

Title	言語進化シミュレーションにおけるピジンとクレオール の創発に関する研究
Author(s)	中村, 誠
Citation	
Issue Date	2004-03
Type	Thesis or Dissertation
Text version	author
URL	http://hdl.handle.net/10119/948
Rights	
Description	Supervisor:東条 敏, 情報科学研究科, 博士

博 士 論 文

**言語進化シミュレーションにおけるピジンとクレオール
の創発に関する研究**

指導教官 東条 敏 教授

北陸先端科学技術大学院大学
情報科学研究科情報処理学専攻

中村 誠

2004 年 1 月 8 日

要 旨

本論文では、言語進化シミュレーションにおけるピジンとクレオールの創発について述べる。本研究の目的は、ピジンやクレオールといった言語現象を、言語進化の理論に基づいた計算機シミュレーションによって再現し、これらが創発するための条件を導き出すことである。

人間は、非常に豊かな語彙や統語構造の規則から作り出される言語によって、話者が意図することを表現することができる。同時に、等質なひとつの言語を用いることによって、聞き手はその発話内容を正確に理解することができる。言語、すなわちその文法の獲得は、子供の言語獲得の臨界期において、親などの発話を聞くことによって行われるが、その言語獲得のメカニズムは未だ明らかにされていない。ピジンやクレオールといった社会環境の変化によって起こる急激な言語変化は、この言語獲得のメカニズムに大きく関係していると考えられている。特にクレオールに見られるいくつかの特徴は、人間の言語の生得性を裏付けるいくつかの証拠を示している。このように、ピジンやクレオールの研究は、言語獲得を解明する上で重大な役割を果たしている。

本研究においては、実際のピジンやクレオールが創発する環境に倣い、複数の言語が使用される特殊なコミュニティを仮定し、言語獲得と言語話者の人口変化の関係を調査した。ピジンとクレオールは、言語の連続的な変化の段階であるが、それぞれ別の要因で発現することから、ピジン化、クレオール化に関する実験をそれぞれ独立して行った。まず、LTAG と GA を組み合わせた文法獲得機構を提案し、マルチエージェントによってピジンの創発現象を再現した。次に、普遍文法を仮定し、数理生態学的な理論である言語動力学を修正することによって、より現実に近いモデルを提案し、クレオールの創発を観察した。実験の結果から、クレオールの創発が言語話者の人口構成比と、子供が複数の言語に接触する頻度に依存することを示した。また、クレオールの創発は、言語間の類似性にも依存し、クレオールになるために必要な言語の条件を導き出した。これら一連の実験は、現実世界においてピジンやクレオールの発生を予測するモデルとして言語学の分野への大きな貢献であると考えられる。

目次

1	はじめに	1
1.1	研究の背景と目的	1
1.2	本論文の構成	3
2	ピジン・クレオールに対する言語学および計算機科学的アプローチ	7
2.1	ピジン, クレオール研究からの言語獲得の解明	8
2.1.1	ピジン化とクレオール化のプロセス	8
2.1.2	ピジンとクレ奥ールの文法	11
2.1.3	クレオールと言語獲得の生得性との関係	14
2.2	言語進化研究からの言語獲得の解明	17
2.2.1	計算機シミュレーションによる利点	18
2.2.2	これまでの言語進化研究の流れ	20
2.2.3	言語進化研究における本研究の位置づけ	23
2.3	ピジン・クレオール創発のモデル化	24
2.4	第 2 章のまとめ	26
3	マルチエージェント環境での人工ピジンの生成	27
3.1	人工ピジン生成モデルの提案	28
3.1.1	中間表現の定義	29
3.1.2	文法の定義	30
3.1.3	エージェント間通信	32
3.1.4	学習機構	33
3.2	実験および考察	39
3.2.1	実験	39
3.2.2	実験 1, 2 の結果と考察：優勢言語の学習	41

3.2.3	実験3の結果と考察：言語の混合	42
3.3	第3章のまとめ	48
4	マルチエージェントによる人工クレオール生成モデルとその問題点	50
4.1	クレオール生成モデル	51
4.1.1	人工クレオール生成モデルの提案	51
4.1.2	クレ奥ールの定義	52
4.1.3	実験と考察	53
4.1.4	エージェントの構成とクレ奥ールの生成に関する実験	54
4.1.5	クレオール化の条件に関する考察	56
4.1.6	文法と遷移の関係についての考察	58
4.2	マルチエージェントの問題点	59
5	言語動力学におけるクレオールの創発	62
5.1	文法獲得に関する人口動力学	63
5.1.1	Komarova のモデル	63
5.1.2	Niyogi のモデル	65
5.2	動的遷移行列モデル	69
5.2.1	モデルの改良	69
5.2.2	接触確率 α の導入	70
5.2.3	学習アルゴリズム	71
5.2.4	動的遷移行列 $\overline{Q}(t)$	73
5.3	実験1 - 動的遷移行列モデルの検証	74
5.3.1	人口動力学上のクレオールの定義	74
5.3.2	実験1の言語セットと S 行列	75
5.3.3	実験1-1 - 人口動力学上でのクレオールの創発	76
5.3.4	実験1-2 - 初期条件に見る優勢言語の領域	80
5.3.5	実験1の考察	82
5.4	実験2 - 類似性に関する条件の検証	82
5.4.1	実験2の言語セットと S 行列	82
5.4.2	実験2の初期条件とパラメータ	83

5.4.3	実験 2 - 優勢クレオールが創発する条件	84
5.4.4	実験 2 の考察	88
5.5	第 5 章のまとめ	90
6	考察	92
6.1	関連研究との比較・検討	92
6.1.1	ピジン・クレオール研究に対する考察	93
6.1.2	言語進化研究に対する考察	94
6.2	本研究の問題点と今後の課題	96
6.2.1	マルチエージェント・モデルの問題点と今後の課題	96
6.2.2	人口動力学モデルの問題点と今後の課題	97
7	おわりに	100
	謝辞	104
	参考文献	105
	本研究に関する発表論文	110
A	パラメータに対応する言語群	112

目 次

2.1	ピジン化・クレオール化のプロセス [19]	9
2.2	言語の進化に関わる 3 つの適応システム [22]	23
3.1	LTAG による日本語の文法表現	31
3.2	LTAG による英語の文法表現	31
3.3	エージェント間の会話の流れ	33
3.4	遺伝子の表現型を介した文法の変換	34
3.5	日本語文法から英文法への変換の例	36
3.6	会話モデルへの GA による学習の適用	37
3.7	実験 1 の結果	41
3.8	実験 2 の結果	42
3.9	実験 3 の結果	43
3.10	p_f 値の増加によるヒット率の変化	44
3.11	各世代における中間表現ごとの伝達状況	45
4.1	会話と学習の流れ	53
4.2	実験結果の一例	57
5.1	人口変化の流れ	64
5.2	マルコフ状態図 [32]	67
5.3	接触確率 α	71
5.4	単純な言語獲得アルゴリズムの導入	72
5.5	動的遷移行列モデルの結果 ($0 \leq \alpha \leq 0.5$)	77
5.6	動的遷移行列モデルの結果 ($0.5 < \alpha \leq 1$)	78
5.7	優勢文法の領域の出力	80
5.8	3 言語での優勢クレオールの創発	85

5.9 優勢クレオールとなるための条件 ($\theta_d = 0.9$)	86
5.10 言語空間上のクレオール創発の条件	89
5.11 言語の共存	90

表 目 次

2.1	ピジンとクレオールの比較	12
2.2	英語とハワイ・クレオールの構造の相違 [6]	14
2.3	子供の発話文とクレオールの比較 [6]	16
3.1	各実験のパラメータ	40
3.2	各世代における発話例	46
4.1	チョムスキー標準形で書かれたルール	52
4.2	文法の数に対するクレオールの数と比率	55
4.3	人口比率に依存した遷移行列 Q の例	59
A.1	言語群 (Gibson et al. [16], Niyogi [32] から引用)	113

第 1 章

はじめに

1.1 研究の背景と目的

社会言語学や発達言語心理学などの分野において、ピジンやクレオールといわれる言語現象について盛んに研究が行われている [3, 15, 40]. ピジンとは共通の言語を持っていないが、通商その他の目的で互いにコミュニケーションをとる必要がある人々の間に発達した伝達システムである. その語彙や文法は単純化されており、構造や言語使用の点ではるかに発達している土地の言語と並んで、補助言語として習得される. その後ピジンが発展し、ある共同体の母語となったものはクレオールと呼ばれ、元となったどの言語とも文法的に異なり、それ自身が文法的にも表現能力としても充実した土地の言語となる. 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、非常に似通った語彙体系や文法構造を持っていることが大きな特徴である [43, 46]. これらを含む言語変化は人間の言語獲得に深く関わり、言語学習者を取り巻く社会環境の急激な変化に対応するため、獲得する言語が変化した実例であると考えられる. ピジン話者のコミュニティからわずか一世代でクレオールが発生した例も報告されており、周りに獲得すべき言語がないときに、人間がもともと持っている生得的な言語が発現すると Bickerton [5] は主張している. このように、ピジンやクレオールを研究することは、人間の生得的な言語獲得のメカニズムの解明に直結するため、言語に関連するさまざまな分野から大きな関心を集めているのである.

本研究の目的は、ピジンが創発し、クレオールに至るまでの言語獲得能力と社会、言語環境についての条件を数理科学的に導出し、言語変化の過程を通時的な側面から一般化された形で定式化することである。これまで言語学的な側面からピジンやクレオールは定義されてきた。しかし、その元となった言語間の関係からもしくは各言語話者の人口構成比から、これらの条件が存在することは想像に難くないが、それを実際のピジンやクレオールから厳密に求めることは現実的に不可能であった。近年ではこれらピジンやクレオールを含む言語の変化や言語進化に関する問題を解決する試みとして、言語の構造や人間の学習機構を直接解析するようなアプローチではなく、言語学習者を取り巻く環境を含めたシステムを構築することで理解しようとする構成論的手法が注目されている [11]。すなわち、言語学習に関するある特定の項目に関して仮説を立て計算機に実装する場合、それによって言語を学習する一個体だけでなく、複数の個体と発話環境を含めたシステムを実装することにより、そこから発現するさまざまな事象を観察する。それを実際の言語現象と比較検証することによってその仮説が健全であることを主張する方法である。これにより、これまで言語学の世界で行われてきた実地調査からは得られなかった、ピジンやクレオールが創発するための人口構成比、言語間の類似性、周りの環境に関する条件を計算機上のシミュレーションによって導き出すことが可能となる。

言語の変化を扱った研究は近年盛んに報告されており、その多くは自律的で能動的なエージェントが互いに文を発話し、その発話文を認識することによって文法を学習するというものである [20, 35]。マルチエージェントによるシミュレーションは、各エージェントが持つ機能を自由に設定することが可能であるため、構成論的手法による仮説の検証が行われやすい [11]。しかし有限の数からなるエージェントによるシミュレーションからは、現実世界で発生する現象を一般化して結論づけることは困難であると考えられる。それに対し、Nowak et al.[34] は普遍文法を用いた数理生態学的な言語動力学（Language Dynamics Equations）を提案している。これは、個々のエージェントの能力に注目するのではなく、コミュニティ全体の言語話者人口の遷移を人口動力学（Population Dynamics）として扱っており、上記に挙げたマルチエージェントモデルとは大きく異なるアプローチである。

本研究では、実際にピジンやクレオールが創発する環境に倣い、複数の言語が

使用される特殊なコミュニティを仮定し、そこで子供が獲得する言語と各言語話者の人口変化の関係を調査する。これらの相互作用によって言語獲得に与える影響が言語変化の本質であり、これを促す環境を計算機上に実装することは、ピジン、クレオール研究に対して大きな貢献になると考えられる。本稿において、我々は次の2つの目的から、それぞれの実験を行った。ひとつは、異なる言語を母語とする言語集団の接触により、ピジン化を経てクレオールを獲得するに至るまでの過程を再現するモデルを計算機上に実装することである。このとき、言語の変化を通時的に捉え、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ことが必要である。もうひとつは、言語動力学についての提言である。Nowak et al. [34] の言語動力学は、クレ奥ールの調査に有用であると考えられるが、現実の状況と比較して非常に単純化されている。したがって、次の目的は、この言語動力学をより現実なものに修正することにより、クレオールが創発する過程を示すことである。また、そのモデルからクレオールが創発するための言語間の類似性に関する条件を導き出すことである。

1.2 本論文の構成

本論文の構成および各章の関係は以下の通りである。まず第2章では、ピジン、クレオール研究に関わる従来の言語研究を、言語学的、進化言語学的な2つの側面から概観することにより、本研究の立場を明らかにする。第3章では、言語接触における文法の変化に重点をおいたマルチエージェント・モデルを提案し、ピジン化の過程を観察する。ここでは、大きく2つの問題について論じている。ひとつは、既存の自然言語技術を用いることによって、計算機上でピジンやクレオールが創発するのかということである。もうひとつは、計算機上で発生させた混合言語に関して、何をもってそれがピジンであるとみなすのかという問題についてである。第4章では、第3章に引き続いて、マルチエージェント・モデルによるクレオール化の過程を観察することを目的としたモデルを提案する。これら2つのモデルは、自律的で能動的に振る舞うエージェントによる共通な文法を獲得する際の自己組成の過程をモデル化したものである。したがって、これらを通じて、異なる言語話者集団の接触に始まり、ピジン化を経てクレオール化に至るまでの

連続的な変化を観察することができる。第 5 章では、言語動力学をより現実的に改良したモデルを提案し、これによりクレオールの創発現象を観察する。さらにこのモデルから、クレオールが創発するための、言語話者の人口構成比と言語の類似性に関する条件を導き出すことを目的とした実験を行う。さらに第 6 章では、本稿で提案したモデルについての考察を行い、第 7 章では、本稿のまとめを行う。以下に、各章の概要を簡潔に述べる。

第 2 章 ピジン・クレオールに対する言語学および計算機科学的アプローチでは、これまでのピジン、クレオールに関わる研究を概観し、本稿の立場を明らかにする。この章では大きく 2 つの節に分けられる。ひとつがピジン、クレオールに関する言語学的な知見であり、もうひとつが計算機科学からの言語進化に関するこれまでの研究成果についてである。まず、ピジンやクレオールとは一体どのようなものなのか、これまでに言語学の世界で報告されてきた成果を背景に、ピジン、クレ奥ールの定義と、それぞれの文法的な特徴を解説する。また、なぜピジンやクレオールといった言語変化に関する研究が注目されているのか、人間言語の生得性に関する議論から、その重要性について述べる。次に、言語獲得問題に関する計算機シミュレーションの有効性について論ずる。言語進化に関する研究は、近年、人工生命の分野から発展を遂げて来た。すなわち、人間がその進化の過程において、どのようにして言語を話すようになったのか、どのような過程を経て複雑な文法を持つように発展していったのか、という進化の過程を研究する分野である。この言語進化について論じる上で、ピジン、クレオールはやはり重要なキーワードとして扱われているが、厳密にいうと、現在観察することができるピジンやクレオールは、人間の言語獲得能力の進化によって発現する現象ではない。これらが言語進化に関する研究となぜこれほど密接に関わるのか、また、これまでどのような研究成果があり、本稿において何を論じる必要があるのかをここで説明する。

第 3 章 マルチエージェント環境での人工ピジンの生成では、これまでの人工知能や自然言語処理で培われた技術を基に、マルチエージェントの枠組でピジン化の過程を再現する。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、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。

第4章 マルチエージェントによる人工クレオール生成モデルとその問題点 では、前章から引き続き、マルチエージェントの枠組でクレオール化の過程を再現することを目的とする。本モデルにおいて、言語学的な知見である原理とパラメータ理論を導入し、これに基づいてエージェントが持つことができる文法のセットと、クレオールとみなすことができる文法の定義を行った。本章の位置づけとして、マルチエージェントの枠組でピジン化を取り扱った前章の続きであると同時に、次章の人口動力学モデルにおける問題を提言するための、前段階の章であると捉えることができる。

第5章 言語動力学におけるクレオールの創発 では、人口動力学を用いた言語進化に関する理論である言語動力学を用いて、クレオールの創発現象を観察する。また、クレオールが創発するための条件を言語使用者の人口比率および言語の類似性から導き出す。マルチエージェントによる動力学を考えると、エージェントが持つ機能の精度に関して二極性が考えられる。一方では Briscoe [8] に代表される、エージェントに非常に洗練された文法構造と学習機構を持たせることである。この場合、いかに人間の言語獲得メカニズムに近づけることができるかという、メカニズムそのものに焦点が当てられる。もう一方では、Kirby [21] のモデルのように、その構造は単純ではあるが、実験結果から得られることから言語獲得に関するある特定の機能について言及する手法である。第3章、第4章で述べたモデルは、既存の技術を組合せることによって人工ピジンおよび人工クレオールを生成するという目的は達成したが、人間の学習機構との等価性や、その後のピジンの拡張やクレオールへの遷移など、モデルの応用性に限界が感じられ、この位置付けに関して大きな疑問が生じるという結論に達した。この問題を省みて、第5章では構成論的手法によって新たなモデルを提案し、クレオールの創発について述べる。言語使用人口と子供の言語獲得に与える影響に関して議論する場合、人口動力学は、これらの相互作用を数学的な理論として導き出されたモデルであると考えることができる。ここでは、この人口動力学に基づいて提唱された言語

動力学について，第 4 章で得られた結論をもとに問題を指摘し，より現実的なものに修正することから始め，そのモデルを用いた実験を行う．

第 6 章 考察 では，本稿で述べた計算機実験の結果を踏まえ，本研究の方法論や位置付けに関する考察を行う．特に，第 2 章で取り上げた従来の研究との相違や，言語学的な研究分野への貢献について述べる．また，本研究の結果からは言及できない問題については，今後の研究課題として述べる．最後に，第 7 章では，本論文のまとめを行う．

第 2 章

ピジン・クレオールに対する言語学的 および計算機科学的アプローチ

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) [3] がそれである。これらは社会的な理由から起こった言語接触によって発生した混合言語であるが、これらを文法的に未発達で、その基盤となった英語などの優勢言語とは異なることから、「墮落した混合語」として片付けてしまうことは近視眼的である [40]。社会言語学の分野では、言語獲得の生得性とクレオールとの関係について調査が行われており、これらの研究は人間の言語獲得メカニズムを解明するアプローチとして注目を浴びている [5]。本研究ではこれを受けて、子供の言語獲得とクレオールの創発との関係を計算機科学の立場から解明することを目的とする。近年の計算機の発達により、シミュレーションによる言語の起源と進化に対する研究が活発に行われており [10]、ピジンやクレオールの研究は、この分野においても注目されている。本章では、これまでの言語変化や言語進化に関する研究成果を紹介するとともに、本研究との関連を述べる。ここから、計算機科学からのアプローチとして本研究に求められることと、本研究の立場を明らかにする。

2.1 ピジン、クレオール研究からの言語獲得の解明

言語学者のあいだで、かつて好奇心の対象でしかなかったピジンとクレオールは、今や言語の発達に関する研究分野の中心的な存在となっている [14]。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。ピジンはそのもとに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。ピジンとクレオールは、異なる言語話者集団の接触によって、言語が変化していく過程における連続的な段階として定義される。本節では、この言語接触のプロセスから、ピジン、クレオールがどのように定義されているかについて述べる。また、実際の発話例を比較することにより、その違いを示す。次に、ピジン、クレオールに関して議論されている、子供の言語獲得との関連について述べ、本研究に求められること、言語学の分野にどの点において貢献できるのかについて述べる。

2.1.1 ピジン化とクレオール化のプロセス

ピジンやクレオールは言語変化における通時的な変化の段階として定義される [40]。本節では図 2.1 に見られるような、ピジンの発生からクレオール化、また、その後の脱クレオール化に至る過程を各段階ごとに詳しく見て行くことにする。

ピジン化は、まず複数からなる言語の接触から始まる。ひとつの上層言語を話す少数の人を中心に、ひとつもしくは複数の土着言語や、他の言語を母語として持つ人々の間で言語の混合が発生する。商取引、奴隷売買、プランテーションにおける年季労働などの理由から、このような状況が作りだされるケースがほとんどである。彼らは互いにコミュニケーションをとる必要があるにもかかわらず、共通の言語を持っていないため、必然的に見ぶり手ぶりを交えた簡単なことばのやりとりでしか方法がない。このような、接触言語形成の最も初期の状態、まだ言語コードとしての安定性、規範性が確立していない段階を混合語もしくは**ジャーゴン (jargon)** と呼ぶ。ジャーゴンは、その定義から言語獲得の最終目的ではなく、そこに関与する人全員にとって、不明瞭な言語である [3]。アメリカ北

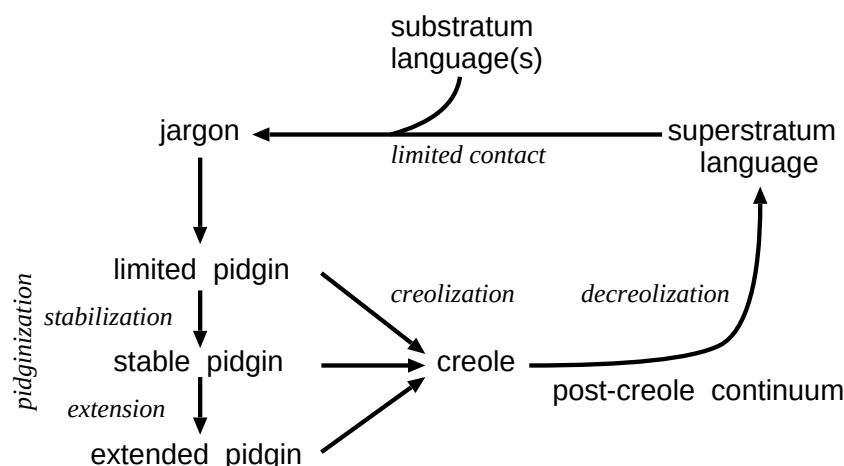


図 2.1: ピジン化・クレオール化のプロセス [19]

西海岸の有力な交易言語，広域共通語として発達したチヌーク・ジャーゴンなどが有名であるが，その名称に関わらず，実質的には後に解説する安定ピジンの段階に達していたと考えられる [46]．

その後，構造的な規則性を持った言語が開発され，その地域において広く使われるようになる．この構造的に単純化された言語がピジンである．ジャーゴンとピジンの違いは，主としてそのコードとしての安定性と語形成における規則性にある．このうち，ごくわずかな取り引きといったような最低限の接触の結果生まれたものは**限定ピジン (limited pidgin)** と呼ばれ，これを生み出した接触の場がなくなると，すぐ消滅する傾向がある [40]．例として朝鮮戦争時に韓国人とアメリカ人の間においてごく狭い範囲に限って使われた，韓国のバンブー・イングリッシュ (Korean Bamboo English) や，ベトナム戦争時のベトナム・ピジンなどがあげられる．

安定ピジン (stable pidgin) の段階になると，ピジン語の使用者と使用環境は異なる言語話者間のみにとどまらず，社会生活基盤の言語として使用されるようになる．このとき，語彙は拡張され，ある一定の規則で造語が生産されたり，文法緻密化のためのルールが作り出される．例えば，音節反復による形容詞の強調，名詞を動詞化する接辞の定着，複合語形成のパターンや，複数形の派生，前置詞の意味拡張，否定や疑問の形式の確立などである．したがって，**ピジン化 (pidginization)** の本質は，言語体系の縮小 (reduction) や劣化 (corruption, impoverishment) と

いう側面において把握するよりも、構造再編（restructuring）の過程として理解する必要がある。

安定化したピジンを話すことが、その使用者たち、特に地元住民にとって、生活上、なんらかの利益をもたらす場合、使用場面や使用地域が広がることになる。使用場面が増大することによって、語彙が拡張され、文法構造も複雑さを増し、また、使用者も急増する。そのような段階に至ったものを、**拡張ピジン（extended pidgin）**といい、パプア・ニューギニアで話されている tok・ピシン（Tok Pisin）などが有名である。拡張ピジンが広域共通語としての地位を占めるようになると、一方で、ほぼ必然的に、地域的な変種（dialects）、社会階層や使用目的等に応じた変種（sociolects）、また、場合によっては、母語を異にする民族集団ごとの変種（ethnolects）、などが分化するようになる。広域で使用される拡張ピジンは、都市部などで母語化が始まると同時に、一方では、ジャーゴンないし初期ピジンとして遠隔地に浸透していく。つまり、「ジャーゴン → ピジン → クレオール」という通時的連続相が、共時的に展開されることになる。

拡張ピジンが日常生活で使用されるようになると、ピジン話者である子供が第一言語としてそのピジンを獲得する、**クレオール化（creolization）**がおこる。すなわち、社会的にはその言語を初めて母語として使用する集団の登場であり、構造的に完成度の高い文法の獲得と語彙の全面的な拡張を意味する。前述の tok・ピシンなどは、都市部でクレオールとして使用されている。実際には、不安定なジャーゴンの段階から安定ピジンの段階を経ず、ハワイ・クレオールのように一気にクレオール化が起こることもある。また、トレス海峡クレオールのように、拡張ピジンの段階を経ずにクレオールに至る場合もある。

いったんクレオール化した言語が、語彙供給言語と継続的に接触あるいは再接触すると、一般に**脱クレオール化（decreolization）**と呼ばれる過程が進行する。接触、併存する語彙供給言語は、大抵の場合、社会的な権威と強制力を伴う上層語（公用語、国家語、宗主語、教養語）であるがゆえに、語彙形式上の連続性を持ったクレオールを同化、吸収する求心力を発揮しやすい。しかし、上層語である語彙供給言語との併存は、ただちにクレ奥ールの消失という事態をまねくものではなく、むしろ、**クレオール形成後の連続相（post-creole continuum）**と呼ばれる状態になるケースが多い。ハワイの英語クレオール、ハイチのフランス語

クレオールなどでは、脱クレオール化がかなり進行しており、これらは衰退する傾向にある。

本節では、ピジンとクレオールの定義を言語接触のプロセスに沿って述べた。ピジンやクレオールは世界中で発見されており、異なるピジンの間、また、異なるクレオールの間で文法的に共通な特徴を持っていることが注目されている。しかし、これらの変化過程は、接触状況の違いなどからピジンで消滅してしまうもの、安定ピジンとして残るものから、クレオール化に至るものまでさまざまである。また、チヌーク・ジャーゴンのように、ジャーゴンと名付けられていても実質的にはピジンと同等の文法構造を持っていたり、トク・ピシンのように、ピジンであるのに実際にはクレオール化が進行しているように、厳密な定義は難しいとされている。

2.1.2 ピジンとクレオールの文法

ここでは、ハワイで実際に使用されたピジンとクレオールについて、その経緯を紹介するとともに、ピジンとクレオールの具体的な発話例を基に、文法的な違いを見ていく。

19世紀末から20世紀初頭にかけて、ハワイの砂糖プランテーションの発展とともに、中国、フィリピン、日本、朝鮮、ポルトガル、プエルトリコなどから契約労働者が集められた [6]。このときの上層言語は英語であり、ハワイ語を始めとする、労働者がもともと使用していた言語は下層言語となる。労働者は、英語話者であるプランテーションのオーナーとコミュニケーションをとるため、英語を習得する必要があった。しかし、目標言語である英語の話し手よりも、潜在的学習者の数の方がはるかに多く、その比率は1対30にも及んだ。また、そこは厳格に階層化された社会であり、労働者とオーナーとの間の接触は最小限に限られていたことから、目標言語の習得はほぼ不可能であった。しかし、これらの移民者たちは共通の言葉を持っていなかったにもかかわらず、互いに意志疎通を図る必要があり、その結果、ハワイ・ピジンが発生したのである。その後、ピジン・コミュニティの労働者に子供が生まれると、その子供たちが接する言語のほとんどがピジンであった。英語話者との接触は限られており、両親の間で異なる母語を持つ

表 2.1: ピジンとクレオールと比較

	ピジン	ハワイ・クレオール	英語
(1)	Now days, ah, house, ah, inside, washi clothes machine get, no? Before time, ah, no more, see? And then pipe no more, water pipe no more.	Those days bin get [there were] no more washing machine, no more pipe water like get [there is] inside house nowadays, ah?	In those days there were no washing machines and no piped water like there is in houses nowadays.
(2)	Good, this one. Kaukau [food] any kind this one. Pilipin island no good. No more money.	Hawaii more better than Philip-pines, over here get [there is] plenty kaukau [food], over there no can, bra [brother], you no more money for buy kaukau [food], 'a'swhy [that's why].	Hawaii is better than the Philippines. Here there is plenty of food but over there is none. You have no money to buy food, that's why.

