



Title	幼児の言語理解と産出における文法カテゴリ獲得の計算論モデル
Author(s)	河合, 祐司
Citation	大阪大学, 2017, 博士論文
Version Type	VoR
URL	https://doi.org/10.18910/61747
rights	
Note	

The University of Osaka Institutional Knowledge Archive : OUKA

<https://ir.library.osaka-u.ac.jp/>

The University of Osaka

博士学位論文

幼児の言語理解と産出における
文法カテゴリ獲得の計算論モデル

河合 祐司

2017 年 1 月

大阪大学大学院工学研究科

タイトル: 幼児の言語理解と産出における
文法カテゴリ獲得の計算論モデル

主査: 浅田 稔 教授
副査: 石黒 浩 教授
細田 耕 教授

Copyright © 2017 by 河合 祐司
All Rights Reserved.

要旨

幼児は名詞や動詞といった文法カテゴリを明に教わることなく、入力言語から母語に適した文法カテゴリとそれらの間の遷移規則を獲得していく。本論文では、その幼児の文法カテゴリ獲得の計算論モデルを提案し、このモデルが幼児の言語理解と産出における発達的变化を再現できることをシミュレーション実験により示す。そして、そのときのモデルの内部表現を解析することにより、再現した発達的变化に潜在する文法カテゴリ表現を明らかにする。提案モデルは、隠れマルコフモデル (hidden Markov model: HMM) を用いて文法カテゴリを獲得する。HMM は入力された語列に対し、隠れ状態として文法カテゴリを推定する。ここで、幼児の文法カテゴリ発達段階をこの HMM の隠れ状態の数でモデル化する。隠れ状態が少ないと、曖昧な文法カテゴリが獲得され、隠れ状態が十分に多いと、精緻な文法カテゴリが獲得される。このモデルが幼児の行動的実験結果を再現できるかどうか、言語理解と産出の二つの課題で検証する。

言語理解に関して、幼児は語意学習のために文法カテゴリを利用し、この語意学習には三歳から五歳にかけての発達的变化と言語間差異が存在することが報告されている。この現象の計算論的メカニズムを理解するために、HMM に基づく語意学習モデルを提案する。HMM の隠れ状態と観察した視覚特徴とを連合学習することで、モデルは新奇語の指す特徴を、その隠れ状態（文法カテゴリ）を介して推定できるようになる。従来の行動実験と同様に、英語、日本語、および、中国語の三つの言語環境で提案モデルをシミュレートした結果、各言語条件において、モデルと幼児の成績が類似する傾向にあることがわかった。HMM の隠れ状態が多くなるにつれて、精緻な文法カテゴリが獲得され、新奇語の指す特徴を推定可能になる。そして、次のような言語固有の文法カテゴリが獲得されることによって、語意推定に幼児と同様の言語間差異が現れることがわかった。日本語は主語や目的語の欠落が多いため、日本語を学習

した HMM は形態的な手がかりを用いて文法カテゴリを形成していた。一方、英語と中国語の言語入力と比較的安定な語順を有しているため、それらを学習した HMM は統語的な手がかりを用いていた。

言語産出に関して、幼児の発話する多語文には、体系的な誤り（過剰生成）が含まれる。英語を母語とする幼児では不規則変化動詞に形態素「ed」を付けたり、日本語を母語とする幼児では形容詞の後ろに格助詞「の」を配置したりする過剰生成が報告されている。本研究では、文法カテゴリの不十分な増加がこのような過剰生成を引き起こし、それはさらに文法カテゴリが増加することによって解消されるという仮説に基づいて、その計算モデルを提案する。すなわち、規則変化動詞と不規則変化動詞、あるいは、名詞と形容詞が混在する文法カテゴリが獲得されることによって、それぞれ、「規則変化動詞 + ed」の規則が不規則変化動詞へも、あるいは、「名詞 + の」の規則が形容詞へも適用される。言語理解モデルと同様の HMM に基づいたモデルのシミュレーションにより、上記の仮説を部分的に支持する結果が得られた。英語の場合、カテゴリの増加に従って、過剰生成が出現し、減少する傾向が観察された。しかし、多くのカテゴリを持つモデルであっても規則変化動詞と不規則変化動詞を分けられなかったため、過剰生成が完全に消失することはなかった。日本語の場合、名詞と形容詞の分化に従って、過剰生成の出現と消失がみられた。ただし、「の」の過剰生成は初期値依存性が大きく、過剰生成が生じないモデルには、すでに名詞と形容詞が分化している場合と、「の」を含む助詞が獲得されていない場合があることを明らかにした。

謝辞

本研究を進めるにあたり、多くの方々のご支援を賜りました。ここに感謝の意を示します。まず、指導教員である大阪大学大学院 工学研究科 知能・機能創成工学専攻 浅田稔 教授に深く感謝いたします。浅田先生には学部生の頃からご指導、ご鞭撻を賜りました。また、筆者が研究に悩み、辛い日々を送っていたときでも、浅田先生は決して見放すことなく、多大な便宜を図ってくださいました。浅田先生のエネルギーで意欲的な研究態度は、筆者の憧れであり、目標です。

ご多忙の中、本論文の副査をしてくださった、大阪大学大学院 基礎工学研究科 システム創成専攻 石黒浩 教授、および、同専攻 細田耕 教授に感謝いたします。両先生には厳しいですが励みになるご質疑とご助言を賜りました。先生方のご期待に添える研究をやり遂げられるように精進いたします。

知能・機能創成工学専攻 長井志江 特任准教授、石原尚 助教、および、セルギーポントワーズ大学 森裕紀 先生からは、日頃より本研究に関して議論していただき、的確なご助言をいただきました。特に、長井先生にはモデル研究の面白さ、研究生活態度、および、研究のプレゼンテーション方法など、研究者としての基本事項を教えてくださいました。深く感謝いたします。

富士通研究所 笹本勇輝 氏（当時、大阪大学大学院 工学研究科 知能・機能創成工学専攻 博士後期課程学生）からは、本研究の確率モデルの立式に関して多大なご助言をいただきました。また、笹本氏から、工学的・数学的な困難に立ち向かう学びの態度を教わりました。笹本氏の「数式から逃げるな」の叱咤激励のお言葉は今でも筆者のスローガンの一つになっています。本論文はNTT イノベーションセンタ（当時、同上専攻 博士前期課程学生）大嶋悠司 氏との共同の研究成果です。本論文の一部は大嶋氏の修士論文に基づいています。大嶋氏との研究議論は非常に面白い、かつ、建

設的で、彼なしでは本研究の成果を挙げられませんでした。両氏に心から感謝いたします。

大嶋氏の修士論文にも関連して、慶應義塾大学 環境情報学部 今井むつみ 教授、および、日本学術振興会 安西祐一郎 先生から研究のご助言と激励のお言葉を賜りました。今井先生には言語発達研究者が期待するモデル研究像を、安西先生には認知科学におけるモデル研究像をお教えいただきました。2014年9月の日本認知科学会サマースクールにご招待いただき、両先生から貴重なご進言をいただけたことは、筆者の研究の大きなモチベーションになっています。深く感謝いたします。

大阪大学 大学院工学研究科 知能・機能創成工学専攻 浅田研究室の学生各位、特に、同級生である堀井隆斗 氏と朴志勲 氏に感謝いたします。お互いに助け合い、研究意欲を刺激し合うことで、本論文の成果に至ることができました。今後とも良き仲間・同僚として共に研究領域を牽引できることを願います。

最後に、「好きなことを好きなだけやりなさい」と筆者の研究生生活を常に応援してくださった両親に深く感謝いたします。両親の自由、誠実、かつ、他利的な精神を筆者も引き継いでいきます。

2017年 1月 10日

河合 祐司

目次

第1章	序論	1
1.1	背景	1
1.1.1	言語理解における文法カテゴリ	2
1.1.2	言語産出における文法カテゴリ	3
1.2	本論文の目的	5
1.3	論文の構成	7
第2章	従来研究と本論文のキーアイデア	9
2.1	文法カテゴリ化の手がかり	9
2.2	文法カテゴリを用いた語意の推定	11
2.3	英語の「ed」の過剰生成モデル	12
2.4	本論文のキーアイデア	13
第3章	文法カテゴリに基づく語意推定モデル	15
3.1	はじめに	15
3.2	モデル	16
3.3	実験設定	19
3.3.1	文・特徴の対	19
3.3.2	評価	22
3.4	実験結果	23
3.4.1	新奇語の意味の推定	23
3.4.2	文法カテゴリの表現	25
3.4.3	文法カテゴリの獲得過程	28

3.5	議論	30
3.5.1	他動詞と自動詞の統語ブートストラッピング	33
3.5.2	名詞優位性の統語論的メカニズム	34
第4章	未分化な文法カテゴリによる過剰生成モデル	37
4.1	はじめに	37
4.2	モデル概要	39
4.3	実験設定	40
4.3.1	学習コーパス	41
4.3.2	コーパス条件	42
4.4	実験結果	44
4.4.1	英語での過剰生成	44
4.4.2	日本語での過剰生成	49
4.5	議論	53
4.5.1	英語と日本語の過剰生成に共通するメカニズム	54
4.5.2	英語での過剰生成	56
4.5.3	日本語での過剰生成	56
第5章	討論	59
5.1	HMMの妥当性	59
5.2	文法カテゴリ発達のトリガ	60
5.3	意味カテゴリの利用と獲得	61
5.4	言語理解と産出の相互関係	62
第6章	結論	63
	関連図書	65
付録A	付録	79
A.1	行動データの引用	79

A.2	語意推定モデルの学習法	80
A.2.1	文法カテゴリの獲得	80
A.2.2	語と特徴の連合学習	82
A.3	学習コーパス	84
A.3.1	語意推定における学習コーパス	84
A.3.2	過剰生成における学習コーパスの語彙	87

表 目 次

A.1	Vocabulary sizes of word categories in each corpus.	86
A.2	Vocabulary sizes of word categories in the English corpus.	88
A.3	Vocabulary sizes of word categories in the Japanese corpus.	88

目 次

1.1	An outline figure of the proposed model for language comprehension and production in young children from two to five years old. In the early stage of development, words are not categorized into grammatical categories. Grammatical categories begin to be acquired at around three years old, but they are not clear enough to comprehend a category of novel verbs. The undifferentiated grammatical categories also cause overproduction in utterances that children produce. Acquiring accurate grammatical categories enables children to identify a category of novel verbs and to fix the overproduction.	6
3.1	A graphical model for inferring word meanings via grammatical categories. This indicates the inference process of a target feature o designated by a word w_t through the grammatical category s_t and meaning category m . The dashed arrow represents a direct induction from a word to a feature, but the model cannot learn the relation for a novel word. The grammatical category s_t is a clue for inferring the feature o from the observed features $\mathbf{f}$. The gray circles indicate that the variables are observable in the environment.	18

3.2	Inference of word meanings depending on the number of grammatical categories. (a) The model acquires obscure grammatical categories because the number of grammatical categories S is smaller, resulting in the failure to infer a target feature designated by a novel word. (b) A sufficiently large S enables the model to acquire well-differentiated grammatical categories. The larger S improves the accuracy with which target features are inferred.	20
3.3	Estimation of the target feature corresponding to a novel word by the model (bars) and children (points). Results for the model trained with the (a) English, (b) Japanese, and (c) Chinese corpus. The bars indicate the average percentage of model estimates in which the inferred feature was a novel action. The white and black bars denote the results for the models with fewer hidden states (three-year-old case) and more hidden states (five-year-old case), respectively. The stars above the bars denote significant differences (*: $p < 0.05$, **: $p < 0.01$) between the model estimation and the chance level (0.5), shown as a dashed line. The error bars indicate the standard error. The black and white points denote the results achieved by three- and five-year-old children, respectively, as reported by Imai et al. (2008).	24
3.4	Part-of-speech representation for each hidden state of the HMM that was learned the English corpus. The graphs on the left indicate a typical representation of the three-year-old cases, whereas the graphs on the right indicate a typical representation of the five-year-old cases.	26
3.5	Part-of-speech representation for each hidden state of the HMM that was learned the Japanese corpus.	26
3.6	Part-of-speech representation for each hidden state of the HMM that was learned the Chinese corpus.	27

3.7	Different part-of-speech representation of the hidden states with the number of hidden states varying from two to six in the English case. .	29
3.8	Different part-of-speech representation of the hidden states with the number of hidden states varying from two to six in the Japanese case.	29
3.9	Different part-of-speech representation of the hidden states with the number of hidden states varying from two to six in the Chinese case.	30
4.1	The proposed statistical model. The transition rule of words is a tri-gram for words, and the transition rule of grammatical categories is a trigram for grammatical categories. A probability for generation of the next word is given as the product of these rules.	39
4.2	Generation rules for artificial corpora.	43
4.3	Results of overproduction in English. Markers indicate rates of overproduction of “ed” for (a) irregular verbs and (b) words other than verbs. Error bars indicate standard deviation.	45
4.4	Averaged ratios of a probability of the transition rule of grammatical categories to a probability of the transition rule of words in English and Japanese. The probabilities were for “ed” when “ed” was produced in the English case, while they were for “no” when “no” was produced in the Japanese case. Error bars indicate standard deviation.	46
4.5	Typical representation of three grammatical categories acquired from the English corpus.	47
4.6	Typical representation of ten grammatical categories acquired from the English corpus.	48
4.7	Typical representation of nineteen grammatical categories acquired from the English corpus.	48
4.8	Results of overproduction in Japanese. Markers indicate rates of overproduction of “no” for (a) adjectives and (b) words other than nouns and adjectives. Error bars indicate standard deviation.	50

4.9	Typical representation of four grammatical categories acquired from the Japanese corpus, where the overproduction appeared (its rate was 0.41).	51
4.10	Typical representation of four grammatical categories acquired from the Japanese corpus, where the overproduction did not occur (its rate was 0.00074).	51
4.11	Typical representation of five grammatical categories acquired from the Japanese corpus, where the overproduction disappeared.	52
4.12	Rates of overproduction of “no” in several corpus conditions. “Regular” denotes the corpus used the experiment shown as Fig. 4.8. Error bars indicate standard error.	53
4.13	A typical representation of grammatical categories which was acquired from the Japanese corpus including more “no”.	54
A.1	A graphical model for the Bayesian hidden Markov model. W indicates the number of words in a corpus, i.e., the length of a corpus.	82
A.2	Transition rules between word categories for the production of the English corpus.	85
A.3	Transition rules between word categories for the production of the Japanese corpus.	85
A.4	Transition rules between word categories for the production of the Chinese corpus.	86

第1章 序論

1.1 背景

幼児の言語発達の理解は、人の言語処理メカニズムの解明へ向けた重要なアプローチとなるため、これまでに多くの観察実験やモデルが報告されてきた。人は生後およそ一歳後半から二歳頃にかけて多語文を発話し始め、その初期の多語文にも統語規則があることが報告されている [1]。その一方で、文中での語の役割（動詞など）を表す文法カテゴリは未発達であるとされる [2, 3]。特に、動詞のカテゴリは一般化されておらず、Tomasello [4, 5, 6] は、言語獲得初期の幼児の統語規則は項目依拠的に構成されており、特定の語間の遷移の規則であるという仮説を提案している。そして、三歳頃から動詞一般に適用可能な抽象的な動詞カテゴリを獲得し、それを含めた文法カテゴリとそれらの遷移規則、いわゆる、抽象スキーマを獲得する。

一方で、統語ブートストラッピング仮説は、幼い幼児であっても、意味的な表現と連合した動詞の抽象的な文法的知識を有していることを主張している [7, 8, 9, 10]。例えば、他動詞は通常、何かにはたらきかける行為を指す。このような意味と結びついた文法知識によって、幼児は新奇動詞の意味を推定できるようになる。例えば、Naigles [8] は英語を母語とする幼児は、動詞の下位カテゴリ（他動詞と自動詞）を新奇動詞に適用でき、他動詞の意味（因果的な行為）、あるいは、自動詞の意味（非因果的な行為）を推定できることを示している。

上記の対立を含め、幼児がどのように文法カテゴリを獲得し、それがどのように言語産出や言語理解（語意の推定）に反映されるかは今でなお議論の余地がある。本論では、計算論モデルの見地から、文法カテゴリの獲得メカニズムに新しい洞察をもたらすことを目指す。そのためには、言語理解と産出の双方の、また、複数の言語にお

ける発達的变化を再現できるモデルであることが望ましい。本章の以降では、幼児の文法カテゴリの獲得が言語理解・産出に及ぼした影響を示した研究を概説する。そして、それらの現象を説明できる計算論モデルを提案する。

1.1.1 言語理解における文法カテゴリ

Imai et al. [11] は、英語、日本語、または、中国語を母語とする幼児が、新奇語の指す対象が物体なのか、動作なのかをその文法カテゴリを用いて推定できるかどうかを調査した。ここでの文法カテゴリとは、名詞や動詞といったカテゴリであり、本論文でもこのような基本的な文法カテゴリを扱う。彼女らは、この語意推定において、三歳から五歳にかけての母語依存な発達的变化を発見した。まず、女性が新奇な物体を用いて新奇な動作をしている標準刺激を幼児に提示する。この映像刺激と同時に、英語を母語とする幼児に対しては、意味のない新奇語「*dax*」を含む文を大人が読み上げる。新奇語を提示する文の構成に三つの条件がある。名詞条件、動詞（項省略）条件、および、動詞（項明示）条件ではそれぞれ、「this is a *dax*」, 「*daxing*」 および、「she is *daxing* it」のように新奇語を与える。次に、物体のみが標準刺激と異なる刺激（物体変化刺激）と、動作のみが異なる刺激（動作変化刺激）を同時に幼児に提示する。そして、新奇語「*dax*」がどちらの刺激に対応するか幼児に尋ね、幼児はどちらかを選択する。名詞条件の場合は動作変化刺激、動詞条件の場合は物体変化刺激が正しい選択となり、それぞれ、名詞般用と動詞般用と呼ばれる。

その実験の結果、母語に関係なく、三歳児は名詞般用可能だが、動詞般用に失敗することがわかった。五歳児では、動詞般用も可能になるが、一部の条件で、動詞般用に失敗し、言語間の差異が現れる結果となった。日本語を母語とする五歳児は、両方の動詞条件で動詞般用に成功した。一方で、英語を母語とする五歳児は動詞（項明示）条件では動詞般用できるが、動詞（項省略）条件ではその推定がチャンスレベルになり動詞般用に失敗した。中国語を母語とする五歳児も同様に、項明示条件で般用可能だが、項が省略されると、新奇動詞を物体に対応付ける傾向にあることがわかった。

この実験は、幼児の語意推定がその幼児の母語の影響を受けることを示している。

Imai et al. [12, 11] の課題を達成するためには、幼児は新奇語の文法カテゴリを推定しなければならない。すなわち、新奇語が名詞カテゴリなのか、動詞カテゴリなのかを推定する必要がある。そして、名詞カテゴリと動詞カテゴリは一般的に、それぞれ物体と動作に対応付くという知識を用いることで、幼児は新奇語を刺激中の正しい特徴に対応付けられる。ここで重要なメカニズムは文法カテゴリの獲得である。三歳児の文法カテゴリ化の能力は未熟であり、その結果、新奇語の般用課題の成績の精度が低い。五歳児は母語に基づいて異なる文法カテゴリを獲得するため、その課題の成績に言語間差異が生じると考えられる。しかし、どのような言語的特性が言語固有の文法カテゴリの獲得を導き、この言語間差異を生じさせたのかという問いが残る。行動学的・観察的手法では、この情報処理に潜在するメカニズムを直接明らかにすることは困難である。

1.1.2 言語産出における文法カテゴリ

幼児の文法発達と言語産出の側面からも研究されている。ただし、発話された言語が、幼児の文法知識に基づく生産的な発話であるのか、単に聞いた発話を繰り返しているだけなのかを区別することは困難である。そのため、幼児が聞くはずのない文の生成、すなわち、幼児の産出文の誤用に注目した研究が多く報告されている。このような誤用の一部には規則性があるため、このような誤りは単なる間違いではなく、幼児の言語規則の獲得過程を反映しているとされる。したがって、幼児の誤用に注目することで、幼児の言語獲得メカニズムを明らかにできる可能性がある。

典型的な幼児の誤用は、英語における過去形の規則変化の過剰生成（または、過剰一般化）である [13, 14, 15]。通常、英語の規則変化動詞を過去形にするときには、形態素「ed」を付加する。しかし、およそ三歳の幼児は、例えば、「I buyed an apple」のように、この規則を不規則変化動詞に誤って適用することがある。英語の「ed」の過剰生成について、以前は正しく過去形にできた不規則変化動詞でも、三歳以降に誤用が生じることがある。すなわち、U字発達がみられる。また、この誤用が九歳頃まで続いた例もあり、これは長期間に及ぶ現象といえる [16, 17]。

日本語を母語とする幼児は、およそ一歳後半で、形容詞による連体修飾ができるようになる。その一、三ヶ月後に、格助詞の誤用がみられ [18, 19], 特に「の」の過剰生成が多く報告されている [18, 20, 21, 22, 23]。例えば「ちいさい の わんわん。」のように、二歳頃の幼児は「の」を名詞だけでなく、形容詞の後にも配置することがある。2歳後半になると、この過剰生成は消失する [21, 23]。

過剰生成された「の」の素性に関して、それが格助詞であるという仮説 [18, 21, 22] だけでなく、形容詞や動詞などの体言化のための準体助詞（「綺麗な の が欲しい」の「の」）[24] や補足節のための補文標識（「彼が帰ってくる の を待つ」の「の」）[20] であるという仮説がある。しかし、伊藤 [22] の報告によると、「形容詞 + の + 名詞」の出現に先立って（または、同時に）「名詞 + の + 名詞」が産出されるようになり、この時期には準体助詞や補文標識の「の」の産出は観察されなかった。よって、伊藤 [22] はこの「の」が格助詞であると結論づけている。したがって、本論文でも過剰に生成される「の」は格助詞であると仮定し、「名詞 + 格助詞」の文法規則が形容詞へと過剰に適用されたとみなす。この過剰生成のメカニズムについて、伊藤 [22] は、名詞と形容詞の統語素性の理解が不十分である結果、「名詞 + 格助詞」の規則を形容詞にまで適用しているという仮説を提案している。

しかし、「幼児の助詞ノの誤用は Ervin [14] が規則動詞の過去時制形態素の過度の一般化について述べているほど絶対的なものではなく」[21]、また、形態論と統語論の違いもあり、英語の過去形の過剰生成との共通メカニズムについては今までほとんど議論されていない。ただし、Words and Rules 理論におけるグラマーの処理 [25, 26] のように、英語の過去形の語形変化規則を動詞語基と接尾辞の並びの規則として捉えれば、英語での「ed」の過剰生成と日本語での格助詞の過剰生成との類似性を見出すことができる。

文法カテゴリの発達の観点から、英語の「ed」の過剰生成と日本語の「の」の過剰生成に共通したメカニズムが考えられる。規則変化動詞と不規則変化動詞、あるいは、名詞と形容詞が未分化な文法カテゴリであると、それらを区別しない規則が習得され、過剰生成が生じる。そして、精緻な文法カテゴリが獲得されると、それらの誤用は解消される。しかしながら、全く同じモデルで英語と日本語の過剰生成を再現できるか

は不明である。また、なぜ、英語の「ed」の過剰生成に比べて、日本語の「の」の過剰生成の頻度が少なく、個人差が大きいのだろうか。

1.2 本論文の目的

本研究は幼児の文法発達メカニズムを明らかにするためのモデルの提案が最終目標であるが、その第一段階として、本論文では、上記の言語理解と産出における発達的变化を現象論的に説明できる計算論モデルを提案する。そして、計算機シミュレーションにより、上記の現象を再現し、その内部構造を解析することによって、それらの現象の背後にある文法カテゴリを理解することを目的とする。Fig. 1.1 に本研究の概要図を示す。今回、語遷移規則（下側のブロック）と文法カテゴリ遷移規則（上側のブロック）の二つの確率的規則が言語入力から獲得されると仮定する。これらの規則はそれぞれ、Tomasello のモデルでは、項目依拠的規則と抽象スキーマに対応する。そして、上記の発達的变化は文法カテゴリとその遷移規則の精緻化によってもたらされるという仮説を提案する。なお、今回、二歳ごろから精緻な語遷移規則が獲得されていると仮定し、その発達的变化は考慮しない。二歳頃では、有意な文法カテゴリ遷移規則が獲得されておらず、語の遷移規則が優位にはたらく。したがって、文法カテゴリに基づいた言語理解はできないが、語の遷移規則により、文法的に正しい文を産出できる。三歳頃から入力言語に基づいた文法カテゴリが獲得され始めるが、その文法カテゴリは未分化であるとする。例えば、名詞カテゴリとそれ以外の語カテゴリといった曖昧な文法カテゴリとその遷移規則が獲得される。この場合、言語理解の課題では、明瞭な名詞カテゴリによって、未知名詞のカテゴリ分類とそれに基づく名詞般用は可能だが、動詞カテゴリは曖昧であるため、未知動詞の般用は困難である。また、このような曖昧な文法カテゴリの獲得によって、言語産出における過剰生成が生じる。例えば、規則変化動詞と不規則変化動詞を混同したカテゴリが形成されると、これらの動詞を区別せずに過去形として「ed」を付加する規則が獲得される可能性がある。そして、五歳頃に、入力言語に適した手がかりを用いた精緻な文法カテゴリとその遷移規則が獲得されることで、その言語固有な動詞般用を示すようになる。例えば、英

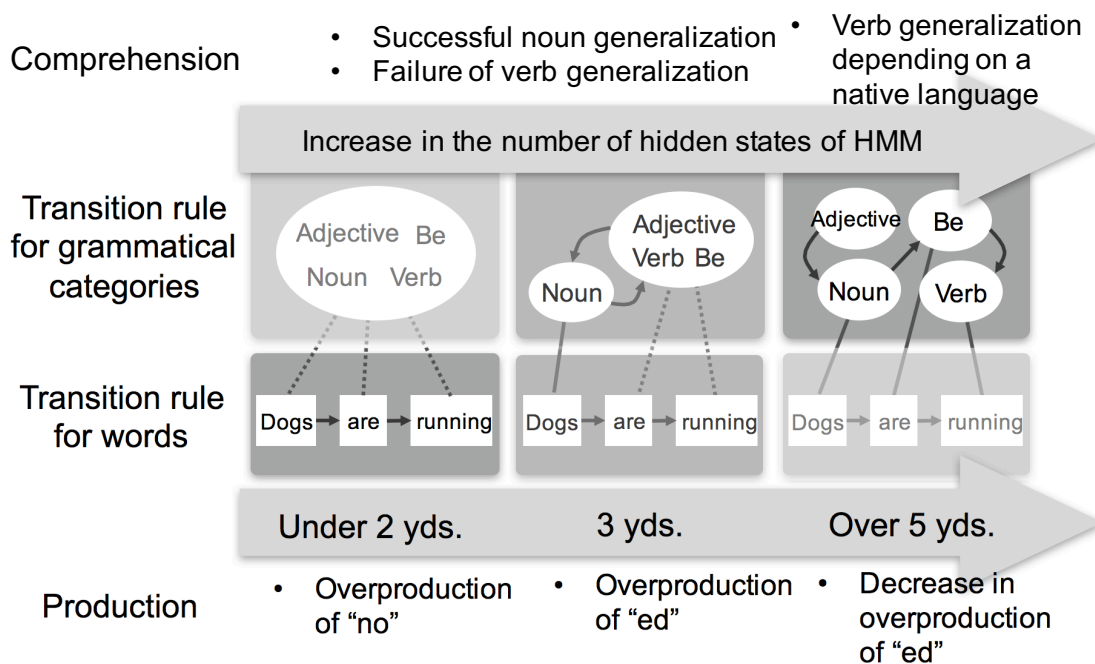


Figure 1.1: An outline figure of the proposed model for language comprehension and production in young children from two to five years old. In the early stage of development, words are not categorized into grammatical categories. Grammatical categories begin to be acquired at around three years old, but they are not clear enough to comprehend a category of novel verbs. The undifferentiated grammatical categories also cause overproduction in utterances that children produce. Acquiring accurate grammatical categories enables children to identify a category of novel verbs and to fix the overproduction.

