

Title	二者関係の概念化に基づく構文交替の文化進化シミュレーション
Author(s)	岩村, 入吹
Citation	
Issue Date	2025-03
Type	Thesis or Dissertation
Text version	author
URL	http://hdl.handle.net/10119/19714
Rights	
Description	Supervisor: 橋本 敬, 先端科学技術研究科, 修士 (知識科学)

修士論文

二者関係の概念化に基づく
構文交替の
文化進化シミュレーション

岩村 入吹

主指導教員 橋本 敬

北陸先端科学技術大学院大学
先端科学技術研究科
(知識科学)

令和7年3月

Abstract

This study explores the condition of emergence of more concrete grammatical phenomena such as “alternations” through cultural transmission of language, based on insights from previous constructive research on language evolution and cognitive linguistics. Alternations refer to linguistic phenomena where different forms are used to express different conceptualized meanings of the same situation, such as the alternation between active and passive voice.

To this end, I conducted a simulation of the cultural evolution of language, modeling intergenerational transmission in which speaker(a parent) produces linguistic forms through conceptualization of external events, while hearer(a child) infers the speaker’s conceptualization based on their utterances and the shared events. Developing this study-specific evaluation metric, I examined the correspondences between linguistic forms that represent different conceptualizations of the same situation. Through these simulations, I investigated the conditions under which languages capable of encoding different conceptualizations through alternations can emerge.

The results indicate that languages that reflect conceptualization through word order or redundant forms can only emerge when the hearer is able to perfectly infer the speaker’s conceptualization. In contrast, when inference is even slightly imperfect, different conceptualizations tend to be expressed using the same word order and vocabulary, effectively neutralizing conceptual distinctions.

These findings imply that in cases where language transmission does not involve perfect inference, ignoring conceptual distinctions may have contributed to reducing linguistic ambiguity and enabling more faithful intergenerational transmission. Furthermore, linguistic simplification due to inference failure is not limited to “exoteric communication” among unfamiliar interlocutors, as suggested by previous studies, but also occurs in vertical transmission within generations, as demonstrated in this study.

And, these findings suggest that in the earliest stages of language emergence, it was unlikely that languages encoding speakers’ subjective conceptualizations would arise. The emergence of grammatical phenomena that reflect conceptual distinctions is likely to have occurred at a much later stage in cultural evolution. Furthermore, the cognitive linguistic principle that ‘language reflects cognition’ may not have applied in the earliest stages of language evolution.

目次

第1章	はじめに	1
1.1	研究背景	1
1.1.1	言語進化	1
1.1.2	言語進化研究における2大勢力	2
1.2	研究の目的と手法	3
1.2.1	研究目的	3
1.2.2	言語の文化進化モデル	4
1.2.3	構成性・体系性：意味と形式の対応関係について	6
1.2.4	構成性・体系性：その評価指標	8
1.3	本論文の構成	8
第2章	関連分野と本研究の位置付け	9
2.1	認知言語学と生成文法における意味の所在	9
2.1.1	認知言語学における「概念化」	9
2.1.2	生成文法における「意味の二重性」	12
2.1.3	構文交替	13
2.2	言語進化における認知言語学と生成文法の対立点	14
2.2.1	生物進化：領域固有の言語能力	15
2.2.2	文化進化：領域一般の認知能力	15
2.2.3	共進化	17
2.3	計算機シミュレーションによる言語の文化進化の構成論的手法	17
2.3.1	構成性の創発	17
2.3.2	再帰性の創発	18
2.3.3	格や一致の創発	18
2.3.4	本研究の対象とする言語現象：構文交替	19
2.4	言語の性質を測る評価指標	21
第3章	概念化搭載繰り返し学習モデルと評価指標の改変	29
3.1	本モデルにおける「意味」の定義	29
3.2	繰り返し学習モデルの設計	29
3.2.1	概念化搭載繰り返し学習モデルの構成要素	37
3.2.2	概念化搭載繰り返し学習モデルにおける推論	41

3.3	シミュレーションの設計	41
3.3.1	概念化を含めた意味空間	41
3.3.2	シミュレーションプロセス	43
3.4	評価指標の改変と作成	44
第 4 章	概念化搭載繰り返し学習モデルのシミュレーション	62
4.1	パラメータ設定	62
4.2	仮説	63
4.3	実験結果	64
4.3.1	安定世代の恣意的な定義と安定世代の出現数	64
第 5 章	考察と議論	73
5.1	考察	73
5.1.1	概念化の推論と構文交替	73
5.1.2	概念化の無視を「文法の単純化」と捉えてみるとむしろ自然	74
5.1.3	作成した評価指標についての更なら改善の余地	75
5.2	議論	76
5.2.1	学習アルゴリズムから予測できること	76
5.2.2	本モデルの設計から予測できること	77
5.3	現行モデルの課題	78
第 6 章	結論	80
6.1	まとめ	80
6.2	結論	80
6.3	残された課題	81

目 次

1.1	本研究の注目する現象	4
1.2	「自転車が建物の横にある」・「建物が自転車の横にある」	6
1.3	意味空間と形式空間の地形 (Culbertson and Kirby (2016) を元に作図)	7
2.1	ルビンの壺	9
2.2	above と below の例 (Langacker (2008) を改変)	11
2.3	トラジェクターとランドマークの異なる選択	20
2.4	再掲：意味空間と形式空間の地形 (Culbertson and Kirby (2016) を元に作図)	22
2.5	完全に構成的な言語知識	25
2.6	部分的に構成的な言語知識	27
2.7	全く構成的でない言語知識	28
3.1	繰り返し学習モデル	30
3.2	二者関係の概念化	38
3.3	概念化を区別して計算する	49
3.4	従来の評価指標と作成した評価指標の比較：実線と点線の区別は、「類似な順序」と「異なる順序」に対応する．なお，青実線と緑実線は重なっている．また，評価指標の 4・5 に関しては相関係数を計算するものではないため，グラフには描画しない．	61
4.1	安定世代をもつ典型例（推論成功確率 0.98）	65
4.2	安定世代をもつ典型例（推論成功確率 0.98）における学習アルゴリズムの適用状況	65
4.3	再び安定しなくなる場合（推論成功確率 0.97）	66
4.4	安定世代が一時的に途切れる例（推論成功確率 0.98）	66
4.5	安定世代をもつ seeds における各評価指標の平均：推論成功確率 0.99	68
4.6	安定世代をもつ seeds における各評価指標の平均：推論成功確率 1.0	68
4.7	評価指標 4・5 による概念化間の言語の性質	72
6.1	推論成功確率=0.1 の言語知識 100seeds 分	91
6.2	推論成功確率=0.3 の言語知識 100seeds 分	92
6.3	推論成功確率=0.5 の言語知識 100seeds 分	92

6.4	推論成功確率=0.7 の言語知識 100seeds 分	93
6.5	推論成功確率=0.9 の言語知識 100seeds 分	93
6.6	推論成功確率=1.0 の言語知識 100seeds 分	94
6.7	安定世代をもつ seeds における各評価指標の平均 (推論成功確率 0.1)	95
6.8	安定世代をもつ seeds における各評価指標の平均 (推論成功確率 0.3)	95
6.9	安定世代をもつ seeds における各評価指標の平均 (推論成功確率 0.5)	96
6.10	安定世代をもつ seeds における各評価指標の平均 (推論成功確率 0.7)	96
6.11	安定世代をもつ seeds における各評価指標の平均 (推論成功確率 0.9)	96

表 目 次

2.1	意味要素	24
2.2	1つ目の言語知識	24
2.3	意味要素	26
2.4	2つ目の言語知識	26
2.5	意味要素	27
2.6	3つ目の言語知識	27
3.1	概念化搭載繰り返し学習モデルにおける意味空間の設定	42
3.2	疑似コード：概念化搭載繰り返し学習モデルのシミュレーション . .	43
3.3	概念化のない言語知識	46
3.4	概念化間で「異なる順序」の言語知識	46
3.5	疑似コード：Segmented TopSim アルゴリズム	50
3.6	疑似コード：Segmented Sequential TopSim アルゴリズム	51
3.7	疑似コード：Segmented Bag TopSim アルゴリズム	52
3.8	疑似コード：Separated Sequential Average アルゴリズム	53
3.9	疑似コード：Separated Bag Average アルゴリズム	54
3.10	形式長の異なる言語知識の例	55
3.11	「類似な順序」である言語知識の例	56
3.12	「異なる順序」である言語知識	56
3.13	「類似な順序」である言語知識	57
3.14	「異なる順序」である言語知識	57
3.15	「類似な順序」である言語知識	58
3.16	「異なる順序」である言語知識	59
3.17	「類似な順序」である言語知識	59
3.18	「異なる順序」である言語知識	60
3.19	従来の評価指標と作成した評価指標の比較	60
4.1	疑似コード：概念化搭載繰り返し学習モデルのシミュレーション . .	63
4.2	推論成功確率と安定世代の出現頻度	67
4.3	異なる CV と異なる右辺	69
4.4	多義な単語ルールの抜粋	70
4.5	異なる CV と共通の右辺	71

第1章 はじめに

本研究は、人間言語における意味と形式の対応関係について議論する。特に言語使用者の解釈する意味と、解釈に応じた形式とがどのように対応するのか、という対応関係の性質に注目する。具体的には、ありうる意味と形式の対応関係のうち、構成的・体系的である言語システムがどのような条件によって創発しうるのかという問題について、進化言語学と理論言語学の知見に基づき、計算機シミュレーションを用いた構成論的手法によって明らかにする。最後に、上記の結果から、言語進化の一つのシナリオを提示する。

本章ではまず研究背景として、

- 解くべき問いはなにか
- どの理論的基盤に立脚するのか
- どのような手法を使うか
- 解くことの学問的意義や社会的意義はなにか

という点について述べたあと、本論文の構成を示す。なお、「形式」や「意味」という言葉は言語学の専門用語として使われる。「形式」は他者にとって知覚可能な音声や文字を指し、「意味」は他者にとって知覚的にアクセス不可能な心の状態である。但し、伝統的に言語学で「形式」と言えば、音声を指示する。言語進化・起源の議論に際してはなおさら音声を指す。

1.1 研究背景

1866年、パリ言語学会は、言語起源に関する議題を禁じた(山内, 2012)。この会則は改定されたものの、これにより1世紀近く言語起源の研究は停滞した。しかし、自然主義的アプローチで言語研究を目指す、Noam Chomsky を嚆矢とする生成文法の登場により、言語起源・言語進化研究は学際性を帯びて復活した。

1.1.1 言語進化

言語進化とは、まさに「言語の進化」のことであるが、言語とは何か、進化とは何か、という問いが重要である。そして、これらの問いは、立脚する理論的基

盤のバリエーションを生む。これはちょうど Jackendoff (2010) のタイトル “Your theory of language evolution depends on your theory of language (「言語進化の理論は言語の理論に依存する」)” によく示されている。

大雑把であるが、言語をヒトがもつ生物学的形質として捉えるならば、その内在的言語能力の生物進化が主な研究対象となる。但し、言語能力が運用されて歴史的に言語が変化する現象は「言語変化」という研究分野であり、言語進化の問題としては中心的ではない。一方で、言語をコミュニケーションの道具として捉えるならば、その外在的な言語運用つまり発話された言語の文化進化が研究対象となる。但し、運用の基盤となる能力や骨格の生物進化の側面は発達心理学や人類学など言語学の周辺の研究であり、進化言語学者の出る幕はそうそうない。いずれにせよ、思考の道具としてもコミュニケーションの道具としても使われる言語の諸性質を明らかにする必要がある。特に、なぜ人間言語にだけ、なぜ人間のコミュニケーションシステムにだけ、他の動物とは異なる諸性質があるのかということ明らかにする必要がある。

本研究の主題は、人間言語のもつ意味と形式の対応関係に関するものである。この主題に対しては、上記の2つの立場のいずれか一方を採用したならば、即座に他方を否定するような二者択一が迫られるべきではなく、両者の共創を目論むのが妥当だろう。しかしながら、進化の速度を考えれば、遅い生物進化と早い文化進化の区別もまた必要である。他の性質の進化同様、言語進化もこうした緊張関係にさらされている。そのため、本研究の裏テーマとしては、言語能力の生物進化と言語の文化進化の貢献範囲を、どのように線引きできるか、という対立する両者における説明範囲の棲み分けを促すこともまた1つの目的である。つまり、文法現象のうち、どれは生物進化の産物で、どれは文化進化の産物であるのかを区別し、両理論の説明責任を分担することを目指す。

1.1.2 言語進化研究における2大勢力

進化言語学における2大勢力は、ちょうど前述した「言語とは何か」という探究に関する理論的違いに対応する。その違いとは、生成文法が仮定するところの生物進化で獲得した生得的な生物学的形質としての言語能力の有無である。生成文法に基づく進化言語学は生物言語学 (Biolinguistics) (Chomsky, 2005) と呼ばれ、離散無限性と階層性を可能にする計算論的メカニズムとして *Merge* (併合) と呼ばれる再帰的な集合形成演算を想定し、その進化を議論する。この併合の進化については、突然変異による脳の再配線が併合の跳躍的進化 (大躍進, “*Great Leap Forward*”) をさせたとする仮説 (Chomsky, 2005) や、前駆体として非再帰的な集合形成演算を仮定する運動制御起源仮説 (藤田, 2012) や、自己家畜化による生態学的制約の解除に伴い併合の再帰的な無限性も漸進的に解放されていったとする仮説 (Benítez-Burraco and Progovac, 2020)、主に構文処理を担うブローカ野と主に意味処理を担うウェルニッケ野を接続する弓状束 (Arcuate Fasciculus) において、神経伝達を効率化する

ミエリン化 (Myelination) が生じたとする神経言語学からの仮説 (Friederici, 2009) など、様々提案されてきた。

一方、認知言語学・社会言語学 (いずれも機能主義言語学に属する) に基づく進化言語学は、コミュニケーションが生じる使用事象 (usage event) における一般認知能力の進化を議論する。しかし、一般認知能力それ自体の生物進化を扱うというよりは、その能力が使用されるコミュニケーション様式やコミュニケーションとして使われる記号体系の文化進化が主である。特に、共同注意と意図共有に基づく言語発達と使用基盤モデルに基づく漸進的文法化による記号的コミュニケーションの起源 (Tomasello, 2005) や、再帰的読心能力の進化に伴う意図明示推論コミュニケーションへの移行 (Scott-Phillips, 2014) や、社会的な文化伝達による構成性の創発 (Kirby, 2002) など、提案されてきた。

進化言語学における2大勢力は異なり理論的基盤に立ってはいないものの、本質的には同じことをモデル化しようとしている。それは、言語の「意味と形式の対応関係」である。生成文法においては、階層構造によって意味と形式が接続されているとする分析に基づき、言語能力を司る言語機構 (the Faculty of language) の中枢に Syntax と呼ばれる階層構造構築システムを据える。すなわち、意味と形式が階層的な対応関係を持つための処理について研究している。

一方で、認知言語学においては、意味と形式が言語使用者の認知処理を媒介して接続しているとする分析に基づき、意味が形式に反映される認知メカニズムを探究する。すなわち、意味と形式が体系的な対応関係を持つことを前提に、その対応を作り出す認知処理について研究している。

以上より、意味と形式の対応関係がどのように作られているのかということを問うことは、進化言語学において重要である。

1.2 研究の目的と手法

1.2.1 研究目的

本研究の目的は、これまでの言語の文化進化研究の構成論的手法と認知言語学の知見を組み合わせ、より具体的な言語現象の創発条件を解明することにある。具体的には親は外部事象を主観的に概念化して発話し、子は推論して学習するという世代間継承において、概念化された意味が構成的かつ体系的に形式に反映し、同一事象に異なる構文を使う構文交替 (図 1.1) の創発条件をシミュレーションにより明らかにする。上記の目的を達成することは、言語創発当初に近い状況における言語進化シナリオの示唆を与えることになるだろう。

これまでの言語の文化進化研究において文法の核となる構成性や体系性の創発が探究されてきた。しかし、そこには「意味」構造の設計に大きな単純化があり、人間言語の特徴を捉えきれていない。認知言語学では、言語使用者が概念化プロセスを通じて「意味」を形成し、その概念化された意味が言語形式に反映されるこ

とを主張してきた。言語使用者がどのように事象を概念化したかに応じて、構文（例：能動文と受動文など）が使い分けられることを示唆してきた。これが正しいとすれば、同一事象に異なる構文を使う構文交替は異なる概念化によってもたらされたと言える。ただし、コミュニケーションにおいて聞き手が話し手の真意を完全に理解することは難しい。話し手が密かに概念化を構文交替を使って反映させていたとしても、聞き手は話し手の概念化を知ることは難しいように思う。特に、言語創発初期の段階で仮に言語が構成的でも体系的でもなかったならば、話し手の概念化が聞き手に正しく伝わる可能性は低くなる。このような状況下で、構文交替がどのように創発するのかを明らかにすることが、本研究の中心的な目的である。

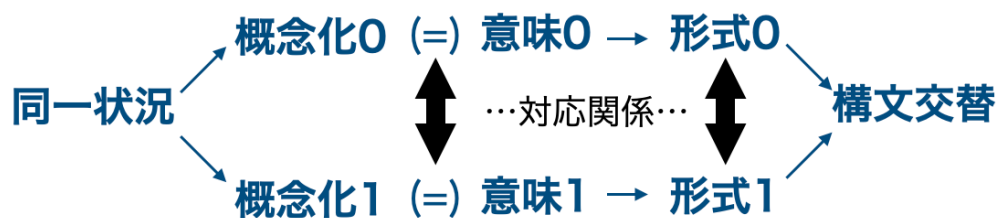


図 1.1: 本研究の注目する現象

1.2.2 言語の文化進化モデル

人間言語の起源・進化・創発という現象を説明するためには、言語化石なる人類学的証拠は得にくく、また一回性の (singular) 出来事であることから、計算機上で言語現象を模倣することで理解を進める構成論的なアプローチが有効である。言語という記号システムを長い歴史の中で人々が相互作用して作り上げた社会文化的構築物としてみた時、言語を文化進化研究の俎上に載せることができる。

Kirby は、言語の文化進化研究において代表的な「繰り返し学習モデル (Iterated Learning Model; ILM)」を提唱した (Kirby, 2001)。これは前世代の個体の発話データを学習した個体が今度は、次世代の個体が学習するためのデータを発話をするという、学習と発話が連鎖する文化進化のモデルである。言語創発最初期には意味と形式の対応関係がたとえランダムだとしても、その対応関係が学習と発話の連鎖を通じて規則的になり、複合的な意味をより単純な部分の意味の組み合わせで表現できる構成性 (Compositionality) が創発する。鳥類のいくつかは生得的でない形式を獲得する発声学習を行う (Jarvis, 2004) ため、発声パターンが世代間継承

されることで、形式の分節化と複雑化が生じる (岡ノ谷, 2010). しかし、ここに意味論はない (Berwick et al., 2011; 米納, 2022)¹. Kirby らの研究は、予め分節化されている意味表現と構造のない形式表現のペアを言語として定めている. そしてこうした言語が文化進化を経ることで、分節化された意味に基づいて形式が分節化されていき、学習が容易なように分節化された意味と形式の要素同士が1対1対応の関係を持っていく. 鳥類の例から分かることは、形式は意味とは独立に分節化される方向にあるということである. 人間言語の場合は、そうした方向に加えて、分節化された意味に駆動されて形式も分節化し、意味の構成要素が分節化した形式要素と一対一対応の関係を持ち、構成性が創発する²、と解釈できる. そして Kirby らにとっての「言語起源・進化・創発」とは、「言語らしきもの (原型言語)」が徐々に構成性などの普遍的性質を帯びて「完全な言語 (現存する言語)」になる過程である.

これまで、構成論的な方法に絶えず寄せられる恣意的なモデル化への懐疑を避けるべく、様々な学習機構 (ルールベース, ニューラルネット, ベイズなど) による感度分析により ILM の頑健性が検証されきたものの、意味構造の設計には議論の余地がある. 例えば述語論理 (確定節文法) (Kirby, 2002), 素性 (feature) の順序集合 (Kirby et al., 2015), 属性 (attribute) の順序集合 (Ren et al., 2020) が意味構造の記述に使われているが、こうした設計は、外部事象の集合が意味であることを前提にし、また、意味の構成要素同士の“関係性”を操作する自由度はないという制約を内包する.

機能主義言語学に属する認知言語学では、言語は認識の反映であるというテーゼ³を基に、言語使用者内部の意味に注目する. ここでの意味は、真理値が決まる客観的な外部事象ではなく、言語使用者が外部事象をある特定の「観点 (perspective)」からの「捉え方 (construal)」を採用して「概念化 (conceptualization)」することで形成される (Langacker, 2014). ここには事象の構成要素同士の“関係性”を言語使用者がある程度主観的に操作する自由度がある.

¹ただし、シジュウカラの地鳴きという生殖に関与しない警戒コールなどの歌に構成性と統語論があるとする研究もある (Suzuki et al., 2016, 2018).

²このこと自体、今後検証を必要とする仮説である. しかしながら、上田他 (2024) においても同様の指摘がされていることから、一定程度合意のある仮説であると言える.

³このテーゼに関する哲学的批判は酒井 (2013) を参照. 特に、客観的世界と主観的世界の二元論が成り立たないことや、言語の慣習性を指摘し、客観的世界を知り得ないのに言語が主観的な認知を反映していると前提にできるのか、また言語を他者から学習する以上、純粹に個人の認知を反映することはできないし、言語に反映された個人の認知は無限後退してしまうと指摘する.

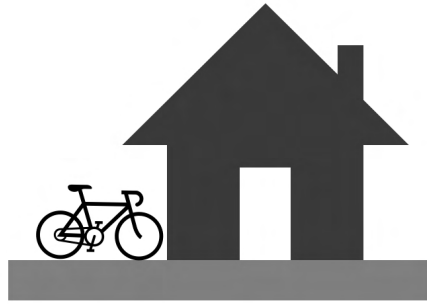


図 1.2: 「自転車が建物の横にある」・「建物が自転車の横にある」

例えば、後で詳しく述べるが、図 1.2 であるように、自転車に注目するか建物に注目するかによって、言語表現は異なる。どちらに焦点を当てるかという figure/ground や trajector/landmark の選択は、「際立ち (salience あるいは prominence)」といった認知的動機づけ (Langacker, 1990) こそあれ、本来言語使用者に選択の自由がある⁴。ただし、意味が言語使用者内部にあるならば、外部世界の共有だけでは他者には不可知であり、意図共有 (Tomasello, 2005) に基づいた推論が必要となる。特に、構成的な文法がなかったと想定される言語創発最初期において、意味と形式の対応関係が構造化されていないことは、言語使用者の概念化を推論することを難しくさせるはずだ。

1.2.3 構成性・体系性：意味と形式の対応関係について

意味と形式の対応関係は無限に考えられる。ここでは、図 1.3 の 3 種類について説明する。左図の対応関係であれば、ある意味を表現する形式がどれなのか、予測が効きづらい。一方、右図の対応関係であれば、全ての意味が 1 つの形式で表現されており、これだと形式を聞いてもどの意味を表すのかが曖昧である。真ん中の図の対応関係は規則的であり、相互に予測が可能である。自然言語での例を出せば、“He Kicked the bucket(彼は死んだ)”と “He kicked the ball(彼はボールを蹴った)”と比較した時、形式は非常に近いにもかかわらず、意味はほとんど関係がないように思われる。こうした 2 つの文だけの言語を想定すると、左図に当てはまる。なぜなら、“He Kicked the bucket(彼は死んだ)”に関しては、“kick the bucket”の意味が “kick”にも “the bucket”にも還元できず、全体として「死んだ」を意味するが、“He kicked the ball(彼はボールを蹴った)”に関しては、それぞれの形式が意味に還元できてくるからだ。右図の例は、多義語だ。“bank”という英単語は「銀行」の意味もあれば「土手」の意味もある。一見無関係に思われる意味

⁴筆者には「自転車が建物の横にある」という言語表現の方が、日常的な感覚としてはより自然であるけれども、文脈を補えば、「建物が自転車の横にある」という言語表現も自然に感じられる。

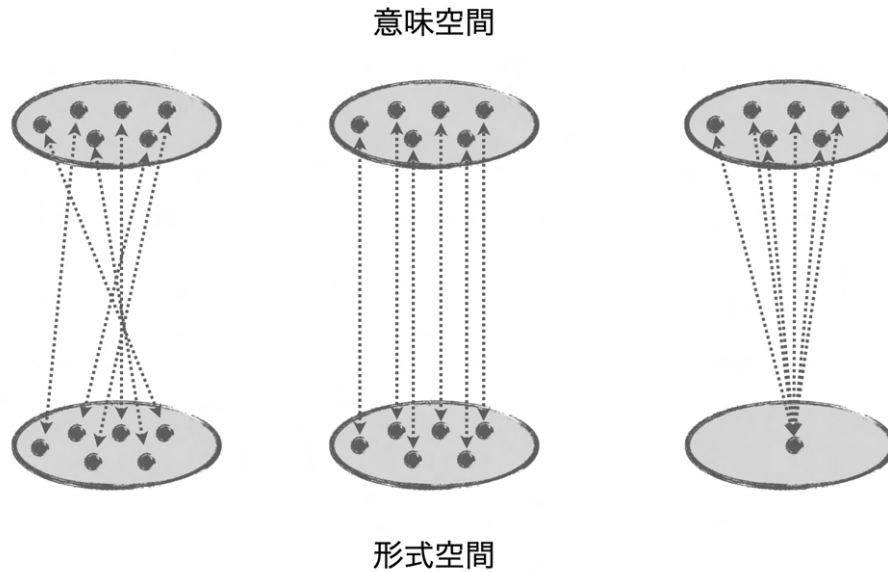


図 1.3: 意味空間と形式空間の地形 (Culbertson and Kirby (2016) を元に作図)

が1つの形式で表されている．真ん中の図の例は，“He kicked the ball(彼はボールを蹴った)”と“He kicked the pole(彼はポールを蹴った)”である．“He”という形式は「彼は」に対応し，“kicked”は「蹴った」に対応する．そして，“the ball”と“the pole”はそれぞれ「ボール」と「ポール」に対応する．

また、上の図1.3では表現されないが，“the stars visible”と“the visible stars”(Cinque et al., 1994; Culbertson and Kirby, 2016)は、形式の組み合わせ方によって異なる意味を表現できる．このように、意味と形式それぞれが対応関係を持つだけでなく、その組み合わせ方においても対応関係を持つ場合がある．

以上のように、意味と形式の対応関係はいくつも考えられ、それぞれ自然言語に存在する．ただし、自然界の他動物のコミュニケーションと比べると、意味と形式の対応関係が真ん中のように綺麗な規則をもつ性質は、人間言語に特徴的である⁵．そのため、進化言語学では、人間言語の特徴的な性質に着目してきた．

ここでは、構成性と体系性について定義を示しておく．しかし、非常に密接に関わる概念であるため、明確に区別して定義を与えることは難しい．但し、大まかに言えば、構成性は局所的な範囲でも成立可能であるが、体系性は全体的に成立する必要がある．

- 構成性：全体の意味が部分の意味とその組み合わせ方で表現される性質
- 体系性：形式の組み合わせ方が一貫して予測可能である性質

⁵Culbertson and Kirby (2016) において “arguably” という断言を避ける弱い表現が使われていることから、人間言語固有であるかどうかは議論の余地があることが窺える．

1.2.4 構成性・体系性：その評価指標

計算機を使った言語進化シミュレーション研究において、構成性と体系性の評価指標を作成することは重要である。なぜならば、多くの言語の文化進化研究において、言語創発当初は構成性や体系性などの性質が存在しない言語であったと仮定し、文化進化を経ることで、それらの性質が創発してくる可能性を探究してきた。そうであるならば、諸性質の創発過程を観測する指標が必要である。ある対象にある性質があるという判断や評価は自明ではないので、それらを（できれば没個人的に）判断する基準が必要である。これまでに構成性を測る指標は提案されてきた（詳しくは2.4節）。言語の文化進化研究と近年接続が期待される創発コミュニケーション (Emergent Communication; EmCom)(上田他, 2024)においても、同様の評価指標が使われることがある。しかしながら本研究でモデル化する同一状況を異なる概念化して形成した意味と、それらを表現する形式との対応関係を評価することが難しい。そこで本研究における意味と形式の対応関係を観測するためには新たに評価指標を作成する必要がある。よって、言語の文化進化研究においても、創発コミュニケーション分野においても、新しい評価指標が必要となる。

1.3 本論文の構成

本論文の構成を概説する。本章に続く第2章では、言語の文化進化研究と認知言語学における基礎タームを解説する。言語の文化進化研究で見過ごされてきた意味の側面に対して、認知言語学の知見に基づき、「言語使用者の意味解釈（本研究での二者関係の概念化）」を取り込む。第3章では本研究のモデルと本シミュレーションについて、2章で説明したことを踏まえて提示する。第4章ではシミュレーション結果を提示する。第5章では、4章での結果を考察し、自然言語の現象とどのような関係にあるのかについて議論する。第6章で本研究の結論と限界、今後の展望を行う。

第2章 関連分野と本研究の位置付け

本章では、理論言語学と進化言語学の交差点に本研究が位置していることを明確にするため、主に認知言語学の研究と言語の文化進化研究を、意味とコミュニケーションの観点から概説し、構文交替に着目する意義を述べる。言語の生物進化と文化進化、共進化の研究においてどのような仮定のもと、「何」の進化が問題とされ、どのような仮説が提唱されているのかについて述べる。続いて、言語の文化進化に焦点を当て、これまで何が明らかになってきたのかを述べ、本研究で明らかにする言語現象を再び述べ直す。最後に、これまで創発言語の性質を観測するための評価指標を解説する。

2.1 認知言語学と生成文法における意味の所在

2.1.1 認知言語学における「概念化」

認知言語学では、言語使用者が知識や能力を使って外界を概念化することを重視する。例えば、「ルビンの壺」が有名である(図 2.1)。絵に描かれていることは同じでも、言語使用者がどのように概念化したのかによって、「壺」が見える場合と「2つの顔」が見える場合があり、それに動機づけられて命題も異なる。いわゆる「図と地の反転」である。中央のくびれた部分に注目して、その周辺(=地)を背景化すれば「壺(=図)」が浮かび上がってくるし、逆に周辺の凹凸部分に注目し

