

1. 知的情報の表現と利用

—自然言語処理を中心として—

田中 穂積

東京工業大学工学部教授

はじめに

人間と機械とのコミュニケーションの形態としてもっとも素直なものは、自然言語を用いることです。それが人-機械インターフェースの究極の姿であると考えられます。そのためには、機械の側で自然言語の意味を理解する必要があり、機械による意味理解を想定した意味情報の表現形式を研究する必要があります。

私どもは言葉(自然言語)を使って他人とコミュニケーションすることにより、文化を築き上げてきました。自然言語の本質に関しましては未解明の問題も多いのですが、最近、対話理解システムの構築、機械翻訳、音声理解システム、オフィスにおいては文書の作成・校正、管理・検索、要約の問題など、コンピュータで自然言語を処理する技術(自然言語処理技術)の研究が重要になってきました。

この自然言語処理技術につきましてはさまざまな問題が山積といった状況にあります。以下では、意味理解を行う場合の知識の表現形式と、それらを利用した意味理解の方式につきまして、私どもがこれまで行ってきた研究の一端を述べさせていただきたいと思えます。具体的には、コンピュータによる意味理解がもっとも困難であるとされている比喩理解のための意味情報の表現形式とその利用に関する研究と、文を読み進むにつれて意

味的な曖昧性を漸進的に解消する、新しい意味情報の表現形式の提案に焦点を当てたいと思います。具体的なお話をする前に、自然言語処理がどのようにして行われているかを概観してから、私どもの研究を紹介することになります。

コンピュータによる自然言語処理

自然言語を処理する全体のシステム構成は、3種類ほどのエンジンから成り立っています(図1)。

一つは自然言語を解析するエンジンです。ここで、与えられた自然言語の文の文法的な正しさを解析し、意味・文脈を解析します。解析エンジンを動かすために、知識ベースエンジンとのコミュニケーションが必要です。文法的な知識、辞書的な知識、百科辞典的知識、生活するうえで必要な常識を蓄積した知識ベースがあり、それらを利用するために知識ベースエンジンを働かせるわけです。

解析する場面では、いろいろな推論を行いますが、解析した結果をうけて、何らかの応答を行うために推論を行うこともあります。後者の場合も、知識ベースを参照することになります。推論を行う部分を推論エンジンと呼ぶことにします。

最後に、応答すべき答えを文章として提示する、生成のエンジンがありますが、これは解析エンジンとあわせて言語処理エンジンと

呼ぶことにします。要約すると、自然言語を処理するシステムは、①言語処理エンジン、②知識ベースエンジン、③推論エンジンから構成されています。それぞれのエンジンをどう構築するかは、大きな長期の研究課題であります。

言語処理エンジン中の統語解析は、与えられた文が文法にかなっているかどうかをチェックするものですが、文法をどのように定義するかという点にもまだ問題があります。日本語の文法体系についての決着がついていないため、言語学を研究している方が大勢いるわけです。

統語解析につきましては、最近さまざまな技術が開発され、高速な解析が可能になりつつあります。私どもが研究してきました解析技術は統語解析ではなく、意味解析技術です。文の解析が進むにつれて漸進的に意味的曖昧性を解消する技術です。この技術と統語解析技術とを融合させますと、意味的な情報を早期に用いて曖昧性を解消し(探索空間を刈り込み)、組み合わせ的爆発を防ぐことができます。一般に、文が長くなると統語解析結果に多数の曖昧性が生じますが、意味的に異常な解析結果を早期に捨てて、効率のよい解析が可能になります。

自然言語処理の問題点

自然言語処理のもっとも大きな問題は、意味の問題です。意味解析や文章生成においては、意味的な曖昧性を解消することが特に重要です。「さくらを食べた」と「さくらが咲いた」とでは「さくら」の意味が異なります。「さくら」のもつ意味的曖昧性を解消する必要があります。その他に省略語(句)の補強も問題になります。元来、言葉には同じことの繰り返しを避ける性質があるため、頭のなかでそれらを補わないと文の意味が理解できません。

