈自由変換に関して、二つの単純な一般化がある：ここでは、例を通して、それを紹介する。また、それらについての完全な精密化演算子の定義は、読者にまかせる。このような精密化演算子は、次の例で示される公理化を推論するために、モデル推論システムにより利用された。最初の一般化では、多重文脈自由のホーン構文がある。この構文で、稠密な半順序をもつ述語を公理化できる。すなわち、この述語は、0と終端点1をもち、XとYの間の点として解釈する $m(X, Y)$ をもつ。

$$\begin{aligned} 0 &\leq X \\ X &\leq 1 \\ X &\leq X \\ m(X, Y) &\leq X \leftarrow Y \leq X \\ m(X, Y) &\leq Y \leftarrow X \leq Y \\ m(X, Y) &\leq Z \leftarrow X \leq Z, Y \leq Z \end{aligned}$$

2進木のisomorphismについては、

$$\begin{aligned} &\text{isomorphic}(X, X). \\ &\text{isomorphic}(t(X1, Y1), t(X2, Y2)) \leftarrow \\ &\quad \text{isomorphic}(X1, X2), \text{isomorphic}(Y1, Y2) \\ &\text{isomorphic}(t(X1, Y1), t(X2, Y2)) \leftarrow \\ &\quad \text{isomorphic}(X1, Y2), \text{isomorphic}(Y1, X2) \end{aligned}$$

また、trueとfalseをもつ真偽公理の証明については、

$$\begin{aligned} &\text{satis}(\text{true}) \\ &\text{satis}(\text{and}(X, Y)) \leftarrow \text{satis}(X), \text{satis}(Y) \\ &\text{satis}(\text{or}(X, Y)) \leftarrow \text{satis}(X) \\ &\text{satis}(\text{or}(X, Y)) \leftarrow \text{satis}(Y) \\ &\text{satis}(\text{not}(X)) \leftarrow \text{unsatis}(X) \\ \\ &\text{unsatis}(\text{false}) \\ &\text{unsatis}(\text{or}(X, Y)) \leftarrow \text{unsatis}(X), \text{unsatis}(Y) \\ &\text{unsatis}(\text{and}(X, Y)) \leftarrow \text{unsatis}(X) \\ &\text{unsatis}(\text{and}(X, Y)) \leftarrow \text{unsatis}(Y) \\ &\text{unsatis}(\text{not}(X)) \leftarrow \text{satis}(X) \end{aligned}$$

文脈自由変換の第2の一般化は、補助述語をもつ、項自由変換である。このクラスの範囲内で、かけ算とべき乗算を公理化できる。

```
times(O, X, O)
times(X', Y, Z) ← times(X, Y, X), plus(Y, X, Z)

exponent(O, X, O)
exponent(X', O, O')
exponent(X, Y', Z) ←
    exponent(X, Y, W), times(W, X, Z)
```

リスト反転については

```
reverse([], [])
reverse([A | X], Y) ← reverse(X, Z), conc(Z, A, Y)
```

補集合関係については

```
subset([], [])
subset([A | X], Y) ← subset(X, Y), member(A, Y)
```

挿入ソートについては

```
sort([], [])
sort([A | X], Y) ← sort(X, Z), insert(A, Z, Y)

insert(A, [], [A])
insert(A, [B | X], [A, B | X]) ← A ≤ B
insert(A, [B | X], [B, Y]) ← B ≤ A, insert(A, X, Y)
```

[Angluin & Shapiro 81] では、これらにクラスの公理の表現能力について、十分な特徴づけが報告だろう。

5. 2 精密化グラフ

任意の1階述語 L に対して、 L 上の精密化演算子 ρ は、精密化グラフ $G_\rho (\rho^*, E)$ を導く。これは、直接的な巡回グラフであるが、頂点は、 ρ の構文である。また、 $\text{edge}(p, q) \in E$ が存在するための必要十分条件は、 $q \in \rho(p)$ 、 $p \in \rho^*$ となることである。 $(p, q) \in E$ が、 $\text{size}(p) < \text{size}(q)$ を導出するので、このグラフは、巡回的である。 p と q を ρ^* における二つの構文としよう。 p から q までの G_ρ のパスは、 $(p_i, p_{i+1}) \in E$ とする
 $(0 \leq i < n) p = p_0, p_1, \dots, p_n = q$ の有限例である。 G_ρ のパスが、有限の全 ρ チェーンにである事に注意しよう。 G_ρ で p から q へのパスが存在するとき、 q は、 G_ρ で p から到達可能であるという；もし $(p, q) \in E$ であり、 q が p の後続要素であれば、 p は q の先行要素である。

以下は、 G_ρ に対するいくつかの特性である。

1. 構文 q が構文 $p \in \rho^*$ から G_ρ で到達可能であるための必要十分条件は、 $p \leq q$ である、その結果 $p \vdash q$ に限る。
2. 空文 $\square$ は、任意の $p \in \rho^*$ から到達可能でない。
3. すべての $p \in \rho^*$ は、 $\square$ から到達可能である。
4. もし、構文 $p \in \rho^*$ が L のあるモデル M で真であれば、 G_ρ のすべての後続要素は、 M でも真である；もし p が M で偽であれば G_ρ のすべての後続要素は、 M で偽である。

L で、任意に与えられたモデル M に対し、 M で真なる $\leq$ に関する最小構文の集合を定義する。

定義5. 7: M を L のモデルとし、 ρ を L 上の精密化演算子とする。 p が M とすべての構文 $q \in \rho^*$ で真になれば、構文 $p \in \rho^*$ は、 ρ に関してモデル M のソース構文 (source sentence) と呼ばれる。ここで、 $q \leq p$ は M で偽である。 ρ に関する M のソース構文の集合は、モデル M のソース集合 (source set) と呼ばれ、 $\mathbb{S}$ と示される。

ソース構文の定義により、 G_ρ で空節口から p へのすべてのパスでのすべての構文が M で偽になるならば $p \in \rho^*$ は、ソース構文である。

$p_0, p_1, \dots$ の無限な構文列は、 $i > 0$ に対して $p_i \leq p_{i+1}$ (又は、 $p_{i+1} \leq p_i$) が成立すれば、無限の昇順列 (又は降順) ρ -チェーンと呼ばれる。 $q \in \rho(p)$ は、 $\text{size}(p) < \text{size}(q)$ を導出するので、 L で無限な降順 ρ -チェーンは、存在しない。

ゆえに、すべての $p \in L$ に対し、集合 $\{q \mid q \leq p\}$ が有限であり、ソース構文でない M で真なるすべての構文 $p \in \rho^*$ に対し、 $q \leq p$ を満すソース構文 q が存在する。あるモデルのソース集合は、次の特性をもつ：

定理 5・8： L を 1 階述語の言語とし、 M を L のモデルとする。

ただし $\square \in L_0 \subset L_h \subset L$ が成立する。また、 ρ を L_h に完全な L 上の精密化演算子とする。もし、 $L_h(k_0)$ が、ある $k_0 \geq 0$ に対し、 M の有限な L_0 -完全公理を含むとするならば、 $L_h(k_0)$ は、また、そのような公理化である。

証明： $T \in L_h$ を M の有限な L_0 -完全公理化であるとし、 k_0 を T の任意の構文の最大数とする。もし、 T が $L_h(k_0)$ の部分集合でないとすると、ソース構文でない構文 $p \in T$ が存在する。モデルに関するソース構文の定義により $q \leq p$ となるソース構文 $q \in L_h(k_0)$ が存在する。

T の p を q で置き換えてみよう。 $q \leq p$ であるので、 $q \vdash p$ が結論づけられる。 q が、 M で真で、 T が、 M の L_0 -完全公理化であるので、構文の集合は、 M の L_0 -完全公理化でもある。公理化 T' のすべての構文がソース構文になるまで、置換のステップを繰り返す。 T' は、 M に関する有限の L_0 -完全公理化であり、 $L_h(k_0)$ の部分集合である。 $L_h(k_0)$ は有限であり、 M で真なので、 $L_h(k_0)$ は、また M の有限な L_0 -完全公理化である。■

もし、 T_1 と T_2 が $T_2 \leq T_1$ であるような構文の有限集合であれば、 $T_1 \vdash T_2$ となるような、ある j が存在する。これにより、 M が L の h -easy モデルであるならば、 M のソース集合は、 M に関する $(h'(n) = h(n) + j$ に対し) h' -easy で有限な L_0 -完全公理化を含む。しかしながら、以下の帰納推論アルゴリズムの正当性に対し、 ρ は、比較的強い要求を満足しなければならない：

定義 5・9： 任意の 2 組の有限集合 T_1, T_2 に対し、 $T_2 \vdash_n p$ が $T_1 \vdash_n p$ を導出するならば、精密化演算子 ρ は $\vdash$ に関して、保存されるといわれる。ただし、構文 T_1 と T_2 の関係は、 $T_1 \leq T_2$ であり、 $p \in L$ である。

もし、証明手続きが、分解証明 [Robinson 65] であり、 T_1 が T_2 を吸収 (subsume) すれば、 $T_2 \vdash_n \square$ は、 $T_1 \vdash_n \square$ を導出する。また、導出は、反ばくの途中にあるので、任意の $p \in L$ に対して、 $T_2 \vdash_n p$ は、 $T_1 \vdash_n p$ を導出する。

本論文に記述されているすべての精密化演算子 ρ に対し、関係 $\leq$ は、吸収関係の部分集合であるので、それらは分解証明に関してすべて保存的 (conservative) である。

系 5. 10: 定理 5. 8 と同じ仮定の下で、 ρ が $\vdash$ に関して保存されていれば、 h は、全帰納的であり、 $T \subseteq L_h(k_0)$ は、ある $k_0 \geq 0$ に対し、 M について有限な L_0 -完全で、 h -easy な公理化である。

5. 3 最も一般的な精密化演算子

本節では、任意に与えられた 1 階述語 L に対し、演算子 ρ_0 を定義し、それが L に対する精密化演算子になることを証明する。即ち $\rho^* = L$ である。最も一般的な精密化演算子の存在は、理論的な意味で最も興味がある。実用的に、推論される仮説の集合がもつ構造についての情報は、通常、知られている。そのような場合、このクラスに対して完全で、さほど一般的でない精密化演算子を、いつも選んでいる。これは、興味をもっているモデルについて、比較的効率の良い推論を保証している。

以下では、半順序 $\leq$ が任意の 1 階述語に対して定義されている。 L は、最小要素として $\square$ をもつ。精密化演算子 ρ_0 が定義され、それを計算するためのアルゴリズムのスケッチが与えられる。定理 5. 14 は、 ρ_0 が L について完全であることを示している。また、それは、本節の主な結論である。この結論をのべるために、いくつかのローカルな道具を開発し、 $\leq$ が $\leq$ に等しいことを証明する (定理 5. 17)。

