

TM-0280

読み返しによる学習

岡 夏樹

February, 1987

©1987, ICOT

ICOT

Mita Kokusai Bldg. 21F
4-28 Mita 1-Chome
Minato-ku Tokyo 108 Japan

(03) 456-3191~5
Telex ICOT J32964

Institute for New Generation Computer Technology

読 み 返 し に よ る 学 習

岡 夏樹

新世代コンピュータ技術開発機構
oka@icot.junet

[概 要]

我々が日常よく行なっている読み返しを模擬する学習方式として、「読み返しによる学習」(Learning by Rereading)を提案する。まず、一般的に読み返しによる学習の枠組みを述べ、つぎに例からの学習に焦点を絞って検討する。incremental な例からの学習において問題となりやすい次の点を、読み返しにより改善することを目標とする。

- ・ 探索空間が大きいときの学習効率
- ・ 例の呈示順の影響
- ・ 論理和の形を含む概念の学習

読み返しによる学習の実現のための基本的な機能として、「概要の把握→詳細化」の枠組みを提案する。学習すべき概念が 属性 = 属性値 を基本命題とする命題論理式で表現され、十分な数の例が与えられる場合の概要の把握法として、例の統計的な性質から推測する方法を示し、これに続く詳細化法についても述べる。概要の把握の段階で、いつも正確ではなくても、正しいことが多いと言えることは重要である。また、この統計的推測の確からしさの評価を同時に行なうことにより、例の数が少ない場合にも対応できるので、incremental に与えられる例に対し、確からしくなったところから順次判断していくアルゴリズムを実現できる。

最後に読み返しによる学習の認知モデルとしての意味について補足する。

1. はじめに

我々は本を読むとき、次のようなことをよく経験する。ていねいに一度通読するより、手早く何度か読み返す、あるいは、分からないところにはこだわらずにとりあえず読み飛ばしておくほうが理解しやすいし、結局効率的であるように思われる。このような読み返しを利用した学習を「読み返しによる学習」(Learning by Rereading)と呼ぶことにする。本研究は人の学習過程の理解のためにそれを模擬する事を目標とするが、人間にとって読み返すことが効果的であれば、計算機にとってもそうである可能性がある。すなわち、読み返しにより、学習の効率を上げ、学習できる対象の範囲を広げられる可能性がある。なお、本論文において「読む」(read)とは、かならず

しも本の read だけでなく一般的にデータの read を意味する。

読み返しが有効である原因は次のように大きくは二つ考えられる。

① テキストの性質による（情報の提示順が悪い、あるいは、相互に関連している）

ある部分の理解に必要な情報が後になって初めて出てくることがしばしばある。これは著者がいいかげんで書く順序になんの配慮もしなかった場合はもちろん、そうでなくても、著者は書こうとする内容全体についてすでによく知っているから、読者にも分かるつもりで書いてしまったが、その理解のためには実は後に書かれる知識が必要だったということが起こりがちである。また、十分注意したとしても、本に書かれる知識は前後で複雑に関係していたり、互いに補い合う関係であったりすることが多く、その場合理解に必要な情報が後から出てくるとい状況は避けられない。このような時は、分かりにくいところにはこだわらず読み進み、適当に読み返すことが有効である。

② 人間の情報処理の特性による

短期記憶の容量の制限により、保っていることのできる情報量には限りがあるし、長期記憶にすべて整理されて格納されるわけではないので、外部記憶に頼ること、すなわち読み返しが必要になる。また、人はアルゴリズムに突き詰めて処理するのは不得意だが、全体の様子をつかむのは得意らしいということもある。