語の場合，比較的強い統語的手がかりに基づいた文法カテゴリが獲得され，その統語的手がかりを有する条件でのみ，動詞般用が可能になる．また，文法カテゴリの曖昧性が解消されることで，過剰生成が消失する．

提案モデルでは，文法カテゴリを学習するために，教師なしの機械学習手法である，隠れマルコフモデル (hidden Markov model: HMM) を用いる [27]．語列が与えられたとき，HMM はその語の文法カテゴリを隠れ状態として推定し，そのカテゴリ間の遷移確率を計算する．提案モデルによる文法カテゴリ獲得は語列の統語的・形態的統

計規則のみに基づき、知覚や意味はそれに関与しないとする。言語理解と産出における発達的变化を説明する重要なパラメータが、文法カテゴリ数、すなわち、HMMの隠れ状態数である。多くの隠れ状態を持つモデルは、より精緻な文法カテゴリ遷移規則を獲得でき、隠れ状態数が幼児の発達の目安となる。上記の仮説を検証するために、言語理解の実験では、英語、日本語、および、中国語の三つの言語環境で提案モデルをシミュレートし、Imai et al. [11] の実験を再現することを試みる。そして、獲得された文法カテゴリの内容を解析することで、般用成績の言語間差異を引き起こした文法カテゴリを明らかにする。言語産出の実験では、英語と日本語の二つの言語環境で提案モデルをシミュレートし、「ed」と「の」の過剰生成の再現を試みる。そして、そのときの文法カテゴリの解析により、過剰生成の発生メカニズムを議論する。さらに、学習コーパスを操作することで、過剰生成に寄与する入力要因を明らかにし、提案モデルの有効性を示す。ただし、提案モデルでは、HMMの隠れ状態数はあらかじめ設定される固定値であり、シミュレーション中に連続的に変化することはない。各隠れ状態数に対してそれぞれシミュレーションを実行し、そのときのモデルのパフォーマンスと内部表現の変化を議論する。

1.3 論文の構成

本論文は六つの章で構成される。第1章では、背景とモデル化の対象、および、本論文の目的を述べた。第2章では、従来研究を紹介し、本研究の新規性とアイデアを再度述べる。そして、第3章と第4章で、文法カテゴリの獲得モデルを用いて言語理解と産出を説明する実験を行う。第3章ではImai et al. [11] の実験結果を再現するシミュレーションを行い、その成績の言語間差異と発達的变化の背後にある文法カテゴリ構造を明らかにする。第4章では、前章と同じHMMモデルを用いて、「ed」と「の」の過剰生成が現れるかを検証し、その発生要因について検討する。各実験についての考察は各章で述べる。以上の実験結果を踏まえて、第5章で提案モデルの限界と今後の展望について述べる。最後に、第6章で本研究をまとめる。

第2章 従来研究と本論文のキーアイデア

2.1 文法カテゴリ化の手がかり

幼児の文法カテゴリの獲得を扱ったモデル研究や解析的な研究は多くある。それらは幼児が文法カテゴリを推定するために、統語的（分布的）手がかり、形態的手がかり、および、音韻的手がかりの三つを利用していることを示している。

文法カテゴリ化に関するほとんどの計算論モデルは、語の順序パターンである統語的、あるいは、分布的情報を用いている（例えば、[28, 29, 30, 31, 32, 33, 34, 35, 36]）。例えば、「they use balls」と「they throw balls」という文が与えられたとすると、「use」と「throw」に近接した語は同じであることがわかる。同じ文法カテゴリに属す語どうしは、その近接した語を共有する傾向にあり、これが語を分類する手がかりとなる。また、この分類の精度は語順の規則だけではなく、文法カテゴリ間の規則（例えば、他動詞カテゴリは主語カテゴリの後ろ、目的語カテゴリの前）でも向上させられる[29, 33]。すなわち、語の文法カテゴリは、近接した語と近接した文法カテゴリの両方の規則から決められる。このような統語規則に基づくモデルでは、カテゴリ化のための教師なし学習法が用いられている。Elman [28] は英語の単純な三語文を用いて、単純リカレントニューラルネットワーク (simple recurrent neural network: SRN) が文法カテゴリを表現できることを示した。SRN は現在の語から次の語を予測するように学習し、その結果、そのSRNの文脈層ニューロンで分布的表現を獲得した。そして、そのニューロンの表現をクラスタリングすることによって、文法カテゴリが得られる。したがって、そのSRNの出力は入力だけでなく過去の文法カテゴリに基づいて決まる。今回用いるHMMも次の語を文法カテゴリによって予測するため、同様の

方法で文を扱っているといえる。しかしながら、HMMのような確率的手法は信号を確率変数とその確率分布としてシンボル化し、一般的なニューラルモデルよりもそのシンボルの処理を明に定義する。そのため、確率的手法の方が計算的に扱いやすく、また、原理的な情報処理の理解や内的表現（今回では、文法カテゴリ）の視覚化と解釈が容易である。

対幼児発話を解析した近年の研究では、文法カテゴリ獲得における形態的・音韻的手がかりの重要性が強調されている。Onnis and Christiansen [37] は多くの言語において、対幼児発話には文法カテゴリの分類に寄与できる形態的情報が多くふくまれていることを報告した。例えば、英語の動詞はよく接尾辞「ing」や「ed」を伴う。また、実験的な研究は英語を母語とする幼い幼児でも動詞を拘束形態素と自由形態素に分けられることを示している [38, 39, 40]。したがって、幼児はこのような形態の手がかりを用いて文法カテゴリを推定している可能性がある。さらに、いくつかの音韻的要素が文法カテゴリの識別に有用であることが示されている [41]。例えば、英語において、動詞の音節の平均長よりも名詞の方が大きい。また、モデル研究からも、統語的な手がかりだけでなく、形態的・音韻的な手がかり（例えば、音節の数や語尾の形態素）に基づいて文法カテゴリを学習すると、その分類性能が向上することが示されている [42]。

まとめると、従来の計算論的研究は、文法カテゴリを推定するいくつかのアルゴリズムを提案してきた。そして、これらの研究は統語的な手がかりだけでなく、語内部にある手がかりを用いて、幼児は語を分類している可能性を示している。しかし、ほとんどの既存研究はそれらの手がかりによる分類可能性を示したのみであり、幼児の言語発達の限定的な側面しか検討されていない。計算モデルの妥当性を検証するためには、言語間の差異といった言語発達の様々な側面を再現する必要がある。特に、Imai et al. [11] の実験は文法カテゴリに基づく語の般用における発達的变化と言語間差異を示しており、この実験を再現することは提案モデルの妥当性を強く支持する。例えば、英語を母語とする五歳児は、単に「*daxing*」とだけ聞くとその動詞を動作に対応付けることができないが、その項を省略しない場合では、その対応付けが可能になる。これは、英語を母語とする幼児は必ずしも形態の手がかりを用いて動詞カテゴ

りを推定しないことを示している。一方で、日本語を母語とする五歳児が「*dazing*」と同様の「ネケってる」と聞いた場合では、形態的手がかりにより、その動詞語意の推定ができる。上記のような幼児のできないこと、および、矛盾を含めて説明できるモデルとして妥当性を検証することが望ましい。さらに、上記の言語理解だけでなく、言語産出における発達の現象を再現できることもそのモデルの妥当性を強く保証する。

2.2 文法カテゴリを用いた語意の推定

語意の学習は、語 w が与えられたときにシーン中にある特徴 f の確率 $P(f|w)$ の計算としてモデル化される。通状況的学習と呼ばれる語意学習法では、多くの文・シーンの対からこの確率を計算する（例えば、[43, 44, 45, 46]）。

本研究における語意推定の基本的なアイデアは、 w だけでなく、その語の文法カテゴリ s を用いて f を推定すること（ $P(f|w, s)$ ）である。Imai et al. [12, 11] の実験のように、たとえ w が未知語であっても、 s がその語と適切な特徴との対応付けを可能にする。Alishahi and Fazly [47] は s を持たない文法カテゴリ獲得モデルと比べて、 s を含めたモデルを用いることで、語意学習の精度と速度が向上することを示した。しかし、その研究では、 s が完全に正しく、あらかじめ与えられるものと仮定されていた。したがって、そのモデルでは s がどのように獲得され（発達し）、 s の言語依存性がどのように $P(f|w, s)$ に影響するのかを説明できない。一方で、コーパスに基づく教師なし機械学習により s を獲得するモデルもある [48, 49, 50, 36, 51]。Toyomura and Omori [49] は s を獲得するために SRN を用い、その文脈層ニューロンの文法表現を観察した特徴と連合させた。そして、彼らのモデルはその文法表現を介して、語の指す標的物体を推定した。しかし、彼らは名詞と動詞のカテゴリ化についての言語間差異や発達的变化を考慮していない。Alishahi and Chrupala [51] は統計的な教師なしモデルを用いて s を学習することで、教師ありの文法カテゴリ化 [47] と同様の学習効率となることを示した。

上記のモデルのほとんどは、そのモデルの学習効率を評価したに過ぎず、言語依存

性や発達的变化を議論していない。これらの研究でも、前節の文法カテゴリ化と同様の問題がある。システムの効率はずしも認知メカニズムとしての妥当性を保証しない。

2.3 英語の「ed」の過剰生成モデル

英語の「ed」の過剰生成を説明する計算論モデルには、コネクショニストモデルがある（例えば、[52, 53, 54]）。これは一つの語形変化モデルで過剰生成の説明を試みている。ニューラルネットワークの入力を動詞原形の音素とし、その出力を入力動詞の過去形の音素とする。そして、ニューラルネットワークは入出力間の語形変化規則を学習する。学習初期では、全ての動詞に「ed」の変化規則が適用されるため、過剰生成が生じるが、学習後期では、規則変化と不規則変化の両方が学習されるため、過剰生成が消失する。しかし、これらのモデルは一語の屈折の規則の学習であり、これらを明らかに統語規則の誤りである日本語の過剰生成と直接的に関連づけることはできない。

対照的に、「ed」の過剰生成を説明する Words and Rules 理論では、語形変化は二つの処理、すなわち、グラマーとレキシコンの相互作用・競合によってなされるとされる [25, 26]。規則変化の規則は、グラマーの処理により「動詞語基 + ed」の手続き記憶として習得される。不規則変化の場合、レキシコンの処理がグラマーの手続きを抑制し、「ed」を付加することなく、レキシコンの処理により記憶していた別の語に変化する。この抑制がはたらかないと、グラマーの手続きが不規則変化動詞にまで適用され、「ed」の過剰生成が生じる。

一方、日本語の「の」の過剰生成を説明する計算論モデルは未だ提案されておらず、また、それと英語の「ed」の過剰生成との関連もほとんど議論されていない。

2.4 本論文のキーアイデア

本論文では、幼児の言語理解・産出の両方の発達的变化，すなわち，新奇名詞・動詞の般用課題成績の母語依存性の出現 [11]，および，過剰生成の出現と消失 [13, 18] を再現することによって，提案モデルの妥当性を評価する．提案モデルでは，形態的・統語的手がかりに基づいて文法カテゴリを推定するために，HMMを用いる．この教師なしの HMM は語入力 of 隠れ状態として文法カテゴリを学習する．言語理解のシミュレーションでは，この隠れ状態と観察した特徴との対応関係を通状況的に計算する．提案モデルが学習を通して，入力言語の文法カテゴリ化のための適切な手がかりを選択し，その結果，動詞般用の言語間差異が生じることが期待される．また，HMM の表現容量の増大，すなわち，隠れ状態数の増加が，幼児の言語発達の一側面を説明できるという仮説を提案する．言語理解については，名詞・動詞の般用の言語間差異と発達的变化 [11] をそれぞれ，学習する言語と隠れ状態数の増加により説明できることをシミュレーション実験により示す．言語産出について，英語の「ed」と日本語の「の」の過剰生成が，HMM の隠れ状態数の増加により出現・消失することを示す．さらに，それぞれの場合で，獲得した隠れ状態を解析することによって，文法カテゴリがどのように言語入力 of 特性を表現し，言語理解や産出に影響を及ぼしたのかを明らかにする．

第3章 文法カテゴリに基づく語意推定モデル

3.1 はじめに

本章では、言語理解に関する課題として、1.1.1 項で述べた Imai et al. [11] の名詞・動詞般用の計算論モデルを提案する。この課題では、幼児は未知語の文法カテゴリを推定することによって、その文法カテゴリを介して、語に対応する視覚特徴を推定しなければならない。この行動実験は、幼児の有する文法カテゴリを明らかにするだけでなく、その文法カテゴリが幼児の語意学習を促進することを示している。

語意学習は新奇語と観察した視覚特徴の対応関係の獲得とみなせられる。しかし、ここで、その新奇語の指す特徴に不確かさがあることが問題になる。複数の視覚特徴（動作主、物体、動作、属性など）のあるシーンを観察しているときに新奇語が与えられると、幼児はその新奇語が何を指しているのかを推測しなければならない。この問題は Quine [55] によって提議されて以来、幼児の語意獲得の中心的な研究対象となっている。その一つの解法は通状況的学習である [43, 44]。これは、多くの学習経験を通した、語と視覚特徴の間の共起頻度の統計規則の獲得である。一方で、幼児は一回の学習経験であっても、何らかのバイアスに基づいて語意を推定していることがわかっており、これは即時マッピングと呼ばれる（例えば、[56]）。その手がかりの一つに相互排他性バイアスがあり、これは新奇語を新奇特徴に対応させる傾向をいう [57]。また、幼児は新奇語を対象全体の名前であると解釈する傾向にあり、これは事物全体バイアスと呼ばれる [58]。最近の研究では、音象徴と呼ばれる語の音響的特徴と視覚特徴のアモダルな共通性も幼児の語意推定の手がかりになるが示されている [59]。これらのバイアスや手がかりに加えて、Imai et al. [11] の実験結果は、文法カ

テゴリに基づいて幼児が語意推定していることを示している。したがって、この課題をモデル化することは、幼児の文法カテゴリ獲得だけでなく、語意推定メカニズムの解明としても意義深い。

提案モデルでは、言語入力から HMM によりその文法カテゴリを学習し、その文法カテゴリと視覚特徴の対応関係を通状況的に学習する。これらの学習の結果、モデルは文法カテゴリを介して、新奇語の指す視覚特徴を推定できるようになる。今回、Imai et al. [11] の幼児の実験と同様に、英語、日本語、および、中国語の三言語で実験し、その語意推定の言語間差異を比較する。また、HMM の隠れ変数（文法カテゴリ）の数の少ない場合と多い場合とで課題成績を比較し、三歳から五歳の差異をもたらす文法カテゴリについて議論する。

3.2 モデル

Fig. 3.1 に、語系列 $\mathbf{w} = (w_1, w_2, \dots, w_t, \dots)$ と視覚特徴 $\mathbf{f} = (f_1, f_2, \dots)$ が与えられたときに、語 w_t の指す特徴 o の確率を推定するグラフィカルモデルを示す。矢印は円内の変数間の条件付き依存性を示す。今回、 $\mathbf{w}$ は動詞語基と分離した動詞語尾（接尾辞）を含むと仮定する。また、 $\mathbf{f}$ は名詞、動詞、および、形容詞に対応する概念的な視覚特徴の集合とする。本モデルの目的は $P(o|w_t, \mathbf{f})$ を適切に計算することである。モデルは与えられたシーン中の候補特徴 $\mathbf{f}$ から、与えられた語 w_t の指す標的特徴 o を選択する。この確率は $P(o|w_t)P(o|\mathbf{f})$ に比例し、ここで、 w_t と $\mathbf{f}$ は独立と仮定する。現在、観測しているシーンに存在する特徴に対して、 $P(o|\mathbf{f})$ は一様分布として与える。一方、多くのシーンと語の対から、モデルは $P(o|w_t)$ を学習する。しかし、モデルは新しく与えられた語 w^{novel} とそれが指す特徴 o の直接的な関係 $P(o|w_t = w^{\text{novel}})$ (Fig. 3.1 中の破線矢印) を獲得していないので、その語意を推定することはできない。したがって、 w_t の文法カテゴリ s_t を通して間接的に o を推定する。すなわち、 $P(o|w_t = w^{\text{novel}}) = \sum_{s_t} P(o|s_t)P(s_t|w_t = w^{\text{novel}}, \mathbf{s}_{-t})$ である。ここで、 $\mathbf{s} = (s_1, s_2, \dots, s_t, \dots)$ は $\mathbf{w}$ の文法カテゴリ系列であり、 $\mathbf{s}_{-t}$ は s_t 以外の $\mathbf{s}$ を意味する。提案モデルでは、文法カテゴリ $P(s_t|w_t = w^{\text{novel}}, \mathbf{s}_{-t})$ を HMM

の隠れ状態として推定する．この s の推定過程を Fig. 3.1 の破線で囲まれたグラフに図示する．今回，文法カテゴリ間の遷移に，トライグラムのマルコフ連鎖を仮定する．すなわち，現在の文法カテゴリの確率は，近接する四つのカテゴリに基づいて推定される： $P(s_t|w_t, \mathbf{s}_{-t}) \equiv P(s_t|w_t, s_{t-2}, s_{t-1}, s_{t+1}, s_{t+2})$ ．ただし，Fig. 3.1 には $P(s_t|w_t, s_{t-1}, s_{t+1})$ を示す矢印のみが示されている．HMM は文法カテゴリ間の遷移確率を学習し，これが簡単な統語規則に対応する．例えば，英語を学習したモデルでは， $P(s_t = \text{“verb stem”} | s_{t-1} = \text{“be-verb”}, s_{t+1} = \text{“verb suffix”}, \dots)$ が比較的大きな値を持つことが期待される．

そして，モデルは文法カテゴリ s_t から $P(o|s_t)$ を用いて o を推定する．しかし，新しく与えられた新奇特徴については，この確率的関係を学習する機会がない．なお，Imai et al. [11] の実験でも，幼児に提示される物体と動作は新奇なものであった．この問題を扱うために，物体や動作，属性のような抽象的な特徴カテゴリである意味カテゴリ m を導入する．今回，このカテゴリは特徴 o と曖昧性なく対応し，この対応関係は事前に与えられるものとする．すなわち， $P(o|m = m_i)$ の分布は，カテゴリ m_i に属す特徴に対して一様とする．また，モデルは新奇な o の m を認識可能であると仮定する．モデルは新奇特徴 o と s の直接的な写像を学習できないが， m と s の関係は学習できる．このようにして，新しく与えられた o は， s から m を介して間接的に推定される： $P(o|s_t) = \sum_m P(o|m)P(m|s_t)$ ．

以上から， w_t の指す o の確率は，直接経路 $P(o|w_t)$ と間接経路 $\sum_{s_t} \sum_m P(o|m)P(m|s_t)P(s_t|w_t, \mathbf{s}_{-t})$ の積で得る．

$$P(o|w_t, \mathbf{s}_{-t}, \mathbf{f}) \propto \sum_{s_t} \sum_m P(o|m)P(m|s_t)P(s_t|w_t, \mathbf{s}_{-t})P(o|w_t)P(o|\mathbf{f}), \quad (3.1)$$

ここで，右辺の $P(o|w_t)$ は語と特徴の直接写像を表す．もし， w_t が新奇語であれば，この確率は一様になる．さらに，新奇語を新奇特徴に対応付けしやすくするために，式 (3.1) の右辺を $\max_o(P(o|w_t))$ で除算する．この仮定は，語学習における相互排他性の制約 [57] に基づく．

提案モデルは式 (3.1) における二つの確率を学習しなければならない．一つは， $P(s_t|w_t, \mathbf{s}_{-t})$ であり，文法カテゴリ間の遷移規則を基に，新奇語の文法カテゴリを推

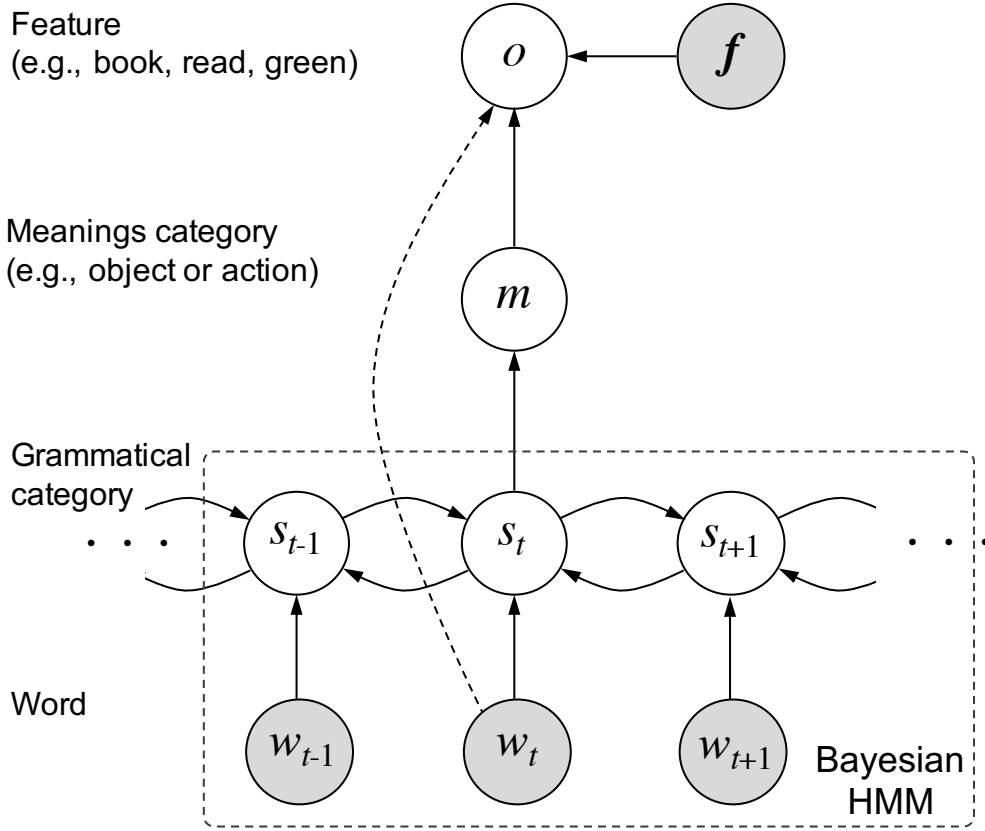


Figure 3.1: A graphical model for inferring word meanings via grammatical categories. This indicates the inference process of a target feature o designated by a word w_t through the grammatical category s_t and meaning category m . The dashed arrow represents a direct induction from a word to a feature, but the model cannot learn the relation for a novel word. The grammatical category s_t is a clue for inferring the feature o from the observed features $\mathbf{f}$. The gray circles indicate that the variables are observable in the environment.