いくつかの定義から始めよう。 $A_1 \theta < A_2$ かつ $B_1 \theta < B_2$ が成立するか、要約して $(A_1 \leftarrow B_1) \theta < A_2 \leftarrow B_2$ が成立すれば構文 $p = A_1 \leftarrow B_1$ が、置換 θ に対して構文 $q = A_2 \leftarrow B_2$ を包含する。 p が q を吸収し、 q が p を吸収するとき $p \approx q$ と書く。関係 $\approx$ は、等価関係である；その下に p の等価クラスを示すために $[p]$ を使う。

包含関係は、空節をもった L の等価クラスの集合に半順序を導く。空節 $\square$ は、最小要素である。また、もし、 p が q を吸収するならば、 $p \vdash q$ である。

例えば、構文 $\{plus(X, Y', Z')\} \leftarrow$ と $\{plus(X, Y', Z'), plus(X, Y', W)\} \leftarrow$ は等価である；それらは、共に $\{plus(X, Y', Z')\} \leftarrow \{plus(Z, Y, Z')\}$ と $\{plus(X, O', X')\} \leftarrow$ を吸収する。

また、これらは、 $\{plus(X, Y, Z)\} \leftarrow$ により吸収される。

$size(p) \leq size(q)$ を導出しないならば、 p は q を吸収する。

これより、 $p \sqsubseteq q$ のような精密化演算子が存在するのは、 p が q を吸収するときに限る。Plotkin [Plotkin 71] は、 $\sqsubseteq$ が成立しないという条件下で、無限に完全な降順チェーンの例を与えている。この包含関係は、次の集合に拡張される：構文 S と T の任意の二つの集合に対し、 S が T を包含するためには、すべての $q \in T$ に対し、 p が q を保存するような $p \in S$ が存在することである。

以下の定義と結果は、Plotkinによる [Plotkin 70, Plotkin 71b, Plotkin 71a]。 $S \subset S'$ と $S \approx S'$ が $S = S'$ を導出すれば、アトム集合 S は、縮小されるといわれる。

いいかえると、 S がそれ自身固有の部分集合に等しくないならば、 S は、縮小される。もし、 A と B が縮小されれば、構文 $p = A \leftarrow B$ は、縮小される。構文は反転するのでもし p と q が縮小され $p \approx q$ であれば、 p と q は変数の名前をつけかえることに等しい。

Plotkinは、与えられた構文 p が $p \approx q$ を満たす縮小された構文 q を計算するアルゴリズムを示している。このアルゴリズムは、それらの等価性に対する構文をしらべるために使用され得る。本節では、‘構文’は、‘この構文に関する等価クラスの縮小代表’を意味する。

定義 5.11: 置換 θ がアトム集合を縮小するというのは、 $|S\theta| < |S|$ を満足するときである。置き換えが構文 $p = A \leftarrow B$ を縮小するというのは、 A 又は B のどちらかが縮小しているときである。

定義 5.12: p と q が L の 2 つの構文であるとする。このとき、 p を縮小しない置換 θ が存在し、 $p\theta \subset q$ が成立するとき $p \leq_0 q$ と書く。

$\leq$ の定義により、 $p \leq q$ は、 p が q を包含するが、その逆は真でない。

空の置換に対しては、 $p \leq p$ であり $p \leq q$ かつ $q \leq p$ は、 p が変数名の付け換えに匹敵して q に等しい。関係 $\leq$ は、推移的である。

それは、非縮小の置換に θ_1 に対し $p_1 \theta_1 \subset p_2$ および非縮小の置換 θ_2 に対する $p_2 \theta_2 \subset p_3$ が、次のことを導出するからである。

$\theta_1 \circ \theta_2$ は、 p_1 と $p_1 \theta_1 \circ \theta_2 \subset p_3$ に対して非縮小の置換である。

また、空節の置き換えに対するすべての $p \in L$ に対して $\square \leq p$ である事に注意しよう。そのため、 $\leq$ は、 L の構文（等価クラス）上の半順序であり、 L は、最小要素として $\square$ をもっている。

上の定義から、二つの縮小構文 p, q が、 $\text{size}(p) < \text{size}(q)$ 導出していることが明確になる。

ここで、 P と Q を L のアトムとし、 U をアトム又は構文の集合とする。 P が U に関して、 Q よりも一般的であるとは、 $P \theta = Q$ かつ $U \theta = U$ が成立する置換 θ が存在することである。 $A \leftarrow B$ を、縮小構文とする。 $A \leftarrow B$ に完全して P よりも一般的であるような任意の Q に対して、構文 $A \leftarrow B \cup \{Q\}$ が縮小されなければ、 $A \leftarrow B \cup \{P\}$ は縮小される。このとき P は、最も一般的であるという。同じことが、 $A \cup \{P\} \leftarrow B$ に対してもいえる。

定義 5.13: $p = A \leftarrow B$ を、 L の縮小構文とする。このとき、以下の 1 つが正しければ、 $q \in \rho_0(p)$ である。

1. $q = p \theta$ 、ここで $\theta = \{V/W\}$ は、 p を縮退しない。また、 V と W は、 p の独立変数である。
2. $q = p \theta$ 、ここで、 $\theta = \{V/f(X_1, \dots, X_n)\}$ は、 p を縮退しない。 f は n 引数の関数であり、 V は、 p の独立変数である。
 $X_i, 1 \leq i \leq n$ は、 p に存在しない独立変数である。
3. $q = A \cup \{P\} \leftarrow B$ 、ここで P は $A \leftarrow B$ に関する最も一般的なアトムである。 $A \leftarrow B$ に対して、 $A \cup \{P\} \leftarrow B$ は縮小される。
4. $q = A \leftarrow B \cup \{P\}$ 、 P は、 $A \leftarrow B$ に関して最も一般的なアトムである。 $A \leftarrow B$ に対して $A \leftarrow B \cup \{P\}$ は、縮小される。

定理5. 14: ρ_0 は、 L 上で完全な精密化演算子である。

証明: $q = \rho_0(p)$ であれば、 p は q を吸収する。
そのため、 $p \vdash q$ が成立する。1と2の場合での置換は、非縮退的であり、それらは、1つずつ p のサイズを増加している。
3と4の場合、あるアトムが、その残りの部分を変えずにその構文に追加されている。これより、 $q \in \rho_0(p)$ は、 $\text{size}(p) < \text{size}(q)$ を導く。
したがって、 ρ_0 は、構文からそれらの精密化への写像である。

1と2は、有限に多くの構文 q にだけ成立するが、一般に、3と4を満足する無限に多くのアトム p が存在する。以下では、いくつかの定められた範囲よりも小さなサイズのアトムをすべてを見つけるために効率的な手続きを記述する。これより $\rho(p)(n)$ の集合は、計算可能である。したがって、 ρ_0 は、 L 上の精密化演算子である。

定理5. 17では、 $p \leq q$ ならば、 p から q への有限の全 ρ_0 -チェーンが存在することを示した。任意の構文 $q \in L$ に対して、 $\square \leq q$ なので、 ρ_0 が L に対して計算可能であることを導く。■

$\rho(p)(n)$ の集合を計算するためのアルゴリズムをスケッチする。
 $\rho(p)(n)$ には、構文 p と $n \leq 0$ が与えられている。このとき、 ρ_0 が完全であることの証明を続ける。1と2の場合に、非縮退置換の計算は、簡単である。また、有限に多くの構文が、この演算により生成される。3と4の場合、 $A \leftarrow B$ に関する最も一般的なアトム P を見つけなければならない。 $A \leftarrow B$ は、精密化される構文であり、結果として得られる構文 $A \cup \{P\} \leftarrow B$ (又は $A \leftarrow B \cup \{P\}$) は縮小される。

この集合は、一般に無限である。例として、次の構文を考える。

$\{p(X, Z)\} \leftarrow \{p(X, f(Y)), p(f(Y), Z)\}$

アトムに関する以下の無限リストは、

$p(f(X), f(Y)), p(f(f(X)), f(Y)),$

$p(f(X), f(f(Y)), p(f(f(X)), f(f(Y))), \dots$

ρ_0 の定義で、4の場合を満足する。しかしながら $\text{size} \leq k$ のすべてのアトムは、以下の方法でシステムチックに生成される:

次の条件を満たすアトム $p(X_1, X_2, \dots, X_n)$ を選択してみよう。

p は、 L の n 引数の述語記号であり、 $X_1, X_2, \dots, X_n$ は、 $A \leftarrow B$ に存在しない独立変数である。

連続的に、 P に存在する独立変数を選択し、以下の演算の1つを実行する：

1. P に存在するが、 $A \leftarrow B$ に存在しない独立変数 U を選択し、 P を、 $P \{V/U\}$ とする。
2. $n \geq 0$ に対する n 引数の関数記号 f を選択し、 P を $P \{V/f(Z_1, Z_2, \dots, Z_n)\}$ とする。
ここで、 Z_i ($1 \leq i \leq n$) は、 P 又は $A \leftarrow B$ に存在しない変数である。
3. $A \leftarrow B$ に存在する変数 U を選択し、 P を $P \{V/U\}$ とする。

$A \cup \{P\} \leftarrow B$ が縮小されるか $\text{size}(P) > k$ になるまで、行なわれる。

もし、2番目の条件が成立すれば、failする。

そうでなければ、構文 $A \cup \{Q\} \leftarrow B$ が任意のアトム Q に対して縮小されないことが示される。このアトム Q は、 $A \leftarrow B$ を変化させない置換に対して $Q \rightarrow P$ を成立させている。もし、以上が示されれば P を返す。

アトム P に演算1及び2を適用したそれぞれの結果は、 $P \rightarrow P'$ でありその場合、 $\text{size}(P') = \text{size}(P) + 1$ である。

3番目の演算は、 $A \leftarrow B$ の変数内で、有限に多くの時間を要す。

これより、このアルゴリズムは、すべての計算パス上で終了される。

演算1と2は L のすべてのアトムを生成するのに十分である。これは、変数名の付け換えに等しい。以上から、この手続きにより $\text{size} \leq n$ のすべてのアトムが生成されることが保証されている。演算3は、 P で変数の付け換えを適用するためのメカニズムを与えている。そのため、このアルゴリズムの連続計算により生成されるアトムの集合は、次のように要求された集合である。

以下の2つの補題と1つの定理は、 ρ_0 が L に対して完全であることを示している。

補題5. 15: p と q をいくつかの置換に対して、 $p\theta = q$ である二つの構文とする。ただし、 θ は、 p を縮小しない。このとき、 p から q への有限な全 ρ_0 チェーンが存在する。

証明：この補題は、Reynold の論文における定理4の一般化である。この定理の証明の吟味は、有限の全チェーンを超えるために P と $P\theta$ に、その定理が適用されることで示される。■