さて、本からの学習の実現のためには自然言語理解研究の成果を待たねばならないが、本論文では、自然言語の問題を含まない、例からの学習に焦点を絞って検討する。一般的には、例からと直接明示的に教えられることからとの複合による学習が考えられ、この場合には、上述の読み返しが有効になる原因に相当する状況がいろいろ起こりうるのであるが、以下では最も簡単な場合である例だけからの学習を扱う。例からの帰納では前後の例が相互に関係を持っていることは本質的であり、この意味で、ある部分の理解に必要な情報が後になって出てくるとい状況は、避けられない。さらに例の提示順が悪いときは、いっそう読み返しが有効であると考えられる。

以下、2 節では、例からの学習における従来の手法とその欠点を述べる。3 節では、読み返しによる学習に必要な基本的な機能としての「概要の把握—詳細化」について述べる。4 節では、さらに推測の「確からしさを評価」することによりどのように学習が進められるかを述べる。5 節では、読み返しによる学習が認知モデルとしてどのように位置付けされるかを示す。

2. 関連する研究

例からの学習法には例を一つずつ順に処理する方法（逐次法）と一括して処理する方法（一括法）とがある。逐次法の利点は、学習しやすいように本質的な例を先に与えることができることと、学習の途中でもそれまでに示された例によりそれなりに学習されていることとである。一方欠点は、学習結果が例の提示順の影響を大きく受けることと、論理和の形の概念の学習やノイズを含む例による学習が難しいこととであ

る。一括法はこれとほぼ反対の性質を持つ。

逐次法の代表的な例であるT.Mitchellの候補消去アルゴリズム[1]やこれと類似の他の学習アルゴリズムは、大きな規則空間に対しては遅過ぎて使い物にならないと言われている。また、G集合より一般的な正の例やS集合より特殊な負の例が呈示されたら、論理和を含む概念であると仮定して、論理和に分割して処理する方法[1][2]が提案されているが、うまく学習できるかどうかは例の呈示順に強く依存している。Far missに対応して論理和に分割する方法[3]も提案されているが、これも例の呈示順に強く依存する。

これらの方法に対して、人の学習の仕方は次の点で異なると思われる。人は大きな規則空間では、通常は候補消去アルゴリズムには従わない。例の数が少ないときや、規則空間上で離れた正の例や負の例が呈示されたときのように探索すべき空間が大きいときにはその例の処理を保留するのが普通である。この後、他の例が集まって、次の項目で述べる概要の把握により、あるいは明示的に教えられたことにより、探索空間が狭まったところで保留しておいた例を読み返す。このとき、探索空間が十分小さくなっていれば、候補消去アルゴリズムに従えばよい。このようにして、論理和の形の概念やノイズを含む例が存在しても、例の呈示順にあまり依存せず、無駄な探索や後戻りを避けて効率的に学習をしていると思われる。候補消去アルゴリズムと人の通常のやり方との違いは、人の扱える探索空間が狭いことや、普通は規則空間が完全に与えられていないことにも起因する。

一方、一括法の代表的な例としては、J.R.QuinlanのID3[4]がある。これは与えられた例を一括して見て、例の正負を判別する木（情報量の多い属性から順にテストする）を作る。この方法の欠点は、ノイズのある環境では、無意味に枝の繁った木を作ってしまうことと、例が増えたときには初めからやり直さないといけないこととである。前者の欠点を緩和するために、カイ2乗検定によりある属性と学習しようとする概念が独立であるという仮説が高い信頼性で棄却されない限り、その属性によるテストはしないという方法が報告されている[5]が、例が少ないときの分析は不十分である。後者の欠点を緩和するためには、ID3を逐次型に修正したアルゴリズムが報告されている[6]。

J.C.Schlimmer[7]は、2つの重みLS, LNを持った分布的な表現の学習において、予測に失敗したときにヒューリスティックにそれまでの記述のAND, OR, NOTをとったものを作ってみる方法を提案しているが、ヒューリスティックの検討が不十分であると思われる。

3. 概要の把握→詳細化

読み返しによる学習にとって最も基本的な機能は、初めに全体の概要をつかんでおき、続いて順次詳細化していくことである。この「概要の把握→詳細化」の枠組みにより、次のことが期待される。