定する確率である。もう一つは、 $P(m|s_t)$ であり、文法カテゴリと意味カテゴリの対応関係である。上述の通り、 $P(s_t|w_t, \mathbf{s}_{-t})$ を求めるためにHMMを用いる。このHMMでは、学習はコーパス $\mathbf{w}$ のみに基づき、 o や m は用いない（詳細は付録 A.2.1 を参照のこと）。そして、 $P(m|s_t)$ は通状況的に獲得する。多くのシーンとそれを説明する文が与えられたときに、各文・シーン対における文法カテゴリと意味カテゴリの共起

確率を計算する（詳細は付録 A.2.2 を参照のこと）。

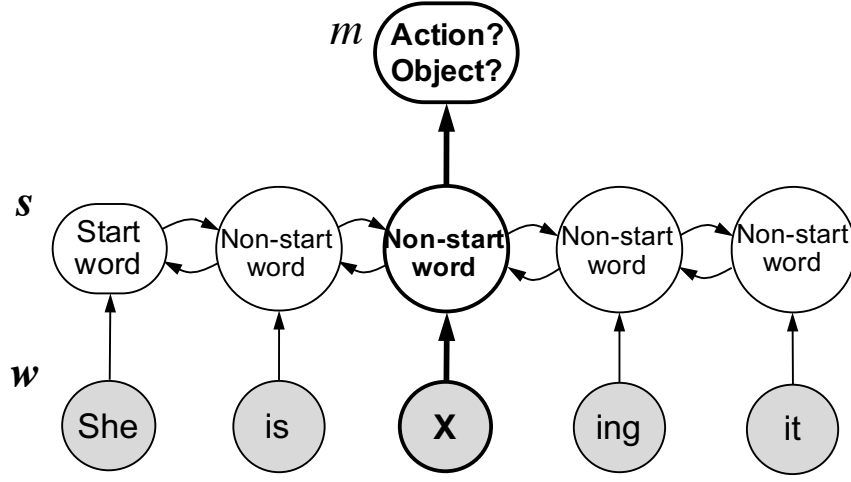
提案モデルにおいて、文法カテゴリの推定 ($P(s_t|\mathbf{w}, \mathbf{s}_{-t})$) が重要となる。モデルはこの確率をコーパスから学習する（学習に用いるコーパスの詳細は次節と付録 A.3.2 を参照のこと）。したがって、文法カテゴリとそれらの遷移規則は、与えられた言語の統語構造に依存する。また、HMM の隠れ状態数 S は各年齢のシミュレーションにおいて異なる定数として与える。この数は文法カテゴリの精度に重大な影響を与える。もし、 S が小さければ、例えば 2 であるとき、Fig. 3.2 (a) に示すように、文法カテゴリは文頭の語のカテゴリと、それ以外の語のカテゴリのみを表現する可能性がある。このような未分化な文法カテゴリでは、 o と m の確率の推定は不正確になる。一方で、モデルが十分に大きな S を有していれば、動詞カテゴリが名詞カテゴリから適切に分化するため、 m と o を正しく推定できる (Fig. 3.2 (b))。本研究では、 S に異なる値を設定することで、Imai et al. [11] が報告した、異なる年齢での新奇名詞・動詞の語意推定を再現できることを示す。大きな S と小さな S を持つモデルはそれぞれ、三歳と五歳の振る舞いを再現できることを以降の実験で示す。

3.3 実験設定

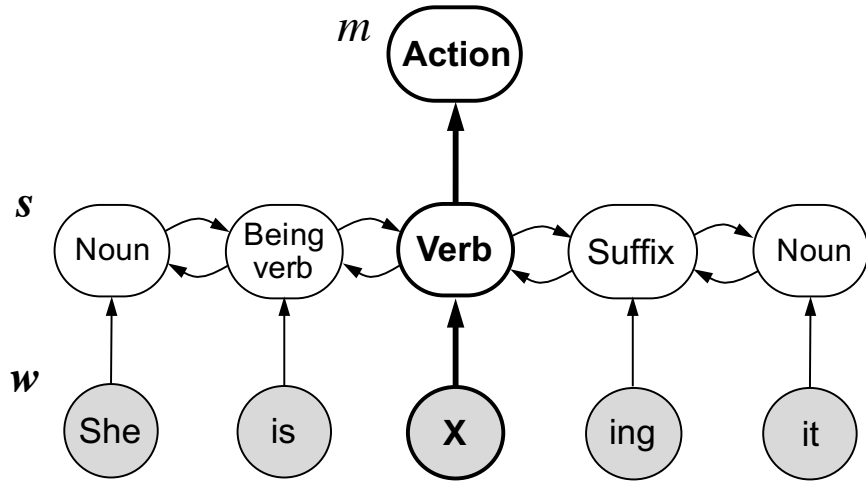
本実験の目的は、提案モデルが幼児の名詞・動詞般用 [11] を再現し、異なる年齢と母語での般用の差異の原因を探ることである。モデルを人工的な文・シーン対で学習させ、その後、新奇語が指す特徴を推定することによって、モデルを評価する。そして、そのモデルの隠れ状態の内的表現を解析することによって、その新奇語の語意推定に寄与した文法カテゴリを明らかにする。以降では、モデルの学習と評価の方法について述べる。

3.3.1 文・特徴の対

学習文として、日本語、英語、および、中国語の各言語構造を有する人工コーパスを作成する（詳細については付録 A.3.2）。今回、意味論的な影響を除いた文法カテゴ



(a) Less grammatical categories



(b) More grammatical categories

Figure 3.2: Inference of word meanings depending on the number of grammatical categories. (a) The model acquires obscure grammatical categories because the number of grammatical categories S is smaller, resulting in the failure to infer a target feature designated by a novel word. (b) A sufficiently large S enables the model to acquire well-differentiated grammatical categories. The larger S improves the accuracy with which target features are inferred.

りの獲得を検討するため、実コーパスではなく、人工コーパスを用いた。語カテゴリ（例えば、主語、目的語、動詞、または、形容詞のカテゴリ）間の確率的な遷移規則を各言語について設計し、その語カテゴリの持つ語彙リストから語を一つをランダムに選択することで文を作成した。なお、これらの人工コーパスにおいて、入れ子構造文はない。それぞれの言語のコーパスに含まれる文法構造とその割合は、CHILDES データベース [60] における二歳から五歳の幼児に対する大人の発話コーパス [61, 62, 63] やコーパス解析研究 [64] を基に決めた。

一般に、日本語は英語や中国語よりも主語や目的語が欠落することが多く、また、主語や目的語の語順が入れ替わることも多い。今回作成した日本語コーパスでも、主語や目的語の欠落や入れ替えが頻繁にある。また、英語では、「動詞-ing」単体で発話されることは極端に少なく、英語の人工コーパスでもそのような単体文の割合は 0.2% と非常に少ない。一方で、日本語コーパスでは、動詞単体文が比較的多い。

全てのコーパスにおいて、語はその語基と接尾辞に分離されていることを仮定した。例えば、英語コーパスでは、進行形の動詞は二つの形態素（語基と接尾辞「ing」）に分離される。この符号化法により、モデルは統語的な手がかりだけでなく、形態的な手がかりを用いて文法カテゴリを推定できるようになる。幼い幼児でも動詞を自由形態素と拘束形態素に分けられる [38, 39, 40] ことから、この仮定は妥当といえる。

文が与えられると同時に、その文中の名詞、動詞、形容詞のそれぞれに対応する特徴がモデルに入力される。さらに、全ての物体や生物にはなんらかの属性を持つと仮定し、その文中では形容詞で修飾されていない名詞にも、属性を割り当てる。この属性はランダムに選ばれた形容詞に対応する。例えば、「お姉さんが赤い本を読んでいる」という文であれば、「お姉さんの属性」「お姉さん」「赤い」「本」「読んで」に対応する特徴が与えられる。今回、三つの意味カテゴリ、すなわち、物体、動作、および、属性のカテゴリを仮定し、これらはそれぞれ、名詞、動詞、および、形容詞に曖昧性なく対応する。この意味カテゴリと特徴の対応関係は既知とする。今回、各言語コーパスとその指示対象の 10,000 セットを用いてモデルを学習させた。

3.3.2 評価

モデルが上述の各コーパスを学習した後に, Imai et al. [11, 12] の実験に倣って, 学習コーパスに含まれない新奇語 X を以下の文でモデルに与える.

日本語

- 名詞条件 「 X がある .」
- 動詞 (項省略) 条件 「 X ってる .」
- 動詞 (項省略) 条件 「お姉さん が 何か を X ってる .」

英語

- 名詞条件 「There is a X .」
- 動詞 (項省略) 条件 「 X ing .」
- 動詞 (項省略) 条件 「She is X ing it .」

中国語

- 名詞条件 「有 X .」
- 動詞 (項省略) 条件 「 X .」
- 動詞 (項省略) 条件 「彼女 (tā) 在 X 某物 .」

この文と同時に「女性」と「女性の属性」(いずれも既知特徴), 「新奇物体」, 「新奇物体の属性」, および, 「新奇動作」に対応する五つの特徴をモデルに与える. 女性の属性と新奇物体の属性は, コーパスの形容詞語彙からランダムに選ばれた語に対応する属性である. そして, モデルは式 (3.1) を用いて X から特徴 o を推定する. なお, 女性や形容詞に対応する特徴が o として推定された場合, モデルは新奇物体と新奇動作を等確率で選択する. この仮定は, 名詞・動詞般用実験 [11, 12] で幼児が新奇語を物体や動作以外と推定していたとしても, 物体・動作変化刺激のどちらかを選択していたことに対応する.

各条件で X に対する o を推定する実験を，HMM の初期隠れ状態を変えて 20 回実施し， X を新奇動作と推定する確率の平均値でモデルを評価した．一標本 t 検定により，その評価値がチャンスレベル（0.5）を有意に上回れば，モデルは X を新奇動作に対応付けたと結論付ける．また，評価値が有意に下回れば，モデルは新奇物体に対応付けたとする．隠れ状態数 S が大きい場合と小さい場合をそれぞれ三歳児と五歳児の場合とみなし，これらの値を比較する．Imai et al. [11] の実験結果を再現できるように，隠れ状態数 S を三歳児モデルでは 3，五歳児モデルでは 5～7 に設定した．

さらに，獲得された隠れ状態（文法カテゴリ）の持つ品詞表現を解析した．学習コーパスの各語に品詞タグを割り当て，各隠れ状態についてその品詞の占める割合を求めた．例えば，動詞の表現についての数値は次式で与えられる．

$$\sum_{i \in \text{"verbs"}} P(s_i | \mathbf{w}, \mathbf{s}_{-i}) \quad (3.2)$$

ここで，“verbs” はコーパス中での動詞のインデックス集合である．全ての品詞についてのこの値の割合を，各隠れ状態に対して計算した．

3.4 実験結果

3.4.1 新奇語の意味の推定

Fig. 3.3 に (a) 英語，(b) 日本語，(c) 中国語のそれぞれの言語条件での結果を棒グラフで示す．縦軸は新奇語を動作と推定した割合を示す．名詞条件ではチャンスレベル 0.5 に対してその値が有意に小さく，また，動詞条件ではその値が有意に大きいと正しい般用といえる．棒グラフはモデル結果であり，点は Imai et al. [11] により報告された幼児の実験結果である（このデータの引用については，付録 A.1 を参照せよ）．白と黒の棒はそれぞれ， S の小さい，および，大きな場合の結果を示す．また，黒い点と白抜きの点はそれぞれ，三歳児と五歳児の結果である．

全ての条件において，モデルによる推定は幼児の成績と類似した傾向を示しており，モデルが各年齢・各言語についての幼児の新奇名詞・動詞の語意推定を再現できてい

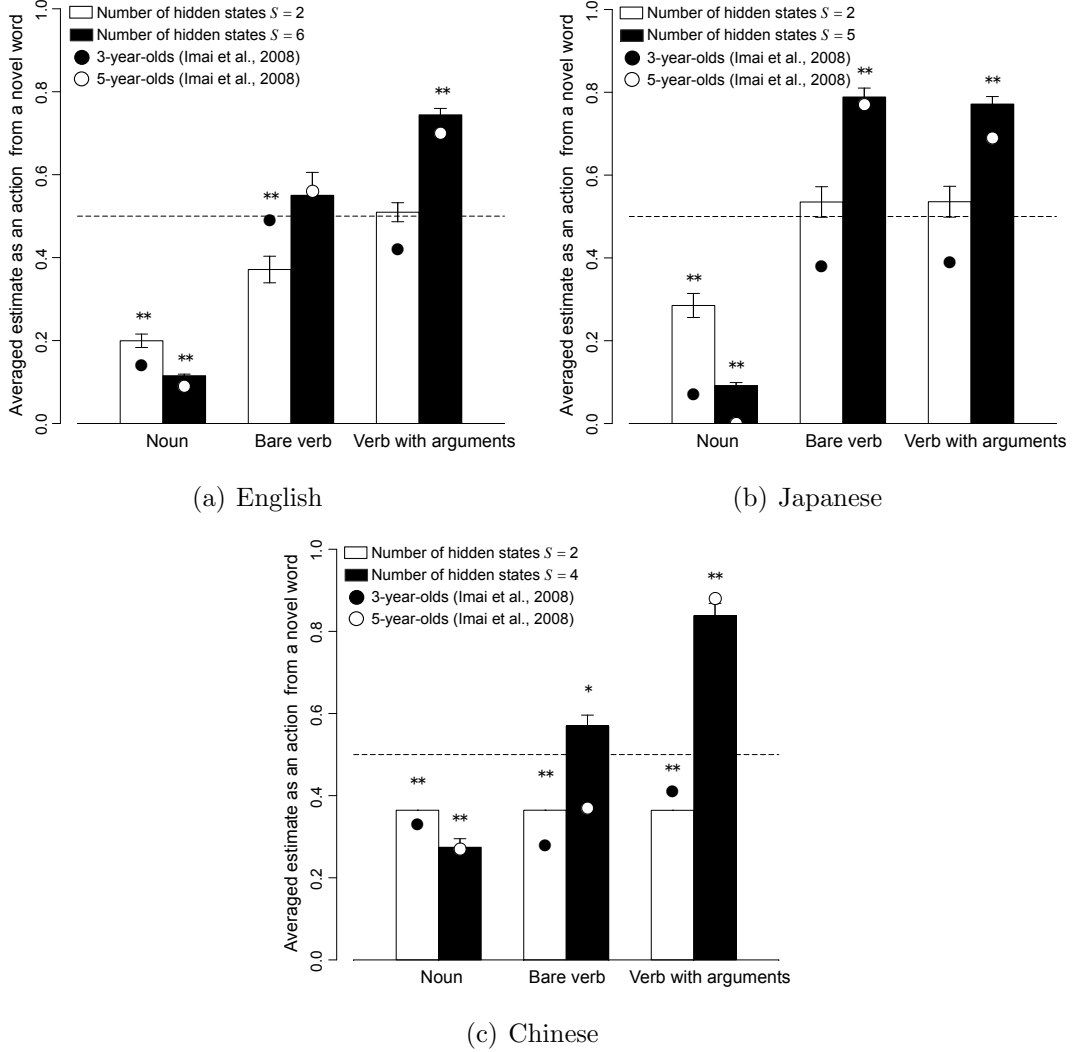


Figure 3.3: Estimation of the target feature corresponding to a novel word by the model (bars) and children (points). Results for the model trained with the (a) English, (b) Japanese, and (c) Chinese corpus. The bars indicate the average percentage of model estimates in which the inferred feature was a novel action. The white and black bars denote the results for the models with fewer hidden states (three-year-old case) and more hidden states (five-year-old case), respectively. The stars above the bars denote significant differences (*: $p < 0.05$, **: $p < 0.01$) between the model estimation and the chance level (0.5), shown as a dashed line. The error bars indicate the standard error. The black and white points denote the results achieved by three- and five-year-old children, respectively, as reported by Imai et al. (2008).

るといえる。モデルと幼児はどちらも、母語や年齢に関係なく、名詞条件で名詞を正しく選択できている（全て $p < .01$ ）。三歳児（ $S = 2$ ）の場合、どちらの動詞条件でも、チャンスレベルと同程度、または、それより小さい。一方で、五歳児の場合、学習した言語に依存して、異なる動詞条件の結果を示している。隠れ状態を六つ持ち、英語を学習したモデル（Fig. 3.3 (a)）では、項明示動詞条件で、新奇動作を正しく選択しているが（ $p < .05$ ）、項省略動詞条件では、その選択に失敗している（ $p > .05$ ）。隠れ状態を四つ持ち、中国語を学習したモデル（Fig. 3.3 (c)）も英語を学習したモデルと同様の傾向を示しているが、動詞（項省略）での値はチャンスレベルより有意に大きい（ $p < 0.05$ ）。それに対して、五つの隠れ状態を持ち、日本語語を学習したモデル（Fig. 3.3 (b)）では、どちらの動詞条件でも新奇動詞を動作へ正しく対応付けできている（ $p < 0.01$ ）。この場合、平均値で比較するとモデルと幼児の推定は同様の傾向であることがわかるが、実際には、動詞（項明示）条件での幼児の選択は統計的に有意ではない [11]。なお、今回は五歳児の名詞・動詞般用課題成績を最もよく再現する S を選択し、言語間で異なる S となったが、三つの言語条件で同じ S （ $= 6$ ）を持つ場合でも上記と同様の結果が得られた。

3.4.2 文法カテゴリの表現

前項で示した言語間、および、年齢間の差異の原因を調査するため、式 (3.2) を用いて HMM の隠れ状態の持つ品詞表現を解析する。Figs. 3.4–3.6 に英語、日本語、および、中国語を学習した場合の品詞表現の典型例を示す。左のグラフは三歳（小さな S ）、右のグラフは五歳（大きな S ）の場合である。図中のブロックが品詞表現の割合を表す。これらの品詞は、それぞれのコーパスの文法構造により決めた（品詞の詳細な説明については、付録 A.3.2 を見よ）。

二つの隠れ状態を持ち、英語を学習したほとんどのモデルは、名詞（1 番）と動詞（2 番）が分離した隠れ状態を獲得している（Fig. 3.4 の左図）。この名詞カテゴリにより、モデルは新奇名詞を正しく物体に対応付けられた。しかし、動詞カテゴリは形容詞や名詞の一部を含んでおり、これによって、動詞条件の推定がチャンスレベルに

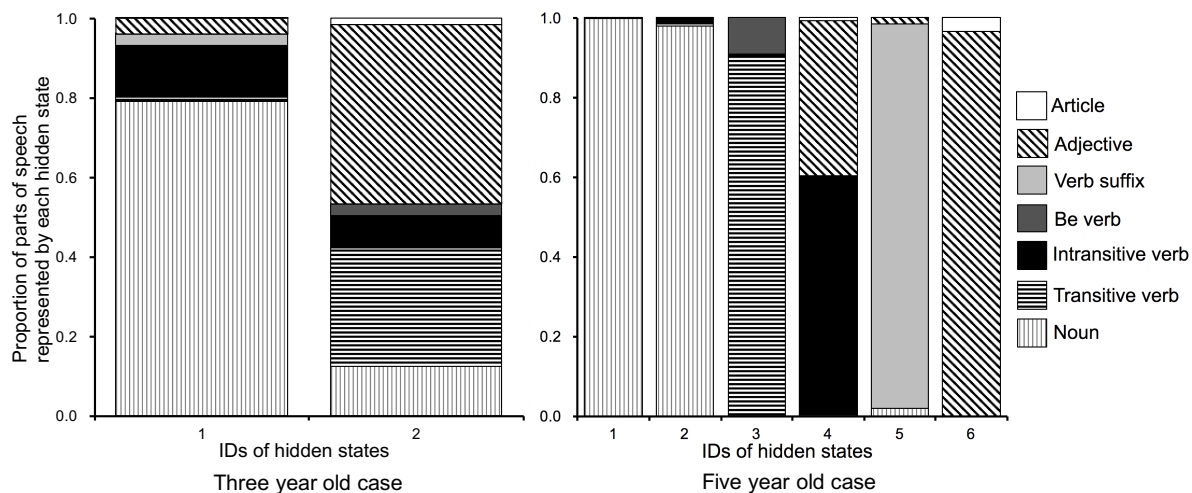


Figure 3.4: Part-of-speech representation for each hidden state of the HMM that was learned the English corpus. The graphs on the left indicate a typical representation of the three-year-old cases, whereas the graphs on the right indicate a typical representation of the five-year-old cases.

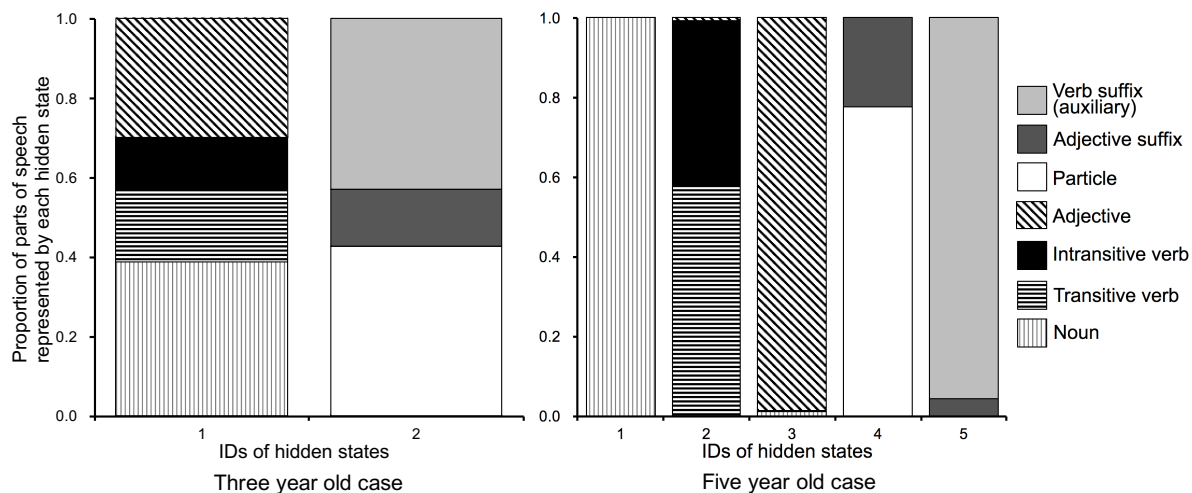


Figure 3.5: Part-of-speech representation for each hidden state of the HMM that was learned the Japanese corpus.

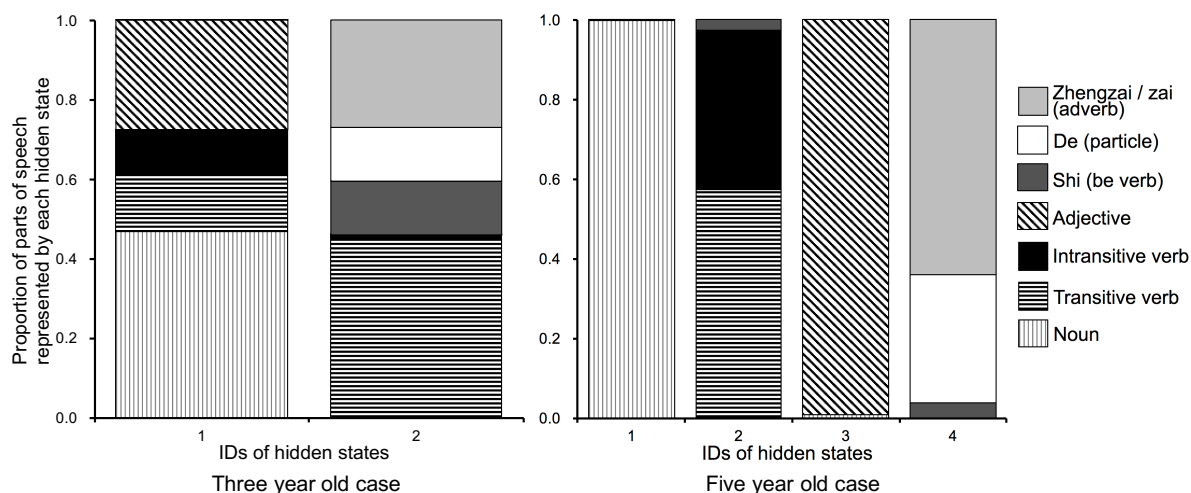


Figure 3.6: Part-of-speech representation for each hidden state of the HMM that was learned the Chinese corpus.

なったと考えられる。一方で、六つの隠れ状態を持つ場合 (Fig. 3.4 の右図)，モデルはよく分化した品詞表現を獲得している。特に動詞カテゴリ (3・4 番) が分化することによって、項を明示した場合に新奇動詞を動作に対応付けられた。名詞カテゴリは目的語カテゴリ (1 番) と主語カテゴリ (2 番) に分かれ、つまり、これらの文法カテゴリは語順，すなわち、統語的手がかりに依存しているといえる。したがって、動詞単体文が与えられた場合、文頭は名詞だが、進行形の接尾辞の前は動詞という矛盾した推定をモデルは余儀なくされ、動詞 (項省略) 条件での推定がチャンスレベルになったと考えられる。英語は名詞の欠落が少なく、動詞単体で文が与えられることがほとんどないため、モデルは安定した統語的手がかりに依拠することになり、上記の現象が現れたといえる。

二つの隠れ状態を持ち、日本語を学習した結果、二つの場合に分かれた。一つは、Fig. 3.5 の左図に示すような、自立語カテゴリ (名詞，動詞，および，形容詞) (1 番) と付属語カテゴリ (2 番) が獲得される場合であり、20 回中 11 回はこの表現となった。このような未分化なカテゴリにより、どの条件でもおよそチャンスレベルの推定になる。もう一つの場合では、名詞と動詞が分かれて表現された。しかし、動詞カテ

ゴリは形容詞を含んでいるため、動詞条件でのそのモデル推定はチャンスレベルとなった。一方で、五つの隠れ状態を持つ場合、動詞のみを表現するカテゴリが一つ現れる (Fig. 3.5 の右図の 2 番)。このモデルは動詞語尾カテゴリ (5 番) の前は動詞カテゴリであるという規則を獲得し、その結果、新奇動詞はこの動詞カテゴリに分類され、項が省略されたとしても、新奇動詞を動作に対応付けられる。日本語は名詞がよく欠落し、語順も入れ替わることがあるため、モデルは統語的手がかりよりも形態的手がかりに基づく文法カテゴリを獲得したと考えられる。したがって、英語で発生したような手がかり間の競合が、日本語の場合では起きず、動詞の両条件において、正しく動作特徴を推定できた。

二つの隠れ状態を持ち、中国語を学習した場合 (Fig. 3.6 の左図)、名詞や形容詞などの文頭の語のカテゴリ (1 番) とそれ以外の語のカテゴリ (2 番) が表現された。このその他カテゴリは、助動詞と、助動詞のない動詞を含む。したがって、新奇動詞は文頭にある、あるいは、助動詞の後にあるため、1 番のカテゴリに分類され、チャンスレベルに近い推定となった。四つの隠れ状態を持つ場合 (Fig. 3.6 の右図)、モデルは動詞のみで構成されるカテゴリを獲得する。しかし、中国語は英語や日本語のような名詞と動詞を区別する接尾辞を持たないため、項が省略されると、動詞般用ができない。一方で、助動詞カテゴリ (4 番) の後に新奇動詞を配置して、統語的手がかりを明示すると、その動詞を動作特徴に正しく対応付けられる。

3.4.3 文法カテゴリの獲得過程

Figs. 3.7–3.9 に S を 2 から 6 まで増加させたときに獲得した隠れ状態の品詞表現を示す。これらの図より、隠れ状態数の増加に伴う表現の変化が、幼児の文法カテゴリの発達に対応するという仮説を提案できる。ただし、各条件について S はあらかじめ設定されており、現在のモデルでは学習中に S が連続的に変化することは考慮されていない。

Fig. 3.7 より、英語の場合、まず名詞カテゴリが出現することがわかる。そして、この名詞カテゴリは主語 (文頭の語) と目的語 (文頭以外の名詞) に分化する。この結

$S = 2$	3	4	5	6
Noun	Noun	Noun	Noun	Noun
Noun (subject)	Noun (subject)	Noun (subject)	Noun (subject)	Noun (subject)
Transitive verb	Transitive verb	Transitive verb	Transitive verb	Transitive verb
Be verb	Be verb	Be verb	Be verb	Be verb
Intransitive verb	Intransitive verb	Intransitive verb	Intransitive verb	Intransitive verb
Verb suffix	Verb suffix	Verb suffix	Verb suffix	Verb suffix
Adjective	Adjective	Adjective	Adjective	Adjective
Article	Article	Article	Article	Article

Figure 3.7: Different part-of-speech representation of the hidden states with the number of hidden states varying from two to six in the English case.