補題5.16: $p = A_1 \leftarrow B_1$ と $q = A_2 \leftarrow B_2$ を $p \subset q$ となる2つの縮小構文とする。このとき、 p から q への有限の全 ρ_0 -チェーンが存在する。

証明: この証明は、 n の帰納による差集合 $A_2 - A_1$ 及び $B_2 - B_1$ のアトム
の総和である。もし、 $n=0$ ならば、 $p=q$ であり、空節チェーンは、補題を満
足する。この補題は、いくつかの $n \leq 0$ に対して、真であると仮定している。ま
た、 $A_2 - A_1$ と $B_2 - B_1$ に $n+1$ 個のアトムが存在している。

P を $A_2 - A_1$ のアトムとし、 A_2 を $A_2 - \{P\}$ とする (対称的引数は、
 $A_2 = A_1$ の場合、 $P \in (B_2 - B_1)$ に適用する)。帰納的仮定により、
 p から $A_2 \leftarrow B_2$ までの全 ρ_0 -チェーン

$p = p_0, p_1, \dots, p_j = A_2 \leftarrow B_2$ が、存在する。

もし、 P が $A_2 \leftarrow B_2$ に関して、最も一般的なアトムであれば、
即ち、 $A_2 \cup \{P\} \leftarrow B_2$ が縮小すれば、 ρ_0 の定義により、

$(A_2 \leftarrow B_2) \leftarrow \rho_0 (A_2 \cup \{P\} \leftarrow B_2)$ であり、この補題が証明される。
それ以外は、 $A_2 \leftarrow B_2$ に関して P よりも一般的であるようなアトム Q が存在する。
また、 $A_2 \cup \{Q\} \leftarrow B_2$ は、縮小される。 P' を最も一般手なアトムとする。
 ρ_0 の定義により、 $(A_2 \cup \{P'\}) \leftarrow B_2 \in \rho_0 (A_2 \leftarrow B_2)$ が成立する。
 P' の選択により

$(A_2 \cup \{P'\} \leftarrow B_2) \theta = (A_2 \cup \{P\} \leftarrow B_2) \theta = A_2 \leftarrow B_2$ が
成立するような置換 θ が存在する。これは、 θ が $A_2 \cup \{P'\} \leftarrow B_2$ を縮退し
ないことを示している。これより、補題5.15より有限の全 ρ_0 -チェーン

$p_{j+1}, p_{j+2}, \dots, p_k$ が存在する。

故に、 $p_0, p_1, \dots, p_j, p_{j+1}, \dots, p_k$ は、補題を満足する。 ■

定理 5.17: $q \leq p$ が成立する為の必要十分条件は、 $p \leq_{\rho_0} q$ である。

証明: もし、 $p \leq_{\rho_0} q$ ならば、 $0 \leq j \leq n$ に対して、 $p_{j+1} \in \rho_0(p_j)$ が成立するような有限の ρ_0 -チェーン $p = p_0, p_1, \dots, p_n = q$ が存在する。 ρ_0 の定義により、すべての i ($0 \leq i < n$) に対し、 $p_i \theta \subset p_{i+1}$ のような非縮退の置換 (あるいは、空節) が存在する。

これにより、 $0 \leq i < n$ に対して、 $p_i \leq p_{i+1}$ が成立し、 $\leq$ の推移性により $p \leq q$ を導く。

もし、 $p \leq q$ のとき、 $p \theta \subset q$ を満たす、置換 θ が存在する。

補題 5.15 により、有限の全 ρ_0 -チェーン $p = p_0, p_1, \dots, p_k = p \theta$ が存在する。したがって、 $p_0, p_1, \dots, p_k, p_{k+1}, \dots, p_n$ は、 p から q への有限の全 ρ_0 -チェーンであり、 $p \leq_{\rho_0} q$ が成立する。■

これは、 ρ_0 が最も一般的な精密化演算子であり、1 階述語 $\mathcal{L}$ に対して完全であることを示す。

6. 一般的な増殖的モデル推論アルゴリズム

枚挙的アルゴリズム 1 に使用されている方法は、最も簡単な説明的推測を選択し、いくつかの事実と一致しないことが発見されるまでそれを保持するための Popper のアイデアに関するだまされやすいインプリメンテーションとして見ることができる。

これが生ずると、この推測に対するある修正を捜さなければならない。即ち、この事実と一致する次の比較的簡単な推測を、選択する。

しかしながら、アルゴリズム 1 に存在するこのアイデアのインプリメンテーションは、たしかにもっともらしくない。さらに良いインプリメンテーションを実行するためには、古い推測を捨て、取り消しから新しい推測を探す代わりに、仮説の古い集合に対する局所的修正を行い、新しい事実に対抗する仮説を加えることが必要である。

この節では、矛盾探索アルゴリズムと精密化演算子の利用方法を記述する。これにより、新しい事実と一致する推測を局所的に修正し、増殖的方法で、せまい範囲で進化するモデル推論アルゴリズムを得る事ができる。

Figure 6-1 は、そのようなアルゴリズムをスケッチしている。

Figure 6-1：増殖的モデル推論アルゴリズムの概要

Tに{□}を、設定する。

repeat.

次の事実を読む。

repeat.

while 推測Tが強すぎる do.

矛盾探索アルゴリズムを適用し、

Tから反ばく仮説を除去する。

while 推測Tが弱すぎる do.

以前に反ばくされた仮説に対する、精密な仮説を、

Tに付け加える。

until 推測Tが強すぎることもなく弱すぎることもない。

(事実の読み込みが続く限り関係する)

forever.

推測が、強すぎるか弱すぎるかどうかという問題は、一般的に非決定的である。また、我々は、いくつかの有限資源をもつ計算により、それを近似することができる。

導出の複雑さという制限の下で、いくつかの固定的な帰納関数を混合すると、任意のh-easyモデルの極限で、推論可能なアルゴリズムを生ずる。この概要に基づくアルゴリズムは、充足的である(定義3.3を見よ)ので、定理3.4は、そのようなアルゴリズムがいくつかの固定的な帰納関数hに関するh-easyモデルだけの極限で推論できると保証している。従って、それは、この集合の最もバワフルなアルゴリズムである。

このアルゴリズムは、他の問題を提出している：理論が弱すぎる時、どこに精密化が付加されるべきか？ このアルゴリズムがいくつか自然な方法で徹底されていれば、このアイデアがインプリメントされる正確な方法は、このアルゴリズムの正当性に対する精神論である。

実用的に、この問題は、重大であるが、この推測に精密化の追加を順序づける良いヒューリスティクスは、このアルゴリズムの収束をかなり高速化くすることになるかも知れない。

本節の残りで、このアルゴリズムのさらに詳細な説明を行う。また、このアルゴリズムが、任意のh-easyモデルの極限で推論されることを示す。いくつかのインプリメンテーション上の論点は、以下で示す。

6. 1 あるモデルに関するソース集合の近似

系5. 10では、あるモデルの L_0 -完全公理化を探すための正しい立場は、ソース集合の中に存在すると述べている。あるモデルのソース集合は、一般に無限である。我々は、それを近似する方法を、任意に要求される程度で示す。

$G_\rho(L_h, E)$ を、いくつかの精密化演算子 ρ の精密化グラフとする。 ρ は、 L_h に関して、完全である。 G_ρ のマーク m は、 L_h の構文のある集合である。 L_h は、'false'とマークされようとしている。 G_ρ のマーク m が、無矛盾マークであるためには、 $p \leq q$ の関係がある'false'とマークされた任意の q と、任意の $p \in L_h$ に対し p は、また'false'とマークされていることである。それが無矛盾なマーク付けであり、'false'とマーク付けられているすべての構文が、 M で偽になれば、マーク付け m は M に無矛盾である。

定義6. 1: m を G_ρ に関し、無矛盾なマーク付けであるとする。構文 $p \in \rho^*$ がマーク付け m のソース構文と呼ばれるためには、 p が'false'とマーク付けされなくかつ $q \leq p$ となるすべての構文 $q \in \rho^*$ が m で'false'とマークされていることである。マーク付け m のソース構文の集合は、マーク付け m のソース集合と呼ばれる。また、この集合は、 L_m と示される。

もし、あるマーク付け m が、モデル M で無矛盾であれば、そのとき、そのマーク付けのソース集合 L_m は、 L の近似である。 L は、そのモデルのソース集合である。近似の概念の量を定める前に、 m がある無矛盾なマーク付けであるとする、集合 $L_m(k)$ が任意の $k \geq 0$ に対し、計算可能になることを示す。これを調べるために、いくつかの構文 $p \in L$ を考えよう。

もし、 p が'false'とマーク付けされているか $\text{size}(p) > k$ であれば、 $p \notin L_m(k)$ である。それ以外は、 $(q, p) \in E$ となるすべての構文 q を考えよ。それらに関して、有限に多くの構文が存在する。もし、それらのすべてが'false'とマーク付けされているならば、 $p \in L_m$ である。それ以外は、 $p \notin L_m$ である。

補題 6. 2 : ρ を L_h に対し、完全な L 上の精密化演算子であるとする。 $G_\rho(L_h, E)$ を ρ の精密化グラフとする。 M を、 L のモデルとする。 m を、 M で無矛盾な G_ρ のマーク付けであるとする。このとき、任意の $k \geq 0$ に対し、 $\models(k) \leq \models_\rho(k)$ が成立する。さらに、 M で偽なる $L_h(k)$ のすべての構文が 'false' とマーク付けされているならば、 $\models(k) = \models_\rho(k)$ である。もし、 L_h に関するすべての偽の構文が、'false' とマーク付けされているとき、 $\models = \models_\rho$ である。

証明：この補題の仮定の下で、 M で真なる任意の構文 $p \in L_h$ に対して $\models \leq \rho(p)$ が成立する。これは、 $\models$ の定義による。特に、いくつかの $k \geq 0$ に対し $\text{size}(p) \leq k$ が成立すれば、そのとき、 $\models(k) \leq \rho(p)$ が成立する。 $\models$ は、 M で真であるので、 $\models(k) \leq \models(k)$ が成立する。

さらに、 M で偽なるすべての構文 $L_h(k)$ は、'false' とマーク付けされていると仮定する。 $\models$ の定義により、 $\models(k)$ は、 M で真である。 M で真である任意の p に対し $\models \leq \rho(p)$ が成立するので、 $\models(k) \leq \models(k)$ が成立する。 $\models(k) \leq \models(k)$ という事実と共に、等式 $\models(k) = \models(k)$ が成立すると、主張する。明らかに、すべての $k \geq 0$ に対し、 $\models(k) = \models(k)$ であれば、 $\models = \models$ が、成立する。