- ・ 探索空間が大きい場合の効率的な学習
- ・ 例の呈示順の影響の緩和

- ・ 論理和の形の概念の学習
- ・ ノイズを含む例からの学習

3. 1. 概要の把握

一度の通読で概要を把握しておく、次に読むときそれを利用して効率良く学習できる。あるいは、初めの方を読むときに概要を把握して、続きはそれによって例を処理しながら読んでもよく、読み進めるうちにその把握内容が適当でないことが分かったら、その時点で変えればよい。

把握すべき項目には次のようなものがある。

- ・ 規則空間中で重要な属性とそうでない属性
- ・ abnormality
- ・ 論理和の分割の仕方
- ・ ノイズを含む例

また、どのくらい例が集まったところで帰納するかの判断のために

- ・ どのような例がどのくらい与えられるか

も把握すべきであるが、これについては4節で述べる。

ここでは、与えられる例と学習すべき概念の表現形式が下記の通りである場合に、多数の例の統計的性質から、上記項目を推測する方法を提案する。以下に挙げるのは説明のための簡単な例題である。考慮すべき属性の候補とそのとる値（規則空間）が与えられているとする。たとえば、

$a=\{a1,a2\},$
 $b=\{b1,b2\},$
 $c=\{c1,c2,c3\},$
 $d=\{d1,d2\}.$

ここに、 $a=\{a1,a2\}$ は属性 a は属性値 $a1$ または $a2$ をとりうることを示す。学習すべき概念は、たとえば

$a=a1 \text{ and } (b=b1 \text{ or } c=c1)$

のような 属性 = 属性値 を基本命題とする命題論理式であり、連言標準形で考えることにする。次のような例とその正負が任意の呈示順で与えられる。

$(a=a1, b=b1, c=c1, d=d1)$	正,
$(a=a1, b=b1, c=c1, d=d2)$	正,
:	
$(a=a1, b=b2, c=c2, d=d1)$	負,
:	
$(a=a2, b=b2, c=c3, d=d2)$	負,

これらが片寄りなく呈示されるとすると、十分な数の呈示の後では（簡単のため、各属性が等確率でその属性値をとる場合を示すと）、次のような統計量が得られる。

$r(\text{pos} a=a1)=0.66,$	$r(\text{pos} a \neq a1)=0.0,$
$r(\text{pos} b=b1)=0.5,$	$r(\text{pos} b \neq b1)=0.17,$
$r(\text{pos} c=c1)=0.5,$	$r(\text{pos} c \neq c1)=0.25,$
$r(\text{pos} d=d1)=0.33,$	$r(\text{pos} d \neq d1)=0.33.$

ここに、 $r(\text{pos}|a=x)$ は $a=x$ の例が正の例である割合を表わす。なお、ここでまず、常に一定の属性値しかとらない属性は無関係なものとして排除される。

さて一般に、ある属性（たとえば a ）に注目すれば連言標準形の論理式は、次のように表わせる。

$$(a=x \text{ or } G) \text{ and } H.$$

ここに、 G はリテラルの選言からなる任意の論理式、 H は連言標準形の任意の論理式である。 $a=x$, G , H が真であるかどうかは、それぞれ独立であると仮定すると、

$$P(G=\text{true}) = P(\text{pos}|a \neq x) / P(\text{pos}|a=x),$$

$$P(H=\text{true}) = P(\text{pos}|a=x)$$

が成立することは容易に証明できる。ここに、 $P(G=\text{true})$, $P(H=\text{true})$ はそれぞれ G , H が真である例が呈示される確率、 $P(\text{pos}|a=x)$ は $a=x$ の例が正の例である確率を表わす。

