

Title	マルチエージェント環境での人工ピジンの生成
Author(s)	中村, 誠; 東条, 敏
Citation	認知科学, 10(2): 193-206
Issue Date	2003-06
Type	Journal Article
Text version	author
URL	http://hdl.handle.net/10119/7930
Rights	Copyright (C) 2003 日本認知科学会. 中村 誠, 東条 敏, 認知科学, 10(2), 2003, 193-206.
Description	

マルチエージェント環境での人工ピジンの生成

中村 誠, 東条 敏

A grammar of natural language is a result of statistical analyses, and is different from that of formal language. Thus, the grammar should be represented in a flexible, changeable formalism. In this paper, we propose a model of common language acquisition in multi-agents. Agents are divided into two groups, each of which has its own language. They try to communicate each other, especially with those which are in the different language group. Because the grammars of the languages are represented in the mutable form, the agents change their own grammars and temporarily build up a common grammar, that is of an artificial 'pidgin,' that enables the communication. We applied Genetic Algorithm to the statistical mutation of grammar, and utilize Lexicalized Tree Adjoining Grammar that tolerates the change of word-order. In this formalism, we could observe the phenomenon that the majority group assimilates the minority group, preserving grammatical features of each language.

1. はじめに

自然言語の文法は、形式言語のそれと違い、統計的な観察結果として求められるものであり、従ってその文法を厳密に、静的に表現するのは難しい。なぜなら言語は自然科学が対象とする外界の現象ではなく、人間の頭脳の内部で受理、生成される現象であり、きれいな法則性を持っているとは言い難い。必然的に言語の理論は言語の第一次近似を与えるものという位置付けにならざるをえない(長尾, 1996)。また、言語は時間とともに絶えず変化しており、ある時点での言語の文法を静的に表現したとしても、それが恒久的にその言語に対応できるものではない。つまり自然言語の文法とは文の生成力を規制するものではなく(小野・東条, 1998)、周りの環境の変化から柔軟に対応するように表現されるべきである。ここで我々はいかに柔軟な文法表現をするかという問題を解決するために、どのようにし

て統計的に自然言語文法が発生するかという点に着目した。

現実世界において、言語変化の本質に対する証拠を示すような言語現象が存在する。ピジン (pidgin) とクレオール (creole)(Crystal, 1987)(Pinker, 1994) がそれである。ピジンとは共通の言語を持っていない言語話者集団が接触した際に発生する混合言語を指す。ピジンはそのもとなった言語に比べて、語彙は限られており、文法構造は単純化され、機能する範囲ははるかに狭い。ピジンを母語とするものはだれもないが、その後ピジンが発達し、その話者の子供たちによって母語化されるとクレオールとなる。ピジンやクレオールは世界中で発見されており、それぞれが独自に発達した言語体系であるにも関わらず非常に似通った特徴と文法構造を持っていることから、言語学の世界で盛んに研究が行われている(Crystal, 1987)(亀井・河野・千野, 1996)。特にピジンは、ある限られたコミュニケーションの必要を満たすための言語であり、接触の初期の段階においては、思想の詳しい交流を必要としないやりとり、そして、ほとんど1つの言語だけから選ばれる、小さな語彙で十分足りるような取引に限られることが多い(Todd, 1974)。しかしながら、

Generating artificial pidgins by a multi-agent model by Makoto Nakamura (Graduate School of Information Science, Japan Advanced Institute of Science and Technology) and Satoshi Tojo (Graduate School of Information Science, Japan Advanced Institute of Science and Technology).

一人ないしはごく少数の人間が、その環境から優位となる言語話者集団と接触する際、日常的で基本的な少数の語を身につけて満足するか、優位言語をそのまま、もしくは不完全に身につけてしまう場合がほとんどであると考えられる。前者の場合、それは言語とは言えないし、後者の場合もそれはピジンではない。ピジン、クレオールに見られるような言語変化は、相互に理解することが不可能な言語の話し手が接触するところであればどこでも発生する可能性がある (Bickerton, 1990) が、そのためには、それなりの条件が整っていなければならない。また、ピジンを形成するにあたって、どの言語がより多くの語彙素材を提供するか、どの言語の音声組織が土台になるか、語順はどうするか等を左右するパラメータとなるのは、接触当事者の人口構成比、集団間の政治的、経済的、軍事的、あるいは宗教的な力関係、そして、その地域の元々の言語的多様性の度合いなどである (亀井他, 1996)。今日話されている言語は、長い時間をかけて分化と他言語との接触による融合を繰り返してきた結果であり、現在公用語として使用されているような言語でさえ、世界各地に現存するピジン、クレオールに見られる過程と似た過程を経て発展してきたのではないかと推測されている (Todd, 1974)。

本研究の目的は計算機上で人工ピジンを実現することである。ここで人工ピジンを生成するという問題は、同時に動的な文法表現と、環境に応じたその柔軟な変化をモデル化するという問題を含んでいる。本研究では、自然言語の複雑さはピジン化、クレオール化といった現象の結果もたらされたものであると考え、その変化の過程を通時的な側面から抽象化することによって、より単純な文法構造を表現することを目的とする。その初期モデルとして、本稿では動的に文法を変化させるモデルを提案することによって計算機上でピジン化の過程を観察するのが狙いである。

以上を基に、本研究では言語変化のモデル化を行い、自然言語のより自然な文法表記を提案する可能性を追求する。本研究の位置付けとしては、自らメッセージを発信できる、能動的で、自律的なエージェント間の協調、競合または組織化であると考えられる。エージェントを人間として捉えることにより、自律的な会話とそれに伴う学習を繰り返し、共通文法を獲得する過程をシミュレートする。

本稿では以下の構成をとる。まず2節では本研究を始めるにあたって調査した背景について述べる。3節は本稿で提案するモデルについてであり、4節で本モデルを計算機上に実装した際の実験の設定および実験結果を示し、考察を行う。最後に5節において、本稿のまとめを行う。

2. 言語の進化モデル

本研究においては、ピジンのような言語現象を計算機上に実装するため、そのモデル化にあたっては、具体的な表現が不可能な項目について多くの人工的な設定を必要とする。その際、それらの設定が認知科学の観点から出来る限り不自然にならず、またそれらが計算量の観点から見て現実的な定義であることが要求される。これらに関しては、すでに計算機上に実装された言語獲得モデルが参考になる。本節では、エージェントを用いた言語獲得について実験が行われた、次の4つのモデルを説明する。

共通言語、もしくは共通規約をエージェント間に構築しようという研究は、人工生物の進化論的な興味から行われてきている。Werner and Dyer (1991) は、仮想空間に配置された人工生物の集団内で、交配の相手を効率よく探し出すために共通の通信規約を遺伝的に自己組織化することを試みた。この研究は、一種の言語や通信規約を多数の自律的なエージェント群に自己組織化させようとした基本的なモデルを提供した。ただしここで導き出される言語とは、そのエージェントの動作に直結した単純なプロトコルであり、文法というほどの構造を持たない。また、Hashimoto and Ikegami (1995) によるモデルでは、各々のエージェントが文法を持ち、ほかのエージェントとコミュニケーションをすることを試みた。ここではエージェントは相手にいかに長く、複雑な文を話し、理解できるかを評価の基準として学習を行っており、これが満たされることを文法の進化と呼んでいる。この実験の結果として、当初はチョムスキー階層の正規文法しか持たなかったエージェントが、文脈自由のクラスの文法を持つようになったとしている。ここにあげた研究例は、自然言語のような高水準の言語とは大きく異なるものの、コミュニケーションをするために言語を変化させていく能力を持つマルチエージェント・モデルとして興味深い。