主張の証明： p を $\models(k)$ の構文とする。 $\models(k) \leq \models(k)$ であるので、 $q \leq p$ となるいくつかの $q \in \models(k)$ が存在する。 $\models(k) \leq \models(k)$ であるので、 $p' \leq q$ なるいくつかの $p' \in \models(k)$ が存在する。 $\leq$ の推移性より、 $p' \leq p$ が成立する。あるモデルのソース集合の定義より、 $p' \leq p$ はあり得ない。故に $p' = p$ である。ただし、 p は、 $p = q$ を導出する。この対照的議論を、 $\models(k)$ の構文 p に適用すると、またそれらもまた、この主張を証明する。■

6. 2 増殖的なアルゴリズム

これらの概念を用いて、このアルゴリズムを述べてみよう。
アルゴリズム 3 は、 $\langle L_0, L_h \rangle$ が言語に許される組みであると仮定している。 ρ は、 L_h に対して完全な精密化演算子であり分解証明を保存する。 $\alpha_1, \alpha_2, \alpha_3, \dots$ は、 L_0 のすべての構文の固定的に効果的な枚挙である。 $F_1, F_2, F_3, \dots$ は、 L におけるいくつかのモデル M の枚挙である。 h は、全帰納関数である。

我々は、これらの過程のもとに、以下の定理が成立することを証明する：

定理 6.3： もし、 M が L に対する h -easyモデルならば、アルゴリズム 3 は、極限で M と一致する。

モデル推論アルゴリズムが、極限であるモデルに一致と言われるためには、いくつかの $n > 0$ に対して事実 F_n を読んだ後、 M の L_0 -完全公理化を出力し、異なる推測を2度と再び出力しないことである。

我々は、アルゴリズム 3 が、以下のことを示すことによって、その極限で一致することを証明する。もし、 $\mathbb{L}(k)$ が M に関し、有限な L_0 -完全で、 h -easy公理化である最小の k としての k_0 をもつならば、アルゴリズム 3 は、実際に、 k を k_0 に至るまで増加するがさらにそれ以上、増加しない。

また、 $\mathbb{L}(k_0)$ は、Pointwise に $\mathbb{L}(k_0)$ に収束する。そのようにするために、そのアルゴリズムについて、以下の事実を、定めるが、それは h -easyモデルの枚举に適用されていると仮定している：

- ・ $\mathbb{L}(k_0)$ が L_0 -完全で、かつ M の h -easyモデルであるならば、それは、 k_0 以上の k を増やさない。
 - ・もし、 $\mathbb{L}(k_0)$ が M の h -easyな L_0 -完全公理化となる最小の k を k_0 とすると、少なくとも k_0 に至るまで、 k を増加する。
 - ・もし、 k が k_0 まで増やされ k_0 よりも大きくしないならば、 $\mathbb{L}(k_0)$ は、Pointwise に $\mathbb{L}(k_0)$ に収束する。
- ともに、これらの事実は、定理 6.3 を証明する。

証明：マークの数に対する帰納により、その補題を証明する。最初のマーク付けで、 $\square$ は、‘false’とマークされる。 $\text{edge}(q, \square)$ は、 E の中に1つも存在しないので、このマーク付けは、無矛盾である。また、 $\square$ は、任意のモデルで偽であるので、このマーク付けは、 M で無矛盾である。構文が、矛盾探索アルゴリズムによって反ばくされるならば、構文がこのアルゴリズムによって‘false’とマーク付けされる。いくつかのステージで、 G_0 のマーク付けが M で無矛盾であり、構文 $p \in \mathbb{L}$ が矛盾探索アルゴリズムによって反ばくされると仮定する。 $\mathbb{L}$ の定義により、構文 p が $\mathbb{L}$ で真になるためには、 $(q, p) \in E$ となるすべての構文 $q \in L_h$ がまた‘false’とマークされることである。それ故、無矛盾のマーク付けの定義により、‘false’とし G_0 における p のマーク付けは、無矛盾である。矛盾探索アルゴリズムの正当性により、このマーク付けは、 M で無矛盾である。 ■

アルゴリズム3：増殖的なモデル推論アルゴリズム。

k に0を、 S_{false} に $\{\square\}$ を、 S_{true} に $\{\}$ を、セットする。
 $\square$ を‘false’とマークする。

repeat.

次の事実 $F_n = \langle \alpha, V \rangle$ を読み、 α を S_V に加える。

repeat.

while $\alpha \in S_{false}$ に対し、 $L(k_0) \vdash_n \alpha$ である do.
矛盾探索アルゴリズムを適用し、反ばく仮説‘false’を
マークする。

while $\alpha_i \in S_{true}$ に対し、 $L(k_0) \vdash_{h(i)} \alpha_i$ である do.
 k を1増加する。

until while ループのどちらにもはまらない。

output $L(k_0)$.

forever.

補題6. 5： もし、アルゴリズム3があるモデル M の枚举に適用され、
 $L(k_0)$ が M に関し、 h -easyな L_0 -完全公理化であれば、このアルゴリ
ズムは、多くて k_0 まで k を増やす。

証明：補題6. 4により、このアルゴリズムにより作成されるグラフ G_ρ の任
意のマーク付けは、 M に無矛盾である。

補題6. 2と共に、これは、任意の $k \leq 0$ に対して、 $L(k_0) \leq L(k)$ が
成立することを示している。 ρ は、導出に関して保存されるので、 k を、 k_0 ま
で増加後、第2の条件のwhile ループは、二度と満すことがない。ゆえに、 k は、
 k_0 を超えることはない。■

補題6. 6： h -easyモデル M の枚举に適用されたアルゴリズム3は、実際
にその事実のすべてを読み込む。

証明： M がある h -easyモデルとするならば、内部のrepeatループは、すべての
新しい事実の読み込みをさえぎっている。

新しい事実 F_n がいくつかの $n \leq 0$ に対して読み込まれると仮定する；
2つのwhile ループが実行される全回数を数え上げている。

Mは、 L のh-easyモデルであるので、いくつかの有限な L_0 -完全で、h-easyな公理化をもつ。Tを、Mの公理化とする、即ち、任意の構文に対して、 $T \vdash_{h(i)} \alpha_i$ が成立する。 K_0 をTに関する任意の構文の最大サイズとする。

ρ は、導出に関して保存されるので、系5.10より、 $L^b_h(K_0)$ は、またMのh-easyで、 L_0 -完全公理化でもある。

補題5.6より、 k は、 K_0 を越えて増加しない。すでに、第2のwhile ループは、有限に多くの時間でだけ実行されている。 $L^b_h(K_0)$ は、有限であるので、任意に与えられた時間で $L^b_h(K_0) \subset L^b_h(K_0)$ が成立し、最初のwhile ループのすべての反復は、少なくとも $L^b_h(K_0)$ の構文を'false'とマークする。それは、最初のwhile ループが有限に多くの時間だけで実行され超えることを示す。それ故、内部のrepeatループの反復が終了する。またそれは、有限に多くの反復だけで実行される。■

補題6.7: いくつかの $K_0 \geq 0$ に対して、アルゴリズム3が k を K_0 まで増加し、 K_0 を越えないならば、 $L^b_h(K_0)$ は、Pointwise に、極限で $L^b_h(K_0)$ に収束する。

証明: 補題6.2により、 $\text{size} \leq K_0$ の任意の構文に対して、 D がMで偽ならば、実際にそれは、'false'とマークされる。
もし、 $L^b_h(K_0)$ が偽の構文(即ち、 $L^b_h(K_0) \neq L^b_h(K_0)$)を含めば、実際に矛盾探索アルゴリズムが適用され、いくつかの構文が'false'とマークされる。

任意のマーク付け m に対し、 $L^b_h(K_0) \subset L^b_h(K_0)$ が成立し、その最後は、有限であるので、この過程は、有限に多くの時間に生じ得る。また、 $L^b_h(K_0)$ で、実際にすべて偽の構文が実際マークされる。

k が K_0 まで増加されてから、いくつかのマーク付け m がそのアルゴリズムにより行なわれたとしよう。もし、 $L^b_h(K_0) = L^b_h(K_0)$ ならば、補題が成立する。即ち、背理法によって、 $L^b_h(K_0)$ は、偽の構文を含み、 m は、これ以上変化しないと仮定している。 $L^b_h(K_0) \leq L^b_h(K_0)$ かつ $L^b_h(K_0)$ はMの L_0 -完全公理化であるので、 $L^b_h(K_0) \vdash L^b_h$ が成立する。ここで、 L^b_h は、Mで真なる観測文の集合である。許容される要求(admissibility requirement)により、 $L^b_h(K_0)$ は、いくつかの $i \leq 0$ に対するその虚偽性 $\alpha_i \in L_0$ の証拠をもつ。即ち、いくつかの $j > 0$ に対し、 $L^b_h(K_0) \vdash_j \alpha_i$ が成立し、 α_i はMで偽である。

$n = \max \{i, j\}$ としよう。このとき、 $L_b(k_0) \vdash_n \alpha_i$ が成立し、 n 番目の事実 F_n を読み込んだ後、集合 S_{true} は、 α_i を含む。最初の while ループの条件が満足された時点で、矛盾探索アルゴリズムが呼び出され、いくつかの $p \in L_b(k_0)$ を、反ばくする。この p は、 m が変化しないという仮定と矛盾した状態で、'false' とマーク付けされている。したがって、 $L_b(k_0)$ は、Pointwise に $L_b(k_0)$ へ収束する。

補題 6. 8: $L_b(k_0)$ が M の h -easy で、 L_0 -完全公理化であるとする。最小の k を k_0 とするならば、アルゴリズム 3 は、少なくとも k を k_0 まで増加する。

証明：背理法により、 h -easy モデル M の枚举に適用されているアルゴリズム 3 は k を k' まで増加し、 k' を越えさせていないとする。この k' は $L_b(k_0)$ が M の h -easy でないとしている。補題 6. 7 により、 $L_b(k_0)$ は Pointwise に $L_b(k_0)$ に収束する。

k' の選択により、 $L_b(k')$ は、 M の h -easy でも、 L_0 -完全公理化でもないで $L_b(k') \not\vdash_{h(j)} \alpha_j$ となるいくつかの $j \leq 0$ に対して、 $\alpha_j \in L_0$ が存在する。 n を、最小インデックスとする。

アルゴリズムは、 F_i を読み込む前に k を k' まで増加している。また、 n_0 を $\max \{n, j\}$ とする。このとき、 n_0 番目の事実を読み込んだ後、 $\alpha_i \in S_{true}$ かつ $k = k'$ である。故に、第 2 番目の while ループの条件は、満足され、 k は、 k' まで増加される。これは、我々の仮定に反する。これにより、 k は少なくとも k_0 まで増加し、この補題が証明されている。 ■

定理 6. 3 の証明： M を L の h -easy モデルとする。 k_0 を $L_b(k_0)$ が M の h -easy、 L_0 -完全公理化であるような最小の $k \leq 0$ であるとする。そのような k_0 は、系 5. 10 により存在する。アルゴリズム 3 は、 M の枚举に適用されると仮定する。補題 6. 5 と 6. 8 により。