十分な数の例が呈示されると、

$$P(\text{pos}|a=x) \sim r(\text{pos}|a=x)$$

等が成立するので、以上のことから、次のように推測できる。

- ・ $(a=a1 \text{ or } G) \text{ and } H$ とすると、

$$P(G=\text{true}) \sim r(\text{pos}|a \neq x) / r(\text{pos}|a=x) = 0.0,$$

$$P(H=\text{true}) \sim r(\text{pos}|a=x) = 0.66,$$

したがって、 $a=a1 \text{ and } H$ の形であろう。

- ・ $(b=b1 \text{ or } I) \text{ and } J$ とすると、同様に、

$$P(I=\text{true}) \sim 0.33,$$

$$P(J=\text{true}) \sim 0.5.$$

- ・ $(c=c1 \text{ or } K) \text{ and } L$ とすると、同様に、

$$P(K=\text{true}) \sim 0.5,$$

$$P(L=\text{true}) \sim 0.5.$$

- ・ $(d=d1 \text{ or } M) \text{ and } N$ とすると、同様に、

$$P(M=\text{true}) \sim 1.0,$$

$$P(N=\text{true}) \sim 0.33.$$

したがって、 d は無関係であろう。

なお、ここには示さなかったが、たとえば $(a=x \text{ or } G) \text{ and } H$ に対して $P(G=\text{true}) \gg 1$ となれば、 $\text{not}(a=x)$ の形を含むと推測できる。

3. 2. 補足

3. 1 節の議論の中で、幾つかの仮定や単純化をしていたが、ここではそれらを 3. 1 節での順序にしたがってひとつひとつ検討する。

- ① 学習すべき概念の表現形式を 属性 = 属性値 を基本命題とする命題論理式に限定

する

現状のエキスパートシステム中に表現されている知識は大部分命題論理の範囲内に収まっていると思われる。従って、この範囲での学習を考えることは現実的に意味がある。ただし当然、将来は属性間の関係やリスト表現を扱える方向に拡張すべきである。この際、概要の把握法を新たに考える必要があるように思われる。

② 学習すべき概念は説明のため簡単なものとする

これは何も問題はない。かえって複雑な場合のほうが読み返しの効果がいっそう大きくなる。ここで述べた概要の把握法は、

○ (候補となる属性の数 × 例の数)

の手間しかかからない。

③ 考慮すべき属性の候補が与えられている

人間が普段行なっている学習と、計算機による学習についてのこれまでの大部分の研究との大きな違いの一つは、この点にある。ただし、5節で述べるように、統計情報による概要の把握を、例をたくさん眺めていると関係する属性が見えてくるという発想のモデルとみなすことができる。

④ 学習すべき概念を連言標準形で表わす

上述の概要の把握およびその後の詳細化はすべて、選言標準形でも同様に議論できる。これら標準形以外の形での学習については4節で述べる。

⑤ 例は片寄りなく呈示される

例がある確率分布に従って呈示されることを前提として推測しているのであるから、これは本質的な仮定である。逆に、このような推測をする立場からの例の良い呈示とは、その本来の確率分布に従った呈示であるということになる。

⑥ 十分な数の例が呈示される

この仮定が十分成立していないときには、推測が不正確になるが、重要な属性とそうでない属性、論理和の分割の仕方を大まかにつかむことはでき、これを利用して効率的に学習できる。つまり、概要の把握の段階では最終的には必要だがそれほど重要でない属性までとらえる必要はないということである。

なお、4節では、推測の確からしさを合わせて評価することにより、例が少ないときでも自然な無駄のない対応ができることを示す。

⑦ 各属性は等確率でその属性値をとる

これは説明を簡単にするための仮定であり、まったく問題はない。

⑧ $a=x$, G , H が真であるかどうかは、それぞれ独立である

この仮定により、一般の命題論理にさらに制限を加えたことになる。しかし、たとえば、

$(a=a1 \text{ or } a=a2 \text{ or } G') \text{ and } H$ や、

$(a=a1 \text{ or } G') \text{ and } (a=a1 \text{ or } G'') \text{ and } H'$

などの場合は、近似的には概要を把握できる。問題となるのは、

$(a=a1 \text{ or } G') \text{ and } (a \neq a1 \text{ or } G'') \text{ and } H'$ (G, G' は空でない)