エージェントを人に見立てて、自然言語を対象に

したモデルを提案したのは Ono, Tojo, and Sato (1996), 小野他 (1998) である。Ono et al. (1996) は子供の文法の精緻化という問題を扱っている。これは文法が未熟な子供が大人との会話を聞くことにより、素性を精緻化していく過程をモデル化したものである。また小野他 (1998) では推論機能を有するエージェントからの共通文法の組織化を行うモデルを提案している。ここでは自然言語の特徴として適応性 (adaptability) と頑健性 (robustness) を取りあげ、エージェントの持つ推論機構から共通文法の組織化を行った結果、言語の融合と分化の過程を模擬的に実現することができたと述べている。

言語獲得問題とは、生まれたばかりの子供が言語を獲得する際の問題を扱う第一言語獲得と、第二言語獲得に大きく分けられる (Archibald, 2000)。上記における小野の2つのモデルに関しては前者が第一言語獲得問題、後者が第二言語獲得問題に関するものと考えられる。筆者らはビジン化とは、目的言語が不明瞭なまま (Bickerton, 1990)、集団が互いに相手の言語を学ぼうとする第二言語獲得問題であり、クレオール化に関してはそのビジンを元にした第一言語獲得問題であると捉えている。本研究では上記モデルを参考にし、次節において人工ビジンを形成するモデルを提案する。

3. 人工ビジン生成モデルの提案

本節では、これまでに述べてきたことをもとに、ビジン化、すなわち共通文法獲得モデルをマルチエージェントの枠組みで提案する。

本モデルは、日本語の文法を予め獲得している日本語エージェント群と、英文法を獲得している英語エージェント群から構成される。これらが相互作用することにより、各エージェントが共通の文法を獲得する過程をシミュレートする。

なお、本稿において次のようにエージェントに関する用語を定義する。

同言語エージェント：日本語エージェントに対する
その他の日本語エージェント・英語エージェント
に対するその他の英語エージェント。

異言語エージェント：日本語エージェントに対する
英語エージェント・英語エージェントに対する
日本語エージェント。

本研究で取りあげる日本語と英語の差異は、日本語は助詞をマーカとして、英語は名詞句の表層位

置と代名詞の活用により、格を指定しているということである。最初のうちは文法の違いによりコミュニケーションできなかったエージェント同士が、試行錯誤を重ねるうちに、これらの特徴を併せ持つように文法を学習することによりコミュニケーションがとれるようになることを期待している。ここで学習によって得られるものは、メタレベルの文法書き換え規則であり、エージェントは共通文法を獲得するために既存の文法を書き換えていく仕組みを持つ。

また、本モデルの制約として以下のことを仮定する。

- 同言語エージェント群だけで会話を行う場合、学習規則により文法が変化してはならない。
- 日本語話者のエージェントが完全に英語を理解し、話すことができるだけの、文法のメタレベルでの書き換え規則が存在する。また逆も同じ。

前者は言語の安定性に関する制約であり、外的要因が無い限りは文法が変化すべきではないことを意味する。また後者は言語変化の可能性についてであり、エージェントが完全な第二言語獲得の能力を持っていなければ意味がない。状況によっては相手の言語を習得しなければならないという条件のもとで、お互いが歩み寄っていった場合、どのような言語変化を起こすのかが本研究における最大の関心である。

3.1 中間表現の定義

エージェント間で会話が成立するということは、発話された内容が聞き手のエージェントにとって有意味であり、話し手の意図が伝達されたということである。本来、ことばの通じない相手に意思を伝えるには身振り手振りなど、言語とは別のモダリティを用いることが考えられる。たとえば食卓で「塩を取って」と言ったときに、それを聞いた人が「理解した」とは、その人が塩を取って渡してくれるという動作に及んでくれることである。しかしながら別モダリティの情報を画像として計算機上で処理し、知覚することを、複数のエージェントが行うことは現実的ではないし、その知覚処理自体は言語研究の枠内ではない。

本研究では言語の論理意味論、特に機械翻訳との関連においての研究成果 (長尾, 1996) (Nirenburg, 1987) (田中・辻井, 1988) を踏まえ、意味記述方式として中間表現を考える。すなわち、本研究において「意思が伝達された」とは、ある言語表現を受け

取ったエージェントが (i) それを言語として認知し, (ii) それに相当する意味表現を内部に作ることができることとする。よってこれに呼応して, (i) 相手の言語表現をパースし, (ii) それが well-formed な形式表現に変換されることを本研究での“理解”の形式化と定義し, 聞き手が話し手に自分の構成した中間表現を返すことで理解の正しさを検証するという方法を取る。この定義は古くは生成文法 (Chomsky, 1981) の深層の論理構造を意識した定義であるとともに, 工学上も機械翻訳のピボット方式で用いられる手法であることから, 伝統的な言語学から工学的応用の広い幅で受容可能な「理解」の定義と考える。

以上より, 中間表現は Fillmore の深層格 (Bruce, 1975) を基に定義する。具体的には“私はあなたにペンをあげる”という中間表現は,

[*give:(agt:*I)(co-agt:*you)(obj:*pen)]

と表現される。ここで*が先頭についた語は, その動作や事物に直結したもっとも単純な意味表現である。また agt, co-agt, obj はそれぞれ同一括弧内の意味要素が動作主格, 対行為者格, 目的格であることを意味する。先頭に記述されている動詞がこれらの格を持つことにより, 一つの中間表現を形成する。

3.2 文法の定義

3.1 節において, 中間表現を定義した。ここではこの中間表現からの文生成および, 文からのパースによる中間表現の抽出が可能となる文法の定義を行う。これらにおける処理は, その文法を用いることにより, パースと生成の両方の操作が可能となることが望ましい。また, 中間表現にも示されるように, 動詞となる語が中心となって処理されるため, 中間表現からのトップダウンによる文生成, ボトムアップによるパースを行う際, その動詞に対する下位範疇化項目が, 構文木中の根ノードからも, またその動詞が位置する葉ノード, すなわち単語からも参照が可能であることは, 処理上有利である。そのためには語彙化がなされている文法であり, ユニフィケーションによって文生成とパースが行われる文法が望まれる。

以上の要件から本研究では, LTAG (Lexicalized Tree Adjoining Grammar) (長尾, 1996) (Abeillé & Rambow, 2000) を文法の枠組みとして用いる。LTAG は CFG のように記号列を書き換える文法ではなく, 構文木を直接書き換える文法である。また,

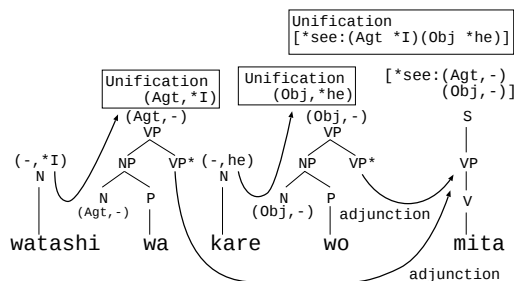


図 1 LTAG による日本語の文法表現

Fig. 1 Expressions of Japanese grammar in LTAG.