「旅行にいきたい人」と誰かがいって「僕も

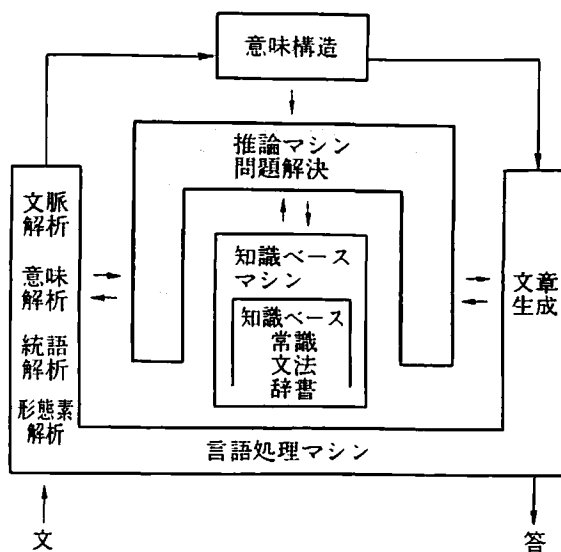


図1 自然言語処理システムの構成

いきたい」と答えたとします。この答は、「いつどこへ何のために」という情報が欠落していますから曖昧です。こうした省略されたものを補うためには、文法的・辞書的な知識だけでは不十分で、常識を使わなければならないこともあります。これはコンピュータにとって難しい問題となっています。

自然言語には曖昧性があることを前提として、自然言語処理の研究に取りかかる必要があるといえます。曖昧性があるからこそ言語はここまで発達したともいえますが、この問題を何とかして解決しなければなりません。

自然言語の曖昧性の種類にはいろいろあります。第1に、語彙レベルでの曖昧性があります。私がこの研究分野にはいるとき、ccという物理単位をコンピュータに理解させることは簡単だと思っていました。しかし、「1000ccのビーカーに10ccの水をいれる」という文の場合、1000ccは内容積であり、ビーカーの体積ではありません。「1000ccの自動車」というと、自動車の体積が1000ccではなく、排気量です。物理単位ですら、その意味理解は一見簡単そうですが、そうではないことがわかります。

文脈に依存して曖昧性を解消しなくてはな

らないこともあります。「小さな大学の門」といった場合、大学が小さいのか、門が小さいのかという、解釈の問題があります。この文がおかれた周囲の状況・文脈を正しく解析しなければなりません。例えば、機械翻訳システムでは、small が university にかかるのか、gate にかかるのかで翻訳結果が異なったものとなるので重要です。

指示詞の場合、“A taxi is coming over there.” “Let’s take it.” の it は、taxi を指すと考えられます。it は何かわからなくても「それ」と訳せばよいのではないかという考え方もありますが、take それ自体は多義語で曖昧性がありますから、it が taxi を指すことがわかって初めて「乗る」と訳すことができます。この点をきちんとしなければなりません。研究が進んでいるとはいえ、まだまだ完全ではありません。

文の生成にもさまざまな問題があります。例えば、曖昧性のない文を生成するときの語順の問題があります。

①「花子と太郎が結婚する」

②「太郎が花子と結婚する」

は似ていますが、①の場合、「花子と太郎がそれぞれ別の相手と結婚する」とも解釈することができます。②の場合、そのような解釈はありません。「花子と桃子が結婚した」を逆にして、「桃子が花子と結婚した」とすると、おかしいことになります。このような問題には常識的、社会的な規範などが関係していることがわかります。

残念ながら上記したすべての問題が扱えるレベルにはまだありません。そこで以下では、「さくら」や「1000cc」といった語彙レベルの曖昧性解消に限定した私どもの研究について、(比喩理解の研究を紹介した後で)述べたいと思います。

ここまでの話で、自然言語処理については問題山積で、いささか悲観的になった方もおられるのではないのでしょうか。そこで、自然

言語をコンピュータで処理し、完全に理解できなくても、コンピュータの側を人間が支援して、全体として役立つシステムができていることを述べたいと思います。その一つはワードプロセッサです。

平仮名を漢字に変換するには、本来大きな問題があるはずですが。例えば、「にわとりがいる」と入力したとき、「2羽鳥がいる」「鶏がいる」という二つの意味があり、それに応じて漢字を選択しなければなりません。どちらが正しいかを、コンピュータできめることはとても無理だといえます。どちらが正しいかを決定するために、私どもは文脈をみます。コンピュータが文脈をみることはたいへんな仕事で、適切な方法がみつからないのが現状だということをすでに述べました。しかし、実際使われているワードプロセッサは、かぎられた自然言語処理の技術しか利用していないにもかかわらず、必要なときに人間に判断を仰ぐことで、道具として立派に役立ち手放せません。このことは、自然言語処理の研究成果の一部は、すでに役立つシステムを産みだしており、実用に遠い研究ではないということを示唆しています。

比喩の理解

私どもが使用する言語には、比喩的な表現が必ずといってよいほど含まれています。例えば、「リンゴのような頬」「アイディアが育つ」「人工知能は石頭」「鍋を食べる」などを翻訳させる場合、いろいろな問題が生じます。比喩は、言語学的にも十分に明らかにされていない面がありますが、ある言語学者は「我々の表現のなかには比喩が満ち満ちており、その比喩によって、言葉はこれだけ豊かな能力をもつようになった」と述べています。