などの場合である。

ただし、以上のような統計情報に基づく判断としては、これは当然の結果である。人間も例をざっと見て概要をつかむときには、この程度の情報しか使わず、同じような誤りを犯すことはよくあると思われる。

3. 3. 詳細化

概要の把握結果をどのように使って学習を進めるかについては、幾つかの方法が考えられる。

① 他の学習アルゴリズムの前処理（枝刈り）としての概要の把握

上記の例で言うと、概要の把握の結果、無関係でありそうな属性 d を除いて、属性 a, b, c に絞って何らかの学習アルゴリズム（従来から知られているものでよい）を効率良く走らせることができる。上記の例は簡単のため候補となる属性が少ない場合を示したが、多数の無関係なものを含む何千、何万という属性に対してこのような枝刈りをするには大いに意味がある。また、その際、関係の強そうなものに絞って（関係があるかもしれないという程度のものは洩れてもよい）、これを行なえば目標とする概念の概要（例外的なところまでは分からないが、大体正しい）を得ることができる。

② 基本命題の組み合わせを試みる

上記の例で言うと、単独で概念の必要条件であるらしい基本命題($a=a1$)や、無関係でありそうな基本命題($d=d1$)を除いた基本命題の組み合わせについて、統計情報を調べてみる。これが望ましい値であればその命題も候補に含めてこれを繰り返す。上記の例の場合には($b=b1 \text{ or } c=c1$)について調べることになり、それが必要条件になっていることが分かる。

一般的には多くの組み合わせの候補が存在することになるが、ヒューリスティックとして、必要条件（連言標準形のとき）に近いものから試みていくことが有効である。これは、例外はあるがほぼ必要条件であることが分かっているものに対し、その例外的な条件を詳細にしていくという人のやり方に相当する。

もう一つのヒューリスティックは、それぞれの候補の統計値を照らし合わせてもっとも望ましいものから試みていくことである。上記の例では、 $P(J=\text{true})$ の値と

$P(L=true)$ の値に近いことは $(b=b1 \text{ or } c=c1)$ のもっともらしさを増している。また、 $b=b1$ である割合や $c=c1$ である割合も求めておけば、これらの値と $P(L=true)$ や $P(K=true)$ の値の比較も同様に $(b=b1 \text{ or } c=c1)$ のもっともらしさの尺度となる。ただし、例が十分与えられていないと十分な精度は得られない。

属性間の階層構造が与えられている場合には、これを利用してもっともらしい組み合わせを求めることができる。逆に、or による組み合わせを作っていくことは、階層構造を学習していることになる。

4. 確からしさの評価

以上では、例が十分多く呈示されたという条件のもとで議論してきたが、ここではそうでない場合も扱う。少ない例だけから例全体の集合の様子を決め付けてしまうような学習の仕方は、無理があるように思える。したがって例えば P.H.Winston のアーチの概念の学習 [8] のように near miss というまい呈示順に頼ったりすることになる。そこで、ある与えられた例から何かを推測するとき、統計的にどの程度の確からしさで言えるかを合わせて評価することにする。これは人間がどのくらい例が集まったら確からしい帰納ができるかを判断していることのシミュレートである。

概要の把握について言えば、ある命題が必要条件や十分条件であるらしいこと、無関係であるらしいことが、どの程度の確からしさで言えるかを評価する。このことにより、ある少数の例だけから無理に仮説を立てて袋小路にはいるような不自然さ、効率の悪さを避けることができる。このようにして、以下に述べるように、incremental に与えられる例に対して、システムが能動的にタイミングをとりながら、確からしくなったところから順次判断していくアルゴリズム（言い替えば、incremental に与えられる例を適宜一括処理していくアルゴリズム）を実現できる。また、このように確からしさを評価しておく、統計的に見てほかにどのような例があるとより確かな判断が下せるかが分かるので、システム側からの例の呈示（質問）が可能になると思われる。