各ルールは必ず一つの単語を持ち, 文法の語彙化がなされている。

この文法を採用することは, 逆に, 言語共有の問題において語彙をどのように共有するかという根幹の問題の解決にも寄与する。すなわち, 語彙は各エージェントが手元に持っている辞書で word-by-word で翻訳を与えればよいというわけではなく, 語彙に依存して構文も変化させるべきものである。LTAG はこの要請にも同時に応えるものである。

本モデルではこの LTAG に深層格を持たせるように拡張した (図 1 参照)。これにより, 中間表現からの文生成, またはパースによる概念の抽出が, 深層格素性のユニフィケーションにより, 容易に行うことができる。

3.2.1 日本語の文法の表現

本モデルで日本語エージェントが所有する文法例を図 1 に示す。掛川・神田・藤岡・伊丹・伊藤 (2000) の日本語文法の定義を基にしており, 図 1 のように, 名詞を格助詞に修飾させた名詞句を接合木とすることにより, 語順の自由を表現することが可能となる。

3.2.2 英文法の表現

本モデルで英語エージェントが所有する文法例を図 2 に示す。日本語と同様, 名詞句を接合木としており, また主格に相当する語と目的格に相当する語を, footnode の位置, すなわちその単語の修飾の方向を変えることにより, 語順を確定させることとする。

本稿ではビジン化を前提としているため, 英文法として定義する素性を最小限にとどめた。従って本モデルでは, 計算量の削減のため, 自然言語処理の

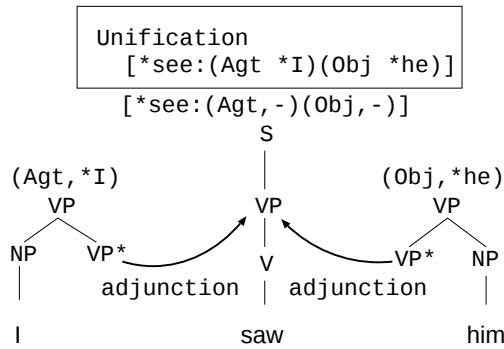


図2 LTAGによる英語の文法表現

Fig. 2 Expressions of English grammar for this model in LTAG.

分野で作成されている既存の文法 (XTAG Research Group, 2001) を使用するのではなく、単純化された文法を定義することとした。

3.3 エージェント間通信

同言語エージェント同士の会話モデルを図3に示す。図中には文法の学習機構は含まれていない。学習を含めたエージェント間での文のやりとりを図に沿って以下に説明する。

● 文の発話手順：

- (1) ランダムに中間表現を発生させる。
- (2) 中間表現と文法を基に文を生成する。その文法を使用して文を生成した結果が適応度に反映される。
- (3) あるエージェントに対し、生成した文を発話する。
- (4) 発話を受け取ったエージェントから中間表現が返ってくる。自分の中間表現と相手の中間表現を照合した結果が適応度に反映される。

● 文の理解と中間表現での応答手順：

- [1] エージェントが文を受け取る。
- [2] その文を自分の文法を用いてパースする。
- [3] パースした結果から、中間表現を得る。
- [4] 中間表現が得られたら、パースに用いた文法の適応度に反映される。
- [5] 発話したエージェントへ中間表現を返す。

以上のように、発話と中間表現による応答が一組となって会話を形成する。

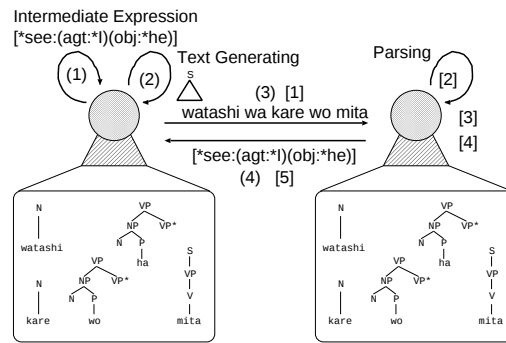


図3 エージェント間の会話の流れ

Fig. 3 The communication between agents.

3.4 学習機構

このモデルでは、遺伝的アルゴリズム (Genetic Algorithm, GA) (伊庭, 1994) による学習を行い、会話から得られた適応度によって、自分の文法に変化を加えたものを評価する。

以降、GAの用語で「世代」という言葉が出てくるが、この世代と、ビジンを学習する人間の世代とは異なる。ビジンとはわずか一世代で形成されるのが特徴のひとつであり、従って何世代にも渡って行われるGAの学習というのは、一人の人間の中で行われるものであるとする。

以下に学習機構の詳細を示す。

3.4.1 遺伝子からの変化文法への表現

エージェント m は世代 L において文法 $G_{L,m}$ を保有している。文法 $G_{L,m}$ は n 個から成るルール $g_i (0 \leq i < n)$ の集合であり、すなわち $G_{L,m} = \{g_0, g_1, \dots, g_{n-1}\}$ である。以降特に明示が必要ではない場合、添字は適宜省略する。ここで例えば日本語エージェント J が世代 0 で持っている文法 $G_{0,J}$ は日本語の文法そのものである。またエージェントは染色体を保有しており、エージェント m は $G_{L,m}$ からこの染色体の表現型 (phenotype) として $G'_{L,m}$ を得る仕組みを持つ。このように、1個の染色体から独立した文法セット $G'_{L,m}$ を得ることによって、染色体の中の遺伝子の操作によって新たな文法セットの導出が可能であるということと、 c 個の染色体から $G'_{L,m}, G''_{L,m}, \dots, G^{(c)}_{L,m}$ のように、それぞれが $G_{L,m}$ を基にしている且つ、 c 個のお互いに独立した文法セットの導出が容易であるということから、GAを用いた文法の変換につ

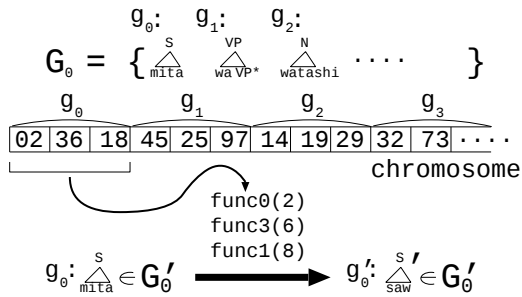


図4 遺伝子の表現型を介した文法の変換

Fig. 4 The method of changing own grammar by phenotype of gene.