アルゴリズム 3 は、 k が k_0 まで増加させるが k_0 を越えさせていない。補題 6. 7 により、 $L_b(k_0)$ は、Pointwise に $L_b(k_0)$ に収束する。 $L_b(k_0)$ は有限であるので、アルゴリズム 3 が n 番目の事実を読み込んでから $L_b(k_0) = L_b(k_0)$ となる $n \leq 0$ が存在する。また m と k は、決して変化しない。ゆえに、 n 番目の事実を読み込んだ後、アルゴリズム 3 は、 M について h -easy で、 L_0 -完全公理化を出力する。また、決して再び異なる推測を出力しない。即ち、アルゴリズム 3 は、極限で M と一致する。 ■

定理 6. 3 を証明することにより、アルゴリズム 3 は、アルゴリズムと同じぐらい有能であることがわかる。また、アルゴリズムは、充足的なので、定理 3. 4 はその種の中で最も有能であるといえる。

6. 3 インプリメンテーション上の論点

ホーン理論の推論のためのアルゴリズム 3 の特殊化と、Prolog でのインプリメンテーションは、将来の論文で議論する予定である。ここでは、そのような努力の中に含まれる主な論点を、指摘する。

効率の良いアルゴリズムをインプリメントするために取りかからなければならない、1 つの論点は、推測に対し、ある仮説を追加又は除去するという結果をいかにすばやく計算するかである。

この問題は、推測（それを 'false' とマークすることにより）から、ある仮説の除去が、 S_{true} のある構文を証明不能性にするか否かをテストするための効率的方法を見つけることにある。 S_{true} は、以前、その仮説から証明できていた。また、この問題は、その推測（ k を増加することによる）に対するある仮説の追加が、 S_{false} のある構文を結論づけるかどうかをテストすることにもある。 S_{false} は、以前、仮説から証明不能であった。いまや、この複雑な制限で証明可能になっている。

これらの問題の中で、最初の問題は、仮説と、事実の間の論理的な従属性（logical dependency）を維持することにより、解かれる。即ち、仮説を記録し、事実の導出として利用した（この仕事に対するインプリメンテーション技術の具体的な効果については、[Charniak et al. 80] を参照せよ）。第 2 の問題についての一般的に効果的である解を知らないが、この問題は、 S_{false} でのすべての構文を二度と証明しないことを避けている。モデル推論システムは、仮説がホーン節である場合に適用されているいくつかのヒューリスティクスを混合している。

推測における仮説の数を最適化することは、他の問題を解くことにある。モデル推論アルゴリズムは、推測の中のある仮説の最大のサイズを最小化することを保証している。この推測は、その論理的能力を増加しない多くの不必要な仮説を含むかもしれない。while ループの条件のテストがもつ複雑度が推測中の仮説の数を大きくするので、そのふるまいは、又、アルゴリズムの効率を左右する。

推測を局所的に最適化するための単純なアルゴリズムがある。これは、仮説がホーン節のときである。統計（statistics）が作成されたモデル推論システムのバージョンでは、十分な最適化をまだ実施していない。

即ち、読者は、それにより作成される最終的な推測の中に不必要な仮説をいくつか見つけることができる。

指摘される他の論点は、Oracle (コーザ) がこのシステムに与える情報の量である。

例えば、推論される述語 $\text{sort}(X, Y)$ において、変数のタイプが、整数のリストであるとシステムが認識するならば、構文的には正しいが意味的には正しくない公理の生成は、精密化演算子によりさけることができる。

タイプの記述は、このモデル推論システムの中に、まだ組み込んでいない。

モデル推論アルゴリズムの効率を改善する多くのヒューリスティクスがある。それは、収束のために必要とされる時間と事実の数という意味の効率である。

例えば、仮説言語と観測言語の組立が決定的であれば、即ち、 $\Gamma \vdash \alpha$ であるか否かという問題が任意の $\Gamma \in \mathcal{L}_h$ と任意の $\alpha \in \mathcal{L}_0$ に対して決定的であれば、複雑度の制限 h は、無視される。この場合、 $\mathcal{L}_h$ は、制限された形式 (5 節でのべたすべてのクラスのような形式) のホーン節を、含む。また、この証明手続きは、後方チェーン (Prolog で使われる分解証明についてのある制限された形式) である。即ち、このスタック上でゴールの重複オカレンスに対する簡単なテストが行なわれるだろう。これらに関係するトピックスは、将来の論文でもっと十分に議論するつもりである。

7. 結論

本論文は、事実から理論を推論するための一般的で増殖的なアルゴリズムを与えた。理論的な分析は、複雑な理論的アプローチから帰納推論に至るまでに知られた最も有能なアルゴリズムのいくつかに同等であることを示した。そのインプリメンテーションは、帰納推論と例によるプログラム合成に関して、既存のシステムをしのいでいる。これらの勇気づけは、基礎的なモデルの計算として 1 階述語を使って可能にさせると信じている。

帰納推論の手段としてロジックが成功したいくつかの理由がある：

- ・ロジックは、自然な意味をもつ。もし、チューリング機械が、ある入力で、正しくない結果を計算すれば、その有制限制御の中で、ある推移が無意味な“誤り”になっている。“誤り”の推移を示す候補に対し、この推移を変えずに、そのチューリング機械をいつも直すことができる。これにより、その機械は、この入力の上で、正しくふるまうようになるだろう。

一方、もし、論理的な公理の集合が偽の結論をもつならば、少なくとも、その公理の一つは、自然な意味で、完全に偽である。この事実は、矛盾探索アルゴリズムのようなエラー検出アルゴリズムの存在を示している。

・ロジックは、その構文と意味の間に親密な関係をもつ。例えば、変数のある項で置き換えるか、ある構文の2つの変数をユニファイするならば、いつも、論理的に比較的弱い（又は等しい）構文を結論づける。同じことは、アトムを構文の条件又は結論に付け加えることによっていえる。これは、反ばく仮説の論理的能力（計算能力）を弱くする自然な方法が存在する理由になる。あるいは、別のいかたをすると、これは、自然でかつ計算可能な精密化演算子が存在する理由である。

・ロジックは、単調でかつモジュラーである。公理を追加又は除去することにより公理体系を変えることは、この公理体系に関する表現（計算）能力に明白な効果をもつ；もし、公理を付け加えるならば、さらに多い帰結をもつ（入出力関係）；もし、公理を除去すれば、より少ない帰結をもつ（入出力関係）。あるプログラムに対し、そのような構文上の改造は計算による予言可能な効果をもつが、このような実用的プログラミング言語は、多くない。

・プログラミング言語としてのロジックの特徴は、ロジックと制御の分離である。モデル推論システムの効率化の方法は、プログラム [Kowalski 79b] の“ロジックの成分”だけを推論し、“制御の成分”を特に記述しないで残すことにあると思える。プログラムのロジック成分は、その記述より以上のものを含む。それは、quicksort のロジックプログラムとinsertion sortのプログラム間の差が示している。ロジックプログラムでの可能な制御は、プログラムの最適化の仕事に似ている。プログラムの最適化と例によるプログラム合成の問題は、それらを同時に解くことを止めているように、それ自身、非常にむずかしい。

例による効率的なプログラム合成の仕事は、次の二つの仕事に分離する事を提案する：例によるプログラムの推論（ときどき、効率が悪い）と、プログラムの最適化。その最適化の仕事が、オブジェクトに関し良く定義されたクラス間の写像を行う事であれば、仕事の分離は、遂行を容易にする。それは、ロジックプログラミングに制御を負わせることにより導びかれる写像をもつ場合である。

・ロジックプログラミング言語として、ロジックの特徴は、プログラムとデータの等価性にある。ロジカルな事実は、直ちに、推論されるロジックプログラムの一部分になり得る。この等価性の意味は、リスプのような他のプログラム言語を記述するのに使われる意味よりも強い。この等価性は、モデル推論アルゴリズム又は、モデル推論システムで、帰納推論のさらに理論的論点にまどをしぼるために使用していない。

しかしながら、ユーザと、より優雅に対話するシステムでは、即ち、McCarthy's の Advice Taker [McCarthy 63] の精神では、そのような特性は、非常に有効であるかもしれない。そのようなシステムは事実による理論の推論だけの指向性ばかりでなく、人間のアドバイスによる対話的な理論の創成、獲得、虫取りの指向性もある。“低レベル”の記述（真のグランドアトム）と、“高レベル”の助言（提案された仮説）を記述する均質な言語は、すべての人間の助言を均質に取り扱うアルゴリズムとシステムを考案する力を与えるだろう*。

〔謝辞〕

省略。

1. モデル推論システムの効率

第5回目のインプリメントにチャレンジ中。

- ・コンカレントな推論。
- ・意味的にまちがった仮説の生成を排除するために、タイプの記述を、精密化演算子で利用する。
- ・推測を局所的に最適化するアルゴリズムと、仮説の追加と問合せ数の最小化に対する improved なヒューリスティクス。

以下、原文38ページ ～ 50ページまで省略。

- ・実行例。
- ・参考文献など。

* Truth Maintenance System, Belief System を包含する知識獲得システムとして、まとめあげられる。

Acknowledgements.

The development of the ideas in this paper spans into the past a bit further than NSF grant MCS8002447, which I acknowledge.

While in high school, Moshe Kroy introduced me to Logic and Philosophy of Science. His exciting lectures on Popper's philosophy were a traumatic experience in my intellectual development. Dr. Kroy suggested that if we could test epistemological theories and theories of the mind, then they would become science, rather than metaphysics, and that it might be possible to carry out this task with a computer. This may be the reason for my leaving the homeland, philosophy, for the promised land of computer science.

Years later, by then an undergrad in Tel Aviv University, Amiram Yehudai told me that I did not have to invent Inductive Inference, since such a discipline already existed. He pointed me to the Blums' paper, which Jonathan Stavi helped me to decipher. Prof. Stavi also helped me relate Gold's and Blums' results to inductive inference in logic.

At Yale, I could not have been luckier having Dana Angluin as my advisor. Truly thanking her would make this acknowledgement longer than the paper, so I reserve the honor to for a forthcoming, lengthier document. I will just say that there is nothing more trustable than a proof Dana says is OK. If anything in this paper appears to be to the contrary, it is probably because I rewrote it afterwards.

Drew McDermott provided crucial insights into logic hacking, in moments of baffle and bewilderment. Although he preferred my system learning about arches and non-arches, rather than inferring atomic-complete axiomatizations of dense partial orders, he was, nevertheless, encouraging all the way.