一つ例を示す。属性 a, b, c, d, e はそれぞれ属性値 1 または 0 を等確率でとるものとする。学習すべき概念は、

$$a=1 \text{ and } (b=1 \text{ or } c=1)$$

とする。次のように例がその確率分布に従って、incremental に呈示される。

	a	b	c	d	e	
1st:	0	0	0	1	0	負,
2nd:	1	0	1	1	1	正,
3rd:	0	1	0	1	0	負,
4th:	0	1	0	1	0	負,
5th:	0	1	0	0	0	負,
6th:	1	1	1	1	0	正,
7th:	1	0	1	1	0	正,
8th:	1	1	0	0	0	正,
9th:	0	1	0	0	1	負,
10th:	1	0	0	1	0	負,

11th:	1 0 1 1 1	正,
12th:	0 1 1 0 1	負,
13th:	0 0 1 0 1	負,
14th:	0 0 1 1 0	負,
15th:	1 1 1 1 1	正,
16th:	1 0 0 0 0	負,

:

例えばこの時点で、

a=1の例は8個、そのうち正の例が6個、
a=0の例は8個、そのうち正の例が0個、
b=1の例は8個、そのうち正の例が3個、
b=0の例は8個、そのうち正の例が3個、
c=1の例は8個、そのうち正の例が5個、
c=0の例は8個、そのうち正の例が1個、
d=1の例は10個、そのうち正の例が5個、
d=0の例は6個、そのうち正の例が1個、
e=1の例は6個、そのうち正の例が3個、
e=0の例は10個、そのうち正の例が3個、

であり、確からしいのは、a=1 が必要条件であるらしいということだけである。

ここで、a=1 and H の形であると仮説を立て、a=1 である例についての統計情報をとりなおす。

b=1の例は3個、そのうち正の例が3個、
b=0の例は5個、そのうち正の例が3個、
c=1の例は5個、そのうち正の例が5個、
c=0の例は3個、そのうち正の例が1個、
d=1の例は6個、そのうち正の例が5個、
d=0の例は2個、そのうち正の例が1個、
e=1の例は3個、そのうち正の例が3個、
e=0の例は5個、そのうち正の例が3個、

b=1, c=1, e=1 が十分条件である可能性があることが分かるので、それらの or による組み合わせを試みることもできるが、それほど確かではないので、より多くの例が集まるまで判断を保留することもできる。

一般的に、必要条件または十分条件になるものからこのように確からしくなっていくので、標準形で推測することにこだわらず、確からしくなった条件から括りだしていく方法が有効であると思われる。こうして得られる論理式は、簡潔で人間にとって自然なものになっていると予想される。

5. 認知モデルとしての読み返しによる学習

5. 1. 関係する属性の発想

人間の情報処理を、

意識して行なう逐次処理

+

無意識に行われる大規模並列処理

ととらえる。逐次処理からは、長期記憶のうち活性化されているものと短期記憶をアクセスできる。長期記憶は並列処理により、活性化／不活性化され、また、意識的処理や、感覚器からの入力により活性化され、ある時定数をもって不活性化されるとする。

例からの学習とこの枠組みとの対応は次の通りである。人間は例からの学習を意識的に行なうこともできるし、無意識に行なうこともある。従来研究されてきた学習アルゴリズムの大部分は逐次処理に対応する。例えば、候補消去アルゴリズム[1]、ID3[4]、説明に基づく学習等である。一方、connectionism の学習研究は大部分が並列処理に対応する。