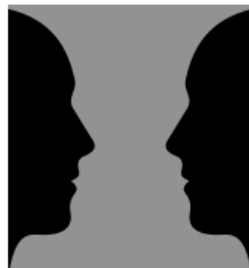


図 2.1: ルビンの壺

て、中央を単なる何もない空間 (=地) として背景化すれば、「2つの顔 (図)」が浮かび上がる。

こうした動的な認知過程としての概念化は、「捉え方 (construal)」に本質的に依存する。Langacker (2008, Ch.3) は、捉え方を4つに分類した。

- 詳述性 (specificity)
- 焦点化 (focusing)
- 際立ち (prominence)
- 観点 (perspective)

ただし、これらは綺麗に棲み分けできるものではなく、相互に綿密に関わっている。また、本研究では3つ目の際立ちについて主に注目するため、他の要素に関しては、詳細には立ち入らず、代表的な例による手短な説明に止める。

1つ目の詳述性について。これは個別具体的な表現と抽象的な表現の連続に見られる記述の粒度に関わる。例えば、ある人を見て「誰かさん」と言っても良いし、「〇〇さん」と固有名詞を使って良い。事象レベルであれば、ある出来事を見て、「誰かさんが何かした」と情報量の少ない発言も可能であるし、「ジョンがメアリーを10発も激しく叩いた」と発言することも可能である。このように、同一の状況に対して異なる詳細さで捉えることは可能であり、そこで捉えられたことが表現として現れる。

2つ目の焦点化について。これは全体の中から一部分を選択する過程に関わる。例えば、図2.1で見た「ルビンの壺」だけでなく、ドーナツの「穴」や身体部位のうちの「ひじ」という表現に見られる。これらはどれも、明示的に表現するかはともかく、ある事物や事象の中から特定の部分に注目して、その部分を捉え、そこで捉えたことが表現として現れる。

4つ目の観点について (3つ目は後述する)。これは、どの視点から事物・事象を捉えるかということに関わる。例えば、上り坂と下り坂である。端的に、言語使用者が (実際に立っているかはともかく) 坂の上から見なのか下から見るのかに依存する。

3つ目の際立ちについて。これは焦点化と深く関わる。焦点化との関係は、焦点が当てられたものは当てられなかったものに比べて、意味論的にも語用論的にも際立ちを持っていると説明できる。つまり、際立ちとは、特定の要素が他の要素よりも認知的に目立つ程度である。際立ちを持つから、認知的に目立つからこそ焦点が当てられるのである。そして、これは本研究において、もっとも重視する分類である。なぜならば、意味論だけでなく文法構造にも関わる概念であるからだ。際立ちには、さらに、

- プロファイル (profile)

- トラジェクター (Trajector) とランドマーク (Landmark)

という下位分類がある。つまり、目立つ程度の順位と言える。本研究では「トラジェクター/ランドマーク」に注目するが、プロフィールに関しても説明を加える。

プロフィールについて。これは、事物や事象の性質を抜き出す過程である。事物は単体で存在しているというより他の事物との関係性のネットワークに埋め込まれている。事物のプロフィールとは、関係性のネットワークから事物単体を抜き出すことであるが、背後にはネットワークが隠れている。例えば、「おば (Aunt)」や「弧」などは代表的である。「おば」単体は単なる女性でしかないが、親族関係の親戚ネットワークのなかでのみ位置付けられる女性である。また、「弧」も単体では単なる曲線であるが、ある図形において位置付けられる曲線である。このように、ある表現は、概念化の対象として全体のなかの一部を際立たせる効果を持っている。また、概念化の対象全体が同じ場合でも、言語使用者が異なる表現を使うことで、異なる一部が際立ちを持つことになる。例えば、「肘 (elbow)」と「手 (hand)」はどちらも腕全体における異なる身体部位を際立たせている (Langacker, 2008)。

また、関係性のプロフィールもある。一般に動詞表現がこれを担う。「親がいる (have a parent)」と「子供がいる (have a child)」状況は親子関係にある 2 人という意味において等価であるが、どちら側から関係性を記述するかによって異なる表現となっている。

そして、いよいよ「トラジェクター/ランドマーク」の説明に入る。

この道具立ては、関係性の記述における非対称性を扱う。トラジェクターもランドマークも際立ちを持つが、トラジェクターの方が主要な (primary) 対象で、ランドマークはその他の (other) 対象である (Langacker, 1990)。

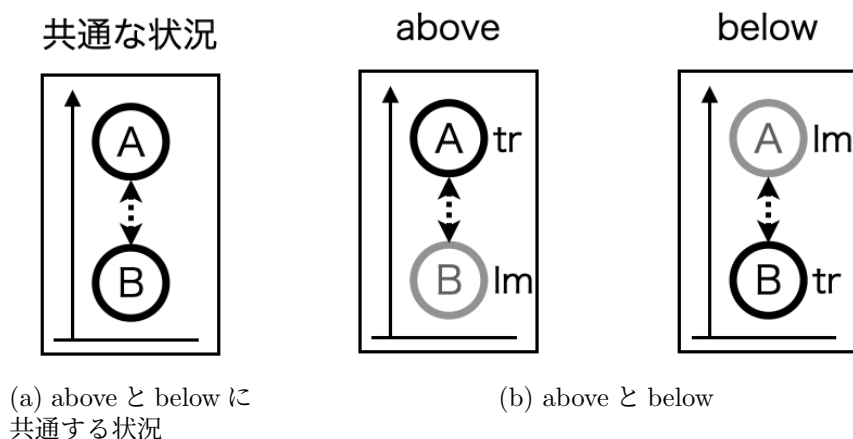


図 2.2: above と below の例 (Langacker (2008) を改変)

例えば、図 2.2 にあるように、“A above B” と “B below A” は同じ相対的な関係にある A と B の状況をプロフィールしているが、どちらに焦点を置き際立ちを

持たせるのかによって異なる表現を持つ。つまり、このトラジェクターかランドマークかを選択する言語使用者は、捉え方の自由度を持つ¹。典型的には、動作動詞における動作主がトラジェクターで、被動作主がランドマークに対応するが、より広く関係性の記述に適用可能である。

Langacker (2008, p. 73) の脚注に興味深い言及がある。

This is not to deny that the world imposes itself in particular ways, thereby constraining and biasing our apprehension of it. Likewise, normal patterns of conceptualization constrain and bias the conceptualizations invoked for linguistic purposes. We nevertheless have enormous conceptual flexibility, and the biases at each level are easily overridden. (Langacker, 2008)

これ [捉え方の他の側面と同様、際立ちは概念的現象であり、世界を認識する際に内在するのであって、世界自体の中にあるのではないこと] は世界が特定の仕方で世界自体を提示することで世界の認識の仕方を制約したりバイアスをかけたりすることを否定するわけではない。同様に、概念化の標準的なパターンが概念化を言語目的²に沿うように制約したりバイアスをかけたりする。それにも関わらず、我々は概念化 [の仕方] に非常に大きな柔軟性を持ち、各レベルでバイアスが覆される。(拙訳)

どのような観点からどれほどの詳細さでもって対象を概念化したとしても、対象ひとつひとつをどのような関係性（主要かその他か）をもつものとして概念化するかは、言語使用者に任されている。コミュニケーション場面における語用論的制約があったとしても、我々が持つ概念化の柔軟性を持つことで、依然としてどのように概念化するかという自由度は残されている、と言える。

2.1.2 生成文法における「意味の二重性」

生成文法において、言語運用での意味の取り扱いはどちらかというと軽視される傾向にある。生成文法における言語は、理想的言語話者の内在的言語 (I-言語) である。そのため、生成文法の解くべき課題は、意味と形式を階層構造によってつなぐ言語システムの計算論である。言語システムは、意味を解釈する概念・意図システムと音声形式を解釈する感覚・運動システム、それら両者をつなぐ統語演算システム (syntax)、そして統語演算システムに入力を与えるレキシコンからな

¹しかしながら、コミュニケーションにおいて常にある事物が話題化されて進行しているため、原理的には言語使用者の捉え方には自由度が存在するが、コミュニケーション場面における語用論的制約・情報構造の制約も存在することが指摘されている (Langacker, 2008)

²おそらく動詞が取れる項の数などを指す

る (Chomsky, 2014, Ch. 3). 特に, 統語演算システムの解明が生成文法の中心的な課題であり, その周辺のシステムの研究はやや先送りの状態である. こうした状況に対して, レキシコンの中身である語彙項目の進化過程に取り組む研究がある (藤田, 2022).

ただし, どのように語彙項目が存在するものとして, 統語演算システムに入力が与えられたときに, 概念・意図システムと感覚・運動システムにとって解釈可能な構造を生成できるのかという課題に際して, 言語形式にどのように意味が符号化されているのかについては, 生成文法の最初期から整理がされている. 近年では「意味の二重性 (Duality of semantics)」という概念が使われている.

At the semantic interface, the two types of Merge correlate well with the duality of semantics that has been studied within generative grammar for almost 40 years, at first in terms of “deep and surface structure interpretation” (and of course with much earlier roots). To a large extent, EM yields generalized argument structure (theta-roles, the “cartographic” hierarchies, and similar properties); and IM yields discourse-related properties such as old information and specificity, along with scopal effects. (Chomsky, 2008, p. 140)

意味インターフェイスにおいて, 2種類の併合は意味の二重性によく対応する. これは約 40 年間の生成文法で研究されてきた概念で, 最初期における「深層・表層構造の解釈」という形であった (ただし, もっと古い起源を持つ). 大部分において, EM(External Merge; 外的併合) は一般化された項構造 (意味役割, 「カートグラフィー」における階層, 類似の性質) を, IM(Internal Merge; 内的併合) は旧情報や特定性といった談話に関係する性質, さらにスコープ効果をもたらす. (拙訳)

すなわち, 項構造の意味役割は外的併合が与え, 談話・語用論の意味は内的併合により与えられるという分業体制である. 藤田 (2017, p. 83) は「外的併合＝深層意味, 内的併合＝表層意味」と定式化する³

2.1.3 構文交替

野中 (2024, p. 13) によれば, 構文交替 (alternation)⁴とは「ある種の動詞が基本的な意味を保ったまま異なる形式の構文に現れる現象」と定義する. 例えば, 場所

³しかし, 現行のミニマリストにおいて, 外的併合と内的併合は区別せずに統合する方向に研究が向かっている. しかしながら, 内的併合はコミュニケーションにおける談話的情報に関わるのだとすれば, 生物言語学における「思考のための言語」に特徴づけられる概念・意図システムへの最適化という方針は崩れ, 統語演算システムに「コミュニケーションのための言語」を導入することになる. この方針はいかに担保されるのか. 今後の動向が気になる. (というポエムを書いている暇はない)

⁴同書において, 生成文法では伝統的に一方から他方への「変形」という概念を使って構文交替

格交替 (“John loaded hay into the truck” ・ “John loaded the truck with hay”) や与格交替 (“Sally sent a letter to Harry” ・ “Sally sent Harry a letter”) や身体部位所有者上昇交替 (“Mary struck John’s arm” ・ “Mary struck John on the arm”) などを挙げ、認知言語学において、2つの構文の形式上の違いが「捉え方」を含む意味の違いに起因すると説明してきたと述べる。但し、構文 (construction) を「形式と意味の慣習的な結びつきのうち、複合的なパターン (複数の語または複数の形態素から成るもの)」(野中, 2024, p. 13) と定義することから、「構文交替」の「構文」よりも一般的な意味で使っている。これを踏まえれば、「ある種の動詞」に限定させる必要はなく、構文交替における構文を「基本的な意味を共有する異なる形式」として本研究では定式化する。このことは以下の引用 (Langacker, 1987, p. 39) から補強される。

These sentences [“He sent a letter to Susan”, “He sent Susan a letter”] have the same truth value and can be used interchangeably to describe the same event, but I [=Langacker] suggest that they nevertheless differ semantically[reference omitted].[...] The differences in grammatical structure therefore highlight one facet of the conceived situation at the expense of another; I will say that the two sentences present the scene through different images. Langacker (1987, p. 39)

これらの文 (“彼は手紙をスーザンに送った” と “彼はスーザンに手紙を送った”) は同じ真理値であり、同一事象を描写するのに相互に交換可能であるが、それでもなお意味的に異なると私ラネカーは提案する。(中略)。このように、文法構造の違いは、想定する状況の一面を強調し、他方を犠牲にする。つまり、2つの文は異なる捉え方で事象をしていると言える。(拙訳)

これより、構文交替ないし構文選択とは、真理条件的意味が同じ同一事象を異なる捉え方で概念化して符号化された形式のミニマルペアであると言える。本研究では上記に示した意味で構文交替という用語を使うこととする。

2.2 言語進化における認知言語学と生成文法の対立点

言語進化において問うべきことは、ヒト以外の動物コミュニケーションには見られない人間言語固有の性質かつ普遍的な性質は何に由来するのか？ ということである。人間言語固有で普遍的な性質を直接反映するような生物学的特質が存在

が説明されていたと述べている。これを考えると「交替」という用語は多分に生成文法の香りがするものであり、認知言語学・構文文法の想定に反する。むしろ一方ではなく他方を「選択」という方が、認知言語学・構文文法の想定に適う用語のように思われる。

すると仮定する方が有益なのか、それとも人間言語進化で普遍的な性質を間接的に反映していくような文化的仕組みが存在すると仮定する方が有益なのか。大きく分けたこの2つの立場は理論言語学におけるアプローチの違いに対応する。

2.2.1 生物進化：領域固有の言語能力

生成文法に基づく生物言語学において、概念意図システムと感覚運動システム、そしてそれらをインターフェイスでつなぐ統語演算システム全体の進化が問題となる。中核である統語演算システム内で働く併合 (Merge) により再帰的な階層構造を生成することから、併合の進化について説明する研究がある (Berwick and Chomsky, 2017; Chomsky, 2005)。まず、併合の定義を示す。

$$\text{Merge}(X, Y) = \{X, Y\}$$

併合とは、2つの離散的な要素を入力として、それらを元とする1つの無順序集合を生成する二項演算である。また、出力を入力として呼び出すことも可能である。これにより、併合を繰り返し適用すれば、再帰的な階層構造を生成することができる。

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

ただし、併合は結合律は満たさない。

1.1.2 節で言及した Chomsky (2005) や藤田 (2012) や Benítez-Burraco and Progovac (2020) はどれも「併合の進化」に関するものである⁵。

2.2.2 文化進化：領域一般の認知能力

言語の文化進化研究において、大きなテーゼがある。それは、文化進化を通して人間の心の性質が文化的継承物に染み出してくる、というものだ。例えば、Griffiths and Kalish (2007) は、継承される情報が最終的に構成的な言語を好むといった学習者の心的性質を反映するようになるという示唆する。Kirby et al. (2007) は、文化継承が弱いバイアスを増幅させて構成性という強くてロバストな言語普遍性を導いた可能性を主張する。Culbertson and Kirby (2016) は、単純性バイアスが言語という認知処理ドメインに固有な形で関わることで、構成性・規則性・語順秩序性・同型写像性といった言語現象を導いた可能性を主張する。以上のように、文化継承を通じて人間の認知的特性（心の性質）が言語の構成性や規則性などの普遍的性質を形作ることが示唆されている。

これらの研究が示すことは、文化が人間の認知的特性を反映するということである。文化継承には、その文化の学習と生成が不可欠である。言語の文化進化研

⁵こうした併合中心の進化研究に対して松本 (2022) は批判が寄せられている

究においては、認知的特性の反映過程である文化継承に必要な言語の学習と生成の能力を、言語領域固有に内在する特殊な能力を前提とせず、言語に限らない一般的な認知能力を仮定している。Culbertson and Kirby (2016) の主張から伺えるように、文化進化を経ることで、能力が使用される対象に徐々に人間側の性質が増幅したり累積したりして、結果的にその性質を帯びていく。そして言語の文化進化は人間普遍的な現象であるが故に、人間の普遍的な認知的特性が言語に反映され、言語に普遍的な性質が見つかるのである。

しかしながら、Kirby et al. (2015) はコミュニケーション時の表現力と学習時の圧縮力が拮抗する (trade-off) 状態でのみ構成的な言語が創発することをシミュレーションで示した。Culbertson and Kirby (2016) においては「単純性バイアス」が様々な言語現象において異なるような性質 (構成性・規則性・語順秩序性・同型写像性) を見せることを理論的に説明したが、これは Kirby et al. (2015) の結果とは幾分食い違う。というのも、コミュニケーションない学習だけの垂直伝播の場合だと、同じ形式で様々な意味を表現するような縮退的言語 (degenerate language) が頻発したことためだ。この結果は、必ずしも単純性バイアスが規則性や構成性を促進するわけではなく、コミュニケーションの表現力が必要であるということだ。

吉川 (2017) も言語の規則性は「どこにあるのか」と問う。そして、個人の脳内ではなく、個人間の相互作用の場としての社会にあると説明する。つまり、個人知ではなく社会知であるとする。鴨川の等間隔にカップルが並ぶ現象を引き合いに出し、この等間隔という規則性は、全体を俯瞰した規則的な並び方を実現しようとする各個人の脳内の計算処理ではなく、「できるだけ隣と離れたい」という原理を持った各個人たちの相互作用の結果創発した現象であるという。これを言語の規則性の説明についてアナロジーを適用させる。文化進化における未学習者への継承は、学習者にとっては「できるだけ同じ表現を使う」という一種の制約を課すことであると見做す。その上で、最小限の編集原理を取り入れたシミュレーションの結果 (Batali, 1999) や人工言語の伝言ゲームをさせる実験室実験の結果 (Kirby, 2002) や見知らぬ人と話すことで文法が単純化するという記述研究 (Wray, 2007; Wray and Grace, 2007) などの成果を踏まえ、言語の規則性は「できるだけ同じ表現を使う」という社会的圧力 (social pressure) が働く下で創発する社会知であると説明する。このように、文法という言語の規則性とは、社会的圧力が働く下で産出される言語運用に見られる現象的な性質であると同時に、その現象的な性質を規範として推論した結果の共同幻想的な性質であると述べる。規則が社会から創発する「生成」過程と、規則が社会に定着し「存在」するようになる実体化を繋ぎ、文法のありかを個人の言語機構なり認知なりによって規定せず、社会まで拡張するアイデアである。この説明においても、言語領域固有な能力は仮定していない。

2.2.3 共進化

言語進化研究には、ヒトの言語能力の生物進化と言語の文化進化の両過程が互いに影響を与え合うことを重視した共進化研究がある。Kirby (2017) はこれまでの研究事例を整理し、学習と文化継承と生物進化がループする複雑適応系として言語進化の描像を提示する。Smith (2020) は言語継承において学習と表現が存在する場合、バイアスを伴う学習能力の生物進化と、その結果強くなったり弱くなったりしたバイアスを反映した言語の文化進化が相互に影響しあうことで言語が進化しうることを計算機シミュレーションと実験により示した。

また、Azumagakito et al. (2013) では、突然変異で構造の複雑な言語を発話する個体が、その発話を理解できない大多数の周囲とコミュニケーションできず、適応度が低く集団に広まっていけない「孤独な突然変異体問題」を取り上げる。そして、この問題を表現型可塑性（学習の柔軟性）によって解決する生物進化と、世代間・内継承によってコミュニケーション可能なように言語自体が変化する文化進化とが、循環的に共進化 (cyclic coevolution) するシミュレーションをした。その結果、構造の複雑な言語を生成できる言語能力が遺伝的に固定化するボールドウィン現象が生じ、その言語能力と言語が共進化する描像を提示した。

2.3 計算機シミュレーションによる言語の文化進化の構成論的手法

前節までは、言語に関する能力を扱った。ここでは、言語に関する性質を扱う。言語の特定の性質に着目した計算機シミュレーションの代表的な研究を紹介する。これまでの言語のどんな普遍的な性質や文法現象が注目され、その創発条件が解明されてきたのかについて、簡単に示す。

2.3.1 構成性の創発

構成性とは言語の普遍的性質の一つであり、2つの方向から定義づけが可能である。全体を部分に還元できる性質（全体→部分）という意味を強調すれば、「複合的な意味をより単純な意味の規則的な組合せとして表現できる性質」（上田他, 2024）という「分節化可能な程度」として定義でき、複数の部分から全体を構築する性質（部分→全体）という意味を強調すれば、形式全体の意味が部分の意味から決定される性質 (Cann, 1993) という「合成可能な程度」として定義できる。

この構成性が文化進化によって創発した可能性が提示されている。まず、文化進化の過程を、親の発話と子の学習が循環する世代間継承の過程としてモデル化した繰り返し学習モデル (Iterated Learning Model; ILM) (Kirby, 2001) がある。初めの世代の言語が仮に形式と意味がランダムな対応関係にある総体的言語 (holistic

language)であったとしても、子が汎化学習できるのであれば、徐々に形式と意味の対応関係に体系性が生じ、構成的言語が創発する。エージェントベース・シミュレーション (Kirby, 2001, 2002; Kirby et al., 2015) やニューラルネットワーク (Ren et al., 2020) や数理モデル (Brighton, 2002; Griffiths and Kalish, 2007; Kirby et al., 2007), 実験室言語進化実験 (実験記号論) (Kirby et al., 2008, 2015; Scott-Phillips and Kirby, 2010) によって例証されてきた。

2.3.2 再帰性の創発

再帰性とは言語の普遍的性質の一つであり、一般的な定義としては、計算の出力が入力となる性質である。言語表現を例にとれば、“John thinks that Mary believes that …” といった埋め込み表現に現れるし、生成文法における再帰性は、埋め込み表現だけでなく、集合形成の繰り返しでもある (2.2.1 節)。

Kirby (2002) では、あくまでも、無限の表現を作る埋め込み表現の創発をシミュレーションで再現した。ここでは意味構造において埋め込み表現を加えている。そして、最初の世代では

$$“S/\text{believes}(\text{heather}, \text{praises}(\text{gavin}, \text{mary})) \rightarrow \text{uic}”$$

といった総体的言語であっても、

- (1) “S/p (x, q) $\rightarrow$ S/q C/p gp B/x d”
- (2) “S/p (x, y) $\rightarrow$ stlw A/p B/y B/x”

といった埋め込み表現を表現できる汎化したルールが創発した。これらのルールから再帰的な埋め込み表現が可能である。自然言語における、“I think that John believes that Mary says hello” を想定してみる。これは、

$$S/\text{think}(I, \text{believes}(\text{John}, \text{says}(\text{Mary}, \text{hello})))$$

と書ける。ルール (1) を適用するが非終端記号 “q” が右辺に存在するため、再度ルール (1) が呼び出され、最後にルール (2) が呼び出され、展開される。つまり、ルール (1) が存在する言語知識は再帰的埋め込み表現を可能とする。

2.3.3 格や一致の創発

van Trijp (2010) では、代名詞の格システムの変遷について研究し、スペイン語の代名詞に残された格 (case) の消失を調べた。格とは、主格や対格のように、事態参与者の役割を表示するものである。かつてのスペイン語の 3 人称代名詞で主格・対格・与格などで異なる形態を付与する格システムであったが、現在は廃れ、性と有生性を区別する参照システム (referential system) に移行しつつあるという

記述研究に基づく．3人称代名詞において性 (Gerner) を区別できない場合に，性を区別できるような格システムが新たな戦略として用いられるように成ることを，マルチエージェントシミュレーションを用いて再現した．

また，Beuls and Höfer (2011) では，一致 (agreement) の創発を研究した．一致とは，数 (Number) ・ 性 (Gender) ・ 人称 (person) を持つ要素が他の要素ともその性質を共有する現象である．例えば，英語の現在形において主格である名詞が持つ3人称単数という性質が，動詞という他の要素において，-s を付加されているという現象である．この研究では，内在的一致⁶(internal agreement) を取り扱う，スペイン語が例に挙げられているが，“決定詞＋形容詞＋名詞”といった名詞句において，名詞の性や数に，決定詞と形容詞が一致する現象を指す．世界 W の部分集合であるコンテキスト C を共有する2人は，2つの対象を参照する「記述ゲーム (description game)」を行う．話し手エージェントの「記述」に聞き手エージェントが同意すれば，コミュニケーションが成功したとみなす．2つの対象とは，例えば，「tall man(上背のある男性)」と「old cheese(古いチーズ)」のようなものである．結果として，2人のエージェントが記述する必要がある2つの対象を区別するのにかかる認知負荷を下げることができる手段として，内在的一致の現象が創発する．言語処理の負荷を考慮する場合，性質を共有する要素同士が一致により対応関係を持つことで，より処理がしやすいために創発した可能性を示した．

2.3.4 本研究の対象とする言語現象：構文交替

本研究では，概念化という人間普遍的な認知現象と，それを構成的かつ体系的に表現する文法現象に着目する．言語使用者の異なる概念化によって，使われる表現が異なる構文交替という現象がある．しかしながら，2つの表現において完全に異なる形式が使われているわけではない (図 2.3)．上から見ていくが，“A resemble B” と “B resemble A” という表現は，類似性を共有する二つの事物，例えば2枚のTシャツを記述する点で，外界が同じである．しかし，言語使用者がAを基準にBとの類似性を報告するのか，Bを基準にAとの類似性を報告するのかによって，異なる表現が使われている．“A do B” と “B be-done-by A” に関しても，一方から他方への行為を記述している点で外界は同じであるが，例えば，言語使用者がAの行為者性を強調して報告するのか，Bの被害者性を強調して報告するのかによって，異なる表現が使われている．また，一辺にまとめてしまったが，“A on/above/after B” と “B under/below/before A” においても，AとBの上下前後というある特定の空間的關係を記述する点で，外界は同じである．しかし，言語使用者がどちらを基準とするのかによって，異なる表現が使われる．最後の “load A onto B” と “load B with A” という表現も，一方が他方に積む作業をしている点で，記述する外界は同じである．しかしながら，言語使用者がどちらに焦点を当てているのかによって異なる表現となる．

⁶定訳を見つけられなかった．

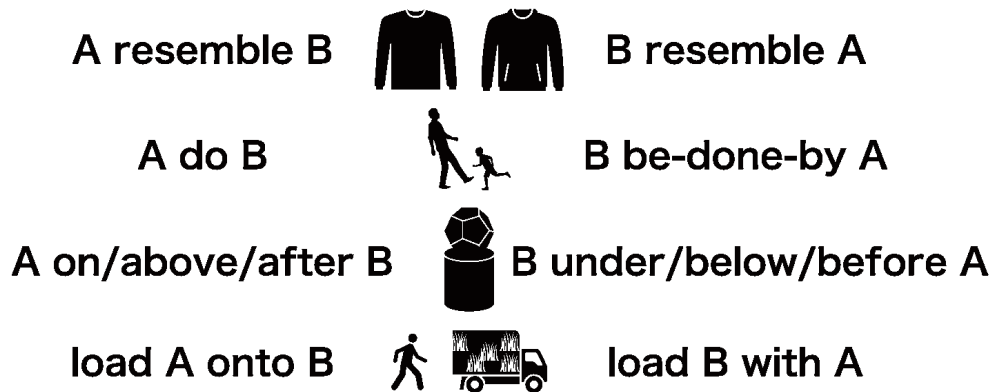


図 2.3: トラジェクターとランドマークの異なる選択

但し、図 2.3 で見たように、同一状況に対して異なる概念化をしたとしても、形式要素が部分的に共通している場合があるからだ。本来、異なる概念化が異なる形式で表現されていることになんの問題もない。なぜなら例えば、日本語と英語において異なる語彙と文法に基づいて全く異なる形式で表現されているからだ。異なる概念化によって全く異なる形式で表現されることは、言語間差異を見れば明らかである。同一言語内において見つけることは難しいが、例えば、食事行為において「食べる」と「召し上がる」などの敬語表現や、「教師」と「先生」という表現などが挙げられる。

いずれにせよ、本研究では、図 2.3 で見たように、同一状況に対して異なる概念化をしているが、依然としてある程度の「同じさ」を持つ形式で表現を使い分けるような構文交替がどのように創発するのかを検討する。この問いに対して、シミュレーションを用いて、聞き手が話し手の概念化を正しく推論できる確率についての条件を明らかにする。「同一状況に対して異なる概念化」することを明確にするため、同一状況には2つの要素とその関係性のみが含まれることとし、そこにおける概念化とは、2つの要素を「どのような関係性」として概念化するのかという自由度のみを許すこととする。その上で、

- 「二者関係」を「概念化」する言語使用者にとって、その概念化を「反映」することができる言語体系は創発するのか？
- 創発のメカニズムはどのようなものか？
- 仮に反映できるならば、学習時にどのような条件が必要であるか？
- 仮に反映できるならば、どのような「反映」のさせ方なのか？
- 「反映のさせ方」を捉える評価指標はあるか？

という問いに分解する。なお、1つ目の問いはモデルの設計上、どのような反映のさせ方を問わなければ、「反映できる」状態から開始される。ただし、これは全て

をランダムな形式で発話しているだけであり、概念化の「違い」を反映できていないわけではない。

2.4 言語の性質を測る評価指標

シミュレーション上で創発する言語は人間にとって理解できないため、創発言語が人間言語の諸性質を持っているのかは自明ではない。それは進化言語学にとって大問題である。というものの人間言語の進化を構成論的に探求しているのか、計算機上の呪文を解読しているのか判別できないからだ。そこで、創発言語と実際の人間言語とどれほど近いものなのか評価する指標が提案され、使用されてきた。進化言語学の文脈で最も代表的な評価指標は構成性に関するもので、*Topographic Similarity*(*TopSim*) と呼ばれる (Brighton and Kirby, 2006; Chaabouni et al., 2020)。逐語訳すれば「地形的類似性」である。3.4 小節において詳しく問題点を指摘するが、この評価指標は概念化を含めた意味構造を考慮に入れた計算ができないものの、様々な研究で使われてきた。創発言語と人間言語の共通性を考える上で、構成性は特に重視されてきた。構成性とは、全体の意味が部分の意味とその組み合わせで表現されるという性質であった。これの定義に従うならば、部分の意味とその組み合わせが「似ている」ならば全体の意味も似ており、それらの意味に対応する形式同士も「似ている」ことが予測される。逆に、部分的な意味とその組み合わせが「似ていない」ならば全体の意味の似ず、それらの意味に対応する形式同士も「似ていない」ことが予測される。こうした、「似ている」「似ていない」を n 次元空間上の点同士の距離として捉えたのが図 2.4 である。構成的な言語は、ちょうど真ん中の図のような、意味と形式の対応関係にあるときである。