家庭が多く、その場合、単一の言語で話されることはなかった。また、両親が長時間に渡ってプランテーションで労働をしているあいだ、子供たちは数人の年配の女性によって世話をされていたようである。そのころになると、労働者たちが話すピジンは、単語の混合と構造の簡素化によって広まっており、クレオールが創発したときのピジンは、言語話者の母語に依存せず著しく似通った文法拡張がなされていた。しかし、それでも子供たちはクレオールを話すようになっていったのである [8]。

以降、実際のピジンとクレオールの発話例をあげ、それを基にそれぞれの文法を説明し、比較を行う。

表 2.1 に同じ意味を表す発話例をピジンとクレオールで示し、それぞれの文の構造を比較する。例文中の角括弧（[]）の中は、語や句の英語訳である。例文 (1) は「当時は、今家にあるような洗濯機も水道もなかった」、例文 (2) は「ハワイはフィリピンよりいい。ここにはたくさんの食べ物があるが、むこうにない。買う金がないからだ」という意味をそれぞれハワイ・ピジンとハワイ・クレオールで記述したものである。例文中の “washi” や “Pilipine” は、それぞれ “wash” と “Philippine” に対応するが、これらは誤記ではなく、発話者の発音に対応した綴りである。これら 2 種類の文は同じ意味を表しており、それぞれの構造の違いがわかる。移住者だけが話すピジンは、語順や使用単語などが話者ごとに大きく異なっているのが大きな特徴である。

ピジンとクレオールを明確に分ける基準のひとつとして、ピジン話者の発話を

理解するために、不明瞭な項目を聞き手が文脈や発話者の背景から補完する必要があるのに対し、クレオールに関しては構文規則から曖昧さを生じることなく空要素を同定できるということが挙げられる。したがって、英語やクレオールで言える内容はすべて、ピジンでも言うことができるにも関わらず、これらの例文を見てわかるようにピジンの構造は非常に未発達である。まず最初に、語彙に関しては多くの場合ひとつの言語が語彙供給言語となるが [3]、ハワイ・ピジンのような言語接触の初期の段階においては、混在ジャーゴン (macaronic jargon) と呼ばれ、さまざまな言語から単語が用いられる。実際には、上層言語である英語がほとんどで、あとは使用人口の最も多い言語であるハワイ語と、話し手の母語から借用される [15]。例文 (2) における kaukau は中国語系ピジンの “chowchow” から来ており、「食べ物」を表す一般的な単語である。次に、統語構造に注目すると、ピジンには冠詞、前置詞、補文標識はなく、文法要素によって表される構造は見られない。ピジンとクレオールの、文法面での端的な差異は、動詞の体系において最もよく現れる [46]。とりわけ限定ピジンの動詞からは、時制 (tense)、法 (mood)、相 (aspect) (これらはまとめて TMA と呼ばれる) を明示する要素がことごとく欠落しているのが大きな特徴とされる。例文 (1) のピジンの発話例においては、now days や before time などによって時間関係を明示してはいるものの、動詞の活用は存在しない。つまり、古い情報を文の始めに、新しい情報を終わりに置く、というだけの原則に基づいた、名詞、動詞、形容詞の単なる連続に過ぎない。それに対し、クレオールは移住者の子供たちの間からのみ発生し、ピジンに比べ、はるかに豊かな文法構造を持っている。例えば、ピジンにはなかった TMA を明示する体系が確立していること、埋め込み文による表現、空要素の同定などであり、クレオールは完全な言語 (full-fledged language) のための基準の全てを満たしているのである [5]。

ピジンが発話者によって語彙、語順などの文法的特徴が大きく異なっているのに対し、クレオールは話者間で一貫した、等質的なひとつの言語である。この等質性は優勢言語であった英語から由来するものではなく、クレオールの構造は、クレオール話者が接触したであろうと考えられるどの言語の構造とも大きく異なっている。英語とクレオールの構造の相違を表 2.2 に示す。同じ意味を表すハワイ・クレオールと英語の文の相違は、ハワイ・クレオールの文法が英語から取り入れ

表 2.2: 英語とハワイ・クレオール of 構造の相違 [6]

	英語	ハワイ・クレオール	英語の意味
(1)	The two of us had a hard time raising dogs.	Us two bin get hard time raising dog.	私たち 2 人には犬を飼うのは大変だった.
(2)	John and his friends are stealing the food.	John-them stay cockroach the kaukau.	ジョンと彼の友達は何物も盗んでいる.
(3)	If I had a car, I would drive home.	If I bin get car, I go drive home.	もし車があれば, 車で家に帰るのに.

た文法に基づいたものではないことを示している. 例文 (1), (2) のように, ハワイ・クレ奥ールの動詞の過去時制は, 接尾辞 “-ed” の代わりに, 主動詞の前に不変化詞 “bin” をおいて表し, 非瞬間相または進行相は接尾辞 “-ing” の代わりに, 語 “stay” をおいて表す. また, 例文 (1) の “dog” のように, 名詞の単数か複数を明示しない. 例文 (2) においては, “cockroach” が動詞として使われていたり, ピジンでもあった “kaukau” が使用されている. 例文 (3) は, 法 (mood) を表す助動詞の使用例である. 非現実世界を “go” という不変化詞によって表し, これを主動詞の前に置いて使用する. また, クレオールの構造は, ここで示した英語だけでなく, 中国語, ハワイ語, 日本語, 朝鮮語, ポルトガル語, スペイン語, フィリピン諸語のそれとも異なる.

ここでは, 具体的な発話例を基に, ピジンとクレオールの文法を比較した. クレオールの文法は, 他の諸言語と同じような複雑さを持っているにも関わらず, そのうちのどの言語とも似ていない. この特徴は, 他のクレオール諸語にも見られるうえに, これらクレオール諸語間の文法は, 非常に似たものとなっている. ここから, クレオールと言語の生得性についての関係が議論されている. 2.1.3 節で, クレオールと普遍文法との関係について述べる.

2.1.3 クレオールと言語獲得の生得性との関係

第一言語を獲得する際, その学習期間は限られており [26], 学習者はそのコミュニティで話されている発話を聞き, そこからもっともらしい文法を身につけなければならない. このとき, 有限の文から文法ルールを完全に特定することは不可能

であることが数学的に示されているにも関わらず [17], それでも同じコミュニティで育った子供は潜在的な文法のルールを正しく推論し, 同じ言語を矛盾なく獲得するのである. 文法獲得におけるこの困難性 (プラトン問題 [12]) を解決する考え方として, 原理とパラメータ理論が提唱されている [13]. 原理とパラメータ理論とは, 普遍文法はすべての人間言語に共通な原則の体系, すなわち原理 (principle) とそれに付随するパラメータ (parameter) からなると仮定し, 子供の文法獲得は普遍文法の原理に組み込まれたパラメータの値を言語経験により固定する過程と捉える考え方である [56]. 近年ではこの普遍文法がクレオール発生の創発と大きく関係していると考えられている [15].

2.1.2 節において, ハワイ・クレオールと上層言語である英語との構造的な相違を示し, また, 他の接触言語の文法とも異なることを述べた. このような文法の相違は, 世界中の他の地域で生じたクレオール諸語にも見られ, さらに, これらには, ハワイにおいて観察されるものと同じ等質性, そして同じ文法構造さえも見られるということが明らかになっている [6]. クレオール諸語はそれぞれのコミュニティの中で発生したものであり, 人間に共通な言語獲得に関する普遍的な能力が, この言語的類似点の原因となっていると考えられる. Bickerton はこのような現象について, 言語学習の臨界期において特定のモデルがない場合には, 言語を再創造すると主張しており, この仮説を **バイオプログラム仮説 (bioprogram hypothesis)** と呼んでいる [7]. このバイオプログラム仮説から, クレオール諸語が類似しているのは, 適用される生得的な言語能力は普遍的であり, クレオールの構造が単純であるのは, TMA などの, いかなる言語にとっても欠かせない特徴のみを備えた最小の体系についてバイオプログラムが作用するためであると考えられる [3]. クレオール諸語は, この仮説に基づいて, 他に例を見ないほど直接的な形で具現させた言語である [5]. パラメータの値は, 子供が周りの言語に遭遇することにより, つまり, 経験により変更される [51]. このとき, パラメータの値を決定する経験資料が入ってこない場合は, パラメータが最初に決まっている値, すなわちデフォルト値になり, この場合を無標値と呼ぶ. バイオプログラム仮説は, この無標値から導き出されるものとして, 普遍文法の考えに整合している [4].

クレオール諸語間の類似点と, これらの諸言語が互いに無関係に発生したとい

表 2.3: 子供の発話文とクレオールと比較 [6]

	子供の言語	英語系クレオール諸語	英語
(1)	Where I can put it?	Where I can put om? (Hawaii)	Where can I put it?
(2)	I go full Angela bucket.	I go full Angela bucket. (Guyana)	I shall fill up Angela's bucket.
(3)	Johnny big more than me.	Johnny big more than me. (Jamaica)	Johnny is bigger than I.
(4)	Nobody don't like me.	Nobody no like me. (Guyana)	Nobody likes me.

う可能性は、モデルとなるべき適当な母語が存在しない場合に、子供たちの間にクレオール諸語が発達することを示唆している。すなわち、クレオールは、ピジン話者の子供の第一言語獲得過程においてパラメータの値が定まらないままその学習期間を終えてしまい、無標値として固定したと考えることができる。クレオールが発生する地域ではないひとつの言語が使用されているコミュニティにおいて、子供が言語獲得の臨界期で行うことのひとつが、デフォルト値を示しているパラメータの値を経験によって変更することである。したがって、2歳から4歳くらいの、まだパラメータ値が定まっていない時期においては、子供の発話はクレオール話者のそれと非常に多くの類似点があげられる [5]。

英語を話す両親から生まれた子供の発話文と、英語に基づいたクレオールの発話文との比較を表 2.3 に示す。正しい文法を用いた英語の発話文と比較すると、これらの間には次のような類似点がある。

- 疑問文において、主語と助動詞が倒置しない（例文 (1)）。
- “go” という形が未来時制の標識として使われている（例文 (2)）。
- 形容詞 “full” が他動詞のように使われている（例文 (2)）。
- 形容詞の活用が省略されている（例文 (3)）。
- 否定の主語と否定の動詞が共存している（例文 (4)）

これらは双方の話者が持つ、パラメータがデフォルトの値を示しているために現れた現象であると考えられる。

ここではクレオールと普遍文法との関係について述べた。原理とパラメータ理論におけるパラメータについて、ピジン話者のコミュニティという特殊な環境下では、その値を決定するために必要な刺激が十分に得られず、クレオールは創発

する。パラメータの種類は、空主語パラメータ (*pro-drop parameter*)¹や、語順の決定に関与する主要部パラメータ (*head-parameter*) のように、いくつかのパラメータの候補があげられているが [51]、世界中の言語に見られる多様性を区別するための全てを満たすほどのパラメータは確定されておらず、また、それらを検証することは難しい。クレオールは、言語獲得の過程にある子供の言語と同様、パラメータの値にバイアスが掛けられた特殊な状態を表していると考えられていることから、普遍文法を研究するにあたり、クレオール研究は大変重要な位置にあるといえる。

2.2 言語進化研究からの言語獲得の解明

ピジン、クレオール研究の重要性を 2.1 節で述べた。本節ではこの取り組みに対し、計算機科学からのアプローチとして、何が求められるかということについて、これまでの研究成果を交えて述べる。

ピジンやクレオールは言語の変化過程の段階に過ぎないため、これまで簡単な辞書や文法書などは作成されて来たものの [44]、近年の自然言語処理技術による統計的な解析 [27, 45] を行うことを考えると、そのために必要となる十分な量のコーパスや音声のデータベースといった類は存在しない。Roberts [39] はハワイ・ピジンからハワイ・クレオールへの変遷を記録した大規模なデータベースを利用して、クレオール化の理論の確立に貢献しているが、ある共時態において使用された言語に注目した場合、そのデータ量はやはり十分ではない。言語学的な分野から求められていることは、これまでに提唱されてきた、ピジン、クレオールの創発に関するさまざまな理論の検証、および、環境と創発現象との因果関係を求めることである。したがって、ピジンやクレオールを静的な言語として共時的に扱うのではなく、むしろこれらの通時的な変化過程を動的に計算機上に表現する方法を考えなければならない。この変化過程を観察するために、シミュレーションによる実験が有効である。以降、計算機シミュレーションによる利点と、その具体的な手法を述べ、本研究でどのようなモデルを組むべきかについて論じてい

¹イタリア語のように、時制文の主語の位置に音形を持たない要素が現れることを許す言語と、英語のようにそれを許さない言語がある。このような言語の多様性を区別するパラメータが空主語パラメータである [56]。

くことにする。

2.2.1 計算機シミュレーションによる利点

ピジンやクレオールは移民者やその子供など、各個人がおかれている環境と、そこから生まれる必要性から創発した言語である。したがって、これらの言語の混合が発生した環境と、個人の言語学習能力等を計算機上に実装し、シミュレーションによる実験を行うことで、言語の変化過程を再現することが可能となる。計算機によるシミュレーションは、さまざまな実験を支援するツールであり、現実的な設定をすることや、これまで不可能であった実験を行うことが可能となる仮想的な実験室であると考えられる [10]。この計算機上の実験室において、一度システムを組むだけで、さまざまな条件の下で現象を観察することができ、またその現象を左右する条件やパラメータ変数を操作することが可能となる。したがって、実際にピジン化やクレオール化の過程についてシミュレーションを行うことは、理論の実証、ピジン化、クレオール化に関するさまざまな条件の特定をすることに等しい。

ピジン化の過程を例に挙げて説明する。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 [40]。しかしながら、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が [5]、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 [46]。現実のピジンが発生した環境を模倣し、これらのパラメータを変化させることにより、ピジンが創発するために必要な条件を特定することが可能である。

これらの条件の特定が、計算機シミュレーションに委ねられる理由は2つある。ひとつは、これまでに発見されてきたピジンやクレオールの種類は限られており、これらの実地調査をもとに仮説を提唱しても、その妥当性を検証する方法は限られていることである。例えば、ハワイ・クレオールはハワイの特殊な環境のもとで創発したが、もし上層言語である英語話者との接触機会が多くあったならば、ピジンやクレオールは創発しただろうか。もし英語話者との接触機会が増えるとクレオールが創発しないというのであれば、どの程度の接触頻度を境にしてクレオールが創発するケースとしないケースに分けられるだろうか。これを調査するために、実際にクレオールが創発したときと、他の地域においてクレオールが創発しなかったときの状況を比較する手段が考えられるが、優勢言語との接触を除いて、使用言語や人口構成比など、考えられる他の要素が全く同じ状況であった地域を見つけることは現実的に不可能である。計算機によるシミュレーションでは、この接触頻度をパラメータとして連続的に値を変えて繰り返し実験を行うことにより、その条件を特定することが可能である。

もうひとつは、ピジンやクレオールの創発を決定づけるであろうさまざまな要因は複雑に絡み合っており、これらの創発に関するある条件は、他の要因によって成り立たないことがある。例えば、異なる言語を話す複数の言語集団が接触したとき、その接触の度合いなどからピジン化に至ると考えられているが、このときピジン化ではなく、コイナー化と呼ばれる現象が起こることがある。コイナー化とは、接触する言語間の系譜的な距離が近い場合に起こる、単純化を伴う方言融合である [46]。ここから、集団間の人口構成比や接触頻度に加え、言語の類似性もピジンが創発するための条件に含まれ、これらが互いに影響を及ぼすことは想像に難くない。

これらの複雑性や、依存関係を解明するために、構成論的手法といわれる統合的な解析手法が提案されている。構成論的な取り組みは、複雑系、特に人工生命の分野においてよく用いられ [20]、従来の科学的方法である分析的なトップダウンによるアプローチとは全く異なる。構成論的手法は、構成要素からのボトムアッ

プによって統合的に問題を捉える手法である．まず最初にシステムの基本的な構成要素を決め，そのシステムによって作用するルールや環境といったものを設定する．次に計算機によるシミュレーションにより，構成要素のあいだに起こる相互作用によって系全体の振る舞いを観察し，これらの因果関係を求める方法である．シミュレーションの結果として得られるものは，システムの振る舞い，すなわち創発現象である．創発（emergence）とは，あるシステムにおける特定の機能や全体の秩序が，構成要素間の局所的な相互作用によって自発的に発生する現象として定義される [20, 55]．この結果を現実世界のものと比較することにより，システムを組むときに仮定した構成要素や相互作用のルール等に関する実行可能性や有効性を検証することが可能となる [10]．

ここでは，計算機によるピジン，クレオール研究に対するアプローチについて述べた．マルチエージェントなどによってモデルを設計し，構成論的手法によって構成要素と結果との因果関係を得ることが有効であると考えられる．本研究においてピジン化およびクレオール化の過程を再現させるモデルを計算機に実装し，これらが創発するために必要な条件を導出するにあたり，その要因が何であるのかを仮定する必要がある．2.2.2 節では，これまでの言語進化，言語変化を対象にした研究を紹介し，モデルを構築する際に必要とされる構成要素の詳細を述べる．

2.2.2 これまでの言語進化研究の流れ

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，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6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 [55]．Werner et al. [42] 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。ここで使用された学習機構は、ニューラルネットワークとGAを組み合わせたもので、シミュレーションの結果、人工生物の雄と雌の間に共通な言語を開発する能力が見られた。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 et al. [18] によるモデルでは、各々のエージェントが文法を持ち、他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、また複雑な文を話し、それを理解できるかということの評価の基準として学習を行っており、文法の進化によってこれらが満たされること主張している。この研究では、形式文法をあらかじめ明示的に持ったエージェント間の会話を繰り返すモデルを扱っている。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。

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 et al. [35]、小野ら [52] である。Ono et al. [35] 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。また小野ら [52] は、高度な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、生得的な言語獲得能力が大きく影響しない第二言語獲得に大きく分けられる [2]。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。

子供の第一言語獲得に関するモデルを考える場合、言語学習者として子供を、その学習対象として大人の言語をそれぞれ定義する。子供は学習期間において大人の言語に接することで第一言語の文法を獲得する。その後、ある世代における子供

が次世代の大人になり、繰り返し学習をすることで系全体の言語の振る舞いや、大人が獲得した文法を観察するという手法が一般的である [8, 21]. Kirby et al. [22] は特にこの学習法を繰り返し学習モデル (Iterated Learning Model; ILM) と呼んでいる.

近年においては、普遍文法が存在を仮定し、人間の言語獲得について直接的に言及したモデルが提案されている. その代表的なものが Briscoe [8] のモデルである. 各エージェントはカテゴリ文法における各カテゴリがノードとなる階層的なラティスと、語順を決定するためのパラメータを持つ. このモデルは、無標、動詞 (V)、目的語 (O)、および主語 (S) によって構成される言語について、語順の同定に関する問題を扱っており、自然言語に近い言語の進化を表現した. なお、このモデルからクレオール創発が観察されたことが報告されている.

また、Kirby [21] は、親から子供への対話を何世代にも渡って繰り返すことにより、文法構造の進化が必然的に起こったと主張している. このモデルにおいて、子供は言語を理解することに対する特別な報酬を与えられていないのに、コミュニケーションにおけるボトルネックを解消するために、自発的に構文の組成とルールの再帰的な適用を行うことを学習している. 人間が話すことができる言語は、原理上考え得る言語よりもその範囲は狭く、それを制限しているのが普遍文法であると考えられている. 従来の研究では、人間には普遍文法があらかじめ備わっているという仮定によって、そこから導出される言語を人間の言語とみなしてきたが、Kirby はこのモデルを通じて、人間の学習バイアスによって必然的にこのような言語が形成されたと主張している.