Bob Kowalski warned me, just in time, that there will be Sirens in my way, singing "I am a cute, good-looking, interesting, quick-shot, irrelevant theoretical question", and that I should never succumb. So whenever I have imagined hearing one, I plugged my ears and did some more Prolog hacking, until it was over. I have failed only once (twice?), and this is a lesson I will never forget. At least until next time.

Even if I have disturbed him in midst of the most intense hack, Bob Nix was always willing to answer any question, be it on English grammar or combinatorics, or listen to my latest new idea. He read the final draft very carefully, and made sure that there were no errors left.

Ernie Davis and Chris Riesbeck read a previous draft, and provided audible feedback. The fact that even Ernie got lost in Section 6 was a conclusive evidence to me that I should rewrite it. I hope it passes Ernie's Test this time.

Without the outstanding implementation of Prolog we have received from David Warren at Edinburgh, this work would have remained yet another theoretical exercise. The pace with which their system matured into a practical and efficient programming tool can not be attributed to the wonders of Logic alone.

Finally, our Tools group provided a comfortable environment for both hacking and writing, and I am thankful for that. Special thanks to Steve Wood, whom I bugged until Z did (almost) exactly what I wanted, and to John Ellis, who wrote the Scribe interface to the typesetter, and then answered a hundred questions about how to use it.

I. Performance of the Model Inference System

The Model Inference System is implemented in the programming language Prolog [Pereira et al. 78]. It was developed hand in hand with the algorithms, problems in the former motivating improvements in the later. The system has had four stable incarnations so far. The first inferred ground complete axiomatizations of regular sets; the second inferred atomic complete axiomatizations of such sets; and the third inferred atomic complete axiomatizations of a first order language with an arbitrary alphabet, provided the axioms were context-free transformations. The fourth incarnation of the system enabled the inclusion of more general refinement operators, and could infer multiple context-free Horn sentences and term-free transformations with auxiliary predicate. So far it has succeeded in inferring atomic complete axiomatizations of dense partial order with endpoints, binary tree isomorphism, list reversal and list subset, multiplication and exponentiation, and satisfiability of boolean formulas.

A fifth implementation is on its way, featuring the possibility of concurrent inference of arbitrarily many programs; typing specification, to be used by the refinement operators to restrict the generation of semantically wrong hypotheses; an algorithm that locally optimizes the conjecture, and improved heuristics for adding hypotheses and minimizing the number of queries (experiments).

Following are examples of application of the different incarnations of the system, with some relevant statistics. In the first two incarnations the programs were all interpreted under the Edinburgh Prolog-10 interpreter version 1.32. In the third and fourth the programs were compiled using the in-core compiler of the Edinburgh Prolog-10 version 3. The computer is a DECsystem-2060, running Prolog-10 under the compatibility package PA1050. In the following, the running time in CPU seconds measured the time the system took to converge to a correct and sufficient axiomatization of the intended model. The decision to terminate the inference process was made by me, since it is an inherent property of an inductive inference algorithm that it can never tell whether it has converged. Note that since this decision was subjective, there may be some errors in the axiomatization below (I am aware of one). The reader is encouraged to find them. *Facts* are the number of facts read therein, and *hypotheses generated* is the number of hypotheses generated by the refinement operator during the inference.

In the earlier versions of the system it would simply go through all possible atoms, in increasing order of complexity, asking the user (or a built-in oracle for the model, when the user was tired) whether they are true in the model. The inefficiency of this approach became more evident as the models got more and more complex, since true atoms are scarce in such models. Therefore in later versions of the system the user is capable of specifying some initial set of facts about the models, hence these implementations were more "fact efficient".

I.1 Inferring Regular Sets.

The first order language for this task consisted of the two predicates $\text{in}(X)$ and $\text{out}(X)$. The terms of the language were the null string Λ , and strings constructed from a variable or Λ and the successor functions $0(X)$ and $1(X)$. For convenience these were simulated by lists of 0's and 1's. The hypotheses language contained Horn clauses, where the only structural restriction on them was that the size of the terms in the condition would not exceed the size of the term in the conclusion. This restriction releases us from a need to impose complexity bound on the size of the proofs. In the following axiomatizations, the number associated with an axiom is its sequential number in the order the axioms were generated by the refinement operator.

Inferring Ground Complete Axiomatizations of Regular Sets.

Strings of even number of 1's and even number of 0's.

157 CPU seconds, 64 facts, 613 hypotheses generated.

(in(Δ) $\rightarrow$ true), 3
 (out(0) $\rightarrow$ true), 9
 (out(1) $\rightarrow$ true), 23
 (out(0X) $\rightarrow$ in(X)), 12
 (in(00) $\rightarrow$ true), 37
 (out(0X) $\rightarrow$ in(1X)), 14
 (in(11) $\rightarrow$ true), 76
 (out(1X) $\rightarrow$ in(X)), 26
 (in(00X) $\rightarrow$ in(X)), 40
 (out(00X) $\rightarrow$ out(X)), 52
 (out(01X) $\rightarrow$ in(X)), 171
 (out(01X) $\rightarrow$ in(0X)), 173
 (out(11X) $\rightarrow$ out(X)), 91
 (in(01X) $\rightarrow$ in(10X)), 191
 (out(1X) $\rightarrow$ in(0X)), 28
 (in(10X) $\rightarrow$ in(01X)), 70
 (out(10X) $\rightarrow$ out(01X)), 289
 (in(11X) $\rightarrow$ in(X)), 79
 (in(010X) $\rightarrow$ in(001X)), 219
 (in(011X) $\rightarrow$ in(0X)), 305

Strings of even parity.

16 CPU seconds, 20 facts, 102 hypotheses generated.

(in(Δ) $\rightarrow$ true), 3
 (in(0) $\rightarrow$ true), 9
 (out(1) $\rightarrow$ true), 23
 (in(0X) $\rightarrow$ in(X)), 12
 (out(0X) $\rightarrow$ out(X)), 20
 (out(1X) $\rightarrow$ in(X)), 26
 (in(1X) $\rightarrow$ out(X)), 34

Strings that contain 11 as a substring.

41 CPU seconds, 50 facts, 141 hypotheses generated.

(out(Δ) $\rightarrow$ true), 3
 (out(1) $\rightarrow$ true), 6
 (in(11X) $\rightarrow$ true), 18
 (out(0) $\rightarrow$ true), 23
 (out(0X) $\rightarrow$ out(X)), 27
 (in(0X) $\rightarrow$ in(X)), 45
 (out(10) $\rightarrow$ true), 72
 (out(10X) $\rightarrow$ out(X)), 76
 (in(1X) $\rightarrow$ in(X)), 19

Strings that contain 11 or 00 as a substring.

25 CPU seconds, 31 facts, 188 hypotheses generated.

(out(Δ) $\rightarrow$ true), 3
 (out(0) $\rightarrow$ true), 6
 (in(00X) $\rightarrow$ true), 17
 (out(1) $\rightarrow$ true), 23
 (in(11X) $\rightarrow$ true), 32
 (in(0X) $\rightarrow$ in(X)), 19
 (out(10) $\rightarrow$ true), 67
 (in(1X) $\rightarrow$ in(X)), 33
 (out(10X) $\rightarrow$ out(0X)), 79
 (out(01X) $\rightarrow$ out(1X)), 49

Strings that are of even number of ones or contain 00 as a substring.

64 CPU seconds, 69 facts, 383 hypotheses generated.

```
(in( $\Lambda$ )  $\rightarrow$  true), 3
(in(0)  $\rightarrow$  true), 9
(in(00X)  $\rightarrow$  true), 10
(out(1)  $\rightarrow$  true), 23
(in(1X)  $\rightarrow$  out(X)), 34
(in(0X)  $\rightarrow$  in(X)), 12
(out(10)  $\rightarrow$  true), 73
(in(100X)  $\rightarrow$  true), 93
(out(01)  $\rightarrow$  true), 37
(out(01X)  $\rightarrow$  out(1X)), 49
(in(11X)  $\rightarrow$  in(X)), 53
(out(11X)  $\rightarrow$  out(X)), 111
(in(10X)  $\rightarrow$  in(01X)), 98
(out(101X)  $\rightarrow$  out(X)), 175
```

Inferring Atomic Complete Axiomatizations of Regular Sets.

The inefficiency of inferring ground complete axiomatizations of regular sets was due not so much to the algorithm as to the fact that I was looking at the wrong problem. Logical theories are essentially non-deterministic computational devices. A non-deterministic Turing machine has the privilege of succeeding in a computation only when it is accepting a string. For rejecting a string it merely has to fail to accept it, and there is no need for it to explicitly succeed in a rejecting computation, a requirement that a ground complete theory of a set has to satisfy. This consideration naturally led to the reformulation of the regular set inference problem: find an *atomic complete* axiomatization of the intended model, that is an set of axioms that prove $\text{in}(s)$ for every s in the intended regular set, and fail to prove $\text{in}(s)$ for strings that are not in that set. The first order language for this task consists of only one predicate $\text{in}(X)$, with the same terms as before.

The results of the previous application of the system suggested that regular sets can be axiomatized with transformations only. Therefore the hypothesis language was restricted to such axioms. Applying the system to the inference of atomic-complete axiomatizations of the regular sets as before led to remarkable improvements in the running time of the system, the number of facts needed and the number and complexity of the axioms found.

Strings of even number of 1's and even number of 0's.

28 CPU seconds, 35 facts, 75 hypotheses generated.

```
(in( $\Lambda$ )  $\rightarrow$  true), 4
(in(00X)  $\rightarrow$  in(X)), 13
(in(11X)  $\rightarrow$  in(X)), 19
(in(010X)  $\rightarrow$  in(1X)), 41
(in(011X)  $\rightarrow$  in(0X)), 48
(in(100X)  $\rightarrow$  in(1X)), 61
(in(101X)  $\rightarrow$  in(0X)), 67
```

Strings of even parity.

8 CPU seconds, 17 facts, 29 hypotheses generated.

```
(in( $\Lambda$ )  $\rightarrow$  true), 4
(in(0X)  $\rightarrow$  in(X)), 5
(in(11X)  $\rightarrow$  in(X)), 24
(in(10X)  $\rightarrow$  in(1X)), 15
```

Strings that contain 11 as a substring.

11 CPU seconds, 35 facts, 12 hypotheses generated.

(in(11X) — true), 7
(in(0X) — in(X)), 9
(in(1X) — in(X)), 5

Strings that contain 11 or 00 as a substring.

14 CPU seconds, 52 facts, 12 hypotheses generated.

(in(00X) — true), 6
(in(11X) — true), 11
(in(0X) — in(X)), 5
(in(1X) — in(X)), 9

Strings that are of even number of ones or contain 00 as a substring.

10 CPU seconds, 18 facts, 24 hypotheses generated.

(in(A) — true), 4
(in(0X) — in(X)), 5
(in(00X) — true), 6
(in(100X) — true), 16
(in(11X) — in(X)), 19
(in(10X) — in(1X)), 15

I.2 Inferring Atomic Complete Axiomatizations of Simple Models.