このように意味と形式を空間的に把握することで、TopSim は言語全体の統計的性質を見る。言語が構成的ならば、意味が類似していると形式も類似し、意味が類似していないと形式も類似しない (上田他, 2023) という直感を TopSim は表現できる。具体的には、2つの意味の距離と各々に対応する形式の距離の相関係数を計算する。相関係数の計算方法にはバリエーションがあり、ピアソンの積率相関係数 Brighton et al. (2005) とノンパラメトリックなスピアマンの順位相関係数 Chaabouni et al. (2020) とがある。要点を先取りして述べておく。相関係数 ρ, ρ_s の定義から $-1 \leq \rho, \rho_s \leq 1$ であり、0に近いほどその言語知識は構成的ではないし、1に近いほどその言語知識は構成的である⁷。

ここで：

- J は意味表現のインデックス集合、または形式表現のインデックス集合であり、ペアの一意性を確保するため $j, k \in J$ には順序 $j < k$ がある。

⁷もちろん -1 に近い値を取ることはあるが、いわば一義語と多義語のみ占められる言語知識であり、一義語とが発生しない設定にしているため、ここで言及は止める。

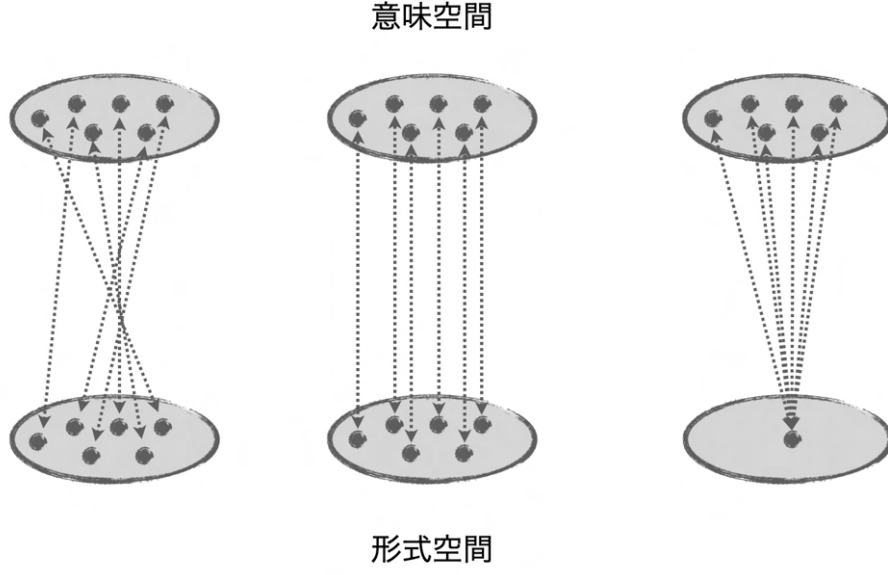


図 2.4: 再掲：意味空間と形式空間の地形 (Culbertson and Kirby (2016) を元に作図)

- Δm_i は2つの意味 $m_j, m_{k \neq j}$ の距離であり, Hamming 距離として定義される:

$$\Delta m_i = d_H(m_j, m_k)$$

- Δs_i は2つの形式 $s_j, s_{k \neq j}$ の距離であり, Levenstein 距離として定義される:

$$\Delta s_i = d_L(s_j, s_k)$$

- n は意味あるいは形式の距離のペアの総数は, 次の式で定義される:

$$n = |\{(m_j, m_k) \mid j, k \in J, j < k\}| = |\{(s_j, s_k) \mid j, k \in J, j < k\}|$$

- $\overline{\Delta m}$ と $\overline{\Delta s}$ は, それぞれ Δm_i と Δs_i の平均値を表す:

$$\overline{\Delta m} = \frac{1}{n} \sum_{i=1}^n \Delta m_i, \quad \overline{\Delta s} = \frac{1}{n} \sum_{i=1}^n \Delta s_i$$

ピアソン積率相関係数 ρ を使った TopSim は, 次の式で定義される:

ピアソン積率相関係数の TopSim

$$Topsim = \rho(\Delta m, \Delta s) = \frac{\frac{1}{n} \sum_{i=1}^n (\Delta m_i - \overline{\Delta m})(\Delta s_i - \overline{\Delta s})}{\sqrt{\frac{1}{n} \sum_{i=1}^n (\Delta m_i - \overline{\Delta m})^2} \cdot \sqrt{\frac{1}{n} \sum_{i=1}^n (\Delta s_i - \overline{\Delta s})^2}}$$

スピアマン相関係数 ρ_s を使った TopSim は, 集合に基づく次の式で定義される:

スパマン相関係数の TopSim

$$Topsim = \rho_s(\Delta m, \Delta s) = \{ (d_m(j, k), d_s(j, k)) \mid j, k \in J, j < k \}.$$

ここで重要なのが、どのようにして意味の距離と形式の距離を測るかという問題である。意味の距離と形式の距離は、それぞれ、Hamming 距離と Levenshtein 距離で計算するのが伝統的であるようだが、なぜ採用されるかという根拠は明示されていない。まず、意味は意味要素の順序付き集合として表現される。順序付き集合 $(a_1, a_2, \dots, a_n)$ は、次のように定義される：

$$\text{順序付き集合 } m_i = (a_1, a_2, \dots, a_n), \quad a_i \in S,$$

そして、意味の距離を測る Hamming 距離とは、2つの意味要素の順序付き集合を数え上げ差異を数え上げることである。Hamming 距離 d_H は、同じ長さを持つ2つのベクトル $m_1 = (m_{1,1}, m_{1,2}, \dots, m_{1,n})$ と $m_2 = (m_{2,1}, m_{2,2}, \dots, m_{2,n})$ の間の相違の数として定義される：

Hamming 距離

$$d_H(m_1, m_2) = \sum_{i=1}^n \delta(m_{1,i}, m_{2,i}),$$

ここで：

$$\delta(a, b) = \begin{cases} 0, & \text{if } a = b, \\ 1, & \text{if } a \neq b. \end{cases}$$

例えば、 $\{kick, john, mary\}$ と $\{kick, carol, bob\}$ 比較すると、*kick* のみが同じであり、それ以降の2つの意味要素が異なる。よって、hamming 距離は2となる。また、Levenshtein 距離とは、Edit 距離とも言われ、一方の形式から他方の形式に近づけていく時に何回「編集」を行うかを計算する。形式の距離を測る Levenshtein 距離の計算方法には最大形式長で標準化 (Normalized) しないもの (Brighton and Kirby, 2006) とするもの (Kirby et al., 2008) とがある。Levenshtein 距離は以下で定義される：

Levenshtein 距離

$$d_L(s_j, s_k) = \min\{k \mid k \text{ 回の編集操作で } s_j \text{ を } s_k \text{ に変換できる}\}.$$

編集回数が少なければ2つの形式の類似度は高く、逆に編集回数が多ければ2つの形式の類似度は低い、というのが直感的な理解である。ここで、「編集操作」は以下の3つから成り立つ：

- 挿入： s_j の任意の位置に1文字を追加する操作。

- 削除： s_j の任意の位置から 1 文字を削除する操作.
- 置換： s_j の任意の位置の 1 文字を別の文字に置き換える操作.

この定義では、2つの文字列 s_j と s_k は 1 文字単位で比較され、編集距離 $d_L(s_j, s_k)$ は、これらの操作を最小限で実行するために必要な回数として計算される．以下で編集操作の例を示す

- 挿入： $ab \rightarrow abc$ ならば、“c” を挿入（編集距離=1）
- 削除： $abc \rightarrow c$ ならば、“a” と “b” を削除（編集距離=2）
- 置換： $ab \rightarrow ad$ ならば、“b” を “d” に置換（編集距離=1）

となる．ただし、この計算方法のままだと、“ $ab \rightarrow cd$ ” の編集距離と “ $ab \rightarrow abcd$ ” の編集距離がともに 2 となり、“ab” という形式要素を共有する後の方が似ているという人間の直感を取りこぼしている．この問題を解決するために、2つの形式のうちの最大形式長で標準化する方法が提案されている．これが、Normalized Levenshtein distance である (Kirby et al., 2008)．これにより、形式長によるスケールを相対化する．

Normalized Levenshtein distance

$$d_{NL}(s_j, s_k) = \frac{d_L(s_j, s_k)}{\max(|s_j|, |s_k|)}.$$

ここで、 $|s_j|$ および $|s_k|$ は、それぞれ文字列 s_j と s_k の長さである．この標準化により、形式長に依存しないスケールで類似度を比較できる．

具体的にさまざまな性質の言語知識をそれぞれ計算し、構成性がどのように相関係数として現れるのかを直感的に説明する．

まず、非常に規模の小さい言語知識について検討する．意味と形式の対応関係が非常に綺麗な言語知識を用意した．

表 2.1 の意味要素から作れる言語知識は、表 2.2 の 6(= 2 × 3 × 1) つである．見やすさのために、形式表現は、意味表現の頭文字を使う．

表 2.1: 意味要素

意味要素			
i	関係性要素 P_i	j	対象要素 X_j
1	kick	1	john
2	love	2	mary
3	help		

表 2.2: 1 つ目の言語知識

言語知識	
1	kick(john, mary) → jkm
2	kick(mary, john) → mkj
3	love(john, mary) → jlm
4	love(mary, john) → mlj
5	help(john, mary) → jhm
6	help(mary, john) → mhl

TopSim の直感は、類似の意味なら形式も類似で、異なる意味なら形式の異なる、というものである。

- 言語知識 1 の $kick(john, mary) \rightarrow jkm$ と、言語知識 2 の $kick(mary, john) \rightarrow mjk$ は、意味表現において、 $(john, mary)$ と $(mary, john)$ のように引数の順序が異なる。意味の距離は 2 と「異なる」と判断する。もしこの言語知識全体が構成的でならば、形式表現においても、「異なる」という判断がされるほど距離が遠くあってほしい。すると、 jkm と mkj というように、順序が異なり、距離は 2 である。よって、このペアにおいては TopSim の直感に沿う。
- 言語知識 1 の $kick(john, mary) \rightarrow jkm$ と、言語知識 3 の $love(john, mary) \rightarrow jlm$ は、意味表現において、関係性要素 P の $kick$ と $love$ だけが異なる。意味の距離は 1 である。もしこの言語知識全体が構成的ならば、形式表現においても、異なりが最小限であってほしい。すると、 k と l だけが異なる。形式の距離は 1 である。よって、このペアにおいては TopSim の直感に沿う。
- 言語知識 1 の $kick(john, mary) \rightarrow jkm$ と、言語知識 4 の $love(mary, john) \rightarrow mlj$ は、意味表現において、全てが異なる。意味の距離は 3 である。もし、この言語知識全体が構成的ならば、形式表現も全く異なるようになってほしい。すると、 jkm と mlj と全く異なる。形式の距離は 3 である。よって、このペアにおいても TopSim の直感に沿う。

このようなことを全てのペアで行う。この過程で、意味の距離と形式の距離という 2 変数を計算しているので、それらの相関係数を計算することになる。

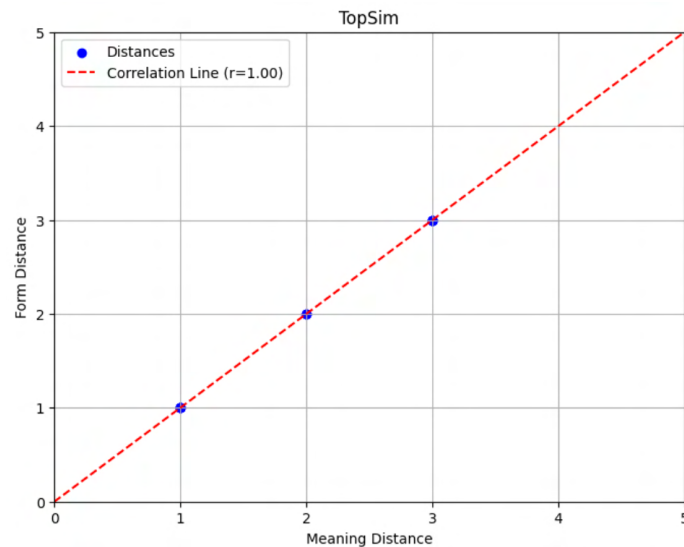


図 2.5: 完全に構成的な言語知識

この結果は人間の直感だけでなく、部分の組み合わせから全体を構築するという構成性の定義にも沿っていることがわかる。というのも、始めに「見やすさのた

めに、形式表現は、意味表現の頭文字を使う。」と述べたが、これはつまり、*john* という意味が形式 *j* に対応し、*mary* という意味が形式 *m* に対応していることが示されており、かつ、その順序も意味と形式の距離ともに考慮されている。順序が異なればそれだけ距離も遠くなるように計算されている。

次に、意味と形式の対応関係が、先ほどよりは整理されていない言語知識 (表 2.4) について見てみる。

表 2.3: 意味要素

意味要素			
i	関係性要素 P_i	j	対象要素 X_j
1	kick	1	john
2	love	2	mary
3	help		

表 2.4: 2 つ目の言語知識

言語知識	
1	kick(john, mary) $\rightarrow$ jkm
2	kick(mary, john) $\rightarrow$ mkj
3	love(john, mary) $\rightarrow$ <u>abc</u>
4	love(mary, john) $\rightarrow$ mlj
5	help(john, mary) $\rightarrow$ <u>def</u>
6	help(mary, john) $\rightarrow$ mhl

見やすさのために、「整理されていない」言語知識に下線を引いた。

- 言語知識 1 の *kick(john, mary) $\rightarrow$ jkm* と、言語知識 3 の *love(john, mary) $\rightarrow$ abc* は、意味表現において、関係性要素 P の *kick* と *love* だけが異なる。意味の距離は 1 である。もしこの言語知識全体が構成的ならば、形式表現においても、異なりが最小限であってほしい。しかし、*jkm* と *abc* と全く異なる。形式の距離は 3 である。よって、このペアにおいては TopSim の直感に沿わない。
- 言語知識 5 の *help(john, mary) $\rightarrow$ def* と言語知識 6 の *help(mary, john) $\rightarrow$ mhl* は、意味表現において、引数の順序が異なる。意味の距離は 2 である。もしこの言語知識全体が構成的ならば、形式表現においても、全体的に同じような「異なり」となってほしい。形式表現において、*def* と *mhl* と全く異なり、形式の距離は 3 である。よって、これだけでは TopSim の直感に沿っているかは分からない。

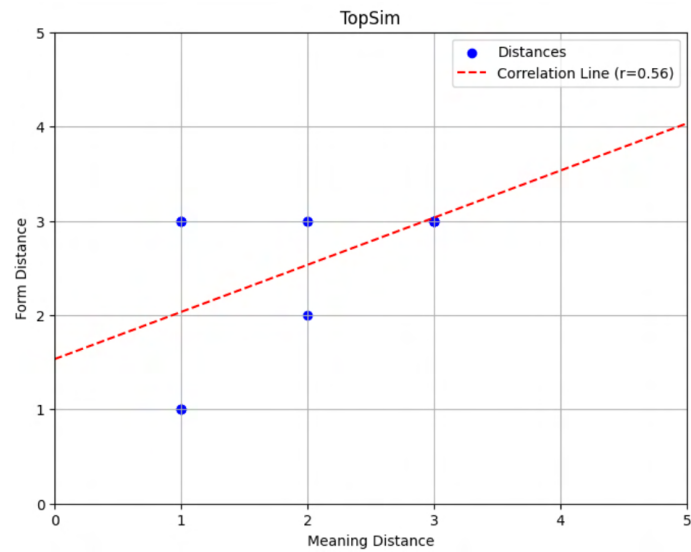


図 2.6: 部分的に構成的な言語知識

最後に、直感的にも、全く整理されていないと判断できるような言語知識 (表 2.6) の構成性をみる。形式表現はアルファベットを 3 つ単純に並べただけである。

表 2.5: 意味要素

意味要素			
i	関係性要素 P_i	j	対象要素 X_j
1	kick	1	john
2	love	2	mary
3	help		

表 2.6: 3 つ目の言語知識

言語知識	
1	kick(john, mary) $\rightarrow$ abc
2	kick(mary, john) $\rightarrow$ def
3	love(john, mary) $\rightarrow$ ghi
4	love(mary, john) $\rightarrow$ jkl
5	help(john, mary) $\rightarrow$ mno
6	help(mary, john) $\rightarrow$ pqr

すでに説明が不要のように思えるが、意味の距離に関わらず、形式の距離は常に 3 である。

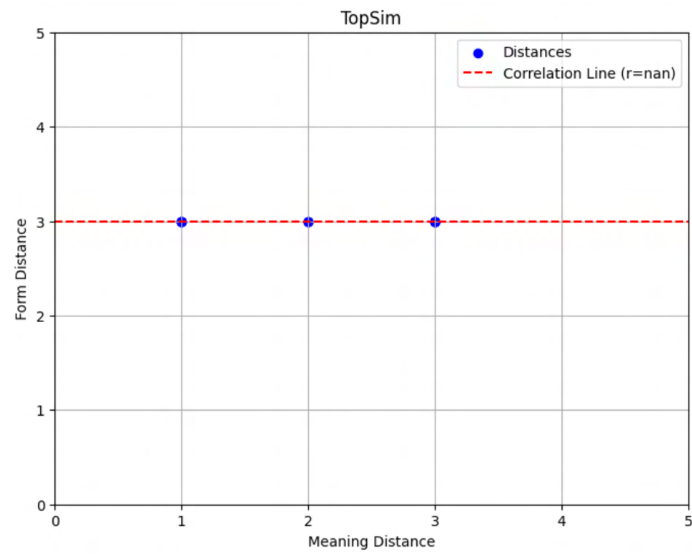


図 2.7: 全く構成的でない言語知識

これまで見てきたように，構成性の評価指標として代表的な TopSim は，人間が手作業で観察して素朴に感じる構成性への直感を捉えることができているように見える．

第3章 概念化搭載繰り返し学習モデルと評価指標の改変

本章では、言語使用者の「概念化」を言語の文化進化研究に取り入れた「概念化搭載繰り返し学習モデル」について説明する。そのために、本研究で明示的に取り入れる「意味」の構造を説明し、次に、本研究のモデルの基盤である言語の文化進化モデルである繰り返し学習の設計と構成要素について詳細に説明し、本研究でどのように拡張するのかを述べる。

3.1 本モデルにおける「意味」の定義

本モデルにおける意味は、3.2.1 節で述べた認知言語学が依拠するところの概念化という認知過程を含んだ意味とする。すなわち、真理値が決まる外部世界の構成要素ではなく、外部世界を言語使用者が主観的に解釈した結果形成される意味内容とする。というのも、これまでの意味の設計から前提とされる意味の定義には議論の余地があるからである。例えば、述語論理（確定節文法）(Kirby, 2002)、素性 (feature) の順序集合 (Kirby et al., 2015)、属性 (attribute) の順序集合 (Ren et al., 2020) が意味の設計に使われているが、こうした設計は、外部世界の集合が意味であることを前提にし、また、意味の構成要素同士の“関係性”を概念化を通して操作する自由度はないという制約を内包する、と考えられるからだ。

3.2 繰り返し学習モデルの設計

Kirby et al. (2014) では、「繰り返し学習 (Iterated Learning)」を再帰的な表現を使って、文化進化一般に当てはまる定義を与える。

Iterated learning:

The process by which a behaviour arises in one individual through induction on the basis of observations of behaviour in another individual who acquired that behaviour in the same way. (Kirby et al., 2014)

繰り返す学習とは、他個体の行動を観察と帰納推論して獲得した行動を、また他の個体が同じく観察と帰納推論して獲得した行動をするという過程である。(拙訳)

すなわち、内在的な能力や知識は直接継承（遺伝など）されることはなく、外在化した行動の観察と学習を通じて間接的に継承される文化進化の過程である。外在化した行動を言語に当てはめたのが図 3.1 である。各世代の構成人数を 1 人に限定したならば、他個体とは前世代の個体であり、それすなわち「親エージェント」である。またその学習する個体は「子」ということとしてモデル化されている。親の脳内にある言語知識は子供に「直接継承されない」ので、その言語知識は発話される。そして子供はそれを受けとって、親の言語知識を間接的に継承する。つまり、言語の世代間継承としての繰り返し学習モデルでは、親の発話と子の学習が循環する。

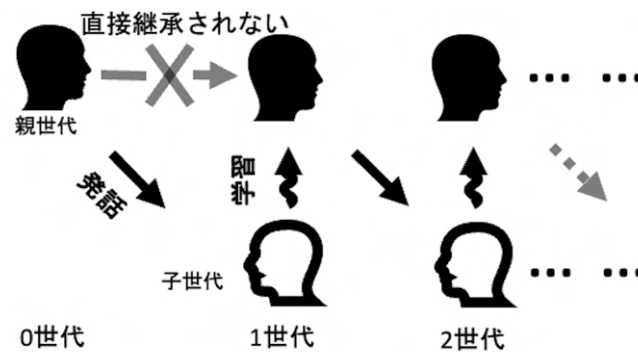


図 3.1: 繰り返し学習モデル

つまり、繰り返し学習モデルの設計には、親が発話した言語を、子が受け取って学習し、学習した知識に基づいて発話し、その発話された言語を受けとって、次世代の子供が学習することをモデル化する必要がある。構成要素を整理すると、

- 世代構成員としての「親」と「子」
- 親が発話を求められる「意味空間」
- 親の「発話アルゴリズム」
- 親の実際の「発話」
- 子の「学習アルゴリズム」
- 世代交代：言語を継承した親は死ぬ
- 新たに誕生する「子」

が必要である。特に重要なことは、発話アルゴリズムと学習アルゴリズムである。ここでは、Kirby (2002) とそれを文法化現象へと拡張させた橋本・中塚 (2007) をもとにして説明する。

言語知識の表し方：確定節文法

多くの言語の文化進化研究と同様、Kirby (2002) においても言語は意味と形式の対応表であるとしてモデル化され、その対応表全体を言語知識と呼ぶ。意味と形式を対応させるため、意味から形式を導出するルールとして文脈自由文法を拡張させた確定節文法を用いる。確定節文法とは、非終端記号・終端記号・開始記号・導出規則に加え、非終端記号に付与される条件 C がある。

非終端記号とは、文法規則において左辺に現れ、右辺を導出する記号を指す。

まず、大文字アルファベットの集合 N を以下で定義する：

$$N = \{X \mid X \in \{A, B, C, \dots, Z\}\}$$

非終端記号の集合 N^* は、 N から C と P と S を除いて、以下で定義される：

$$N^* = \{X \mid X \in N, X \neq C, P, S\}.$$

ここで、 C は条件で、 P は導出規則の集合で、 S は文法の開始記号、である。

終端記号とは、小文字アルファベット Σ の要素からなる形式列 $\sigma \in \Sigma^L$ の要素を指す。小文字アルファベット Σ は、非終端記号の集合 N と区別される別の有限集合であり、以下で定義される：

$$\Sigma = \{x \mid x \in \{a, b, c, \dots, z\}\}.$$

ただし、 $L \geq 1$ を満たす有限長の形式列であり、終端記号の集合 Σ^* は以下で定義される：

$$\Sigma^* = \bigcup_{L \geq 1} \Sigma^L,$$

ここで、空列は除外される。

開始記号とは、文全体を導出するために特別に指定された非終端記号であり、文法規則において左辺に登場する記号を指す。通常、開始記号には S が用いられ、その役割は以下で定義される：

$$S \in N,$$

ここで、 N は非終端記号の集合である。

導出規則とは、非終端記号 N 、終端記号 Σ 、および条件 C^* を用いて、左辺から右辺への導出を規定するルール全体を指す。各導出規則 P は以下で定義される。

$$P = \{(N_{i_1}/C_{i_2}^* \rightarrow V_{i_3}^*) \mid N_{i_1} \in N, C_{i_2}^* \in C^*, V_{i_3}^* \in (N^* \cup \Sigma^*) \cup (N^* \cap \Sigma^*)\}.$$

ここで、「 $(N_{i_1}/C_{i_2}^* \rightarrow V_{i_3}^*)$ 」という導出規則に関して、左辺に開始記号と非終端記号を含む集合 N の要素 i_1 が1つと条件の集合 C^* の要素 i_2 があり、右辺に終端記号の集合 Σ^* または非終端記号の集合 N^* との、組み合わせ集合 V^* の要素 i_3 を持つ導出規則である。文脈自由文法では無条件に N から V^* を導出するが、確定節文法では、条件 C^* の要素を満たす場合に限り導出が可能である。

条件 C^* は一階述語論理を用いた述語項構造で表される。述語 *predicate* は2つの項 *argument* を引数として受け取る。

$$predicate_i(argument_j, argument_{k \neq j})$$

モデルでは、この条件 C^* を「意味」としている。ここで留意すべきことは、“predicate” は「述部」を表す単語であるが、これは「他動詞」に限定されるわけではない。あくまで「2つの事物を特定の方法で関係づける」意味要素である。

人間言語において、述語の第1引数 $argument_j$ と第2引数 $argument_{k \neq j}$ は直積ではなく、順序対 $\langle argument_j, argument_{k \neq j} \rangle$ であることは重要である (田中, 2016, p. 41)。つまり、 $(argument_j, argument_{k \neq j})$ と $(argument_{k \neq j}, argument_j)$ は順序の観点から同一ではない。それ故、もちろん、 $predicate_i(argument_j, argument_{k \neq j})$ と $predicate_i(argument_{k \neq j}, argument_j)$ も異なる状況を描写している。なぜなら、直積であれば、引数は集合として扱われるため、自然言語における主語と目的語の区別が完全になくなる。一方、順序対であれば、第1引数と第2引数が主語なのか目的語なのかは定かではないものの、しかしそのどちらかであるという区別は明瞭に示してあるためだ。

言語知識のうち、例えば、

- 文ルール： $S/kick(john, mary) \rightarrow abc$
- 単語ルール： $N/john \rightarrow gh$

と表現される。単語ルールの非終端記号は、語彙範疇 (lexical category) に相当する。

ここで、文ルールは変数 x を持つことができる。これは「汎化ルール」と呼ばれる。変数には述語位置が p 、第1引数の位置が x 、第2引数の位置が y とする。次の学習アルゴリズムに進む前に、4種類の文ルールを定義しておく。

4つのルール：文ルール3つと単語ルール1つ

- 総体論ルール (Holistic Rule)：文ルールのうち、条件 C^* に変数がないルール (例： $S/kick(john, mary) \rightarrow abc$)
- 1スロットルール (Generation Rule (1 var))：文ルールのうち、条件 C^* に変数が1つあるルール (例： $S/p(john, mary) \rightarrow aQ/pc$, $S/kick(x, mary) \rightarrow W/xbc$, $S/kick(john, y) \rightarrow abR/y$)

- 2スロットルール (Generation Rule (2 var)): 文ルールのうち, 条件 C^* に変数が2つあるルール例 (例: $S/p(x, mary) \rightarrow W/xQ/pc$, $S/kick(john, y) \rightarrow aQ/pR/y$, $S/kick(x, y) \rightarrow W/xbR/y$)
- 3スロットルール (Generation Rule (3 var)): 文ルールのうち, 条件 C^* に変数が3つあるルール例 (例: $S/p(x, y) \rightarrow W/xQ/pR/y$)
- 単語ルール (word rule): 条件 C^* に要素が1つのもの (例: $N/john \rightarrow gh$, $N/like \rightarrow nm$)

学習アルゴリズム: Chunk • Category-Integration • Replace

まず, 学習の大前提として, 学習によって発話可能な意味の集合, つまり導出できる終端記号の集合を減らすことはなく, 維持されるか拡張するかのどちらかである. それでは, 学習アルゴリズムについて説明する. 学習に必要な材料は, 親の言及対象としての「状況」と発話の「形式」である. この学習アルゴリズム全体としては汎化学習するためのものであり, 3つのルールベースのアルゴリズム Chunk, Category-Integration¹, Replace からなる.

• Chunk

2つの文ルールを比較し, 意味表現と形式表現においてそれぞれ1つのみ差異部分が存在する場合, 2つの文ルールを汎化したスキーマルールと, 差異部分にあたる単語ルールを生成する (差異部分には分かりやすくするため下線を引いた).

Rule 1: $S/kick(john, \underline{mary}) \rightarrow abc\underline{jk}de$

Rule 2: $S/kick(john, \underline{anya}) \rightarrow abc\underline{st}ude$

↓

Rule 3: $S/kick(john, y) \rightarrow abcQ/ydef$

Rule 4: $Q/mary \rightarrow jk$

Rule 5: $Q/anya \rightarrow stu$

このように, Rule 1 と Rule 2 から, 汎化した Rule 3 と, 差異部分を抜き出した単語ルール Rule 4 と Rule 5 とが作られる. Rule 3 の変数 y に付いたものと Rule 4 と Rule 5 における非終端記号 (この場合では Q) は, Rule 1 と Rule 2 の2つのルールから chunk によって作られたことが分かるよう, Rule 1 • Rule 2 ペアに固有な非終端記号が選ばれる. すなわち, 異なる文ルールの

¹Kirby and Christiansen (2003); 橋本・中塚 (2007) では, merge というタームで指示されているが, Hashimoto et al. (2010) で “a kind of category integration” であることが指摘され, 橋本 (2012) において, “category integration” という訳語が採用された. なお, この訳語の選択は生成文法における「Merge(併合)」との区別を明確にするためである.