これまでに述べたモデルは、マルチエージェントの枠組で提案されている. しかし、言語進化における諸問題を数理的に議論する場合、有限数からなるエージェントによるモデルでは、一般性を論じることが困難であると考えられる. Nowak et al. [34], Komarova et al. [24] は、文法獲得に関して人口動力学を用いた数学モデルを提案した. この人口動力学モデルは原理とパラメータ理論の考え方に基づいている. すなわち、その原理によって与えられた文法の探索空間は有限であると仮定することにより、言語話者が用いる文法は $\{G_1, \dots, G_n\}$ のどれかひとつであるとして、考えられる言語の全てをあらかじめ定義する. すると言語の変化とは、言語話者が所有するパラメータの変化を示しており、その変化はそれぞれのパラメータ値

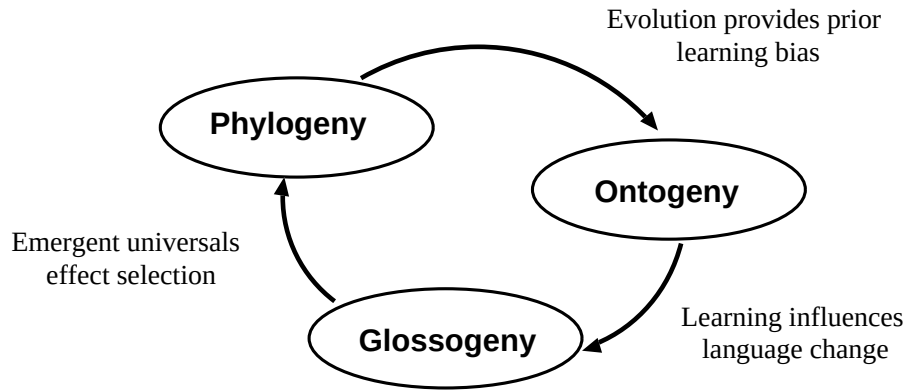


図 2.2: 言語の進化に関わる 3 つの適応システム [22]

に対応する言語間の人口遷移によって表現される．ここで，あるコミュニティにおいて文法 G_i を持つ言語話者の人口比率を x_i ，すなわち $\sum_{i=1}^n x_i = 1$ であるとすると，ある世代 t における文法 G_i の人口比率 x_i の変化は微分方程式を用いた動力学系として表される．このように人口動力学モデルは，各言語話者の人口比率の変化を追跡するものである．結果として，Komarova et al. は，言語学習者が持つ学習メカニズムと，言語進化の間に起こる均衡を導き出している．

2.2.3 言語進化研究における本研究の位置づけ

2.2.2 節において，これまでの言語進化の研究から代表的なモデルを紹介した．言語進化の研究は，主に言語の誕生と，文法が創発し，人間の言語が動物のコミュニケーションよりも複雑化した過程について議論される．それに対し，ピジンやクレオールは，人間がこれまで扱えなかった複雑な文法を獲得した結果として得られた言語ではないため，進化言語学の分野で主に研究対象とする言語とは異なっている．特にクレオールは，言語獲得の臨界期においてパラメータを決める際に圧力がかかった結果として得られる言語変化の現象であると考えられる．文法構造の複雑さに関する進化がないこれらの現象が，言語進化の分野でどのように捉えられるのか，その位置づけについて説明する．

言語は，次のような 3 つの複雑な適応システムの組合せによって進化すると考えられている（図 2.2 参照） [22]．

学習 (Learning) 個体発生論 (ontogeny) 的な考え方から，子供は他人の発話を

理解したり、理解できる発話を生成する能力を最大にすることによって、言語の知識を環境に適応させる。

文化的進化 (Cultural evolution) 言語に関する長く歴史的な (glossogeny) タイムスケールの上で、言語は変化している。使用される単語の移り変わり、意味の変遷や音韻や構文のルールはその都度順応する。

生物的進化 (Biological evolution) 人間が言語を身につけるために必要とされる学習メカニズムは、環境からの淘汰圧に応じて適応することにより、系統発生 (phylogeny) する。

学習に必要な能力は、生物的進化によって与えられる。文化的進化を通じて存続し得る言語とは、その時点での人間が持つ言語能力で学習可能な言語である。そして、人間の言語獲得能力が進化する際、より複雑な構造を表現する言語を話すことができるように、または、これまでと同程度の複雑な構造を持った言語を、より簡単に話すことができるように淘汰される。このことから、これらの3つは互いに影響しあっていると考えられる。ピジンやクレオール創発は、非常に短いタイムスケールで行われる言語進化である。すなわち、生物的な進化は伴わないが、言語話者の急激な発話環境の変化は、文化的進化によって単語の混合や、構造の簡略化が促され、ピジンとなる。また、クレオールは、そのピジン環境において子供の言語獲得能力を最大にした結果得られた言語であると考えられる。これらの相互作用を実際に観察できるものはピジンとクレオールに限られるため、言語進化の枠組でこれらを取り扱うことは、実地検証との比較ということに関して大きな優位性を持っている。また、英語などの主要な言語にも見られる、長い時間をかけた言語の変化と比較して、その変化の要因が特定しやすい。そのため、ピジン、クレオール研究は、言語進化研究においても重要な研究課題とされている。

2.3 ピジン・クレオール創発のモデル化

2.1 節では言語学の分野で問題とされていること、2.2 節では計算機科学の分野における課題をそれぞれ述べ、ピジンやクレオールを計算機上で創発させることの重要性を説明した。本節では、本研究の実験計画について述べる。

本研究では、次の実験を行う。

- 動的な文法表現と環境に応じたその柔軟な変化を起こす学習メカニズムを持ったエージェントによる、異言語集団の接触と、限定ピジンの発生（3章）からクレオール化（4章）にかけての観察
- 普遍文法を仮定したモデルによる、クレオール化の観察と、その条件の特定（4章，5章）

まず最初に、計算機シミュレーションでピジンのような言語混合が発生するかどうかを確認する。2.2.2 節で説明した Briscoe [8] のモデルが報告されるまで、ピジン、クレオールを対象とした言語進化研究はあまりされていなかった。したがって、言語の混合の過程を模倣するシステムをどのように組むのが課題となる。ここでは、LTAG で記述された文法と、その各ルールを改編するために作用する関数群を与え、その組合せを GA によって学習するモデルを提案する。エージェントが持つ文法の変化の様子と、集団間で広まる言語を観察し、実際のピジンの発話と比較することにより、文法と学習メカニズムの検証を行う。

次に、クレオールの創発を観察する。2.1.3 節で述べたように、クレオールの創発は、普遍文法が大きく関係していると考えられる。したがって、これを仮定し、クレオールが創発するための条件を求める。まずピジン化のモデルと同様、マルチエージェントの枠組でクレオール化のモデルを提案する。このモデルは、原理とパラメータ理論を仮定しており、各エージェントが普遍文法を持っている点でピジン化のモデルとは大きく異なる。原理として与えられる文脈自由で記述された文法群を用いてエージェント間で会話をを行い、EM アルゴリズムの一種である Inside-Outside アルゴリズムを用いた学習によってパラメータ推定を行う。

最後に、既存の人口動力学モデルをより現実的なモデルに修正し、構成論的アプローチによって本モデルの妥当性を検証する。実験を行うにあたり、計算機上で発生するピジンとクレオールの定義、ピジンやクレオールの創発の条件を特定するパラメータ変数等を定義する必要があるが、これらは各章で説明する。

2.4 第 2 章のまとめ

2.1 節ではこれまでのピジン、クレオール研究について述べた。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ことを示した。2.1.1 節ではピジンとクレオールが発生する過程を示し、その文法的特徴を 2.1.2 節で述べた。また、クレ奥ールの文法が生得的な普遍文法に基づいて形成されることを 2.1.3 節で述べたが、クレ奥ールの創発には、ピジンが形成される土壌が必要である。すなわち、ピジン化とは、目的言語が不明瞭なまま [5]、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ピジンを元にした第一言語獲得問題であると考えられる。

2.2 節では、計算機科学の側面からどのようにこの問題に取り組むべきかという問題について述べ、特に言語進化研究からもピジン、クレ奥ールの研究が重要であることを示した。2.2.1 節では、ピジン、クレ奥ールの研究を、計算機シミュレーションで解析することについて、どのようなメリットがあるのかについて述べ、構成論的手法の有効性を示した。2.2.2 節では、これまでの言語進化研究の流れを示し、モデルを構築する際の指標を与えた。また、2.2.3 節では、この研究を行うことによる言語進化研究に対する貢献について言及した。

2.1 節と 2.2 節を通じて、次のことがいえる。

- 言語変化の本質は、子供の言語獲得と、それを含めた周辺環境の変化によって言語獲得に与える影響との相互作用である。
- 言語獲得と言語変化の関係を調査するにあたり、実際に発生した現象と比較検証できる研究対象はピジンとクレ奥ールの創発についてである。

これらが、計算機上にピジン、クレ奥ールのモデルを実装するに至った理由である。

最後に 2.3 節では、本研究の実験計画を示した。本研究において求められることは、ピジン、クレオールが創発するための、あらゆる面における条件を求めることである。

第 3 章

マルチエージェント環境での人工ピジンの生成

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、したが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 [48]。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 [52]、周りの環境の変化から柔軟に対応しう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。

本章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。本モデル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、ピジンが発生する過程をシミュレートする [49]。

以降、3.1 節は本章で提案するモデルについて述べる。3.2 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 3.3 節

においてまとめを行う。

3.1 人工ピ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、ピジン化，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，本モデル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．

本モデル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ーとして，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，話すことのできるだけの，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。

3.1.1 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。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。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。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ける研究成果 [31, 53, 54] 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 [13] 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 [9] 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，一つの間表現を形成する。

3.1.2 文法の定義

3.1.1 節において，間表現を定義した。ここではこの間表現からの文生成および，文からのパースによる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，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，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) [1] 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，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逐語的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 3.1 参照）。これにより，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。

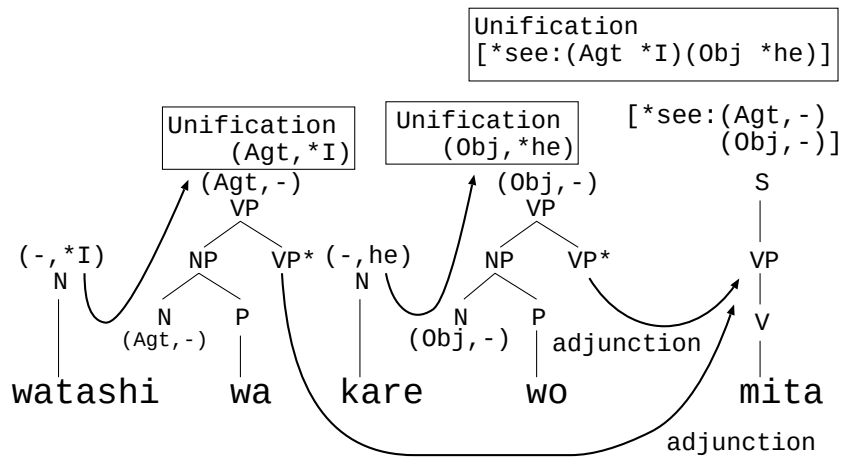


図 3.1: LTAG による日本語の文法表現

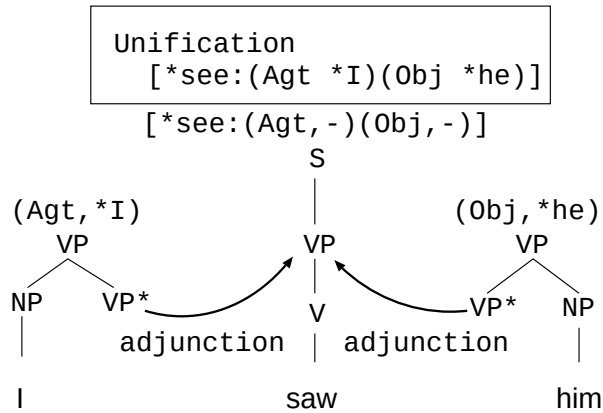


図 3.2: LTAG による英語の文法表現

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 3.1 に示す。掛川ら [47] の日本語文法の定義を基にしており、図 3.1 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。

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 3.2 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。

本モデルではピ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。したがって計算量の削減のため、自然言語処理の分野で作成されている既存の文法 [38] 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。

3.1.3 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 3.3 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。

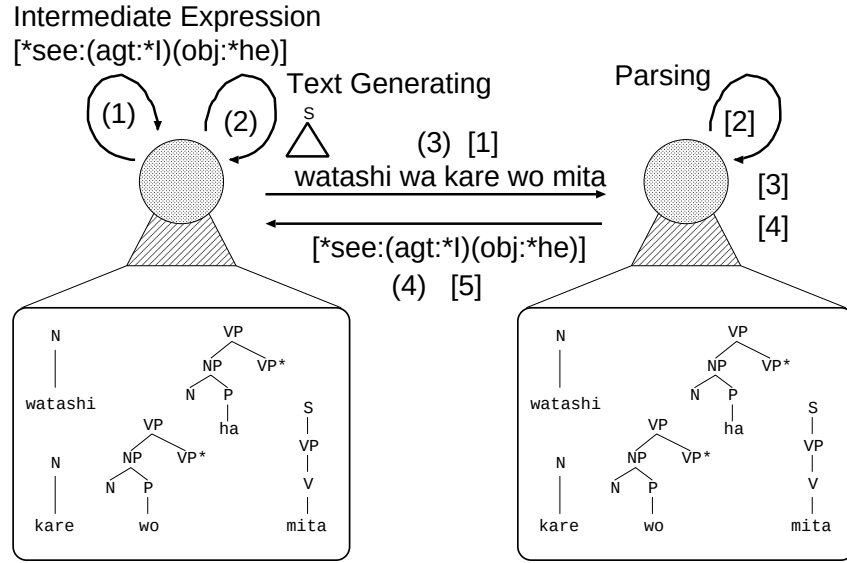


図 3.3: エージェント間の会話の流れ

3.1.4 学習機構

このモデルでは、遺伝的アルゴリズム (Genetic Algorithm, GA) [50] 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、ピジンを学習する人間の世代とは異なる。ピ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、したが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。

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、

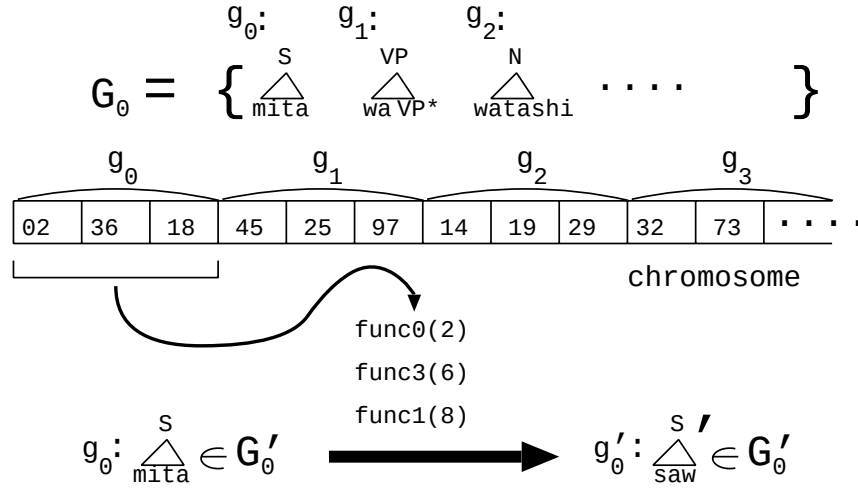


図 3.4: 遺伝子の表現型を介した文法の変換

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 3.4 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．3.1 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*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.

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 3.5(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 3.5(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. したが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 3.4 についての説明を行う. 図 3.4 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.

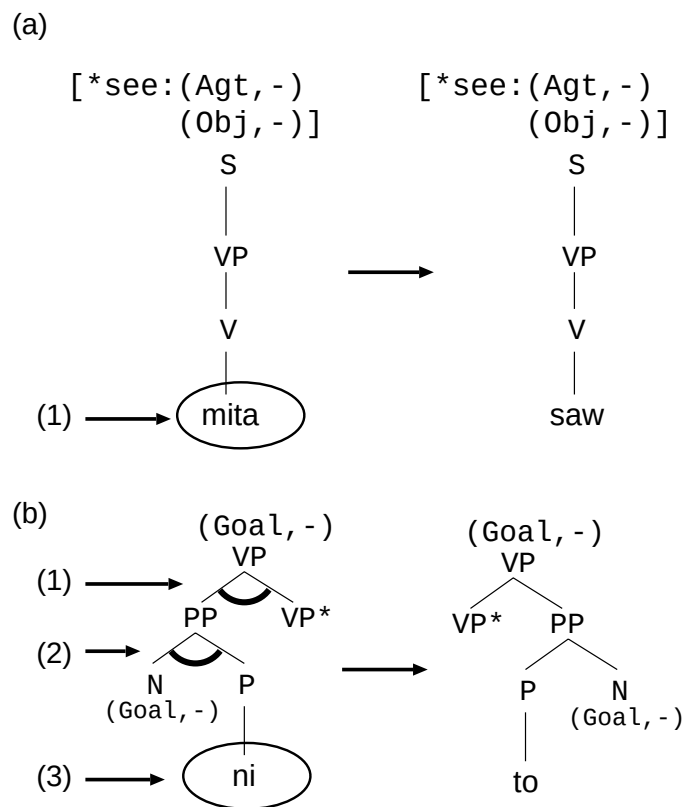


図 3.5: 日本語文法から英文法への変換の例

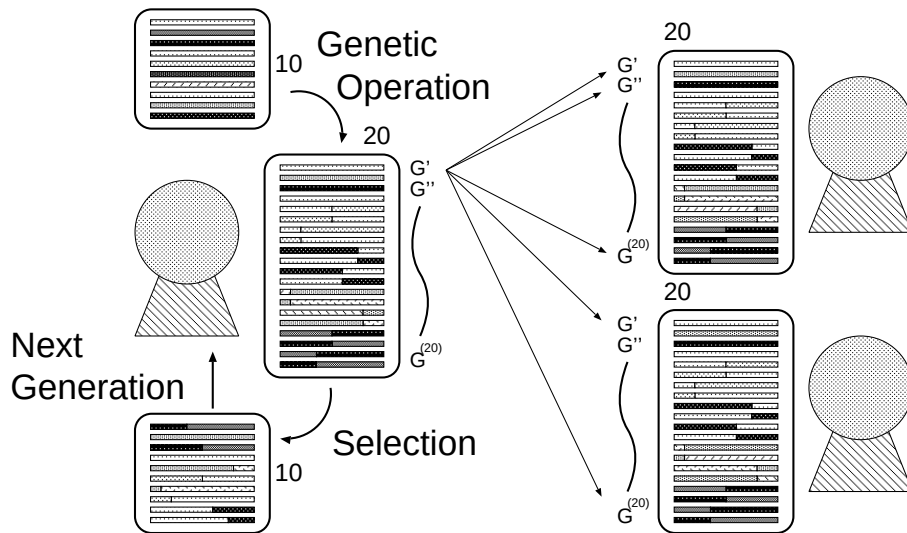


図 3.6: 会話モデルへの GA による学習の適用

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 3.6 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，会話をするという設定は，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つの染色体が生成される．
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．

- 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3.1.4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、 $G', G'', \dots, G^{(20)}$ が作られる。これらの文法セットを用いて、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。3.1.3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。

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 $G_{L1,m} \equiv G_{L2,m}$ ($L1, L2$ 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．

3.2 実験および考察

3.1 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．

3.2.1 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．

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

表 3.1: 各実験のパラメータ

	$ 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

- $[*go:(Agt *I)(Goal *tokyo)]$
- $[*go:(Agt *he)(Goal *tokyo)]$

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 3.1 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 3.7 から図 3.10 までに示しているが, これらは全て R 世代毎,

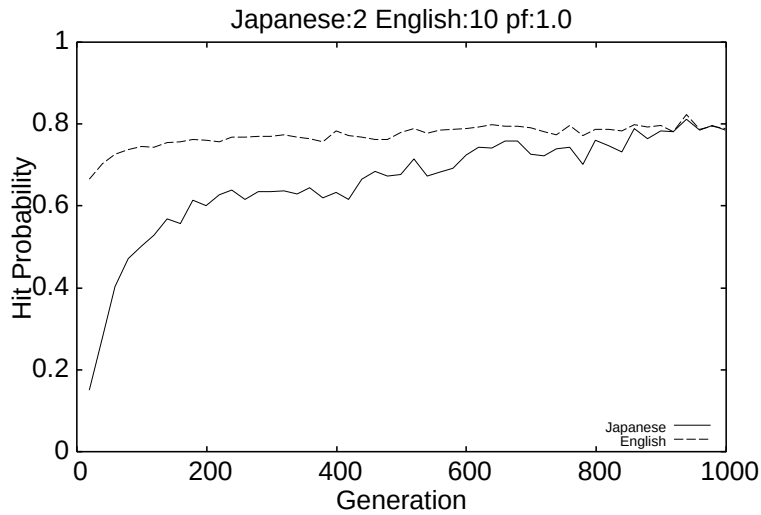


図 3.7: 実験 1 の結果

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。

3.2.2 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 3.7, 実験 2 の結果を図 3.8 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。

3.2.1 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

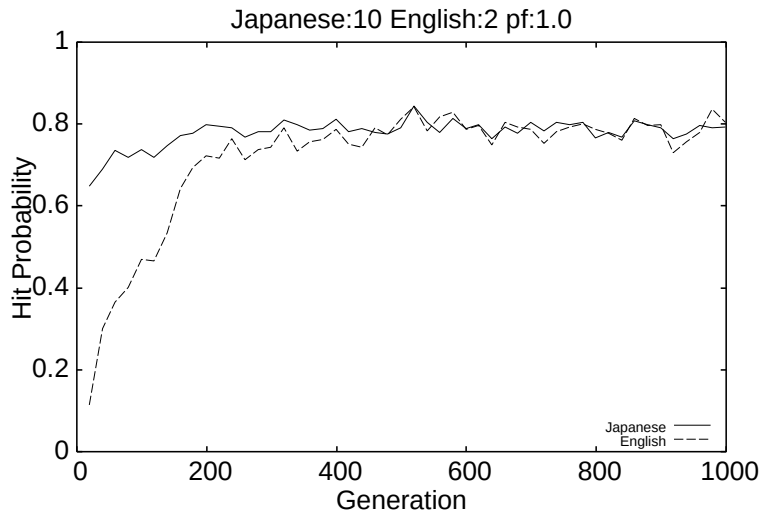


図 3.8: 実験 2 の結果

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もいえる。

3.2.3 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 3.9 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

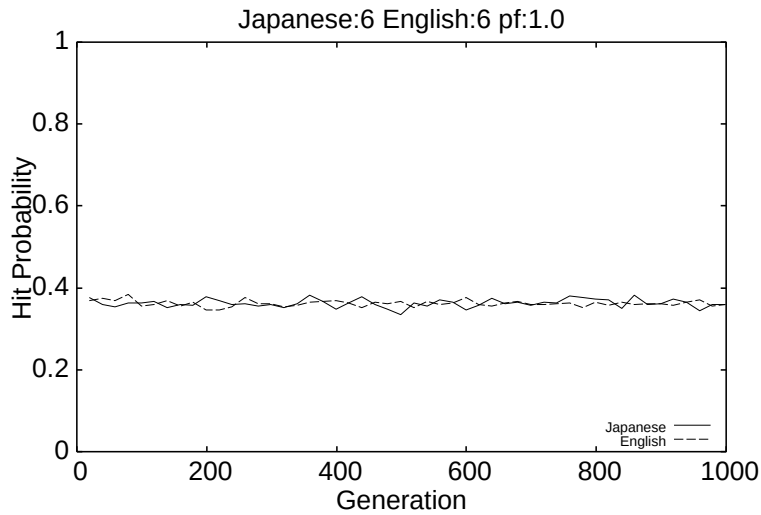


図 3.9: 実験 3 の結果

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 $5/11 \simeq 0.4545$ 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 3.10 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 値の上昇と共に減少した．

実験 3 の結果の 1 つ， $p_f = 5$ について，図 3.10 および図 3.11 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

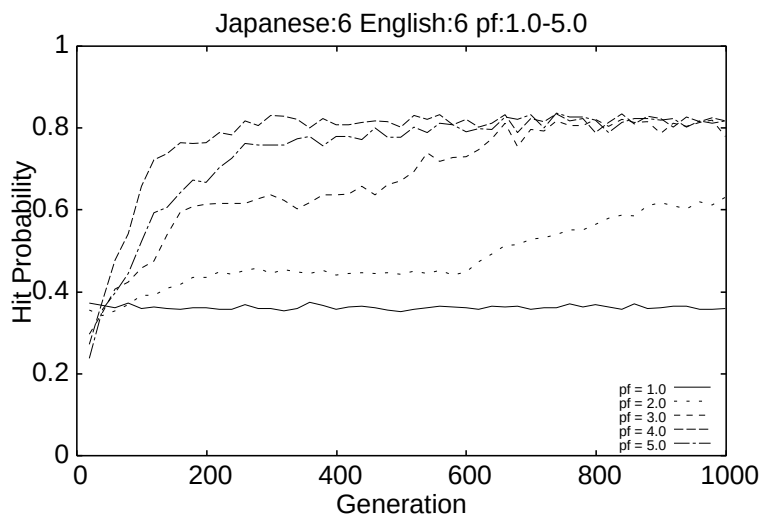


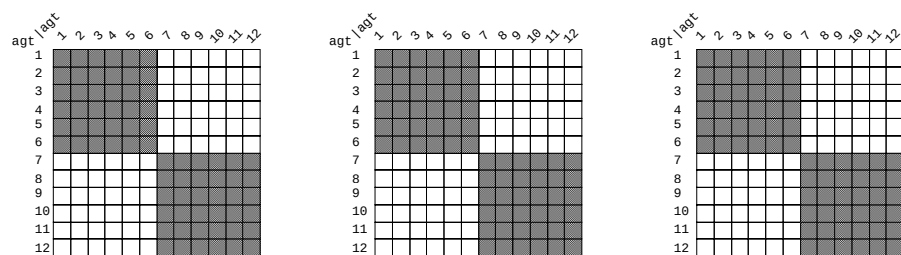
図 3.10: p_f 値の増加によるヒット率の変化

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 3.11 で示している。図中の agt1 から agt6 まだが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

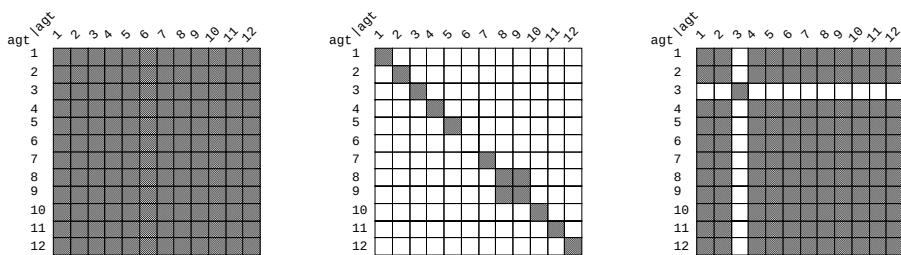
1. [*run: (Agt *I)]
2. [*see: (Agt *he)(Obj *I)]
3. [*go: (Agt *I)(Goal *tokyo)]

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。

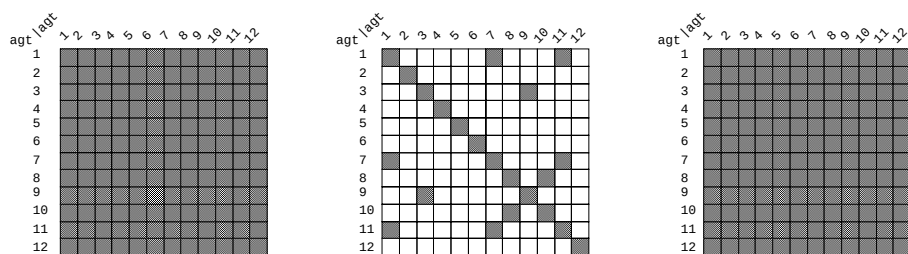
ここでは、文法 $G_{L,m}$ を用いた、特に文法に大きな変化があった第 0, 40, 80,



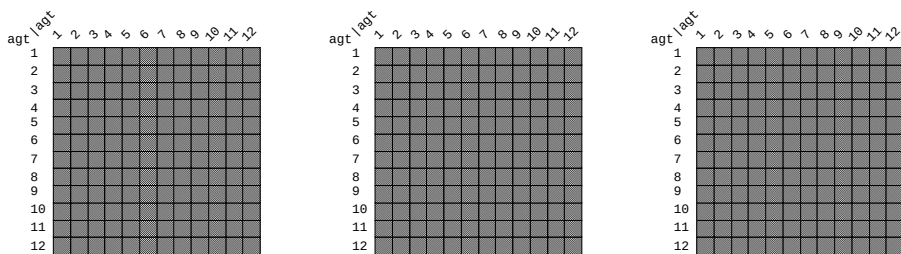
(a) Generation 0



(b) Generation 40



(c) Generation 80



(d) Generation 180

図 3.11: 各世代における中間表現ごとの伝達状況

表 3.2: 各世代における発話例

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

180 世代でのエージェントの代表的な発話例を表 3.2 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 3.11(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う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 $G'_{0,m} \cdots G_{0,m}^{(20)}$ 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、

[*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 3.2 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 倍の得点が適応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 3.11(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 3.11(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はなく、「(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。同様にして、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 [3]。次の文はトク・ピジンの例であり、

斜体が現地語である [46].