他人とコミュニケーションする場合、比喩的な表現は、相手に伝えたい内容をよりよく伝達する作用があることをもっと研究する必要があります。コンピュータで自然言語を解

析する場面で、今後、比喩理解の問題は避けて通れない問題となるでしょう。この問題に私たちがどう対処したかということを紹介しします。

簡単な直喩の理解からはじめることにします。たとえばものが明らかで「リンゴのような頬」といった比喩をコンピュータでどのように処理をするかという問題です。リンゴは種々の性質をもっていますが、それらのうちどれが頬に写像されるか、その計算が難しいわけです。この点に関してこれまで形式的な方法がなく、最初からリンゴのなかのある性質が写像されることをコンピュータに入力しておくという方法をとっていました。ここをより柔軟に行う方法を考えてみました。

比喩の視点表現による理解

さて、比喩は、人間の高次レベルでのコミュニケーションにおいて重要な役割を果たしていますが、比喩自体の言語学的な検討が十分でないことがあげられます。コンピュータで直喩を理解することも、さまざまな問題を抱え容易ではありません。そこで私どもは、直喩の理解にあたって、どのような情報構造を設定したらよいかを、まず考えてみました。

従来の直喩についての考え方は、リテラルな(文字通り)読みを中心におくものです。「彼女の頬はリンゴのようだ」の「ようだ」ととると、頬がリンゴであるはずがないので、これはリテラルな意味でおかしいというところから処理をはじめのが従来の方法です。それとは異なった観点から、以下のような方法でコンピュータに直喩を理解させることを考えてみました。

「少女の頬はリンゴのようだ」といった場合、リンゴの視点から頬をながめる、つまり言語表現に含まれる比喩的な表現の理解を、たとえば概念の特徴的な性質が、たとえられる概念に移行する過程としてとらえるようにしました。この過程を視点表現による理解と呼ん

でいます。この視点表現は、「概念Tを概念Sという視点でみる」ことを表現する形式で「 $*(T)/*(S)$ 」と表します。そして、 $*(T)$ をターゲット概念、 $*(S)$ をソース概念と呼びます。例えば、「リンゴのような頬」の場合、リンゴがソース概念となり、リンゴの性質が写像される頬の側がターゲット概念となり、視点表現「 $*(頬)/*(リンゴ)$ 」が対応することになります。

このような視点表現を用いますと、比喩表現ではありませんが「望遠鏡をもった少女をみた」という文章の場合、具体物の視点から望遠鏡をみているということで、統一的に扱えます。つまり、比喩的表現も、そうでない表現も、ある意味では同じ形式で表現することが可能になり、視点表現による理解を介すことにより、比喩理解のために特別な仕掛けを考えなくてもよい、と私どもは考えました。

言い換えますと、視点表現による理解とは「ソース概念の性質を、可能な場合にかぎりターゲット概念に移す」ことによって、視点表現で表された概念のもつ性質を計算することであるということができます。私どもはこのようなシステムを構築して動かしています。それを比喩(実際には直喩ですが)理解実験システム(A Metaphor Understand System for computer Experiment: AMUSE)と呼んでいます。リンゴのもつ性質のなかから特徴的なものを選び、それと頬のもついろいろな性質とを対応させて、頬のもつ性質を変化させ、変化した結果が「リンゴのような頬」の意味になるという戦略を考えたわけです。そこで次に、頬のもつ性質を変化させる方法について紹介します。

顕現性の計算とエントロピー

視点表現による比喩理解において、

- ①ソース概念のどの性質がターゲット概念に移されやすいか
- ②性質が移された結果、ターゲット概念は

どのように変化するかが問題となります。この問題を私どもは、顕現性という指標によって扱うことにしています。

顕現性とは、ある概念の性質の典型度を表す指標で、ソース概念において、顕現性が高い性質ほどターゲット概念に移されやすく、移された性質はターゲット概念において顕現性が高くなります。例えば、「*(頬)/*(リンゴ)」では、顕現性が高い*(リンゴ)の性質「赤い」「丸い」などが*(頬)に移された結果、*(頬)の性質「赤い」「丸い」の顕現性が高くなります。

類推や比喻理解に、概念の性質の顕現性を用いた研究はいくつかありますが、顕現性の計算方法について明確な指針が与えられていませんでした。私どもは、概念を、属性名と属性値集合の対の集合として表現し、属性値集合のエントロピーとして顕現性を定量的に計算する方法を考えました。