$S = 2$	3	4	5	6
Noun	Noun	Noun	Noun	Noun
Transitive verb	Transitive verb	Transitive verb	Transitive verb	Transitive verb
Intransitive verb	Intransitive verb	Intransitive verb	Intransitive verb	Intransitive verb
Adjective	Adjective	Adjective	Adjective	Adjective
Particle	Particle	Particle	Particle	Particle
Adjective suffix	Adjective suffix	Adjective suffix	Adjective suffix	Adjective suffix
Verb suffix	Verb suffix	Verb suffix	Verb suffix	Verb suffix

Figure 3.8: Different part-of-speech representation of the hidden states with the number of hidden states varying from two to six in the Japanese case.

$S = 2$	3	4	5	6
Noun	Noun	Noun	Noun	Noun
Transitive verb	Transitive verb	Transitive verb	Transitive verb	Transitive verb
Intransitive verb	Intransitive verb	Intransitive verb	Intransitive verb	Intransitive verb
Adjective	Adjective	Adjective	Adjective	Adjective
Shi (be verb)	Shi (be verb)	Shi (be verb)	Shi (be verb)	Shi (be verb)
De (particle)	De (particle)	De (particle)	De (particle)	De (particle)
Zhengzai/zai (adverb)	Zhengzai/zai (adverb)	Zhengzai/zai (adverb)	Zhengzai/zai (adverb)	Zhengzai/zai (adverb)

Figure 3.9: Different part-of-speech representation of the hidden states with the number of hidden states varying from two to six in the Chinese case.

果より、英語の文法カテゴリを獲得する過程は語順，すなわち，統語的手がかりに依拠することがわかる．一方で，日本語の場合，まず付属語と自立語が分離し，付属語から助詞や助動詞が分化していく（Fig. 3.8）．したがって，日本語の文法カテゴリ獲得過程は助詞や接尾辞に基づいていることがわかる．また，中国語を学習する場合も同様に，付属語が現れる（Fig. 3.9）．しかし，中国語には名詞と動詞を区別できる接尾辞がなく，その付属語カテゴリは助動詞カテゴリと助詞カテゴリに分化する．同一の単語が動詞にも名詞にもなりえる中国語にとって，この助動詞が重要な統語標識であるといえる．このように，本モデルによって，言語構造を反映した文法カテゴリの獲得過程を表現することができた．

3.5 議論

本研究では，幼児の新奇名詞・動詞の意味を推定するモデルを提案した．HMMが入力文からその語の文法カテゴリを推定し，これらのカテゴリは観察した特徴と対応

付けられる。そして、提案モデルは獲得した文法カテゴリを用いて新奇語の指す適切な特徴を推定する。今回、モデルによる推定結果と、Imai et al. [11] が報告した英語、日本語、または、中国語を母語とする幼児の実験結果とを比較した。その結果、提案モデルは、各年齢についての幼児の推定における言語間差異を再現することができた (Fig. 3.3)。さらに、学習した HMM の隠れ状態を解析することによって、その言語間差異を生じさせた文法カテゴリ表現の詳細を明らかにできた (Figs. 3.4–3.6)。

これらの結果より、三つの事項を提案する。一つ目は、提案モデルは、文法発達初期の文法カテゴリ獲得を説明する計算論的枠組みとなることである。そして、このような抽象的な文法カテゴリが、幼児の語意学習に影響を与えることを示した。一つのモデルによって、三つの言語環境と二つの年齢群の複数の傾向を再現できたことは、提案モデルの妥当性を示す。統語的知識を新奇語へ一般化するには、抽象的な文法カテゴリやスキーマが必要とされ、それは二歳から三歳ごろに発達すると考えられている [4, 5, 6]。Tomasello はまた、その幼児の統語知識の構成における言語入力的重要性を強調している。さらに、コネクショニズム研究でもこの見方を支持しており、すなわち、言語獲得は、幼児が耳にする言語の統計的・確率的性質の学習とみなされる [65, 66]。

二つ目に、異なる年齢での語意推定法の違いは、HMM の隠れ状態数の違いで説明できることである。この数が小さいと、モデルは曖昧な文法カテゴリを獲得する。獲得した名詞カテゴリは比較的曖昧でなく、一方で、動詞カテゴリは他の品詞と混同されることが多い。名詞は文頭に來ることが多く、この統語的手がかりによって、言語に依存せず、名詞カテゴリを獲得しやすい。したがって、モデルは新奇名詞は物体に正しく対応付けられた。一方で、隠れ状態数が十分に大きくなると、文法カテゴリが分化し、新奇動詞でも般用できるようになる。今回、英語は $S = 6$ 、日本語は $S = 5$ 、および、中国語は $S = 4$ といったように、各言語の五歳児モデルとして、異なる S の値を設定した。これらの値は、モデルの結果が Imai et al. [11] の結果とできるだけ一致するように選ばれた。この言語間での S の違いは、学習コーパスの文法的複雑度の差を反映している可能性がある。したがって、これらの値に重要な意味があるというよりは、異なる年齢の幼児による語意推定が異なる隠れ状態数で説明できることを強

調したい。

三つ目に、五歳児の動詞般用の言語間差異が、言語入力に対応する文法カテゴリを獲得した結果といえる点である。英語コーパスの統語規則は欠落や入れ替えがなく安定しており、また、英語の動詞はよく拘束形態素を伴うため、文法カテゴリを推定するための統語的・形態的手がかりを多く含む。したがって、一語文「X-ing」が与えられたとき、Xは文頭であるため名詞であり、「ing」の前であるため動詞という二人の手がかりによりモデルは文法カテゴリを推定する。その二つの手がかりの競合の結果、動詞（項省略）条件ではチャンスレベルの推定となった。しかし、主語や目的語といった統語的手がかりがよく欠落する日本語を学習したモデルは、形態的手がかりを重視する結果となった。これによって、日本語モデルは項省略条件でも、動作を正しく推定できた。一方、中国語コーパスは動詞接尾辞を持たず、語順（統語手がかり）でのみ動詞と名詞が区別される。したがって、中国語を学習したモデルでは、項が省略され、統語手がかりがなくなると、動詞と名詞を区別できなかった。

以上のことから、提案モデルはその入力言語に対して適切な手がかりを用いているといえる。評価の際に、テスト文がその手がかりを含んでいれば、新奇語を正しいカテゴリに分類でき、語意推定に成功する。従来研究は複数の手がかりを用いて文法カテゴリを計算可能であることを指摘しているのみであるが（例えば、[34, 37, 41]）、提案モデルは、従来の行動実験 [11] が示したような、幼児が母語に適した手がかりを用いていることを示唆する。このメカニズムは、語意は文中の複数の手がかりの競合により推定されるという競合モデル（例えば、[67, 68]）に類似している。ただし、提案モデルでは、競合のメカニズムを明に導入しておらず、学習によって、文法カテゴリの同定に有利な手がかりを自動的に学習している。提案手法は幼児の競合メカニズムの発達についても新しい示唆を与えられると思われる。

以降の節では、幼児の言語発達に対するさらなる洞察や、現状モデルの限界について議論する。

3.5.1 他動詞と自動詞の統語ブートストラッピング

本研究では、新奇名詞と動詞の般用についての問題を扱っているが、他の多くの統語ブートストラッピングに関する研究では、自動詞と他動詞といったより下位のカテゴリを扱っている。それらの研究では例えば、三歳以下の幼児で、新奇な自動詞、または、他動詞が指す動作を適切に選択できるかを調査している（例えば、[8, 69, 70, 71]）。Naigles [8] は、英語を母語とする 25ヶ月児に他動詞文「the duck is *gorping* the bunny」（*gorp* は新奇語）を聞かせたとき、自動詞的な動作よりも他動詞的な動作をより長く注視することを報告した。ここでの自動詞的な動作とは、アヒルがウサギの背中を押して前屈させることであり、また、他動詞的な動作とは、アヒルとウサギが別々に自身の腕を曲げることである。一方、幼児が「the duck and the bunny are *gorping*」という自動詞文を聞いたとき、その結果は逆になり、自動詞的な動作をより見るようになる。これはおよそ二歳の幼児でもすでに自動詞と他動詞についての抽象的な文法カテゴリの知識があることを示し、提案モデルや Imai et al. [11] の実験結果と整合的でない。

提案モデルでは上記の結果を説明することはできないが、そのメカニズムについての示唆を述べる。提案モデルでは、全ての言語条件において、他動詞と自動詞のカテゴリよりも名詞と動詞のカテゴリが早く分化しており（Figs. 3.7–3.9）、これはおそらく他の教師なし分類法でも同様に起こり得る。他動詞と自動詞のカテゴリを獲得するためには、（半）教師あり学習法が必要になる。もし、モデルが語の主題役割（例えば、動作主や被動者）を同定できれば、名詞や動詞のカテゴリを分化させる前に、他動詞と自動詞のカテゴリを獲得できるかもしれない。前言語期の乳児は視覚刺激中の動作主と被動者といった因果関係を知覚できるという研究がいくつかある [72, 73]。幼児にとって、多くの視覚特徴から具体的な動作や物体を同定することよりも、動作主と被動者の因果関係を理解するほうが簡単であるかもしれない。したがって、幼児の言語能力の一部についての学習が、因果関係の認識に影響されることは自然といえる。そして、幼児は与えられた課題に対して適した文法知識を選択的に利用する可能性がある。エージェント間の因果関係を同定する課題であれば、幼児は因果関係を基に学習した文法知識を用いているのかもしれない。一方で、具体的な動作や物体を同

定する課題であれば、幼児は教師なしで学習した文法知識を用いる可能性がある。

3.5.2 名詞優位性の統語論的メカニズム

提案モデルは、三言語全ての場合において、三歳児は名詞般用可能だが、動詞般用不可であることを示した。これは、学習言語が動詞カテゴリの手がかりよりも名詞カテゴリの統語手がかりを多く有しているからである。名詞の学習が動詞の学習より容易であり、早いという傾向は、様々な言語での、幼い幼児の語彙発達に関する多くの研究で普遍的に観察されている（例えば、[74, 75, 76, 77]）。いくつかの観察的・実験的研究では、この名詞優位性は知覚的・意味論的要因によりもたらされると議論している [74, 75, 78, 11]。すなわち、動作の時空間的境界は不明瞭であり、動作の切り出しは物体に比べて難しいことが名詞優位性をもたらす。

対照的に、提案モデルでは、名詞優位性は統語論的な見方で説明できる。他の文法カテゴリに関する計算論的な研究でも、名詞カテゴリの獲得がより早いことを指摘している（例えば、[35]）。また、この優位性は言語入力の影響を受けることを示唆する研究がある（例えば、[79, 80]）。Choi and Gopnik [79] は韓国語を母語とする養育者による発話は、英語を母語とする養育者の発話より多くの動詞を含むことを報告した。そして、およそ 19ヶ月で、韓国語を母語とする幼児の発話における動詞の数が急増したが、この傾向は英語を母語とする幼児にはみられなかった。韓国語では、主語や目的語が欠落することがあり、動詞がよく文尾に配されることから、上記の結果は韓国語が動詞カテゴリを構成するための統語的情報を多く含んでいる結果と解釈できる。さらに、Imai et al. [12, 11] と同様の課題において、英語を母語とする二歳児でも、目的語と標的物体が幼児にとって親しみのあるものであれば、新奇動詞を動作に対応付けられることが報告されている [81]。この結果は、既知の物体の選択肢が除外されたことで、幼児が新奇動詞の指す動作を選択できるようになったと解釈でき、幼児は動作の知覚がすでに可能であることを示す。すなわち、三歳児での動詞般用の失敗は、知覚的・概念的な動作カテゴリの切り出しによる困難さではなく、動詞カテゴリ形成の困難さに起因する可能性がある。

本研究は言語的（形態論・統語論的）性質が語学習に及ぼす影響を調査したが、語彙発達の初期において、言語的要因のみがその学習法を決めると主張するものではない。レキシコンの獲得は、語や特徴のカテゴリを規定する言語的・言語外の要因といった複数の手がかりに支えられる。実際、Imai et al. [11] は新奇語の意味の学習における言語外の手がかりの重要性を強調している（付録 A.1 も見よ）。特に、中国語を母語とする五歳児は、標準刺激動画で物体の存在を強調すると、動詞を物体に対応付ける結果となった [11]。幼児は言語的手がかりと言語外の手がかりをどのように統合し、語意を推定するのかについての実験的・計算論的研究は今後の重要な課題である。

第4章 未分化な文法カテゴリによる過剰生成モデル

4.1 はじめに

前章では、提案モデルによる文法カテゴリ獲得が言語理解（語意推定）に及ぼす影響を検討したが、本章では、それが言語産出に与える影響を調査する。今回、1.1.2項で述べた英語と日本語の過剰生成を、前章と同じHMMを用いてモデル化する。HMMは語のカテゴリ化だけでなく、そのカテゴリに基づいて語を生成することもでき、学習後のHMMから生成された語列を産出文とみなす。しかし、未分化な文法カテゴリからは、過剰生成の誤用だけでない全く非文法的な文章が産出される可能性がある。前章のモデルでは、三歳頃から文法カテゴリが精緻化されていくと仮定したが、実際には、三歳児でも文法的に正しい文章を産出できる。そこで、より一般的な文法発達モデルに目を向ける。

Tomasello [82, 6] は幼児のもつ統語構文には三つの段階があるとしている。

1. 語結合と軸語スキーマ（18ヶ月頃）：頻度の高い語結合から抽出されるパターンであり、少数の語（軸語）が特定の語順で他の語と結びつく構造。
2. 項目依拠的構文（24ヶ月頃）：既知の動詞を正しい語順で使用し、統語標識（語順や格標識）が機能し始める。ただし、これは特定の語（特に動詞）固有の構文パターンであり、新奇の語にこのパターンを適用することはできない。
3. 抽象的統語構文（36ヶ月以上）：動詞一般に適用できる規則を獲得する。新奇動詞であっても、動詞構文として理解・産出することができる。

このように、過剰産出が出現する二歳から三歳にかけて、語レベルでの語遷移規則から、より抽象的な文法カテゴリ（動詞や名詞など）の遷移規則の獲得がなされている。このような規則の抽象化の過程で、その規則が定着していない（明確化されていない）ため、その規則の過剰な一般化による誤用が生じると考えられている [4]。

Tomasello [6] の統語発達段階を簡潔に表現するために、二つの規則を仮定する（ただし、Tomasello は明に二つの処理を言及しておらず、むしろ、項目依拠的なパターンから抽象的な規則が形成されていくことを強調している）。一つは語遷移規則であり、語結合と軸語スキーマに対応する。これは、語の文法カテゴリを考慮しない語ごとの遷移確率で表され、入力文により獲得される。したがって、過剰生成のような誤用はほとんど生じない。もう一つは文法カテゴリ遷移規則であり、抽象的統語構文に対応し、HMM を用いてモデル化する。これは HMM で推定された文法カテゴリ列に対する、カテゴリ遷移確率で表される。そして、この文法カテゴリからそれに属す語が生成される。この文法カテゴリが適切に獲得されていないと、非文法的な語列が生成される。これらの二つの規則を競合させて、語列を生成させる。この競合は、複数の基本的な認知メカニズムを競合させて言語習得や文処理を行うという競合モデル [67, 68] の考えに基づく。今回の場合、語のコロケーション規則をそのまま記憶するメカニズムと、その抽象的な文法カテゴリ間の規則を用いるメカニズムとの競合である。この二つの処理と競合により、Words and Rules 理論 [26] でいうグラマーの処理を実現する。

このモデルにおいて、文法カテゴリ数の増加によって過剰生成が出現・消失することが期待される。カテゴリ数が少ない場合（発達初期）では、有意な文法カテゴリ遷移規則が形成されないため、語遷移規則が優位となり、過剰生成は生じない。カテゴリ数が中程度の場合（発達中期）では、以前に比べると文法カテゴリ遷移規則が優位になるが、そのカテゴリ内容が曖昧であり、過剰生成が出現する。そして、カテゴリ数が十分に多くなると（発達後期）、精緻な文法カテゴリが形成され、過剰生成が消失する。

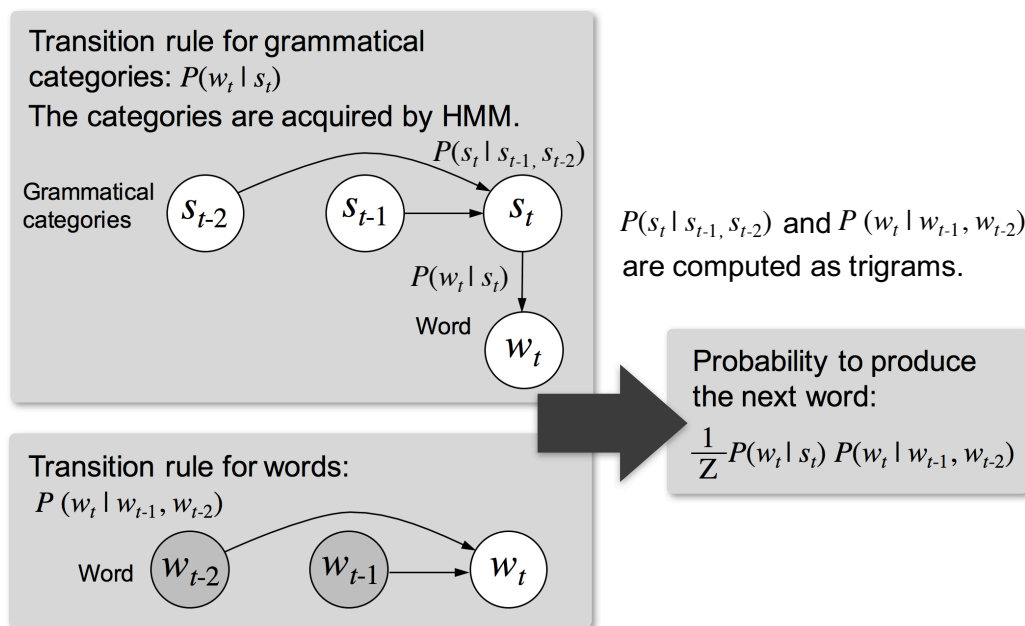


Figure 4.1: The proposed statistical model. The transition rule of words is a trigram for words, and the transition rule of grammatical categories is a trigram for grammatical categories. A probability for generation of the next word is given as the product of these rules.

4.2 モデル概要

Fig. 4.1 に提案するモデルを示す. このモデルは語遷移規則と文法カテゴリ遷移規則の二つの規則により構成される. それぞれの規則により次に生成する語の確率を求め, それらを競合 (乗算) させて次の語の生成確率を得る.

語遷移規則は語列 $\mathbf{w} = \{w_1, w_2, \dots\}$ に基づく. 今, 時刻 t における語 w_t の生成確率を, 以前の語列から求めたいとする. 今回, 二時刻前までの語 (w_{t-1} と w_{t-2}) に基づいて w_t の確率 $P(w_t | w_{t-1}, w_{t-2})$ が定まるとする. この確率を次式のトライグラムによって求める.

$$P(w_t | w_{t-1}, w_{t-2}) = \frac{C(w_t, w_{t-1}, w_{t-2}) + \alpha}{C(w_{t-1}, w_{t-2}) + W\alpha} \quad (4.1)$$

ここで, $C(*)$ は $*$ の並びが学習コーパス中にある頻度であり, W は全体の語彙数で

ある。また、 α はゼロによる除算を避けるための定数であり、小さな値が選ばれる。したがって、語遷移規則には学習コーパスにない語の並びの確率は非常に低くなり、過剰生成はほとんど生じない。

文法カテゴリ遷移規則は、 $\mathbf{w}$ に対応する文法カテゴリ列 $\mathbf{s} = \{s_1, s_2, \dots\}$ に基づく。 $\mathbf{s}$ はベイジアン HMM [83] により計算される（付録 A.2.1 を参照）。 w_t を生成するとき、まず、 w_t の文法カテゴリ s_t を二時刻前までの文法カテゴリに基づいて求める ($P(s_t|s_{t-1}, s_{t-2})$)。そして、 $P(w_t|s_t)$ により、 w_t の生成確率を得る。 $P(s_t|s_{t-1}, s_{t-2})$ と $P(w_t|s_t)$ はそれぞれ、次式で表される学習コーパスの $\mathbf{s}$ でのトライグラムと、カテゴリと語の共起確率で求める。

$$P(s_t|s_{t-1}, s_{t-2}) = \frac{C(s_t, s_{t-1}, s_{t-2}) + \beta}{C(s_{t-1}, s_{t-2}) + S\beta} \quad (4.2)$$

$$P(w_t|s_t) = \frac{C'(s_t, w_t) + \gamma}{C'(s_t) + S\gamma} \quad (4.3)$$

ここで、 $C'(*)$ は学習コーパス中で $*$ の要素が共起する頻度であり、 S はあらかじめ設定される文法カテゴリ数である。また、 β と γ は α と同様の小さな値の定数である。文法カテゴリには複数の語が属するため、そこから生成される語は、学習コーパスにはない語順である可能性がある。

以上の二つの語生成確率 ($P(w_t|w_{t-1}, w_{t-2})$ と $P(w_t|s_t)$) を乗算して、次の語の生成確率 $\frac{1}{Z}P(w_t|w_{t-1}, w_{t-2})P(w_t|s_t)$ を得る。ここで、 Z は正規化定数である。

4.3 実験設定

前節のモデルに人工コーパス 10,000 文を入力して学習させ、式 (4.1) – (4.3) の $C(*)$ と $C'(*)$ を求めた（コーパスの詳細は次節を参照のこと）。学習後のモデルは、式 (4.1) – (4.3) に従って、学習コーパスに続いて一語を生成する。そして、その語に続いて、次の一語を生成する。これを 5,000 文になるまで繰り返し、これらの生成文を解析する。なお、式 (4.1) – (4.3) 中の定数 α 、 β 、および、 γ はいずれも 0.0001 とした。また、学習後の語生成ではモデルは学習しない。カテゴリ（隠れ状態）数 S は学習中固定であり、1 から 20 まで変えて実験した。学習に用いたコーパスが持つ語カテゴリ

(Fig. 4.2 の円の数, 次項で詳細を述べる) は高々15であり, モデルカテゴリの最大数 20 はこれらのコーパスを表現するために十分大きいといえる.

まず, 生成文の過剰生成割合を計算した. 英語の場合, 不規則変化動詞の直後に「ed」がある頻度を, 「ed」の全体の頻度で除算した値を過剰生成割合とした. また, 幼児にはみられない過剰生成として, 動詞以外の語の直後に「ed」がある頻度を, 「ed」の全体の頻度で除算した値も求めた. 日本語の場合, 形容詞の直後に「の」がある頻度を, 「の」の全体の頻度で除算した値を過剰生成割合とした. また, 名詞・形容詞以外の語の直後に「の」がある頻度を, 「の」の全体の頻度で除算した値も求めた. HMM の初期カテゴリ割り当てをランダムに変えて, 各 S の値について 10 回試行し, その平均を求めた. そして, 獲得された文法カテゴリを図示するために, 第 3 章の語意推定実験と同様に式 (3.2) を用いて, それらが表現する品詞を算出した. また, 語遷移規則と文法カテゴリ遷移規則のどちらが優位にはたらいて過剰生成が生じたかを示すために, 英語なら「ed」, 日本語なら「の」が生成されたときの各規則の確率 ($P(w_t = \text{「ed」 or 「の」} | w_{t-1}, w_{t-2})$ と $P(w_t = \text{「ed」 or 「の」} | s_t)$) からそれらの比を計算した.

4.3.1 学習コーパス

英語コーパスの生成規則を Fig. 4.2 (a) に示す. 文は BOS (begin of sentence) から始まり, EOS (end of sentence) で終わる. BOS と EOS も語 (ダミーワード) としてモデルの学習の対象となる. ただし, 次節では, BOS と EOS のカテゴリを省略して結果を示す. 各円は, それぞれ定義された語彙の中からランダムに一語生成する. 語彙数は付録 A.3.2 に記載する. 灰色の円上の小数は省略確率, 矢印上の小数は遷移確率を示す. 他動詞文と自動詞文の割合 (10 : 3) は, 英語の対幼児発話文を解析した研究 [64] における主語述語文での割合に依拠する. 下部の四角形は自動詞の内容であり, 現在形と過去形がランダムに選ばれる. 過去形の場合は, 規則変化, または, 不規則変化が選ばれる. 一般に, 英語の不規則変化動詞の語彙数は規則変化動詞のものよりも少ないが, 特に対幼児発話文においては, 不規則変化動詞 (gave や took など)

のほうが多用される傾向にある [6]. したがって、過去形における不規則変化動詞と規則変化動詞の割合を 3 : 2 に設定した. なお, この割合は Plunkett and Marchman (1993) の研究でのモデルへの初期の言語入力とおよそ一致している. この遷移規則は他動詞の場合でも同様である.

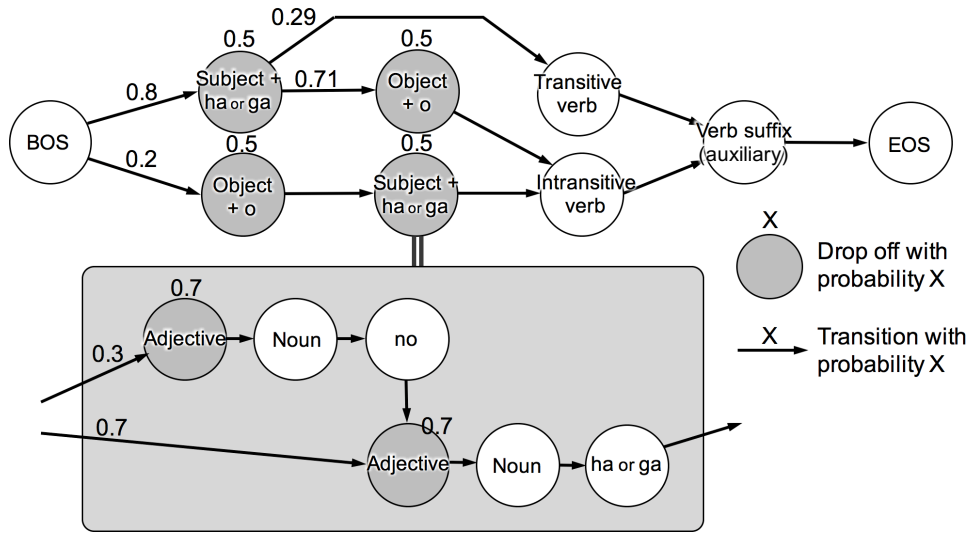
日本語コーパスの生成規則を Fig. 4.2 (b) に示す. 日本語の場合, 主語-目的語-動詞 (SOV) と目的語-主語-動詞 (OSV) の語順が存在する. また, 主語と目的語, および, それらに付随する格助詞が 0.5 の確率で省略される. これは日本語の対幼児初話文における格助詞の省略割合がおよそ 0.4 から 0.7 であることに基づく [84]. 下部の四角形は「主語 + 助詞」の内容であり, 0.3 の確率で「名詞 + の + 名詞」, 0.7 の確率で「名詞」が選ばれる. また, 名詞の前には 0.3 の確率で形容詞が配置される. 自動詞と他動詞の割合は英語コーパスと等しくした.