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、活用の欠落：時制、法、相 (TMA) といったものに関しては、ピジンにはほとんど見られない [3]. 下の文は 92 歳の日本人移民が発話したハワイ・ピジンである [37].

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：ピ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 [3]. 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、本実験ではピ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.

3.3 第 3 章のまとめ

本章の目的は言語学上実際に存在する混合言語であるピジンを計算機上で実装することであった. 本モデルではこの問題を、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ピ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，ピ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 3.9 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ピ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ることができる．

第 4 章

マルチエージェントによる人工クレオール の生成モデルとその問題点

本章の目的は、第 3 章に引き続き、マルチエージェントの枠組でクレオールが発現する過程のモデル化を行い、その条件を導き出すことである。第 3 章において、ピジンの創発を観察するためのマルチエージェント・モデルを提案した。しかし、そのモデルでは、異なる言語話者集団の初期の接触から、ジャーゴン (jargon) もしくは初期ピジンに至るプロセスを観察することができたが、その後の安定ピジンの特徴を示す文法拡張が困難であった。また、普遍文法を仮定していないことから、計算機上でクレオールを定義することが困難であった。この問題を踏まえ、本章で提案するモデルでは、普遍文法の原理とパラメータ理論を導入する。これにより、エージェントが話すことができる言語の種類は有限となる。したがって、本章で扱うモデルは、第 3 章で行ったような共通文法を開発していく過程をモデル化したものではなく、エージェントが有限の文法集合の中の最も獲得するにふさわしい文法を選択するという問題をモデル化したものであると捉えられる。本章で論じることは、前章から取り組んできたマルチエージェントによる文法獲得モデルのまとめであるとともに、第 5 章で論じる言語動力学への問題提起である。

以降、4.1 節でマルチエージェントモデルによる人工クレ奥ールの生成モデルの提案、実験および考察を述べる。次に 4.2 節でこれまでのマルチエージェントのまとめと問題点について指摘する。

4.1 クレオール生成モデル

4.1.1 人工クレオール生成モデルの提案

本モデルにおいては普遍文法における原理とパラメータ理論を導入する．各エージェントが持つ文法は，原理として記述される文法集合のうち，パラメータによってひとつが選択される．原理として与えられる全ての文法ルールはチョムスキー標準形で表される．また，2値をとる3つの独立したパラメータをここで仮定することにより，探索空間には8種類の文法が存在する．各パラメータは文法中の特定の複数のルールに割り当てられており，ルール中の右辺の非終端記号の順序を特定する．例えばあるパラメータの値が0であると，それに対応する文法中のルール“ $S \rightarrow NP VP$ ”について，右辺はそのままの順序が与えられるのに対し，パラメータの値が1である場合，このルールは“ $S \rightarrow VP NP$ ”と変化する．8つの文法はそれぞれパラメータのセットの値で名前が付けられている．すなわち，例としてパラメータ (b_2, b_1, b_0) の値が $(1, 0, 0)$ であったならば，その文法の名前は $G_{(b_2, b_1, b_0)_2}$ ，つまり G_4 となる．この文法集合を $\{G_0, \dots, G_7\}$ とする．

本モデルにおいて15個の再帰的ではない文脈自由文法のルールを原理として設定し，パラメータ (b_2, b_1, b_0) の各値は文法中の特定のルールの右辺の並びを決定することとした．表4.1で示されているルールにおいて，1桁目に書かれているパラメータが2桁目のルールに影響を及ぼすことを示している．3つのパラメータの各値は互いに独立しているが， b_2 は b_1 よりも構文木上の上位のルール，すなわち根ノードに近いルールに関係を持っている．同様に b_1 は b_0 よりも上位のルールに割り当てられている．したがって，もし2つの文法間の距離と言うものを考えたとき， b_2 の値の違いは b_1 以上に距離に影響を及ぼす． b_1 と b_0 の関係も同様である．

各エージェントは Inside-Outside アルゴリズム [27] を各ルールの確率推定に用いる．これは，統計的手法によってモデル中の隠れたパラメータの値を推定する手法である EM アルゴリズムの一種であり，自然言語処理においてタグ無しコーパスからの教師なし学習による文法獲得に用いられる手法である [25, 36]．Inside-Outside アルゴリズムは文脈自由文法にのみ適用され，確率推定されるルールには初期確率が与えられる．そして訓練コーパスによる推定確率の変化が最小になる

表 4.1: チョムスキー標準形で書かれたルール

b_2	$S \rightarrow NP \quad VP$	b_2	$S1 \rightarrow VP1 \quad NP1$
	$S \rightarrow VP$		$S1 \rightarrow VP1$
	$VP \rightarrow Vi$		$VP1 \rightarrow Vi$
b_1	$VP \rightarrow Vt \quad NP$	b_1	$VP1 \rightarrow Vt \quad NP1$
b_0	$VP \rightarrow VP1 \quad PP$	b_1	$NP1 \rightarrow Det \quad N$
	$NP \rightarrow NP1$		$NP1 \rightarrow N$
b_0	$NP \rightarrow NP1 \quad PP$		
b_0	$NP \rightarrow NP1 \quad S1$		
b_1	$PP \rightarrow Prp \quad NP1$		

まで繰り返し確率計算が行われる．表 4.1 のような単純なルールに対し，高い精度での文法推定能力を持ち，自然言語処理の分野において広く用いられる手法であることから本モデルにおいて Inside-Outside アルゴリズムを採用した．すなわち，エージェントがそれぞれ所有する文法を用いて会話を行い，その発話を蓄える．その文のセットをコーパスとみなすことにより，文法の解析を行うという方法である．

世代 t において，エージェントは自分自身を含めた全てのエージェントに等しく発話し，それを聞いたエージェントは受け取った文を全て記憶しておく．その後，各エージェントは文法の各ルールの適用確率を推定し，それにふさわしいパラメータ値を決定する．図 4.1 に会話と学習の流れを示す．

4.1.2 クレオールの変義

ここで本モデルにおけるクレオールの変義を行う．普遍文法によって制限される文法集合 $\{G_0, \dots, G_7\}$ において，その文法使用者の人口比率を考える．すなわち文法 G_i の話者の比率を x_i とする．このとき $\sum_i x_i = 1$ である．もし人間の言語の数が有限であるとすれば，クレオールもその中に含まれているはずである．ここでクレオール文法 G_c を考える．最初はどのエージェントもその文法を持っ

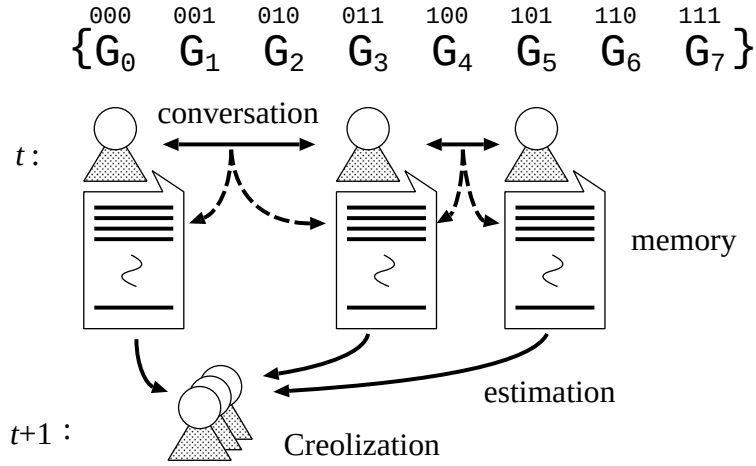


図 4.1: 会話と学習の流れ

ていない．すなわち $x_c(0) = 0$ である．クレオール化とはその文法を用いた言語がエージェント間で話され，最終的にそのエージェント群の中で優勢な文法となることである．本モデルにおいて，クレオールは次のような条件を満たす文法 G_c として定義される．

$$x_c(t) = 0, \quad x_c(t+1) \geq \theta_{io} \quad (4.1)$$

ここで θ_{io} は，その文法がどれだけの人口比率で話されれば優勢であるとみなすかということを示す閾値を表している．

4.1.3 実験と考察

10人のエージェントからなるコミュニティを仮定し，実験を行った．各エージェントは8種類ある文法のいずれかをもち，会話と学習を行うことによって文法の変化を調査する．エージェントが世代 t において，文を生成，および獲得し，文法を習得するプロセスは次の通りである．

1. エージェントは自身の文法を用いてそれぞれのルールの適用確率にしたがい，ランダムに文を生成する．左辺の品詞が同じルールの適用確率は均一である．その文を他のエージェントに対して発話する．それを聞いたエージェントはその文を記憶する．
2. 全てのエージェントが処理1を行う．つまり全てのエージェントは一回発話

をする。

3. 全エージェントが合計 1,000 文を蓄えるまで処理 2 を繰り返す。
4. 各エージェントは記憶した文から、各ルールの確率を Inside-Outside アルゴリズムで推定する。その結果から、パラメータ (b_2, b_1, b_0) の値を確定し、次世代でエージェントが所有する文法を決める。

以降、実験の結果と考察について述べる。

4.1.4 エージェントの構成とクレオールの生成に関する実験

世代 t において、10 人のエージェントがそれぞれ 8 種類の文法のなかからひとつを持つ。このとき、各文法使用者の人口に注目すると、取り得る可能な組み合わせは $17!/(10!7!) = 19,448$ 通りである。この全ての可能な組み合わせについて、エージェントの文法間の遷移を調査するための実験を行った。1 回の実験は、わずかに世代のあいだに起こる文法変化に関するものである。しかし、それぞれの言語集団の人口について、全ての組合せの実験を行うことにより、次世代の人口構成比に対応する組合せはその中に必ず存在する。したがって、この実験は、何世代にも渡って繰り返される文法変化を見ることに等しく、そのうち、一世代でクレオールが発現する組合せに着目したものである。

クレオールが発現する例を図 4.1 から説明する。このコミュニティで話されているのは G_0, G_3, G_5 の 3 言語である。10 人のエージェントはこの 3 つの文法のうちのひとつを所有している。このため、クレオールとなり得る言語は G_1, G_2, G_4, G_6, G_7 である。このコミュニティのエージェント同士でそれぞれ会話をを行い、学習によって文法の各ルールの使用確率を推定する。その結果、各エージェントはパラメータ値を変更し、それぞれが所有する文法を最も受理しやすい文法へ入れ換える。この文法の変化が世代の入れ替わりを意味する。すなわち、子供の第一言語獲得の結果、親の文法と異なる文法を身につけた状況である。多くの場合、既存のひとつの言語に集中するか、既存の複数の言語を共有し、人口の均衡を保つ。しかし、ある特定の人口の組合せについてクレオール化が観察された。図は、世代 t で 3 言語が話されていたコミュニティにおいて、世代 $t+1$ では全て

表 4.2: 文法の数に対するクレオールの数と比率

G	D	Creolization	
		$\theta_{io} = 0.5$	$\theta_{io} = 1.0$
1	8	0 (0%)	0 (0%)
2	252	1 (0.37%)	0 (0%)
3	2016	33 (1.64%)	9 (0.45%)
4	5880	112 (1.90%)	27 (0.46%)
5	7056	56 (0.79%)	10 (0.14%)
6	3528	6 (0.17%)	0 (0%)
7	672	0 (0%)	0 (0%)
8	36	0 (0%)	0 (0%)
Total	19448	208 (1.07%)	46 (0.24%)

G : The number of grammars.

D : The number of possible distributions.

のエージェントが G_1 を話すようになった状態を表している．これがクレオール化である．

(4.1) 式において $\theta_{io} = 1.0$ とした場合，クレオール化の現象が見られたのは合計 46 通りであった．また $\theta_{io} = 0.5$ の場合は合計 208 通りであった．ここでクレオール化と世代 t におけるエージェントが構成する言語集団の数との関係について考えてみる．使用された文法の数に対するクレオールが発現したコミュニティの数と比率を表 4.2 に示した． $\theta_{io} = 1.0$ であるとき，エージェント群が全体で所有する文法の種類が 3 種類から 5 種類のときにおいてのみクレオールは発現した．また， $\theta_{io} = 0.5$ で考えると，それ以外の場合もクレオールは現れたが，特に 3 種類から 5 種類のときが最も多いことがわかる．このとき，エージェント群が所有する文法が 3 種類のときに限ってみると，その可能な組み合わせが 2,016 通りあるのに対し，そのうちの 1.64% である 33 通りのケースでクレオール化が現れている．文法が 4 種類の場合では 5,880 通り中 112 通りでクレオールが発現し，1.90% であった．同様に 5 種類では 7,056 通り中 56 通りであり，0.79% となっている．この 3 つ

のケースが全てのクレオール約 97%を占めていた。6 種類以上のケースについては新言語が誕生する可能性としては低いと考えるのが妥当だろう。なぜなら世代 t において、ほとんどの言語が少なくとも 1 人のエージェントに話されているため、クレオールとなり得る文法に限られてしまうためである。

この結果はエージェント間の接触言語が 3 種類以上のときにクレオール化が起こりやすいということを示しており、2 言語での接触は、その内の一方の言語に収束しやすいという傾向にある。

4.1.5 クレオール化の条件に関する考察

4.1.4 節で行った実験結果から、クレオール化の条件を考察する。全ての組み合わせを行った各実験では、全てのエージェントは固定数、ここでは 1,000 文を学習のためにそれぞれ蓄える。これらの文は他のエージェントが、それぞれの文法を用いて生成した文であるため、記憶された 1,000 文は、生成された文法ごとに分割することが可能である。ここで、各文法ごとの文の数というのは、発話したエージェントが持っている文法の比率に比例する。つまり、記憶中にある文に使用されている文法の比率は、その文法使用者の人口比率そのものとなるのである。この比率は全ての組み合わせで行った各実験において 10 人のエージェントに共通であり、さらにエージェントはこの記憶から次世代の文法を推論するため、ほとんどの場合、10 人のエージェントは同じ文法を習得することになる。その他のケースとして、エージェントの文法が共通の文法に収束しなかったケースが挙げられるが、世代 t で優位であった文法の人口は世代 $t+1$ においてほとんどが増加したか、もしくは同数であった。

ここで図 4.2 において、(4.1) 式の条件を基に、クレオール化の条件を知ることが出来る。図中全てのエージェント群は世代 t における使用文法が 3 種類または 4 種類となっているものを取り上げた。例えば図 4.2(a) では G_0 , G_5 , G_6 だけが使用されている。この時、世代 t でどのエージェントも使用していなかった新言語 G_4 が、世代 $t+1$ ですべてのエージェントによって話されているのがわかる。

ここで各パラメータ (b_2, b_1, b_0) の視点から、新言語が話される条件を考えてみる。図 4.2(a) のパラメータ b_2 だけに注目してみると、世代 t において 10 人のエー

		G_0	G_1	G_2	G_3	G_4	G_5	G_6	G_7
(a)	t	2	0	0	0	0	4	4	0
	$t+1$	0	0	0	0	10	0	0	0

$b_2 b_1 b_0$	$b_2 b_1 b_0$	Value
0 # #: 2	1 # #: 8	1
# 0 #: 6	# 1 #: 4	0
# # 0 : 6	# # 1 : 4	0

 $\Rightarrow G_4$

		G_0	G_1	G_2	G_3	G_4	G_5	G_6	G_7
(b)	t	1	0	0	3	0	3	3	0
	$t+1$	0	0	0	1	0	0	0	9

$b_2 b_1 b_0$	$b_2 b_1 b_0$	Value
0 # #: 4	1 # #: 6	1
# 0 #: 4	# 1 #: 6	1
# # 0 : 4	# # 1 : 6	1

 $\Rightarrow G_7$

		G_0	G_1	G_2	G_3	G_4	G_5	G_6	G_7
(c)	t	0	1	4	0	3	2	0	0
	$t+1$	6	0	0	0	4	0	0	0

$b_2 b_1 b_0$	$b_2 b_1 b_0$	Value
0 # #: 5	1 # #: 5	-
# 0 #: 6	# 1 #: 4	1
# # 0 : 7	# # 1 : 3	1

 $\Rightarrow \begin{cases} G_0 \\ G_4 \end{cases}$

図 4.2: 実験結果の一例

ジェントのうち2人が0の値を持っており、残りの8人は1である。それゆえ次世代の言語では b_2 に関して優勢である $b_2 = 1$ である傾向が見られる。同様に b_1 は0であり、 b_0 も0となる。図中の‘#’はワイルドカードを示しており、その値は問わない。このような傾向から、次世代の文法の各パラメータは $(b_2, b_1, b_0) = (1, 0, 0)$ を持ち、すなわちその文法とは G_4 となった。その他の例 (b) においても (a) と同様の理由により、次世代においてほとんどのエージェントが習得した文法は、 $(b_2, b_1, b_0) = (1, 1, 1)$ から G_7 である。(c) は全てのエージェントが共通の文法を獲得できなかった例であり、この理由は b_2 の値が固定されていないからである。

このように、クレオール化は異なる文法間の人口配分に大きく依存していることが確認された。パラメータは、 b_2, b_1, b_0 の順に文法に与える影響が大きくなるように定義されているにも関わらず、これらはほとんどクレオールの発現に影響を及ぼさなかった。この結果はクレオール化を含めないケースにも一般的に論じることが可能であり、もしエージェント群の中に優勢な文法を持った集団がいると、次世代においてもその文法は生き残る上、他の文法を持つエージェントはその文法に引き寄せられるように習得すると考えられる。このように、次世代にどの言語が流行するかということを決めるのは、各パラメータの多数決で求められるということが観察できた。

4.1.6 文法と遷移の関係についての考察

ピジン化が、言語学習者と発話環境との相互作用によって現れる現象であるのに対し、クレオール化は、親から子供へ文法を継承する際にその文法が正常に継承されずに変化していく現象であるという点で異なる。何世代にも渡ってくり返し文法を学習するモデルを提案する際、文法の遷移確率を考えることができる。ここで、ある文法 G_i を持つエージェントの子供が文法 G_j を獲得する確率を q_{ij} と定義する。この確率は全ての文法から全ての文法への遷移について考えることができるため、遷移行列 $Q = \{q_{ij}\}$ として表される。このとき任意の i について $\sum_j q_{ij} = 1$ である。この遷移行列 Q は、各言語間の類似度と学習アルゴリズムに依存することが既に知られている [23, 33, 34] (5.1.1 節に後述)。しかし、本実験ではこれらに加え、世代 t における各言語の人口比率に大きく影響を受けることを示唆して

表 4.3: 人口比率に依存した遷移行列 Q の例

	G_0	G_1	G_2	G_3	G_4	G_5	G_6	G_7
G_1	0.85	0	0	0	0.15	0	0	0
G_2	1.00	0	0	0	0	0	0	0
G_3	0.85	0	0	0.05	0.10	0	0	0
G_4	0.80	0	0	0	0.20	0	0	0

いる.

表 4.3 に人口比率に依存した遷移行列 Q の例を示す. このときの比率は世代 t において $x_1 = 0.2$ ($2/10$), $x_2 = 0.2$ ($2/10$), $x_3 = 0.2$ ($2/10$), $x_4 = 0.4$ ($4/10$) である. 表中の各値は 4.1.4 節と同様の方法で 10 回計算したものの結果である. この表は世代 t において文法 G_1, G_2, G_3, G_4 を持ったそれぞれのエージェントが, 世代 $t+1$ において G_0 から G_7 のうちのいずれかの文法に遷移する確率を表したものである. 例えば一行目に注目すると, 世代 t で G_1 を持っているエージェントは, 会話と学習を行った結果, 世代 $t+1$ において獲得する文法は 0.85 の確率で G_0 , 0.15 の確率で G_4 である. すなわち, $q_{10} = 0.85, q_{14} = 0.15$ であることを表している.

各パラメータ値はそれぞれの多数決で求まるとした場合, 世代 $t+1$ で優勢になると予想される文法は $G_{(0,0,0)_2}$, すなわち G_0 である. 実験結果では, $q_{i0} (i = 1, 2, 3, 4)$ の値が他の文法への遷移確率よりも高いことを示しているが, このとき常に G_0 に収束するのではない. 実験を通して, 遷移行列 Q は, 世代 t における x_i の配分に大きく依存することを示したが, それが次世代の文法を決める要因のすべてではなく, 言語間の類似性なども関係していると考えられる.

4.2 マルチエージェントの問題点

本節では, 本章で提案したモデルのまとめを行うとともに, 第 3 章と第 4 章で提案したマルチエージェント・モデルを通じて発生した問題点をここに挙げる.

4.1 節では, マルチエージェントの枠組で, クレオール創発のためのモデル化を

行った。このモデルでは原理とパラメータを仮定したため、エージェントが話す言語は、原理によって与えられた8種類の言語の中から選択される。文法の学習には、自然言語処理の分野において大規模コーパスから文法を獲得する標準的手法の一つである EM アルゴリズムを用いた。したがって、第3章のピジン化のモデルのように、エージェントが所有する文法の構文規則の変化を議論するものではなく、あらかじめ原理として定義された文法集合のうちのひとつをパラメータによって選択する過程を扱うものとなった。このため、クレオールの定義は、エージェントが適用する文法ごとの人口比率の変化として与えた。これらを基に、すべての文法に可能な組合せで人口を割り当てて実験を行った。パラメータごとの多数決で決定されるという単純な結果であったが、クレオールの創発には文法そのものの性質だけではなく、言語話者の集団の人口構成比に大きく影響を受けるということが導かれた。

本章の実験は、第3章に引き続いてマルチエージェントの枠組みでクレオール化の過程を扱ったものである。これまでの2つモデルによって行った実験の目的は、ピジン化やクレオール化の過程をシミュレーションによって再現し、その創発条件を得るということであった。しかし、マルチエージェント・モデルでこれらの言語進化モデルを提案するにあたり、いくつか問題が発生した。そのひとつとして、計算量の問題からエージェントの数が10人程度に制限されたことが挙げられる。エージェントの人口は、計算機上に言語変化のモデルを扱う上で、最も直接的に現実世界を模倣することができる要素のひとつである。したがって、現実的な時間内で実験を遂行するためにこの数が制限されたことは、モデルの拡張なども含めて大きな問題であると考えられる。このため、ピジン、クレオールのための人口構成比に関する条件について実験から導き出された結果は、一般的性を欠くものとなったことは否めない。また、自然言語処理や人工知能の分野で培われた技術を適用することにより、ピジン化やクレオール化のモデルを与えてきたが、これらの学習機構や文法は、言語学的に根拠が乏しく、モデルの妥当性の評価が難しいという問題も考えなければならない。

これらの問題を解消するために、次章において人口動力学に基づいたモデルを展開する。本章で提案したクレオール化のモデルは、その人口動力学モデルの起点となる重要な役割を果たしていると考えられる。特に、4.1.6 節の考察

が，次章で述べる動的遷移行列モデルの提案に大きく影響を与えている．

第 5 章

言語動力学におけるクレオールの新発

本章の目的は、人口動力学モデルによってクレオールの新発を観察し、そのための条件を導き出すことである。これまでに、マルチエージェントの枠組で第 3 章でピジン化、第 4 章でクレオール化の過程を計算機上に再現し、その条件について論じてきた。しかし、ピジンおよびクレオールの新発と、それぞれの言語話者人口の関係について求めようとしたとき、有限数からなるマルチエージェント・モデルの枠組では、人口構成比に関する条件について言及する際、一般性を欠いたものとなる。これは、マルチエージェント・モデル全体が抱える重大な問題である。

これらの問題を解決するために、2.2.2 節で述べた人口動力学を用いてクレオールの新発を観察する。本章では、言語動力学とよばれる既存の人口動力学による数学モデルを基に、クレオールの新発に必要であると考えられる項目を取り入れたモデルを提案する。この人口動力学モデルは、同時に計算量の問題も解消することが可能である。これまでに提案したマルチエージェント・モデルでは、論理形式で記述された意味から文生成を行い、発話文をパースするという処理を各エージェントについて行う必要があったため、一回の実験で膨大な時間を要した。このため、必然的にエージェントの数や処理内容が限られていたのである。それに対し、人口動力学においては個々のエージェントの処理内容に言及する必要がないことが利点とされる。

本章では、まず最初にモデルの提案と構成論的手法によるそのモデルの検証を行う。その後、クレオールが新発する原因と考えられるさまざまな要因について、新発のための条件を求めることを目的とした実験を行う。

以降、5.1 節において言語動力学について説明し、その後我々の修正点について述べる。5.2 節では修正した動力学に基づいた具体的なモデルを示す。このモデルの健全性を示す実験およびその結果を 5.3 節で述べ、次にクレオールの創発条件を求める実験とその結果を 5.4 節に示す。最後に 5.5 節でまとめる。

5.1 文法獲得に関する人口動力学

本節ではまず最初に、Nowak et al. [33, 34], Komarova et al. [23] によって提案された言語の動力学について説明し、その問題点を議論する。そして動力学モデルをより現実的なものにするために、これらのモデルの改善点を考える。

以降、文法 G_i から導出される言語を $L(G_i)$ とし、文法 G_i を所有して言語 $L(G_i)$ を発話する大人を G_i 話者と呼ぶことにする。

5.1.1 Komarova のモデル

Nowak と Komarova は文法獲得に関する人口動力学の数理的な理論を発展させることを目的として、進化ゲームの動力学 [41] を応用したモデルを提案している [23, 33, 34]。これを言語の動力学 (Language Dynamics Equations) と呼ぶ。彼らは子供の第一言語獲得が系全体で話される言語にどのように影響を与えるかについてモデル化を行った。このモデルでは大人を言語話者、子供を言語学習者として定義している。子供は親からのみ言葉を聞き、そこから文法を推定する。