いてこのような手法を採用した．以下にエージェント m が1個の染色体から $G'_{L,m}$ を得る手法について述べる．

エージェントが持つ染色体の遺伝子型 (genotype) は, 2桁の10進値を1組にした, 3組の遺伝子から構成される (図4参照)．文法 $G_{L,m}$ に含まれる各ルールと遺伝子とが対応づけられており, ルールに対応した3組の遺伝子が, 文法ルール $g_i \in G_{L,m} (0 \leq i < n)$ に作用する関数およびその引数として適用される．

本モデルにおいて, 1つのルールに対して3組の遺伝子を割り当てたが, 一般的には各ルール毎に適用される遺伝子の数については, 適用される関数群に依存する．3節の最初の箇条書きで述べた本モデルの二つ目の制約は, 日本語話者のエージェントが完全に英語を理解し, 発話することができるだけの, 文法のメタレベルでの書き換え規則が存在することである．次に挙げる関数群がそれにあたり, 適用する関数の組み合わせによってお互いの文法に変換可能である．

指定された文法ルール $g_i \in G_{L,m} (0 \leq i < n)$ に対して適用される関数の内容を以下に定義する．

func0 何もしない．

func1 適用されるルールにおける木構造上の, 引数で指定された中間ノードの, 子ノードの位置を反転する．

func2 適用されるルールの, 引数に応じた単語を変換する．例えば $*\text{see}$ のルールの語彙は, 引数に応じて “見た” または “saw” に割り振られる．このルールを適用しても, 語彙が変化しないケースもありえる．

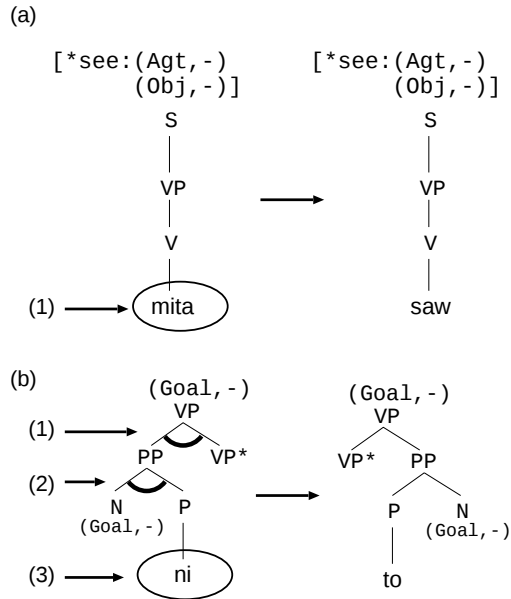


図5 日本語文法から英文法への変換の例

Fig. 5 Examples of transformations from a Japanese rule to an English rule.

func3 もし適用されるルールが, 接合ノードを含んでいなければ, 接合ノード (VP) を根ノードおよびその子ノードとなるフットノードを追加する．

func4 もし適用されるルールが, 接合ノードを含んでいるならば, 接合ノードを削除する．

この関数群を用いて英語から日本語, もしくは日本語から英語への文法書き換えを行う場合, 3つの関数の適用が必要となるルールが存在する．例えば, 図5(a)では “見た” に関する日本語文法のルールを英語のルール “saw” に書き換えるために, 単語の変換 (func2) のみ, すなわち関数を1回適用するだけでルールの書き換えが可能であるが, 図5(b)では終点格を示す日本語の “に” を英語の “to” に書き換えるためには, (1) VPの子ノードの位置を反転 (func1), (2) PPの子ノードの位置を反転 (func1), および (3) 単語の変換 (func2) という3回の関数の適用が必要となる．従って各ルールにつき, 3組の遺伝子が割り当てられる必要がある．

以下に G_0 から G'_0 を導出する例として, 図4についての説明を行う．図4ではルール g_0 が “見た” についてのルールであり, g_0 に適用される遺伝子

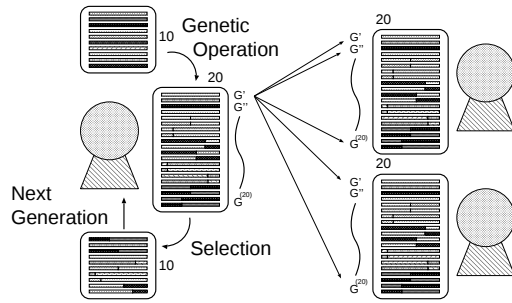


図 6 会話モデルへの GA による学習の適用

Fig. 6 The application of GA learning into the communication model.

の値は、それぞれ 02, 36, 18 となる。これはそれぞれ関数番号 0, 引数 2, 関数番号 3, 引数 6, 関数番号 1, 引数 8 の 3 つの作用関数を $g_0 \in G_{L,m}$ に対して適用することを意味する。これらを適用した結果、得られたルールを $g'_0 \in G'_{L,m}$ とし、同様に全ルールに適用した集合を G' とする。すなわち $G'_{L,m} = \{g'_0, g'_1, \dots, g'_{n-1}\}$ 。これにより文法 $G_{L,m}$ を用いることによって、1 個の染色体の表現型として $G'_{L,m}$ を導出することが可能となる。

3.4.2 GA を用いた学習

ここではエージェントが行う学習を、本モデルに適用した実際の値を交えて説明する（図 6 参照）。

各エージェントは各々 10 個の染色体を所有している。1 個の染色体の表現型が 1 つの文法セットであるから、この時点でエージェントは元々所有している 1 つの文法セットに加え、10 種類の文法セットを持っていることになる。さらにこの 10 個の染色体に遺伝的操作を加え、20 種類の文法セットに増やした後にエージェント同士で会話を行う。ここで、たった 1 つの共通言語を得るために、人間が 20 個もの独立した文法セットを所有し、会話をするとする設定は、単に実装上の問題を解消するためであり、文法を逐一変化させて評価する様子をまとめたものである。

各世代の最初に次のような遺伝的操作を加えることにより、染色体の数を 10 個から 20 個に増やす。

- (1) 10 個の染色体から非復元抽出により、2 つの染色体からなるペア 4 組をランダムに抽出する。
- (2) これらのペアに対して 1 回交叉を行う。2 通りの交叉をしたものを新たな染色体に加え

る。この場合、1 ペアに対して 4 つの染色体がつくられる。

- (3) 16 個の染色体中の各遺伝子に対し、4% の確率で突然変異をさせる。
- (4) 最後に元の染色体セットから、前世代で適応度が高かった染色体 4 つを加え、新たな 20 個の染色体セットとする。

これらのパラメータ（4 組、4% など）の設定は実験を繰り返しチューニングした結果、意図する効果が発現するところに固定したものである。

ここで各染色体の表現型は 3.4.1 節で説明した通りであり、エージェントが持っている文法 G に対し、 $G', G'', \dots, G^{(20)}$ がつくられる。これらの文法セット $G', G'', \dots, G^{(20)}$ を用いて、エージェント間で会話を 1 世代につき N 回行い、適応度を計算する。1 エージェントが 1 世代で会話を行う回数は、(染色体の数) $\times$ (発話対象) $\times$ (発話回数) = $20 \times ((|Agents| - 1) \times 20) \times N$ である。ここで $|Agents|$ はエージェントの数となる。