例えば、リンゴの概念の表現形式をまず検討してみます。リンゴにはいろいろな属性がありますが、その属性値に確率値が付加されると考えるのは自然です。このような確率付きの属性の集合として、リンゴの概念の表現形式が図2のようであったとします。これまでは、確率が最高値の属性が頬に写像されるとしていましたが、より精密に顕現性を計算するアルゴリズムを開発したいわけです。

ここでちなみに、色の性質で「赤」が選ばれた場合と、値段で「高い」が選ばれた場合

概念の例*(リンゴ)

$$*(リンゴ) = \left\{ \begin{array}{l} \text{色:} \left\{ \begin{array}{l} \text{赤} \#0.8 \\ \text{緑} \#0.15 \\ \text{黄} \#0.05 \end{array} \right\} \\ \text{値段:} \left\{ \begin{array}{l} \text{高} \#0.4 \\ \text{中} \#0.3 \\ \text{低} \#0.3 \end{array} \right\} \\ \text{種:} \left\{ \begin{array}{l} \text{有} \#0.95 \\ \text{無} \#0.05 \end{array} \right\} \end{array} \right\}$$

図2 リンゴの概念表現

とを比較します。「赤」に比べて「緑」や「黄色」は非常に確率が低いことに注意してください。色が「赤」という選択が、リンゴに特徴的な性質になっているかという点においては、値段で「高い」を選んだ場合より、前者がシャープな分布をしているので、色が「赤い」という性質は、値段が「高い」と比較してより特徴的な性質であるといえます。

このような違いは、情報理論で用いるエントロピーの違いとして表せます。リンゴの色や値段についての確率分布から、エントロピーを計算するわけです。その結果、確かにリンゴの色は値段よりエントロピーが大きいことは明らかです。これを計算しておけば、リンゴのどの性質を頬に写像するかがきめられます(図3)。

顕現性と差異度

基本的な考え方は、以上に述べた通りなのですが、その後、エントロピーを計算し、利用するだけでは不都合ではないかという問題がでてきました。例えば、リンゴには種がありますが、種があるという性質はほとんど正しいことになり、分布の観点からは非常にシャープな分布をしているといえます。ところが、頬には種はありません。また、リンゴと似たものはほかにもあるのに、なぜ直喩の対象としてリンゴが選ばれたのかも考慮する必要があります。それにはリンゴと似た兄弟概念の

リンゴの色: 赤#0.8, 緑#0.15, 黄#0.05

$$1 - \frac{0.8 \times \log_2 \frac{1}{0.8} + 0.15 \times \log_2 \frac{1}{0.15} + 0.05 \times \log_2 \frac{1}{0.05}}{\log_2 3} = 0.4421$$

リンゴの値段: 高#0.4, 中#0.3, 低#0.3

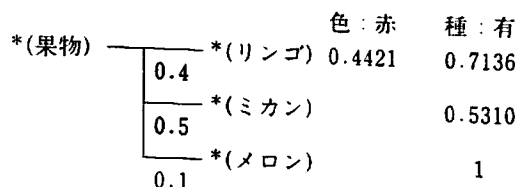
$$1 - \frac{0.4 \times \log_2 \frac{1}{0.4} + 0.3 \times \log_2 \frac{1}{0.3} + 0.3 \times \log_2 \frac{1}{0.3}}{\log_2 3} = 0.0089$$

図3 情報量の計算

存在を考慮することが考えられます。兄弟概念を考慮した指標として、私どもは差異度を導入しています。今の場合、リンゴと兄弟の概念、すなわちいろいろなフルーツを考えることになります。いろいろなフルーツがあるのに、なぜリンゴが選ばれたかを考えるわけです。

フルーツを考えますと、ほとんどのものが種をもっていますから、「種がある」という性質のエントロピーは大きいものの、リンゴに特有の性質とは考えられないことがわかります。次に、他のフルーツ、例えば「メロンのような類」ではなく、なぜ「リンゴのような類」という表現が選ばれたかを考慮する必要があります。そのために差異度を使います。差異度の計算例を図4に示します。差異度とエントロピーを掛け算した形で顕現性、つまりリンゴを特徴づける性質を決定します。このような方法で、直喩理解モデルをつくりました。

リンゴのいろいろな属性について差異度を計算して顕現性を計算した結果を図5に示します。そして、顕現性と、類とリンゴが共通にもつ性質を使って、類の性質を変化させます。リンゴにかぎって言えば、種があるのはほとんど確かなことで、エントロピーは大き



$$d(*(\text{リンゴ}), \text{色：赤}) = \frac{0.4 \times 0.4421}{0.4 \times 0.4421} = 1$$

$$d(*(\text{リンゴ}), \text{種：有}) = \frac{0.4 \times 0.7136}{0.4 \times 0.7136 + 0.5 \times 0.5310 + 0.1 \times 1} = 0.4385$$