Komarova et al. が提案した言語の動力学方程式は、言語が持つ性質を次のように与えている：

- **類似性**：言語間の類似度を表す。任意の 2 つの文法間の類似性を示す確率行列 $S = \{s_{ij}\}$ で表される。要素 s_{ij} は、 G_i 話者がランダムに文を発話したとき、文法 G_j を持つ聞き手に理解される確率として求められる。
- **遷移性**：言語間の遷移度を表す。 G_i 話者である親の子供が文法 G_j を獲得する確率は行列 $Q = \{q_{ij}\}$ で表される。

言語話者である大人は子供を産み、その数は各言語に与えられる適応度に比例す

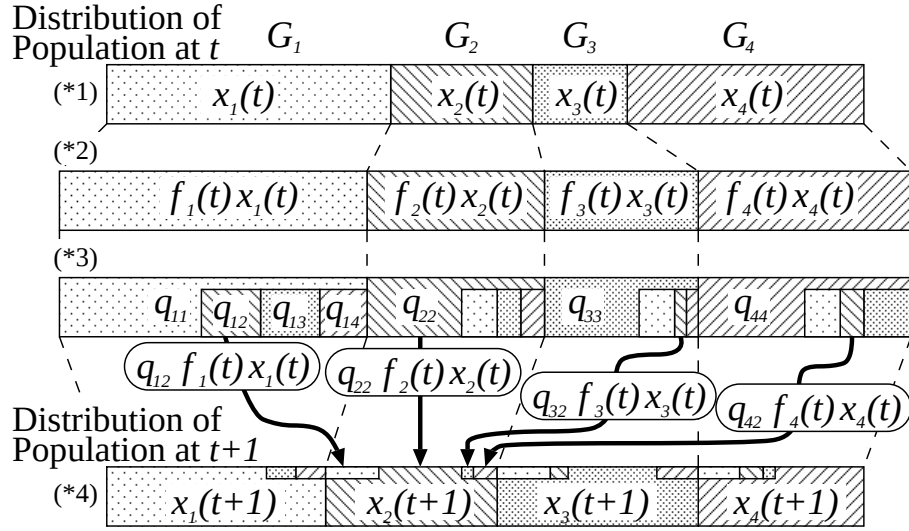


図 5.1: 人口変化の流れ

る．ここで G_i の適応度は $f_i(t) = \sum_{j=1}^n (s_{ij} + s_{ji})x_j(t)/2$ である．すなわち，より言葉を理解しかつ，理解された者が，より多くの子孫を残せると仮定している．これは G_i 話者の発話が，コミュニティで理解される確率を表している．これらを用いて言語の動力学方程式は次のような微分方程式で与えられる：

$$\frac{dx_j(t)}{dt} = \sum_{i=1}^n q_{ij} f_i(t) x_i(t) - \phi(t) x_j(t) \quad (j = 1, \dots, n). \quad (5.1)$$

ここで $-\phi(t)x_j(t)$ は世代を通じて総人口を一定に保つための項であり， $\phi(t) = \sum_{i=1}^n f_i(t)x_i(t)$ としている．

この式は次のような状況を描写していると解釈される（図 5.1 参照）：

- (*)1 あるコミュニティにおいて $\{G_1, \dots, G_n\}$ の言語が話されている．各個体はその中のひとつを話すことができる． G_i 話者人口の割合を x_i とすると，この図は各言語話者の人口構成比を表す．
- (*)2 言語話者は各言語の適応度に比例して子供を産む．すなわちコミュニケーション能力が次世代に子孫を残すための条件であることを表している． G_i 話者の人口比率 x_i に対して，産まれる子供の総数の比率は $f_i x_i$ となる．この比率は大人の全人口に対するものであり，一般には $\sum_i f_i x_i \neq 1$ である．

- (*3) 産まれた子供たちは親からのみことばを受け取り、文法を推定する． G_j 話者の子供は q_{jj} の確率で正常に G_j を獲得し、その人口比率は $q_{jj}f_j(t)x_j(t)$ である．その他の子供は q_{ji} の確率で G_i を身につける．このとき q_{jj} の値は学習アルゴリズムに依存するため、言語学習の精度（accuracy）と呼ばれる．
- (*4) 次世代の G_i 話者の人口比率 $x_j(t+1)$ は親の言語を正しく継承したものと、他の言語話者からの流入の総和となる．すなわち $t+1$ 世代の G_j 話者の人口比は $\sum_{i=1}^n q_{ij}f_i(t)x_i(t)$ である．ここで ϕ により総人口を一定に保つように調整する．

上の解釈から、言語話者である親は多言語コミュニティの中で他の言語話者と会話をした結果、子供を産むのに対し、子供は親以外からことばを聞くことはない．しかし、このような状況において、子供が親の持つ文法を獲得するのに失敗し、他の文法を身につけるということは現実的に考えにくい．この文法獲得に失敗する可能性について、Niyogi [32] が提案したモデルを用いて次節で論じることとする．

5.1.2 Niyogi のモデル

Niyogi [32] は、言語学的な根拠に基づいた文法を基に、言語学習者がトリガー学習アルゴリズム（Trigger Learning Algorithm; TLA） [16] を用いて文法獲得を行うモデルを提案した．TLA は、原理とパラメータ理論を仮定した文法獲得アルゴリズムである．学習者が入力文からパラメータの値を推定し、それに対応する文法を獲得する．学習の際、入力文を蓄える必要がないことから、最も単純なメモリーレス・学習アルゴリズム（memoryless learning algorithm） [33] として知られている．

アルゴリズム (TLA)： n 個からなる二値をとるパラメータが初期状態として与えられると、学習者は、そのパラメータに対応した文法によって入力文 S に対し構文解析を試みる．もし S の解析に成功すると、学習者は、現在仮定した文法が目標文法であるとみなし、その仮定を覆さない．しかし、もし学

習者が S を解析できなかったとき、 n 個のパラメータのうち、等確率でパラメータの要素 P を選択する。このとき、各パラメータについて選択する確率は $1/n$ である。そしてパラメータ P の値を反転させ、変更したパラメータ値に対応する文法によって、再度 S の解析を試みる。もしそれで解析することができたら、その変更したパラメータを採用するが、さもなくば、元のパラメータを維持する。

Niyogi は、この TLA による学習過程をマルコフ状態図で表し、学習者が目標文法を正常に獲得できない可能性を指摘した。付録 A の文法を例として、説明する。この文法は、二値を持つ 3 つのパラメータに対応し、8 つの独立した文法を持つ。これらの文法は、N, V, O, O1, O2, Adv, Aux からなる構文カテゴリについて、それぞれ異なった組合せを生成する。

ここで、 G_5 （対応するパラメータは [010]、語順は S V O）が目標文法であると仮定する。すると、学習過程を図 5.2 のようなマルコフ状態図で表すことができる。数字ラベルが付いた円は、パラメータに対応した状態、また、図の中心にある、二重丸で囲まれた状態が目標文法である。各状態が遷移することができる近隣の状態とは、TLA にしたがってひとつのパラメータ値だけが異なっており、状態間を結ぶ有向な弧は、それぞれ遷移が可能であることを表している。薄い線で描かれた大きな円は、目標文法と各状態の間にあるパラメータ値の差を表している。例えば、目標文法のパラメータ値が [010] であるのに対し、最も内側にある状態のひとつである G_1 の値は [110] であり、 p_1 がひとつ反転すれば目標文法に届く状態にあることから、この 2 つの文法の間にある距離は 1 である。

学習者は初期状態としてランダムにパラメータ値を与えられ、これに対応した文法が目標文法であると仮定する。学習者は一度に目標文法によって生成される文を一文受け取り、円の内側または外側に進むか、同じ状態にとどまる動作をする。TLA のアルゴリズムから、それぞれの動作は次の条件によってなされる。

- 1) 仮定した文法で受理できる文を受け取ったとき、同じ状態に留まる。
- 2) 仮定した文法で受理できない文を受け取ったとき、ランダムにひとつ選んだパラメータの値を変えた文法を用いて再度その文を解析する。もし受理できたら、仮定した文法を変更し、それに対応する状態に遷移する。遷移するの

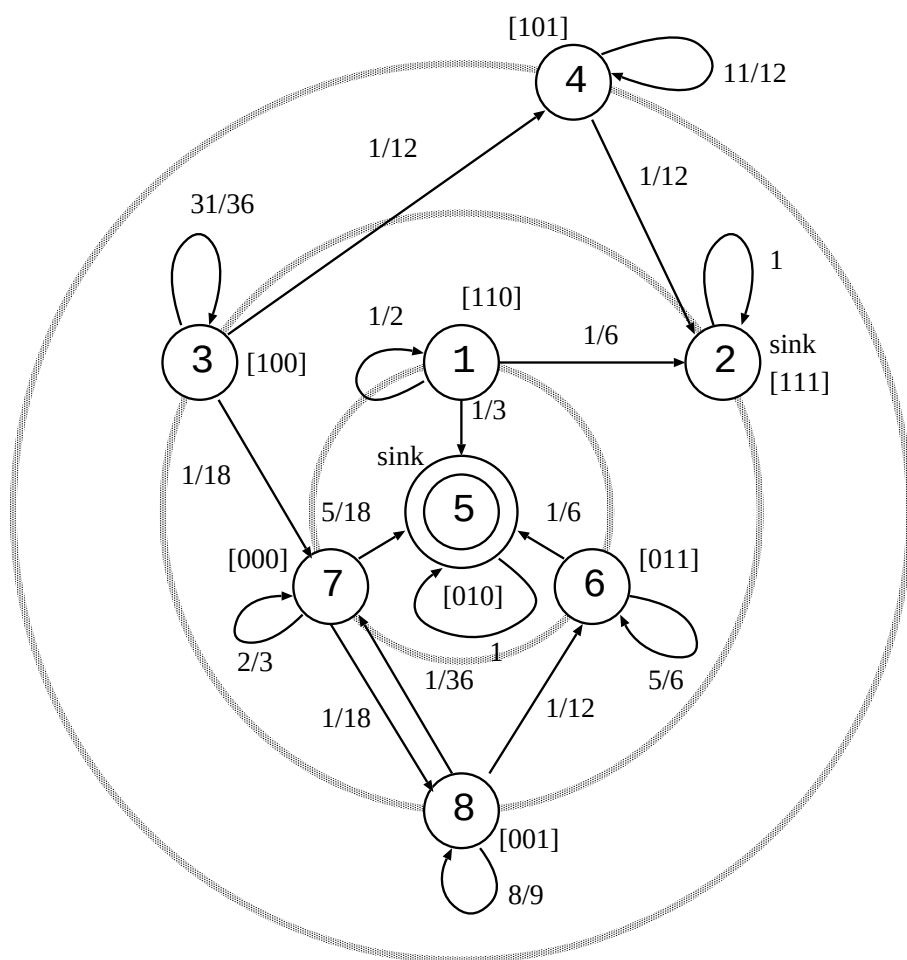


図 5.2: マルコフ状態図 [32]

は、現在の状態から円のひとつ内側か、ひとつ外側のどちらかである。

3) 2) において、再度受理できなかった場合、同じ状態に留まる。

図中の弧に付いている値は、学習者が目標文法であると仮定する文法の遷移確率を表している。パラメータ値の違いがひとつしかない状態間においても、遷移できない状態、すなわち確率が 0 の弧は図から省いている。例えば、 G_5 から導出され、 G_1 で受理できない任意の文（例えば “S V”）は、同様に G_3 においても受理することができないため、 G_1 から G_3 への遷移は存在しない。また、逆も同じである。

学習者は、全て G_5 から導出される入力文に対応して、目標文法の仮説を変えていき、いくつか入力文を受け取った後、最終的に目標文法にたどり着く。例えば、初期状態として G_8 を目標文法であると仮定していた学習者は、目標状態にたどり着くのに最低で 2 文聞く必要がある。また、 G_6 から学習を始めると、目標状態にたどり着くために聞く必要がある文の数の期待値は 6 文である。目標文法にたどり着くと、入力文は必ず受理できるため、その文法から他へ遷移することはなくなり、学習は終了する。しかし、学習を進めて行く過程において、学習者が G_2 を仮定すると、それ以降目標文法にたどり着くことができなくなる。 $G_2([111])$ で受理できない文を受け取ったとき、近隣の状態で受理できる場合、学習者は仮説を変える可能性があるが、 $G_4([101])$ 、 $G_1([110])$ 、 $G_6([011])$ のいずれの文法においても、 $G_2([111])$ で受理できない文を受理することができないためである。そのため、一旦 G_2 を目標文法と仮定してしまうと、 G_5 にたどり着くことはできなくなってしまうのである。 G_2 のような状態を吸収状態 (Absorbing State) とよぶ。

このモデルから、5.1.1 節で我々が指摘した子供が親以外の文法を間違えて獲得する可能性、すなわち $q_{ij} > 0$ ($i \neq j$) の具体例を考えることができる。

- A) 子供がマルコフ過程のある状態、すなわち目標文法とは異なるある文法を仮定している状態に陥ると、親が持つ目標文法から導出された正しい例文のみを受け取っているにも関わらず、その状態から逃れられないという状態が存在する。この状態を吸収状態 (Absorbing State) と呼ぶ。
- B) 子供は現在の状態では受理できない文を受け取った際、これを学習における刺激として学習によって状態を遷移させ、目標状態に近づいていく。しかし

目標文法に到達する前に、子供に十分な刺激が与えられないまま学習期間が終了してしまうと、学習過程にある状態に対応した文法を誤って身につけてしまう。

S 行列は文法間の推移性に関わるため、 Q 行列は S 行列に依存する [23]。また言語獲得における精度もまた学習アルゴリズムに影響を受ける。そのため、Niyogi が用いた学習アルゴリズム (TLA) は q_{ij} ($i \neq j$) の確率を不自然に高くする可能性がある。

5.2 動的遷移行列モデル

本節では 5.1 節で述べた言語動力学の問題を解消するために、新たなモデルを提案する。

5.2.1 モデルの改良

ここで、5.1.2 節で示した Niyogi が述べている 2 つの記述について考えてみる。まず A) に関して、現実世界の子供は、親からいくら言葉を聞いても学習過程にある自分の文法を修正できなくなるという状態に陥ることは考えにくい。一般的に、子供は任意の言語について、その文を聞いて学習する限り、目標文法を誤りなく獲得できると考えられている。それゆえ、吸収状態は存在しないと考えるべきである。また B) については、たとえ子供が言語学習期間において十分な刺激を得られなかったとしても、その子供は他の言語を身につけるということは考えにくく、Genie¹ のように言葉を身につけることが出来なくなるだろう。このような理由から、現実世界において子供が親の文法の獲得に失敗し、他の文法を身につけるといふ確率は非常に小さいと考えられる。すなわち $q_{ij} \simeq 0 (i \neq j)$ と考えるべきだろう。

上記を踏まえ、より現実世界に近いモデルを提案するため、我々は Q 行列の定義を変えることから始める。子供が習得する言語は、その言語学習期間において接

¹ 誕生以来 13 年間父親によって小部屋に閉じ込められて育った子供で、第一言語を正常に獲得できなかった例としてしばしば取り上げられる [37]。

触した言語とその接触頻度に大きく影響されることは第 4 章でも述べている [30].
 よって、言語学習者である子供は親からのみ発話文を受け取り文法を獲得するという Komarova et al. のモデルを修正し、子供はコミュニティに属するさまざまな言語話者と接触し、そこから文法を学習すると考える. このとき、他言語話者との接触の結果、親の言語を正常に身につけられない可能性が考えられ、その確率を $\bar{Q}$ 行列として新たに定義することを提案する. ここでコミュニティの言語話者ごとの人口比率は世代によって変遷するため、 $\bar{Q}$ 行列は時間に関するパラメータを持つようになる. それゆえ $\bar{Q}(t) = \{\bar{q}_{ij}(t)\}$ となる. 我々はこれを**動的遷移行列** (Dynamic Transition Matrix) と呼ぶ. したがって (5.1) 式は次のように修正される [29].

$$\frac{dx_j(t)}{dt} = \sum_{i=1}^n \bar{q}_{ij}(t) f_i(t) x_i(t) - \phi(t) x_j(t) \quad (j = 1, \dots, n). \quad (5.2)$$

これを**動的遷移行列モデル** (Dynamic Transition Matrix Model) と呼ぶことにする.

5.2.2 接触確率 α の導入

次に我々は、子供が親以外の言語話者と接触する確率を表すパラメータ α を導入する. これを**接触確率** (Exposure Probability) と呼ぶ. ここで子供が親の言葉を聞く確率は $(1 - \alpha)$ である. このとき α は親の言語以外の言語と接触する確率ではなく、親の言語も含めた多言語との接触確率である. 例を図 5.3 に示す. G_p はある子供の親の文法である. その子供は確率 α の割合で他の言語話者と集団の言語話者の比率に応じて接触する. すなわち、図中の影がかかった部分の割合で子供は親の言語を聞くことになる. ここで $\alpha = 0$ のとき、親からしか言語を学習することがないため、Komarova et al. が想定した状況と同じである. また逆に、 $\alpha = 1$ のとき、各言語話者の人口構成比に完全に比例した割合で言語と接触するため、どの言語話者の子供も獲得する言語の条件は等しくなる.

以上をまとめると、新たに定義した $\bar{Q}(t)$ 行列は、接触確率 α および各言語話者の人口構成比 $X(t) = (x_1(t), x_2(t), \dots, x_n(t))$ に依存する.

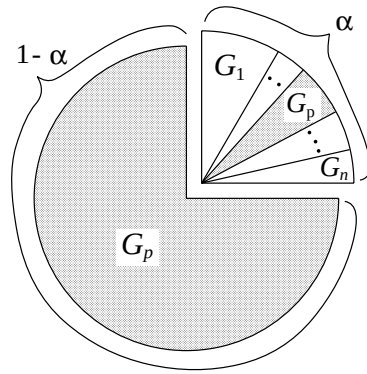


図 5.3: 接触確率 α

5.2.3 学習アルゴリズム

我々は Niyogi のモデルの問題を踏まえ、学習アルゴリズムに次のような制約を与えた：

- a) 言語学習者である子供は生まれた時点で特定の文法を持たない．すなわちパラメータの初期値を与えない．これに対し，Niyogi のモデルでは初期値としてランダムにパラメータ値を与えるため，生まれてすぐになんらかの文法を持っていると仮定している．
- b) 子供は親からしかことばを聞かずに学習した場合，必ず親の文法を獲得する．これは Niyogi のモデルでは保証されず，子供の文法の獲得過程を示す状態遷移に依存する．また Komarova et al. のモデルでは，この状況における文法獲得の失敗確率を Q 行列として定義している．
- c) 学習期間中は，目標文法の推定に十分な時間と例文が与えられる．

ここで上記制約を満たす単純な学習アルゴリズムを導入する．図 5.4 に示した学習の様子を以下に解説する：

- 1) 子供は言語話者によって発話された一文を聞く．この図では G_8 話者から “S V O” という一文を受け取っている．
- 2) 子供は頭の中で文法の数だけカウンタを持っており，もしその文がある文法によって受理されるなら，その文法に対応したカウンタの値をひとつ上げ

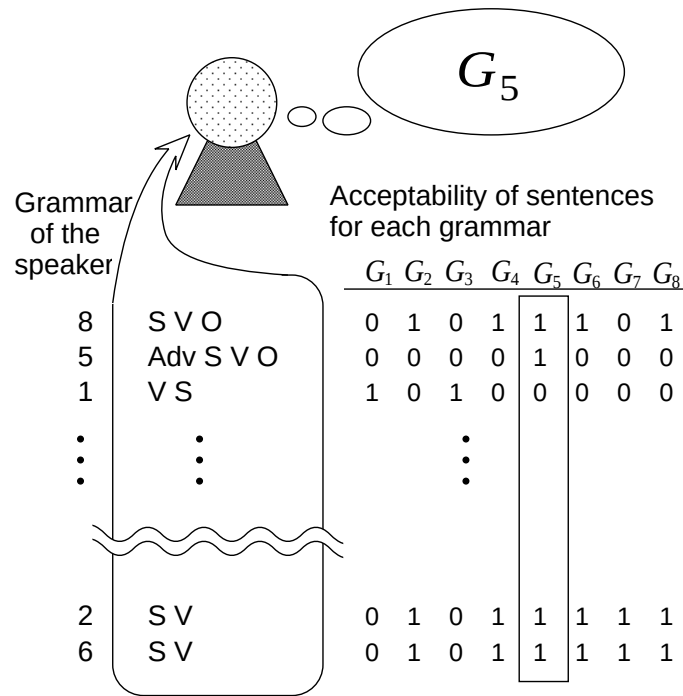


図 5.4: 単純な言語獲得アルゴリズムの導入

る．これを全ての文法について行う．この図は “S V O” を受理可能な文法が G_2, G_4, G_5, G_6, G_8 であることを表している．

- 3) 文法の推定に十分であると考えられる数の文を受け取り，その間， 1) と 2) を繰り返す．この図では “S V O” 以降 “Adv S V O” “V S” ... “S V” の順に文を受け取っている．

- 4) 最も高い値を示したカウンタに対応した文法を子供は採用する．この図は受け取った文を最も多く受理した文法が G_5 であることを表している．

このアルゴリズムを定式化することを考える．学習対象が親の言語だけであった場合，上記の制約 b) から，子供が獲得する文法は，次のような G_{j^*} となる：

$$j^* = \operatorname{argmax}_j s_{pj} \quad (= p),$$

ここで p は親の文法のインデックスを意味する．

また子供は，コミュニティの各言語話者の人口に比例してそれぞれの言語を聞く機会がある．その場合，子供が獲得すると予想される文法は，次の式を満たす

G_{j^*} となる：

$$j^* = \operatorname{argmax}_j \left\{ \sum_{k=1}^n s_{kj} x_k(t) \right\}.$$

ここで 5.2.2 節で定義した，親以外の言語話者と接触する割合を表す接触確率 α を導入する．これにより，文法の選択は上記 2 式の線形結合となり，子供が推定する文法は次のような G_{j^*} となる：

$$j^* = \operatorname{argmax}_j \left\{ \alpha \sum_{k=1}^n s_{kj} x_k(t) + (1 - \alpha) s_{pj} \right\}. \quad (5.3)$$

5.2.4 動的遷移行列 $\overline{Q}(t)$

動的遷移行列 $\overline{Q}(t) = \{\overline{q}_{ij}(t)\}$ の定義は， t 世代における各言語の話者に影響を受けながら文法を学習した結果， G_i 話者の子供が G_j に遷移する確率である．したがって (5.3) 式を確率関数に変換する必要がある．ここで (5.3) 式から， $P^n(i, j) = \alpha \sum_{k=1}^n s_{kj} x_k(t) + (1 - \alpha) s_{ij}$ とする．これは n 種類ある言語のうち， G_i 話者の子供が G_j によって受理することができる文を受け取る確率である．まず最初に 2 つの文法 G_1 と G_2 しか存在しない場合を考える． G_1 を持っている言語話者の子供は，次のような条件を満たした場合 G_1 を獲得する：

$$P^2(1, 1) \geq P^2(1, 2)$$

両辺の値はそれぞれ独立して 0 から 1 までの範囲で値をとる．このとき子供の学習前の初期状態で，どちらの値もわからない場合の文法の採択確率を考える．ここで両辺の値が 0 から 1 までの範囲で一様に分布すると仮定すると， G_1 を採用する確率は左辺の値そのもの ($0 \leq P^2(1, 1) \leq 1$) である．同様に n 個の文法 $\{G_1, \dots, G_n\}$ のケースを考える． G_1 を持っている言語話者の子供が G_1 を獲得するためには，

$$P^n(1, 1) \geq P^n(1, i) \quad \text{for all } 2 \leq i \leq n$$

という条件を満たさなければならない．すなわち $n - 1$ 個の文法と比較するため， G_1 の採択確率は $(P^n(1, 1))^{n-1}$ となる．同様に， G_i を持つ言語話者の子供が G_j を獲得する確率を，それぞれの文法が受理する確率から求めたものは次のように

なる：

$$(P^n(i, j))^{n-1} = \{\alpha \sum_{k=1}^n s_{kj} x_k(t) + (1 - \alpha) s_{ij}\}^{n-1}. \quad (5.4)$$

これを j に関して正規化することによって $\bar{q}_{ij}(t)$ を得る：

$$\bar{q}_{ij}(t) = \frac{(\alpha \sum_k s_{kj} x_k(t) + (1 - \alpha) s_{ij})^{n-1}}{\sum_l (\alpha \sum_k s_{kl} x_k(t) + (1 - \alpha) s_{il})^{n-1}}. \quad (5.5)$$

このとき $\sum_{j=1}^n \bar{q}_{ij}(t) = 1$ である．

この節で論じたモデル，すなわち動的遷移行列 $\bar{Q}(t)$ の有効性を検証するために実験を行い，人口構成比と接触確率の変化においてクレオール創発を見る．また先行研究のモデルを修正する際に Q 行列と並んで重要な役割を負っていた S 行列（類似性）についてもクレオール創発の条件を検証する．このため，本研究では，以下のように実験計画を立てる．

実験 1 動的遷移行列モデルの検証

実験 2 優勢クレオールが創発する条件の検証

次節以降では，それぞれの実験について節を分けて実験の方法と結果について論じる．

5.3 実験 1 - 動的遷移行列モデルの検証

ここでは 5.2 節で提案した動的遷移行列を含む言語動力学の振る舞いを，クレオールの創発に注目して分析する．

5.3.1 人口動力学上のクレオールの定義

ここで我々は，これまでの言語学的な定義 [43, 46] とは大きく異なるが，人口動力学の視点から見たクレオールの定義をする．クレオールとは，新しい言語の創発現象である．すなわち，あるコミュニティである時点に存在しなかった言語が，後に存在するようになる現象と考えることができる．よって，次のように定義できる．