さて、統計情報からの概要の把握は、「例を眺めていると、共通のあるいは異なる属性が思い浮かぶ」ことに相当し、大規模並列処理の一つの機能のモデルになっているとみなすことができる。したがって、読み返しによる学習の認知モデルへの一つの可能な位置付けは、まず並列処理により関連する属性が活性化され、それを使って逐次アルゴリズムにより学習することに対応するということである。

なお、一般的に、このような統計的な処理は、大規模並列処理の特徴のひとつをとらえたモデルになっていると言える。

5. 2. 例外のある規則に対する overgeneralization

人間が帰納する場合は overgeneralization がしばしば起こる。例えば、母国語の獲得過程において、overgeneralization が観測されることはよく知られている。英語では、初め幾つかのよく使われる不規則動詞の過去形の正しい使用が観測され（例：went, broke）＜段階Ⅰ＞、次の時期には、それらの正しく使われた不規則動詞も含めて、規則動詞であるかのように過去形を作ってしまう overgeneralization が観測され（例：jumped, goed, brokeed）＜段階Ⅱ＞、やがて規則動詞、不規則動詞両方の正しい過去形が使われるようになる（例：jumped, went）＜段階Ⅲ＞[9]。

読み返しによる学習の枠組みでは、＜段階Ⅰ＞は、例がまだ少なく統計的にあまり確かでないので、規則動詞の過去形を作る規則（語幹+ed）がまだ帰納されない状態であると説明できる。＜段階Ⅱ＞は、一応その規則が帰納できる程度に例の呈示が増えたが、例外（不規則動詞）までは処理できない、つまり概要を把握した状態である。＜段階Ⅲ＞は、さらに例外まで把握できるだけの、多数の例が呈示された状態に対応する。

この overgeneralization は、「いつも正確ではないが、これまでに得た情報から正しいことが多いと思われることは言う」という従来の帰納システムに欠けていた機能であると言える。

6. おわりに

読み返しによる学習 (Learning by Rereading) の一般的な枠組みを述べ、さらに例からの学習に焦点を絞って検討した。読み返しによる学習の実現のための基本的な機能として、「概要の把握→詳細化」の枠組みを提案し、例の統計的な性質から概要を把握する方法を示した。また、同時にこの統計的推測の確からしさの評価を行なうことを提案した。

今後の主な課題は次の通りである。

- ・ 学習できる概念の表現形式の拡張
- ・ 全体のアルゴリズムの効率の見直し

【参考文献】

- (1) Mitchell, T.M. et al., Learning by Experimentation: acquiring and modifying problem-solving heuristics, in: R.S. Michalski et al (Eds.), Machine Learning (Tioga,1983) 163-190.
- (2) Bundy, A. et al., An Analytical Comparison of Some Rule-Learning Programs, Artificial Intelligence 27 (1985) 137-181.
- (3) Langley, P., Language acquisition through error recovery, CIP Working Paper 432, Carnegie-Mellon University, June 1981.
- (4) Quinlan, J.R., Learning Efficient Classification Procedures and Their Application to Chess End Games, in: R.S. Michalski et al (Eds.), Machine Learning (Tioga,1983) 463-482.
- (5) Quinlan, J.R., The Effect of Noise on Concept Learning, in: R.S. Michalski et al (Eds.), Machine Learning volume 11 (Morgan Kaufmann,1986) 149-166.
- (6) Schlimmer, J.C. et al., A Case Study of Incremental Concept Induction, AAAI-86, 496-501.
- (7) Schlimmer, J.C. et al., Beyond Incremental Processing: Tracking Concept Drift, AAAI-86, 502-507.
- (8) Winston, P.H., Learning Structural Description from Examples, in: P.H.Winston (Ed.), The Psychology of Computer Vision (Mcgraw-Hill,1975) 157-209.
- (9) Clark, H.H. et al., Psychology and Language (Harcourt Brace Jovanovich, 1977) 333-373.
- (10) 岡 夏樹, 読み返しによる学習, 情報処理学会第33回全国大会 (1986) 1247-1248.