適応度の算出方法は次の通りである。3.3 節で説明した会話の各部分において、用いられた文法ルール（すなわち染色体）の適応度について加点する。最初の中間表現はランダムで求められ、それはエージェントが持っている初期の文法 G_0 において、文生成およびパースが可能なものとする。1 世代の会話数 N 、第 L 世代のエージェント m について、 i 番目の染色体の表現型 $G_{L,m}^{(i)}$ の適応度は次の算出方法を全エージェントの全染色体に対して N 回適用して求まる。ここで p_c, p_f はそれぞれ同言語エージェントによる会話の評価点および異言語エージェントとの会話の評価点である。

- ランダムに選んだ中間表現から、 $G_{L,m}^{(i)}$ を用いて文を生成することができたらその遺伝子に対して $+p_c$ 。
- 生成した文をエージェントに渡し、返ってきた中間表現が合っていたら、会話相手が同言語エージェントのとき $+p_c$ 、異言語エージェントのとき $+p_f$ 。
- 任意のエージェントから文を渡され、それを $G_{L,m}^{(i)}$ によってパースして中間表現を求めることができたなら、会話相手が同言語エージェントのとき $+p_c$ 、異言語エージェントのとき $+p_f$ 。

ここではやはりチューニングの結果、 $p_c = p_f = 1$ とした。

その世代において会話が終了すると、適応度が求まるため、これら 20 個の染色体セットから単純選択により適応度の高い上位 10 個の染色体を次世代に残す。

3.4.3 自己の文法の更新

これまでの学習では、世代毎に各エージェントが持つ文法を書き換えた文法を生成しているものの、元の文法に変化はない。すなわちエージェント m が持っている文法は、常に $G_{L1,m} \equiv G_{L2,m}(L1, L2 \text{ は任意の世代})$ である。ここに一定の世代毎に文法を書き換える操作を加える。 R 世代毎に、最も適応度が高い遺伝子の表現型によって変化させた文法を次世代の文法とする。すなわち $G_{L,m} \equiv G_{L-1,m}^{(h)}$ (ただし $\text{mod}(L, R) = 0$)、ここに h は最も適応度が高かった染色体とする。

なお、元の文法を書き換えてしまうと、遺伝子群は元の文法を対象とした関数群を表しているため、その表現型である文法は全く異なったものになってしまう。このことから、文法を書き換えた直後に遺伝子群を乱数で初期化している。

本モデルにおいては、全ての実験について書き換えを行う世代を $R = 20$ 世代とした。

4. 実験および考察

3節で述べたモデルを実装し、実験を行った。本節では実験と実験結果、およびその考察について述べる。

4.1 実験

ここでは実験のためのパラメータ設定を行う。

ランダムに発生させる中間表現では次の 6 種類の中から常に選択するようにした。各エージェントは第 0 世代において、これらの中間表現から文を生成できるだけの文法 $G_{0,m}$ を持っているものとする。

- $[\text{*run:}(\text{Agt *I})]$
- $[\text{*run:}(\text{Agt *he})]$
- $[\text{*see:}(\text{Agt *I})(\text{Obj *he})]$
- $[\text{*see:}(\text{Agt *he})(\text{Obj *I})]$
- $[\text{*go:}(\text{Agt *I})(\text{Goal *tokyo})]$
- $[\text{*go:}(\text{Agt *he})(\text{Goal *tokyo})]$

パラメータを変化させることにより、3 つの実験を行った。実験内容を表 1 に示す。

実験 1, 2 では 2 つのエージェント群の人口構成

表 1 各実験のパラメータ

Table 1 Parameters of each experiment.

	$ Agt_J $	$ Agt_E $	R	N	L	p_c	p_f
実験 1	2	10	20	1	1000	1	1
実験 2	10	2	20	1	1000	1	1
実験 3	6	6	20	1	1000	1	1-5

$|Agt_J|$: 日本語エージェント数

$|Agt_E|$: 英語エージェント数

R : 辞書書き換え世代数

N : 1 世代あたりの繰り返し会話数

L : 実験世代数

p_c : 同言語エージェントとの会話の得点

p_f : 異言語エージェントとの会話の得点

比に偏りが見られる場合、どのような文法に収束するかについての実験である。本モデルの制約から、一方の言語話者は他方の言語を習得するだけの文法書き換え能力を持っているため、人口の多くを占める言語（優勢言語）に吸収されることが予想される。

また、実験 3 ではエージェント群の人口構成比が同じ場合、2 つのエージェント群が話す文法が、1 つに収束するかどうかについての実験である。以降に実験結果とその考察を述べる。

実験結果は、ヒット率を各世代毎の出力とした。ここで

$$\text{ヒット率} = \frac{\text{発話先エージェントが認識した数}}{\text{発話数}}$$

である。

なお実験結果を図 7 から図 10 までに示しているが、これらは全て R 世代毎、すなわち自己の文法を書き換える直前の世代のヒット率を抽出したグラフである。各実験結果は、同条件での実験を 5 回行い、その結果の平均をとっている。

4.2 実験 1, 2 の結果と考察：優勢言語の学習

実験 1, 2 について説明するにあたり、人口が多い言語を持つ方のエージェントをメジャーエージェント、少ない方をマイナーエージェントと呼ぶことにする。例えば実験 1 では英語エージェントがメジャーエージェント、日本語エージェントがマイナーエージェントである。実験 1 の結果を図 7、実験 2 の結果を図 8 にそれぞれ示す。図中、若い世代においてヒット率が格段に低い値を示している線が、マイナーエージェント群、すなわち実験 1 では日本語エージェントであり、実験 2 では英語エージェントを指している。マイナーエージェント群のヒット率



図 7 実験 1 の結果
Fig. 7 Result of Exp1.

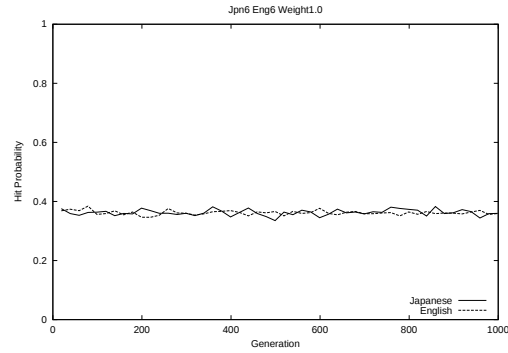


図 9 実験 3 の結果
Fig. 9 Result of Exp3.

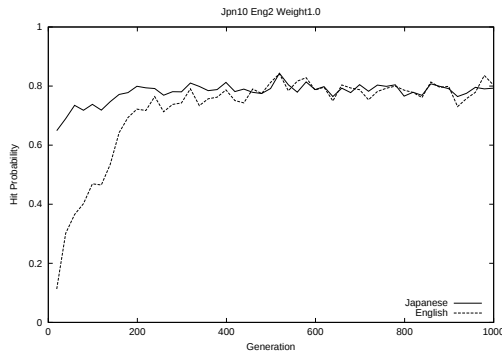


図 8 実験 2 の結果
Fig. 8 Result of Exp2.