定義 1 (共存クレオール) 他の言語と共存するクレオールとは、次のような文法 G_c である：

$$x_c(0) = 0, \quad x_c(t) > \theta_c.$$

定義 2 (優勢クレオール) 優勢言語となるクレオールとは、次のような文法 G_c である：

$$x_c(0) = 0, \quad x_c(t) > \theta_d.$$

ここで $x_c(t)$ は人口動態が収束し、安定した t 世代における G_c 話者の人口比率、 θ_c と θ_d はそれぞれ共存（coexistent）クレオールと優勢（dominant）クレオールであるとみなすための人口比率の閾値を示す．本稿では $\theta_c = 0.1$, $\theta_d = 0.9$ としている．これらの定義は、初期状態では誰も話していなかった言語が、最終的にはある割合の話者を獲得することを表している．定義 1 は、少数ではあるが、一定数の個体が世代を通じて文法を維持することを意味し、定義 2 はそのコミュニティ内で使用される言語のほとんどがクレオールによって占有される状態を表している．

5.3.2 実験 1 の言語セットと S 行列

実験 1 では Niyogi のモデル [32] と同様、Gibson et al. [16] によって導出された 8 つの文法を採用した（付録 A 参照）．各文は、文法項目、すなわち品詞の順序で指定されると考える．

一般的に、各文法とそこから導出される文の生成確率が求まれば S 行列を求めることができる．ここでは、 G_i を持つ言語話者がそれぞれの文型の適用確率（またはそれぞれの文を発話する確率）が均一であると仮定し、 S 行列の要素 s_{ij} は $L(G_i)$ と $L(G_j)$ にある共通な文の種類数を $L(G_i)$ にある文の数で割った値として求めた．表 A.1 から、静的に S 行列を求めることが可能であり、次のように

なる。

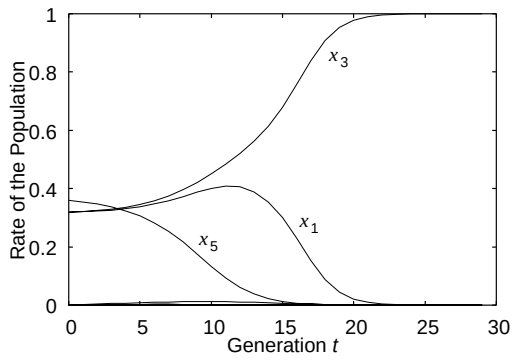
$$S = \begin{pmatrix} 1 & \frac{6}{12} & \frac{2}{12} & \frac{3}{12} & \frac{0}{12} & \frac{1}{12} & \frac{0}{12} & \frac{1}{12} \\ \frac{6}{18} & 1 & \frac{2}{18} & \frac{10}{18} & \frac{6}{18} & \frac{8}{18} & \frac{1}{18} & \frac{5}{18} \\ \frac{2}{12} & \frac{2}{12} & 1 & \frac{2}{12} & \frac{0}{12} & \frac{2}{12} & \frac{0}{12} & \frac{2}{12} \\ \frac{3}{18} & \frac{10}{18} & \frac{2}{18} & 1 & \frac{3}{18} & \frac{5}{18} & \frac{1}{18} & \frac{8}{18} \\ \frac{0}{12} & \frac{6}{12} & \frac{0}{12} & \frac{3}{12} & 1 & \frac{6}{12} & \frac{2}{12} & \frac{3}{12} \\ \frac{1}{18} & \frac{8}{18} & \frac{2}{18} & \frac{5}{18} & \frac{6}{18} & 1 & \frac{1}{18} & \frac{10}{18} \\ \frac{0}{12} & \frac{1}{12} & \frac{0}{12} & \frac{1}{12} & \frac{2}{12} & \frac{1}{12} & 1 & \frac{1}{12} \\ \frac{1}{18} & \frac{5}{18} & \frac{2}{18} & \frac{8}{18} & \frac{3}{18} & \frac{10}{18} & \frac{1}{18} & 1 \end{pmatrix} \quad (5.6)$$

対角要素である s_{ii} は常に 1 であり，また s_{12} を求める場合， $L(G_1)$ に含まれる文の総数は 12 なのに対し， $L(G_2)$ と共通な文は 6 個であるため， $s_{12} = 6/12 = 0.5$ となる。

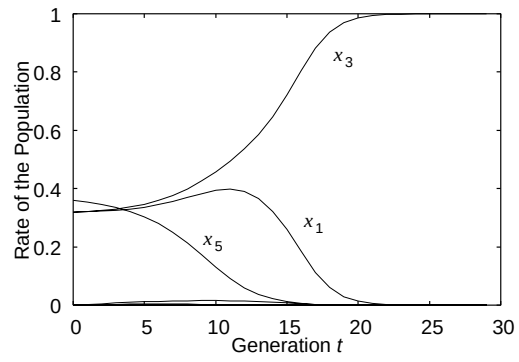
実際に話されている言語を考えると，異なる言語間で共通の文を理解しあう確率は全く無いに等しい．すなわち $s_{ii} = 1$ ， $s_{ij} \simeq 0$ ($i \neq j$) と考えるのが妥当である．それに比べ，(5.6) 式で与えられた S 行列の各要素の値は，大きな値となっているものが多い．これは，付録 A で与えられた言語セットは，それぞれの言語の語彙について言及されておらず，埋め込みがない文型についてのみ列挙していることが原因である．この言語セットは，言い替えれば，語彙が共有され，埋め込みがない文型のみを使用するというピジンの特徴を持っているとみなされる．したがって，(5.6) 式として与えられた S 行列はピジン化の過程を経てもたらされたものであると考えることができる．

5.3.3 実験 1-1 - 人口動力学上でのクレオールの創発

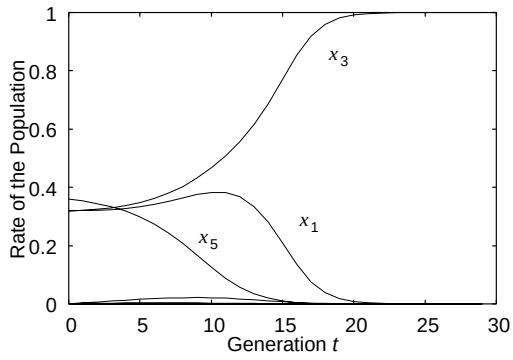
本実験の目的は，接触確率 α をパラメータとすることによって，我々が提案した動的遷移行列モデルの振る舞いを観察することである．予備実験において最も顕著に特徴が現れたところとして，初期状態 $x_1(0) = 0.32$ ， $x_3(0) = 0.32$ ， $x_5(0) = 0.36$ ，その他は $x_i(0) = 0$ の値を与えたときの実験を行った．この初期値を用いて接触確率 α を 0 から 1 の範囲で与え，モデルの振る舞いを観察する．動的遷移行列モデルの結果を図 5.5 および図 5.6 に示す．



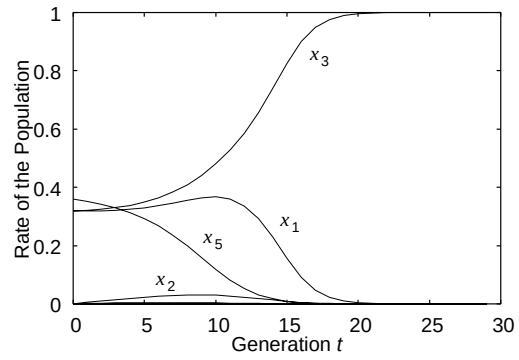
(a) $\alpha = 0$ Not Creolized



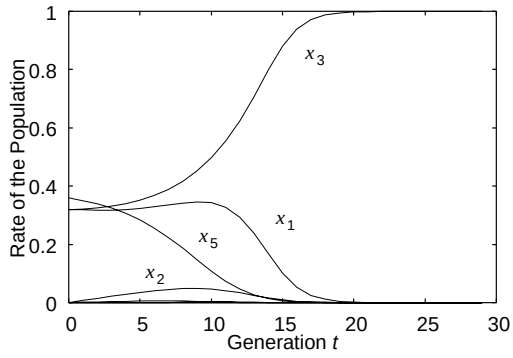
(b) $\alpha = 0.100$ Not Creolized



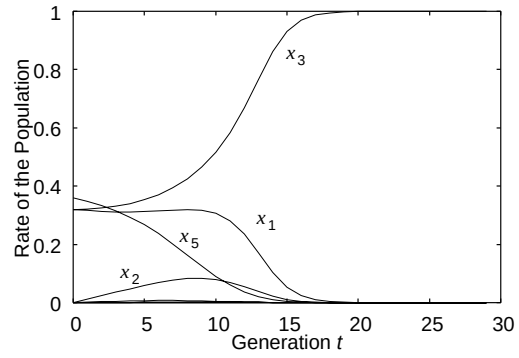
(c) $\alpha = 0.200$ Not Creolized



(d) $\alpha = 0.300$ Not Creolized

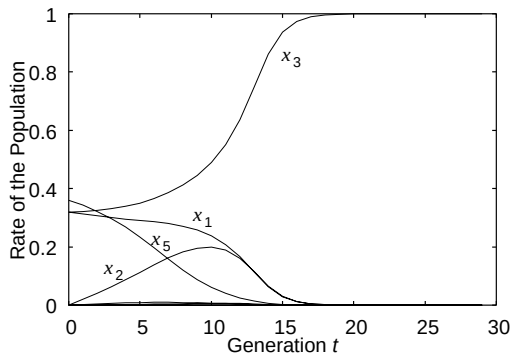


(e) $\alpha = 0.400$ Not Creolized

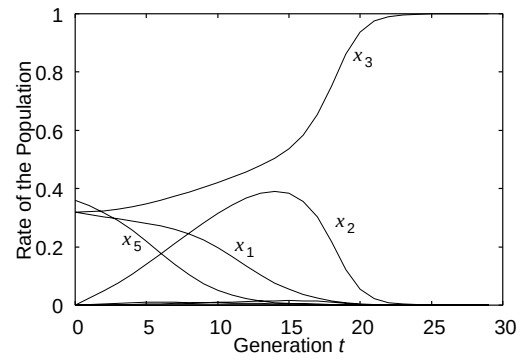


(f) $\alpha = 0.500$ Not Creolized

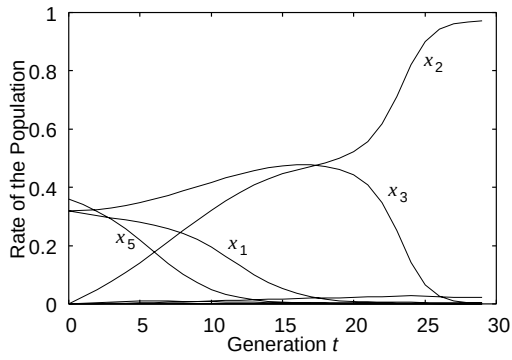
図 5.5: 動的遷移行列モデルの結果 ($0 \leq \alpha \leq 0.5$)



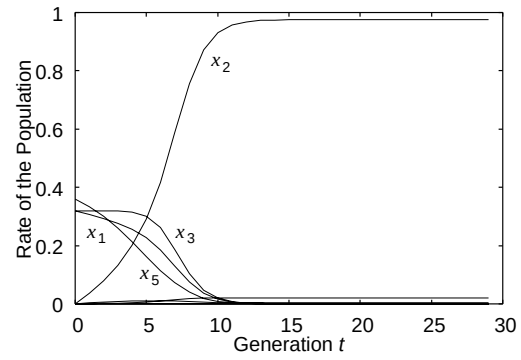
(a) $\alpha = 0.600$ Not Creolized



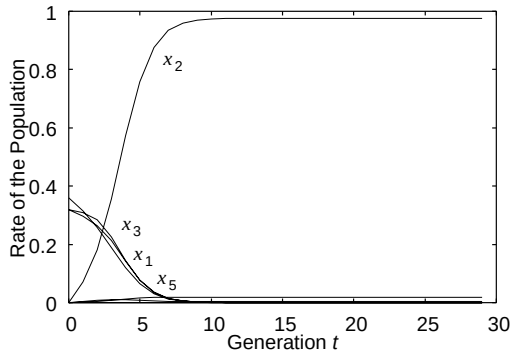
(b) $\alpha = 0.627$ Not Creolized



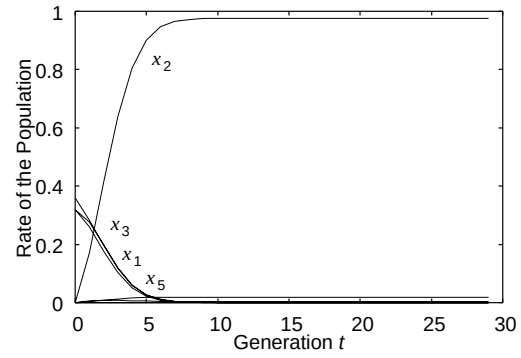
(c) $\alpha = 0.628$ Creolized



(d) $\alpha = 0.700$ Creolized



(e) $\alpha = 0.800$ Creolized



(f) $\alpha = 1$ Creolized

図 5.6: 動的遷移行列モデルの結果 ($0.5 < \alpha \leq 1$)

図 5.5(a) は $\alpha = 0$ を与えたときの結果である．このとき言語学習者である子供は親からしか言語を学ばない．すなわち，(5.2) 式の動的遷移行列 $\overline{Q}(t) = \{\overline{q}_{ij}(t)\}$ は定数であるため，この動力学は (5.1) 式と同じ振る舞いをする． G_3 話者の人口比率は世代を経るごとに増加し，最終的にはコミュニティのほとんどが G_3 を獲得する．これは文法の適応度に応じてそれぞれの文法を持つ人口が増減した結果であり， x_3 が増加するのに対して， x_1 と x_5 は適応度の減少によって人口比率を減らしてしまうのである．図 5.5 と図 5.6 は， α の値を増やしていったときの結果を表している．

図 5.6(b) は $\alpha = 0.627$ を与えたときの結果である． α が増加するに従い，初期人口を与えられていない x_2 が徐々に増加している．これは初期状態において誰も話していなかった G_2 が，世代を経ることによってその話者を増加させていることを表している．この原因は， α が増加することによって (5.2) 式中の特に q_{12} の値が増加し， G_1 話者の子供が G_2 を獲得するようになったためであると考えられる．我々はこの現象を**クレオール**の創発であると捉えている．これは，言語学習者である子供が，さまざまな言語話者と頻繁に接触することによって，親の言語よりもそのコミュニティにおいて最もコミュニケーション能力の高い文法を選んだ結果である．しかし，この図 5.6(b) の場合， x_1 と x_5 のほとんどが x_2 へ流出した後， x_2 はそれ以上人口を増加させることができず， x_3 に吸収され，最後には消失してしまう．

$\alpha = 0.627$ と $\alpha = 0.628$ の間には大きな境目が存在することが図 5.6(c) からわかる． $\alpha = 0.627$ では G_3 が最終的な優勢言語であったのが， $\alpha = 0.628$ において優勢言語は G_2 に転じていることが観察された．すなわち定義 2 の優勢クレオールが創発したことを示している．これまでの α の変化による系全体の振る舞いと比較して， α が 0 から 0.6 付近までは緩やかな変化だったのに対し，この境目付近における振る舞いの変化は，それまでと比べて非常に急激なものであった．

その後さらに α の値を増加していったところ， G_2 は優勢クレオールであることを保ち続け，より安定していった（図 5.6(f) 参照）．そのうえ， α が増加すると，収束世代が短くなっていることが図 5.6(c) と図 5.6(f) を比較するとわかる． α の定義から， $\alpha = 1$ であるときに最もクレオールが創発しやすい条件であることが認められる．

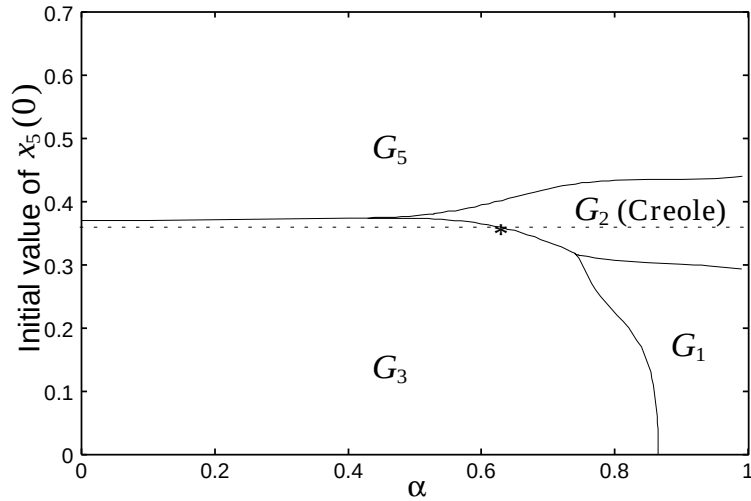


図 5.7: 優勢文法の領域の出力

本実験において、特定の S 行列と各言語の人口比率を初期値として与えたところ、クレオールの新規発生を観察した。またその創発は接触確率 α に依存することが確認された。

5.3.4 実験 1-2 - 初期条件に見る優勢言語の領域

5.3.3 節で見た動的遷移行列モデルの振る舞いから、どの言語が優勢となるかは、初期人口比率と接触確率 α に依存することがわかる。このとき、初期人口を与えられていない言語が優勢言語となると、クレオールと認識されるのである。次の実験では上記 2 つの初期条件と、優勢言語となる言語の関係を示し、そこからクレオール化の条件を調査することを目的とする。

優勢言語となる言語間の境界を明白にするため、 G_5 の初期人口比率をパラメータとする。すなわち $x_5(0)$ を 0 から 1 の範囲で与え、 $x_1(0)$ と $x_3(0)$ には次のようにその残りを均等に配分する：

$$x_1(0) = x_3(0) = (1 - x_5(0))/2.$$

この $\alpha - x_5(0)$ 平面において、図 5.6(b) と図 5.6(c) のような、優勢言語が入れ替わる境界を調べた。その結果を図 5.7 に示す。図中、実線は優勢言語が入れ替わる境目を表しており、それぞれの領域内に示されている文法 G_i が優勢言語と

なった文法を表している．その中にはクレオールとなった G_2 も含まれている．破線は 5.3.3 節で行った一連の実験を表し，アスタリスク (*) が描かれている点は図 5.6(c) の初期値 ($\alpha = 0.628, x_5(0) = 0.36$) に対応する．

さらに図 5.7 を詳しく見ていくと， G_2 が図の右半分以降に現れていることから，クレオールが創発するための条件として， α の値がある程度大きくなければならないことがわかる．

図 5.5(a) では， G_5 の初期人口比率が最も多いにも関わらず，初期人口が割り当てられている他の言語との類似度の低さから G_5 の適応度は G_3 のそれ以下となっている．そのため， G_5 は優勢言語になっていないが，それよりもさらに初期人口比が増えた $x_5(0) = 0.371$ になると，類似性の低さを初期人口比率で補い，優勢言語となる． $x_5(0)$ の値が低い場合， G_1 ， G_3 の人口比が増す． $t = 0$ における G_1 と G_3 の適応度は同じであるが， α の値によって優勢言語は異なる．例えば $\alpha = 0$ のとき，動的遷移行列は S 行列だけで求まるため， $\bar{q}_{11}(t) \simeq 0.99219$ であるのに対し， $\bar{q}_{33}(t) \simeq 0.99998$ である．したがって G_1 話者の子供は G_3 話者の子供に比べ他の言語に遷移しやすい．また，他の言語話者の子供は G_1 よりも G_3 に遷移しやすいため， G_3 が優勢言語となる．一方， G_1 は G_3 よりもクレオール G_2 と類似度が高い．すなわち $s_{13} = s_{31} = 2/12$ であるのに対し， $s_{12} = 6/12, s_{21} = 6/18$ である．そのため， α が大きいと図 5.6(b) のように x_2 が上昇し， G_2 に伴って G_1 も適応度が増える．同時に G_1 に関する動的遷移行列の値も影響を受け，特に $\bar{q}_{21}(t)$ と $\bar{q}_{12}(t)$ の値が増加する．このとき， $x_1(0)$ の値が小さいと， $\bar{q}_{21}(t) < \bar{q}_{12}(t)$ となり， G_1 話者の子供は G_2 に遷移しやすくなり， G_2 が優勢クレオールとなる．反対に， $x_1(0)$ の値が大きいと， $\bar{q}_{21}(t) > \bar{q}_{12}(t)$ となり， G_1 が優勢言語となる．このようにして， α の値が小さいところでは G_3 ，大きいところでは G_1 が優勢言語となるのである．

初期の人口構成比に関してさまざまな組合わせで実験を行ったが， $\alpha = 0$ のときにクレオールが創発したケースは確認されていない． $\alpha = 0$ では我々のモデルは Komarova et al. のモデルと一致する．すなわち，動的遷移行列 $\bar{Q}(t)$ と接触確率 α を導入したことにより，クレオールの創発現象を実現することができた．

5.3.5 実験1の考察

本節の実験においてクレオールが創発する現象を観察したが、初期の人口構成比や使用される言語によっては優勢クレオールは創発しない。例えば図 5.7 から、初期人口を $(x_1, x_3, x_5) = (0.4, 0.4, 0.2)$ として与えると、 α を増加させることによって優勢言語は G_3 から G_1 に変化するが、クレオールは優勢にならないことがわかる。

また、実験 1-1 の環境において、 G_1 から G_8 までさまざまな人口比率を初期値として与えたところ、クレオールとして創発した言語は G_2 だけであった。クレオールとは、その定義から初期において誰も話していなかった言語である。したがって、このモデルにおいては単純に G_2 話者の人口が初期値として与えられる場合、すなわち $x_2(0) > 0$ のとき、クレオールは創発しない。クレオールとなる可能性は G_2 として与えられた文法だけではない。本実験で G_2 だけがクレオールとなった原因は、 G_2 と他の言語との類似度を表す S 行列の要素、 s_{i2} または s_{2i} ($1 \leq i \leq n$) が他の値と比べて比較的高いためであると考えられる。

したがって、実験 1 から得られた結論は、言語話者の人口遷移は、言語間の類似性、言語話者の人口構成比、および接触確率 α に大きく依存し、クレオールが出現するためには α が十分大きく、他言語との類似度がある程度大きい必要がある、ということである。また、既存のモデルでは見られなかったクレオールの創発を観察することで、本章で導入したモデルがより現実的であることを示した。

5.4 実験2 - 類似性に関する条件の検証

実験 1 で、クレオールが創発するためにはクレオールと他の既存の言語間の類似性について条件があることが示唆された。本節ではこの条件を一般的に導出することを目的とする [28]。

5.4.1 実験2の言語セットと S 行列

ここでクレオールとして創発する言語とはどのようなものであるか考える。動的遷移行列モデルから動力学を導出するために必要な言語情報は S 行列だけである。

実験1では、現実的な言語セットを仮定してモデルを検証することが必要であったことから、それぞれの言語は原理とパラメータ理論から導かれる8言語から S 行列を求めた。しかし、優勢クレオールが創発するための言語間の類似性に関する条件を導き出すことを目的としたこの実験では、原理とパラメータ理論とは独立に導き出すことができる。 S 行列は本来文法セットから計算されるものであるが、本実験では逆にクレオール創発の条件を S 行列に求めるため、 $s_{ii} = 1$ ($1 \leq i \leq n$) や、 $s_{ij} = 0$ ならば $s_{ji} = 0$ ($i \neq j$) といった制約のもとに S 行列の各要素を連続的に値を変え、その条件を満たす S 行列を求める。ここで導き出される言語間の条件とは、それぞれの言語から導出される表層の類似性に関するものであり、構文規則の条件ではない。

5.4.2 実験2の初期条件とパラメータ

クレオールが創発するための最も単純なモデルとして3言語の場合を考える。すなわち2つの言語集団が接触した結果、第3の言語としてクレオールが創発する可能性がある環境である。 G_1 と G_2 を既存の言語とし、 G_3 をクレオールになり得る言語とする。人口構成比に対する影響を避けるため、既存の2言語の初期人口を等配分する。したがって初期に与える人口は $(x_1(0), x_2(0), x_3(0)) = (0.5, 0.5, 0)$ である。また、接触確率 α の値はクレオールが最も創発しやすい値をとることとした。すなわち、実験1-2から明らかのように、 $\alpha = 1$ である。この初期値は次のような状況であると考えられる。

- それぞれ異なる言語を話す2つの集団の接触であり、その人口構成比は1:1である ($x_1 = x_2 = 0.5$)。
- コミュニティで生まれた子供は、両親から完全に離れて育てられ ($\alpha = 1$)、それぞれの言語に接触する機会は均一である。
- それぞれの言語について、子供が獲得する確率は、親の言語に依存しない ($\alpha = 1$ のとき、 $q_{ij} = q_{kj}$ ($i \neq k$))。
- クレオールとなり得る言語はひとつだけである (G_3)。

これらの値は以降の実験を通じて共通である。

ここで最も単純な場合として、 S 行列を次のような対称行列とした：

$$S = \begin{pmatrix} 1 & a & b \\ a & 1 & c \\ b & c & 1 \end{pmatrix}. \quad (5.7)$$

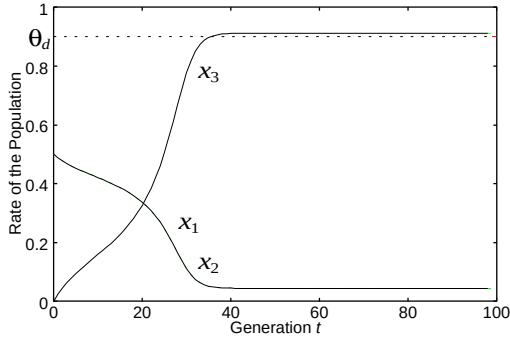
要素 $a (= s_{12} = s_{21})$ は既存の 2 言語間の類似度を表しており、 $b (= s_{13} = s_{31})$, $c (= s_{23} = s_{32})$ はそれぞれの言語とクレオールとの類似度である。