ペアから生成されたスキーマルールには、それぞれ異なる非終端記号が割り当てられる。

- **Category-Integration**

2つの単語ルールを比較して、同じ意味要素を同じ形式表現で導出するルールであるならば、どちらか一方の非終端記号カテゴリーに統合する学習である。どの2つの文ルールから導かれた単語ルールなのか分かるために固有のカテゴリーが付与されているが、同じ意味を同じ形式で表している単語ルールであれば、どちらか一方を削除する。

Rule 1: $\underline{W}/anya \rightarrow jk$

Rule 2: $\underline{U}/anya \rightarrow jk$

Rule 3: $\underline{U}/kirk \rightarrow ld$

Rule 4: $S/kick(john, y)/0 \rightarrow abc\underline{U}/ydef$

↓

Rule 1: $W/anya \rightarrow jk$

Rule 5: $W/kirk \rightarrow ld$

Rule 6: $S/kick(john, y) \rightarrow abcW/ydef$

このように、Rule 1 と Rule 2 から、統合され Rule 1 だけが残る。そして、非終端記号の統合で選択されなかった Rule 2 の非終端記号 U をもつ他の Rule も全て、その非終端記号が書き換えられる。Rule 3 の非終端記号 U が W として Rule 5 では書き換えられている。注目すべきは、Rule 4 から Rule 6 への変化である。Rule 4 で使われている変数に付与された非終端記号（この場合では U ）が Rule 6 では統合された Rule 1 の非終端記号（この場合では W ）に書き換えられている。このように、統合の際に選択されなかった非終端記号を選択された非終端記号に、言語知識全体で書き換えるというものである。この学習アルゴリズムは導出可能な終端記号の集合を維持・拡張する (橋本・中塚, 2007)。なぜならば、単語ルールの非終端記号が統合されるということは、汎化したルールのスロットに代入できる単語ルールの数が増えるからである。今回の場合で言えば、Rule 6 の非終端記号 W に代入できる単語ルールは、統合前では Rule 1 のみであったのが、統合後では Rule 1 と Rule 5 の2つに増えている²。

- **Replace**

単語ルールと文ルールを比較し、単語ルールの意味と形式のペアが文ルール

²仮に U に統合されたならば、導出可能な終端記号の集合を維持される。統合前では Rule 2 と Rule 3 の2つで、統合後も Rule 2 と Rule 3 の2つであるからだ。

の中にも存在するならば，文ルールからその単語ルールを抜き出し，汎化したルールを生成する学習である．

Rule 1: $Q/\underline{mary} \rightarrow \underline{jk}$

Rule 2: $Q/anya \rightarrow mnm$

Rule 3: $S/kick(john, \underline{mary}) \rightarrow abc\underline{jk}de$

↓

Rule 1: $Q/mary \rightarrow jk$

Rule 2: $Q/anya \rightarrow mnm$

Rule 4: $S/kick(john, y) \rightarrow abcQ/yde$

このように，Rule 1に見られる意味 (*mary*) と形式 (*jk*) が Rule 3 にも存在することから，Rule 3 のその部分を Rule 1 で置き換えることができるので，Rule 4 を導くことができる．この時，Rule 4 に付与される非終端記号 Q は，もともと Rule 1 が持っていた非終端記号 Q が付けられる．この学習は，導出可能な終端記号の集合を拡張する可能性がある．というのも，Rule 1 の他に同じ非終端記号をもつ Rule 2 を，汎化した Rule 3 に代入することで，新しく $S/kick(john, anya) \rightarrow abcmnmde$ という導出規則が生成される，これまでにない終端記号を作ることができる．

発話アルゴリズム

次に，発話アルゴリズムについて説明する．発話とは「状況」を「言語化」する過程であり，つまり，意味を形式に対応させることである．モデルでは「状況」は意味を表す条件 c^* として，「言語化」は意味と形式を対応させる導出規則の生成を指す．付け加えると，生成される導出規則の右辺には1つも非終端記号が存在してはいけない，という制約がある．繰り返し学習モデルにおいては，親エージェントに対して発話すべき条件として意味表現 $S/predicate_i(argument_j, argument_{k \neq j})$ が与えられ，それに対応する右辺の終端記号として形式表現を導出することである．この過程は親の言語知識に基づいて行われる．

発話アルゴリズムには，親の言語知識の状態に応じて，以下のように分類できる．

- 言語知識にある導出規則をそのまま使用するパターン

発話条件: $kick(john, mary)$

言語知識: $S/kick(john, mary) \rightarrow abcjkde$

↓

発話過程: $S/kick(john, mary) \rightarrow abcjkde$

つまり、発話要求されている意味表現とそれに対応する形式表現を導出する文ルールが言語知識に「丸ごと」存在している場合がこれに当てはまる。

- 言語知識にある汎化ルールに単語ルールを代入するパターン

発話条件： $kick(john, mary)$

言語知識： $Q/mary \rightarrow jk$

$$S/kick(john, y) \rightarrow abcQ/Yde$$

↓

発話過程： $S/kick(john, mary) \rightarrow abcjkde$

- 複数の方法があれば、ランダムにどちらかを選択する

発話条件： $kick(john, mary)$

言語知識： $Q/mary \rightarrow jk$

$$Q/mary \rightarrow stu$$

$$S/kick(john, y) \rightarrow abcQ/Yde$$

$$S/kick(john, y) \rightarrow twqlQ/y$$

↓

発話過程： $S/kick(john, mary) \rightarrow twqlstu$

このように、言語知識において同じ条件 C^* で異なる終端記号を持つ導出規則が存在し、どちらを選んでも発話条件に応えることができる場合、いずれかの導出規則を選択する。

- 単語ルールを1つ創作 (invention) して、それを汎化ルールに代入するパターン

発話条件： $kick(john, mary)$

言語知識： $Q/anya \rightarrow qwr$

$$S/kick(john, y) \rightarrow abcQ/yde$$

↓

発話過程： $S/kick(john, mary) \rightarrow abcmnmde$

$$Q/mary \rightarrow mnm \quad (\text{創作した導出規則})$$

この発話において、発話条件を満たす導出規則が既存の言語知識には存在していない。この場合、単語ルールに限り、1つだけ新しく作る「創作 (invention)」を行う。その際、同然、発話条件を満たすように非終端記号は代入される文ルールが持つ非終端記号を使用する。具体的に、非終端記号が Q の意味表現 $mary$ が存在しないので、即席で単語ルール $Q/mary \rightarrow mnm$ を新しく作り出す。この時、終端記号はのアルファベットから長さ L で生成される形式列

である。なお、この時作り出した単語ルールは後の発話場面において使用できるように言語知識に新しく組み込む。

- 発話条件を満たす文ルール全体を創作するパターン

発話条件： $kick(john, mary)$

言語知識： $S/love(susan, anya) \rightarrow ghn m$

↓

発話過程： $S/kick(john, mary) \rightarrow kapkd$

このように、発話要求に対して、既存の知識をそのまま使用したり代入したり、最低限の単語ルールを創作しても応えることができない場合、完全にあたらしく形式表現全体を作り出す。

3.2.1 概念化搭載繰り返し学習モデルの構成要素

本研究では、言語使用者が外界世界を「概念化」することを通して形成した概念内容を意味と捉える。そのため、先ほどと対比的に表現すれば、発話とは「状況」を「概念化」して「言語化」することである。また、学習は発話場面における言及対象としての「状況」と親の状況解釈としての「概念化」と発話「形式」である。そこで、本モデルにおいても、「概念化」を既存の繰り返し学習の意味構造に組み込むことをしている。概念化の本質は関係性の自由度であることから、そのみを捉えるため二者の関係性のみ発生するようにした。

以上を踏まえ、エージェントは、意味と形式の対応を確定節文法で記述した導出規則を言語知識として持つ。言語知識の大半は、前節の説明の基づく。まず、意味構造の設計について、概念化の本質は、関係性の自由度である。関係性の自由度を単純化させ、述語 ($predicate_i$) の取る項 ($argument$) を2つ ($j, k \neq j$) に限定する。これは、2つのうちどちらかに注目する非対称性という概念化の本質を捉えているためだ。本研究では、前節で条件 C^* として表されていた意味構造を大きく2つに分け、外部事象を述語項構造で、概念化を $CV(\text{Conceptualization Value}) = \{0,1\}$ で表す。二者の関係性しか概念化できる余地は残されていないため、2つのうちいずれかを概念化することを示す値である。そして、この概念化する2つの可能性を組み込んだ意味構造を以下のように定式化する：

$$predicate_i(argument_j, argument_k \neq j)/CV.$$

次に、言語の「形式」として表現される終端記号は、小文字アルファベット Σ の要素からなる長さ $L (L \geq 1 \text{ を満たす})$ 形式列 $\sigma \in \Sigma^L$ である。

図3.2を参照し、本モデルにおける導出規則を説明する。条件 C^* として「外部世界」がある。ここには、星形の物体 A と四角い物体 B が上下関係によって関係づけられている状況が表されている。この1つの状況 ($S/vertical(a, b)$) に対して、

a と b のどちらに注目を集めるかによって、異なるように概念化することができる。そして、異なる概念化によって形成された意味を、異なる方法 (“A above B, B under A” や “a 上 b, b 下 a”) で発話している場面である。

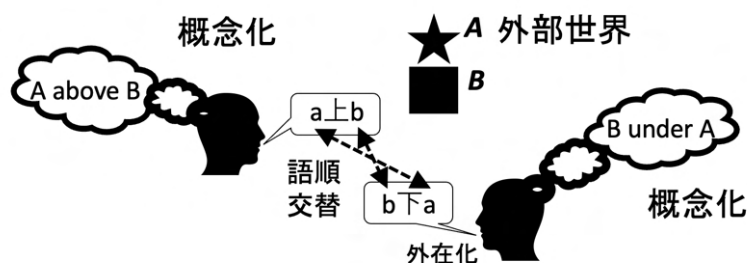


図 3.2: 二者関係の概念化

すなわち、文ルールは以下で表現する (分かりやすくするため、概念化の違いを単純な形で示した CV に下線を引いた。).

$$S/\text{vertical}(a, b)/\underline{0} \rightarrow abc$$

$$S/\text{vertical}(a, b)/\underline{1} \rightarrow def$$

学習アルゴリズム **Chunk** の修正 概念化の値を意味表現に組み込んだことで、学習アルゴリズムの chunk と発話アルゴリズムには少しの修正を加えた。

- **Chunk** : あらかじめ親がどのような概念化をしたのかに基づいて学習が行われる。そのため、概念化の値が 0 同士、1 同士の文ルールの比較だけ行われ、概念化の値が 0 の文ルールと概念化の値が 1 の文ルールの比較は行われない。

学習が行われる例を以下で示す。

$$\text{Rule 1: } S/\text{kick}(\text{john}, \underline{\text{mary}})/0 \rightarrow \text{abcjkde}$$

$$\text{Rule 2: } S/\text{kick}(\text{john}, \underline{\text{anya}})/0 \rightarrow \text{abcstude}$$

↓

$$\text{Rule 3: } S/\text{kick}(\text{john}, Y)/0 \rightarrow \text{abcQ/Ydef}$$

$$\text{Rule 4: } Q/\text{mary} \rightarrow \text{jk}$$

$$\text{Rule 5: } Q/\text{anya} \rightarrow \text{stu}$$

このように、学習前に持っている言語知識である Rule 1 と Rule 2 から、汎化した Rule 3 と、差異部分を抜き出した単語ルール Rule 4, Rule 5 とが作られる。Rule 3 の変数に付いたものと Rule 4 と Rule 5 における非終端記号

(この場合では Q) は、Rule 1 と Rule 2 の 2 つのルールから chunk によって作られたことが分かるよう固有なものが選ばれる。

学習が行われない例を以下で示す。

Rule 1: $S/kick(john, \underline{mary})/0 \rightarrow abcj\underline{k}de$

Rule 2: $S/kick(john, \underline{anya})/1 \rightarrow abcstude$

↓

Rule 1: $S/kick(john, \underline{mary})/0 \rightarrow abcj\underline{k}de$

Rule 2: $S/kick(john, \underline{anya})/1 \rightarrow abcstude$

このように、学習前に持っている言語知識である Rule 1 と Rule 2 は、概念かの値 CV が異なっている。そのため、学習はそもそも起こらない。

学習アルゴリズムの最小限の修正 学習アルゴリズムにおいて、なぜ chunk 以外の学習アルゴリズム category-integration と replace において変更がないのか、より正しく言い換えれば、なぜ他 2 つの学習アルゴリズムでは概念化の値を推論する必要がないのか、という点について理由を述べる。そのために、なぜそもそも Chunk 学習において概念化の値を推論する必要があるのかをまず述べる。それは開始記号 S を持つ「2 つの文ルールだけ」を学習材料としているからである。2 つの文ルールを比較するということは、2 つの異なる状況を比較することである。それはつまり、2 つの状況から共通部分と差異部分を検出し、その区別に対応する共通部分と差異部分を、2 つの形式表現を比較して、検出することであった。 $S/love(john, mary) \rightarrow abc$ と $S/love(john, anya) \rightarrow abst$ が Chunk 可能なルールペアであるのは、状況の差異部分 (mary と anya) と形式の差異部分 (c と st) を検出されるからだ。つまり、状況と形式の両者の差異部分を「対応づけられる」ことで初めて Chunk 学習ができるのだ。それでは、状況を概念化した結果を意味に含める本研究の立場からすれば、 $S/love(john, mary)/0 \rightarrow abc$ と $S/love(john, anya)/1 \rightarrow abst$ は Chunk 可能なルールペアではない。それは概念化も異なるからである。つまり、形式の差異部分を状況の差異部分に「対応づけられる」ことができない。なぜなら形式の差異部分 (c と st) が状況の差異部分 (mary と anya) に対応するのか、CV の差異部分 (0 と 1) に対応するのが判別できないからである。すると、判別するためには、状況に対する特定の解釈の仕方である CV が共通でなければならないのだ。

一方で、category-integration と replace では「文ルールだけ」で学習が行われるわけではない。「単語ルール」を学習材料の一部もしくは全部で使用する。「単語ルール」と呼んでいる言語表現の性質は、意味要素単体に関するものであると考えられる。本研究で使用する概念化とは、「関係性の自由度」に限定している。つまり、節で取り上げた記述の粒度に関わる「詳述性 (specifity)」は除外している。もしこれを組み込むならば、意味要素自体を概念化することになってしまい、ジョンは男性にも人間にも生き物にも記述としては可能である。しかし、それはちょ

うどイヌイットにおいて「雪」に関する単語が日本語よりもはるか豊富であるといった世界の分節化、語の分節化の話になってしまい、文法の話に辿り着けない。そのため、本研究では概念化を関係性の自由度に限定し、意味要素は概念化しないとする。そうすれば、どのような概念化をしたとしても、それに依存せず、ジョンはジョンであるし、メアリーはメアリーであるし、蹴るは蹴るである。このように対象 (object) や行為 (action) 単体の認識は概念化には依存しないと考える。つまり、単語ルールを持つということは、概念化とは独立に、行為それ自体や対象それ自体といった単一の意味要素を特定していることと見做す。そのため、単語ルールを使用する category-integration と replace ではアルゴリズム自体の挙動は変更していない。

発話アルゴリズムの修正 次に、発話アルゴリズムに関しては、chunk と同様、概念化の値を区別して発話するため、発話条件としての意味表現には CV が付与されている。但し、主要なアルゴリズムは先行研究に基づく。以下で、例を示す。

- 言語知識にある汎化ルールに単語ルールを代入するパターン

発話条件： $kick(john, mary)/1$

言語知識：Rule 1: $Q/mary \rightarrow jk$

Rule 2: $S/kick(john, y)/0 \rightarrow abcQ/Yde$

Rule 3: $S/kick(john, y)/1 \rightarrow nmqnqwQ/y$

↓

発話過程： $S/kick(john, mary)/1 \rightarrow nmqnqwjk$

このように、発話条件に含まれる CV と Rule 2 の CV は異なる値であるため、発話条件に答えるための導出規則を作ることができない。一方で、発話条件に含まれる CV と Rule 3 の CV は同じ値であるため、Rule 1 を代入することで発話条件に答える導出規則を生成できる

ここで留意すべきことは、「概念化の値を区別して発話する」ということと「概念化した意味を形式に反映する」ということを含意するが、「概念化した意味を“このように”形式に反映する」ことまでは含意していない、ということである。つまり、異なる概念化によって形成された意味をそれぞれ区別して異なる形式に対応させた導出規則を作っているのも、モデルの中で「概念化を反映する」メカニズムが組み込まれていることは自明である。しかしながら、異なる概念化によって形成された意味を「どのような仕方・手段・方法」によって形式に対応させているかということには、なんらの仮定も組み込まれていない。

3.2.2 概念化搭載繰り返し学習モデルにおける推論

意味とは客観的外部世界の構成要素を言語使用者が概念化することによって形成される概念内容が含まれる．発話される言語形式には言語使用者の概念化した意味が対応する．このように意味が言語使用者内部にあると考えれば，外部世界の共有だけでは発話される言語形式に対応する意味内容を知ることができず，意図共有 (Tomasello, 2005) に基づいた推論が必要となる．特に，構成的な文法がなかったと想定される言語創発最初期において，意味と形式の対応関係が構造化されていないことは，言語使用者の概念化を推論することをよりいっそう難しくさせるはずだ．

そこで，ナイーブな想定として，意味表現に含まれる概念化の値を確率的に推論できることとする．つまり，子は親がどのような概念化をしたのかを何らかの方法で理解できるということである．これにより，学習において概念化の値 (Conceptualization Value; CV) を区別した学習ができるようになる．しかし，推論の成功確率が低ければそれだけ2通りの概念化を誤って推論してしまい，親の発話とは異なる学習材料を作ってしまうことになる．

例えば，親は発話条件 C^* として $kick(john, mary)/1$ が与えられ，はつわかていとしての導出規則 $S/kick(john, mary)/1 \rightarrow nmnqwjk$ を生成する．但し，親の純粋な発話は導出規則の終端記号である形式 “nmnqwjk” のみである．親と子は外部状況を共有しているので，条件 C^* は共有できるため，子供は不完全な導出規則として $S/kick(john, mary)/? \rightarrow nmnqwjk$ を受け取るとする．子供が行うべき推論内容は，親がどのような概念化を行なって発話したかということである．すなわち，CV の値を推論しなくてはいけない．確率的に正しい CV を推論することができると仮定し，最終的に $S/kick(john, mary)/1 \rightarrow nmnqwjk$ と推論し，学習材料として使う．

子供の概念化値 CV_{child} は，親の概念化値 CV_{parent} を，推論成功確率 P_{cv} に従って，正しく推論できる．ここで， $CV \in \{0, 1\}$ である．

$$CV_{child} = \begin{cases} CV_{parent}, & \text{with probability } P_{cv}, \\ 1 - CV_{parent}, & \text{with probability } 1 - P_{cv}. \end{cases}$$

3.3 シミュレーションの設計

3.3.1 概念化を含めた意味空間

意味空間とは，概念化を含む確定節文法の導出規則における左辺の条件の集合 C^* である．各条件 $C \in C^*$ は一階述語論理を用いた述語項構造と概念化値 CV は

$$C = \text{predicate}_i(\text{argument}_j, \text{arguments}_{k \neq j})/CV$$

によって表される．

- predicate_i は述語の集合 I から取られる要素（述語）であり， $i \in I$.
- $\text{argument}_j, \text{arguments}_k$ は引数の集合 J から取られる要素であり， $j, k \in J$ かつ $j < k$.
- 再帰表現（e.g., John loves John）を除外するために，引数の組み合わせとして $j \neq k$ かつ $j < k$ のみを考える．
- $CV \in 0, 1$

この時，条件の個数，つまり意味空間の大きさは以下で定義する：

$$C^* = |I| \times |J| \times (|J| - 1) \times |CV|$$

また，改めて注意を促すが，3.2 節の「言語知識の表現：確定節文法」で言及したように， predicate はあくまで「2つの事物を特定の方法で関係づける」意味要素であるため，意味要素のうち「関係性要素」と呼ぶことにする．また2つの項は何かしらの事物であるため，意味要素のうち「対象要素」と呼ぶ．それぞれの集合を以下で定義する：

関係性要素の集合 $P = \{\text{eat}, \text{follow}, \text{get}, \text{help}\}$

対象要素の集合 $X = \{\text{alice}, \text{bob}, \text{carol}, \text{david}\}$

ここで， P は関係性要素の集合であり， X は対象要素の集合である．また，実際に英単語を使用しているが，その意味を扱うわけではなく，あくまでも人間が読み取りやすくした便宜上のラベルの集合である．以下の表 3.1 に関係性要素の集合 P と対象要素の集合 X および，概念化値 CV を整理する．

表 3.1: 概念化搭載繰り返し学習モデルにおける意味空間の設定

i	関係性要素 P_i	j	対象要素 X_j	概念化値 CV
1	eat	1	alice	0 1
2	follow	2	bob	
3	get	3	carol	
4	help	4	david	

以上より，本モデルの意味空間の大きさは以下と定義される．

$$96 = 4 \times 4 \times (4 - 1) \times 2$$

3.3.2 シミュレーションプロセス

世代構成員は、各世代1人であると見做す。厳密には親エージェントと子エージェントは同時に存在するが、親世代には1人の親エージェントが、子世代には1人の子エージェントがいる。親世代の親エージェントが次世代と見做せる子世代の子エージェントに言語知識を継承していく。これは、同世代内のコミュニケーション（水平伝播）はなく、世代間継承（垂直伝播）しかないことを意味する。

継承過程に関しては、最初の世代には親エージェントが存在する。最初の世代の親は何の言語知識も持たないため、言語知識に基づかずに発話するものとする。その発話を受け取る子エージェントは親の発話を材料に学習し、学習が終わったところで親となり、自身の子エージェントに発話を行う。なお、発話を完了した親は世代交代を表現するため亡くなる。

加えて、継承過程には「ボトルネック」と呼ばれる現象がある。これはビンの首が細く狭くなっているところを使って比喻表現であるが、文化進化研究の文脈では、文化の総量が継承により減った状態で継承される現象である。本シミュレーションに適用するならば、親の発話した意味空間よりも、子が受け取る意味空間のほうが小さいということである。これの直感的な説明は、親が頭の中にある言語知識を全て子供に向けて発話することはないという「機会の貧困」であったり、親が仮に言語知識を全て子に発話したとしても子がその全てを忘れずに覚えて学習の材料とすることはないという「記憶の限界」である。本モデルでは親の言語知識を評価する必要があるという研究上の都合から、実装としては、親に意味空間全てを発話してもらい、子にはその一部を継承するという設定にしている。

以下に、概念化搭載繰り返し学習モデルのシミュレーションの疑似コードを示す。

表 3.2: 疑似コード：概念化搭載繰り返し学習モデルのシミュレーション

Algorithm 1 概念化搭載繰り返し学習モデルのシミュレーション

```

1:  $Production_1 \leftarrow \text{FIRSTPRODUCE}(C^* = \text{AllSemanticSpace})$ 
2: for  $i=2 \in \text{Generations}$  do
3:    $Samples_i \leftarrow \text{BOTTLENECK}(Production_{i-1})$ 
4:    $LearningSamples_i \leftarrow \text{INFERENCECV}(Samples_i)$ 
5:    $LearnedRules_i \leftarrow \text{LEARNBY3ALGORITHMS}(LearningSamples_i)$ 
6:    $Production_i \leftarrow \text{PRODUCE}(LearnedRules_i, C^*)$ 
7: end for

```

表 4.1 の疑似コードに示したように、概念化搭載繰り返し学習モデルのシミュレーションにおける主な計算過程は以下の5つである：

1. FIRSTPRODUCE：最初の導出規則を意味空間全体に対して生成する。

2. BOTTLENECK：親の生成した導出規則を一部削除する。
3. INFERENCECV：導出規則にある CV を推論する。
4. LEARNBY3ALGORITHMS：汎化学習を行い，言語知識を構築する。
5. PRODUCE：言語知識に基づき意味空間全体に対して導出規則を生成する。

3.4 評価指標の改変と作成

2.4 節において，創発言語の性質を測る評価指標，特に構成性の評価指標である TopSim について説明し，本モデルで扱う概念化を含めた言語を考慮に入れた計算ができないという点を指摘した．つまり，同一状況を異なる概念化によって「どのように」形式に反映させているのかを測ることは TopSim ではできない，という指摘である．この指摘について，「意味の距離の計算」・「形式の距離の計算」・「概念化間の対応関係の計算」の3つに絞って，詳しく説明していく．

従来評価指標の3つの問題

まず，Hamming 距離で測る「意味の距離の計算」において，従来の評価指標では，概念化だけが異なるルールのパアと意味要素 ($predicate_i, argument_j, argument_k \neq j$) のどれか1つだけが異なるルールのパアを区別することができない．例えば，以下の3つの文ルールを考える．

Rule 1: $S/kick(john, mary)/\underline{0} \rightarrow abcjkde$

Rule 2: $S/kick(john, mary)/\underline{1} \rightarrow dejkabc$

Rule 3: $S/kick(anya, mary)/\underline{0} \rightarrow mnjkde$

Rule 1 と Rule 2 のパアを比較した時，意味の距離である Hamming 距離は，0 と 1 のみ異なるので，1 である．Rule 1 と Rule 3 のパアを比較した時，Hamming 距離は *john* と *anya* のみ異なるので，1 である．これはつまり，意味の要素の違いと概念化の違いを区別しないことにあたる．

次に，「形式の距離の計算」において，従来の評価指標で使われている Levenshtein 距離の計算は，2.4 節で定義したように，1 文字ずつ順序つきで計算するので，形式長が意味要素間で異なる場合や，順序が異なる場合を考慮できない．特に，順序が異なる場合は「異なる概念化間の対応関係の計算」に深く関わるため次に回す．例えば，以下の3つの文ルールを考える．

Rule 1: $S/\underline{k}ick(john, mary)/0 \rightarrow a\underline{j}kdew$

Rule 2: $S/\underline{h}elp(john, mary)/0 \rightarrow a\underline{m}nlgdew$

Rule 3: $S/\underline{l}ike(john, mary)/0 \rightarrow a\underline{t}udew$

Rule 1 と Rule 2 のペアを比較した時、形式の距離は 4 である。Rule 1 と Rule 3 のペアを比較した時、形式の距離は 2 である。Rule 1, 2, 3 の終端記号の形式列はどれも、左端に a と右端に dew があり、これらに挟まれた形式列 (jk, mnl, g, tu) が異なっているだけである。にも関わらず形式の距離が異なるのは、形式長によるスケールを制御できていないのである。この問題は、終端記号全体の形式長で正規化 (Normalize) しても解決されない。

最後に、「異なる概念化間の対応関係の計算」において、従来の評価指標では、そもそも概念化を区別していないわけだが、それよりも深刻な問題がある。それは異なる概念化を“異なる順序”で形式に反映するような言語を考慮に入れられていない、ということである。例えば、以下の 3 つの文ルールを考える。

Rule 1: $S/kick(john, mary)/\underline{0} \rightarrow abc$

Rule 2: $S/kick(john, mary)/\underline{1} \rightarrow cda$

Rule 3: $S/kick(john, mary)/\underline{1} \rightarrow adc$

ここでは、3 つのルールは意味に含まれる CV のみが異なる。Rule 1 と Rule 2 のペアを比較した時、形式の距離は 3 である。Rule 1 と Rule 3 のペアを比較した時、形式の距離は 1 である。Rule 2 と Rule 3 の形式は同じ要素 a, d, c で構成されている。Rule 1 と比較すれば、同一状況に対して、異なる CV 間を「異なる順序」で形式に反映しているのか、「類似な順序」で形式に反映しているのかの違いに見える。自然言語でいうところの $vertical(a, b)$ という同一状況を、異なる概念化間「異なる順序」で “A on B” と “B below A” というように形式に反映している場合が当てはまる。上のように CV の区別を順序によって形式に反映させる言語知識がシミュレーションで創発した場合、従来 TopSim のままでは、それを捉えられない。つまり、形式表現が異なる概念化間でどのような対応関係にあるのか、を計算できない。そうであるならば、いっそのこと、2 つの CV を区別して計算する方がよいことになってしまう。ただし、これが何を意味するかというと、概念化の値 0 と 1 を統合した知識が 1 エージェントの言語知識であるにもかかわらず、2 つの言語知識へと分離してしまっているため、真に 1 個体の言語知識を評価することができない。

評価指標の改善

もう少し分かりやすくするために、文ルールを増やした言語知識の具体例を用いて説明し、問題となる場合（「意味の距離の計算」、「形式の距離の計算」、「概念化間の対応関係の計算」）について説明する。そして、それらを包括的に改善する評価指標を説明する。なお、開始記号 S と非終端記号 N^* は省略する。

まず、従来の評価指標で概念化のない言語知識を観測する。次に、従来の評価指標で概念化を含む言語知識を観測する。その後、改めて、改善すべき箇所に言及する。

表 3.3: 概念化のない言語知識

言語知識		語彙目録
1	kick(john, mary) → abcdef	john → ab
2	kick(mary, john) → efcdab	mary → ef
3	love(john, mary) → abloef	
4	love(mary, john) → efloab	kick → cd
5	help(john, mary) → abhief	love → lo
6	help(mary, john) → efhiab	help → hi

表 3.3 の言語知識を TopSim で計算した結果が以下だ.

- Pearson 積率相関係数の TopSim:1.0
- Spearman 順位相関係数の TopSim:0.9999...
- Normalized Spearman TopSim:0.9999...

結果が示すことは, どの計算手法においても表 3.3 の言語知識の構成性を高く評価しているということである. 人間の目からも, 語彙目録において意味と形式が一対一対応していることがわかり, 言語知識の終端記号の形式も, < 第 1 引数→述部→第 2 引数 > という一貫した組み合わせ順序によって表されている. つまり, 人間の直感や構成性の定義に沿う計算ができていることになる.

それでは, 概念化を含む言語知識のうち, 概念化間で「異なる順序」が使われている言語知識を見てみる.

表 3.4: 概念化間で「異なる順序」の言語知識

言語知識			
	概念化の値=0 の言語知識	概念化の値=1 の言語知識	語彙目録
1	kick(john, mary) → abc <u>def</u>	kick(john,mary)/1 → eft <u>u</u> ab	john → ab
2	kick(mary, john) → efcd <u>ab</u>	kick(mary,john)/1 → abt <u>e</u> ef	mary → ef
3	love(john, mary) → ablo <u>ef</u>	love(john,mary)/1 → ef <u>m</u> nab	kick → cd, kick → tu
4	love(mary, john) → eflo <u>ab</u>	love(mary,john)/1 → abm <u>n</u> ef	love → lo, love → mn
5	help(john, mary) → abh <u>ief</u>	hit(john,mary)/1 → efr <u>g</u> ab	help → hi, help → rg
6	help(mary, john) → efhi <u>ab</u>	hit(mary,john)/1 → abr <u>g</u> ef	

CV の区別をせずに, 3.4 の言語知識の構成性を評価する.

- Pearson 積率相関係数の TopSim:0.1785...
- Spearman 順位相関係数の TopSim:-0.1554...

- Normalized Spearman TopSim:-0.1554...