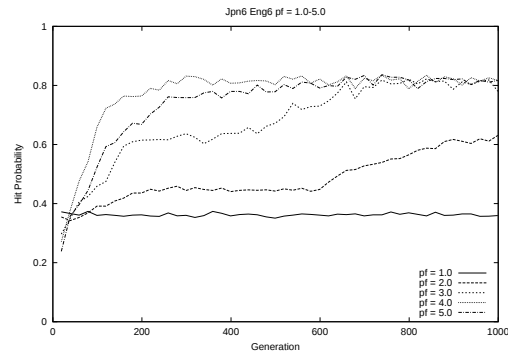


図 10 p_f 値の増加によるヒット率の変化
Fig. 10 The hit probabilities by p_f 's

を上昇させるためには、メジャーエージェント群の文法の理解が必須となる。

4.1節で述べた通り、学習によって世代を重ねる毎にヒット率は上昇し、その後、安定する。理論上、 $G'_{L,Minor} \equiv G_{L,Major}$ となるような、すなわちマイナーエージェントが持つ文法から、1世代でメジャーエージェントが持つ文法と全く同じになるような文法書き換え規則を表現型とする遺伝子型の存在を許しているにも関わらず、1世代での劇的な文法の変化、すなわち全く話が通じていなかった状況から完全にメジャーエージェントの文法を獲得するような文法の変化は確認していない。これは、初期状態の乱数値を持つ遺伝子から、そのような表現型が現れる確率が非常に稀であることと、世代を重ねるごとに適応度が極大もしくは最大に近づくように遺伝子が収束していくといった、GAを用いた学習の特徴から考えられることである。収束した文法から得られる言語のほとんどは、メジャーエージェ

ントが第0世代で持っていた文法、 $G_{0,Major}$ と同じであった。なお、実験1では、日本語エージェントが最終的に完全な英文法を獲得するまでに平均で616世代、実験2において、英語エージェントが日本語文法を獲得するまでに平均で248世代かかっていた。

実験1, 2は、本モデルの健全性を証明する実験であるといえる。

4.3 実験3の結果と考察：言語の混合

この実験は2つのエージェント群の人口構成比を同じ6人ずつにした場合、どのような文法変化が起き、その結果、単一の文法に収束するかどうかの実験である。

実験3の結果の1つ、実験1, 2と同様 $p_c = p_f = 1$ としたときの結果を図9に示す。世代を重ねてもヒット率にほとんど変化は無く、5回の試行のうち文法に変化が起こったものは1つもなかった。ここで

表 2 各世代における発話例

Table 2 Examples of utterance in each generation.

	日本語エージェント発話例	英語エージェント発話例
第 0 世代	watashi wa hashitta kare wa watashi wo mita watashi wa Tokyo ni itta	I ran he saw me I went to Tokyo
第 40 世代	I ran he mita I wo / kare watashi mita I Tokyo went / to Tokyo I itta	I ran he saw watashi / he me saw I Tokyo went (Tokyo I went)
第 80 世代	I ran he I wo mita / he mita watashi I Tokyo went	I ran he watashi saw / he saw me I Tokyo went
第 180 世代	I ran he mita watashi I Tokyo went	I ran he mita watashi I Tokyo went

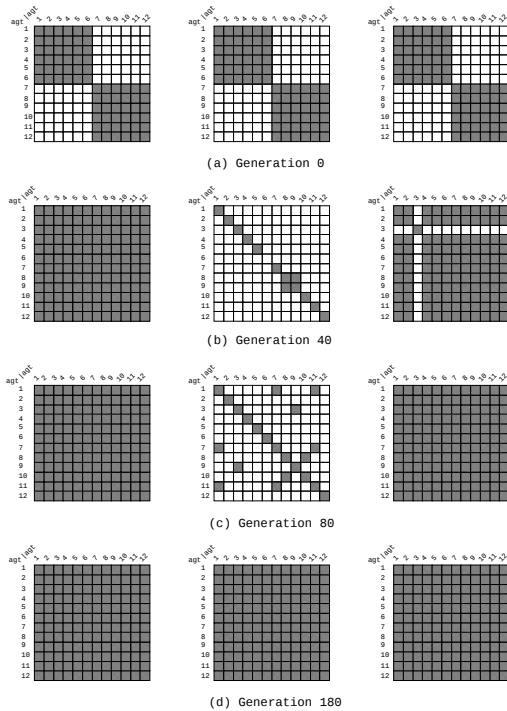


図 11 各世代における中間表現ごとの伝達状況
Fig. 11 The communication rate through generations.

は 1000 世代に渡って実験を行っているが、 $R(=20)$ 世代毎の文法書き換えが無いということは、0 世代から 20 世代の間のヒット率を 1000 世代まで繰り返

しているに過ぎず、このためヒット率は安定していると考えられる。ただし、エージェントが保持している遺伝子には常に突然変異と交叉の操作が加えられるため、同言語エージェント同士ならば必ず発話が認識されるわけではない。 $|Agt_J| = |Agt_E| = 6$ である場合、自分以外の同言語エージェントと会話をする割合が $5/11 \approx 0.4545$ であり、同言語エージェントへの発話が全て認識されるのならば、この値がそのままヒット率になると考えられる。実験ではそれよりも若干低い値を示しているが、それは遺伝子を介した文法を用いて会話することに起因する。

上記実験の結果では文法の変化が見られなかったため、引き続き p_f の値を変更して実験を行った。 $p_f = 1, 2, 3, 4, 5$ とした場合のヒット率の変化を図 10 に示す。図中の各線は、各 p_f の値毎の全てのエージェントのヒット率の平均値を示している。 p_f の値の増加と共にヒット率も上昇し、更に収束するまでの期間が短いことがわかる。 $p_f = 2$ の場合、5 回の試行のうち、 $p_f = 1$ と同様、文法に変化が見られなかったものが 2 回、全エージェントで共通の文法を獲得できたものが 2 回存在した。残りの 1 回は 11 人のエージェントが英語を話すようになった。 p_f の値がそれよりも大きくなると、全ての試行において全エージェントが共通の文法を獲得したことが確認できた。また、 $p_f = 2$ のように一方の言語に収束してしまうという傾向は、 p_f 値の上昇と共に

に減少した。

実験3の結果の1つ, $p_f = 5$ について, 図10および図11で考察していきたい。実験ではヒット率の他に, 文法書き換えの時点で各エージェントが保持する文法 $G_{L,m}$ (ただし $\text{mod}(L, R) = 0$) を出力した。また, その各文法を用いた発話例と, 他エージェントとの会話が可能かどうかの対応も出力しており, その対応を図11で示している。図中の agt1 から agt6 までは日本語エージェントで, それ以降が英語エージェントである。縦軸に示されたエージェントが横軸に示されたエージェントに発話した結果, 理解された場合, その対応箇所を黒く塗り潰してある。各世代に3種類の対応が描かれているが, 左から

- (1) $[\text{*run: (Agt *I)}]$
- (2) $[\text{*see: (Agt *he)(Obj *I)}]$
- (3) $[\text{*go: (Agt *I)(Goal *tokyo)}]$