4.3.2 コーパス条件

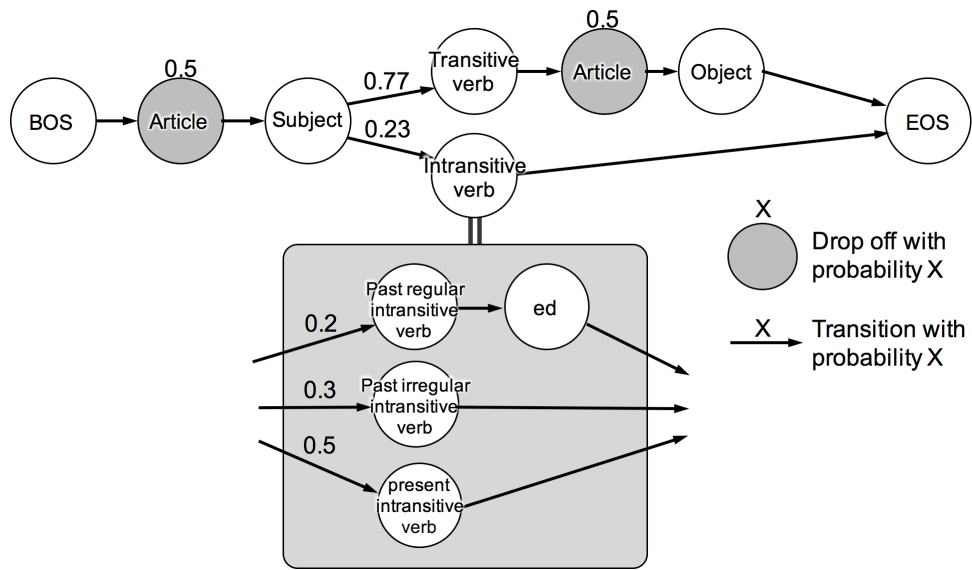
語遷移規則と文法カテゴリ遷移規則のそれぞれの役割を明らかにするために, 二つのモデル条件を設定した. 一つは語遷移規則と文法カテゴリ遷移規則の両方を有するモデル, もう一つは, 文法カテゴリ遷移規則のみを有するモデルである. 後者のモデルは, 式 (4.3) から直接, 語を生成する. なお, 語遷移規則のみの場合, 誤用はほとんど生じないため, 解析対象としない.

また, 学習コーパスが過剰生成に及ぼす影響を調査するために, 前節の「レギュラー」コーパスに加え, 別のコーパス条件を設定した. 英語の場合, 不規則変化動詞に対して規則変化動詞の頻度が大きいと, 些末な不規則変化動詞のカテゴリとその規則が明に獲得されず, 過剰生成がより発生する可能性がある. したがって, 規則変化動詞と不規則変化動詞の相対的な頻度が過剰生成に重要であると考えられる. そこで, 不規則変化動詞と規則変化動詞の割合 (Fig. 4.2 (b)) を 3 : 2 から 3 : 7 に変更した.

日本語の過剰生成の発生要因は「(名詞・形容詞) + 格助詞」のカテゴリ規則の獲得であり, そのためには, 格助詞「の」と「は・が・を」の十分な頻度が必要になると考えられる. さらに, 過剰生成を引き起こす不明瞭なカテゴリは, モデルのカテゴリ



(a) English



(b) Japanese

Figure 4.2: Generation rules for artificial corpora.

り数の限界に対して、入力言語の複雑さ、すなわち、潜在的な文法カテゴリが大きいときに形成される。日本語の場合、この言語の複雑さは「SOV」と「OSV」の二つの語順の存在により増加し、このことが間接的に日本語の過剰生成に寄与している可能性がある。以上のことから、次の三つのコーパス条件を設けた。

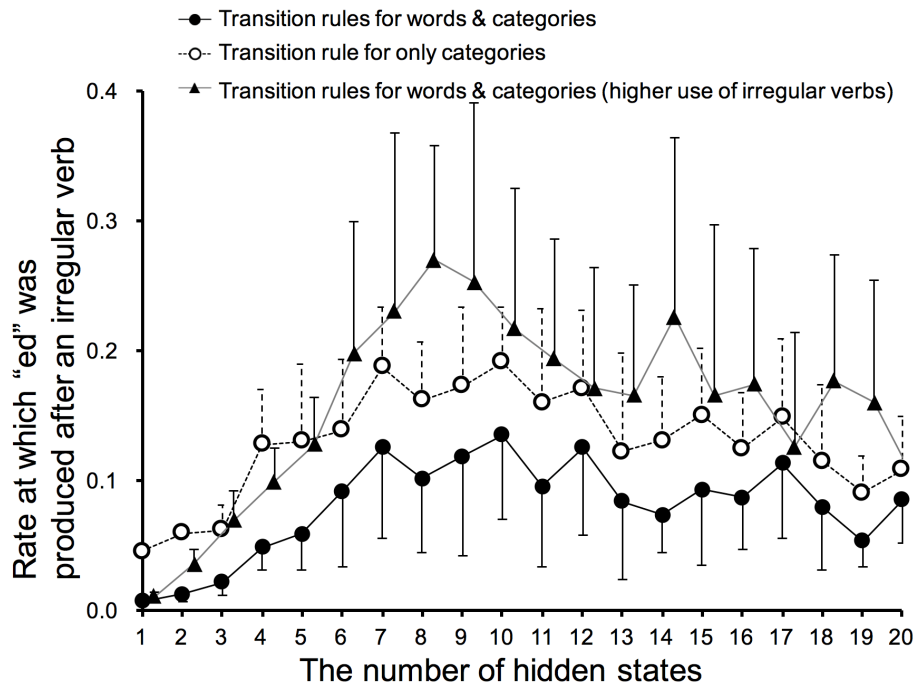
- 「の」コーパス：名詞は必ず「の」で連体修飾され、格助詞「の」の頻度が多い。
- 「は・が・を」コーパス：格助詞「は・が・を」の省略がなく、これらの頻度が多い。
- 「SOV」コーパス：SOVの語順のみであり、OSVの語順は存在しない。

これらのコーパス条件について、語遷移規則と文法カテゴリ遷移規則の両方を持つモデルで実験した。

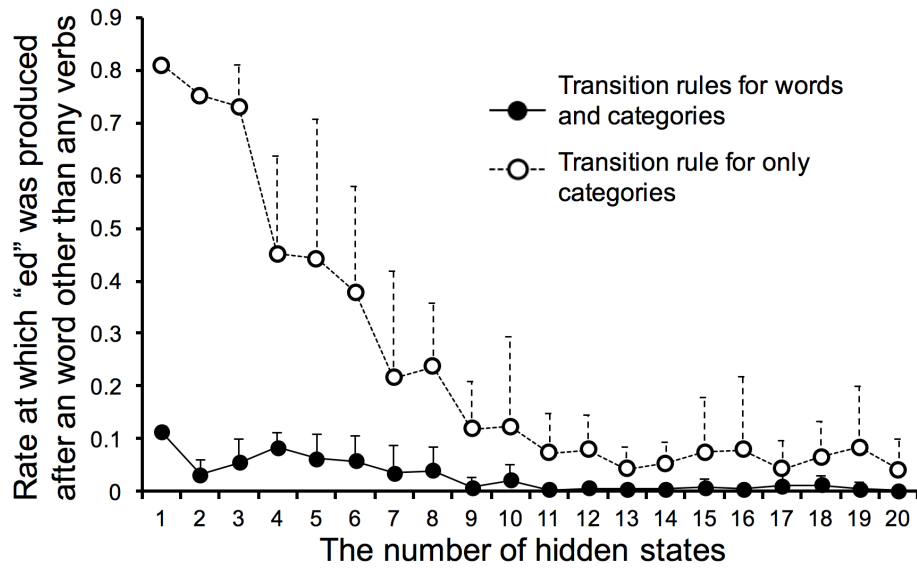
4.4 実験結果

4.4.1 英語での過剰生成

英語での実験結果を Fig. 4.3 に示す。Fig. 4.3 (a) の横軸はカテゴリ数、縦軸は「ed」の過剰生成の割合を示す。白い点は文法カテゴリ遷移規則と語遷移規則の両方を持つモデルの結果、黒い点は文法カテゴリ遷移規則のみを持つモデルの結果である。この図より、カテゴリ数の増加に伴って、過剰生成割合も増加していることがわかる。そして、カテゴリ数 10 をピークにして、その後、過剰生成割合が減少する傾向にある。語遷移規則があることで過剰生成が減少しているものの、この有無についての二つのモデルの結果は類似した曲線を示している。しかし、動詞以外の語に「ed」を付加した割合を示す Fig. 4.3 (b) では、これらの間に大きな違いがみられる。語遷移規則と文法カテゴリ規則の両方を持つモデルでは、そのような誤用がほとんど生じていないのに対して、文法カテゴリ規則のみのモデルでは、特にカテゴリ数が少ないときに大きな割合となっている。したがって、語遷移規則により、動詞以外への「ed」の誤用が抑制されているといえる。



(a) Irregular verbs



(b) Words other than verbs

Figure 4.3: Results of overproduction in English. Markers indicate rates of overproduction of "ed" for (a) irregular verbs and (b) words other than verbs. Error bars indicate standard deviation.

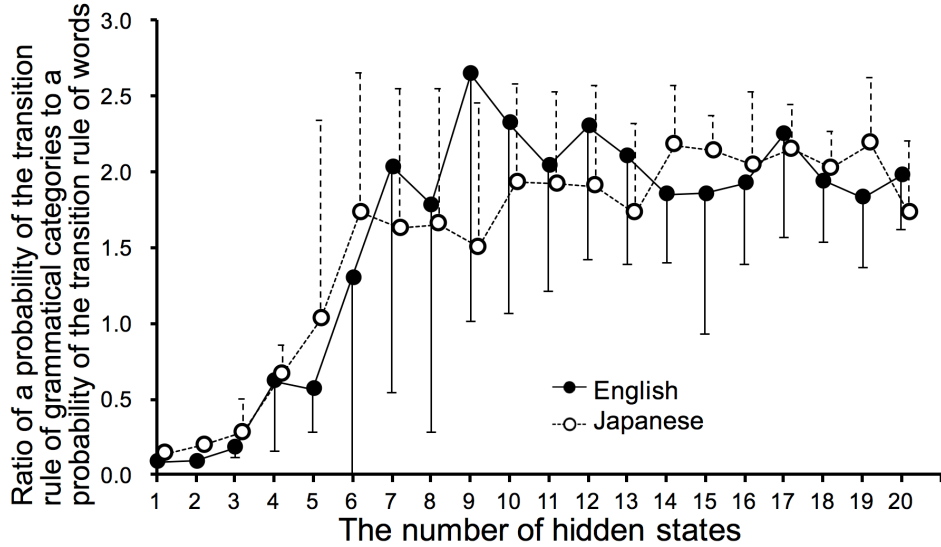


Figure 4.4: Averaged ratios of a probability of the transition rule of grammatical categories to a probability of the transition rule of words in English and Japanese. The probabilities were for “ed” when “ed” was produced in the English case, while they were for “no” when “no” was produced in the Japanese case. Error bars indicate standard deviation.

このような過剰生成の原因を探るため、「ed」を生成したときの語遷移規則の確率に対する文法カテゴリ遷移規則の確率の比を Fig. 4.4 に示す。この値が大きいほど、文法カテゴリ遷移規則が主だって「ed」を生成したことになる。この図より、カテゴリ数が増加し、精緻な文法カテゴリが獲得されるにしたがって、文法カテゴリ遷移確率が「ed」の産出に大きく寄与するようになることがわかる。以上のことより、文法カテゴリ遷移規則が「ed」の過剰生成の出現と消失の傾向の原因であることがわかる。

次に、HMM が獲得した文法カテゴリ構造を調査する。Fig. 4.5 はカテゴリ数が3 のとき獲得されたカテゴリの典型例である。動詞を表現するカテゴリ（2 番と3 番）が「ed」を併せて表現していることがわかる。したがって、「動詞 → ed」の規則が獲得されず、語遷移規則が優位になり、正しく「ed」が生成されたと考えられる。一方、カテゴリ数が10 になると（Fig. 4.6）、動詞カテゴリ（5 番と6 番）と「ed」カテゴ

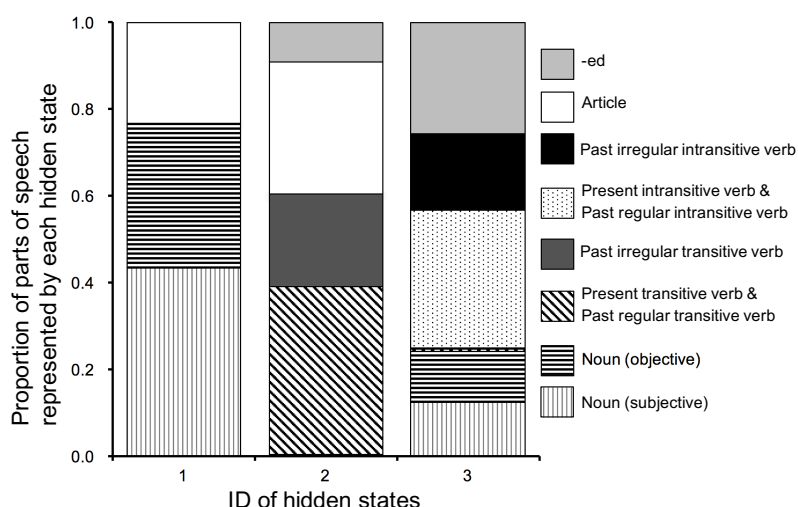


Figure 4.5: Typical representation of three grammatical categories acquired from the English corpus.

り（10 番）が分離して表現される．5 番のカテゴリは他動詞，6 番のカテゴリは自動詞を表現しているが，これらは規則変化動詞と不規則変化動詞を区別せずに含んでいる．したがって，動詞カテゴリから「ed」カテゴリへの有意な遷移を獲得した結果，過剰生成が生じたといえる．そして，Fig. 4.7 はカテゴリ数 19 で，過剰産出が比較的少なかった HMM の例である．12 – 14 番のカテゴリが他動詞，15 番のカテゴリが自動詞を表している．Fig. 4.6 の文法カテゴリに比べると，14 番と 15 番のカテゴリが表現する不規則変化動詞の割合が多いことから，「ed」の過剰生成が減少する傾向がみられた．しかし，完全に不規則変化動詞のみを表現するカテゴリは獲得されず，過剰生成が消失することはなかった．

規則変化動詞の頻度の多いコーパス（「規則変化動詞」コーパス）を用いて学習した結果を Fig. 4.3 (a) の三角形で示す．この図より，上記の「レギュラー」コーパスでの結果に比べ，大きな過剰生成割合となっていることがわかる．規則変化動詞の頻度を増やすことで，カテゴリにおける規則変化動詞表現の割合が増加し，不規則変化動詞がそのカテゴリに取り込まれ，過剰生成が生じやすくなったと考えられる．

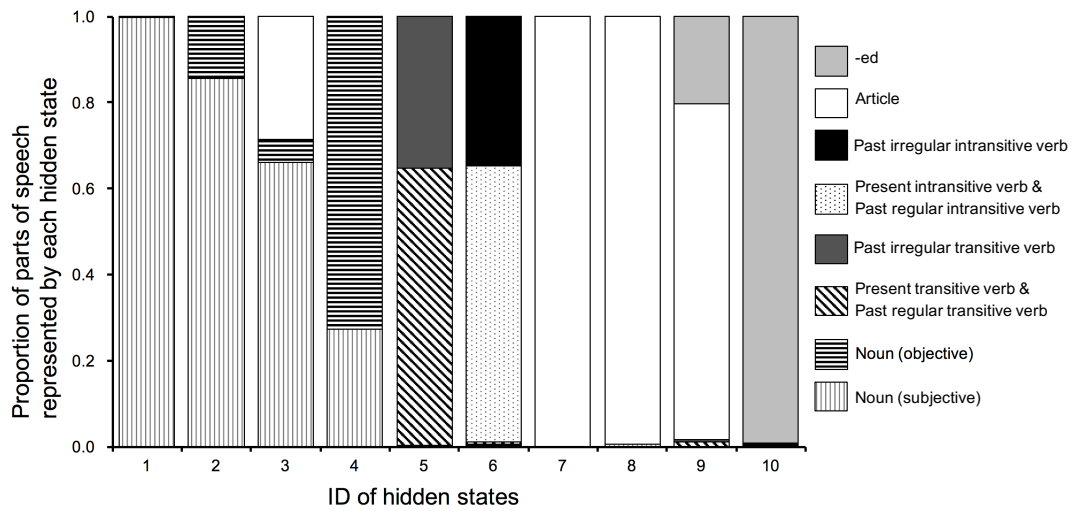


Figure 4.6: Typical representation of ten grammatical categories acquired from the English corpus.

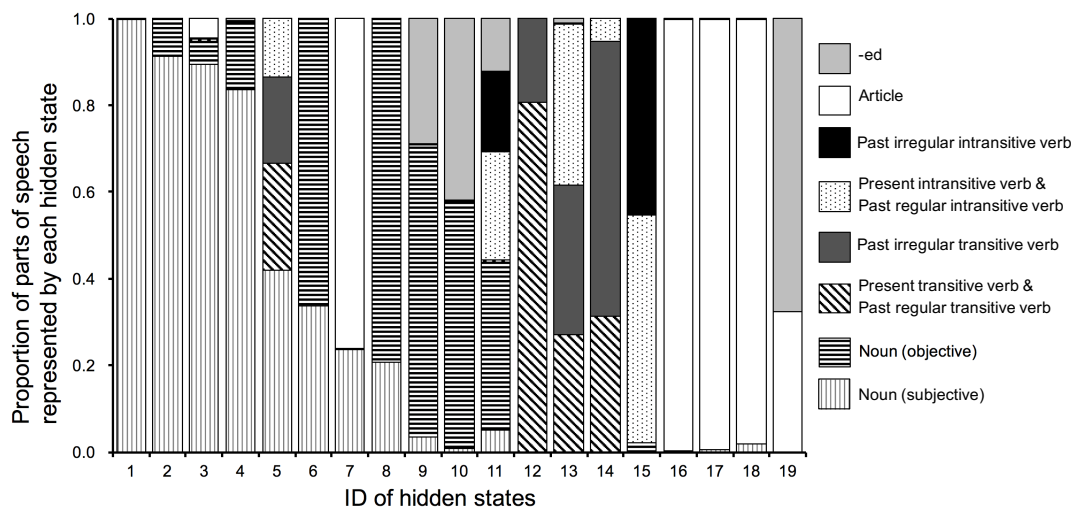


Figure 4.7: Typical representation of nineteen grammatical categories acquired from the English corpus.

4.4.2 日本語での過剰生成

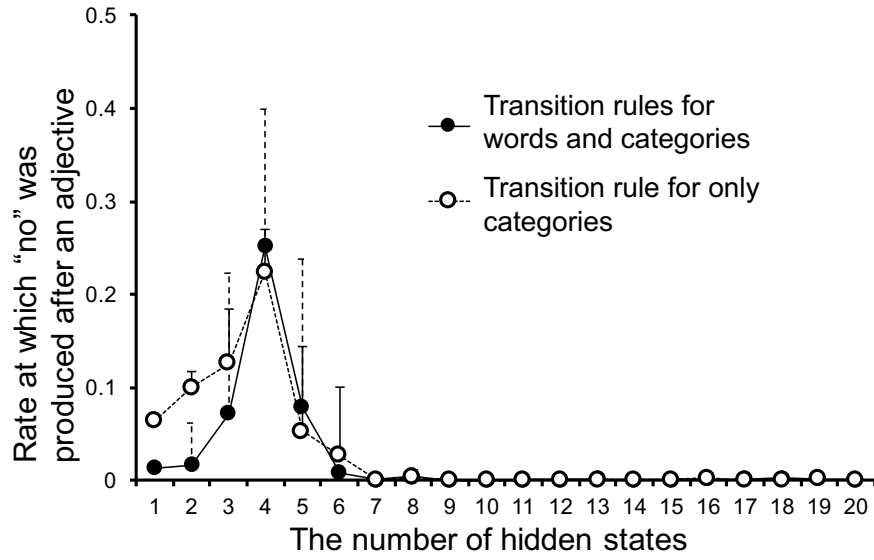
Fig. 4.8 に日本語での過剰生成割合を示す。Fig. 4.8 (a) は各カテゴリ数において形容詞に「の」を付加した割合である。どちらのモデルもカテゴリ数4でピークとなり、カテゴリ数7から全く過剰生成が生じていない。したがって、英語での結果に比べると、少ないカテゴリ数でピークを迎え、すぐに過剰生成が消失するといえる。また、ピーク時の分散も大きい。Fig. 4.8 (b) は名詞・形容詞以外の語の直後に「の」がある割合である。英語の場合と同様に、語遷移規則がカテゴリ遷移規則の誤りを抑制していることがわかる。

Fig. 4.4 に日本語の「の」が生成されたときの、語遷移規則の確率に対するカテゴリ遷移規則の確率の比を示す。英語と同様に、カテゴリの増加にしたがって、「の」の生成にカテゴリ遷移規則が大きく寄与するようになっていくことがわかる。すなわち、カテゴリ遷移規則の変化が「の」の過剰生成の出現と消失の両方に関与しているといえる。

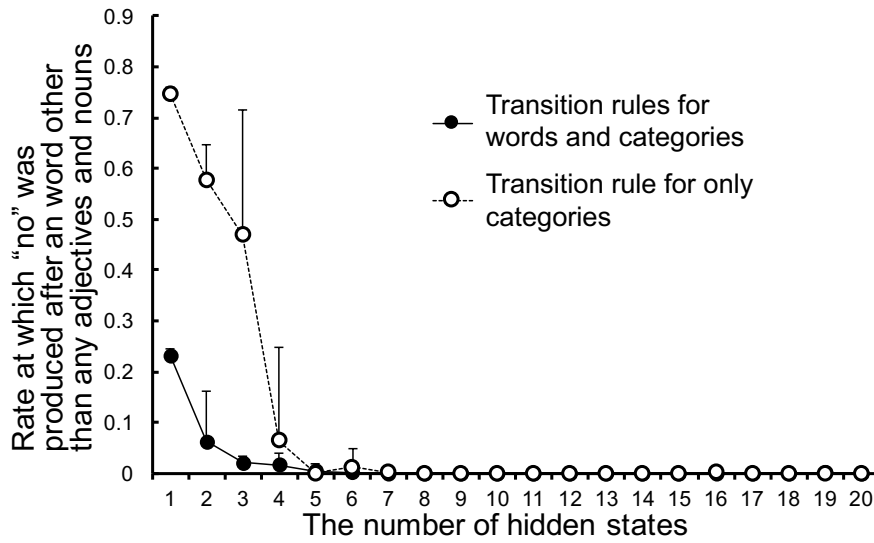
カテゴリ数が4（過剰生成のピーク）のときの10回の試行について、二つのカテゴリ構造に分かれる結果となった。Fig. 4.9 は「の」の過剰生成が生じる構造であり、8/10 試行がこの構造であった。1 番のカテゴリは名詞と形容詞の両方を含んでおり、4 番のカテゴリは格助詞カテゴリである。したがって、1 番のカテゴリから4 番のカテゴリへの遷移により、過剰生成が生じる。一方、2/10 試行では、Fig. 4.10 に示すカテゴリ構造となり、過剰生成はほとんど生じなかった。この場合、1 番のカテゴリが名詞・形容詞と「の」を含み、4 番のカテゴリは「の」以外の格助詞のカテゴリである。よって、「名詞・形容詞 + の」の規則が獲得されず、語遷移規則が優位になり、過剰生成が生じない結果となったことがわかる。

Fig. 4.11 はカテゴリ数が5のときのカテゴリ構造の典型例である。名詞と形容詞が区別され、それぞれ1 番のカテゴリと2 番のカテゴリに表現されたため、過剰生成が生じていない。以上の解析から、「の」の過剰生成が生じるためには、名詞と形容詞が混在したカテゴリと格助詞の明瞭なカテゴリが形成される必要があることがわかった。

次に、過剰生成に寄与する入力を調べるために、コーパスを変化させたときの結果を示す。いずれのコーパス条件でも、カテゴリ数4で過剰生成割合のピークとなった



(a) Adjectives



(b) Words other than nouns and adjectives

Figure 4.8: Results of overproduction in Japanese. Markers indicate rates of overproduction of "no" for (a) adjectives and (b) words other than nouns and adjectives. Error bars indicate standard deviation.

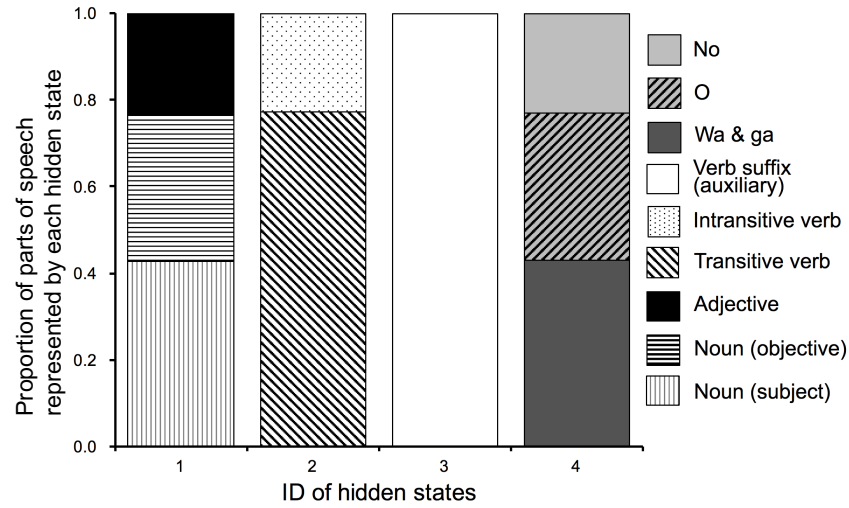


Figure 4.9: Typical representation of four grammatical categories acquired from the Japanese corpus, where the overproduction appeared (its rate was 0.41).