図4 差異度の計算

い値になっていますが、類には種はありませんから考慮する必要はありません。ターゲット概念の性質を変化させる自然な方法は、赤がもっとも顕現値が高いので、類の色のなかの赤の性質を強調し、ほかの性質を抑制することです(図6)。こうして、計算可能な直喩理解のモデルがえられます。

リンゴの性質に確率が付与された情報構造として(例えば図2)、自分の計算モデルに都合のよい値を与えたのでは具合が悪いので、心理学者に実験方法を聞き、実験心理学的にその値を決定し、この比喩は比喩らしい比喩かどうかということも研究しています。現在、

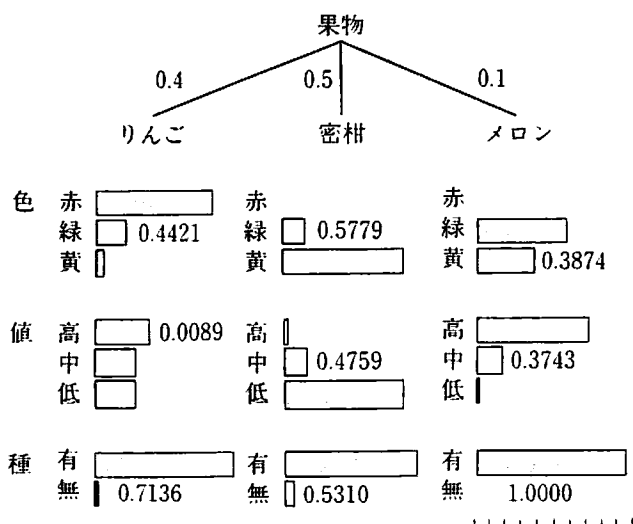


図5 顕現性の計算

	情報量×差異度＝顕現性
りんごの色	0.4421×1.0000=0.4421
りんごの値	0.0089×0.0453=0.0004
りんごの種	0.7136×0.4385=0.3129

被験者にも「これは比喻らしい比喻か」を尋ね、計算結果とよく一致するかどうかを実験的に検討しています。今のところ、人間の行っている判断と近い判断結果が計算できることがわかっています。しかし、ここでいう比喻は、もっとも簡単な直喩理解のモデルで、本当の難しい問題は、まだたくさん残っています。

曖昧性の漸進的解消

人間の言葉にはもともと曖昧性があり、むしろ曖昧性があることが言葉(自然言語)の特質であるといわれています。曖昧性の解消が、自然言語処理の重要な問題であることもすでに述べました。

語彙レベルの曖昧性の解消について考えてみましょう。例えば「東工大にはいった」といった場合、「はいった」が合格したという意味なのか、ただ東工大のキャンパスにはいったという意味なのかで、英語に翻訳すると異なったものとなります。この場合には、文脈が関連し、語彙レベルの曖昧性であっても、文脈的なもの、あるいは常識的なものが関係するので、難しくなります。

しかし、文脈が関係しない語彙レベルの曖昧性解消は、いくらか扱いやすくなります。

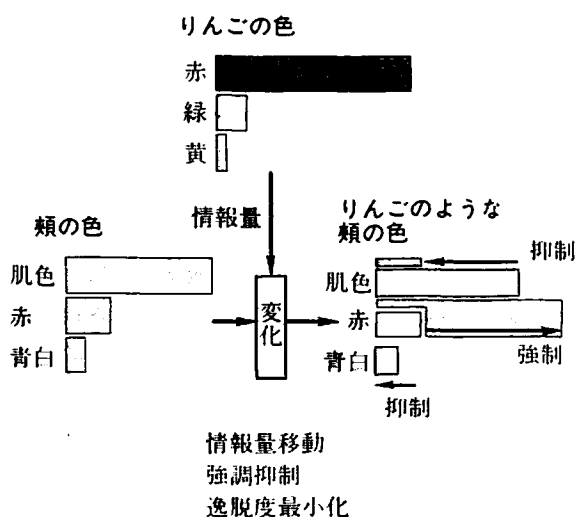


図6 理解の結果

とはいえ、まだ決定打といえるほどのものではありません。私どもは、漸進的に語義的な曖昧性を解消するモデルをつくりたいと考えました。ここで漸進的とは、文を読み進むにつれて徐々に曖昧性解消が進み、意味理解ができていくということです。文を終わりまですべて読み切ってから処理をはじめのではなく、単語を先頭から読み進むと同時に、そこからえた情報を即座に使い、意味解析を行い、インクリメンタルに曖昧性を解消するモデルです。