結果を見ると、どの評価指標も 3.4 の言語知識の構成性を低く評価していることがわかる。語彙目録を見れば、項に入る単語の意味と形式は一対一対応していて、述部に入る単語の意味と形式は一対多の関係である。但し、述部は CV の値に応じて使い分けられていることがわかる。そして言語知識の終端記号の形式も、CV=0 の場合には < 第 1 引数→述部→第 2 引数 > という一貫した組み合わせ順序であり、CV=1 の場合には < 第 2 引数→述部→第 1 引数 > という一貫した組み合わせ順序で表されている。つまり、CV それぞれを独立に計算したら構成性を高く評価されるべき言語知識であるにも関わらず、CV の区別をせずに従来の計算方法を適用すると、構成性は低く評価されてしまう。前述した通り、CV をそれぞれ独立に計算することは 1 個体の言語知識全体を捉えることができていないため妥当ではない。以上より、従来の評価指標は概念化を含む言語知識の観測に適していないことがわかった。「概念化間の対応関係の計算」についても続けてここで指摘する。それは、同一状況で異なる概念化をしている場合、その概念化間の形式の対応関係がどのようになっているのかも捉えられていない。つまり、「類似な順序」なのか「異なる順序」なのかが考慮できていない。そこで、CV のみが異なる導出規則の比較を計算過程に含めることとする。

また「形式の距離の計算」に関して、3.4 で説明したように、“abjkde” と “abstude” の、Levenshtein 距離は標準化しようがしまいが、“jk” と “stu” の部分をまとめて計算するということができない。そこで、予め形式を「分節化」してから、形式の距離を計算する方法を作成した。これにより、1 文字ずつ順序付きで “a, b, j, k, d, e” と “a, b, s, t, u, d, e” の距離を計算するのではなく、“ab, jk, de” と “ab, stu, de” というように分節をもった文字列として距離を計算できる (実際の計算結果としては、前者が 3 で、後者が 1 となる)。いわば、「単語らしき³」分節単位で形式を計算することができる。これにより、“abjkde” と “abstude” の Levenshtein 距離と “abnde” と “abmde” の Levenshtein 距離が、等しく 1 として計算することができる。

また、概念化の違いを形式要素の順序の違いによって表現する言語知識が創発する可能性がある。例えば、自然言語における “A on B” と “B below A” のように、 $vertical(a, b)$ という同一状況に対する概念化の違いを形式要素 (A, B) の順序の違いによって表現する場合である。すなわち、意味の距離は「近い」のに形式の距離が「遠い」と判断されてしまうような言語知識を評価するため、形式の距離を、「1 文字ずつ順序付き」でもなく、「分節単位の順序付き」でもなく、「分節単位のまとまり (Bag)」として計算する手法も必要であると考える。

以下のように、形式の距離計算に関する 3 つの手法を定義する。

³ここで「らしき」という言葉を使っていることから分かるように、「真」の単語ではない。なぜならば、2 つの形式の比較から形式のまとまりを抽出しているだけで意味を参照していないからである。つまり、「なんとなく」単語を抽出しているに過ぎない。

Character-based Editing Distance Character-based editing distance とは、文字列を構成する個々の文字単位で Levenshtein 距離を計算する手法である。文字列を1文字ずつ順序付きで比較し、最小限の編集操作（挿入、削除、置換）によって2つの文字列を一致させる編集距離を求める。

$$d_{\text{char}}(s_j, s_k) = \min\{k \mid k \text{ 回の編集操作で } s_j \text{ を } s_k \text{ に変換できる}\}.$$

Segment-based Sequential Editing Distance Segment-based Sequential Editing distance とは、文字列を分節単位で処理し、各分節を比較して Levenshtein 距離を計算する手法である。分節化された文字列は、事前に適切な分割ルールに従い、「なんとなく単語らしい」まとまりごとに区切られる。距離計算は以下の手順で行われる：

1. 文字列を分節化する：例えば $s_j = \text{"abjkde"}$ を "ab, jk, de" と、 $s_k = \text{"abstude"}$ を "ab, stu, de" と分割する。
2. 各分節ごとに編集距離を計算する：例えば "jk" と "stu" 。
3. 分節ごとの編集距離を合計する： $1 = 0 + 1 + 0$

分節化された形式 $S_j = \{s_{j1}, s_{j2}, \dots\}$ に対して、編集距離は以下で定義される：

$$d_{\text{seg}}(S_j, S_k) = \sum_{i=1}^n d_{\text{char}}(s_{ji}, s_{ki}),$$

ここで s_{ji}, s_{ki} はそれぞれ対応する分節である。

Segment-based Bag Editing Distance Segment-based Bag Editing Distance とは、文字列を分節単位で処理し、分節をまとめた集合 (Bag) として扱い、2つの集合を比較して Levenshtein 距離を計算する手法である。距離計算は以下の手順で行われる：

1. 文字列を分節化して集合とする：例えば $s_j = \text{"abjkde"}$ を $\{\text{ab, jk, de}\}$ と、 $s_k = \text{"destuab"}$ を $\{\text{de, stu, ab}\}$ と分割する。
2. 編集距離として分節集合の対称差を計算する： $\{\text{jk}\}$ と $\{\text{stu}\}$ 。
3. 対称差を $1/2$ する。

分節化された形式 $S_j = \{s_{j1}, s_{j2}, \dots\}$ に対して、編集距離は以下で定義される：

$$d_{\text{bag}}(S_j, S_k) = \frac{1}{2}((s_{ji})\Delta(s_{ki}))$$

ここで s_{ji}, s_{ki} は分節化された形式要素の集合である。

このように, character-based editing distance は文字単位で計算を行い, Segment-based Sequential Editing distance は分節単位で計算を行う. なお, 今後の評価指標としては基本的に Segment-based Sequential Editing distance を使い, character-based editing distance は従来の TopSim でのみ使う.

そこで, CV を区別するが同時に CV 間の対応関係も考慮する評価指標を作成した. 図 3.3 にあるように, 従来では, CV0 と CV1 を独立に計算することが妥当であるのに対し, 作成した評価指標は CV0 と CV1 を独立せずに計算する評価指標を作成した. 大きく分類すると以下である.

- 従来の TopSim 計算に加えて概念化間の Gap を橋渡しするような評価指標
- 概念化間の Gap のみを捉える評価指標

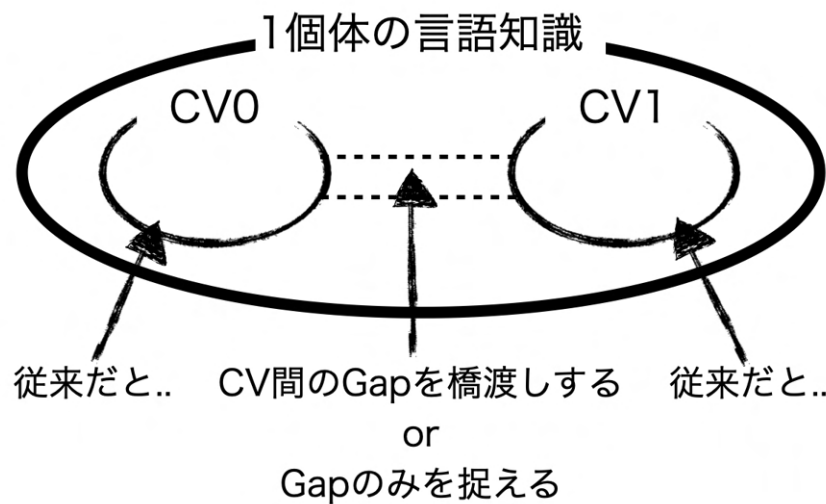


図 3.3: 概念化を区別して計算する

これらにより, CV 別にそれぞれがどのような性質を持っているかを定量的に従来の評価指標を使って分析できるだけでなく, 同一状況を概念化間それぞれでどのように「表現の仕方」を変えているのかという性質を作成する各々の評価指標を使って観察できると考えられる. 計算方法に関しては 2.4 節で説明したことと変わらないか, それよりも単純なため割愛する. 基本的に相関係数を計算するわけなので, 問題は 2 変数 (意味の距離と形式の距離) をどのように所得するかということである.

1. **Segmented TopSim**: 言語知識を概念化の値で区別して, Segment-based Sequential Editing distance を使った標準化 TopSim

この評価指標は, 従来の評価指標の計算過程のうち, 「形式の距離」にのみ修正を施したものである.

表 3.5: 疑似コード：Segmented TopSim アルゴリズム

Algorithm 2 Segmented TopSim

Input: 言語知識 $R = \{ \text{導出規則 } P_i \}_{i=1}^N$

Output: 相関係数 ρ_0, ρ_1

```

1:  $R_0, R_1 \leftarrow \text{SEPARATEBYCV}(R)$ 
2: for  $cv \in \{0, 1\}$  do
3:    $\{m_i^{cv}, s_i^{cv}\}_{i=1}^{|R_{cv}|} \leftarrow \text{SEMFORMPARING}(R_{cv})$ 
4:    $(D_H, D_{seq})^{R_{cv}} \leftarrow \text{DISTPAIRING}((d_H(m_j^{cv}, m_k^{cv}), d_{seg}(s_j^{cv}, s_k^{cv})))$ 
5:    $\rho_{cv} \leftarrow \text{CORRELATION}((D_H, D_{seq})^{R_{cv}})$ 
6: end for
7: return  $\rho_0, \rho_1$ 

```

表 3.5 の疑似コードで示したように，Segmented TopSim の計算過程は以下の 4 つである：

- (a) SEPARATEBYCV: 言語知識を CV 別の集合に分ける.
 - (b) SEMFORMPARING: 導出規則を意味と形式のペアに変換する.
 - (c) DISTPAIRING: 意味と形式のペア同士の距離の組の集合を作る.
 - (d) CORRELATION: 言語知識を CV 別にして，相関係数を計算する.
2. **Segmented Sequential TopSim**: 標準化 TopSim を計算する. ただし，概念化の値のみが異なるルールペアの意味と形式の距離も含める.

概念化の値を区別して個別に 2 変数を取ることに加えて，同一状況を異なる概念化によって表現するペアの意味と形式の距離も 2 変数に取る⁴. すぐ下の評価指標と差別化をするが，形式は分節化しておくが，index 単位で順序付きで計算する. 例えば，*abcjkde* と *destuabc* を分節化して文字を置換させると *abc* と *cda* となる. この形式の距離は 3 (標準化すれば 1) である.

⁴CV のみが異なるルールの意味の距離は，通常の意味要素の距離の半分 (0.5) とした. これに妥当な理由はない. しかし，意味要素が 1 つ異なる，つまり，状況の登場事物や行為が丸ごと異なるより，同じだが捉え方だけが異なる方が，意味の異なり度合い (意味の距離) は小さいだろうという直感に従った.

表 3.6: 疑似コード：Segmented Sequential TopSim アルゴリズム

Algorithm 3 Segmented Sequential TopSim

Input: 言語知識 $R = \{\text{導出規則 } P_i\}_{i=1}^N$

Output: 相関係数 ρ

```

1:  $R_0, R_1 \leftarrow \text{SEPARATEBYCV}(R)$ 
2: for  $cv \in \{0, 1\}$  do
3:    $\{m_i^{cv}, s_i^{cv}\}_{i=1}^{|R_{cv}|} \leftarrow \text{SEMFORMPARING}(R_{cv})$ 
4:    $(D_H, D_{seq})^{R_{cv}} \leftarrow \text{DISTPAIRING}((d_H(m_j^{cv}, m_k^{cv}), d_{seg}(s_j^{cv}, s_k^{cv})))$ 
5: end for
6:  $\text{DiffCV} \leftarrow \text{DIFFCVPAIR}(R_0, R_1)$ 
7:  $\{m_i^{\text{diffCV}}, s_i^{\text{diffCV}}\}_{i=1}^{\text{diffCV}} \leftarrow \text{SEMFORMPARING}(\text{DiffCV})$ 
8:  $(D_H, D_{seq})^{\text{diffCV}} \leftarrow \text{DISTPAIRING}((d_H(m_j^{\text{diffCV}}, m_k^{\text{diffCV}}), d_{seg}(s_j^{\text{diffCV}}, s_k^{\text{diffCV}})))$ 
9:  $(D_H, D_{seq}) \leftarrow \text{CONCAT}((D_H, D_{seq})^{R_{cv=0}}, (D_H, D_{seq})^{R_{cv=1}}, (D_H, D_{seq})^{\text{DiffCV}})$ 
10:  $\rho \leftarrow \text{CORRELATION}((D_H, D_{seq}))$ 
11: return  $\rho$ 

```

表 3.6 の疑似コードで示したように，Segmented Sequential TopSim の計算過程は以下の 8 つである．

- (a) SEPARATEBYCV: 言語知識を CV 別の集合に分ける．
 - (b) SEMFORMPARING: 導出規則を意味と形式のペアに変換する．
 - (c) DISTPAIRING: 意味と形式のペア同士の距離の組の集合を作る．※形式の距離は d_{seg} で計算．
 - (d) DIFFCVPAIR: CV のみ異なる導出規則のペアを作る．
 - (e) SEMFORMPARING: 導出規則を意味と形式のペアに変換する．
 - (f) DISTPAIRING: 意味と形式のペア同士の距離の組の集合を作る．
 - (g) CONCAT: 距離の組の集合を結合する．
 - (h) CORRELATION: 言語知識全体の相関係数を計算する．
3. **Segmented Bag TopSim**: 標準化 TopSim を計算する．ただし，概念化の値のみが異なるルールのペアの意味と形式の距離も含める．

概念化の値を区別して個別に 2 変数を取ることに加えて，同一状況を異なる概念化によって表現するペアの意味と形式の距離も 2 変数に取る．すぐ上の評価指標と対比すると，形式の分節化はするが，index 単位ではなく，bag と

して距離を計算する．例えば、 $abcjkde$ と $destuabc$ を分節化して文字を置換させると abc と cda となる．この形式の距離は、b と d だけが異なるので距離は 1(標準化すれば 1/3) である．

表 3.7: 疑似コード：Segmented Bag TopSim アルゴリズム

Algorithm 4 Segmented Bag TopSim

Input: 言語知識 $R = \{ \text{導出規則 } P_i \}_{i=1}^N$

Output: 相関係数 ρ

```

1:  $R_0, R_1 \leftarrow \text{SEPARATEBYCV}(R)$ 
2: for  $cv \in \{0, 1\}$  do
3:    $\{m_i^{cv}, s_i^{cv}\}_{i=1}^{|R_{cv}|} \leftarrow \text{SEMFORMPARING}(R_{cv})$ 
4:    $(D_H, D_{bag})^{R_{cv}} \leftarrow \text{DISTPAIRING}((d_H(m_j^{cv}, m_k^{cv}), d_{bag}(s_j^{cv}, s_k^{cv})))$ 
5: end for
6:  $\text{DiffCV} \leftarrow \text{DIFFCVPAIR}(R_0, R_1)$ 
7:  $\{m_i^{\text{diffCV}}, s_i^{\text{diffCV}}\}_{i=1}^{\text{diffCV}} \leftarrow \text{SEMFORMPARING}(\text{DiffCV})$ 
8:  $(D_H, D_{bag})^{\text{diffCV}} \leftarrow \text{DISTPAIRING}((d_H(m_j^{\text{diffCV}}, m_k^{\text{diffCV}}), d_{bag}(s_j^{\text{diffCV}}, s_k^{\text{diffCV}})))$ 
9:  $(D_H, D_{bag}) \leftarrow \text{CONCAT}\left((D_H, D_{bag})^{R_{cv=0}}, (D_H, D_{bag})^{R_{cv=1}}, (D_H, D_{bag})^{\text{DiffCV}}\right)$ 
10:  $\rho \leftarrow \text{CORRELATION}((D_H, D_{bag}))$ 
11: return  $\rho$ 

```

表 3.7 の疑似コードで示したように、Segmented Bag TopSim の計算過程は以下の 8 つである．

- (a) SEPARATEBYCV: 言語知識を CV 別の集合に分ける．
 - (b) SEMFORMPARING: 導出規則を意味と形式のペアに変換する．
 - (c) DISTPAIRING: 意味と形式のペア同士の距離の組の集合を作る．※形式の距離は d_{bag} で計算．
 - (d) DIFFCVPAIR: CV のみ異なる導出規則のペアを作る．
 - (e) SEMFORMPARING: 導出規則を意味と形式のペアに変換する．
 - (f) DISTPAIRING: 意味と形式のペア同士の距離の組の集合を作る．
 - (g) CONCAT: 距離の組の集合を結合する．
 - (h) CORRELATION: 言語知識全体の相関係数を計算する．
4. **Separated Sequential Average**: 概念化間のペアの形式の距離の平均である．ただし、index 単位で計算する．これにより上 2 つの評価指標とは異なり、直接的に概念化間の「反映のされ方」を語順を通して見ることができる．

表 3.8: 疑似コード：Separated Sequential Average アルゴリズム

Algorithm 5 Segmented Sequential TopSim

Input: 言語知識 $R = \{\text{導出規則 } P_i\}_{i=1}^N$

Output: 形式距離の平均 AVG_{seq}

- 1: $R_0, R_1 \leftarrow \text{SEPARATEBYCV}(R)$
 - 2: $DiffCV \leftarrow \text{DIFFCVP AIR}(R_0, R_1)$
 - 3: $\{m_i^{diffCV}, s_i^{diffCV}\}_{i=1}^{diffCV} \leftarrow \text{SEMFORMPARING}(DiffCV)$
 - 4: $(D_{seq})^{diffCV} \leftarrow \text{ONLYFORMDIS}(d_{seq}(s_j^{diffcv}, s_k^{diffcv}))$
 - 5: $AVG_{seq} \leftarrow \text{AVERAGEFORMDIS}((D_{seq})^{diffCV})$
 - 6: **return** AVG_{seq}
-

表 3.8 の疑似コードで示したように, Separated Sequential Average の計算過程は以下の 5 つである.

- (a) SEPARATEBYCV: 言語知識を CV 別の集合に分ける.
 - (b) DIFFCVP AIR: CV のみ異なる導出規則のペアを作る.
 - (c) SEMFORMPARING: 導出規則を意味と形式のペアに変換する.
 - (d) ONLYFORMDIS: 形式同士のみの距離の集合を作る. ※形式の距離は d_{seq} で計算.
 - (e) AVERAGEFORMDIS: 平均を計算する.
5. **Separated Bag Average**: 概念化間のペアの形式の距離の平均である. ただし, bag 単位で計算する. これにより上 2 つの評価指標とは異なり, 直接的に概念化間の「反映のされ方」を語順ではなく, 「語彙の共通性」を通して見ることができる.

表 3.9: 疑似コード：Separated Bag Average アルゴリズム

Algorithm 6 Segmented Bag TopSim

Input: 言語知識 $R = \{\text{導出規則 } P_i\}_{i=1}^N$

Output: 形式距離の平均 AVG_{bag}

- 1: $R_0, R_1 \leftarrow \text{SEPARATEBYCV}(R)$
 - 2: $DiffCV \leftarrow \text{DIFFCVP AIR}(R_0, R_1)$
 - 3: $\{m_i^{diffCV}, s_i^{diffCV}\}_{i=1}^{diffCV} \leftarrow \text{SEMFORMPARING}(DiffCV)$
 - 4: $(D_{bag})^{diffCV} \leftarrow \text{ONLYFORMDIS}(d_{bag}(s_j^{diffcv}, s_k^{diffcv}))$
 - 5: $AVG_{bag} \leftarrow \text{AVERAGEFORMDIS}((D_{bag})^{diffCV})$
 - 6: **return** AVG_{bag}
-

表 3.9 の疑似コードで示したように，Separated Bag Average の計算過程は以下の 5 つである．

- (a) SEPARATEBYCV: 言語知識を CV 別の集合に分ける．
- (b) DIFFCVP AIR: CV のみ異なる導出規則のペアを作る．
- (c) SEMFORMPARING: 導出規則を意味と形式のペアに変換する．
- (d) ONLYFORMDIS: 形式同士のための距離の集合を作る．※形式の距離は d_{bag} で計算．
- (e) AVERAGEFORMDIS: 平均を計算する．

それぞれの評価指標の関係を説明する．評価指標 2 と 3，評価指標 4 と 5 は，それぞれパラレルである．構造（語順）が異なる場合でも，形式の距離が小さく計算され，相関係数 R は高く出る．概念化間で語彙が共通するような言語知識を評価指標 2 と 3 で計算した場合，評価指標 2 では index 単位で計算するが故に語順が考慮されることで，語順が異なれば計算結果は低く出るが，評価指標 3 では bag 単位で計算するので語順が異なっていたとしても語順は考慮されず，語彙が共通であるかどうかは考慮されるため，計算結果としては高く出る．ただし各概念化の間で，構造（語順）が類似ならば，index 単位で計算しようが bag 単位で計算しようが計算結果は変わらない．

評価指標 4 と 5 に関しても同様に，語彙が同じで語順が異なるような言語知識を計算した場合，評価指標 4 であれば，index 単位で計算するため，語順が異なることが直接形式の距離に反映されるため，その平均は大きくなる．一方で評価指標 5 であれば，bag 単位で計算するため，語順に関係なく形式の距離を計算し，その平均は小さくなる．また，語彙が同じで語順が類似する言語知識を計算した場

合，評価指標 4 の index 単位であっても，評価指標 5 の bag 単位であっても，形式の距離の平均は小さく出るはずである．

作成した評価指標の適用例

それでは，具体的な言語知識を用意し，上で作成した評価指標がどのように機能するのかについて見ていく．

1. Segmented TopSim：形式を分節化する．但し，概念化の値ごとに計算する．形式を前もって分節化することで，意味要素に対応する形式要素の形式長がバラバラであってもうまく対応できる．表 3.10 の語彙目録を見れば明らかだが，形式長をバラバラにしている．1 文字ずつ計算する場合と比べて，より人間の直感を汲み取ることができる．

表 3.10: 形式長の異なる言語知識の例

言語知識		語彙目録
1	kick(john, mary)/0 → ak b	john → a
2	kick(mary, john)/0 → b k a	mary → b
3	love(john, mary)/0 → al <u>u</u> ht f b	
4	love(mary, john)/0 → bl <u>u</u> ht f a	kick → k
5	help(john, mary)/0 → aow b	love → luhtf
6	help(mary, john)/0 → bow a	help → ow

表 3.10 の言語知識を Pearson 積率相関係数を用いた TopSim で計算すると，0.4564... と出る．Spearman 順位相関係数を用いた TopSim で計算では，0.4121... であり，Normalized Spearman TopSim でも 0.7043... である．

作成した評価指標 1：Segmented TopSim では，ピアソン積率相関係数だと 1.0，スピアマン順位相関係数だと，0.9999... と出る．つまり，人間の直感や構成性の定義に沿う計算ができていることになる．

2：Segmented Sequential TopSim：形式を分節化する．概念化の値を区別せず，言語知識全体の性質を調べることができる．但し，“index” 単位で形式を計算するため，線形順序が形式の距離に反映される．

特に，以下の点で，異なる振る舞いを示す．

- 語彙が共通ならば，概念化間で「類似な順序構造」では「高い」値
- 語彙が共通ならば，概念化間で「異なる順序構造」では「やや低い」値

表 3.11: 「類似な順序」である言語知識の例

「類似な順序」である言語知識			
	概念化の値=0 の言語知識	概念化の値=1 の言語知識	語彙目録
1	kick(john, mary) → abc <u>d</u> ef	kick(john,mary)/1 → abt <u>u</u> ef	john → ab
2	kick(mary, john) → efc <u>d</u> ab	kick(mary,john)/1 → eft <u>u</u> ab	mary → ef
3	love(john, mary) → ablo <u>e</u> f	love(john,mary)/1 → abm <u>n</u> ef	kick → cd, kick → tu
4	love(mary, john) → efl <u>o</u> ab	love(mary,john)/1 → efm <u>n</u> ab	love → lo, love → mn
5	help(john, mary) → abh <u>i</u> ef	help(john,mary)/1 → abr <u>g</u> ef	help → hi, help → rg
6	help(mary, john) → efh <u>i</u> ab	help(mary,john)/1 → efr <u>g</u> ab	

表 3.11 の言語知識を従来のように概念化の区別をせずに「丸ごと」計算すると、Pearson 積率相関係数を用いた TopSim で計算すると、0.8928... と出る。Spearman 順位相関係数を用いた TopSim で計算では、0.9014... であり、Normalized Spearman TopSim でも 0.9014... である。

作成した評価指標 2：Segmented Sequential TopSim では、ピアソン積率相関係数だと 1.0 スピアマン順位相関係数だと 0.9999... と出る。従来の TopSim で測るいずれの数値よりも、相関を高く出力できる。従来 TopSim の問題として、「完全に」構成的な言語を作成したとしてもそれを「完全に」相関ある言語として評価できないことにある。

表 3.12: 「異なる順序」である言語知識

「異なる順序」である言語知識			
	概念化の値=0 の言語知識	概念化の値=1 の言語知識	語彙目録
1	kick(john, mary) → abc <u>d</u> ef	kick(john,mary)/1 → eft <u>u</u> ab	john → ab
2	kick(mary, john) → efc <u>d</u> ab	kick(mary,john)/1 → abt <u>u</u> ef	mary → ef
3	love(john, mary) → ablo <u>e</u> f	love(john,mary)/1 → efm <u>n</u> ab	kick → cd, kick → tu
4	love(mary, john) → efl <u>o</u> ab	love(mary,john)/1 → abm <u>n</u> ef	love → lo, love → mn
5	help(john, mary) → abh <u>i</u> ef	hit(john,mary)/1 → efr <u>g</u> ab	help → hi, help → rg
6	help(mary, john) → efh <u>i</u> ab	hit(mary,john)/1 → abr <u>g</u> ef	

表 3.12 の言語知識を従来のように概念化の区別をせずに「丸ごと」計算すると、Pearson 積率相関係数を用いた TopSim で計算すると、0.1785... と出る。Spearman 順位相関係数を用いた TopSim で計算では、-0.1554... であり、Normalized Spearman TopSim でも -0.1554... である。

作成した評価指標 2：Segmented Sequential TopSim では、ピアソン積率相関係数だと 0.9184...、スピアマン順位相関係数だと 0.8176... と出る。正確に「完全に相関する」ということは言えないが、少なくとも、従来の TopSim で測るいずれの数値よりも、相関を高く出力できる。

3: Segmented Bag TopSim: 形式を分節化する. 概念化の値を区別せず, 言語知識全体の性質を調べることができる. 但し, “bag” 単位で形式を計算するため, 線形順序が形式の距離に反映されず, 共通性を取り出される.

特に, 以下の点で, 異なる振る舞いを示す.

- 語彙が共通ならば, 概念化間で「類似な順序構造」では「高い」値
- 語彙が共通ならば, 概念化間で「異なる順序構造」でも「高い」値

表 3.13: 「類似な順序」である言語知識

「類似な順序」である言語知識			
	概念化の値=0 の言語知識	概念化の値=1 の言語知識	語彙目録
1	kick(john, mary) → abc <u>d</u> ef	kick(john,mary)/1 → abt <u>u</u> ef	john → ab
2	kick(mary, john) → efc <u>d</u> ab	kick(mary,john)/1 → eft <u>u</u> ab	mary → ef
3	love(john, mary) → ablo <u>e</u> f	love(john,mary)/1 → abm <u>n</u> ef	kick → cd, kick → tu
4	love(mary, john) → efl <u>o</u> ab	love(mary,john)/1 → efm <u>n</u> ab	love → lo, love → mn
5	help(john, mary) → abh <u>i</u> ef	hit(john,mary)/1 → abr <u>g</u> ef	help → hi, help → rg
6	help(mary, john) → efh <u>i</u> ab	hit(mary,john)/1 → efr <u>g</u> ab	

表 3.13 の言語知識を従来のように概念化の区別をせずに「丸ごと」計算すると, Pearson 積率相関係数を用いた TopSim で計算すると, 0.8928... と出る. Spearman 順位相関係数を用いた TopSim で計算では, 0.9014... であり, Normalized Spearman TopSim でも 0.9014... である.

作成した評価指標 3: Segmented Bag TopSim では, ピアソン積率相関係数だと 1.0, スピアマン順位相関係数だと 0.9999... と出る. 正確に「完全に相関する」ということは言えないが, 少なくとも, 従来の TopSim で測るいずれの数値よりも, 相関を高く出力できる.