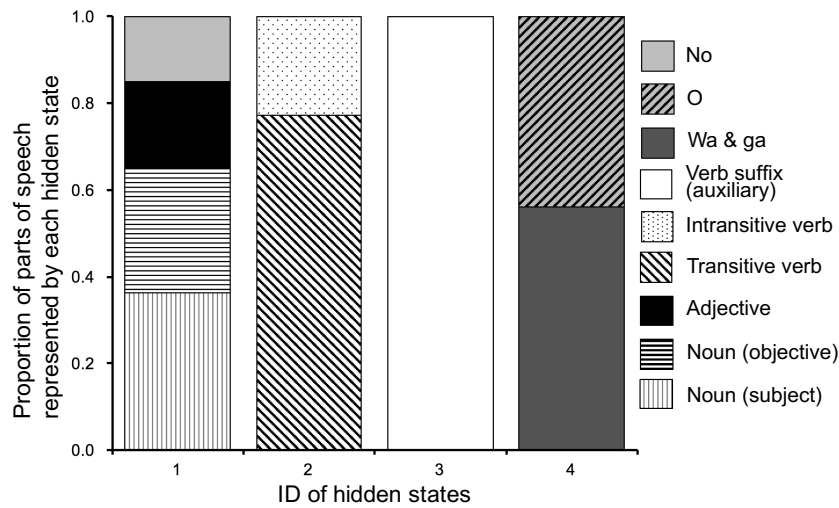


Figure 4.10: Typical representation of four grammatical categories acquired from the Japanese corpus, where the overproduction did not occur (its rate was 0.00074).

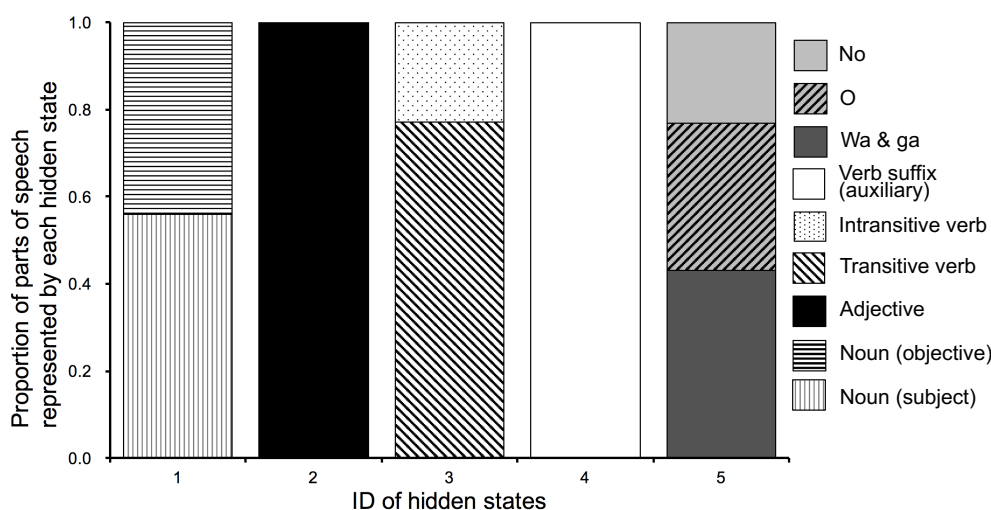


Figure 4.11: Typical representation of five grammatical categories acquired from the Japanese corpus, where the overproduction disappeared.

ため、そのピーク値のみを Fig. 4.12 に示す。左から、「レギュラー」コーパス、「SOV」コーパス、「は・が・を」コーパス、そして、「の」コーパスの結果である。なお、それぞれのコーパスの平均長は 4.7, 4.7, 7.2, 6.1 語である。この結果より、語順が SOV になり、学習文が比較的単純になると過剰生成が増加し、また、格助詞の頻度が多くなり、学習文が比較的複雑になると過剰生成が減少する傾向にあることがわかる。この傾向はコーパスの文長では説明されない。

Fig. 4.13 に「の」コーパスでカテゴリ数 4 のときに獲得されたカテゴリ構造の典型例を示す。「の」の頻度が多いと、「の」をより明にカテゴリ化する必要がある。しかし、モデルのカテゴリ数には限界があるため、格助詞が明瞭に表現されず、名詞や形容詞と混在したカテゴリが獲得される。その結果、有意な規則が学習されず、語遷移規則が優位となったと考えられる。したがって、学習コーパスの文法的な複雑さ、すなわち、表現に必要なカテゴリ数が増加すると、格助詞カテゴリの形成に失敗し、過剰生成が減少したと推測される。以上のことから、「の」の過剰生成は、格助詞カテゴリの形成後に、形容詞・名詞カテゴリが分化するという特定の過程で生じることがわ

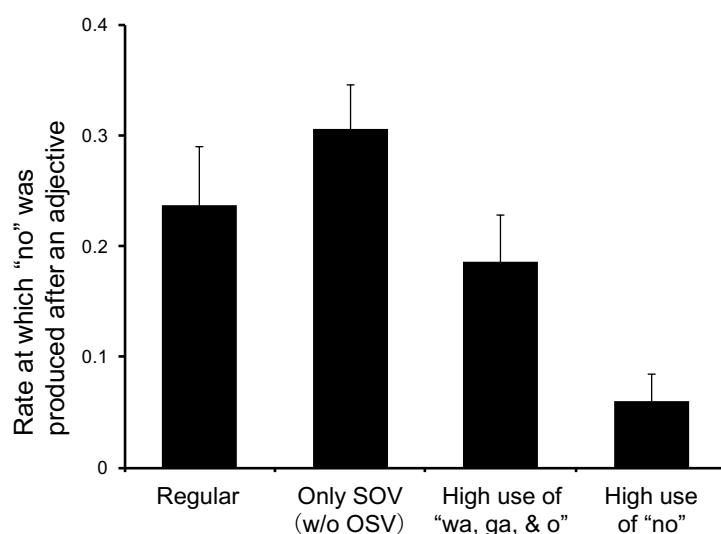


Figure 4.12: Rates of overproduction of “no” in several corpus conditions. “Regular” denotes the corpus used the experiment shown as Fig. 4.8. Error bars indicate standard error.

かった。

Figs. 4.9 – 4.11 からわかるように、「の」は他の格助詞「は・が・を」と同じカテゴリに属す。したがって、「形容詞 + 格助詞」の過剰生成は「の」だけではなく、それ以外の「は・が・を」でも生じ、Fig. 4.8 と同様の曲線が描かれる。

4.5 議論

本研究は、文法カテゴリの増加と精緻化により、英語と日本語における過剰生成が生じるという仮説を、極めて単純な系列学習である HMM を用いて検証した。その結果、全く同一モデルで、英語、または、日本語を母語とする幼児における過剰生成の発生と消失を定性的に再現できた。ただし、英語での過剰生成は減少の傾向を示すものの、完全には消失しなかった。

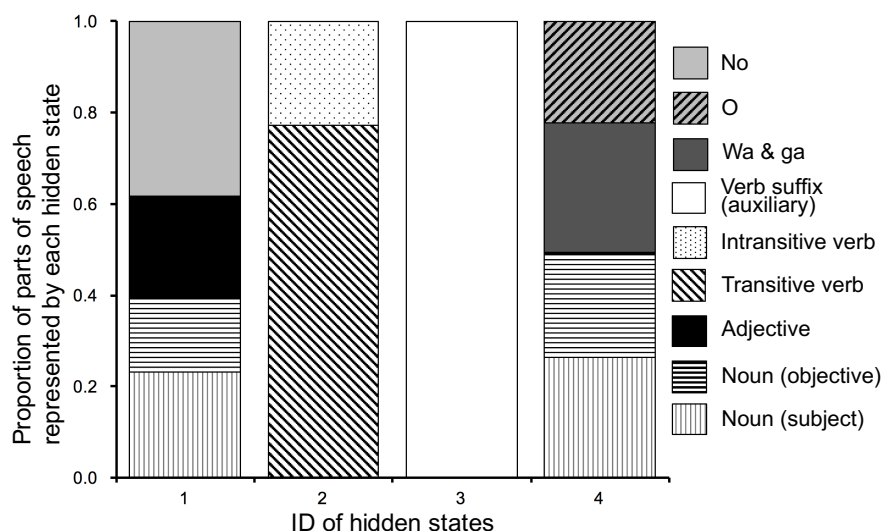


Figure 4.13: A typical representation of grammatical categories which was acquired from the Japanese corpus including more “no”.

4.5.1 英語と日本語の過剰生成に共通するメカニズム

提案モデルは、文法カテゴリ遷移規則と語遷移規則の競合により文を生成する。実験により、それらの規則の役割が明らかとなった。文法カテゴリ遷移規則のカテゴリ数が少ないとき、有意な規則が学習されないため、語遷移規則が優位になる。カテゴリ数が中程度になると、不明瞭なカテゴリから語が生成されるため、入力にない文を創造することがある。これが過剰生成のメカニズムである。一方、語遷移規則は局所的に正しい文法に則った文を生成する。これにより、特にカテゴリ数が少ないときの文法カテゴリ遷移規則から生成される誤りを抑制する。これらの規則が統合されることによって、英語では動詞に対してのみの過剰生成曲線、日本語では形容詞に対してのみの過剰生成曲線が描かれる。

4.1節で述べたように、本モデルは Tomasello [6] の提案する言語発達段階に動機づけられているが、その発達理論は二つの処理を仮定しておらず、提案モデルはそれを反映しているとはいえない。本研究の目的は、この理論の完全なモデル化ではなく、文法カテゴリの獲得がどのように過剰生成に寄与するかの原理の解明であるため、そ

れを明確にするために二つの独立した処理を仮定した。しかし、二つの処理をあらかじめ仮定せずに、言語入力からどのように抽象的な構文が抽出され、語結合から抽象的スキーマへのシームレスな発達的变化をみせる計算モデルは非常に興味深い。近年の言語習得のコネクショニストモデル [85, 86] はこのような仮定を明に設けることなく、複雑な文法を獲得しているが、それらが上記の発達段階を再現できるのかどうかの検証が必要である。

今回、英語の語形変化規則である「動詞語基 + ed」を語列の規則とみなすことで、日本語の統語規則である「名詞 + 格助詞」と同じモデルで再現・比較できるようになった。英語の過剰生成は現在に至るまで、様々なモデルが提案されているが（例えば, [87, 88]），それらのほとんどは英語の過去形における屈折を動詞の音素の変化規則として捉えている。すなわち、モデルは動詞原形の音素列をその過去形の音素列へ変換するように学習する。一方、本研究の結果は、統語規則の習得過程でも「ed」の過剰生成が生じる可能性を指摘するものであり、また、明らかに統語規則の過剰生成と考えられる日本語の「の」の誤用との関連を示唆する。

本実験の結果から、言語入力と過剰生成の新たな仮説を提案できる。英語では、過剰な規則適用の対象（不規則動詞）の頻度を増加させると、その過剰生成割合が低下した。日本語では、文の複雑さを増加させることで同様の結果が得られた。このモデルからの予測を検証するために、実際の養育者・幼児発話コーパスを用いて同様の実験をする必要がある。

今回、完備な行動データの取得は困難であるため、モデル出力を行動データと定量的に比較できておらず、定性的な再現にとどまっている。特に、日本語児の過剰生成の行動データは非常に少なく、さらなる行動データの取得と検証が必要である。本研究の結果は、そのさらなる調査に対する新しい仮説を提供できたことに意味がある。

英語の過去形の規則変化と日本語の格助詞の処理の一部が共通している仮説は、神経科学研究からも支持される。失語症の研究から、左の側頭領域を損傷した患者は規則変化動詞より、不規則変化動詞についての課題の成績が低下することが知られている。一方で、左の前頭領域（シルビウス裂前部）を損傷した患者はその逆の傾向を示す [89]。したがって、「動詞語基 + ed」の規則は前頭領域のいわゆるブローカ野で処

理される。なお、この規則変化動詞と不規則変化動詞の処理の領域的な分離は、2.3節で述べた Words and Rules 理論を支持するものである [26, 90]。日本語についても、fMRI 研究から、同様の前頭領域が格助詞の処理に関わっていることが報告されている [91, 92]。したがって、この脳領域の成熟がそれぞれの言語での過剰生成をもたらしているとし唆される。

モデルの単純さゆえに提案モデルでは再現できなかった現象もある。以降で、それぞれの言語についてのモデルの制限について議論する。

4.5.2 英語での過剰生成

英語の場合、規則変化動詞と不規則変化動詞について未分化なカテゴリが形成された結果、不規則変化動詞にも「ed」が付加された。そして、文法カテゴリを増加させることで、これらが部分的に分化したカテゴリが得られ、「ed」の過剰生成が減少する傾向がみられた。しかし、それらは完全には分離せず、「ed」の過剰生成が完全に消失することはなかった。これは、規則・不規則変化動詞を HMM が完全に区別するほどの統計的エビデンスがなかったためである。Words and Rules 理論では、グラマーの処理を抑制するレキシカルな処理が存在する [26, 90]。レキシカルな処理における語形変化は、「動詞語基 + ed」のような語列の規則ではなく、例えば、「hold → held」のような変化パターンの機械的記憶に基づく。このような語形変化の機械的記憶に基づく処理と、語列規則に基づく提案モデルを競合させることによって、過剰生成が完全に消失する過程を説明できると考えられる。

4.5.3 日本語での過剰生成

日本語の場合、名詞と形容詞が未分化なカテゴリが形成された結果、形容詞の後ろにも「の」が配置された。実験の結果、英語の過剰生成と比較して、日本語の過剰生成は少ないカテゴリ数で生じ、一時的な現象であり、HMM の初期値に敏感であり、また、主に形容詞で発生することがわかった。これは実際の幼児の観察結果である、過剰生成が比較的低月齢（24ヶ月頃）で生じ、数ヶ月で消失し [19, 23]、個人差が大

きく、過剰生成しない幼児もあり [22], また, 形容詞以外への過剰生成はほとんどない [23] ことと一致する.

獲得されたカテゴリの解析から, 「の」の過剰生成には, 名詞と形容詞の混在したカテゴリと, 格助詞の明確なカテゴリが形成されることが必要であることがわかった. このことは, 過剰生成の生じない幼児の有するカテゴリについて, 二つの可能性を示唆する. 一つは, 既に名詞と形容詞の分離したカテゴリを有している可能性である. 名詞と形容詞の混同が過剰生成のメカニズムとする説は伊藤 [23] も議論している. もう一つは, 格助詞カテゴリの獲得に失敗している可能性である. 今回の実験から, 入力の文法要素が多すぎると格助詞カテゴリの獲得に失敗し, 過剰生成が生じないことがわかった. したがって, 見かけ (産出) 上では, 語遷移規則により格助詞を適切に扱えていても, その内部では, 格助詞のカテゴリ化がなされていない可能性が考えられる.

しかし, 提案モデルの産出文では, 「の」だけでなく, 「は・が・を」の格助詞も形容詞の後ろに配置されることがあり, 「の」と同様の過剰生成がみられた. 一方, 実際の幼児では, 「の」の過剰生成が多く報告されている [19]. これは, これらの格助詞の前後には名詞が配置されることが多いということから, これらが同じカテゴリに分類されたことが原因である. しかし, 「の」と「は・が・を」の統語標識としての役割は全く異なる. 今回, 「の」は名詞の連体修飾としての役割, 「は・が・を」は主題役割 (特に, 動作主と被動者) を決める役割がある. このような文での役割・意味を考慮することで, 格助詞の中でさらにカテゴリ化がなされ, 「の」のみの過剰生成が実現できるかもしれない.

意味の考慮に関して, 次のような幼児の誤用も観察されている. 「あかちゃんがつれていく」 [93]. 「しょうぼうしゃがみた」 [94]. どちらも「が」ではなく, 「を」が正しい. 提案モデルでも, 獲得された「が」と「を」は同じカテゴリにあり, これらのような誤用が生じる可能性がある. 主題役割と文法カテゴリの対応の学習の過程で, 上記のような誤用を再現できるかどうかは非常に興味深い.

第5章 討論

提案モデルの限界と今後の展望として、四つの研究の方向性を提案する。

5.1 HMMの妥当性

提案モデルにおいて、文法カテゴリの獲得に重要な構造が隠れマルコフ的状态である。この種の構造は言語関連モデルに特有でなく、運動学習 [95, 96] や心の理論 [97] のような認知モデルや人工知能技術に広く用いられている。このことは提案モデルは言語特異的ではなく、言語特異的構造の存在を仮定する厳格な普遍文法とは対照的であることを示唆する。しかし、提案モデルは非常に単純であり、言語の階層性といった言語の性質の全てを説明できるものではない。ただし、自然言語処理の観点からは、品詞カテゴリ化は言語階層構造の教師なし学習の基礎であるため（例えば, [98]），提案モデルは言語発達初期の基礎的な文法構造とみなせるかもしれない。

また、本研究はHMMで幼児の文法に関する行動的な現象を説明できることを示したが、これが脳の情報処理アルゴリズムとして妥当かどうかを検証する必要がある。近年、生物学的に妥当なスパイクニューラルネットワークがシナプス可塑性によって、HMMと同様に入力系列の隠れ状態を表現できることが示されている [99, 100, 101]。このような神経計算モデルを用いることで、神経科学研究による仮説検証が可能になり、幼児の文法獲得の神経言語的メカニズムの解明につながることを期待される。

5.2 文法カテゴリ発達のトリガ

提案モデルのカテゴリ数は固定であり，試行ごとにその数を変えて実験した．しかし，なぜ，どのように，そのカテゴリ数が増加するのかという疑問は残る．文法カテゴリ発達のトリガについて，入力非依存と入力依存の要因が考えられる．一つは言語入力とはほぼ非依存な脳の成熟である．三歳から五歳の間にも，幼児の脳は急激に発達していく [102, 103, 104]．感覚情報により組織化される神経ネットワークも存在すると考えられるが，脳の成熟については，遺伝的メカニズムも強い要因となるであろう．現状のモデルにおける発達メカニズムは感覚入力に非依存な自動的な脳の成熟とみなせる．前項で述べた，HMM の神経計算モデル [99, 100, 101] において，どのパラメータが隠れ状態に対応し，どのようにすれば隠れ状態が自動的に増加していくのかを発達神経科学の知見を考慮しながら検討する必要がある．

言語入力の要因について，多くの研究が，幼い幼児による発話の動詞構文が，その養育者のものと類似することを示している（例えば，[105, 79]）．よって，養育者の言語入力が，幼児の言語能力の発達に寄与することは広く認められている．幼児の発達段階に応じて，養育者が発話の文法的な難しさを高めているという報告がある [106, 107]．Snow [106] は英語を母語とする二歳児に対する英語の発話と，五歳児に対するものとを比較した．その結果，二歳児に対する発話の平均文長は，五歳児に対するものより短く，複合動詞や従属節の数がより少ないことが明らかとなった．このような文法的に単純化された発話は，幼児にとって扱いやすく，幼児の文法学習を幫助している可能性が議論されている．また，言語入力が徐々に文法的に難しくなっていくことが，文法カテゴリ増加のトリガになる可能性もある．今回用いた HMM を拡張し，入力に適したカテゴリ数を自動決定する無限 HMM [108, 109]（階層ディリクレ過程 HMM [110]）がある．このようなモデルを用いることで，養育者の発話に含まれる文法カテゴリ数を推定でき，上記の仮説の検証が可能になる．

5.3 意味カテゴリの利用と獲得

提案モデルでは、文法カテゴリは語列の統計的性質からのみ学習・推定され、意味の関与を考慮していない。第3章の実験では、入力言語のみから学習した文法カテゴリを介して、新奇語の指す意味カテゴリを推定した。しかし、意味ブートストラッピング仮説は、意味カテゴリによって文法カテゴリが形成される可能性を指摘している [111]。まず、幼児は視覚特徴を物体や動作などの意味カテゴリに分類できることを仮定する。そして、これらの意味カテゴリと語の対応関係から、例えば、語「本」は物体カテゴリ、語「歩く」は動作カテゴリに属することがわかる。このような関係から、ある意味カテゴリに対応する語集合を文法カテゴリとみなすことができる。例えば、物体カテゴリに対応する語集合を名詞カテゴリ、および、動作カテゴリに対応する語集合を動詞カテゴリとすることができる。実際の親子の言語的やりとりにおいて、例えば、全ての行為は動詞により説明されるといったように、対幼児発話中の文法カテゴリと意味カテゴリとの強い共起関係がみられ、意味ブートストラッピングが可能であることが示されている [112]。このようなメカニズムを考慮することで、文法カテゴリと意味カテゴリの双方向的な推定ができ、より頑健で効率的な語意獲得が可能になると考えられる。

ここで、幼児が意味カテゴリをどのように獲得するか、すなわち、どのように連続時間・連続空間の感覚情報から、意味のあるカテゴリを抽出するのかが重要な問いとなる。第3章でのモデルでは、意味カテゴリはあらかじめ設計されており、学習したり、変化したりすることはなかった。それに対して、近年、実世界で動作するロボットがカメラやマイクなどのマルチモーダルなセンサ情報をカテゴリ化することによって、語意を獲得するシステムが提案されている [113]。また、観察したシーンとそれを表す文から、その文を構成する語の意味を学習・推定するシステムも開発されている [114]。なお、この研究でも、文の解析には、HMM とバイグラムが用いられている。また、機械学習分野では、大量の画像とその物体クラスの対応関係を深層学習を用いて学習することで、新奇画像からその物体クラスを推定する課題で、人を超える認識性能が実現できることが報告されている [115]。さらに、画像とそれを説明する

テキストデータを連合する深層学習により、新奇画像を説明する文章を自動生成できるシステムも開発されている [116]. このようなシステムを利用して、意味カテゴリと文法カテゴリを形成しながら、実世界と接地する幼児の言語発達メカニズムを考えていく必要がある.

5.4 言語理解と産出の相互関係

言語理解と産出における実験結果の整合性について述べる. Imai et al. [11] の般用実験は三歳児と五歳児を対象にしており、一方で、英語の「ed」の過剰生成は三歳児頃から現れるとされる. 今回の実験では、言語理解実験での英語を学習する場合の三歳児モデルと五歳児モデルはそれぞれカテゴリ数2と6であった. そのカテゴリ数の間で、言語産出実験では過剰生成の急増がみられた. カテゴリ数6以降でも過剰生成頻度は増加しているが、実際の「ed」の過剰生成が九歳頃まで現れた例もあることから [16, 17], 英語の場合では、言語理解と産出で無矛盾な結果が得られたといえる. 日本語の場合、「の」の過剰生成はより幼い二歳頃にみられる [21, 23]. シミュレーションでは、「の」の過剰生成はカテゴリ数4で最も発生した. このカテゴリ数4は、言語理解実験では三歳から五歳に相当し、言語理解と産出でのカテゴリ数に矛盾が生じている. これはおそらく、言語産出実験で用いた日本語コーパスが、理解実験でのコーパスより複雑であることから、言語産出実験でのカテゴリ数が比較的多くなったためと考えられる. 産出実験でのコーパスは格助詞「の」を含んでいるが、理解実験でのコーパスは「の」を含んでいないため、その表現に必要なカテゴリ数は産出実験でのコーパスのほうが多いといえる. 言語理解と産出を正しく統一的に説明するためには、全く同じコーパスを用いて実験する必要がある. また、実際の対幼児発話コーパスを用いた同様の実験が今後の課題となる. さらに、本来であれば、言語理解によって獲得された語意は言語産出に影響を与え、また、言語産出によって新たな語意獲得につながると思われる. 提案モデルではこのような言語理解と産出が相互に影響を与えることを考慮していないが、この相互フィードバックにより言語を獲得するモデルの構築が今後の課題となる.

第6章 結論

本論文では、複数の課題・言語・年齢についての実験的証拠を現象論的ではあるが、統一的に説明することのできる文法カテゴリ獲得モデルを提案した。文法カテゴリ獲得に関する従来研究では、機械学習としての効率性や、幼児の言語発達の一側面しか考慮されていなかった。それに対し、本論文では提案モデルを言語理解と産出の発達的变化に関して多言語的に検討することで、その妥当性を検証することで、その変化の背後にある文法カテゴリ獲得メカニズムについて新しい洞察を得ることができた。

第3章での言語理解については、Imai et al. [11] の名詞・動詞般用課題に注目し、文法カテゴリの獲得に基づいて、その言語間差異と発達的变化を説明できるモデルを提案した。異なる数の状態数を設定することで、モデルは異なる年齢の幼児の課題成績を表現することができた。また、モデルが言語固有の文法カテゴリを獲得することで、五歳児でみられる言語間差異を再現できた。その結果、提案モデルによる語意推定は幼児の課題成績と同様の傾向を得た。英語は安定な統語規則を有しているため、英語の入力から獲得した文法カテゴリは統語手がかりに強く依存していた。その結果、項が省略されて、統語手がかりがなくなると、英語を学習したモデルは新奇動詞を動作に対応できなくなった。対照的に、日本語は統語手がかりが頻繁に欠落するため、モデルは形態的手がかりに強く依存し、項省略条件でも動詞を般用できた。また、中国語はそのような形態的手がかりを持たないため、モデルは項省略条件では動詞を般用できなかった。さらに、HMM の隠れ状態数を漸次的に増加させて、そのモデルが獲得する文法カテゴリを示すことで、幼児の文法カテゴリの獲得過程に関する仮説を提案した。

第4章での言語産出の実験では、HMM とトライグラムという非常に単純、かつ、言語固有でないモデルを用いて、英語と日本語の過剰生成の現象を説明することを試

みた。この試みは、このような語順しか考慮しない単純な統計モデルで、どこまで幼児の発達を説明可能かという挑戦に言い換えられる。シミュレーションにより、過剰生成現象の核である過剰生成の出現を再現できた。一方で、レキシカルな、および、セマンティックなメカニズムの欠落によって、再現できない現象（英語の過剰生成の完全な消失と、日本語の「の」以外の過剰生成の抑制）も明らかになった。しかしながら、このような単純なモデルから構成し、その限界を慎重に議論していくことによって、言語発達のような複数のシステムが複雑に入り組んだ現象を、構成的な視点から理解できるようになると思われる。特に、日本語の「の」の過剰生成が生じない場合、格助詞が正しく使えているように見えるが、実は格助詞カテゴリが精緻に形成されていないことは興味深い発見といえる。この仮説をさらに検証できるような発達心理学的実験が望まれる。

関連図書

- [1] M. D. S. Braine and M. Bowerman. Children's first word combinations. *Monographs of the Society for Research in Child Development*, Vol. 41, No. 1, pp. 1–104, 1976.
- [2] R. Olguin and M. Tomasello. Twenty-five-month-old children do not have a grammatical category of verb. *Cognitive Development*, Vol. 8, No. 3, pp. 245–272, 1993.
- [3] M. Tomasello and P. J. Brooks. Young children's earliest transitive and intransitive constructions. *Cognitive Linguistics*, Vol. 4, No. 9, pp. 370–395, 1998.
- [4] M. Tomasello. Do young children have adult syntactic competence? *Cognition*, Vol. 74, No. 3, pp. 209–253, 2000.
- [5] M. Tomasello. The item-based nature of children's early syntactic development. *Trends in Cognitive Sciences*, Vol. 4, No. 4, pp. 156–163, 2000.
- [6] M. Tomasello. *Constructing a Language: A Usage-Based Theory of Language Acquisition*. Cambridge: Harvard University Press, 2003.
- [7] L. Gleitman. The structural sources of verb meanings. *Language Acquisition*, Vol. 1, No. 1, pp. 3–55, 1990.
- [8] L. Naigles. Children use syntax to learn verb meanings. *Journal of Child Language*, Vol. 17, No. 2, pp. 357–374, 1990.