一般化弁別ネットワークを用いた曖昧性の漸進的解消

私どもは、一般化弁別ネットワーク(Generalized Discrimination Network: GDN)を提案しました。従来の方式では、曖昧性を表現する際に、一つの単語のもつ複数の語義を、リスト形式で保持します。例えば「はしをかける」という文では、「はし」に「橋」「端」「箸」の3通りの、また「かける」に「掛ける」「欠ける」「書ける」「駆ける」の4通りの語義が少なくとも存在しますが、それらをそれぞれ(橋、端、箸)、(掛ける、欠ける、書ける、駆ける)として表現しておきます。これらの表現を用いた上文の解析(曖昧性解消)では、「はし」の4個の語義の「を」格の情報を、「は

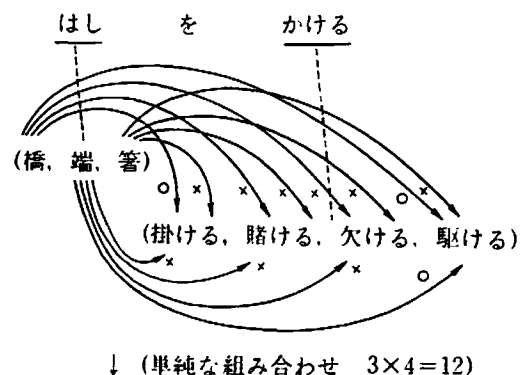


図7 従来の方式(ポラロイド語)

し」の3個の語義に対して単純な組み合わせで適用するため、 $3 \times 4 (=12)$ 通りの組み合わせをすべて検査する方法をとります(図7)。

もちろん、私どもの研究のベースにもそれはありますが、基本的な考え方は、「はし」と「かける」の意味関係によって、両方の単語の意味が決定される計算を、組み合わせ的な方法ではなく、もう少しスマートな方法で行いたいというものです。

そのために、私どもは弁別ネットワークを利用することを提案しています。弁別ネットワークは一種の決定木です。“take”という動詞を例にすると、この語のもつ各語義ごとに別々の語義記述を用意しておきます。“take”の「散歩する」という語義記述には、主語が人間で目的語がアクションであるとか、「はいる」という語義記述には、主語が人間で目的語がbathであるとか書いておくわけです。このような語義記述の羅列には、重複した記述(主語が人間)が随所に現れます。弁別ネットワークを用いるとこうした重複記述を避けることができます。

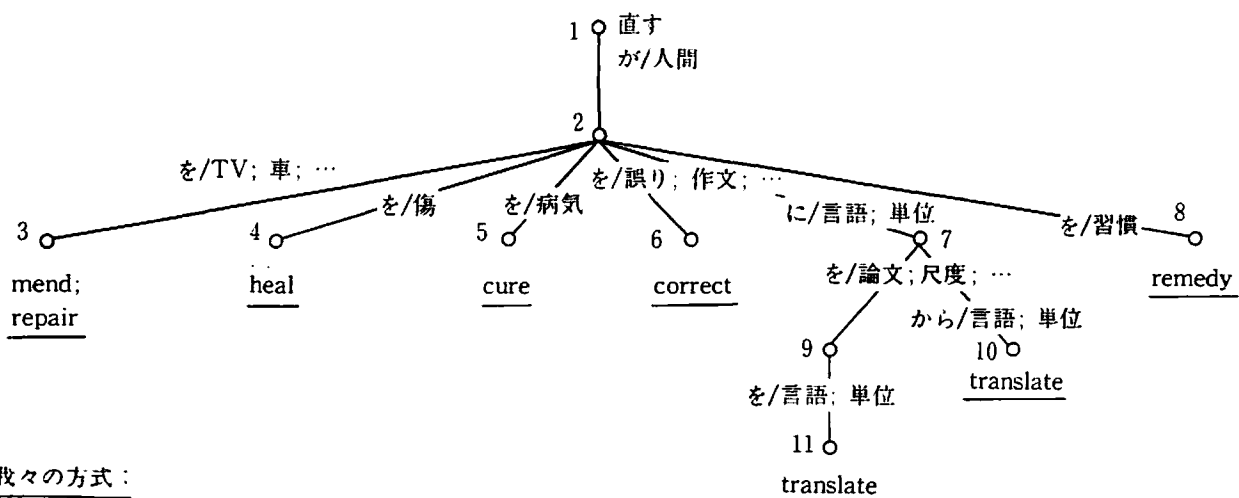
動詞「直す」を弁別ネットワーク(DN)で表現したものを図8に示します。

弁別ネットワークの特徴

弁別ネットワークの枝(リンク)にはラベルとして制約が付与されています。葉ノードには曖昧性が解消された結果(語義)が対応します。葉以外のノードは、それが支配する葉ノードに対応する語義をすべて含む可能な解の集合を表しています。なぜなら、そのノードから枝に沿ってネットワークをたどり、それらの葉ノードに到達することが可能であるからです。したがって、根ノードは、すべての語義を含む集合に対応するため、曖昧性が最大のノードとなります。