表 3.14: 「異なる順序」である言語知識

「異なる順序」である言語知識			
	概念化の値=0 の言語知識	概念化の値=1 の言語知識	語彙目録
1	kick(john, mary) → abc <u>d</u> ef	kick(john,mary)/1 → eft <u>u</u> ab	john → ab
2	kick(mary, john) → efc <u>d</u> ab	kick(mary,john)/1 → abt <u>u</u> ef	mary → ef
3	love(john, mary) → ablo <u>e</u> f	love(john,mary)/1 → efm <u>n</u> ab	kick → cd, kick → tu
4	love(mary, john) → efl <u>o</u> ab	love(mary,john)/1 → abm <u>n</u> ef	love → lo, love → mn
5	help(john, mary) → abh <u>i</u> ef	hit(john,mary)/1 → efr <u>g</u> ab	help → hi, help → rg
6	help(mary, john) → efh <u>i</u> ab	hit(mary,john)/1 → abr <u>g</u> ef	

表 3.14 の言語知識を従来のように概念化の区別をせずに「丸ごと」計算すると、Pearson 積率相関係数を用いた TopSim で計算すると、0.1785... と出る。Spearman 順位相関係数を用いた TopSim で計算では、-0.1554... であり、Normalized Spearman TopSim でも-0.1554... である。

作成した評価指標 3: Segmented Bag TopSim では、ピアソン積率相関係数だと 0.9804..., スピアマン順位相関係数だと 0.9799... と出る。正確に「完全に相関する」ということは言えないが、少なくとも、従来の TopSim で測るいずれの数値よりも、相関を高く出力できる。

4: Separated Sequential Average: 形式を分節化する。言語知識全体ではなく、概念化の値のみが異なる発話のペア、つまり、同一状況を異なる形式で表現するペアのみを計算対象とする。しかも計算することはペア間の形式の距離の平均である。但し、index 単位で計算するため、「類似な順序構造」を持つことを前提として、語彙の共通性を見る。

- 概念化間で「類似な順序構造」ならば、語彙が共通だと「低い」値
- 概念化間で「異なる順序構造」ならば、語彙が共通でも「高い」値

表 3.15: 「類似な順序」である言語知識

「類似な順序」である言語知識			
	概念化の値=0 の言語知識	概念化の値=1 の言語知識	語彙目録
1	kick(john, mary) → abc <u>d</u> ef	kick(john,mary)/1 → abt <u>u</u> ef	john → ab
2	kick(mary, john) → efc <u>d</u> ab	kick(mary,john)/1 → eft <u>u</u> ab	mary → ef
3	love(john, mary) → ablo <u>e</u> f	love(john,mary)/1 → abm <u>n</u> ef	kick → cd, kick → tu
4	love(mary, john) → efl <u>o</u> ab	love(mary,john)/1 → efm <u>n</u> ab	love → lo, love → mn
5	help(john, mary) → abh <u>i</u> ef	hit(john,mary)/1 → abr <u>g</u> ef	help → hi, help → rg
6	help(mary, john) → efh <u>i</u> ab	hit(mary,john)/1 → efr <u>g</u> ab	

作成した評価指標 Separated Sequential Average では、語彙が共通な「類似な順序構造」である表 3.15 の言語知識を計算すると、概念化間のペアの形式の距離は平均は 0.1666... と出る。

表 3.16: 「異なる順序」である言語知識

「異なる順序」である言語知識			
	概念化の値=0 の言語知識	概念化の値=1 の言語知識	語彙目録
1	kick(john, mary) → abc <u>d</u> ef	kick(john,mary)/1 → eft <u>u</u> ab	john → ab
2	kick(mary, john) → efc <u>d</u> ab	kick(mary,john)/1 → abt <u>u</u> ef	mary → ef
3	love(john, mary) → ablo <u>e</u> f	love(john,mary)/1 → efm <u>n</u> ab	kick → cd, kick → tu
4	love(mary, john) → efl <u>o</u> ab	love(mary,john)/1 → abm <u>n</u> ef	love → lo, love → mn
5	help(john, mary) → abh <u>i</u> ef	hit(john,mary)/1 → efr <u>g</u> ab	help → hi, help → rg
6	help(mary, john) → efh <u>i</u> ab	hit(mary,john)/1 → abr <u>g</u> ef	

作成した評価指標 Separated Sequential Average では、語彙が共通な「異なる順序構造」である表 3.16 の言語知識を計算すると、概念化間のペアの形式の距離は平均は 0.5 と出る。

5: Separated Bag Average: 形式を分節化する。言語知識全体ではなく、概念化の値のみが異なる発話のペア、つまり、同一状況を異なる形式で表現するペアのみを計算対象とする。しかも計算することはペア間の形式の距離の平均である。但し、bag 単位で研鑽するため、順序構造は反映せず、純粋に語彙の共通性を見る。

- 概念化間で「類似な順序構造」ならば、語彙が共通だと「低い」値
- 概念化間で「異なる順序構造」ならば、語彙が共通でも「高い」値

表 3.17: 「類似な順序」である言語知識

「類似な順序」である言語知識			
	概念化の値=0 の言語知識	概念化の値=1 の言語知識	語彙目録
1	kick(john, mary) → abc <u>d</u> ef	kick(john,mary)/1 → abt <u>u</u> ef	john → ab
2	kick(mary, john) → efc <u>d</u> ab	kick(mary,john)/1 → eft <u>u</u> ab	mary → ef
3	love(john, mary) → ablo <u>e</u> f	love(john,mary)/1 → abm <u>n</u> ef	kick → cd, kick → tu
4	love(mary, john) → efl <u>o</u> ab	love(mary,john)/1 → efm <u>n</u> ab	love → lo, love → mn
5	help(john, mary) → abh <u>i</u> ef	hit(john,mary)/1 → abr <u>g</u> ef	help → hi, help → rg
6	help(mary, john) → efh <u>i</u> ab	hit(mary,john)/1 → efr <u>g</u> ab	

作成した評価指標 Separated Bag Average では、語彙が共通な「類似な順序構造」である表 3.17 の言語知識を計算すると、概念化間のペアの形式の距離は平均は 0.1666... と出る。

表 3.18: 「異なる順序」である言語知識

「異なる順序」である言語知識			
	概念化の値=0 の言語知識	概念化の値=1 の言語知識	語彙目録
1	kick(john, mary) → abc <u>d</u> ef	kick(john,mary)/1 → eft <u>u</u> ab	john → ab
2	kick(mary, john) → efc <u>d</u> ab	kick(mary,john)/1 → abt <u>u</u> ef	mary → ef
3	love(john, mary) → ablo <u>e</u> f	love(john,mary)/1 → efm <u>n</u> ab	kick → cd, kick → tu
4	love(mary, john) → efl <u>o</u> ab	love(mary,john)/1 → abm <u>n</u> ef	love → lo, love → mn
5	help(john, mary) → abh <u>i</u> ef	hit(john,mary)/1 → efr <u>g</u> ab	help → hi, help → rg
6	help(mary, john) → efh <u>i</u> ab	hit(mary,john)/1 → abr <u>g</u> ef	

作成した評価指標 Separated Bag Average では、語彙が共通な「異なる順序構造」である表 3.18 の言語知識を計算すると、概念化間のペアの形式の距離は平均は 0.1701... と出る。

これまでの、作成した評価指標の性能について、視覚的にわかりやすく表示する。なお、「従来の評価指標」としては Normalized Spearman TopSim とする。「作成した評価指標」としては Spearman で計算した値を出す。

表 3.19: 従来の評価指標と作成した評価指標の比較

完全に構成的な言語知識						
評価指標	概念化の区別なし		概念化の区別あり			
	従来	作成	類似な順序		異なる順序	
			従来	作成	従来	作成
1:Seg_TopSim	0.7043...	0.9999...				
2:Seg_Seq_TopSim			0.9014...	0.9999...	-0.1554...	0.8176...
3:Seg_Bag_TopSim			0.9014...	0.9999...	-0.1554...	0.9799...
4:Sep_Seq_Ave				0.1666...		0.5
5:Sep_Bag_Ave				0.1666...		0.1701...

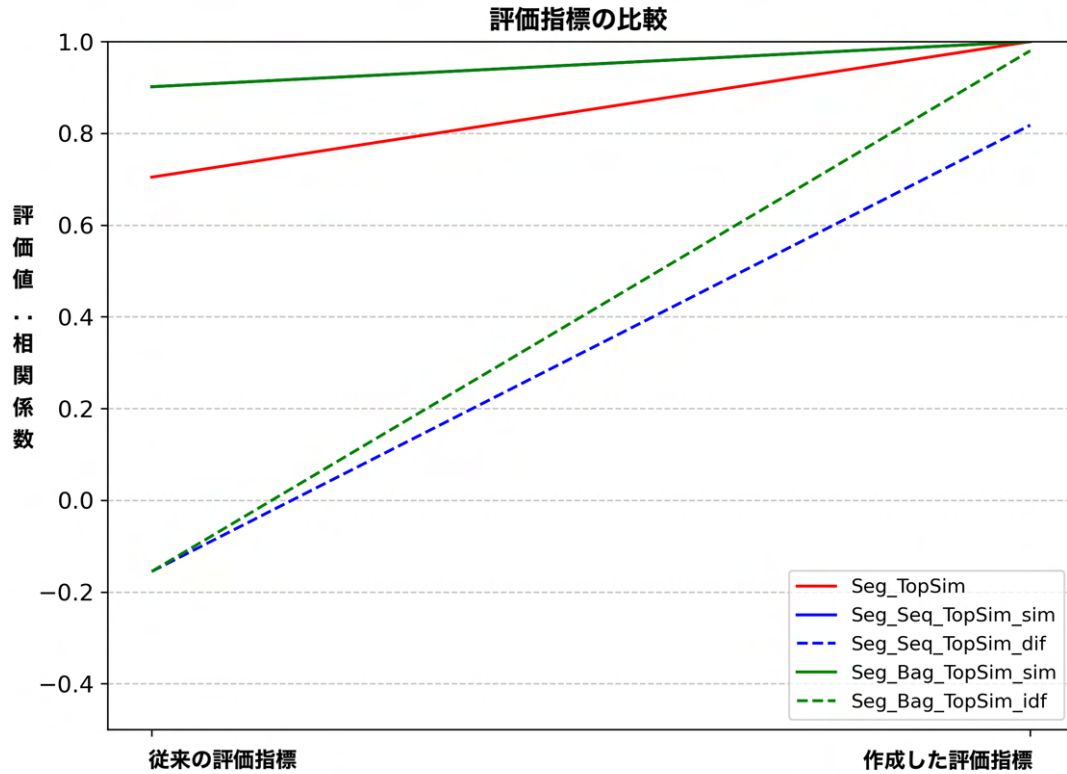


図 3.4: 従来の評価指標と作成した評価指標の比較：実線と点線の区別は、「類似な順序」と「異なる順序」に対応する．なお、青実線と緑実線は重なっている．また、評価指標の 4・5 に関しては相関係数を計算するものではないため、グラフには描画しない．

また、上記の評価指標がまともに機能しなかった場合の打開策として、次の、構成性の間接的な指標を使う．分節化可能性は圧縮できたルールの総数、合成可能性は未知の意味に対応できる表現度を用いることができる．表現度とは、学習後の言語知識に基づいて意味空間全体を、単語ルールや文ルールの創作 (invention) なしにどの程度表現できるかを計算する．これにより、発話アルゴリズムでは invention などして何とか発話していたものを取り除き、純粹の当該世代の言語知識が汎化できているのかを測ることができる．表現度 ϵ は以下で定義される：

$$\text{表現度 } \epsilon = \frac{\text{発話過程で創作のない導出規則の数 } R}{\text{意味空間全体の大きさ } S}$$

留意点として、この表現度は、実際の発話プロセスを観測して計算するものではない．実際の発話の際には新しく創作された単語ルールが言語知識に取り込まれ、それ以降の発話において使用されてしまうため、学習済みの言語知識のみを使った発話ではないためである．

第4章 概念化搭載繰り返し学習モデルのシミュレーション

4.1 パラメータ設定

どのようなパラメータで実験したのかについて改めて整理しておく.

- 世代構成員: 1
- 世代数: 400
- 意味空間の大きさ: $96 = 4 \times 4 \times 3 \times 2$

i	関係性要素 P_i	j	対象要素 X_j	CV
1	eat	1	alice	0
2	follow	2	bob	
3	get	3	carol	1
4	help	4	david	

- 形式長: 3 以上 10 未満
- 発話: 同じ状況に対して, 異なる概念化で 1 つ 1 つ 1 度きりしか発話しない
- ボトルネック: $48 = 96 \div 2$ (各概念化から半分ずつ)
- 推論成功確率: 正しい CV を受け取る確率¹
- 学習: 1 発話ごとに逐次的に学習

既に 3.3.2 小節に貼り付けた表 4.1 の疑似コードを再掲する.

¹推論成功確率が 1.0 未満である場合, 偶然完全に同じ意味表現になる場合がある. 例えば, $S/kick(john, mary)/0 \rightarrow abcjkde$ と $S/kick(john, mary)/1 \rightarrow nmkl$ という発話を親が行って継承された場合に, 子が $S/kick(john, mary)/0 \rightarrow abcjkde$ は正しく推論したが, $S/kick(john, mary)/1 \rightarrow nmkl$ は誤って推論した場合, $S/kick(john, mary)/0 \rightarrow abcjkde$ と $S/kick(john, mary)/0 \rightarrow nmkl$ というように, 全く同じ意味表現を受け取ることがある. しかし, 今回これはあえて省かなかった.

表 4.1: 疑似コード：概念化搭載繰り返し学習モデルのシミュレーション

Algorithm 7 概念化搭載繰り返し学習モデルのシミュレーション

```

1:  $Production_1 \leftarrow \text{FIRSTPRODUCE}(C^* = \text{AllSemanticSpace})$ 
2: for  $i=2 \in \text{Generations}$  do
3:    $Samples_i \leftarrow \text{BOTTLENECK}(Production_{i-1})$ 
4:    $LearningSamples_i \leftarrow \text{INFERENCECV}(Samples_i)$ 
5:    $LearnedRules_i \leftarrow \text{LEARNBY3ALGORITHMS}(LearningSamples_i)$ 
6:    $Production_i \leftarrow \text{PRODUCE}(LearnedRules_i, C^*)$ 
7: end for

```

表 4.1 の疑似コードに示したように，概念化搭載繰り返し学習モデルのシミュレーションにおける主な計算過程は以下の 5 つである：

1. FIRSTPRODUCE：最初の導出規則を意味空間全体に対して生成する。
2. BOTTLENECK：親の生成した導出規則を一部削除する。
3. INFERENCECV：導出規則にある CV を推論する。
4. LEARNBY3ALGORITHMS：汎化学習を行い，言語知識を構築する。
5. PRODUCE：言語知識に基づき意味空間全体に対して導出規則を生成する。

4.2 仮説

2.1.3 小節で述べたように，本研究では，構文交替を「基本的な意味を共有する異なる形式」と定式化し，異なる形式同士は同一事象を異なる捉え方で概念化して符号化された形式であるとした。また本研究では，概念化する対象を二者関係に限定しているので，「異なる捉え方」は 2 つに限定される。事象における行為者性を捉えることで能動文が選ばれ，逆に事象における被行為者性を捉えることで受動文が選ばれるとすれば，一方の捉え方はその捉え方固有な構文を持つ傾向があると考えられる。他方においても同様である。すなわち，同一事象に異なる構文を使う構文交替は異なる概念化によってもたらされる可能性がある。ただし，話し手が独りよがりにより自身の概念化を密かに構文交替を使って反映させていたとしても，聞き手は話し手の内部にある概念化を知るのは難しいように思う，特に，言語創発初期の段階で仮に言語が構成的でも体系的でもなかったならば，より一層難しいはずである。以上を踏まえ，以下の仮説を提示し，それを検証する。

仮説 概念化の値を正しく推論できる確率がある程度高ければ、概念化を異なる語順によって形式に反映する言語が創発する確率が高くなる。一方で、概念化の値を正しく推論できる確率がある程度低ければ、概念化を異なる語順によって形式に反映する言語が創発する確率は低く、また似たような語順、あるいは似たような形式を持つ言語が創発する確率が高くなる。

この仮説を検証するため、推論成功確率を 0.1, 0.3, 0.5, 0.7, 0.9, 0.91, 0.92, 0.93, 0.94, 0.95, 0.96, 0.97, 0.98, 0.99, 1.0 の各段階で設定し、シミュレーションを回した。各々 100seeds ずつ回した。

4.3 実験結果

4.3.1 安定世代の恣意的な定義と安定世代の出現数

安定世代の恣意的な定義 本研究では、「安定世代」を以下と定義する。安定世代とは、表現力が 0.8 以上で、かつ、全ルール数が 24 個未満で、かつ、汎化されていない文ルール数が 1 個未満である状態が 30 世代以上続く世代である。安定世代の典型例を図 4.1 で示す。この図 4.1 において、最初の世代はもちろん学習を何もしていないので、赤で示された総体論 (Holistic) ルール²だけを言語知識としている。しかし、4.2 にあるように学習アルゴリズム (Chunk, Category-Integration, Replace) を適用させていくことで、総体論ルールは数を減らし、徐々にルールが汎化していき、変数の数を増やしていく。およそ 100 世代以降において、汎化されていない文ルール (赤) はなく、単語ルール (青) と一部汎化されたルール (紫) が残っている世代が 200 世代以上連続で続いている。

安定世代を定義する理由は、各 seed, 各条件によって評価指標の値が収束せず、上下に変動する場合があるからだ。それぞれの値に妥当な根拠はなく、あくまで恣意的なものであるが、子の受け取る親からの発話の数は 48 であるため、その半分の数ルールにまで学習で圧縮でき、かつ丸暗記型の総体論的ルールが全くなことは、汎化学習がよくできている状態と見做した。

²ここらへんの用語の説明は 3.2 の「4 つのルール」に書いてある。

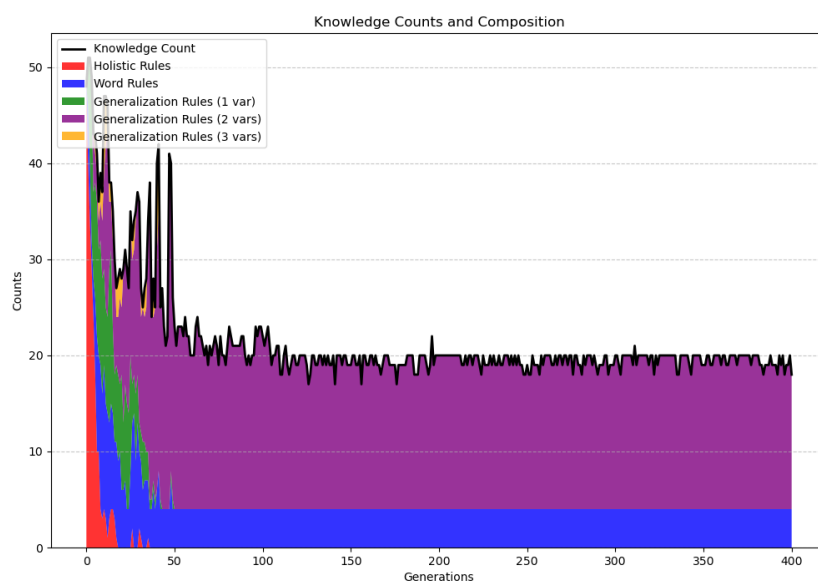


図 4.1: 安定世代をもつ典型例（推論成功確率 0.98）

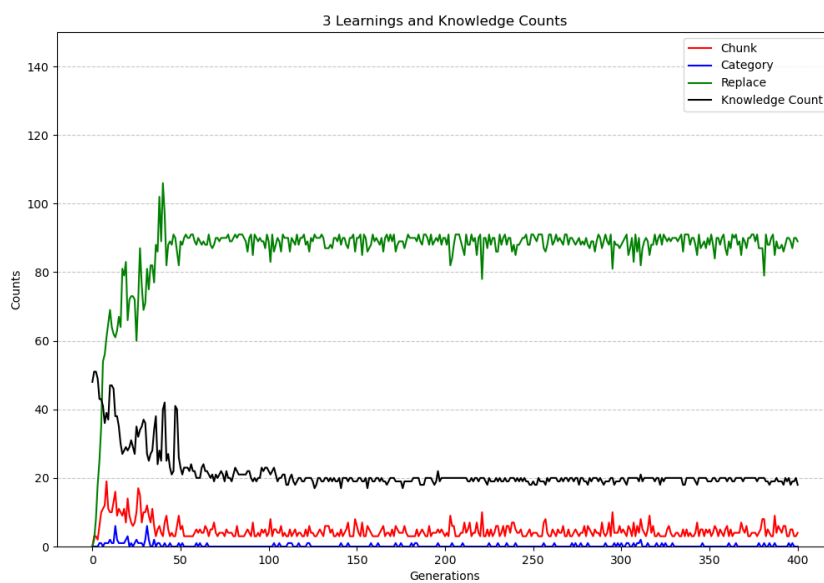


図 4.2: 安定世代をもつ典型例（推論成功確率 0.98）における学習アルゴリズムの適用状況

ただし、安定世代は一度突入すれば必ずその状態が継続するわけではなく、再び安定しなくなる場合 (図 4.3) や、一度急激に安定世代ではなくなるもすぐに安定世代に戻るといった安定世代が一時的に途切れる場合 (図 4.4) も存在する。

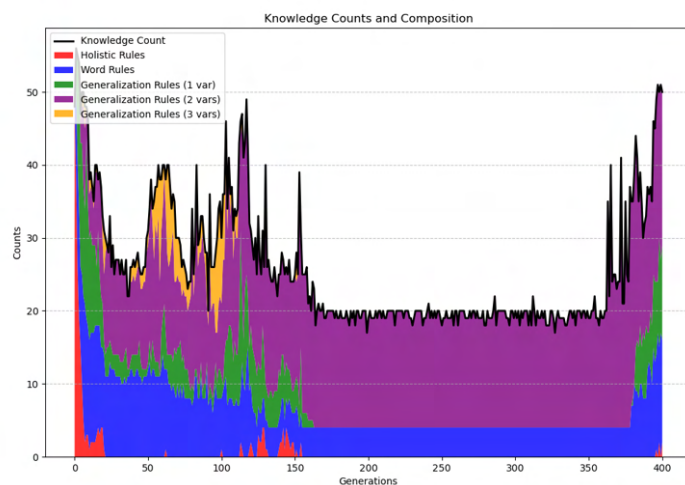


図 4.3: 再び安定しなくなる場合 (推論成功確率 0.97)

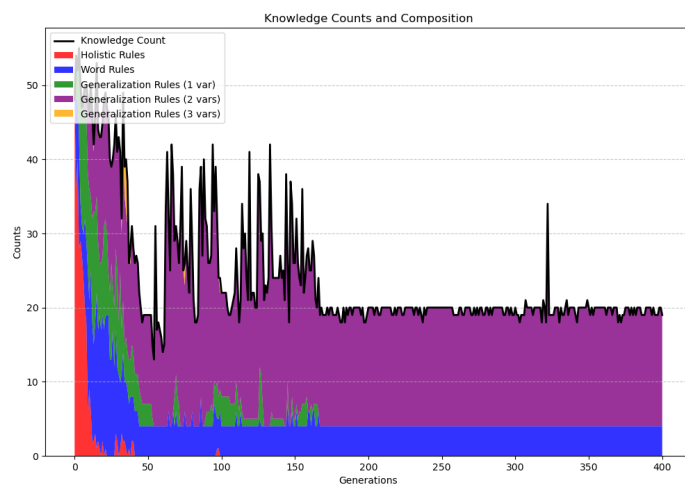


図 4.4: 安定世代が一時的に途切れる例 (推論成功確率 0.98)

安定世代の出現数 推論成功確率が高くなればなるほど、親の発話に含まれる概念化された意味を正しく理解でき、概念化の違いを区別して学習が進められるので安定世代に収束する確率は高く、安定世代になる seed 数も増えるのではないかと考えた。そこで各推論成功確率と安定世代の出現した seed 数の計測した。

表 4.2 に各推論成功確率における安定世代が出現した seed 数をまとめた。推論成功確率と安定世代の数に関して、スピアマン順位相関係数 P_s を計算したが、 $P_s = 0.1194...$ (P 値=0.626..., 有意水準 $\alpha = 0.05$) であり、統計的に有意な相関は確認されなかった。

0.1	0.2	0.3	0.4	0.5	0.6	0.7	0.8	0.9	0.91
.17	.12	.14	.11	.12	.16	.06	.14	.15	.16
0.92	0.93	0.94	0.95	0.96	0.97	0.98	0.99	1.0	
.15	.20	.14	.13	.12	.09	.13	.14	.16	

表 4.2: 推論成功確率と安定世代の出現頻度

付録 A に掲載した言語知識の中身の推移のグラフからも明らかなように、各推論成功確率内においても seed 値ごとにロバストな結果は出なかった。そこで、安定世代を持つ seeds に注目し、それらにおける全評価指標の平均推移を見る。なお、推論成功確率が 0.1 から 0.99 に関してはほとんど違いが見られないため、推論成功確率 0.99 と 1.0 のグラフを提示する（他の推論成功確率のグラフは付録 B に掲載）。ここで、注目すべきことは、「同一状況を異なる概念化」した形式間の比較した指標を表すピンク線 (Separated Sequential Average) とグレー線 (Separated Bag Average) である。ピンク線は index 単位で計算する概念化間における形式距離の平均であるので、「値が小さい」ことは「類似な順序」であることを示し、「値が大きい」ことは「異なる順序」であることを示す。一方、グレー線は bag 単位で計算する概念化間における形式距離の平均であるので、「値が小さい」ことは「必ずしも」「類似な順序」であることを意味しないが、「共通する語彙が多い」ことは示す。また「値が高い」ことは「共通しない語彙」ことを示す。

図 4.5 における縦軸 0.0 から 0.2 の間に描画されるピンク線とグレー線は、最初期の世代では終盤の世代にかけて下降傾向にあるのに対して、図 4.6 では、グレー線のみ緩やかな下降傾向にあるものの図 4.5 のグレー線よりは下がり切っていない。そしてピンク線に関してはほとんど下がらない。以上より、図 4.5 において、ピンク線の「値が低い」とことグレー線の「値が低い」ことは、「同一状況を異なる概念化」した形式間において「共通する語彙」を使って「類似な順序」で表現していることがわかる。一方、図 4.6 において、ピンク線の「値が高い」とことグレー線の「値がやや高い」ことは、「同一状況を異なる概念化」した形式間において「共通しない語彙」を使って「異なる順序」で表現していることがわかる。

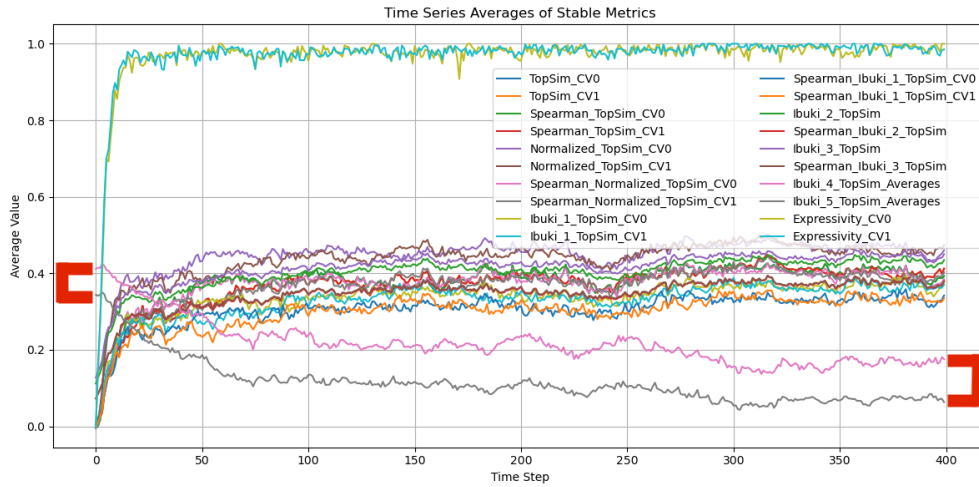


図 4.5: 安定世代をもつ seeds における各評価指標の平均：推論成功確率 0.99

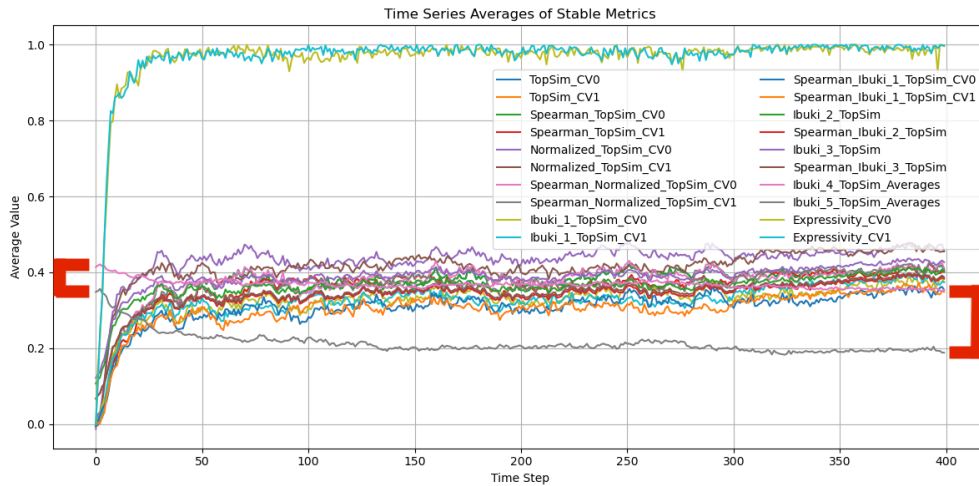


図 4.6: 安定世代をもつ seeds における各評価指標の平均：推論成功確率 1.0

また、図 4.5 と図 4.6 および、付録 B に掲載したその他の推論成功確率における安定世代をもつ seeds の全評価指標の平均を描画したグラフが示す通り、Separated Sequential Average と Separated Bag Average を除く評価指標が非常に弱い相関を示している。つまり、小節で解説したような人工的に作った言語知識の持つ性質を、本シミュレーションにより創発した言語知識は持っていないことがわかる。

安定世代における言語知識の中身

これまで全評価指標の平均から、推論成功確率に応じた言語知識の性質を見てきたが、実際の言語知識の中身を具体的に見ていき、本当に「共通する/しない語彙」や「類似な/異なる順序」といった性質が見られるのか観察していく。その結

果、以下のことが分かった。

- 推論成功確率が 1.0 つまり、推論が完全である実験における安定世代では、異なる語順で概念化を反映していると考えて良い言語が創発している。例えば、推論成功確率=1.0 の安定世代における言語知識を示す。

表 4.3: 異なる CV と異なる右辺

単語ルール	
$K/alice \rightarrow d, K/bob \rightarrow i, K/carol \rightarrow w, K/david \rightarrow k$ $L/eat \rightarrow r, L/fuck \rightarrow m, L/get \rightarrow j, L/hel \rightarrow l$	
CV=0	CV=1
<u>$S/p(x, y)/0 \rightarrow cK/xL/pK/yj$</u>	<u>$S/p(x, y)/1 \rightarrow K/ygK/xL/pwi$</u>
$S/p(x, carol)/0 \rightarrow K/xbL/pn$	$S/p(x, carol)/1 \rightarrow qL/pcnbK/x$
$S/p(carol, y)/0 \rightarrow K/ybL/pn$	
	$S/p(x, bob)/1 \rightarrow K/xazdL/pf$
	$S/p(x, david)/1 \rightarrow L/pgK/xrio$

この表 4.3 において、単語ルールはどれも条件 C^* を区別する終端記号としての形式表現が使われていることから、曖昧さが無い。また、下線を引いた 2 つルールは CV のみが異なるもっとも汎化学習ができたものである。そして、この 2 つのルールを比較すると、「同一状況を異なる概念化」をすることによって「異なる順序」で表現される可能性を持っていることがわかる。なぜならば、CV=0 であれば < 第 1 引数 → 述部 → 第 2 引数 > という順序であるのに対し、CV=1 であれば、< 第 2 引数 → 第 1 引数 → 述部 > という順序となっているためである。しかしながら、単語ルールはどちらのルールでも同じものが使われるため、「共通の語彙」が使われる。そして、「異なる順序」だけでなく、冗長な形式要素を付加することでも「同一状況を異なる概念化」をすることが表現されている。具体的には、CV=0 である $S/p(x, y)/0 \rightarrow \underline{cK/xL/pK/yj}$ において、終端記号の文頭にある c と文末の j で付加されており、CV=1 である $S/p(x, y)/1 \rightarrow K/y\underline{gK/xL/p}\underline{wi}$ においても、途中の g や文末の wi が付加されている。

しかし、全ての安定世代の言語知識が、上の言語知識のように単語ルールが条件ごとに形式表現を区別しない導出規則が頻発した。例えば、以下の言語知識である (単語ルールのみ抜粋)。

表 4.4: 多義な単語ルールの抜粋

単語ルール
$Q/alice \rightarrow f, Q/bob \rightarrow h, Q/carol \rightarrow f, Q/david \rightarrow f$