- [9] C. Fisher. Structural limits on verb mapping: The role of analogy in children's interpretations of sentences. *Cognitive Psychology*, Vol. 31, No. 1, pp. 41–81, 1996.
- [10] C. Fisher, Y. Gertner, R. M. Scott, and S. Yuan. Syntactic bootstrapping. *Wiley Interdisciplinary Reviews: Cognitive Science*, Vol. 1, No. 2, pp. 143–149, 2010.
- [11] M. Imai, L. Li, E. Haryu, H. Okada, K. Hirsh-Pasek, R. M. Golinkoff, and J. Shigematsu. Novel noun and verb learning in Chinese-, English-, and Japanese-speaking children. *Child Development*, Vol. 79, No. 4, pp. 979–1000, 2008.
- [12] M. Imai, E. Haryu, and H. Okada. Mapping novel nouns and verbs onto dynamic action events: Are verb meanings easier to learn than noun meanings for Japanese children? *Child Development*, Vol. 76, No. 2, pp. 340–355, 2005.
- [13] J. Berko. The child's learning of English morphology. *Word*, Vol. 14, No. 2-3, pp. 150–177, 1958.
- [14] S. Ervin. Imitation and structural change in children's language. In E. Lenneberg, editor, *New Directions in the Study of Language*. Cambridge, MA: MIT Press, 1964.
- [15] G. F. Marcus, S. Pinker, M. Ullman, M. Hollander, T. J. Rosen, F. Xu, and H. Clahsen. Overregularization in language acquisition. *Monographs of the Society for Research in Child Development*, Vol. 57, No. 4, pp. 1–178, 1992.
- [16] J. L. Bybee and D. I. Slobin. Rules and schemas in the development and use of the English past tense. *Language*, Vol. 58, No. 2, pp. 265–289, 1982.
- [17] V. Marchman. Rules and regularities in the acquisition of the English past tense. *Center for Research in Language Newsletter*, Vol. 2, No. 4, 1988.

- [18] P. M. Clancy. The acquisition of Japanese. In D. I. Slobin, editor, *The Crosslinguistic Study of Language Acquisition Volume 1: The Data*, pp. 373–524. Hillsdale, NJ: Lawrence Erlbaum Associations, 1985.
- [19] 横山正幸. 幼児による助詞の誤用の出現時期と類型. 福岡教育大学紀要, Vol. 38, pp. 225–236, 1989.
- [20] K. Murasugi. *Noun phrases in Japanese and English: A study in syntax, learnability and acquisition*. PhD thesis, University of Connecticut, 1991.
- [21] 横山正幸. 幼児の連体修飾発話における助詞「ノ」の誤用. 発達心理学研究, Vol. 1, No. 1, pp. 2–9, 1990.
- [22] 伊藤友彦. 過剰生成される「ノ」の統語カテゴリー 幼児一例の縦断研究. 東京学芸大学紀要. 第1部門, 教育科学, Vol. 49, pp. 143–149, 1998.
- [23] 伊藤友彦. 幼児2例に生じた「ノ」の過剰生成の出現・消失メカニズム. 東京学芸大学紀要. 第1部門, 教育科学, Vol. 50, pp. 159–168, 1999.
- [24] 永野賢. 幼児の言語発達 –とくに助詞「の」の習得過程について–. 関西大学国文学会：島田教授古希記念国文学論集, pp. 405–418, 1960.
- [25] S. Pinker and A. Prince. Regular and irregular morphology and the psychological status of rules of grammar. In S. D. Lima, R. Corrigan, and G. K. Iverson, editors, *The Reality of Linguistic Rules*. John Benjamins Amsterdam, 1994.
- [26] S. Pinker. *Words and Rules: The Ingredients of Language*. HarperCollins, 1999.
- [27] C. D. Manning and H. Schütze. *Foundations of Statistical Natural Language Processing*. Cambridge, US: MIT Press, 1999.
- [28] J. L. Elman. Finding structure in time. *Cognitive Science*, Vol. 14, No. 2, pp. 179–211, 1990.

- [29] P. F. Brown, P. V. Desouza, R. L. Mercer, V. J. D. Pietra, and J. C. Lai. Class-based n-gram models of natural language. *Computational Linguistics*, Vol. 18, No. 4, pp. 467–479, 1992.
- [30] S. Finch and N. Chater. Bootstrapping syntactic categories using statistical methods. In W. Daelemans and D. Powers, editors, *Background and Experiments in Machine Learning of Natural Language*, pp. 229–235. Tilburg University, Institute for Language Technology and AI, 1992.
- [31] H. Schütze. Part-of-speech induction from scratch. In *Proceedings of the 31st Annual Meeting on Association for Computational Linguistics*, pp. 251–258, 1993.
- [32] M. Redington, N. Chater, and S. Finch. Distributional information: A powerful cue for acquiring syntactic categories. *Cognitive Science*, Vol. 22, No. 4, pp. 425–469, 1998.
- [33] A. Clark. Inducing syntactic categories by context distribution clustering. In *Proceedings of the 2nd Workshop on Learning language in Logic and the 4th Conference on Computational Natural Language Learning*, pp. 91–94, 2000.
- [34] T. H. Mintz. Frequent frames as a cue for grammatical categories in child directed speech. *Cognition*, Vol. 90, No. 1, pp. 91–117, 2003.
- [35] C. Parisien, A. Fazly, and S. Stevenson. An incremental Bayesian model for learning syntactic categories. In *Proceedings of the 12th Conference on Computational Natural Language Learning*, pp. 89–96, 2008.
- [36] A. Alishahi and G. Chrupała. Lexical category acquisition as an incremental process. In *CogSci 2009 Workshop on Psychocomputational Models of Human Language Acquisition*, 2009.

- [37] L. Onnis and M. H. Christiansen. Lexical categories at the edge of the word. *Cognitive Science*, Vol. 32, No. 1, pp. 184–221, 2008.
- [38] L. M. Santelmann and P. W. Jusczyk. Sensitivity to discontinuous dependencies in language learners: Evidence for limitations in processing space. *Cognition*, Vol. 69, No. 2, pp. 105–134, 1998.
- [39] A. Marquis and R. Shi. Initial morphological learning in preverbal infants. *Cognition*, Vol. 122, No. 1, pp. 61–66, 2012.
- [40] T. H. Mintz. The segmentation of sub-lexical morphemes in English-learning 15-month-olds. *Frontiers in Psychology*, Vol. 4, No. 24, 2013.
- [41] P. Monaghan, M. H. Christiansen, and N. Chater. The phonological-distributional coherence hypothesis: Cross-linguistic evidence in language acquisition. *Cognitive Psychology*, Vol. 55, No. 4, pp. 259–305, 2007.
- [42] F. T. Asr, A. Fazly, and Z. Azimifar. The effect of word-internal properties on syntactic categorization: A computational modeling approach. In *Proceedings of the 32nd Annual Conference of the Cognitive Science Society*, pp. 1529–1535, 2010.
- [43] J. M. Siskind. A computational study of cross-situational techniques for learning word-to-meaning mappings. *Cognition*, Vol. 61, No. 1, pp. 39–91, 1996.
- [44] L. Smith and C. Yu. Infants rapidly learn word-referent mappings via cross-situational statistics. *Cognition*, Vol. 106, No. 3, pp. 1558–1568, 2008.
- [45] M. Frank, N. D. Goodman, and J. B. Tenenbaum. A Bayesian framework for cross-situational word-learning. In *Advances in Neural Information Processing Systems*, pp. 457–464, 2007.

- [46] A. Fazly, A. Alishahi, and S. Stevenson. A probabilistic computational model of cross-situational word learning. *Cognitive Science*, Vol. 34, No. 6, pp. 1017–1063, 2010.
- [47] A. Alishahi and A. Fazly. Integrating syntactic knowledge into a model of cross-situational word learning. In *Proceedings of the 32nd Annual Conference of the Cognitive Science Society*, pp. 2452–2457, 2010.
- [48] 下斗米貴之, 遠山修治, 大森隆司. 文法メタ知識による語彙学習加速のコネクショニストモデル. *認知科学*, Vol. 10, No. 1, pp. 104–111, 2003.
- [49] A. Toyomura and T. Omori. A computational model for taxonomy-based word learning inspired by infant developmental word acquisition. *IEICE Transactions on Information and Systems*, Vol. 88, No. 10, pp. 2389–2398, 2005.
- [50] C. Yu. Learning syntax-semantics mappings to bootstrap word learning. In *Proceedings of the 28th Annual Conference of the Cognitive Science Society*, pp. 924–929, 2006.
- [51] A. Alishahi and G. Chrupala. Concurrent acquisition of word meaning and lexical categories. In *Proceedings of the 2012 Joint Conference on Empirical Methods in Natural Language Processing and Computational Natural Language Learning*, pp. 643–654, 2012.
- [52] D. E. Rumelhart, G. E. Hinton, and R. J. Williams. Learning representations by back-propagating errors. *Nature*, Vol. 323, No. 6088, pp. 533–536, 1986.
- [53] K. Plunkett and V. Marchman. U-shaped learning and frequency effects in a multi-layered perception: Implications for child language acquisition. *Cognition*, Vol. 38, No. 1, pp. 43–102, 1991.

- [54] K. Plunkett and V Marchman. From rote learning to system building: Acquiring verb morphology in children and connectionist nets. *Cognition*, Vol. 48, No. 1, pp. 21–69, 1993.
- [55] W. V. O. Quine. *Word and Object*. MA: MIT Pres, 1960.
- [56] S. Carey and E. Bartlett. Acquiring a single new word. *Papers and Reports on Child Language Development*, Vol. 15, pp. 17–29, 1978.
- [57] E. M. Markman and G. F. Wachtel. Children’s use of mutual exclusivity to constrain the meanings of words. *Cognitive Psychology*, Vol. 20, No. 2, pp. 121–157, 1988.
- [58] E. M. Markman. *Categorization and Naming in Children: Problems of Induction*. MA: MIT Press, 1991.
- [59] M. Imai, S. Kita, M. Nagumo, and H. Okada. Sound symbolism facilitates early verb learning. *Cognition*, Vol. 109, No. 1, pp. 54–65, 2008.
- [60] B. MacWhinney. *The CHILDES Project: Tools for Analyzing Talk (3rd ed., Vol. 2: The Database)*. MahWah NJ: Lawrence Erlbaum Associates, 2000.
- [61] S. Miyata. *MiiPro Nanami Corpus*. PA: TalkBank, 2013. ISBN 1-59642-473-7.
- [62] A. Henry. *English Belfast Corpus*. PA: TalkBank, 2004. ISBN 1-59642-037-5.
- [63] T. Tardif. *Chinese Beijing2 Corpus*. PA: TalkBank, 2007. ISBN 1-59642-287-4.
- [64] T. Cameron-Faulkner, E. Lieven, and M. Tomasello. A construction based analysis of child directed speech. *Cognitive Science*, Vol. 27, No. 6, pp. 843–873, 2003.
- [65] M. S. Seidenberg. Language acquisition and use: Learning and applying probabilistic constraints. *Science*, Vol. 275, No. 5306, pp. 1599–1603, 1997.

- [66] M. S. Seidenberg and M. C. MacDonald. A probabilistic constraints approach to language acquisition and processing. *Cognitive Science*, Vol. 23, No. 4, pp. 569–588, 1999.
- [67] B. MacWhinney. The competition model. In B. MacWhinney, editor, *Mechanisms of Language Acquisition*, pp. 249–308. Hillsdale, NJ: Erlbaum, 1987.
- [68] E. Bates and B. MacWhinney. Functionalism and the competition model. In B. MacWhinney and E. Bates, editors, *The Crosslinguistic Study of Sentence Processing*, pp. 3–73. Cambridge, UK: Cambridge University Press, 1989.
- [69] C. Fisher. Structural limits on verb mapping: the role of abstract structure in 2.5-year-olds’ interpretations of novel verbs. *Developmental Science*, Vol. 5, No. 1, pp. 55–64, 2002.
- [70] J. N. Lee and L. R. Naigles. Mandarin learners use syntactic bootstrapping in verb acquisition. *Cognition*, Vol. 106, No. 2, pp. 1028–1037, 2008.
- [71] S. Yuan and C. Fisher. “Really? She blicked the baby?” Two-year-olds learn combinatorial facts about verbs by listening. *Psychological Science*, Vol. 20, No. 5, pp. 619–626, 2009.
- [72] A. M. Leslie and S. Keeble. Do six-month-old infants perceive causality? *Cognition*, Vol. 25, No. 3, pp. 265–288, 1987.
- [73] R. Saxe, J. B. Tenenbaum, and S. Carey. Secret agents inferences about hidden causes by 10-and 12-month-old infants. *Psychological Science*, Vol. 16, No. 12, pp. 995–1001, 2005.
- [74] D. Gentner. Why nouns are learned before verbs: Linguistic relativity versus natural partitioning. In S. A. Kuczaj, editor, *Language Development: Language, Thought, and Culture*, pp. 301–334. Erlbaum, Hillsdale, 1982.

- [75] D. Gentner and L. Boroditsky. Individuation, relativity, and early word learning. In M. Bowerman and S. C. Levinson, editors, *Language Acquisition and Conceptual Development*, pp. 215–256. Cambridge, UK: Cambridge University Press, 2001.
- [76] M. H. Bornstein, L. R. Cote, S. Maital, K. Painter, S.-Y. Park, L. Pascual, M.-G. Pêcheux, J. Ruel, P. Venuti, and A. Vyt. Cross-linguistic analysis of vocabulary in young children: Spanish, Dutch, French, Hebrew, Italian, Korean, and American English. *Child Development*, Vol. 75, No. 4, pp. 1115–1139, 2004.
- [77] T. Ogura, P. S. Dale, Y. Yamashita, T. Murase, and A. Mahieu. The use of nouns and verbs by Japanese children and their caregivers in book-reading and toy-playing contexts. *Journal of Child Language*, Vol. 33, No. 1, p. 1, 2006.
- [78] M. J. Maguire, K. Hirsh-Pasek, and R. M. Golinkoff. A unified theory of word learning: putting verb acquisition in context. In K. Hirsh-Pasek and R. M. Golinkoff, editors, *Action Meets Word: How Children Learn Verbs*, pp. 364–391. New York: Oxford University Press, 2006.
- [79] S. Choi and A. Gopnik. Early acquisition of verbs in Korean: A cross-linguistic study. *Journal of Child Language*, Vol. 22, No. 3, pp. 497–529, 1995.
- [80] T. Tardif, M. Shatz, and L. Naigles. Caregiver speech and children’s use of nouns versus verbs: A comparison of English, Italian, and Mandarin. *Journal of Child Language*, Vol. 24, No. 3, pp. 535–565, 1997.
- [81] S. R. Waxman, J. L. Lidz, I. E. Braun, and T. Lavin. Twenty four-month-old infants’ interpretations of novel verbs and nouns in dynamic scenes. *Cognitive Psychology*, Vol. 59, No. 1, pp. 67–95, 2009.

- [82] M. Tomasello and P. J. Brooks. Early syntactic development: A construction grammar approach. In M. Barrett, editor, *The Development of Language*. London: University College London Press, 1999.
- [83] S. Goldwater and T. Griffiths. A fully bayesian approach to unsupervised part-of-speech tagging. In *Proceedings of the Annual Meeting on Association for Computational Linguistics*, 2007.
- [84] 池田彩夏, 小林哲生, 板倉昭二. 日本語母語話者の対乳幼児発話における格助詞省略. *認知科学*, Vol. 23, No. 1, pp. 8–21, 2016.
- [85] G. S. Dell and F. Chang. The P-chain: Relating sentence production and its disorders to comprehension and acquisition. *Philosophical Transactions of the Royal Society B*, Vol. 369, No. 20120394, 2014.
- [86] F. Chang. Learning to order words: A connectionist model of heavy NP shift and accessibility effects in Japanese and English. *Journal of Memory and Language*, Vol. 61, No. 3, pp. 374–397, 2009.
- [87] A. M. Woollams, M. Joanisse, and K. Patterson. Past-tense generation from form versus meaning: Behavioural data and simulation evidence. *Journal of Memory and Language*, Vol. 61, No. 1, pp. 55–76, 2009.
- [88] G. Westermann and N. Ruh. A neuroconstructivist model of past tense development and processing. *Psychological Review*, Vol. 119, No. 3, p. 649, 2012.
- [89] M. T. Ullman, R. Pancheva, T. Love, E. Yee, D. Swinney, and G. Hickok. Neural correlates of lexicon and grammar: Evidence from the production, reading, and judgment of inflection in aphasia. *Brain and Language*, Vol. 93, No. 2, pp. 185–238, 2005.
- [90] S. Pinker and M. T. Ullman. The past and future of the past tense. *Trends in Cognitive Sciences*, Vol. 6, No. 11, pp. 456–463, 2002.

- [91] T. Inui, K. Ogawa, and M. Ohba. Role of left inferior frontal gyrus in the processing of particles in Japanese. *Neuroreport*, Vol. 18, No. 5, pp. 431–434, 2007.
- [92] Y. Hashimoto, S. Yokoyama, and R. Kawashima. Neural differences in processing of case particles in Japanese: an fMRI study. *Brain and Behavior*, Vol. 4, No. 2, pp. 180–186, 2014.
- [93] 伊藤克敏. こどものことば – 習得創造. 勁草書房, 1990.
- [94] 横山正幸. 文法の獲得 2 – 助詞を中心に –. 小林春美, 佐々木正人 (編), 子どもの言語獲得, pp. 131–151. 大修館書店, 1997.
- [95] T. Inamura, I. Toshima, H. Tanie, and Y. Nakamura. Embodied symbol emergence based on mimesis theory. *International Journal of Robotics Research*, Vol. 23, No. 4-5, pp. 363–377, 2004.
- [96] T. Taniguchi, K. Hamahata, and N. Iwahashi. Unsupervised segmentation of human motion data using a sticky hierarchical Dirichlet process-hidden Markov model and minimal description length-based chunking method for imitation learning. *Advanced Robotics*, Vol. 25, No. 17, pp. 2143–2172, 2011.
- [97] C. L. Baker, R. Saxe, and J. B. Tenenbaum. Action understanding as inverse planning. *Cognition*, Vol. 113, No. 3, pp. 329–349, 2009.
- [98] D. Klein and C. D. Manning. A generative constituent-context model for improved grammar induction. In *Proceedings of the 40th Annual Meeting on Association for Computational Linguistics*, pp. 128–135, 2002.
- [99] J. Brea, W. Senn, and J.-P. Pfister. Sequence learning with hidden units in spiking neural networks. In *Advances in Neural Information Processing Systems*, pp. 1422–1430, 2011.

- [100] D. J. Rezende, D. Wierstra, and W. Gerstner. Variational learning for recurrent spiking networks. In *Advances in Neural Information Processing Systems*, pp. 136–144, 2011.
- [101] D. Kappel, B. Nessler, and W. Maass. Stdp installs in winner-take-all circuits an online approximation to hidden Markov model learning. *PLoS Computational Biology*, Vol. 10, No. 3, p. e1003511, 2014.
- [102] K. L. Sakai. Language acquisition and brain development. *Science*, Vol. 310, No. 5749, pp. 815–819, 2005.
- [103] P. M. Thompson, J. N. Giedd, R. P. Woods, D. MacDonald, A. C. Evans, and A. W. Toga. Growth patterns in the developing brain detected by using continuum mechanical tensor maps. *Nature*, Vol. 404, No. 6774, pp. 190–193, 2000.
- [104] J. Matsuzawa, M. Matsui, T. Konishi, K. Noguchi, R. C. Gur, W. Bilker, and T. Miyawaki. Age-related volumetric changes of brain gray and white matter in healthy infants and children. *Cerebral Cortex*, Vol. 11, No. 4, pp. 335–342, 2001.
- [105] A. L. Theakston, E. V. M. Lieven, J. M. Pine, and C. F. Rowland. The role of performance limitations in the acquisition of verb-argument structure: An alternative account. *Journal of Child Language*, Vol. 28, No. 1, pp. 127–152, 2001.
- [106] C. E. Snow. Mothers’ speech to children learning language. *Child Development*, Vol. 43, No. 2, pp. 549–565, 1972.
- [107] J. R. Phillips. Syntax and vocabulary of mothers’ speech to young children: Age and sex comparisons. *Child Development*, Vol. 44, No. 1, pp. 182–185, 1973.

- [108] M. J. Beal, Z. Ghahramani, and C. E. Rasmussen. The infinite hidden Markov model. In *Advances in Neural Information Processing Systems*, pp. 577–584, 2001.
- [109] J. Van Gael, A. Vlachos, and Z. Ghahramani. The infinite HMM for unsupervised PoS tagging. In *Proceedings of the Conference on Empirical Methods in Natural Language Processing*, pp. 678–687, 2009.
- [110] Y. W. Teh, M. I. Jordan, M. J. Beal, and D. M. Blei. Hierarchical Dirichlet processes. *Journal of the American Statistical Association*, Vol. 101, No. 476, pp. 1566–1581, 2012.
- [111] S. Pinker. *Language Learnability and Language Development*. Harvard University Press, 1984.
- [112] J. A. Rondal and A. Cession. Input evidence regarding the semantic bootstrapping hypothesis. *Journal of Child Language*, Vol. 17, No. 3, pp. 711–717, 1990.
- [113] T. Nakamura, T. Araki, T. Nagai, and N. Iwahashi. Grounding of word meanings in latent Dirichlet allocation-based multimodal concepts. *Advanced Robotics*, Vol. 25, No. 17, pp. 2189–2206, 2011.
- [114] M. Attamimi, Y. Ando, T. Nakamura, T. Nagai, D. Mochihashi, I. Kobayashi, and H. Asoh. Learning word meanings and grammar for verbalization of daily life activities using multilayered multimodal latent Dirichlet allocation and Bayesian hidden Markov models. *Advanced Robotics*, Vol. 30, No. 11-12, pp. 806–824, 2016.
- [115] K. He, X. Zhang, X. Ren, and J. Sun. Delving deep into rectifiers: Surpassing human-level performance on ImageNet classification. In *Proceedings of the IEEE International Conference on Computer Vision*, pp. 1026–1034, 2015.

- [116] O. Vinyals, A. Toshev, S. Bengio, and D. Erhan. Show and tell: A neural image caption generator. In *Proceedings of the IEEE International Conference on Computer Vision and Pattern Recognition*, pp. 3156–3164, 2015.
- [117] B. Merialdo. Tagging English text with a probabilistic model. *Computational Linguistics*, Vol. 20, No. 2, pp. 155–171, 1994.
- [118] S. Geman and D. Geman. Stochastic relaxation, Gibbs distributions, and the Bayesian restoration of images. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, No. 6, pp. 721–741, 1984.