5.4.3 実験 2 - 優勢クレオールが創発する条件

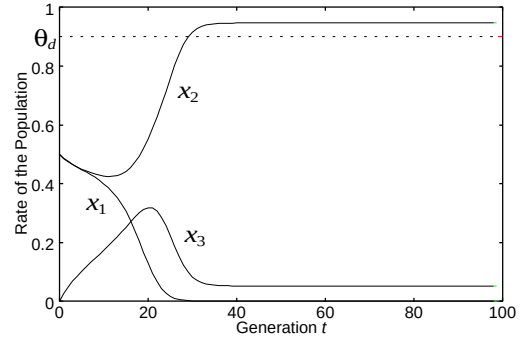
3 言語での優勢クレオールが創発した例を図 5.8 に示す。図 5.8(a) はクレオール G_3 が優勢である、すなわち $x_3(t) > \theta_d = 0.9$ であることを示している。このときの (5.7) 式の各要素の値は $(a, b, c) = (0, 0.174, 0.174)$ である。 b と c が同じ値であるため、 x_1 と x_2 の人口遷移は同じ振る舞いをする。 a の値が 0 であるということは、 G_1 と G_2 に共通の文が全く存在せず、それぞれの言語話者の間の会話は全く成立しないことを意味する。図中、初期状態で G_3 の話者が誰もいなかったが、時間が経つにつれ G_1 , G_2 話者が移行することによって G_3 話者の人口が増加し、最終的に θ_d 以上の比率を占めるようになり収束した様子を示している。このように 3 言語の場合でも実験 1-1 と同様、優勢クレオールが創発することがわかる。

b と c の値を増加させたときの結果を図 5.8(b) に示す。 S 行列の各要素の値は $(a, b, c) = (0, 0.176, 0.182)$ である。このときクレオール G_3 はその話者人口を減少させ、優勢言語は既存の言語 G_2 に変わったことを表している。また、図 5.8(c) はクレオール G_3 が最も人口が多いが、収束点においてその人口比率が優勢クレオールであるとみなす閾値 θ_d よりもわずかに低い、すなわち $x_3(t) < \theta_d = 0.9$ であることを表している。このときの S 行列の値は $(a, b, c) = (0, 0.188, 0.189)$ である。これは優勢言語が存在しない例である。

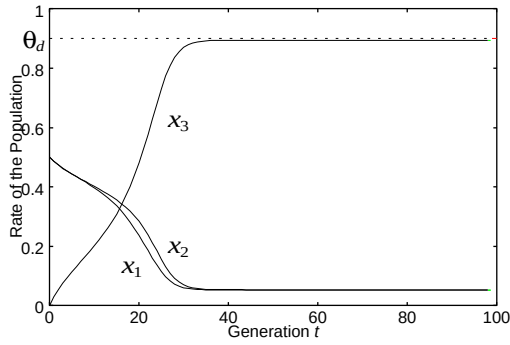
(5.7) 式における S 行列の各要素をパラメータとし、優勢クレオールが創発した値を a, b, c 空間上に表したものが図 5.9(a) である。図を見る限り平面に近く、ク



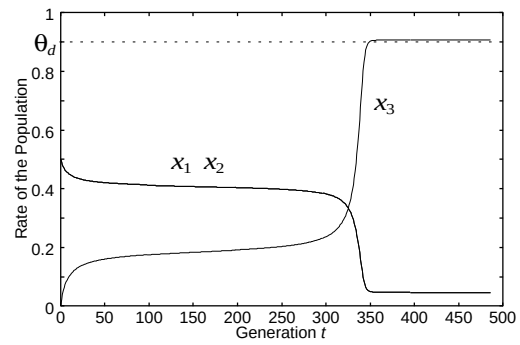
(a) $(a, b, c) = (0, 0.174, 0.174)$, Dominant, Creolized



(b) $(a, b, c) = (0, 0.176, 0.182)$, Dominant, Not-Creolized



(c) $(a, b, c) = (0, 0.188, 0.189)$, Not-Dominant, Creolized



(d) $(a, b, c) = (0.11, 0.174, 0.174)$, Dominant, Creolized

図 5.8: 3 言語での優勢クレオールの新規

レオールが創発する範囲は非常に限られていることがわかる。ここから得られる S 行列のおおよその条件は次のように表される：

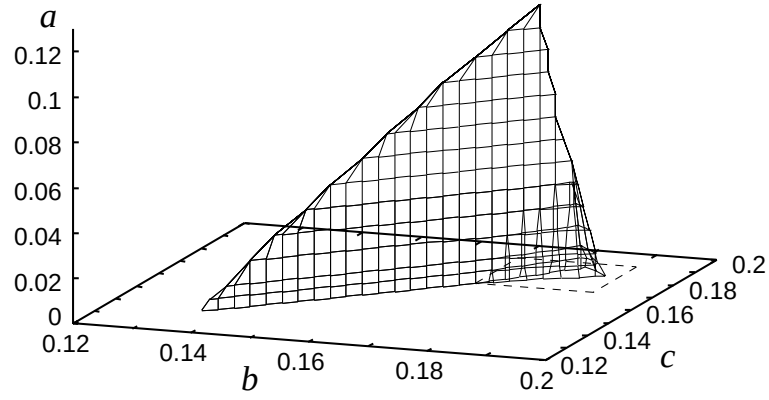
$$a \lesssim 0.12 \quad (5.8)$$

$$0.35a + 0.136 \lesssim b \simeq c \lesssim 0.2 \quad (5.9)$$

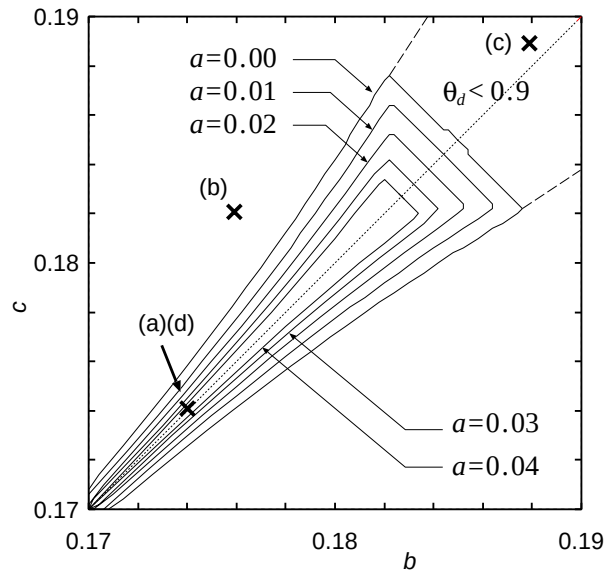
図中の破線で囲まれた部分を、 b - c 空間上に拡大したものを図 5.9(b) に示す。

図 5.8(a) と図 5.8(b) において、それぞれの言語が優勢言語となるような S 行列の値が存在することを示した。それぞれの要素の値が図 5.9(b) 中の $\times$ で示された点 (a) から点 (d) に対応している。

ここで、点 (b) から点 (a) に向けて値を変えていくと、優勢言語が G_2 から G_3



(a) a, b, c -space of the S matrix



(b) Details

図 5.9: 優勢クレオールとなるための条件 ($\theta_d = 0.9$)

に突然入れ替わる境目に遭遇する．その境界を示したのが図 5.9(b) の実線である．実線で囲まれた領域の中の値が S 行列の値として与えられると，優勢クレオールが創発する．最も外側にある実線が (5.7) 式における $a = 0$ のときの境界であり，以降， a の値が大きくなるごとにクレオールが創発する条件が限られて行くことがわかる． $a = 0$ であるということは，既存の 2 言語の話者間で全くコミュニケーションがとれない状況を表しており，このときが最も新言語が創発しやすく，共通言語となりやすい．また， a の値が大きくなるにつれ，既存の 2 言語間でコミュニケーションがとれるようになり，それぞれの言語間で人口の遷移が行われやすくなる．その結果， G_3 への遷移が妨げられるため， a が大きくなるとクレオールが創発しにくくなるのである．図 5.8(d) は図 5.8(a) の初期値から $a = 0.11$ に変更したときのモデルの振る舞いを表しており，図 5.8(a) と比較して収束までの時間が遅くなっているのがわかる．

図 5.9(b) 中，点 (a) から点 (c) に向けて値を変化させたとき， $x_3(t)$ は連続的に人口を減らし，図中の実線部分の短辺を境にして G_3 は優勢クレオールではなくなる． S 行列の要素 b, c の値が大きくなると， G_3 の人口比率が高くなった場合においても， G_1, G_2 への人口の流出が増大するため，それぞれの言語と共存し，クレオールが優勢言語にならないことが原因である．したがって，実線の長辺部分と短辺部分の境界で発生する変化の特徴は大きく異なる．図中の破線で示した境界は， G_1 または G_2 が優勢言語である領域から，優勢クレオールではないが G_3 が最も人口比率の高い言語である領域に突然変化した境目を示している．

b と c の値が大きいと， G_1 と G_2 からそれぞれ G_3 へ多くの人口が遷移する．また逆に G_3 の人口が増加すると G_1 と G_2 へ遷移するため，共存するようになる．しかし逆に b と c の値が小さいと，人口比が安定するまでの世代数が増加する． b と c の値が小さいときの優勢クレオールが創発するケースを考えると， G_3 の人口比率が G_1, G_2 のそれを追い抜くまでに多くの時間を要し，さらに $a = 0$ のときにおいて $b \simeq c \lesssim 0.136$ を境に優勢クレオールが観察できなかった．さらに a の値に比例して収束世代が遅くなり，クレオールが創発しなくなっていく様子が図 5.8(d) と図 5.9(a) をみるとわかる．

5.4.4 実験2の考察

実験2において、クレオールが創発し、優勢言語となるための、言語間の類似性に関する条件を観察した。ここで自然言語を背景に、この条件がどのような意味を含んでいるのか考察する。あるコミュニティにおいて2つの言語が存在したとき、(5.8)式の条件はこれら2言語が互いに似ていないことを表している。もしこれらが十分に類似していれば、それぞれの言語話者はコミュニケーションをとることができるため、ピジンやクレオールといった新言語が誕生する必要がないだろう。また(5.9)式は、クレオールは既存の言語と、既存の言語間の類似度よりも似ていなければならないが、あまり似すぎていてはならないことを表している。もし2言語のうちどちらかが新言語と必要以上に類似していると、クレオールは創発するが、その言語はクレオール話者とコミュニケーションがとれるため、新言語に移行せず、クレオールとその言語とで共存することになる。逆にクレオールと既存の言語とが必要以上に似ていないと、クレオールへの人口の移行に時間がかかり、クレオールは優勢言語になりにくい²。さらに既存の2言語はクレオールとの類似度もほぼ同等でなければならないと制限している ($b \simeq c$)。この均衡が崩れると、クレオールは短期間において創発するものの、話者人口は安定することなく消滅し、クレオールに類似している言語が最終的に優勢言語となる。

これらの条件から、クレオールと既存の2言語を言語空間上に配置したものを図5.10に示す。原理とパラメータ理論から考えると、これらの言語空間上の言語はあらかじめ原理によって与えられており、人間の第一言語獲得はパラメータを設定することによってこの空間上の言語を選択していることとなる。我々はクレオールの創発を観察した結果、この言語空間の言語の配置を導き出したと考えることができる。

パラメータ空間の全体を見ると、優勢言語とクレオールの創発に関して次の4つの領域に分類できることがわかった。

i) 図5.8(a)のように、クレオールが優勢言語となる ($x_3(t) \geq \theta_d = 0.9$)。

ii) 図5.8(b)のように、クレオールが創発しない ($x_3(t) < \theta_c = 0.1$) が優勢言

²ここでいう言語の類似性とは2言語間の文法の類似性について言及しているわけではないため、クレオール文法とは既存の言語の文法が混ざりあった言語ではないという、実際のクレオールの観察結果に矛盾するものではない。

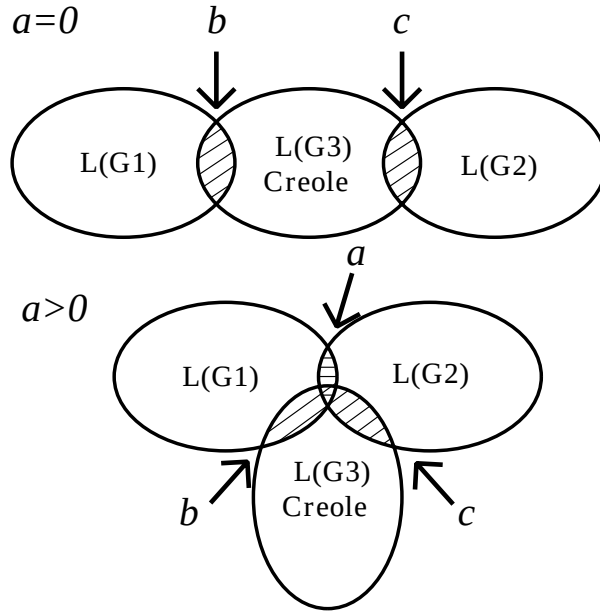


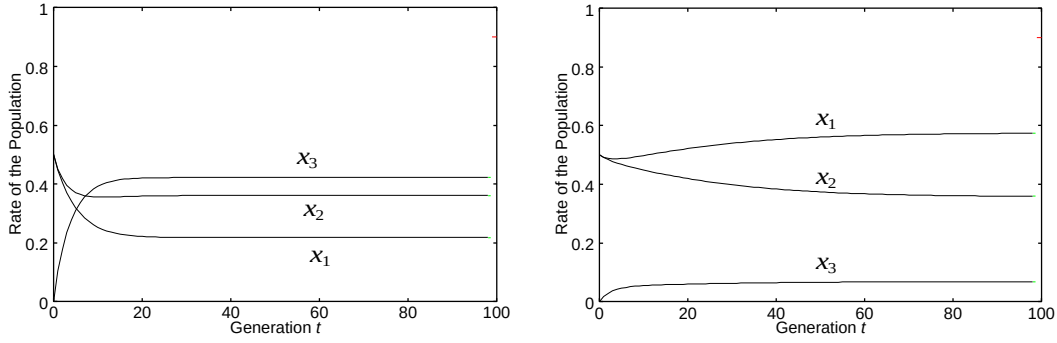
図 5.10: 言語空間上のクレオール創発の条件

語が存在する.

- iii) 図 5.11(a) のように, クレオールが創発するが, 他の言語と共存する ($\theta_c \leq x_3(t) < \theta_d$).
- iv) 図 5.11(b) のように, クレオールが十分な話者を持たず, ($x_3(t) < \theta_c$), 既存の言語が共存する.

実験において, 上記 iii) および iv) の共存カテゴリーが現れるパラメータ領域はおおよそ $a, b, c \gtrsim 0.3$ であった. この値はそれぞれの言語は相対的に類似度が高く, またそれぞれの言語話者は他の言語話者と互いにコミュニケーションをとることができることを表している. この状況から考えると, これらの言語は互いに異なる独立した言語であるとみなすよりも, むしろ方言であると考えた方が良いように思われる. 一般的に言語学的な見方をすると, ある言語が他のある言語の方言であるのか, あるいは異なる言語であるのかということを明白に区別することはできない³. 我々の実験の結果から, 0.3 という類似度がこれらを区別するお

³このような境界線はしばしば政治的に設けられる. 例えばボスニア・ヘルツェゴビナのセルビア語, クロアチア語, ボスニア語などがそれにあたる [14].



(a) $(a, b, c) = (0.3, 0.4, 0.5)$, Coexistent,
Creolized

(b) $(a, b, c) = (0.35, 0.25, 0.1)$, Coexis-
tent, Not-Creolized

図 5.11: 言語の共存

およその基準になると考えることもできる。

5.5 第 5 章のまとめ

言語学の分野におけるクレオールの主な調査対象は、これまでに発見されてきたクレオール間に見られる言語としての類似性と、そこから想像される普遍文法が存在する可能性についてである。ここで我々は普遍文法を仮定し、構成論的手法によってクレオールが創発するための条件を数理的に導き出すことを目的として実験を行った。言語の定義を各言語間の類似性から、クレ奥ールの定義を言語話者の人口比率からそれぞれ与え、これまでの言語学的な視点と全く異なるアプローチで調査を行った。

我々は Komarova et al. [23] および Nowak et al. [33, 34] の言語動力学モデルをより現実的なモデルに改良するため、動的遷移行列モデルを提案した。このとき、現実的であるための条件として次の点を主張した。

- 子供の言語獲得は周りの言語話者が話す言語に影響を受ける。
- 親としか会話を行わなかったとき、子供は親の言語を身につける。

このモデルでは、言語間の遷移を特徴づけるパラメータが、人口比率に依存するので、遷移行列が動的に変化するものとなる。また、ピジン化が既に発生してい

る多言語コミュニティであると仮定した。

このモデルの計算機実験による分析を通じて、以下のことが明らかになった。

- 1) 初期の人口構成比が最終的な優勢言語に影響を及ぼす。
- 2) 言語学習者が、どの言語から入力となる文を受け取るかが最終的な優勢言語に影響を及ぼす。
- 3) クレオールが存在するための、言語間の類似性の条件が存在する。

このうち、3) のクレオールが創発するための条件は次のようなものであった。

- 既存の言語は互いに類似しすぎていないこと。
- 既存のある言語から見て、他のどの既存の言語よりもクレオール言語がその言語と似ていること。
- 既存の言語とクレオールとの類似度は、各既存言語で同程度であること。

普遍文法を仮定した場合、人間が獲得することができる言語は、クレオールも含め、あらかじめ与えられていると考えられている。図 5.10 はその言語間関係をクレオールの創発条件を導いた結果から描くことができたものである。この結果は言語学の分野への大きな貢献であると考えられる。

我々は言語と方言の関係についても述べた。しかし言語と方言の違いは文法的な素性が密接に関係しているため、言語の類似性からこれらを区別することはできない。またクレオールに関しても同じことが言える。これは現在の人口動力学が抱える重大な問題であり、人口動力学の中に何かしらの言語的素性を持たせることが次のモデルの課題となる。より現実的で一般的なクレオールの創発条件を明らかにするために、さらなる研究が必要である。

第 6 章

考察

本章ではまず最初に、これまで述べてきた本研究の結果を踏まえ、第 2 章で論じたピジン・クレオール研究と言語進化研究における課題を中心に考察を行う。次に、第 3 章で扱ったマルチエージェント・モデルと、第 5 章の人口動力学モデルについて、それぞれの問題点と今後の課題について述べる。

6.1 関連研究との比較・検討

本研究の目的は、現実世界で発生しているピジンやクレオールという言語現象について、それらが創発するための条件を数理的に導き出すというものであった。2.4 節で示したように、本稿で主張したことは次の 2 つについてであった。

- 言語変化の本質は、子供の言語獲得と、それを含めた周辺環境の変化によって言語獲得に与える影響との相互作用である。
- 言語獲得と言語変化の関係を調査するにあたり、実際に発生した現象と比較検証できる研究対象はピジンとクレオールの創発についてである。

これらを基本理念として、3つのモデルを提案し、実験を行った。第 3 章から第 5 章にかけて行ったそれぞれの実験によって、当初の目的がどれだけ達成できただろうか。2.1 節および 2.2 節で言及したピジン・クレオール研究と言語進化研究について、それぞれ考察をする。

6.1.1 ピジン・クレオール研究に対する考察

本稿では2.1節において、本研究の題材となったピジン・クレオールについての説明と、そこで議論されている問題について述べた。ピジンとクレオールは言語変化のプロセスにおける各段階として密接な関係を持っているが、これらが創発する要因は異なる。ピジンは、社会的な環境の変化によって、異なる言語を話す人々の間でコミュニケーションをとる必要性から開発された簡易言語である。したがって、人間が試行錯誤をしながら自らがもともと持っていた文法を改編し、第二言語としてピジンを獲得していく過程をモデル化する必要があった。

それに対してクレオールは、ピジン話者の子供がその第一言語獲得過程において、満足以親の言語を聞くことができず、周りに獲得すべき言語がないときに人間がもともと持っている生得的な言語が発現する。したがって、言語獲得のメカニズムとして普遍文法を仮定し、親がもともと持っていた言語から遷移する過程をモデル化することが、クレオール化の条件を求めるために必要であった。

以上を理由に、3つのモデルを提案した。これらはピジン化とクレオール化を扱ったモデルとして2つに分けられる。それぞれについて以下に述べる。第3章においてはピジン化の過程を観察する実験を行った。このモデルは、異なる言語話者集団の接触の初期段階、すなわちジャーゴンもしくは限定ピジンへ移行する過程をマルチエージェントの枠組でモデル化したものである。エージェントの発話内容に意味を持たせ、より多くのエージェント、特に異なる言語話者集団とのコミュニケーションに対して報酬を与え、文法の変化を促した。その結果、それぞれの言語で使用されていた単語の混合、素性の簡略化がみられ、現実のピジンに見られるものと同じような現象が観察された。

また、第4章および第5章では、クレオール化の過程を観察した。普遍文法を仮定したため、原理とパラメータ理論によって人間が話することができる言語の数は有限である。ここから、初期段階において誰も話していなかったが、世代を経て一定の割合で話者人口を獲得した言語をクレオールと定義した。第4章で考察した結果、クレ奥ールの創発には言語獲得期における周りの言語話者の人口比率に大きく依存することを示した。これに基づき、第5章では、言語獲得についてより現実的な設定を考慮することにより言語動力学モデルを修正し、動的遷移行列モデルを提案した。これにより、クレ奥ールの創発現象を観察し、より現実的な

モデルであることを示した。また、クレオール創発は言語集団の初期人口、親との接触確率、および言語間の類似性に依存することを示した。

2つに分けたこれらの実験によって、ピジン化、クレオール化のそれぞれを計算機上で再現することができた。これは、クレオール化のためには必ずしもピジン化のプロセスが必要なのではなく、言語獲得の臨界期において言語を習得することが困難な状況に陥れば、クレオールが創発する可能性があることを示唆している。これについて Bickerton も、英語を話す両親から生まれた子供が、2歳ぐらいで英語を話す社会から隔離されると、大きくなってから話す言語は、語彙は主に英語から取り入れているが、文法はクレオールであるような言語であろうと予測している [6]。しかし、第5章の実験から、クレオールが創発するための言語間の類似度について言及したが、実際には異なる2つの言語間では、語彙の違いなどを考慮すると全くといっていいほど似ていないだろう。ピジン化のプロセスには、異なる言語を持つ集団間で使用される言語の類似性を高める作用があり、本モデルにおいても、既にピジンが浸透した段階であると仮定して各言語間の類似度を設定した。しかし、ピジン化によって、どれほど言語間の類似度を高めることができるかということに関しては言及していない。

6.1.2 言語進化研究に対する考察

これまでの多くの言語獲得モデルにおいて、エージェントは次のような行動を取るように設計されている。

- エージェントは自分自身が持つ文法を元にメッセージを生成し、他のエージェントに発話する。
- 他のエージェントから発生されたメッセージを聞き、自身の文法で理解する。
- 他のエージェントのメッセージから、その文法を評価し、学習する。
- エージェントの文法の理解に応じて報酬を受け取り、その比率によって自分の文法を持った子供を次の世代に残す。

第3章および第4章においてもこれらに基づいたモデルを提案した。言語進化の問題に限らないが、計算機でシミュレーションを行う場合、まず最初に何を仮定

し、何を求めることを目的とするのかを決定しなければならない。第 3 章の目的は、既存の人工知能および自然言語の技術を組み合わせて人工ピジンを生成するとともに、そのピジン化に至るための条件を求めることであった。ここでは文法に LTAG、学習機構に GA を用いたが、これが人間の学習メカニズムと等価であるとはみなすことには疑問が感じられる。上に挙げた項目は、実際の言語獲得の環境を模倣するために、最も重要である共通の要素を指しており、これがモデルの有効性を示す指標ではない。したがって、エージェントが持つ文法構造として LTAG を用いたことに関しては新規性が見られるが、言語獲得モデルとして考えると、特に注目する点はない。第 4 章についても同様のことがいえる。原理とパラメータ理論を導入することにより、ピジン化のモデルと比較して、より現実的なモデルとなったが、原理として与えた文法セットとパラメータが作用するルールについて資意的な設定を行ったことは否めない。また、EM アルゴリズムを採用することにより、文法獲得の効率を追求しようと試みたが、原理とパラメータ理論を採用したモデルにおいては最尤推定で十分であると考えられる。ただし、このモデルに関しては、後に述べる第 5 章のモデルへの足掛かりになったとして、重要なモデルとなっている。

第 5 章では、人口動力学を用いて数理的にクレオールが創発するための条件を導き出した。言語進化の研究において、2.2.3 節で説明したさまざまな進化の相互作用を観察することについてどの方法論が有効であるのかは、まだ明らかではないということが問題として挙げられる [22]。人口動力学を始めとする数学モデルは、マルチエージェント・モデルで問題とされる計算量の問題を解消し、さらに動的なシステムと言語の振る舞いによる相互作用を概観するための有効な手段であると考えられているが、これらはまだ実用的なレベルに達していない。このため、既存の言語動力学をより現実的なものに修正したモデルの提案は、言語進化研究にとって大きな貢献であると考えられる。特に、言語獲得に関する問題であるのに、使用される言語やそれに基づく文法を明示せず、言語間の相対的な類似度による情報からその類似性に関する条件を導き出したのは、これまでにない斬新な手法であると考えられる。

6.2 本研究の問題点と今後の課題

本研究の問題点と今後の課題について、それぞれの実験ごとに述べる。

6.2.1 マルチエージェント・モデルの問題点と今後の課題

第3章ではマルチエージェントの枠組で、ピジンの創発現象を観察したが、ここで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。例えば本モデル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、2.2.2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、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。本モデルにおいてはこの両方を満たすべく単純な学習方法を採択したが、これらの改善については次のモデルへの課題とする。また、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。本モデル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。

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、実験で得られた結果との比較検証をすることが挙げられる。本モデルは現実のピ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。しかし我々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を想定している．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） [46]．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 [5]．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 3.1.2 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 [37, 40]．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これと予測される．