この表 4.4 において、非終端記号である単語ルールのカテゴリーラベル (今回の場合 Q) が共通しているという意味において、学習自体はうまく汎化できている。しかしながら、異なる条件 $\{alice, carol, david\}$ において同じ終端記号 f が対応づいた単語ルールが存在する場合もある。これは、意味要素の区別を表現しなくなった言語知識と言える。つまり、学習としては汎化できているが、表現としては曖昧性の高くなっている。いわゆる「多義語」や「一般名詞」のような単語ルールである。仮に異なる条件 $\{alice, bob, carol, david\}$ が人物を指示していたとした場合、彼ら $\{alice, carol, david\}$ は区別されずに単なる「誰か (someone)」としか指示されず、 bob のみ特権的に他の誰かとは区別された存在として指示される言語であると言える。しかし、このような言語知識を 1 つ 1 つ見ていくと、項に入る単語ルールが条件を区別できなくなっているからといって、必ずしも概念化自体を区別できなくなっているわけではない。というのも、確かに意味要素を 1 つ 1 つ区別することをしなくなっているものの、冗長な形式要素が付与されていたり、中には異なる順序で表現される可能性のある汎化ルールを持っていた。すなわち、意味要素の区別を形式上でも区別して表現できないとしても、概念化された意味は言語使用者にとっては区別できているのであり、それを形式に反映はできるかできないかの問題である、と解釈することができる。

- 推論成功確率が低い (0.1) 実験から非常に高い (0.99) 実験における安定世代での言語知識は、概念化の違いが異なる形式によって表現されることがなく、いわば概念化の違いは「無視」され、同じ語順・構造・形式をもつ言語が創発していた。以下にその例を示す。

この表 4.5 の言語知識において、述語の項に入る単語ルール $\{T/david \rightarrow n, T/carol \rightarrow d, T/alice \rightarrow h, T/bob \rightarrow m\}$ はどれも異なる条件に異なる終端記号が付与されていることから、曖昧性がないことは明らかである。そして、 $S/fuck(x, y)/0 \rightarrow jkT/xuT/yv$ のみ例外的に $CV=0$ のルール群にしか存在しないが、例えば $S/eat(x, y)$ は CV が異なっても同じ右辺 $T/ybT/xm$ を持っているように、それ以外のルールはどれも CV の区別に関係なく同じ右辺を持っている。つまり、異なる概念化であっても同じ形式が発話されてしまう言語知識となっている。このことから、このような言語では概念化の違いを表現する必要がなくなったと解釈される。

以上の結果の観察をまとめると、推論成功確率が 1.0、つまり完全に親の概念化

表 4.5: 異なる CV と共通の右辺

単語ルール	
$T/david \rightarrow n, T/carol \rightarrow d, T/alice \rightarrow h, T/bob \rightarrow m$	
CV=0	CV=1
$S/eat(x, y)/0 \rightarrow T/ybT/xm$	$S/eat(x, y)/1 \rightarrow T/ybT/xm$
$S/fuck(x, y)/0 \rightarrow ulT/yvT/xu$	$S/fuck(x, y)/1 \rightarrow ulT/yvT/xu$
$S/fuck(x, y)/0 \rightarrow jkT/xuT/yv$	
$S/get(x, y)/0 \rightarrow T/xT/yhucv$	$S/get(x, y)/1 \rightarrow T/xT/yhucv$
$S/get(x, bob)/0 \rightarrow T/xcnjii$	$S/get(x, bob)/1 \rightarrow T/xcnjii$
$S/hel(x, y)/0 \rightarrow bT/xT/yeiat$	$S/hel(x, y)/1 \rightarrow bT/xT/yeiat$

を推論できる場合、安定世代において＜「異なる概念化」を「異なる順序」と「冗長な形式要素の付加」＞によって表現する言語知識が存在しうる．一方で、推論成功確率が 1.0 未満、つまり親の概念化を少しでも推論できないと、＜「異なる概念化」を「異なる順序」と「冗長な形式要素の付加」＞によって表現できる言語知識が存在することが（観察の限りにおいて）あり得ず、むしろ、＜「異なる概念化」を「類似な順序」＞によって表現する言語知識が存在しうる．

3.4 節で作成した評価指標 4・5 において顕著に現れた．なお、評価指標 4・5・表現度 (expressivity) 以外のその他の評価指標に関しては、安定世代における平均にほとんど差がなく、分析に使えるものがなかったため、グラフからは外した．評価指標 4・5 は、同一状況で異なる概念化をする言語ペア (Gap 間) のみの形式の距離の平均を、index 単位と bag 単位で測っている．

図 4.7 は、安定世代をもつ run における各評価指標 4・5 の平均である．このグラフから分かることは、推論成功確率が非常に高い (1.0) であれば、概念化間の言語の形式の距離が離れており、それ以外の場合では、ほとんど同じ形式が発話され概念化が形式に反映することができていない言語が創発している、ということである．

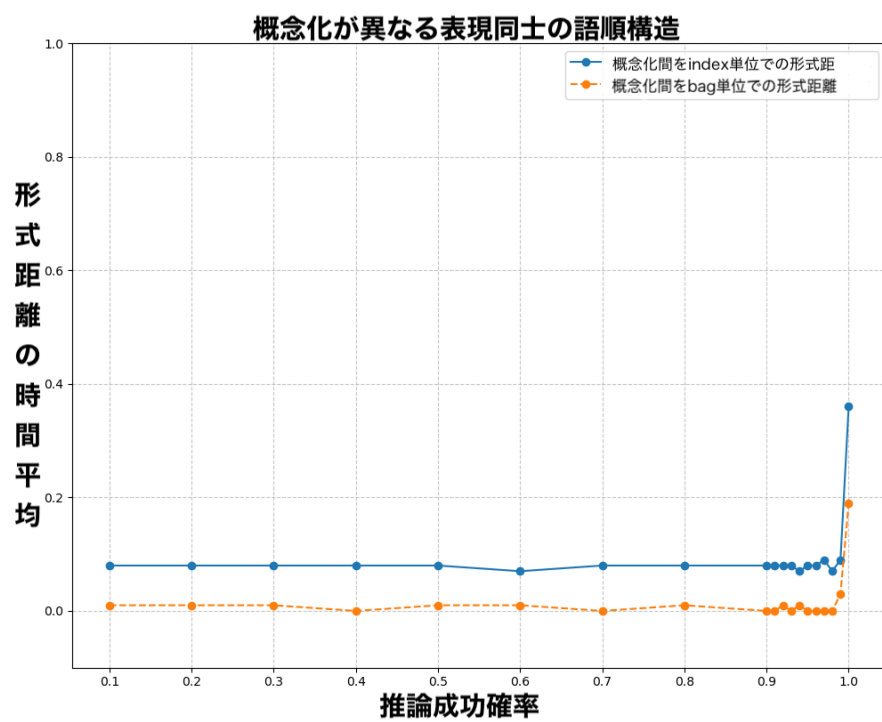


図 4.7: 評価指標 4・5 による概念化間の言語の性質

第5章 考察と議論

本章では、実験結果と評価指標の設計について考察する。次に、その結果を生じさせる要因について、学習アルゴリズムとモデル化した文化進化様式と意味構造の設計の観点から議論し、本モデルの限界を明らかにする。加えて、進化言語学の議論とどのように接続できるのかについて議論する。

5.1 考察

5.1.1 概念化の推論と構文交替

本研究の目的は、構文交替の創発条件の解明、すなわち、同一状況に対して異なる概念化をすることで異なる形式を使い分ける構文交替現象がどのような条件で創発しうるのかを明らかにすることであった。図 4.7 で明らかになったことは以下である。同一状況に対する異なる概念化を表現する形式同士が「どのような対応」にあるのかを、順序と形式要素の観点から観測し距離として計算する評価指標 (Separated Sequential Average と Separated Bag Average) を使うと、推論成功確率が 1.0 はそれ以外の確率と比較して、相対的に形式の距離が長いということである。また、安定世代を持つ run において、親の概念化に対する推論が常に成功する推論成功確率が 1.0 の場合、安定状態にある言語知識は概念化を異なる形式で表現する言語が創発し、逆にそれ以外では概念化を同じ形式で表現する言語が創発する、ことが観察された。とくに、推論成功確率が 1.0 以外の全ての確率の段階 (0.1, 0.3, 0.5, 0.7, 0.9, 0.91, 0.92, 0.93, 0.94, 0.95, 0.96, 0.97, 0.98, 0.99) において 0.1 から 0.99 の範囲で概ね同じ傾向が観察されたが、個々の推論成功確率による微細な違いについてはさらなる分析が必要である。この意味において、仮説は今回の実験条件のもとでは支持されず、仮説と異なる傾向が観察された¹。当初の仮説では、高い推論成功確率であれば構文交替が起こると考えたが、1.0 未満では起こらなかった。そこで、結果から言えることを再度以下で述べ直す。

構文交替が創発しうる条件は、子が親の概念化を完全に正しく推論できる場合である。より詳しく述べ直すと、概念化の値を正しく推論できる確率が非常に高い 1.0 の場合、意味要素単体の語彙は共通するものの、異なる概念化を異なる順序や冗長な形式を付加することによって

¹100seeds 全てにおいて一貫した結果が出たわけではなく、また、その頻度も低い。

形式に反映する言語が創発しうる。一方で、推論成功確率が少しでも 1.0 未満だと、概念化を異なる語順によって形式に反映する言語が創発することはほとんどなく、また全く同じ語順、あるいは全く同じ形式を持つ言語が創発する。

また、実験結果を踏まえれば、本モデルにおいて親がどのように概念化しそれを形式に反映させたのかを子が完全に (確率=1.0 で) 推論できなければ、概念化はたやすく無視されてしまい、語順が1つにまとまり多様性が抑制される結果となっている。つまり、形式だけからはどちらの概念化に基づいていたのかが判別できないということである。形式が同じだからと言って意味も同じであることは意味しないが、本シミュレーションで扱える音声羅列としての形式では概念化された意味の違いを表現することはできない。自然言語においては、制約の多い本シミュレーションとは異なり、感情や態度や声色を含む音声形式の周辺的な要素であるパラ言語情報によって、細かな意味が伝達される。仮に同じ形式を発したとしても、その形式に含まれる意味が異なることがある。今回のシミュレーションの設定では、述語論理で表される意味と音声形式のみを言語として扱っているが、実際の言語伝達場面において、多様な情報がやりとりされている。概念化の違いが無視されてしまったことは、述語論理で表される意味と音声形式のみからしか概念化を推論し学習することができないという制約に起因するかもしれない。しかしながら、いくらパラ言語情報が豊富な実際の言語伝達場面が存在したとしても、意味と形式の対応関係が不明瞭だったと想定される言語創発最初期においては、そのパラ言語情報が意味を推論するのに貢献するかどうかは定かではない。

5.1.2 概念化の無視を「文法の単純化」と捉えてみるとむしろ自然

推論が確実に行われないと構文交替が創発する余地がないという結果は、当初の仮説を立てた段階からするとネガティブな結果であるかもしれない。しかしながら、概念化の違いを構文交替で反映するという言語は反映しない言語よりも単純であると解釈した場合、この単純さは概念化を完璧に推論できないような相手との言語伝達においてむしろ自然であると言えるかもしれない。2.2.2 小節でも引用したように Wray and Grace (2007) は「見知らぬ人々と話すこと (talking to strangers)」によって意味の解釈可能性が高まり、言語が単純化すると述べる。彼らは見知らぬ人々と話すコミュニケーションを「ソト向け (exoteric) コミュニケーション」と呼び、言語に単純な規則体系を持たせる要因であるとする。言語創発最初期において現在あるような言語単位 (日本語や英語など) は存在せず、常に確立していない言語と呼べるか分からないような記号の集合を話していたと仮定すれば、本研究での推論成功確率が少しでも 1.0 未満である場合の言語伝達は「ソト向けコミュニケーション」であったと言える。これはちょうど Wray and Grace (2007) が言うところの「人間言語の自然な初期状態 (the natural default setting for human

language)」に相当する。Kirby et al. (2008)においても、忠実な世代間継承には学習可能でかつ曖昧性がない必要があると述べている。推論が完全には成功しない設定における言語継承においては、同一状況を異なる概念化をすることは、推論が失敗すると言う意味において曖昧性を含むことになる。そこで、曖昧性を減らすこととして概念化の区別を無視することで、忠実な言語継承を促進したという可能性を示唆できる。また、推論が失敗する「ソト向け」コミュニケーションが発生してしまうのは、比較的大規模な集団における言語伝達であると言える。Lupyan and Dale (2010) は 2000 以上の言語を統計分析し、大規模集団で話される言語は格変化や動詞の活用など形態的な複雑さが単純化されていく社会歴史的变化を捉え、言語の文法が使用と学習の場である社会的環境に部分的に依存することとした。本研究の結果においても、推論が必ず成功する場合であれば冗長な形式要素の付加によって異なる概念化を区別していると思われる言語が観察できたが、推論が失敗する場合では異なる概念化の区別に冗長な形式要素²が使われることはなかった。つまり、推論が失敗する大規模集団での「ソト向け」コミュニケーションによって、言語の形式が単純化すると言えるかもしれない。

5.1.3 作成した評価指標についての更なら改善の余地

3.4 節で作成した評価指標は、人工的に「構文交替」させた完全に構成的で体系的な言語知識を対象に、その機能の評価した。そして、従来指標よりもその人工的な言語の性質を高く評価していた。十分かはともかく、相対的には改良版であると言える。それにもかかわらず、表 4.3 で見たように、同一の意味を複数の異なる汎化ルールを使って形式が表現できる場合、意味と形式が 1 体 1 対応しないため、結果として評価指標の値が低いままとなっている (図 4.5, 4.6)。安定世代の定義より言語知識としては同じような言語知識が連続して世代間継承されることはあっても、異なる意味要素が同じ形式要素に対応することや、類似な意味が異なる形式に対応することは不思議ではない。特に、異なる意味要素が同じ形式要素に対応することは構成性の定義に反するわけではない。なぜならば、構成性の定義に部分的な意味要素に対応する形式要素の区別は含まれていないからだ。あくまでも部分的な意味要素が全体を構成し、その部分的な意味要素に対応する形式要素が全体の形式表現を構成していれば良いのである。繰り返しになるが、形式表現において異なる意味が異なる形式に対応しているかは本来問題にならないはずである。

繰り返しになるが、構成性の定義には「部分から全体へ」の過程のみが含まれているのであり、部分同士、全体同士の区別は定義には含まれていない。作成した評価指標が前提としている構成性の定義に、意味と形式の 1 対 1 対応という性質が計算過程において暗黙のうちに取り込まれていた可能性がある。こうした暗黙の前提をどのように排除できるかについては今後の課題である。

²但し、この冗長な形式要素を格や活用の一部と結論づけることはできない。

5.2 議論

実験の結果が示したことは、同一状況に対する概念化の区別を順序や冗長な要素の付加によって形式に反映する構文交替が創発しうる条件は、概念化の値を正しく推論できる確率が1.0の場合のみであるということ。そして、概念化の値を少しでも正しく推論できない場合には、構文交替は生じず、概念化の区別が無視されたように見える類似な順序で表現される言語が生じうる。それでは、なぜこのような条件でない限り、構文交替が創発しないのかについて議論する。

5.2.1 学習アルゴリズムから予測できること

まず、3.2.1小節の中に「学習アルゴリズムの最小限の修正」において確認した通り、学習アルゴリズムにはある前提が組み込まれていた。それはChunkにおいてのみ概念化の値(CV)を区別して学習するということであった。加えて、単語ルールを持つということを、概念化とは独立に、行為それ自体や対象それ自体といった単一の意味要素を特定していることと見做している、ことであった。ChunkにおいてCVを区別していることから、1個体の言語知識の中で学習に使われる「文ルール」のペアはCV別に整理されていると言える。一方で、Category-IntegrationとReplaceはCVの区別ない単語ルールを使うので、言語知識の中でCV別の文ルールとは独立に、存在していると言える。いわば、単語ルールは学習アルゴリズムの性質上、両CVに共通して、つまりCVの区別を無視して、文ルールとペアとなって学習できる。単語ルールを使って学習するということは「状況」の構成要素を認識できることであるので、概念化とは独立である。CVのみが異なる「状況」であっても、その状況が少なくとも「同一状況」であるということは共有できていることを担保していた。表4.3, 4.5で見たように、両概念化における言語知識において共通の語彙が共有される仕組みである。また、Chunkにおいては概念化を区別するが、Chunkによって生成される単語ルールはReplaceによって概念化の区別をせずに転用される。これにより、本来一方の概念化の言語知識だけではReplaceできないルールまでReplaceにより汎化され、ますます共通の語彙が両概念化の間で共有されていくことになる。例えば、以下のChunk学習が起きたとする。

Rule 1: $S/kick(john, \underline{mary})/0 \rightarrow abc\underline{j}kde$

Rule 2: $S/kick(john, \underline{anya})/0 \rightarrow abc\underline{s}tude$

↓

Rule 3: $S/kick(john, Y)/0 \rightarrow abcQ/Yde f$

Rule 4: $Q/mary \rightarrow jk$

Rule 5: $Q/anya \rightarrow stu$

ここでは、Rule 1と2のペアから、Rule 3, 4, 5を学習した。

次に Chunk によって学習した単語ルールの Rule 4 を使った Replace 学習が起きたとする。

Rule 4: $Q/\underline{mary} \rightarrow \underline{jk}$

Rule 6: $S/help(john, \underline{mary})/1 \rightarrow mnmr\underline{jk}t$

↓

Rule 4: $Q/mary \rightarrow jk$

Rule 6: $S/help(john, y)/1 \rightarrow mnmrQ/yt$

ここでは、Rule 6 と、Chunk で学習した Rule 4 のペアから、Rule 4, 6 を学習した。このように、単語ルールは Replace や Category-Integration において CV に関係なく学習ペアを作ることができるし、Rule 4 のように、両 CV で共通の単語ルールとして言語知識の中に存在していることがわかる。

しかしながら、親の発話過程で生成される CV 付き導出規則を、推論成功確率が 1.0 未満である場合だと、子はそのまま受け取ることができない。つまり、親の導出規則とは異なるものを子は受け取ることとなる。このように CV が区別できずに混合することで、CV を区別して学習する Chunk において、親とは異なる文ルールペアを使って学習していくことになってしまう。推論成功確率が 1.0 未満である場合で創発し得た言語が、全く同じ言語形式を持つように至った経緯であると考えられる。

5.2.2 本モデルの設計から予測できること

モデル化した文化継承

今回のモデルにおいて厄介なことは、なかなか収束せずに言語知識の総数が増減を繰り返すことが多かったことだ。だからこそ、安定世代を定義その世代をもつ run における言語知識を分析した。安定世代を持たない run における言語知識の中身を見てみると、同じ意味や複数の類似する意味を表現できるルールが複数存在していた。このような場合、発話の時複数のルールからランダムに選択するため、類似する意味を異なるルールによって生成することで、次世代において汎化しづらい発話となり、学習が進まないことの原因である。このような事態は、おそらく「世代間継承のみ」のモデル化に起因する。というのも、コミュニケーションという世代内の相互作用をモデルに組み込むことで、例えば、相手が持っている単語ルールを優先的に使うといった交渉が成立するかもしれない。このようなことを暗に持ち込むためには、意味要素と形式要素の対応関係を曖昧にさせるような多義語や同義語の単語ルールを削除するアルゴリズムを導入したり、単語ルール同士の弁別性を高めるように形式長を冗長にするアルゴリズムを導入したり、同じ意味を表現できる複数の文ルールを削除する必要もある。但し、文ルールにお

いても表現力に影響しない範囲で削除する必要があるため慎重さを要する。また、こうした1つの規則体系に導くような恣意的な干渉に対して実装上だけでなく、理論的なサポートをしなければならない。

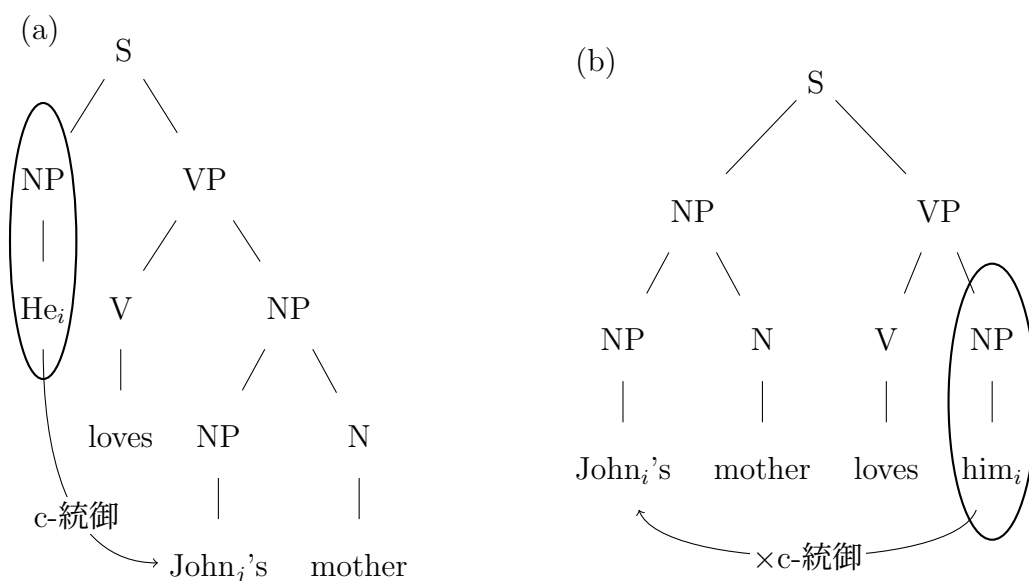
5.3 現行モデルの課題

本モデルの課題として、外部世界の意味要素と概念化の値を同列に扱っていることが挙げられる。概念化がどのように外部世界とかわかり、意味の中で位置付けられているのかは定かではないが、外部世界を構成する1つの要素として追加されているような現状のモデル化は不適切であったかもしれない。というのも、意味の二重性小節で取り上げた「意味の二重性 (duality of semantics)」を研究する生成文法一派において、述語項構造で規定される意味役割とは無関係な談話・語用論的意味に分類され、かつ、階層的に上位にあると分析される。仮に、認知言語学でいうところの概念化された意味が談話・語用論的意味に分類されるとすれば、意味の階層性を本研究における導出規則の条件 C^* に取り込めていないことになる。

こうした意味の階層性を取り込めていないことは、本研究の射程が「能動文・受動文の交替」といった言語現象を包括できないことに対応する。なぜなら、生成文法において「主語と目的語」の非対称性つまり、統語構造における位置関係が異なる点が指摘されているからだ (齋藤, 2024)。詳細は省くが、目的語は動詞句 (VP) 内に生起するのに対し、主語は動詞句と併合する。なぜこのように、構造的な位置が異なると仮定した方が良いのかは、以下の例から導かれる。

(a) He_i loves $John_{j \neq i}$'s mother.

(b) $John_i$'s mother loves him_i .



これらの例は、先行詞 John とその代名詞 (he/him) が同一指示を与えられる条件について構造的位相関係に依存することの例である。(a) の例では、代名詞が John を c-command する構造的位相にあるため、同一指示が得られない。つまり、He が指す人物と John は異なる人物である。それに対して、(b) では c-command する構造的位相にないため、同一指示を許容する。つまり、John と him は同じ人物を指すことが許される。このように、主語位相と目的語位相において同一指示が許されるかどうかはその構造的位相関係に依存するということは、併合の順序が異なることを意味する。

仮にこうした階層構造に依存した主語と目的語の非対称性があるならば、本研究で創発したかに見える、概念化を「語順構造」で形式に反映する言語は、単に線形順序に従っているだけで、階層構造に従ったものではない、ということになる。

それはつまり、図 2.3 で端的に示したうち、上から 3 つめの前置詞としてカテゴリーされる述部の構文交替しか説明対象に据えることができていない可能性がある。この説明対象の射程の推定は、理論言語学からの知見をさらに動員する必要がある課題である。

第6章 結論

6.1 まとめ

本研究の目的は、これまでの言語の文化進化研究の構成論的手法と認知言語学の知見を組み合わせ、より具体的な言語現象の創発条件を解明することであった。詳細としては、同一状況に対して異なる概念化をすることで異なる形式を使い分ける構文交替現象がどのような条件で創発しうるのかを明らかにすることであった。

この目的のため、親は外部事象を主観的に概念化して発話し、子は推論して学習するという世代間継承をモデル化し、言語の文化進化シミュレーションを行った。さらに、同一状況に対する異なる概念化を表現する形式同士の対応関係を観測する、評価指標を作成し、分析に用いた。

シミュレーションの結果、聞き手(子)が話し手(親)の概念化を完全に推論して学習する場合にのみ、概念化を順序や冗長な要素の付加によって反映する言語が創発する可能性が示唆された。また、相手の概念化を少しでも推論できず失敗すると、異なる概念化は同じ順序、同じ語彙によって表現される言語が創発する。つまり概念化の区別は無視されることがわかった。

6.2 結論

本研究の結果が示したことは、推論が完全には成功しない言語継承において、概念化を無視することは言語の曖昧性を減らし、忠実な世代間継承を可能にしたかもしれないということ、また、推論が失敗することで概念化を無視するようになるという言語の単純化は、先行研究で示唆された見知らぬ人々と話す「ソト向けコミュニケーション」に限ったことではなく、本研究の垂直伝播のみの世代間継承にも共通して起こる現象であることである。

これらの知見から、言語創発最初期においては言語使用者の主観的な概念化によって形成された言語が創発する可能性は低く、概念化を形式に反映する文法現象(今回は構文交替のみ)の創発は文化進化のもっと後の方であった。また、「言語は認知の反映である」とする認知言語学のスローガンは言語創発最初期には適用されない可能性がある、ということである。

6.3 残された課題

構文交替をロバストに創発させるメカニズムについて調べることができていないことは、本研究の大きな課題である。それに加え、以下に列挙する諸々の要素について解決することも今後の大きな課題である。

構成論的手法全体の妥当性議論 構成論的手法において、一般的に、恣意的なモデル化に基づく限定的な結果であることを回避する必要がある。そのためには本研究に則せば、異なる学習アルゴリズム（ニューラルネットワークなど）の採用や人間を被験者として実験室実験、意味空間を拡張することや、概念化の自由度を連続的にするなどによる感度分析を行う必要がある。

認知バイアスの欠如 人間言語には「自然な表現」というものがある。今回のモデルに即して日本語の例を挙げれば、「机の上に辞書がある」は自然な表現だが、「辞書の下に机がある」は不自然な表現である。論理的にはどちらの言い換えも可能であるが、日本語母語話者の直感においても統計的な頻度において偏りがある。人間の認知のバイアスである。しかし、これも本モデルでは大雑把に取り扱っているのみであり、本格的に取り扱っていない。これも課題である。

言語の性質の評価指標の作成 構成性の定義は「全体の意味が部分の意味とその組み合わせ方で表現される性質」である。これは「部分から全体」を導く過程についての定義である。そしてこれはある意味と形式の1ペアのみについて当てはまれば良い定義である。しかしながら従来の評価指標 TopSim やその計算過程における亜種は、意味と形式の対応関係のみに注目しその導出過程は考慮されておらず、また言語全体の統計的性質を評価している。つまり、2重に本来の定義とは異なる性質を評価していることになる。本研究ではこの部分について改良することはできず、なかば盲目的に先行研究の評価指標を使い、かつ、大きく従来指標に基づいて本研究の対象である概念化を考慮した評価指標を作成しただけであった。1ペアのみから言語全体の統計的性質へ評価対象を拡張させることは当然ながら必要であるが、その計算手法として従来のものが妥当である根拠を示す必要がある。もっとも大きな課題は「部分から全体」への導出過程を含めた計算手法を考案することである。

謝辞

これまで私は色々なコミュニティに参加し、抜け出し、居座り、浮遊してきた。そのすべてに言及することはできないが、思い付く限り感謝を記す。(橋本研・JAIST卓球部・計算論的精神医学・分析哲学コミュニティ・EmeCom/EvoLang・慶應井上ゼミ・「身体と言語」研究会・共創言語進化若手の会・言語学院生学部生の会・蟹グルサイゼ会・成城会・家族などなど)

JAIST

橋本敬教授(橋本さん・hash)。まずは、本研究を進めるにあたり、ご指導いただき深く感謝します。最初は、言語進化の研究が日本国内で行えて、入学試験が楽チンで、かつ、入学前の面談で「話しが合う」という私の不遜な感想によって、橋本研に入ることを決めました。しかし、入学してすぐに、認知科学の特集「ことばの認知科学：言語の基盤とは何か」で言語創発・言語進化のサーベイ論文(上田他, 2024)の共著メンバーに入れてくださり、様々に刺激的な体験をさせていただきました。そして、橋本さんと船蔵颯(当時京大D1)さんの手を借りてEvoLang XVという言語進化の国際会議にポスター発表でacceptされ、2024年5月に初一人海外・初対面学会を(しかも英語で)無事終えました。つまり、研究面において、とても手厚くご指導していただきました。また、本研究以外のこと、例えば、ジャズやジン、哲学のお話など、幅広い話題を持つことができたのは、本当に嬉しかったです。私の日々の学業が「お勉強偏重(研究逃避)」になっているのを、時にアラートを出し、時に無視してくださり、そしてとてもたま〜に、しかし強烈に研究に引き戻すお言葉をかけてくださいました(胸に刻んでおります...)。正直、自分の研究力(1つのことを粘り強く取り組む集中力や誠実性、得られた結果を考察する思考力)のなさに打ちひしがれる辛い2年間でした。

黒川瞬(当時jaist助教・現在阪大准教授!)さん。黒川さんとはゼミやティーミーティング以外の場所で本当にたくさんの思い出があります。諸々を割愛せざるを得ませんが、とても楽しかったです。

水本正晴准教授には、副テーマ研究において熱心にご指導していただきました。郡司ペギオ幸夫の逆ベイズ推論(Gunji, 2023; Gunji et al., 2017; 郡司, 2016)とFristonの自由エネルギー原理(Friston, 2010; Friston et al., 2006; 乾・阪口, 2020)を無理やり結びつけて、自由エネルギー原理における他者推論を批判的に扱うという私の稚拙なアイデアを、伝統的な哲学の議論と接続していただきました。

垣花元貴さん・星宏侑さん・笹森なおみさん(パンケーキ)は、1年しか被らなかった研究室の先輩方です。入学前、あなたたちとお話しするのを楽しみにしておりました。入った瞬間、マジで「なんでいねえ〜んだよっ!?!」と思いました。僕の図々しさによって結果的にたくさん話すことができ、嬉しかったです。垣花さんは僕の知っていることを全て知っていて、しかも深く理解している上で優しいトーンで穏やかにいつも話してくださり、大変貴重な時間でした。星さん

は僕の研究を理解しようとしてくれて、度々 zoom などでお話を聞いてくださいました。笹森さんはたまに僕の私の研究について励ましたり、厳しい言葉をかけてくださり感謝します。