を発話した結果となっている。自分自身が生成した文は, それが動詞と格要素を持った文であれば, 必ずパースが可能であるが, 図中, 対角要素が白くなっているマスが存在する。これは, そのエージェントの持つ文法の変化により, その格要素が代入もしくは接合することができず, 中間表現から全ての格要素を持った文を生成できなかったことを意味している。その場合, たとえ自分自身が生成した文であっても, その文をパースした結果から格要素を全て満たした中間表現を得ることはできないため, 対角要素は必ず黒く塗り潰されているわけではない。

ここでは, 文法 $G_{L,m}$ を用いた, 特に文法に大きな変化があった第0, 40, 80, 180 世代でのエージェントの代表的な発話例を表2に示す。第0世代においてはそれぞれ日本語, 英語の文法を各グループのエージェントが所有しており, 同言語エージェント同士での会話でのみお互いが理解し合えるようになっている。その様子が図11(a)に示されており, 3つの発話例についての各エージェントの対応の分布がわかるようになっている。しかし, 自身を除いた約半数のエージェントと会話が成立する筈であり, $p_f = 1$ ではそのような傾向を示しているのに対し, $p_f = 5$ ではヒット率のグラフはそのように示していない。これはエージェントが持っている文法 $G_{0,m}$ で直接会話をしているのではなく, 各エージェントが持つ遺伝子の表現型である $G_{0,m}^{(1)} \cdots G_{0,m}^{(20)}$ を用いているため, 乱数で初期化された表現型では, ほと

んどのエージェントに対して会話が成立しないことを表している。 $p_f = 5$ の場合は異言語エージェントとの会話が成り立つことによって高い得点を得られることから, その得点がヒット率に反映されていない。

ヒット率を上げるためには, まず同言語エージェントとの会話を確立すべきであると考えられる。そうすれば $p_f = 1$ と同様, 約40%までヒット率を上昇させることが見込めるからである。しかし, 第40世代の発話例を見ると, 格を1つしか用いていない, 最も単純な中間表現である,

$[\text{*run: (Agt *I)}]$

について, 早々にして共通文法を確立している。その発話状況は表2に示す通りであるが, これらのルールは世代を重ねても変化せずに固定している。これは, 異言語エージェントと話を通じれば, 同言語エージェントのそれと比べて, 5倍の得点が適度に反映されることが効いていると考えられる。代わりに, その他の中間表現に対する会話状況は図11(b)中, 右に示す通りであり, 同言語エージェントとの会話が必ずしも成り立っていないことを意味する。特に図11(b)中に関してはほとんどのエージェントが, 自分以外のエージェントの発話内容を理解できないまでに文法が発散している。

第40世代において確立したのは, 上記中間表現を理解するだけでなく, 「(Agt *I)は動詞の左側に掛る「I」である」というルールを得たことを各エージェントは認識している。このルールはその後, 他の文を生成, 理解する際にも適用され, ヒット率を上昇させることを手伝っている。同様に, その他の格についても単語, 掛る方向を確定していくのである。

このように世代を重ねていき, 最終的に全ての共通文法が得られるまでに180世代, 5回の実験の平均では372世代かかった。ここで得られた共通文法の特徴としては,

- (1) 共通の単語の使用
- (2) 日本語文法として持っていた格助詞の欠落
- (3) 英文法として持っていた代名詞の活用の欠落
- (4) 名詞句が動詞に掛る方向と格の関係

が挙げられる。これらは実際のピジンにも見ることができる。以下にピジンの発話例と, 実験結果との差異を示す。

- (1) 語彙の共有: 一般に語彙の多くは一つの言語

(多くの場合、上層言語が語彙供給言語となる)から採択される(Arends & Muysken & Smith, 1994). 次の文はトク・ビジンの例であり、斜体が現地語である(亀井他, 1996).
Mi kaikai nau. (Tok Pisin)

(私は今食べ始めたところです.)
 本実験では二つの言語が対等であったため、双方の語彙が混ざったものとなった.

- (2) 格助詞、活用欠落: 時制、法、相(TMA)といったものに関しては、ビジンにはほとんど見られない(Arends et al, 1994). 下の文は92歳の日本人移民が発話したハワイビジンである(Pinker, 1994).

Me capè buy, me check make.

(わしコーヒー買う。わし小切手作る.)
 上記2つの例文にある‘mi’または‘me’は一人称単数を指し示すものであるが、格に制限は無い。本実験においても多くの場合、同様の結果が得られたが、格助詞が残るケースも確認された。

- (3) 語順: ビジンの語順は様々であるが、少なくとも一つの接触言語の語順に準じている(Arends et al, 1994). 本実験においては多くの場合SOVまたはSVOの語順に収束したが、希にそうでないケースも確認された。

以上のように、本実験ではビジン化における過程で発生する現象と同様な現象を計算機上で再現することができた。

5. おわりに

本研究の目的は言語学上実際に存在する混合言語であるビジンやクレオールを計算機上で実装することであった。本研究ではこの問題を、エージェントを人間として捉えることにより、マルチエージェントの枠組みでシミュレーションを行った。また、人口構成比を変化させることによりビジンが発生する条件を検証した。ここでいう人口構成比とは、実社会でいうところの力関係を一次元的に投射したものと考えることができる。この実験から、集団間の人口に大きな差がある場合、人口が少ないエージェント群が話していた言語はもう一方のエージェント群の言語に吸収され、優位言語に統一されるようになった。また、人口に差がない場合には、どちらの言語でもなく、双方の特徴を持った混合言語に収束していく様子が観察できた。一般に文法書き換え規

則と、それを取り扱う学習機構さえあれば、ビジンもしくはその他の自然言語変化を表現することが可能というのではなく、図9に示す実験結果のように、全く変化しない例もある。このようにエージェント構成や各種パラメータ設定によってビジンの出現状態を示すことができたことは、本モデルの最大の特徴であると考えられる。

しかしながら、本研究においては現実の人間社会から作り出される言語現象を計算機上で発生させるため、実際の人間の行動と比較して多くの人工的な設定を施してしまったことは否めない。例えば本稿では計算機への実装を考慮し、GAによって学習しやすい設定を行ったが、帰納的学習は行われていない。言語獲得モデルを提案する際、2節に挙げた関連研究のようにいかに単純なモデルで言語変化を発生させるかということと、いかに現実世界の言語現象を模倣するかという、両極端ともいえるこれらの特徴を考慮する必要がある。本稿においてはこの両方を満たすべく単純な学習方法を採用したが、これらの改善については次のモデルへの課題とする。最後に文法に関してであるが、LTAGを用いた文法に関しては、下位範疇化原則を無視した文法定義を行うことによって、言語の互換性という制約を保持するようにした。LTAGとはそのルールにより、動詞に依存する下位範疇化項目を、表層で特定した語彙化が可能であることが大きな特徴となっている。本稿では日本語の語順の自由度を表現するため、また英文法から日本語文法への完全な書き換えを許すように定義するため、結果的には上記のような、TAGの本質的な特徴の一つを失ってしまった。しかしその代わりに、語順が自由な日本語と表層位置によって格が決定する英語の言語間の互換性を得た。今後は、これらの設定をより根拠のあるモデルに改善していく必要がある。