弁別ネットワークを用いた問題解決のプロセスは、単語を読み込み、えられた制約を満足する枝に沿ってネットワークを根ノードから葉ノードに向けて、下向きに1段ずつたどる過程であると考えることができます。この過程で不要な選択枝は排除され、妥当な解が弁別(選択)されます。葉ノードに到達することが、解がみつかり、曖昧性が完全に解消されたことを意味します。

この弁別ネットワークには次の利点があります。第1は、(複数の可能な解が存在する)曖昧性を表現する際に、複数の候補を保持するのではなく、(それらをすべて下位ノードとして包含する)一つのノードを保持することですませられます。こうして、ネットワーク中



我々の方式:

○単語の語義集合を弁別ネットワーク(DN)で表現(決定木もDNの一種)

図8 「直す」の弁別ネットワーク

のノードを下向きに1段ずつたどっていくことによって、徐々に語義に関する曖昧性が解消されます。そのため、この弁別木をたどる過程を、曖昧性の漸進的洗練過程(解消過程)に対応させることができます。

第2は、解を一つ一つ逐次的に調べていく線形探索と比較して、弁別ネットワーク上での探索アルゴリズムは効率がよいことです。これは、解の探索が葉ノードに位置するそれぞれの解(語義)に向けて根ノードからネットワークを下向きにたどる操作になり、この場合、枝に記述されている制約を索引として探索するので、探索空間が徐々に小さくなります。

第3は、冗長性を避けた圧縮表現ができることです。これは、同一の制約がネットワーク中の一つの枝としてまとめて表現できるからです。これにより、制約が満足されるかどうかの検査は、一つの制約につき1回しか行われないので、同じ制約に対して何度も同じ計算を行わなくてもよいことになります。

以上をまとめますと、弁別ネットワークの葉ノードは曖昧性のない語義を表現しており、根ノードは曖昧性のまったく解消されていない状態に対応しています。ネットワークの各枝(リンク)は文を読み進むにつれてえられる名詞句に関する情報を弁別します。これが外から漸進的にえられる(与えられる)制約です。根ノードから出発して、これらの情報(制約)にしたがって枝を葉ノードまでたどることが、漸進的に曖昧性を解消する過程に対応しています。

弁別ネットワークの問題点

ところが、この弁別ネットワークを使う場合、次の問題があります。

第1は、うまく下向きにたどっていくためには、情報は根ノードに近いものから順々に入力されなければならない、制約が適切な順序で与えられなければならない。この順序は

弁別ネットワークの構造に依存してきまってしまう。つまり、弁別ネットワークは、あらかじめ決められた順序で制約が入力されないとはできません。しかし、日本語の場合、語順に比較的自由度があるので、必ずしもえられる制約がいつもきまった順序で漸進的にえられるとはかぎらず、ネットワークを下向きにたどるプロセスは、決められた順序で次々に入力されるはずの制約がえられるまで、しばしば待つ必要があります。

第2の問題は深刻です。ネットワークを下向きにたどる際に、次に入力されるはずの制約がえられない場合もあります。例えば、日本語の場合によくみられる省略語がそれです。日本語では、主語などは特に断わらないかぎり、ほとんど省略されることを考えてみてください。この場合、問題解決プロセスはネットワークを下向きにたどることができず、デッドロックに陥ってしまいます。

一般化弁別ネットワーク

これらの問題を解決するために、私どもは一般化弁別ネットワーク(GDN)を提案したわけ。弁別ネットワークのかわりにGDNを用いることにより、意味的曖昧性解消を漸進的に行うことができます。省略語があっても、語順が自由で情報(制約)がえられる順序が変わっても、必ずネットワークを下向きにたどることが可能になります。

実際には、図8に示す弁別ネットワークを図9に示す表形式に変換しておきます。図9の1行目は、「人間」に「が」がついたものがくれば、ノード2にいけ、ということを表しています。4行目は、「病気」に「を」がついたものがくれば、条件つきでノード5にいけということを表しています。ここで条件は、本来ノード5に到達するためには「人間」に「が」がついたものがくる必要があるのに、それがまだえられていないということを表します。