In the first two stages of the development of the system described above, the first order language used by the system was fixed. In the third stage the language became part of the input to the system. The input is the function and constant symbols of L , and the predicate $p(X_1, X_2, \dots, X_1)$ to be inferred. This version used the refinement operator ρ_2 (page 24), which is complete for atoms and context-free transformations. The fourth version incorporated some more general refinement operators, which were complete for multiple context-free Horn sentences and term-free transformations with auxiliary predicate. In case the system inferred an axiomatization that uses an auxiliary predicate (such as times, which uses plus), the predicate to be used was also specified in the input. The statistics below are on the performance of this version. The two numbers following the axioms are their sequential numbers as they were generated by the refinement operator, and the number of refinement operations needed to generate them.

Arithmetic.

Ordering Relation.

$le(X, Y)$ — X is less than or equal to Y .

2 CPU seconds, 10 facts, 14 hypotheses generated.

(le(X, s(Y)) — le(X, Y)), 10, 2
(le(X, X) — true), 1, 1
(le(0, X) — true), 2, 1

Addition.

$plus(X, Y, Z)$ — X plus Y is Z .

5 CPU seconds, 13 facts, 67 hypotheses generated.

(plus(X, X, 0) — true), 11, 1
(plus(0, X, X) — true), 14, 1
(plus(s(X), Y, s(Z)) — plus(X, Y, Z)), 52, 3

Multiplication, using addition.

times(X,Y,Z) — X times Y is Z.

20 CPU seconds, 13 facts, 205 hypotheses generated.

(times(0,X,0) — true), 12, 1
 (times(s(X),Y,Z) — times(X,Y,W),plus(W,Y,Z)), 61, 2

Exponentiation, using multiplication.

exp(X,Y,Z) — X to the Y is Z.

28 CPU seconds, 18 facts, 252 hypotheses generated.

(exp(X,s(Y),Z) — exp(X,Y,W),times(W,X,Z)), 129, 2
 (exp(X,0,s(0)) — true), 78, 2
 (exp(0,s(X),0) — true), 148, 2

Dense Partial Order with Endpoints.

Partial ordering.

le(X,Y) — X to less than or equal to Y.

10 CPU seconds, 60 facts, 61 hypotheses generated.

Note: m(X,Y) is interpreted as some point between X and Y.

(le(X,X) — true), 7, 1
 (le(X,m(X,1)) — true), 21, 3
 (le(X,1) — true), 5, 1
 (le(0,X) — true), 1, 1
 (le(m(X,Y),Z) — le(Y,Z),le(X,Z)), 57, 4
 (le(m(X,Y),Y) — le(X,Y)), 56, 3
 (le(m(X,Y),X) — le(Y,X)), 58, 3
 (le(m(X,0),X) — true), 49, 3

List Processing Programs.

Member.

member(X,Y) — X is a member of the list Y.

2 CPU seconds, 23 facts, 13 hypotheses generated.

(member(X,[Y|Z]) — member(X,Z)), 6, 3
 (member(X,[X|Y]) — true), 7, 3

Prefix.

prefix(X,Y) — X is a prefix of Y.

4 CPU seconds, 17 facts, 52 hypotheses generated.

(prefix([],X) — true), 1, 1
 (prefix([X|Y],[X|Z]) — prefix(Y,Z)), 50, 5

Suffix.

suffix(X,Y) — X is a suffix of Y.

3 CPU seconds, 17 facts, 26 hypotheses generated.

(suffix(X,[Y|Z]) — suffix(X,Z)), 19, 3
 (suffix(X,X) — true), 5, 1
 (suffix([],X) — true), 1, 1

Subsequence.

$\text{subsequence}(X,Y)$ — X is a subsequence of Y .

8 CPU seconds, 50 facts, 63 hypotheses generated.

($\text{subsequence}(X,[Y|Z]) \rightarrow \text{subsequence}(X,Z)$), 22, 3
 ($\text{subsequence}(X,[Y|X]) \rightarrow \text{true}$), 20, 3
 ($\text{subsequence}([],X) \rightarrow \text{true}$), 1, 1
 ($\text{subsequence}([X|Y],[X|Z]) \rightarrow \text{subsequence}(Y,Z)$), 51, 5
 ($\text{subsequence}([X],[X|Y]) \rightarrow \text{true}$), 45, 5

Subset, using member.

$\text{subset}(X,Y)$ — X is a subset of Y .

31 CPU seconds, 82 facts, 152 hypotheses generated.

($\text{subset}(X,[Y|Z]) \rightarrow \text{subset}(X,Z)$), 34, 3
 ($\text{subset}(X,[Y|X]) \rightarrow \text{true}$), 32, 3
 ($\text{subset}([],X) \rightarrow \text{true}$), 1, 1
 ($\text{subset}([X|Y],Z) \rightarrow \text{subset}(Y,Z), \text{member}(X,Z)$), 113, 3
 ($\text{subset}([X|Y],[X|Z]) \rightarrow \text{subset}(Y,Z)$), 103, 5
 ($\text{subset}([X],[X|Y]) \rightarrow \text{true}$), 97, 5

Append.

$\text{append}(X,Y,Z)$ — Z is the result of appending X to Y .

10 CPU seconds, 33 facts, 106 hypotheses generated.

($\text{append}(X,[],X) \rightarrow \text{true}$), 12, 1
 ($\text{append}([],X,X) \rightarrow \text{true}$), 11, 1
 ($\text{append}([X|Y],Z,[X|W]) \rightarrow \text{append}(Y,Z,W)$), 99, 5

Conc.

$\text{conc}(X,Y,Z)$ — Z is the result of concatenating the symbol Y to the list X .

7 CPU seconds, 29 facts, 66 hypotheses generated.

($\text{conc}([],X,[X]) \rightarrow \text{true}$), 26, 3
 ($\text{conc}([X|Y],Z,[X|W]) \rightarrow \text{conc}(Y,Z,W), \text{true}$), 64, 5

Reverse, using conc.

$\text{reverse}(X,Y)$ — X is the reverse of Y .

6 CPU seconds, 13 facts, 104 hypotheses generated.

Note: The axiomatization of *reverse* using *append* (page 4) is not term-free, hence could not be inferred using the refinement operators currently implemented:

($\text{reverse}([],[]) \rightarrow \text{true}$), 6, 1
 ($\text{reverse}([X|Y],Z) \rightarrow \text{reverse}(Y,U), \text{conc}(U,X,Z)$), 32, 3

Last.

$\text{last}(X,Y)$ — X is the last element in the list Y .

4 CPU seconds, 12 facts, 18 hypotheses generated.

($\text{last}(X,[X]) \rightarrow \text{true}$), 17, 3
 ($\text{last}(X,[Y|Z]) \rightarrow \text{last}(X,Z)$), 8, 3

Double every second element.

$\text{dbl2nd}(X,Y)$ — Y is the list X , where every other element occurs twice.

64 CPU seconds, 49 facts, 229 hypotheses generated.

```
(dbl2nd([],[]) = true), 6, 1
(dbl2nd([X,Y|Z],[X,Y,Y|W]) = dbl2nd(Z,W)), 222, 11
(dbl2nd([X],[Y]) = true), 97, 5
```

Iota, or tails.

$\text{iota}(X,Y)$ — Y is the list of tails of X .

18 CPU seconds, 30 facts, 149 hypotheses generated.

```
(iota([],[]) = true), 6, 1
(iota([X|Y],[X|Y|Z]) = iota(Y,Z)), 100, 7
(iota([X|Y],[X|Z])) = iota(Y,Z)), 97, 7
```

Pack, or one level flatten, using append.

$\text{pack}(X,Y)$ — Y is the list of elements of the lists of X .

9 CPU seconds, 24 facts, 105 hypotheses generated.

```
(pack(X,[]) = true), 3, 1
(pack([X|Y],Z) = pack(Y,W),append(X,W,Z)), 36, 3
```

Oddp.

$\text{Oddp}(X)$ — X is of odd length.

3 CPU seconds, 13 facts, 17 hypotheses generated.

```
(oddp([X,Y|Z]) = oddp(Z)), 14, 5
(oddp([X]) = true), 5, 3
```

Pair.

$\text{pair}(X,Y,Z)$ — Z is a list of pairs of elements of X and Y .

62 CPU seconds, 61 facts, 226 hypotheses generated.

```
(pair(X,[],X) = true), 16, 1
(pair([X|Y],[Z|W],[p(X,Z)|U]) = pair(Y,W,U)), 211, 9
```

Unlabeled Binary tree predicates

Subtree relation.

$\text{subtree}(X,Y)$ — X is a subtree of Y .

9 CPU seconds, 40 facts, 32 hypotheses generated.

Note: $t(X,Y)$ is interpreted as a binary tree whose left subtree is X and whose right subtree is Y .

```
(subtree(X,t(Y,Z)) = subtree(X,Y)), 23, 3
(subtree(X,t(X,Y)) = true), 20, 3
(subtree(X,t(Y,Z)) = subtree(X,Z)), 24, 3
(subtree(X,t(Y,X)) = true), 21, 3
(subtree(X,X) = true), 5, 1
(subtree(0,X) = true), 1, 1
```

Tree isomorphism.

$\text{isotree}(X,Y)$ — X is isomorphic to Y .

20 CPU seconds, 110 facts, 117 hypotheses generated.

Note: $t(X,Y)$ is interpreted as a binary tree whose left subtree is X and whose right subtree is Y .

```

(isotree(X,X) → true), 5, 1
(isotree(t(X,Y),t(Z,W)) → isotree(Y,W),isotree(X,Z)), 108, 6
(isotree(t(X,Y),t(Z,Y)) → isotree(X,Z)), 103, 5
(isotree(t(X,Y),t(Z,W)) → isotree(Y,Z),isotree(X,W)), 41, 6
(isotree(t(X,Y),t(Z,X)) → isotree(Y,Z)), 44, 5
(isotree(t(X,Y),t(Y,Z)) → isotree(X,Z)), 38, 5
(isotree(t(X,Y),t(Y,X)) → true), 35, 5

```

Satisfiability of Boolean Formulas.

Satisfiability, using unsatisfiability.

satis(X) → X is satisfiable.

10 CPU seconds, 40 facts, 66 hypotheses generated.

```

(satis(1) → true), 2, 1
(satis(not(X)) → unsatis(X)), 14, 3
(satis(not(0)) → true), 7, 2
(satis(and(X,Y)) → satis(Y),satis(X)), 54, 4
(satis(and(1,X)) → true), 39, 3
(satis(or(X,Y)) → satis(X)), 33, 3
(satis(or(X,Y)) → satis(Y)), 35, 3
(satis(or(X,1)) → true), 28, 3
(satis(or(1,X)) → true), 23, 3

```

Unsatisfiability, using satisfiability.

unsatis(X) → X is unsatisfiable.