さらに今後の展開として考えられるのは、モデルや文法等を現実的な規模にまで拡張し、本稿で得られた結果との比較検証をすることが挙げられる。本モデルは現実のビジン化の過程を抽象化したという設定であり、実験規模を拡大したときに結果が大きく異なるというのであれば、本モデルの設定が妥当ではないということになる。しかし筆者らは次の理由でエージェントの人口と文法を増やすことで結果は大きく変わらないと予想している。

まずエージェントの人口であるが、本モデルでは

異なる言語集団の初期の接触段階で発生する初期ピジン想定している。このときの接触人口は数人から多くて数百人であると考えるのが妥当であろう(安定したピジンの場合、その人口は数百万人規模になるものもある)(亀井他, 1996)。したがって人口の規模として本実験の10人・20人というのは現実の言語現象を反映できる規模であると考えている。それよりも人口を増やした場合の魅力は、様々な接触状況を設定できることである。例えばハワイで発生したハワイピジンなどは、農園の経営者である英語話者の人口が、全体と比較して非常に少数であり、更に他の言語話者との接触が非常に限られた環境にもかかわらず、英語が上層言語となり、語彙供給言語となった(Bickerton, 1990)。このようなモデル化としてエージェント間の階層を設定し、それに準じて接触機会や会話に対する評価を環境として与えることにより、ピジンが発生する環境のより現実的な境界条件を求めることが考えられる。

次に文法の拡張に関してであるが、本実験で用いた文法は、現実世界で使用されている自然言語と比較して、今回の設定では品詞、格、時制といった素性を欠いたものになっている。これらについて遺伝子を用いて単純に拡張することは困難である。しかし3.2.2節に述べた通り、本実験の文法はピジン化を前提としたものであり、限られた意志疎通を目的としている。また、多くの、おそらくどの言語社会においても、なんらかの理由のために、その社会の正常な話し方をすぐに理解できないと考えられている人たち(例えば赤ん坊、外国人、耳の聞こえない人たち)に対し、人間は、単純化され、反復され、理想化された文を発話するとされている(Todd, 1974)(Pinker, 1994)。本モデルにおけるエージェントが持つ文法とはまさにこれらを想定しており、複雑な素性の共有、欠落は容易に予想される。したがって本モデルからの現実的な文法への拡張を考えても、追加された文法に応じて会話の頻度を増やすことにより、同様なピジン化が起ると予測される。

謝 辞

本研究を行う上で、北陸先端科学技術大学院大学情報科学研究科永田裕一助手から有益なご助言をいただきました。ここに感謝の意を表します。

文 献

- Abeillé, Anne. & Rambow, Owen. (2000). *TREE ADJOINING GRAMMARS*. CSLI Publications.
- Archibald, John. (2000). *Second Language Acquisition and Linguistic Theory*. Blackwell Publishers.
- Arends, Jacques., Muysken, Pieter. & Smith, Norval. (1994). *PIDGINS and CREOLES an introduction*. John Benjamins B.V.
- Bickerton, Derek. (1990). *Language and Species*. The University of Chicago Press: (寛壽雄 訳 (1998). 『ことばの進化論』, 勁草書房)
- Bruce, B. (1975). Case Systems for Natural Language, *Artificial Intelligence*, **13** 327-360.
- Chomsky, N. (1981). *Lectures on Government and Binding*. Dordrecht: Foris.
- Crystal, David. (1987). *The Cambridge encyclopedia of language*. (風間 喜代三・長谷川 欣祐 監訳 (1992). 『言語学百科辞典』, 大修館書店)
- Hashimoto, T. & Ikegami, T. (1995). Evolution of Symbolic Grammar Systems. *Third European Conference on Artificial Life*, 812-823, Springer
- 伊庭 斉志 (1994). 『遺伝的アルゴリズムの基礎』, オーム社
- 掛川 淳一・神田 久幸・藤岡 英太郎・伊丹 誠・伊藤 紘二 (2000). 日本語学習支援システムにおける作文診断処理系の提案と試作. 『電子情報通信学会誌論文誌』, Vol.J83-D-I No.6, 693-701.
- 亀井 孝・河野 六郎・千野 栄一 (1996). 『言語学大辞典』, 三省堂
- 長尾 真 (編) (1996). 『自然言語処理』, 岩波書店
- Nirenburg Sergei. (1987). *Machine translation*. CUP.
- 小野 哲雄・東条 敏 (1998). 推論機能を有するエージェント群による共通文法の組織化. 『人工知能学会誌』, **13** (4), 546-559.
- Ono, Tetsuo. Tojo, Satoshi. & Sato, Satoshi. (1996). Common Language Acquisition by Multi-Agents. *International Computer Symposium (ICS'96), Proceedings on Artificial Intelligence*, 218-223.
- Pinker, Steven. (1994). *THE LANGUAGE INSTINCT*. William Morrow & Company.
- 竹沢 幸一・Whitman, John. (1998). 『格と語順と統語構造』, 研究社出版
- 田中 穂積・辻井 潤一 (編) (1988). 『自然言語理解』, オーム社
- Todd, Loreto. (1974). *Pidgins and Creoles*. RKP: (田中 幸子 訳 (1986). 『ピジン・クレオール入

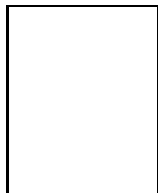
門』, 大修館書店.).

東条 敏 (1988). 『自然言語処理入門』, 近代科学社
Werner, G.M. & Dyer, M.G. (1991) Evolution of
Communication in Artificial Organisms.. in
C.G.Langton (Eds.), *Artificial Life II*, 659–
687 : Addison Wesley.

XTAG Research Group. (2001) A Lexicalized Tree
Adjoining Grammar for English. IRCS, Uni-
versity of Pennsylvania, IRCS-01-03

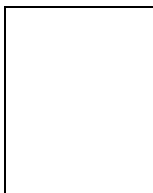
(2002 年 1 月 8 日受付)

(2003 年 1 月 19 日採録)



中村 誠

1972 年生. 1997 年北陸先端科学
技術大学院大学情報科学研究科博
士前期課程修了. 現在北陸先端科
学技術大学院大学情報科学研究科
博士後期課程に在学中. 自然言語
処理, 人工知能などの研究に従事.



東条 敏 (正会員)

1981 年東京大学工学部計数工学
科卒業. 1983 年東京大学大学院工
学系研究科修了. 同年 (株)三菱総
合研究所入社. 1986-1988 年米国
カーネギーメロン大学機会翻訳セ
ンター客員研究員. 1995 年北陸先
端科学技術大学院大学情報科学研究科助教授. 2000
年同教授. 1997-1998 年ドイツ, シュトゥットガルト
大学客員研究員. 博士 (工学). 自然言語理解,
状況推論の研究に従事. 情報処理学会, 日本ソフト
ウェア科学会, 言語処理学会, 日本認知科学会, ACL,
FoLLi 各会員.