Siri-on Umarin さん・成太俊さん（成さん）・周豪特さん（にゃんぱす）・黄文蓮さん（ぶんぶん）・秦慕君さん（ぼくん）・笠野純基さんは、研究室のまだ在籍し（続け）ている先輩方で、私の日頃の失礼極まりない態度を温かく、半分呆れて、見守ってくださり、ありがとうございました。

箕輪朗くん（ミノ）・古川建くん（タケ）は、僕の本当の意味での同期です。精神的にも肉体的にも支えられました。ありがとうございました。箕輪くん、彼女と幸せになれよ！ 古川くん、ちゃんとご飯食べよ！

郝抒妍さん（ハオ・パオ・ピオ）・名嘉琉星くん・松崎由幸くん・王子斌くん（ズビン）・中村都夢くん（トム）は、研究室の直接の後輩であり、不適切な行動を失礼致します。

研究室外では、寮の隣人である二之宮昇くん（ニノ）に、まず感謝します。入学当時、ふとんのこと、水回り・家具・食事のこと、履修登録・wifi のこと、ほとんど学生が必要なものは全てニノに教えてもらいました。ほんとうにニノがいなかったら生きていけなかった。廊下でたまに話せてよかった。本当に感謝する。

中野和久くんは、これまた寮の同階で、研究面でも生活面でもとても助けられた。中野くんがいなければ申請書を書き上げることも、締切を認知することもできなかった。本当に心の底から感謝する。今後もよろしくね。

本名貴喜くんとは、なぜか気が合って、研究のことをたくさん話せる友達だから、今後もよろしくね。

エン・カリンさんとは、最後の1年間、本当にお世話になりました。一緒に大切な時間を過ごさせていただき、また、支えていただきました。

JAIST の事務の方、教職員の方、研究棟と寮の清掃員の方、デイリーヤマザキの従業員の方、などとにかく JAIST の関係者の皆様には、私の気付けない部分で、支えられてきたはずです。この場を借りて、感謝いたします。

つるのや、YuuYuu（ゆうゆう）は、バスと徒歩だけで向かえる JAIST 周辺の飲食店で、頻繁に通わせていただきました。本当にどの料理もおいしく、幸せな気持ちになったことを覚えています。ありがとうございました。

JAIST 外：学部時代からの継続。

井上逸兵教授（慶應）・吉川正人さん（当時慶應非常勤講師、現在群馬大准教授）・古賀裕章准教授は、私に言語学という学問を教えてくださいました。順番が前後しますが、大学1年時に受けた古賀先生の講義によって言語学そのものへの興味を膨らませました。大学3年時の吉川先生の授業で理論言語学・進化言語学へと興味を鋭くしました。井上先生のゼミで興味関心をさらに広げました。

吉川正人先生は、特筆すべき存在です。特に親しくさせていただき、私の興味をホイホイ連絡すると必ず返信をくださり、議論に付き合ってもらいました。大

学4年での講義では受講生が私1人であり、キャンパスを抜け出して、パブやスタバで議論をしてくださいました。分野を超えて僕が関心を持てたのは、吉川先生のおかげです。思い出せるのは、4E 認知・逆ベイズ推論・訂正可能性・Marr の3水準などですが、もっともっとたくさんの話題に付き合ってくださいました。ありがとうございました。

井上ゼミのメンバーにも感謝します。僕は大学院生、彼らは社会人となった今でも、旅行や飲み会で親しくしています。かけがえのない存在です。

中条太聖さん（当時京大 D1）・外谷弦太さん（当時ポスドク）には、僕が大学3年の時に参加した共創言語進化の若手の会でお世話になりました。知識も考え方も言語化も方向性も今よりもっとあやふやな時に、言語進化の読書会やゲーム理論の勉強会で大変お世話になりました。

坂本瑞生さん（東北大 D2; 餅さん）は言語学院生・学部生の会の発起人であり、理論言語学の重要な考え方を教わりました。本当に偉大です。

齋藤諒弥さん（神戸大 M2; りょーや）・長谷川優菜さん（名古屋大 M2; ふぁせ）・山崎俊くん（上智大 M1; しゅん）とは、500 ページ近い生成文法の入門書 Carnie の”Syntax”と一緒に読んだメンバーです。学部4年の9ヶ月間ずっとオンラインで勉強会して、とても仲良くなりました。勉強会以外でも、夜中に電話したり議論したりできて、楽しかったです。

学術的なつながりから離れて

石若駿さん・永見拓也さん・坂井旭（宇奈月温泉の喜泉次男坊）さんたちには、私が金沢でジャズピアノ演奏を続けるきっかけを作ってくださいました。研究の最大の息抜きでした。改めて感謝いたします。

江口一生（えぐり）・矢島理勢は、地元が一緒に、なぜか青山学院中・高等部でも一緒に、青春時代のかげがえのない友人です。話すだけで童心に戻って、ゲラゲラ笑えるのは、僕にとって本当に貴重です。

最後に、家族。

岩村太郎（父）・岩村純子（母）・岩村牧穂（姉）には、精神面・金銭面でとてもお世話になりました。これまで大切に育ててくださりありがとうございました。

参考文献

- Azumagakito, T., R. Suzuki, and T. Arita (2013) “Cyclic behavior in gene-culture coevolution mediated by phenotypic plasticity in language,” in *Artificial Life Conference Proceedings*, pp. 617–624, MIT Press One Rogers Street, Cambridge, MA 02142-1209, USA journals-info ….
- Batali, J. (1999) “The negotiation and acquisition of recursive grammars as a result of competition among exemplars,” *Linguistic Evolution Through Language Acquisition: Formal and Computational Models*, pp. 111–172.
- Benítez-Burraco, A. and L. Progovac (2020) “A four-stage model for language evolution under the effects of human self-domestication,” *Language & Communication*, Vol. 73, pp. 1–17.
- Berwick, R. C. and N. Chomsky (2017) “Why only us: Recent questions and answers,” *Journal of Neurolinguistics*, Vol. 43, pp. 166–177.
- Berwick, R. C., K. Okanoya, G. J. Beckers, and J. J. Bolhuis (2011) “Songs to syntax: the linguistics of birdsong,” *Trends in cognitive sciences*, Vol. 15, No. 3, pp. 113–121.
- Beuls, K. and S. Höfer (2011) “Simulating the emergence of grammatical agreement in multi-agent language games,” in *Twenty-Second International Joint Conference on Artificial Intelligence*, Citeseer.
- Brighton, H. (2002) “Compositional syntax from cultural transmission,” *Artificial life*, Vol. 8, No. 1, pp. 25–54.
- Brighton, H. and S. Kirby (2006) “Understanding linguistic evolution by visualizing the emergence of topographic mappings,” *Artificial life*, Vol. 12, No. 2, pp. 229–242.
- Brighton, H., K. Smith, and S. Kirby (2005) “Language as an evolutionary system,” *Physics of Life Reviews*, Vol. 2, No. 3, pp. 177–226.
- Cann, R. (1993) “Formal Semantics: An Introduction.”

- Chaabouni, R., E. Kharitonov, D. Bouchacourt, E. Dupoux, and M. Baroni (2020) “Compositionality and generalization in emergent languages,” *arXiv preprint arXiv:2004.09124*.
- Chomsky, N. (2005) “Three factors in language design,” *Linguistic inquiry*, Vol. 36, No. 1, pp. 1–22.
- (2008) “On Phases,” in Freidin, R., C. P. Otero, and M. L. Zubizarreta eds. *Foundational Issues in Linguistic Theory: Essays in Honor of Jean-Roger Vergnaud*, pp. 133–166, Cambridge, MA: MIT Press.
- (2014) *The minimalist program*: MIT press.
- Cinque, G. et al. (1994) “On the evidence for partial N-movement in the Romance DP,” in *Paths towards universal grammar*, pp. 85–110: Georgetown University Press.
- Culbertson, J. and S. Kirby (2016) “Simplicity and specificity in language: Domain-general biases have domain-specific effects,” *Frontiers in psychology*, Vol. 6, p. 1964.
- Friederici, A. D. (2009) “Pathways to language: fiber tracts in the human brain,” *Trends in cognitive sciences*, Vol. 13, No. 4, pp. 175–181.
- Friston, K. (2010) “The free-energy principle: a unified brain theory?” *Nature reviews neuroscience*, Vol. 11, No. 2, pp. 127–138.
- Friston, K., J. Kilner, and L. Harrison (2006) “A free energy principle for the brain,” *Journal of physiology-Paris*, Vol. 100, No. 1-3, pp. 70–87.
- Griffiths, T. L. and M. L. Kalish (2007) “Language evolution by iterated learning with Bayesian agents,” *Cognitive science*, Vol. 31, No. 3, pp. 441–480.
- Gunji, Y.-P. (2023) “Connection and Disconnection of Perception and Memory: Déjà vu, Bayesian and Inverse Bayesian Inference,” *Bergson’s Scientific Metaphysics: Matter and Memory Today*, p. 195.
- Gunji, Y.-P., S. Shinohara, T. Haruna, and V. Basios (2017) “Inverse Bayesian inference as a key of consciousness featuring a macroscopic quantum logical structure,” *Biosystems*, Vol. 152, pp. 44–65.
- Hashimoto, T., M. Nakatsuka, and T. Konno (2010) “Linguistic analogy for creativity and the origin of language,” in *The Evolution of Language*, pp. 184–191: World Scientific.

- Jackendoff, R. (2010) “Your theory of language evolution depends on your theory of language,” *The evolution of human language: Bilingual perspectives*, pp. 63–72.
- Jarvis, E. D. (2004) “Learned birdsong and the neurobiology of human language,” *Annals of the New York Academy of Sciences*, Vol. 1016, No. 1, pp. 749–777.
- Kirby, S. (2001) “Spontaneous evolution of linguistic structure — an iterated learning model of the emergence of regularity and irregularity,” *IEEE Transactions on Evolutionary Computation*, Vol. 5, No. 2, pp. 102–110.
- (2002) “Learning, bottlenecks and the evolution of recursive syntax,” *Linguistic evolution through language acquisition: Formal and computational models*, pp. 173–204.
- (2017) “Culture and biology in the origins of linguistic structure,” *Psychonomic bulletin & review*, Vol. 24, pp. 118–137.
- Kirby, S. and M. H. Christiansen (2003) *Language evolution*: Oxford University Press Oxford.
- Kirby, S., M. Dowman, and T. L. Griffiths (2007) “Innateness and culture in the evolution of language,” *Proceedings of the National Academy of Sciences*, Vol. 104, No. 12, pp. 5241–5245.
- Kirby, S., H. Cornish, and K. Smith (2008) “Cumulative cultural evolution in the laboratory: An experimental approach to the origins of structure in human language,” *Proceedings of the National Academy of Sciences*, Vol. 105, No. 31, pp. 10681–10686.
- Kirby, S., T. Griffiths, and K. Smith (2014) “Iterated learning and the evolution of language,” *Current opinion in neurobiology*, Vol. 28, pp. 108–114.
- Kirby, S., M. Tamariz, H. Cornish, and K. Smith (2015) “Compression and communication in the cultural evolution of linguistic structure,” *Cognition*, Vol. 141, pp. 87–102.
- Langacker, R. W. (1987) *Foundations of cognitive grammar: Volume I: Theoretical prerequisites*, Vol. 1: Stanford university press.
- (1990) *Concept, Image, and Symbol: The Cognitive Basis of Grammar*: Mouton de Gruyter.

- (2008) *Cognitive grammar: A basic introduction*, Vol. 535: Oxford University Press.
- (2014) “Conceptualization, symbolization, and grammar,” in *The new psychology of language*, pp. 1–37: Psychology Press.
- Lupyan, G. and R. Dale (2010) “Language structure is partly determined by social structure,” *PloS one*, Vol. 5, No. 1, p. e8559.
- Ren, Y., S. Guo, M. Labeau, S. B. Cohen, and S. Kirby (2020) “Compositional languages emerge in a neural iterated learning model,” *arXiv preprint arXiv:2002.01365*.
- Scott-Phillips, T. (2014) *Speaking our minds: Why human communication is different, and how language evolved to make it special*: Bloomsbury Publishing.
- Scott-Phillips, T. C. and S. Kirby (2010) “Language evolution in the laboratory,” *Trends in cognitive sciences*, Vol. 14, No. 9, pp. 411–417.
- Smith, K. (2020) “How culture and biology interact to shape language and the language faculty,” *Topics in cognitive science*, Vol. 12, No. 2, pp. 690–712.
- Suzuki, T. N., D. Wheatcroft, and M. Griesser (2016) “Experimental evidence for compositional syntax in bird calls,” *Nature communications*, Vol. 7, No. 1, p. 10986.
- (2018) “Call combinations in birds and the evolution of compositional syntax,” *PLoS biology*, Vol. 16, No. 8, p. e2006532.
- Tomasello, M. (2005) *Constructing a language: A usage-based theory of language acquisition*: Harvard university press.
- van Trijp, R. (2010) “Strategy competition in the evolution of pronouns: a case-study of Spanish leísmo, laísmo and loísmo,” in *The Evolution Of Language*, pp. 336–343: World Scientific.
- Wray, A. (2007) 「‘Needs only’ analysis in linguistic ontogeny and phylogeny」, 『Emergence of communication and language』, 53–70 頁.
- Wray, A. and G. W. Grace (2007) “The consequences of talking to strangers: Evolutionary corollaries of socio-cultural influences on linguistic form,” *Lingua*, Vol. 117, No. 3, pp. 543–578.
- 岡ノ谷一夫 (2010) 「言語起源の生物学的シナリオ」, 『認知神経科学』, 第 12 巻, 第 1 号, 1–8 頁.

- 乾敏郎・阪口豊 (2020) 「脳の大統一理論: 自由エネルギー原理とはなにか」, 『岩波科学ライブラリー』.
- 吉川正人 (2017) 「社会統語論の目論見—「文法」は誰のものか」, 『朝倉日英対照言語学シリーズ 発展編 1 社会言語学』, 146–167 頁.
- 橋本敬 (2012) 「繰り返し学習モデルによる文法化の構成論的研究: 創造性と言語の起源における言語的類推の役割」, 『進化言語学の構築: 新しい人間科学を目指して』 藤田耕司・岡ノ谷一夫 (編), 241–258 頁.
- 橋本敬・中塚雅也 (2007) 「文法化の構成的モデル化-進化言語学からの考察」, 『日本認知言語学会』.
- 郡司ペギオ幸夫 (2016) 「知覚と記憶の接続・脱接続: デジャビュ・逆ベイズ推論」, 『書肆心水』, 311–331 頁.
- 山内肇 (2012) 「バリ言語学会が禁じた言語起源」, 『進化言語学の構築: 新しい人間科学を目指して』 藤田耕司・岡ノ谷一夫 (編), 35–53 頁.
- 酒井智宏 (2013) 「認知言語学と哲学—言語は誰の何に対する認識の反映か—」, 『言語研究』, 第 144 巻, 55–81 頁.
- 松本大貴 (2022) 「生成文法にとっての言語進化」, 『言語進化学の未来を共創する』, 9–26 頁.
- 上田亮・石井太河・鷲尾光樹・宮尾祐介 (2023) 「範疇文法導出を用いた創発言語の構成性の評価」, 『言語処理学会 第 29 回年次大会 発表論文集』.
- 上田亮・谷口忠大・鈴木麗璽・江原広人・中村友昭・岩村入吹・橋本敬 (2024) 「言語とコミュニケーションの創発に関する構成論的研究の展開」, 『認知科学』, 第 31 巻, 第 1 号, 172–185 頁.
- 田中拓郎 (2016) 『形式意味論入門』, 第 27 巻, 開拓社.
- 藤田耕司 (2012) 「統語演算能力と言語能力の進化」, 『藤田耕司, 岡ノ谷一夫 (編) 『進化言語学の構築—新しい人間科学を目指して』 . ひつじ書房』, 55–75 頁.
- (2017) 「経済性理論から極小主義まで」, 『理論言語学史』, 58–114 頁.
- 藤田遥 (2022) 「人間言語の漸進進化モデルの構築—レキシコンの成立過程を中心に—」, 博士論文, 京都大学.
- 米納弘渡 (2022) 「外在から内在へ—階層性と意図共有の発生順序に関する考察」, 『言語進化学の未来を共創する』, 177–196 頁.

野中大輔 (2024) 『場所格交替への認知言語学的アプローチ』, 東京.

齋藤衛 (2024) 『生成統語論の成果と課題』, 開拓社.

付録 A：各推論成功確率における 100seeds

全体的な傾向を捉えるために，学習した言語知識がどのように推移するのかを 100seeds 並べた．言語知識には，汎化していない総体論的ルール，変項が 1～3 つのルール，単語ルールに分類される．0.1(図 6.1), 0.3(図 6.2), 0.5(図 6.3), 0.7(図 6.4), 0.9(図 6.5), 1.0(図 6.6) のみを掲載する．

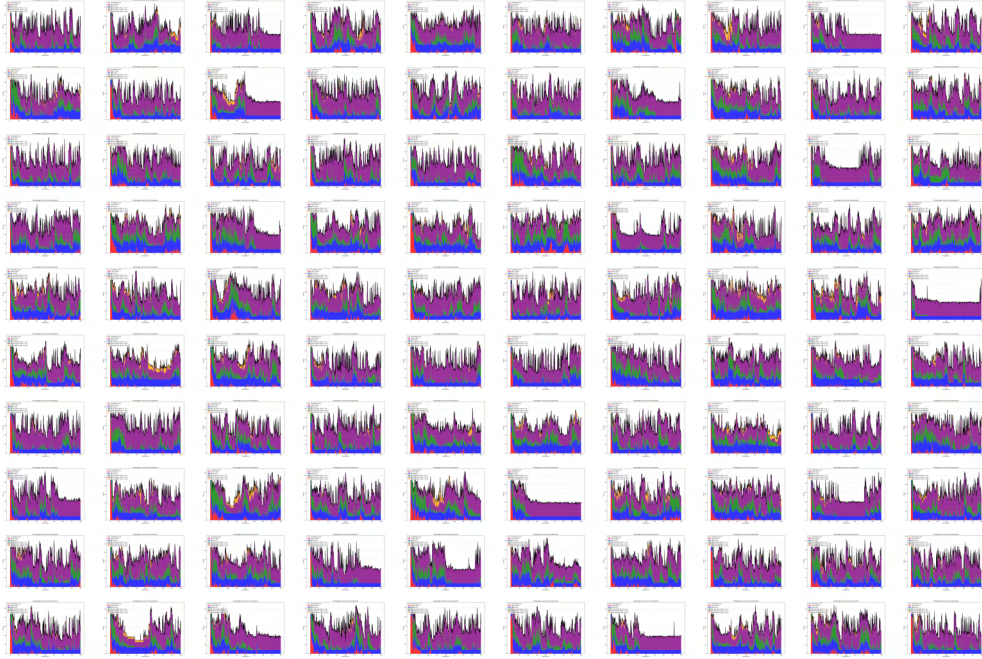


図 6.1: 推論成功確率=0.1 の言語知識 100seeds 分



図 6.2: 推論成功確率=0.3 の言語知識 100seeds 分



図 6.3: 推論成功確率=0.5 の言語知識 100seeds 分

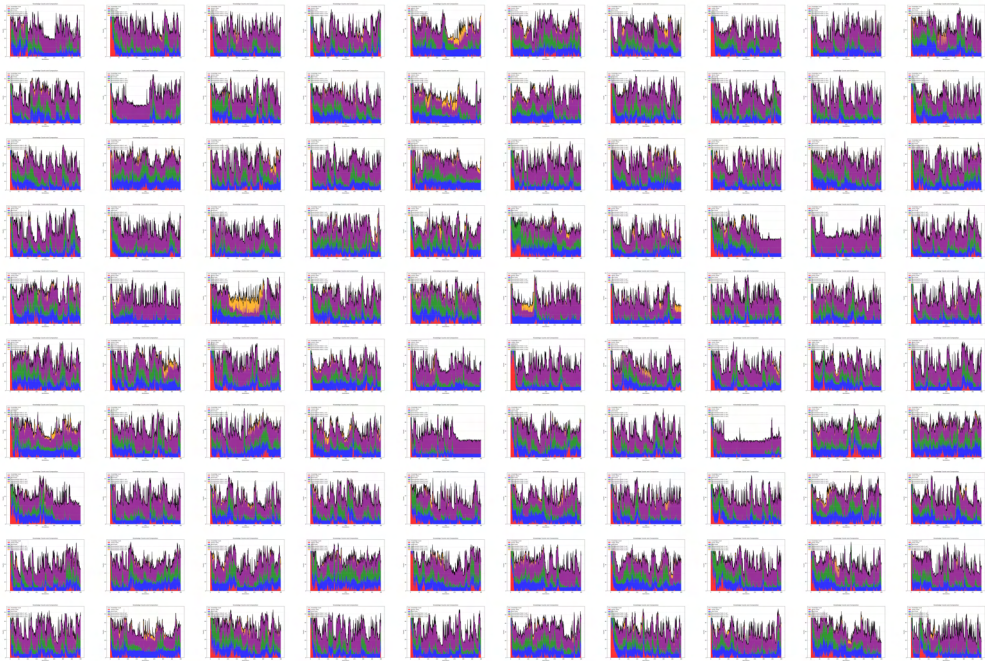


図 6.4: 推論成功確率=0.7 の言語知識 100seeds 分



図 6.5: 推論成功確率=0.9 の言語知識 100seeds 分

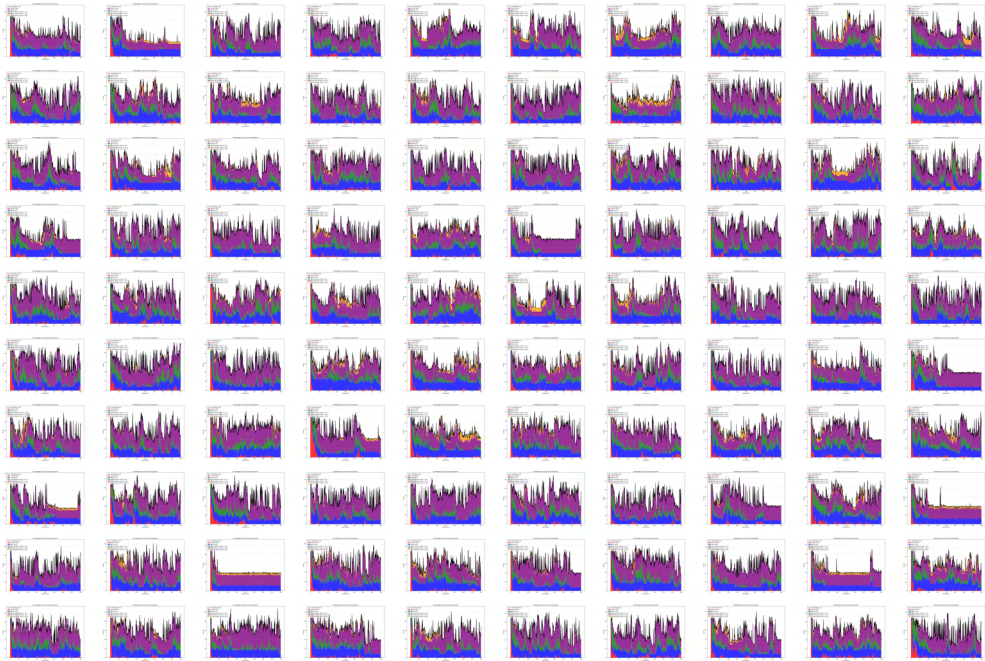


図 6.6: 推論成功確率=1.0 の言語知識 100seeds 分

付録B：各推論成功確率における安定世代をもつseedsの全評価指標の平均

本文でも言及した通り，推論成功確率 0.1 から 0.99 においてほとんど差異がないため，0.1(図 6.7), 0.3(図 6.8), 0.5(図 6.9), 0.7(図 6.10), 0.9(図 6.11) のみを掲載する。

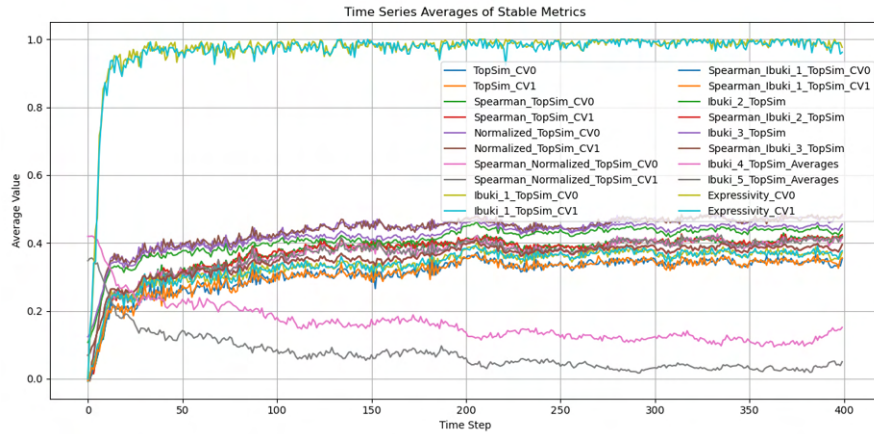


図 6.7: 安定世代をもつ seeds における各評価指標の平均（推論成功確率 0.1）

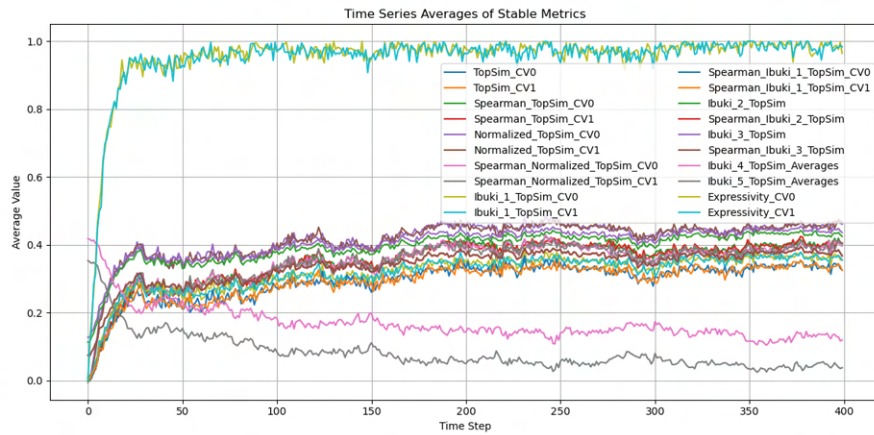


図 6.8: 安定世代をもつ seeds における各評価指標の平均（推論成功確率 0.3）

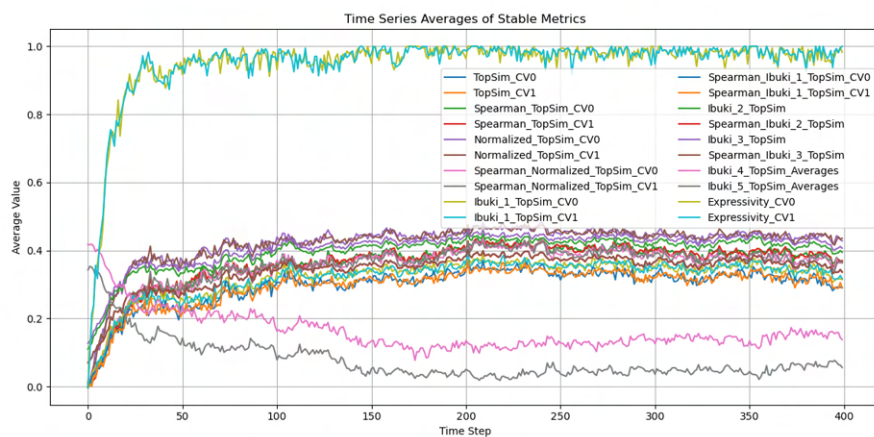


図 6.9: 安定世代をもつ seeds における各評価指標の平均（推論成功確率 0.5）

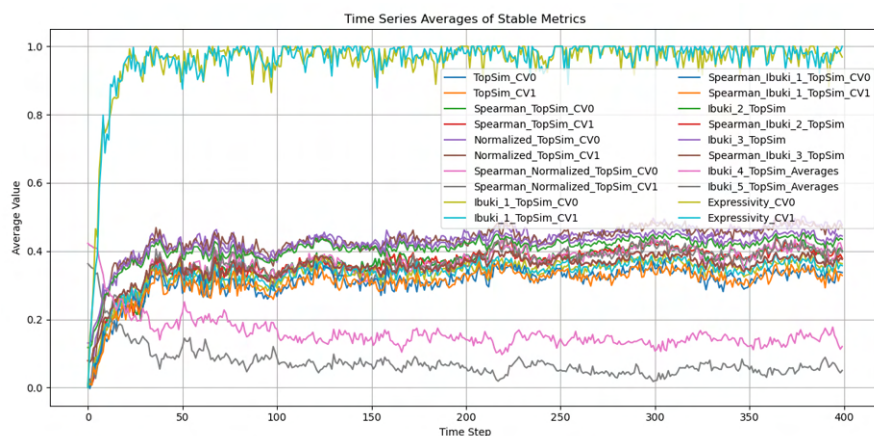


図 6.10: 安定世代をもつ seeds における各評価指標の平均（推論成功確率 0.7）

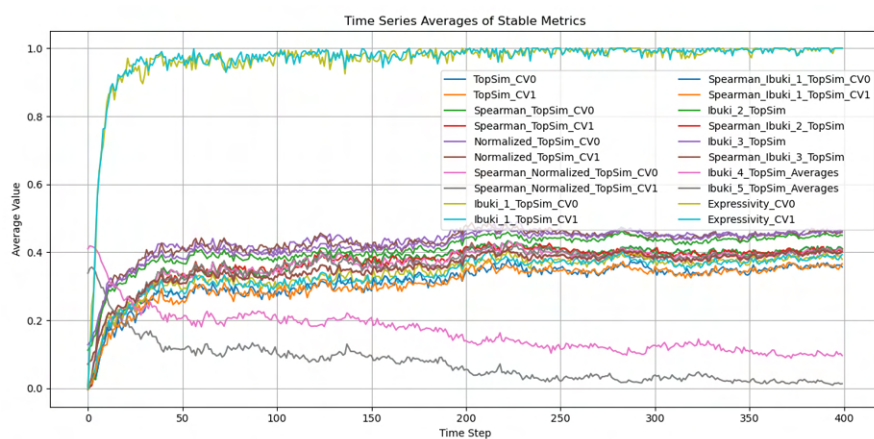


図 6.11: 安定世代をもつ seeds における各評価指標の平均（推論成功確率 0.9）