付 録 A 付 録

A.1 行動データの引用

第3章での語意推定実験では、シミュレーション結果と実際の幼児の行動を比較するために、Imai et al. [11] の実験データを用いた。今回、その論文中の日本語 (Study 1, Table 3)、英語 (Study 2, Table 3)、および、中国語を母語とする幼児の実験結果 (Study 5, Table 5) を引用した。日本語と英語での実験 (Study 1 と Study 2) では、幼児が新奇物体を視認していることを保証するために、刺激動画の開始時に、人が物体をおよそ 0.5 秒保持したまま静止している。その後、その人がその物体を用いた新奇動作を始める。この物体保持期間により、物体が強調される効果が現れ、中国語児は五歳であっても新奇動詞を物体へ対応付ける傾向がみられた (Study 3)。そのため、Study 5 では、この物体保持期間のない刺激動画を提示することで、中国語を母語とする五歳児は新奇動詞を動作へ対応づけられるようになることが観察された。一方、三歳の日本語児と英語児で同様の物体保持期間のない刺激動画で実験 (Study 6) すると、物体保持期間のある刺激動画の結果 (Study 1 と Study 2) と有意な差がないことが明らかとなった。Study 3 と Study 5 の差異に関して、Imai et al. [11] は、中国語は英語や日本語よりも名詞と動詞を区別する形態的な手がかりが少ないため、中国語児は知覚的な (言語外の) 手がかりを重視しているという可能性を議論している。

本研究での提案モデルはその言語外の手がかりではなく、言語的な (統語形態論的な) 手がかりによる語意推定を説明することを目指している。したがって、今回、中国語児のデータとして、物体保持期間のない Study 5 の結果を用いた。可能であれば、日本語児と英語児のデータとして、Study 5 と同じ刺激を用いた Study 6 の結果を引用すべきであるが、Study 6 は三歳児のみの実験であり、五歳児のデータがない。そ

のため、Study 1 と Study 2 の結果をそれぞれ日本語児と英語児の行動データとした。しかし、上述の通り、Study 6 から日本語と英語を母語とする三歳児には刺激動画の違いによる般用の有意差がないことがわかっており、このデータ選択についての大きな問題はないと考えられる。

A.2 語意推定モデルの学習法

提案モデルは式 (3.1) における $P(s_t|w_t, \mathbf{s}_{-t})$ と $P(m|s_t)$ を学習する。これらの確率はそれぞれ、語 w_t の文法カテゴリ s_t の推定と、その s_t と意味カテゴリ m との連合を表す。幼児は例えば文法カテゴリといった文法的知識を明に教わらないにもかかわらず、日常的な言語コミュニケーションから統語的・語意的知識を獲得する。したがって、提案モデルもこれらの確率を教師信号なしに、語系列とそれが説明するシーンの対から獲得するものとする。今回、モデルに与えられるシーン内にある視覚特徴と語は、離散変数に分割されていることを仮定する。

A.2.1 文法カテゴリの獲得

本モデルでは、 $P(s_t|w_t, \mathbf{s}_{-t})$ を計算するために、隠れマルコフモデル (hidden Markov model: HMM) を用いる。これは時系列入力に潜在する隠れ状態列 (マルコフ連鎖) を求める教師なし学習モデルである。その入力を語系列 (文章) にすると、獲得される隠れ状態は結果的に語の品詞を表し、それら状態の遷移規則は簡単な統語構造を表すことが知られている。そのため、HMM は語の品詞タグ付けの教師なし確率モデルとしてよく用いられている [27, 117]。ベイジアン HMM (Bayesian HMM: BHMM) は HMM にベイズ推定を導入したモデルであり、特に小規模なコーパスの場合に、品詞タグ付けの精度が向上することが示されている [83]。ただし、それらのモデルの基本的な機能は同じである。

BHMM のグラフィカルモデルを Fig. A.2.1 に示す。このモデルは、系列的な隠れ状態 $\mathbf{s} = (s_1, s_2, \dots, s_t, \dots)$ における各 s_t が確率的に語 w_t を生成し、語の時系列

$\mathbf{w} = (w_1, w_2, \dots, w_t, \dots)$ を構成することを仮定している．さらに今回， $\mathbf{s}$ がトライグラム構造を有している，すなわち，隠れ状態 s_t は過去の二状態 (s_{t-1} と s_{t-2}) と将来の二状態 (s_{t+1} と s_{t+2}) に依存することを仮定する．そのモデルの詳細な定義を次に示す．

$$s_t | s_{t-2} = s, s_{t-1} = s', s_{t+1} = s'', s_{t+2} = s''', \tau^{(s, s', s'', s''')} \sim \text{Mult} \left(\tau^{(s, s', s'', s''')} \right) \quad (\text{A.1})$$

$$w_t | s_t = s, \omega^{(s)} \sim \text{Mult} \left(\omega^{(s)} \right) \quad (\text{A.2})$$

$$\tau^{(s, s', s'', s''')} | \alpha \sim \text{Dir}(\alpha) \quad (\text{A.3})$$

$$\omega^{(s)} | \beta \sim \text{Dir}(\beta) \quad (\text{A.4})$$

ここで， $\tau^{(s, s', s'', s''')}$ と $\omega^{(s)}$ はそれぞれハイパーパラメータ α と β を持つディリクレ事前分布である．なお，本文中の Fig. 3.1 では，特徴 o の推定のための $P(o|w_t)$ と $P(o|m)P(m|s_t)$ の矢印が描かれているが，BHMM による隠れ状態の学習と推定には， o と $\mathbf{f}$ ，および，意味カテゴリ m を用いない．

このモデル仮定の下，ギブスサンプリング [118] により， $\mathbf{s}$ と $\mathbf{w}$ の事後分布を推定する． $\mathbf{s}$ の分布は次式の条件付き確率の反復計算により得られる．

$$\begin{aligned} & P(s_t | w_t, \mathbf{s}_{-t}, \alpha, \beta) \\ & \propto \frac{n(s_t, w_t) + \beta}{n(s_t) + W_{s_t} \beta} \cdot \frac{n(s_{t-2}, s_{t-2}, s_t) + \alpha}{n(s_{t-2}, s_{t-1}) + S\alpha} \\ & \quad \cdot \frac{n(s_{t-1}, s_t, s_{t+1}) + I(s_{t-1} = s_{t-1} = s_t = s_{t+1}) + \alpha}{n(s_{t-1}, s_t) + I(s_{t-2} = s_{t-1} = s_t) + S\alpha} \\ & \quad \cdot \frac{n(s_t, s_{t+1}, s_{t+2}) + I(s_{t-2} = s_t = s_{t+2}, s_{t-1} = s_{t+1}) + I(s_{t-1} = s_t = s_{t+1} = s_{t+2}) + \alpha}{n(s_t, s_{t+1}) + I(s_{t-2} = s_t, s_{t-1} = s_{t+1}) + I(s_{t-1} = s_t = s_{t+1}) + S\alpha} \end{aligned} \quad (\text{A.5})$$

ここで， W_{s_t} は隠れ状態 s_t が出力できる語数，また， S は隠れ状態数である．カウント関数 $n(*)$ はその隠れ状態と語の系列中に $*$ が起きた（共起した）回数を返す．また， $I(*)$ はその条件 $*$ が真なら 1，偽なら 0 を返す．

学習に用いるコーパス $\mathbf{w}$ には二つのダミー語，すなわち，それぞれ文頭と文尾を表す “BOS” (begin of sentence) と “EOS” (end of sentence) を含む．例えば，英語コーパスは $\mathbf{w} = (\text{“BOS”, she, is, read, -ing, a, book, “EOS”, “BOS,” the, chair, } \dots,$

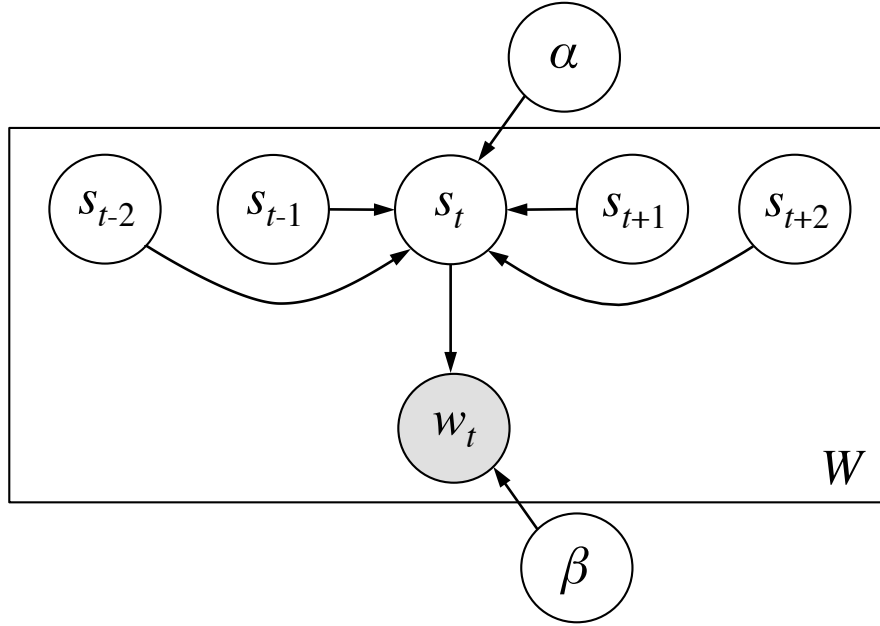


Figure A.1: A graphical model for the Bayesian hidden Markov model. W indicates the number of words in a corpus, i.e., the length of a corpus.

“EOS”)のように表される．これらの“BOS”と“EOS”の隠れ状態にはそれぞれ1と2の番号をあらかじめ割り当て，これらの状態の確率は更新しない．これらのダミー語以外の初期隠れ状態へは，3から S の数をランダムに割り当てる．なお，本文中での S はこれらのダミー語を省略している．そして，式(A.6)を500回反復計算し，そのサンプリングの後半80%を平均し， $P(s_t|\mathbf{w}, \mathbf{s}_{-t})$ を得る．語意推定のすべての実験において，ハイパーパラメータは $\alpha = 0.9$ と $\beta = 1.0$ に経験的に設定した．

A.2.2 語と特徴の連合学習

各語の文法カテゴリの推定の後，モデルは文とシーンの対から，言語的カテゴリ($\mathbf{w}$ と $\mathbf{s}$)と特徴カテゴリ(f と m)の連合を学習する． i 番目の文($i = 1, 2, \dots, N$)は語系列 $\mathbf{w}^{(i)} = (w_1^{(i)}, w_2^{(i)}, \dots, w_t^{(i)}, \dots, w_{T(i)}^{(i)})$ で構成され，それらの文法カテゴリは $\mathbf{s}^{(i)} = (s_1^{(i)}, s_2^{(i)}, \dots, s_t^{(i)}, \dots, s_{T(i)}^{(i)})$ で表される．また，文の提示と同時に，モデルは i

番目のシーンに存在する視覚的特徴 $\mathbf{f}^{(i)} = (f_1^{(i)}, f_2^{(i)}, \dots, f_j^{(i)}, \dots)$ を観察する．この特徴は時系列ではない．そして，ある語 w_q が特徴 f_p と連合する確率を，それらの共起確率として計算する．

$$P(o = f_p | w = w_q) = \frac{1}{Z} \sum_i^N n^{(i)}(w_q, f_p) \quad (\text{A.6})$$

ここで， $n^{(i)}(w_q, f_p)$ は i 番目の文・シーン対 $(\mathbf{w}^{(i)}, \mathbf{f}^{(i)})$ において w_q と f_p が同時に存在する回数であり， Z は正規化定数である．今回，この確率がすべての語変数に適用できることを仮定する．すなわち， $P(o|w_t^{(i)}) = P(o|w)$ ($t = 1, 2, \dots, T^{(i)}$ ，および， $i = 1, 2, \dots, N$) である．

その後に， s_t の下での m の確率 ($P(m|s_t)$) を計算する．この確率は例えば，動詞カテゴリは動作カテゴリと対応しやすいといった傾向を示すことが期待される．このような写像は式 (A.6) と同様の共起確率を用いて計算できる．しかし，実際には， m の確率は一様分布になるため，単純な共起関係から意味のある写像を獲得することは難しい． i 番目のシーンが与えられたとき，モデルは物体や動作で構成される多くの特徴のセット $\mathbf{f}^{(i)}$ を観察し，その結果， m における物体と動作のカテゴリの分布はほぼ一様になる．したがって，モデルは m と s_t の有意な連合を学習できない．これを避けるため， i 番目のシーンにおける意味カテゴリを，語系列に対応する意味カテゴリ系列 $\mathbf{m}^{(i)} = (m_1^{(i)}, m_2^{(i)}, \dots, m_t^{(i)}, \dots)$ に拡張する．意味カテゴリ系列は学習後の $P(o|w)$ を用いて得られる．

$$\begin{aligned} P(m_t^{(i)} | w_t^{(i)}, \mathbf{f}^{(i)}) &= \sum_o P(m_t^{(i)} | o) P(o | w_t^{(i)}, \mathbf{f}^{(i)}) \\ &\propto \sum_o P(o | m_t^{(i)}) P(m_t^{(i)}) P(o | w_t^{(i)}) P(o | \mathbf{f}^{(i)}) \end{aligned} \quad (\text{A.7})$$

ここで，二行目の右辺第一項は既知であり，第二項は一様分布である．また，その第三項は式 (A.6) で得られる．第四項は i 番目のシーンに存在する特徴 $\mathbf{f}^{(i)}$ についての一様分布である．式 (A.7) により，例えば，語 “walk” は動詞カテゴリであり，これは動作カテゴリに含まれる “walk” の視覚特徴を指すという推定が可能になる．そして，

t と i にわたって $m_t^{(i)}$ と $s_t^{(i)}$ の共起回数をカウントすることで $P(m|s)$ を得る.

$$P(m = m_q | s = s_p) \equiv \frac{1}{Z} \sum_i^N \sum_t^{T(i)} P(m_t^{(i)} = m_q | w_t^{(i)}, \mathbf{f}^{(i)}) P(s_t^{(i)} = s_p | w_t^{(i)}, \mathbf{s}_{-t}^{(i)}) \quad (\text{A.8})$$

ここで, Z は正規化定数. この確率はすべての文法カテゴリ変数に適用されることを仮定する. すなわち, $P(m|s_t^{(i)}) = P(m|s)$ ($t = 1, 2, \dots, T(i)$, および, $i = 1, 2, \dots, N$) である.

A.3 学習コーパス

A.3.1 語意推定における学習コーパス

第3での実験では, モデルは英語, 日本語, または, 中国語の簡単な人工コーパスで学習する. Figs. A.3.1–A.3.1 に英語, 日本語, および, 中国語のコーパス生成のための語カテゴリの遷移規則を示す. 図中の各円は語カテゴリを示し, このカテゴリの語彙から一語がランダムに生成される. 今回, 各語は曖昧性なく一つのカテゴリから選ばれると仮定するため, 例えば, 名詞にも動詞にもなり得る語は存在しない. 図中には描かれていないが, いずれのコーパスでも, 名詞 (主語と目的語のカテゴリ) の前に形容詞カテゴリが50%の確率で配される. Table A.1 に各語カテゴリの語彙の大きさを示す. 主語と目的語のカテゴリの語彙は同じ名詞で構成される. これらの規則は “BOS” で始まり, “EOS” で終わる. これらのダミー語の文法カテゴリはあらかじめ割り当てられており, 特徴の推定には用いられない. 矢印上の小数は語カテゴリ間の遷移確率を表し, 灰色の円上の小数はその語カテゴリが欠落する確率を表す. 動詞を含む文では, 三つの時制, すなわち, 現在形, 過去形, および, 進行形のうちのいずれかが等確率で選択される.

英語の規則 (Fig. A.3.1) では, 名詞カテゴリ, または, 形容詞カテゴリの前に, 冠詞カテゴリが配置される. 名詞は単数のみであるが, 現在形の動詞接尾辞の “-s”, “-es”, および, “-ies” は省略する. しかし, 過去形と進行形を表す動詞接尾辞の “-ed” と “-ing” は動詞の後に配置される. 日本語の規則 (Fig. A.3.1) では, 時制は動詞語尾

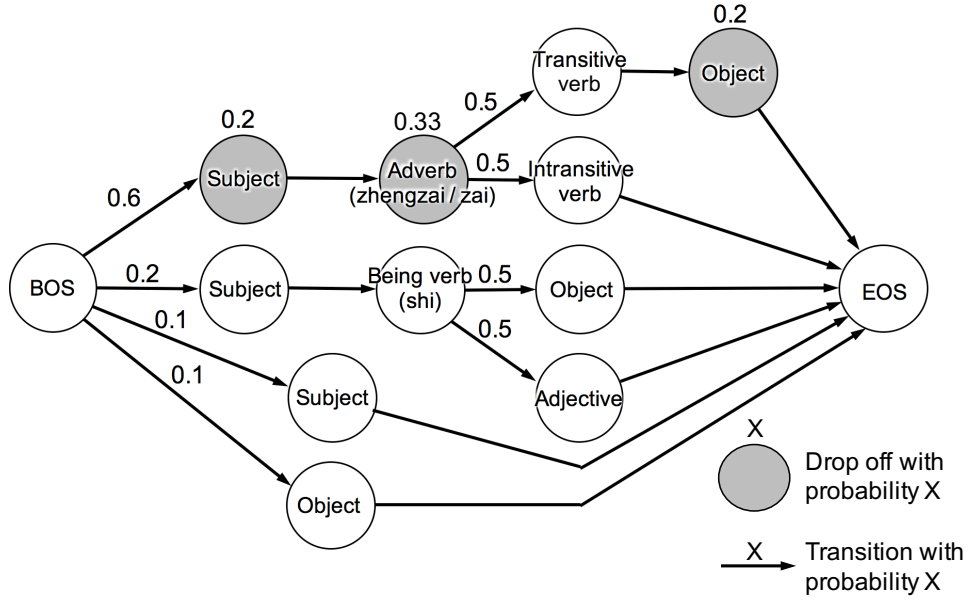


Figure A.4: Transition rules between word categories for the production of the Chinese corpus.

Table A.1: Vocabulary sizes of word categories in each corpus.

Word category	English	Japanese	Chinese
Subject	41	45	56
Object	37	45	56
Adjective	35	34	34
Transitive	20	21	20
Intransitive	14	15	14
Article	2	—	—
Auxiliary	—	3	—
Adverb	—	—	2

(助動詞)で決まる。日本語には三つの格助詞があり、これによって名詞の主題役割(主語か目的語)が決まる。図中の particle1 はと particle3 はそれぞれ格助詞「が」と「は」であり、この前の名詞は主語である。そして、particle2 は対格「を」であり、その前の名詞が目的語であることを示す。日本語の形容詞は語尾「い」を持つことが多いため、今回、形容詞語尾「い」を形容詞語基と分離する。中国語の規則 (Fig. A.3.1) では、助動詞(「正在 (zhengzai)」と「在 (zai)」)が進行形を表す。「是 (shi)」は本来多くの機能を持つが、今回、コピュラ文における be 動詞の機能のみを持つとする。また、格助詞「的 (de)」を形容詞語尾として扱い、形容詞語基の後ろに配置する。

これらの遷移確率は統語的特徴が実際の対幼児発話コーパス [62, 61, 63] と一致するように選択された。まず、実コーパス中の述語動詞、主語、および、目的語の割合を求め、また、名詞句のみの文と動詞句のみの文の割合を求めた。そして、人工コーパスと実コーパスのこれらの割合がおよそ一致するように、適切な遷移確率を設定した。さらに、英語コーパスでは、対幼児発話を解析した研究 [64] から、三つの構文(断片文(名詞句、または、動詞句単体)、コピュラ文、および、主語・述語文)がそれに一致するように遷移確率を設定した。

A.3.2 過剰生成における学習コーパスの語彙

第4章での実験に用いた英語コーパス規則 (Fig. 4.2 (a)) の語彙を Table A.3.2 に、日本語コーパス規則 (Fig. 4.2 (b)) の語彙を Table A.3.2 に示す。

Table A.2: Vocabulary sizes of word categories in the English corpus.

Word category	Vocabulary size
Noun	257
Transitive	45
Intransitive	45
Past tense irregular transitive	10
Past tense irregular intransitive	10
Article	2
“ed” (bound morpheme)	1
Total	372

Table A.3: Vocabulary sizes of word categories in the Japanese corpus.

Word category	Vocabulary size
Noun	243
Transitive	52
Intransitive	52
Adjective	17
“Ha” and “ga” (case particles)	2
“O” (case particle)	1
“No” (case particle)	1
Auxiliary	2
Total	372

研究業績リスト

学術雑誌

1. 河合祐司, 大嶋悠司, 浅田稔. 未分化な文法カテゴリによる幼児発話の誤用：英語の過去形の形態素と日本語の格助詞の過剰生成に共通した計算モデル, 認知科学, 印刷中.
2. 河合祐司, 大嶋悠司, 笹本勇輝, 長井志江, 浅田稔. 幼児の統語発達モデル：日本語, 英語, 中国語の言語構造を反映した統語範疇の獲得過程. 認知科学, Vol. 22, No. 3, pp. 475–479, 2015.
3. 幸田憲明, 河合祐司, 松井伸之. RBF 出力関数を有する RCE ニューロンモデルとその性能評価. 計測自動制御学会論文誌, Vol. 45, No. 11, pp. 620–627, 2007.

国際・国内会議における発表（査読あり）

1. Yuji Kawai, Minoru Asada, and Yukie Nagai. A model for biological motion detection based on motor prediction in the dorsal premotor area in Proceedings of the 4th Joint IEEE International Conference on Development and Learning and on Epigenetic Robotics, pp. 241–247, 2014.
2. Yuji Kawai, Yuji Oshima, Yuki Sasamoto, Yukie Nagai, and Minoru Asada. Computational model for syntactic development: Identifying how children learn to generalize nouns and verbs for different languages in Proceedings of the 4th Joint IEEE International Conference on Development and Learning and on Epigenetic Robotics pp. 78–84, 2014.

3. 河合祐司, 大嶋悠司, 笹本勇輝, 長井志江, 浅田稔. 幼児の統語発達モデル：日本語, 英語, 中国語の言語構造を反映した統語範疇の獲得過程. 日本認知科学会 第 31 回大会, pp. 126–133, 2014
4. Yuji Kawai, Yukie Nagai, and Minoru Asada. Perceptual Development Triggered by its Self-Organization in Cognitive Learning in Proceedings of the 2012 IEEE/RSJ International Conference on Intelligent Robots and Systems, pp. 5159-5164, 2012
5. Jimmy Baraglia, Yukie Nagai, Yuji Kawai, and Minoru Asada. The Role of Robot Embodiment in the Emergence of Mirror Neuron Systems. in Proceedings of the 2nd IEEE International Conference on Development and Learning and on Epigenetic Robotics, accepted, 2012.
6. Yuji Kawai, Jihoon Park, Takato Horii, Kazuaki Tanaka, Hiroki Mori, Yukie Nagai, Takashi Takuma, and Minoru Asada. Throwing Skill Optimization through Synchronization and Desynchronization of Degree of Freedom. in Proceedings of the 16th Annual RoboCup International Symposium, 2012.
7. Yukie Nagai, Yuji Kawai, and Minoru Asada. Emergence of Mirror Neuron System: Immature vision leads to self-other correspondence. in Proceedings of the 1st Joint IEEE International Conference on Development and Learning and on Epigenetic Robotics, 2011.

国際・国内会議における発表（査読なし）

1. 河合祐司, 大嶋悠司, 笹本勇輝, 長井志江, 浅田稔. 幼児の語意学習における統語ブートストラップモデル. 脳と心のメカニズム 第 15 回冬のワークショップ, 2015.

2. Yuji Kawai, Yuji Oshima, Yuki Sasamoto, Yukie Nagai and Minoru Asada. A Bayesian model for syntactic development: How do children generalize nouns and verbs in Japanese, English, and Chinese? in ICIS 2014 Pre-Conference, Computational Models of Infant Development, 2014.
3. 河合祐司, 浅田稔, 長井志江. 背側運動前野の運動予測機能に基づく生体運動検出モデル. ニューロコンピューティング研究会 (NC) , pp. 221-226, 2014.
4. 河合祐司, 大嶋悠司, 浅田稔, 項目依拠的規則に加えた統語範疇遷移規則の獲得がもたらす過去形の過剰一般化の発達的变化: 計算論モデルによる検討. 日本赤ちゃん学会第 14 回学術集会, p. 67, 2014.
5. 大嶋悠司, 河合祐司, 笹本勇輝, 長井志江, 浅田稔. 統語範疇の精緻化に基づく幼児の名詞・動詞般用発達モデル. 日本赤ちゃん学会第 14 回学術集会, p. 74, 2014.
6. 河合祐司, 浅田稔, 長井志江. 新生児期のバイオリジカルモーション知覚モデル: 胎児期での運動による生体運動不変性の獲得. 日本赤ちゃん学会第 13 回学術集会, p.33, 2013.
7. 河合祐司, 長井志江, 浅田稔. 認知機能の学習における感覚空間の自己組織化に応じた発達の制約. 日本機械学会ロボティクス・メカトロニクス講演会 2012 講演論文集, 2A1-M11, 2012.
8. Jihoon Park, Yuji Kawai, Takato Horii, Kazuaki Tanaka, Hihorki Mori, Yukie Nagai, Takashi Takuma, and Minoru Asada. Differentiation within Coordination in Acquisition of Skilled Throwing. 第 35 回人工知能学会 AI チャレンジ研究会, pp. 19-24, 2012.
9. 河合祐司, 長井志江, 浅田稔. 視覚の精緻化が導くミラーニューロンシステムの発達モデル. 第 29 回日本ロボット学会学術講演会予稿集, 1A3-2, 2011.

10. 河合祐司, 長井志江, 浅田稔. 視覚の精緻化が導く他者運動理解能力の発達モデル. 日本赤ちゃん学会第11回学術集会, p. 56, 2011.
11. 河合祐司, 幸田憲明, 松井伸之. RBF関数を出力関数としたRCEニューロンモデルの性能評価. 計測自動制御学会システム・情報部門学術講演会2008講演論文集, pp. 551-554, 2008.

国際・国内シンポジウムにおける発表（査読なし）

1. Tomohiro Takimoto, Yuji Kawai, Jihoon Park, and Minoru Asada. Babbling emerges from rhythmical neural activities through STDP. The 3rd International Symposium on Cognitive Neuroscience Robotics: Toward Constructive Developmental Science, 2016.
2. Koki Ichinose, Jihoon Park, Yuji Kawai, Junichi Suzuki, Hiroki Mori, and Minoru Asada. Dynamical complexity in small world network: emergence of gamma oscillation from the structure of spiking neural network. The 3rd International Symposium on Cognitive Neuroscience Robotics: Toward Constructive Developmental Science, 2016.
3. 河合祐司, 大嶋悠司, 笹本勇輝, 長井志江, 浅田稔. 幼児の統語範疇を手がかりとした語意獲得モデル—統語範疇の母語依存性と発達的变化の再現. 日本赤ちゃん学会若手部会 第2回研究合宿, 2014.
4. 河合祐司, 大嶋悠司, 笹本勇輝, 長井志江, 浅田稔. ベイジアン隠れマルコフモデルを用いた幼児の統語発達モデル. 身体性認知科学と実世界応用に関する若手研究専門委員会 (ECSRA) 第12回研究会, 2014.
5. 河合祐司, 浅田稔, 長井志江. 予測による生体運動の不変性抽出: 生体運動知覚の計算論モデル. 日本赤ちゃん学会若手部会 第1回研究合宿, p. 4, 2013.

6. 河合祐司, 浅田稔, 長井志江. 予測機能に基づく生体運動知覚の計算論モデル. 身体性認知科学と実世界応用に関する若手研究専門委員会 (ECSRA) 第11回研究会, 2013.
7. Yuji Kawai, Yukie Nagai, and Minoru Asada. Maturational Constraints Lifted by Self-Organization in Perceptual Space for Cognitive Learning. Bielefeld-Osaka Workshop 2012, 2012.
8. Yuji Kawai, Yukie Nagai, and Minoru Asada. Self-Other Confusion Caused by Undifferentiated Visual Perception: A Clue to the Understanding of Others' Action. International Workshop for Young Researchers "Knowing Self, Knowing Others", 2011.

書籍（分担執筆）

1. 今福理博, 鹿子木康弘, 河合祐司, 前原由喜夫. ロボットは人の心をもてるのか? –共感性が秘密の鍵. 高校生のための心理学, 第7章, 誠信書房, 2016.
2. Yuji Kawai, Jihoon Park, Takato Horii, Yuji Oshima, Kazuaki Tanaka, Hihorki Mori, Yukie Nagai, Takashi Takuma, and Minoru Asada. Throwing Skill Optimization through Synchronization and Desynchronization of Degree of Freedom. RoboCup 2012: Robot Soccer World Cup XVI Lecture Notes in Computer Science, Vol. 7500, pp. 178-189, Springer, 2013.

受賞

1. 日本認知科学会第31回大会 大会発表賞, 2014
2. 人工知能学会 研究会優秀賞, 2013.
3. ロボカップ日本委員会 ロボカップ研究賞, 2013.

4. 社団法人 生産技術振興協会 宮原国際研究活動助成選抜論文, 2012.
5. 第 35 回人工知能学会 AI チャレンジ研究会 人工知能学会賞, 2012.
6. Microsoft Student Travel Award of the 1st Joint IEEE International Conference on Development and Learning and on Epigenetic Robotics, 2011.
7. 日本機械学会 畠山賞, 2008.