10 CPU seconds, 45 facts, 67 hypotheses generated.

```

(unsatis(0) → true), 1, 1
(unsatis(not(X)) → satis(X)), 14, 3
(unsatis(not(1)) → true), 8, 2
(unsatis(and(X,Y)) → unsatis(X)), 63, 3
(unsatis(and(X,Y)) → unsatis(Y)), 65, 3
(unsatis(and(X,0)) → true), 57, 3
(unsatis(and(0,X)) → true), 52, 3
(unsatis(or(X,Y)) → unsatis(Y),unsatis(X)), 39, 4
(unsatis(or(0,0)) → true), 38, 3

```

II. Glossary of Symbols.

.	function composition
$\vdash, P \dashv Q$	implication, P if Q
$\vdash$	covering relation
$\subset$	subset or equal
$\subsetneq$	strict subset
$\cup$	set union
$\cap$	set intersection
$\geq$	greater than or equal to
$\leq, \not\leq$	{not} less than or equal to
$\&$	and
$\neq$	not equal
$\sim$	not
$\vdash, p \vdash q, \vdash$	p {not} derives q
$\vdash_n, \vdash_n, p \vdash_n q$	p {not} derives q in n derivation steps or less
$\square$	the empty sentence
$\equiv, \not\equiv$	{not} equivalent
α, α_1	observational sentences
θ, θ_1	substitutions
ϵ, ϵ	set membership
Φ_k	the k^{th} step counting function
Δ	the null string
ρ, ρ_1, ρ_0	refinement operators
$\leq_\rho, <_\rho, \not\leq_\rho, \not<_\rho$	partial order induced by ρ
$\rho(p)$	the refinements of p generated by ρ
$\rho^*(p)$	the transitive closure of p under ρ
ρ^*	the transitive closure of $\square$ under ρ
G_ρ	the refinement graph induced on ρ^* by ρ
$\leq_{\alpha_1} <_{\alpha_1} \not\leq_{\alpha_1} \not<_{\alpha_1}$	a partial order defined on sentences
L	a first order language
L_h	an hypothesis language
L_o	an observational language
L_o^M	observational sentences true in M
M	a model
L_o^M	the source set of a model
m	a marking
L_o^m	the source set of a marking
p, p_1, q_1	sentences
T, S, T_1, T_n	sets of sentences
A_1, A_2, B_1, B_2	sets of atoms
P, P_1, Q, Q_1	predicates
t, f, t_1, f_1	terms
X, Y, X_1, Y_1	variables
$h, h(n)$	complexity bound

III. A Useless Lemma.

During the attempts to prove that ρ_* (defined in section 5.3) is a most general refinement operator, I proved the following lemma. This lemma turned out not to be of any use in the proof. However, since I worked so hard to find its right formulation, and then to prove it, I feel the paper will not be complete without it.

Lemma 7.1: Let $S' = S \cup \{P\}$ be a set of atoms such that S is reduced and $P \notin S$. Then S' is not reduced if and only if exactly one of the following holds:

1. There is a substitution θ and an atom $Q \in S$ such that $Q\theta = P$, $P\theta = P$ and $S'\theta \subset S' - \{Q\}$.
2. There is substitution θ such that $P\theta \in S$ and $S\theta \subset S$, or, in other words, $S'\theta \subset S$.

Proof: The *if* direction is a straightforward application of the definition of a reduced set: if one of the cases holds for some substitution θ then $S'\theta \subset S'$, and S' is not reduced.

To prove the *only if* direction, assume that $S \cup \{P\}$ is not reduced. Then there there is a substitution θ such that $S'\theta \subset S'$. There are two cases:

Case 1: $P = P\theta^k$ for some $k > 0$. Since $S'\theta \subset S'$ implies $S'\theta^k \subset S'$, at least one atom $Q \in S$ is mapped into P by θ^k , that is $Q\theta^k = P$, otherwise $S\theta^k \subset S$, in contradiction to the assumption that S is reduced. The substitution θ^k satisfies the first clause of the lemma.

Case 2: $P \neq P\theta^k$ for all $k > 0$. If no $Q \in S$ is mapped into P by θ , that is $Q\theta \neq P$ for all $Q \in S$, then θ satisfies the second clause of the lemma. Otherwise, let $Q_1, Q_2, \dots, Q_n$ be all the atoms in S mapped into P by some multiple of θ , that is, $Q_i\theta^{k_i} = P$ for some $k_i > 0$. For every such atom Q_i , k_i is the only multiple of θ by which it is mapped into P , otherwise there is a k such that $P\theta^k = P$, contrary to the assumption of this case. Let $k = \max\{k_1, k_2, \dots, k_n\} + 1$. Then $P \notin S'\theta^k$, hence $S'\theta^k \subset S$, and the second clause of the lemma holds. ■

References

- [Angluin & Hoover 80]
Dana Angluin and Douglas N. Hoover.
Regular Prefix Relations.
Submitted.
1980.
- [Angluin & Shapiro 81]
Dana Angluin and Ehud Shapiro.
Structural Complexity vs. Expressive Power of Logic Programs.
In preparation.
1981.
- [Angluin 78]
Dana Angluin.
On the complexity of minimum inference of regular sets.
Information and Control 39:337-350, 1978.
- [Bierman 78]
Alan W. Bierman.
The Inference of Regular Lisp Programs from Examples.
IEEE transactions on Systems, Man, and Cybernetics 8, August, 1978.
- [Blum & Blum 75]
Lenore Blum and Manuel Blum.
Towards a Mathematical Theory of Inductive Inference.
Information and Control 28, 1975.
- [Case & Smith 81]
Case J. and Smith C.
Comparison of Identification Criteria for Mechanized Inductive Inference.
to appear.
1981.
- [Charniak et al. 80]
E. Charniak, C. K. Riesbeck, D. V. McDermott.
Artificial Intelligence programming.
Lawrence Earlbaum Associates, Hillsdale, New-Jersey, 1980.
- [Duhem 54]
Pierre Duhem.
The Aim and Structure of Physical Theory.
Princeton University Press, 1954.
Originally published in French in 1906.
- [Gold 65]
E. Mark Gold.
Limiting Recursion.
Journal of Symbolic Logic 30, 1965.
- [Gold 67]
E. M. Gold.
Language identification in the limit.
Information and Control 10:447-474, 1967.
- [Gold 78]
E. Mark Gold.
Complexity of Automaton Identification From Given Data.
Information and Control 37:302-320, 1978.

- [Green et al. 74] Green C. C. et al.
Progress Reprt on Program Understanding Systems.
 Technical Report Stan-CS-74-444, Computer Science Department, Stanford University,
 1974.
- [Green 69] C. Cordell Green.
 Theorem Proving by Resolution as a Basis for Question Answering.
 In B. Meltzer and D. Mitchie, editor, *Machine Intelligence 4*, pages 183-205. Edinburgh
 University Press, Edinburgh, 1969.
- [Harding 76] Sandra G. Harding (ed.).
Can Theories be Refuted? Essays on the Duhem-Quine Thesis.
 D. Reidel Publishing Company, 1976.
- [Klette and Weihagen 80] R. Klette and R. Wiehagen.
 Research in the Theory of Inductive Inference by GDR Mathematicians - A Survey.
Information Sciences 22:149-169, 1980.
- [Kowalski 79a] Robert A. Kowalski.
Logic for Problem Solving.
 Elsevier North Holland Inc., 1979.
- [Kowalski 79b] Robert A. Kowalski.
 Algorithm = Logic + Control
CACM 22, July, 1979.
- [Kuhn 70] Thomas S. Kuhn.
The Structure of Scientific Revolutions.
 The University of Chicago, 1970.
- [Machtey & Young 78] Michaelachtey and Paul Young.
An Introduction to the General Theory of Algorithms.
 Elsevier North Holland, 1978.
- [McCarthy 63] John McCarthy.
 Programs With Common Sense.
 In Marvin Minsky, editor, *Semantic Information Processing*, chapter 7. The MIT Press,
 Cambridge, Mass., 1963.
- [Mitchell 78] Tom Michael Mitchell.
Version Spaces: An Approach to Concept Learning.
 Technical Report STAN-CS-78-711. Stanford Artificial Intelligence Laboratory, December,
 1978.
- [Paterson & Wegman 76] Paterson M. S. & Wegman M. N.
 Linear unification.
 In *8'th Annual ACM symposium on Theory of Computing*, pages 181-186. ACM, 1976.
- [Pereira et al. 78] L. Pereira, F. Pereira and D. Warren.
User's Guide to DECsystem-10 PROLOG.
 Technical Report 03/13/5570, Labortorio Nacional De Engenharia Civil, Lisbon,
 September, 1978.
 Provisional version.

- [Plotkin 70] G. D. Plotkin.
A Note on Inductive Generalization.
In B. Meltzer and D. Mitchie, editor, *Machine Intelligence 5*, pages 153-165. Edinburgh University Press, Edinburgh, 1970.
- [Plotkin 71a] Gordon D. Plotkin.
Automatic Methods of Inductive Inference.
PhD Thesis, Edinburgh University, August, 1971.
- [Plotkin 71b] G. D. Plotkin.
A Further Note on Inductive Generalization.
In B. Meltzer and D. Mitchie, editor, *Machine Intelligence 6*, pages 101-124. Edinburgh University Press, Edinburgh, 1971.
- [Popper 59] Karl R. Popper.
The Logic of Scientific Discovery.
Basic Books, New York, 1959.
- [Popper 68] Karl R. Popper.
Conjectures and refutations: The Growth of Scientific Knowledge.
Harper Torch Books, New York, 1968.
- [Reynolds 70] J. C. Reynolds.
Transformational Systems and the Algebraic Structure of Atomic Formulas.
In B. Meltzer and D. Mitchie, editor, *Machine Intelligence 5*, pages 135-153. Edinburgh University Press, Edinburgh, 1970.
- [Robinson 65] J. A. Robinson.
A Machine Oriented Logic Based on the Resolution Principle.
JACM 12, January, 1965.
- [Summers 76] Philip Dale Summers.
Program Construction from Examples.
PhD Thesis, Yale University, 1976.
Computer Science Dept. research report No. 51.
- [Summers 77] Phillip D. Summers.
A Methodology for LISP Program Construction from Examples.
JACM 24:161-175, January, 1977.
- [Van Emden & Kowalski 76] M. H. Van Emden and R. A. Kowalski.
The Semantics of Predicate Logic as a Programming Language.
JACM 23:733-742, October, 1976.
- [Winston 75] P. H. Winston.
Learning Structural Descriptions from Examples.
In P. H. Winston, editor, *The Psychology of Computer Vision*, . McGraw-Hill, New York, 1975.