実際、図9に示したGDNを使って、「英語に花子が直した」、「花子が英語に直した」という文がどのように漸進的に解析されるかを図10に示します。どちらも解析後にノード7に到達します。図8からわかりますように、ノード7は葉ノードではありませんから、曖昧性が完全に解消されていないことを意味しています。しかし、可能な語義はtranslateに絞られています。

ノード7から最終的な葉ノードにいたるまでの枝に書かれた制約は、これからえたい、あるいはこれからくると考えられる情報を表しています。したがって、この情報を使うと、欠落しているものの探索範囲を定めることができます。例えば、今の場合、「を/言語；単位」や、「を/論文；尺度」、「から/言語；単位」といった情報がまだ欠落していますから、前文からこうした情報がえられるかどうかを調べればよいことになります。この方法は省略語の推定に使えることを示唆しており、現在これについて検討しています。

GDN モデルの展望

GDN モデルは画像理解などにも応用できると考えています。最初ある図形を思い浮かべたとします。ところが、「円がある」という情報がえられれば、最初の図形を円のように変

が/人間	2
を/TV；車；…	3 if が/人間
を/傷	4 if が/人間
を/病氣	5 if が/人間
を/誤り；作文；…	6 if が/人間
に/言語；単位	7 if が/人間
を/習慣	8 if が/人間
を/論文；尺度；…	9 if が/人間
	and に/言語；単位
を/言語；単位	10 if が/人間
	and に/言語；単位
から/言語；単位	11 if が/人間
	and に/言語；単位
	and を/論文；尺度；…

図9 図8からえられたGDN

(太郎が日本語で論文を書き、)

例1:

英語に 花子が 直した。

1 ↓ 7 if が/人間 ↓ 7

ifより後ろはまだえられていない情報を表現

例2:

花子が 英語に 直した。

1 ↓ 2 ↓ 7

図10 図9のGDNを用いた漸進的曖昧性解消の例

形させるでしょう。次に、「立体だ」という情報がえられますと、おそらく「円柱」を思い浮かべるでしょう。こうして、いろいろな情報が加わるにつれて、徐々に初めの像が明確になることがあります。したがって、GDNの考え方は画像理解のモデルとしても使えると考えました。

夢

最後に、自然言語処理に関係している研究者として、夢を語らせていただきたいと思います。

実現は当分難しいとは思いますが、例えばロボットに「喉が乾いた」といったときに水がでて、「お飲物は？」と聞くことのできる気のきいたロボットをつくりたいと思います。この場合、文字通りの意味は「喉が乾いた」というのですが、本当に伝達したい意味は、「飲物が飲みたい」ということでしょうか。これは、スピーチ・アクトといって、哲学者、言語学者がいろいろ研究していますが、いずれ私どもも検討する必要があります。

落語をきいて笑えるようなロボットを開発できれば最高です。落語には、いろいろな比喩や皮肉、時代時代の文化的背景などが含まれており、それらを理解することは容易ではありません。私どもが外国で喜劇を観賞したとき、外国人が笑っているのに自分はまった

く笑えず、つまらなかったという経験をしたことはないでしょうか。落語の理解は、自然言語理解システムの究極であるともいえます。

それがいつ実現するかが問題ですが、それに向けて一步一步研究していきたいと考えています。

Q&A

■Q■

私は裁判官でこの分野は専門外ですが、先ほどの話で、赤や黄色に確率値の重みをつけて、固定されていました。ところが、それが文脈によって、可変でなければならないことがあると思います。例えば、医者が重篤な肝臓疾患患者を診て、「リンゴのような色だ」といった場合、黄色のリンゴを想定して黄疸の症状を指していることもありえます。その場合、肝臓病で重篤の患者だという状況であれば、それによって属性値がかわられて、

いったん学習し、時間的・空間的に隔てるとまた別の値に戻る仕組みが必要だと思います。

●A●

その通りだと思います。個人でも例えば職業などで異なるし、リンゴの場合でも、日本とヨーロッパでは文化的な違いがあります。ヨーロッパでは「バラのような頬」という表現はありますが、「リンゴのような頬」とはあまりいわないようです。おそらく、ヨーロッパのリンゴは小さいし酸だらけで、青いものもあれば形もさまざまであるから

だと思います。性質に付与する確率値をどのようにかえたらよいかについて、私にはよいアイディアがありません。とりあえず、私どもが行った研究は、確率値は情報構造のなかには必要だろうということで付与し、それをどう用いるかについての計算モデルを考えたということです。文化なら文化が、また状況がかわったときに、それらをどのようにかえるかというメカニズムは、いずれどこかで必要になると思いますが、今後の研究課題だと思います。 ●