第 4 章に関しても，同様の問題が挙げられるが，これを解消するかたちとして人口動力学モデルを提案するに至った．

6.2.2 人口動力学モデルの問題点と今後の課題

第 5 章では，人口動力学に基づいた言語動力学をより現実に近いモデルに修正することにより，動的遷移行列モデルを提案した．本モデルの問題点としてまず

最初に挙げられるのが、言語の表現についてである。このモデルにおける言語表現は、言語間の相対的な類似度を確率行列で表したものである。したがって、それぞれの言語が持っている文法的な特徴は、このモデルである数式の中で表現することはできない。実験においては、この問題を逆手に取って、クレオールが創発するための言語間の類似性を求めたが、これに対応するクレオールの文法を予測することは、現実的に不可能である。したがって、言語に関する情報を、動力学の中でより詳しく記述することが今後の課題となる。人口動力学モデルは、それに対応するマルチエージェント・モデルへ置き換えることが可能である。したがって、複雑な文法を持ったエージェントが他のエージェントと会話をするといったモデルを構築することで、人口動力学モデルとの振る舞いを比較し、クレオールと文法との関連について議論することができる可能性がある。

次に、動的遷移行列モデルに含まれている、各言語に与えられる適応度について言及する。このモデルでは、動的遷移行列と各言語話者の人口比に加え、適応度に依存して各言語話者の次世代の人口比率が決まる。しかし、クレオールを話す人が、クレオールを話さない人たちに比べて、より多くの子供を産んだからクレオールが優勢言語になったのではなくて、バイオプログラムのような子供の学習バイアスの影響によって、クレオールが創発すると考えたほうが、より現実的である。したがって、(5.2) 式で与えられている適応度 $f_i(t)$ の項を削除することを考える。これにより、 $\phi(t)$ の項も削除できるため、適応度を考慮しないモデルは次の式として与えられる。

$$\frac{dx_j(t)}{dt} = \sum_{i=1}^n \bar{q}_{ij}(t)x_i(t) - x_j(t) \quad (j = 1, \dots, n). \quad (6.1)$$

このモデルを用いた予備実験を実際に行っている。(5.2) 式と (6.1) 式の系全体の振る舞いとしての大きな違いは、適応度の項目を削除したモデルのほうが、クレオールが創発する接触確率の範囲が、著しく広いことである。第 5 章の実験においては、クレオールが創発するかしないかを定める、接触確率の閾値があった。言い替えれば、クレオールの創発は、子供が他の言語と接触する確率に依存した結果が得られた。対称的に、適応度がその役目を果たさないとき、クレオールの人口比率に差はあるが、接触確率によらずに全ての値でクレオールは現れる。つまり、言語の適応度というのは、子供があまり他の言語と接触しないとき、クレ

オールの新発を抑制する働きがあるのかもしれない。今後はこの関係について、さらに調査する必要がある。

第 7 章

おわりに

本論文では、言語進化シミュレーションにおけるピジンとクレオールの創発について述べた。

まず、第 2 章では、特にピジン、クレオールや言語進化に関するこれまでの言語研究を概観し、本稿における立場を明らかにした。これまでの長い歴史の中で、ピジンやクレオールは、劣った、でたらめの、不完全で価値のないものだと考えられてきた [40]。しかし、現在では人間の言語獲得のメカニズムの解明に欠かせない研究分野である。

ここでは、ピジン化からクレオール化に至るプロセスを説明し、実際の発話例からピジンとクレオールそれぞれの文法的特徴を示した。また、言語獲得の生得性とクレオールとの関連について説明した。ピジンは社会的な環境の変化によって築かれた多言語コミュニティにおいて開発された共通言語であるのに対し、クレオールはピジン話者から生まれた子供が獲得した言語である。したがって、ピジンは目的言語が不明瞭である状態での第二言語獲得、クレオールはピジン・コミュニティという特殊な環境下における第一言語獲得であると考えられる。すなわち、クレオールについて研究することは言語獲得問題や、言語進化の解明に大きく貢献すると考えられている。

人間の言語の起源とその進化に関する疑問に答えることは、古くから、多くの言語研究者にとって大きな課題である。紀元前 7 世紀のエジプトのファラオであったプサメチク I 世 (Psamtik I) によって、人類最初の言葉を探るべく、幼児を隔離して実際に実験を行ったという、古代ギリシャの歴史家ヘロドトス (Herodotus)

の記録すら残っている¹ [6]. また、近代においても、言語の起源と進化に関する理論がいくつも提唱されたが、それらは必然的に推測の域を出ないままにされていた。なぜなら言語の起源も進化の過程も観察することが不可能であることから、実証的証拠がなく、その理論を立証することも反証することもできなかったためである。そのため、19世紀のフランスの言語学会 *Société Linguistique de Paris* では、これらに関する研究と出版を禁止したほどである [11].

この言語進化に関する研究は、それ以降の言語学、進化学、認知科学などの発展と、計算機の発達、特に人工生命の分野からのアプローチ [55] により、大きく進展している。ここでは、言語進化研究に関するこれまでの研究成果について説明し、本研究との関連を述べた。ピジンやクレオールは、人間の生物学的進化を伴わない言語の変化である。現在では、言語の発達、変化に関して実際に観察可能な現象は、ピジン、クレオール、もしくは子供の言語獲得に限られてしまうため、その実地検証との比較ということに関して、大きな優位性を持っている。

これらを基に、ピジン、クレオール研究の立場から言語進化研究に求められていることを述べ、実験計画を示した。

次に、第3章 マルチエージェント環境での人工ピジンの生成 では、計算言語学の研究成果である LTAG (Lexicalized Tree Adjoining Grammar) の文法理論を用いたマルチエージェント・モデルを提案した。本モデルでは、日本語を話すエージェント群と、英語を話すエージェント群からなるコミュニティを仮定した。これらの接触により、エージェントが共通言語であるピジンを生成する過程について述べた。学習機構に遺伝的アルゴリズム (Genetic Algorithm; GA) を用い、各遺伝子に対応する語彙ルールを変形させることによって、演繹的な学習を行った。実験の結果、それぞれのエージェント群の人口に大きな差がある場合は、少数のエージェント群が多数のエージェント群の言葉を獲得したのに対し、それぞれの人口が対等である場合、語彙や文法が混ざりあった混合言語が話されるようになった。さらに、学習対象となる言語の獲得に関して報酬に差を設けた場合、その報酬によってピジンの発現に対して差が現れた。このモデルは、上層言語と下層言語の差をなくし、双方の立場が対等である場面において、人口構成比と言語獲得

¹ちなみにこの実験による結論は、子供たちが発した音声の中で最初に理解できたものは “bekos” であり、古代フリギア語 (Phrygian) の「パン」に相当する単語であった。そこでプサメチクは、人類の最初の言葉はフリギア語であると主張したということである。

の重要性をパラメータにしてピジン化の条件を求める実験であった。

また、第4章 マルチエージェントによる人工クレオール生成モデルとその問題点では、自然言語処理の分野において文法獲得アルゴリズムとして用いられるEM アルゴリズムを学習機構として備えたマルチエージェント・モデルを提案した。結果として、エージェントが持つパラメータの値の多数決によって、次世代の文法が決定されるという非常に単純な結果が得られたが、このモデルによって得られたことが、後述の言語動力学に活用されることとなった。

これら一連のマルチエージェントによるモデルにおいて、この目的の前に、既存の自然言語処理技術を組み合わせることによってピジンやクレオールを発生させることができるか、また、人工的に生成された言語に関して、どのような特徴についてピジンまたはクレオールと定義するのかという疑問を解消する必要があった。結論としては、実験において混合言語が話され、現実のピジンの発話例と比較することにより、ピジンが創発したことを言及した。また、有限の文法セットからなる言語空間において、初期においてどのエージェントも所有していなかった言語が、次世代において優勢な言語となったものをクレオールと定義した。しかし、人間の学習機構との相違や、取り扱う言語の素性に関する拡張性の問題など、モデルの設計自体に大きな問題が発生した。マルチエージェント・モデルを提案する場合、エージェントの機能は自由に与えることが可能であり、またそれがその大きな特徴であることから、ピジン創発のための条件を導出しようとしても、エージェント固有の問題となる危険性を持っている。また、エージェントによる人口構成比の条件を導き出す場合、必然的に離散的になってしまい、さらに計算量の問題からエージェントの構成を十分に配分することが難しいため、一般的な法則を導出しにくい。ここで得られた結論が、次章の人口動力学モデルを導入するきっかけとなった。

第5章 言語動力学におけるクレオールの創発では、言語に関する人口動力学を用いることにより、クレオールの創発について言及した。この創発は、対象となるコミュニティにおける、各言語話者の人口構成比、言語獲得の臨界期での親以外の言語話者との接触、そして言語間の類似性に依存することを示した。これまでのクレオール研究は、さまざまなクレオール間の特性が類似していることが明らかにされているものの、人口動力学や、既存の言語と創発したクレオールの間

にある類似性の視点からは扱われていなかった。第 5 章の実験における我々の貢献は、ある多言語コミュニティにおけるクレオール化の可能性について、人口構成比、接触状況、言語間の類似性から予測を立てることが可能であることである。この予測は現実のさまざまなクレオールから、その文法や上層、下層言語との関係を詳しく調査し、検証する必要がある。

また、この言語動力学から言語と方言についても言及した。しかし、言語と方言の違いは文法の違いに関わるため、これらを実際に使用されている言語間の類似性から区別することは不可能であり、ましてやクレオールに関しても同じことがいえる。これは、現在の人口動力学においてとても重要な問題であり、文法の差異を扱う上での限界である。それゆえ、言語的な素性を人口動力学上に表現することが今後の課題となる。

最後に、第 6 章において、これまで述べてきた本研究の結果を踏まえ、第 2 章で提起した問題を中心に議論を行い、本研究の意義を明確にした。さらに、本研究では達成できなかったことについて、本研究の問題点と今後の課題としてまとめた。

謝辞

本研究を行なうに当たり、大変有益なご助言とご指導を頂きました北陸先端科学技術大学院大学情報科学研究科の東条 敏教授，鳥澤 健太郎助教授，永田 裕一助手，知識科学研究科の橋本 敬助教授に深く感謝致します。

東条教授には指導教官として，本研究を進めるにあたり，全般に渡って大変丁寧なご指導とご助言を頂きました。鳥澤助教授には研究者としての基本的な考え方をご教示いただきました。永田助手には研究室ゼミに限らず適切なご助言をいただきました。また，橋本助教授には副テーマのご指導にとどまらず，それ以降の研究において進化言語学に関するご助言および議論をしていただきました。

最後に，本論文をまとめるに当たって御協力いただいた北陸先端科学技術大学院大学の東条・鳥澤研究室の諸兄に厚く御礼申し上げます。

参考文献

- [1] Anne Abeillé and Owen Rambow, editors. *Tree Adjoining Grammars*. CSLI Publications, California, 2000.
- [2] John Archibald, editor. *Second Language Acquisition and Linguistic Theory*. Blackwell Publishers, Massachusetts, 2000.
- [3] Jacques Arends, Pieter Muysken, and Norval Smith, editors. *Pidgins and Creoles*. John Benjamins Publishing Co., Amsterdam, 1994.
- [4] D. Bickerton. The language bioprogram hypothesis. *Behavioral and Brain Sciences*, Vol. 7, No. 2, pp. 173–222, 1984.
- [5] D. Bickerton. *Language and Species*. University of Chicago Press, 1990.
- [6] D. Bickerton. Creole languages. *Scientific American*, pp. 108–115, July, 1983.
- [7] Derek Bickerton. *Roots of language*. Karoma, Ann Arbor, MI, 1981.
- [8] E. J. Briscoe. Grammatical acquisition and linguistic selection. In Ted Briscoe, editor, *Linguistic Evolution through Language Acquisition: Formal and Computational Models*, chapter 9. Cambridge University Press, 2002.
- [9] Bertram Bruce. Case systems for natural language. *Artificial Intelligence*, Vol. 6, No. 4, pp. 327–360, 1975.
- [10] A. Cangelosi and D. Parisi, editors. *Simulating the Evolution of Language*. Springer, London, 2002.
- [11] Angelo Cangelosi and Domenico Parisi. Computer simulation: A new scientific approach to the study of language evolution. In A. Cangelosi and

- D. Parisi, editors, *Simulating the Evolution of Language*, chapter 1. Springer, London, 2001.
- [12] Noam Chomsky. *Reflections on Language*. Pantheon, New York, 1975.
 - [13] Noam Chomsky. *Lectures on Government and Binding*. Foris, Dordrecht, The Netherlands, 1981.
 - [14] Bernard Comrie, Stephen Matthews, and Maria Polinsky, editors. *The Atlas of Languages*. Quatro Publishing, London, 1996.
 - [15] Michel DeGraff, editor. *Language Creation and Language Change*. The MIT Press, Cambridge, MA, 1999.
 - [16] E. Gibson and K. Wexler. Triggers. *Linguistic Inquiry*, Vol. 25, No. 3, pp. 407–454, 1994.
 - [17] E. M. Gold. Language identification in the limit. *Information and Control*, Vol. 10, pp. 447–474, 1967.
 - [18] T. Hashimoto and T. Ikegami. Evolution of symbolic grammar systems. In Federico Morán, Alvaro Moreno, J. J. Merelo, and Pablo Chacón, editors, *ECAL95*, Vol. 929 of *Lecture Notes in Computer Science*, pp. 812–823. Springer, 1995.
 - [19] Takashi Hashimoto. *Evolution of Code and Communication in Dynamical Networks*. PhD thesis, University of Tokyo, March 1996.
 - [20] Takashi Hashimoto. The constructive approach to the dynamic view of language. In A. Cangelosi and D. Parisi, editors, *Simulating the Evolution of Language*, chapter 14. Springer, London, 2001.
 - [21] Simon Kirby. Learning, bottlenecks and the evolution of recursive syntax. In Ted Briscoe, editor, *Linguistic Evolution through Language Acquisition: Formal and Computational Models*, chapter 6. Cambridge University Press, 2002.

- [22] Simon Kirby and James R. Hurford. The emergence of linguistic structure: An overview of the iterated learning model. In A. Cangelosi and D. Parisi, editors, *Simulating the Evolution of Language*, chapter 6. Springer, London, 2001.
- [23] N. L. Komarova, P. Niyogi, and M. A. Nowak. The evolutionary dynamics of grammar acquisition. *Journal of Theoretical Biology*, Vol. 209, No. 1, pp. 43–59, 2001.
- [24] Natalia L. Komarova and Martin A. Nowak. Population dynamics of grammar acquisition. In A. Cangelosi and D. Parisi, editors, *Simulating the Evolution of Language*, chapter 7. Springer, London, 2001.
- [25] K. Lari and S. J. Young. The estimation of stochastic context-free grammars using the inside-outside algorithm. *Computer Speech and Language*, Vol. 4, pp. 35–56, 1990.
- [26] E. H. Lenneberg. *Biological Foundations of Language*. John Wiley & Sons, Inc., New York, 1967.
- [27] C. D. Manning and H. Schütze. *Foundations of Statistical Natural Language Processing*. The MIT Press, London, 1999.
- [28] Makoto Nakamura, Takashi Hashimoto, and Satoshi Tojo. Creole viewed from population dynamics. In *Proceedings of the Workshop/Course on Language Evolution and Computation in ESSLLI*, pp. 95–104, 2003.
- [29] Makoto Nakamura, Takashi Hashimoto, and Satoshi Tojo. The language dynamics equations of population-based transition – a scenario for creolization. In H. R. Arabnia, editor, *Proceedings of the International Conference on Artificial Intelligence (IC-AI’03)*, pp. 689–695. CSREA Press, 2003.
- [30] Makoto Nakamura and Satoshi Tojo. The emergence of artificial creole by the em algorithm. In Steffen Lange, Ken Satoh, and Carl H. Smith, editors,

- Discovery Science*, Vol. 2534 of *Lecture Notes in Computer Science*, pp. 374–381. Springer, 2002.
- [31] Sergei Nirenburg, editor. *Machine Translation*. Cambridge University Press, Cambridge, 1987.
 - [32] Partha Niyogi. *The Informational Complexity of Learning from Examples*. PhD thesis, Massachusetts Institute of Technology, , 1996.
 - [33] M. A. Nowak and N. L. Komarova. Towards an evolutionary theory of language. *Trends in Cognitive Sciences*, Vol. 5, No. 7, pp. 288–295, 2001.
 - [34] M. A. Nowak, N. L. Komarova, and P. Niyogi. Evolution of universal grammar. *Science*, Vol. 291, pp. 114–118, 2001.
 - [35] Tetsuo Ono, Satoshi Tojo, and Satoshi Sato. Common language acquisition by multi-agents. In *International Computer Symposium (ICS’96), Proceedings on Artificial Intelligence*, pp. 218–223, 1996.
 - [36] F. Pereira and Y. Schabes. Inside-outside reestimation from partially bracketed corpora. In *In Proceedings of the 30th Annual Meeting of the Association for Computational Linguistics (ACL’92)*, pp. 128–135, University of Delaware, 1992.
 - [37] S. Pinker. *The Language Instinct: How the Mind Creates Language*. William Morrow and Company, New York, 1994.
 - [38] T. Research. A lexicalized tree adjoining grammar for english, 1995.
 - [39] S. J. Roberts. The role of diffusion in the genesis of hawaiian creole. *Language*, Vol. 74, No. 1, pp. 1–39, 1998.
 - [40] Loreto Todd. *Pidgins and Creoles*. Routledge & Kegan Paul Ltd, London, 1974.
 - [41] J.W. Weibull. *Evolutionary Game Theory*. The MIT Press, Cambridge, MA, 1995.

- [42] G. Werner and M. Dyer. Evolution of communication in artificial organisms.
In C. Langton, C. Taylor, D. Farmer, and S. Rasmussen, editors, *Artificial Life II*, pp. 659–687, Redwood City, CA, 1992. Addison-Wesley Pub.
- [43] 風間喜代三, 長谷川欣也 (編). 言語学百科事典. 大修館書店, 東京, 1992.
- [44] 守田健一. ピジン語小辞典. 泰流社, 東京, 1990.
- [45] 北研二. 確率的言語モデル. 東京大学出版会, 東京, 1999.
- [46] 亀井孝, 河野六郎, 千野栄一 (編). 言語学大辞典. 三省堂, 東京, 1996.
- [47] 掛川淳一, 神田久幸, 藤岡英太郎, 伊丹誠, 伊藤紘二. 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, pp. 693–701, 2000.
- [48] 長尾真 (編). 自然言語処理. 岩波書店, 東京, 1996.
- [49] 中村誠, 東条敏. マルチエージェント環境での人工ピジンの生成. 認知科学, Vol. 10, No. 2, pp. 193–206, 2003.
- [50] 伊庭斉志. 遺伝的アルゴリズムの基礎. オーム社, 東京, 1994.
- [51] 安藤貞雄, 小野隆啓 (編). 生成文法用語辞典. 大修館書店, 東京, 1993.
- [52] 小野哲雄, 東条敏. 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, Vol. 13, No. 4, pp. 546–559, 1998.
- [53] 東条敏. 自然言語処理入門. 近代科学社, 東京, 1988.
- [54] 田中穂積, 辻井潤一 (編). 自然言語理解. オーム社, 東京, 1988.
- [55] 有田隆也. 人工生命. 科学技術出版, 東京, 2000.
- [56] 井上和子, 原田かづ子, 阿部泰明. 生成言語学入門. 大修館書店, 東京, 1999.

本研究に関する発表論文

- [1] Nakamura, M., Hashimoto, T. and Tojo, S. (2004). The Effect of Fitness in the Emergence of Creole (abstract), *Proceedings of 5th International Conference on Evolution of Language*, Leipzig, (under submission).
- [2] 中村 誠, 橋本 敬, 東条 敏 (2004). 言語動力学におけるクレオールの創発, 認知科学, (投稿中).
- [3] Nakamura, M., Hashimoto, T. and Tojo, S. (2003). Creole Viewed from Population Dynamics, *Proceedings of Language Evolution and Computation Workshop/Course in 15th European Summer School in Logic Language and Computation (ESSLLI2003)*, Vienna, 95–104.
- [4] Nakamura, M., Hashimoto, T. and Tojo, S. (2003). The Language Dynamics Equations of Population-Based Transition – A Scenario for Creolization –, In H. R. Arabnia(Ed.), *Proceedings of the International Conference on Artificial Intelligence (IC-AI'03)*, Las Vegas, CSREA, 689–695.
- [5] 中村 誠, 東条 敏 (2003). マルチエージェント環境での人工ピジンの生成, 認知科学, **10**(2), 193-206.
- [6] Nakamura, M. and Tojo, S. (2002). The Emergence of Artificial Creole by the EM Algorithm, In S. Lange(Eds.), *Proceedings of the 5th International Conference on Discovery Science (DS2002)*, Lübeck, LNCS, Springer, 374–381.
- [7] 中村 誠, 東条 敏 (2002). EM アルゴリズムによる人工クレオールの創発, エージェント合同シンポジウム (JAWS2002) 講演論文集, 71–78.

- [8] 中村 誠, 東条 敏 (1997). マルチエージェント・モデルによる共通文法の獲得, 第 28 回人工知能基礎論研究会資料, 30–35.

第 A 章

パラメータに対応する言語群

Gibson et al. [16] が提案し, Niyogi [32] のモデルで用いられた 8 つの言語 (表 A.1) をここで紹介する.

これら 8 つの言語を導出する文法は, 3 つのパラメータから求められる. 3 つのうち, 2 つが X-バー理論に関するものである. X-バー理論とは, 原理とパラメータ理論における原理に相当する理論のひとつであり, 親のカテゴリに子の情報を残すように句構造規則を制限したものである. X には N や V 等のカテゴリが当てはまる [13]. 2 つのルールは, 句構造規則における Specifier-Head 関係および Head-Complement 関係に相当する. 以下のプロダクションルールは, これら両パラメータと一致する:

$$\begin{aligned}XP &\rightarrow \textit{Spec} \ X' \ (p_1 = 0) \quad \text{or} \quad X' \ \textit{Spec} \ (p_1 = 1), \\X' &\rightarrow \textit{Comp} \ X' \ (p_2 = 0) \quad \text{or} \quad X' \ \textit{Comp} \ (p_2 = 1), \\X' &\rightarrow X.\end{aligned}$$

3 番目のパラメータは動詞の移動 (Verb Movement) に関するものである. これはドイツ語やオランダ語等の平叙文にみられる, 動詞が文中の 2 番目の位置に移動するか否かを指定するパラメータである. この, 動詞が常に 2 番目に移動するルールは, 世界中の言語の中で現れるものと現れないものがあり, この多様性を V2 パラメータとして捉えている. 上記 3 つのパラメータによって得られた, 各文法ごとの埋め込みがない (Degree-0) 言語 ($L(G_1), \dots, L(G_8)$) の一覧を表 A.1 に示す.

表 A.1: 言語群 (Gibson et al. [16], Niyogi [32] から引用)

Language	Spec	Comp	V2	Degree-0 unembedded sentences
$L(G_1)$	1	1	0	“V S” “V O S” “V O1 O2 S” “Aux V S” “Aux V O S” “Aux V O1 O2 S” “Adv V S” “Adv V O S” “Adv V O1 O2 S” “Adv Aux V S” “Adv Aux V O S” “Adv Aux V O1 O2 S”
$L(G_2)$	1	1	1	“S V” “S V O” “O V S” “S V O1 O2” “O1 V O2 S” “O2 V O1 S” “S Aux V” “S Aux V O” “O Aux V S” “S Aux V O1 O2” “O1 Aux V O2 S” “O2 Aux V O1 S” “Adv V S” “Adv V O S” “Adv V O1 O2 S” “Adv Aux V S” “Adv Aux V O S” “Adv Aux V O1 O2 S”
$L(G_3)$	1	0	0	“V S” “O V S” “O2 O1 V S” “V Aux S” “O V Aux S” “O2 O1 V Aux S” “Adv V S” “Adv O V S” “Adv O2 O1 V S” “Adv V Aux S” “Adv O V Aux S” “Adv O2 O1 V Aux S”
$L(G_4)$	1	0	1	“S V” “O V S” “S V O” “S V O2 O1” “O1 V O2 S” “O2 V O1 S” “S Aux V” “S Aux O V” “O Aux V S” “S Aux O2 O1 V” “O1 Aux O2 V S” “O2 Aux O1 V S” “Adv V S” “Adv V O S” “Adv V O2 O1 S” “Adv Aux V S” “Adv Aux O V S” “Adv Aux O2 O1 V S”
$L(G_5)$ (English, French)	0	1	0	“S V” “S V O” “S V O1 O2” “S Aux V” “S Aux V O” “S Aux V O1 O2” “Adv S V” “Adv S V O” “Adv S V O1 O2” “Adv S Aux V” “Adv S Aux V O” “Adv S Aux V O1 O2”
$L(G_6)$	0	1	1	“S V” “S V O” “O V S” “S V O1 O2” “O1 V S O2” “O2 V S O1” “S Aux V” “S Aux V O” “O Aux S V” “S Aux V O1 O2” “O1 Aux S V O2” “O2 Aux S V O1” “Adv V S” “Adv V S O” “Adv V S O1 O2” “Adv Aux S V” “Adv Aux S V O” “Adv Aux S V O1 O2”
$L(G_7)$ (Bengali, Hindi)	0	0	0	“S V” “S O V” “S O2 O1 V” “S V Aux” “S O2 O1 V Aux” “Adv S V” “Adv S O V” “S O V Aux” “Adv S O2 O1 V” “Adv S V Aux” “Adv S O V Aux” “Adv S O2 O1 V Aux”
$L(G_8)$ (German, Dutch)	0	0	1	“S V” “S V O” “O V S” “S V O2 O1” “O1 V S O2” “O2 V S O1” “S Aux V” “S Aux O V” “O Aux S V” “O1 Aux S O2 V” “O2 Aux S O1 V” “Adv V S” “Adv V S O” “Adv V S O2 O1” “Adv Aux S V” “Adv Aux S O V” “S Aux O2 O1 V” “Adv Aux S O2 O1 V”