

自由発話による誘導教示からの言葉の意味の獲得

岡 夏樹[†] 増子 雄哉^{†*} 林口 円^{††} 伊丹 英樹[†] 川上 茂雄^{††**}

[†] 京都工芸繊維大学 大学院工芸科学研究科 情報工学専攻

^{††} 京都工芸繊維大学 工芸科学部 情報工学課程

〒 606-8585 京都市左京区松ヶ崎御所海道町

* 現, コスモスイニシア ** 現, NEC 通信システム

あらまし 迷路内のエージェントを音声で自由に誘導するタスクを設定し, そこでのインタラクションを通じて誘導教示の意味を獲得するアルゴリズムを提案する. 自由発話からの学習では, 誤りを含み, 頻度に偏りがある, 多種類の少数データからの学習が求められる. 本論文では, Fisher の直接確率計算法を用いることにより, そのような性質を持つデータの中から, 共起頻度が有意に高い言葉と状況(意味)の対を精度良く見つける方法を提案し, 先行研究の学習アルゴリズムと比較して, 学習精度が向上することを実験的に示す.

キーワード 言語獲得, 共起, Fisher の直接法, 行動教示, 評価教示

Meaning Acquisition of Everyday Navigation Instructions

Natsuki OKA[†], Yuya MASHIKO^{†*}, Madoka HAYASHIGUCHI^{††},

Hideki ITAMI[†], and Shigeo KAWAKAMI^{††**}

[†] Department of Information Science, Graduate School of Science and Technology

^{††} Department of Information Science, School of Science and Technology

Kyoto Institute of Technology

Matsugasaki Goshokaidou-chou, Sakyo-ku, Kyoto 606-8585, Japan

* Now at Cosmos Initia ** Now at NEC Communication Systems

Abstract Proposed is an algorithm that learns the meaning of action instructions and evaluation instructions through interaction with a navigator. The navigator speaks freely, which requires learning from various kinds of few data with biased frequency and error. Utilizing Fisher's exact test, the algorithm finds the pairs of a phrase and a situation which frequently cooccurs.

Key words language acquisition, cooccurrence, Fisher's exact test, action instruction, evaluation instruction

1. はじめに

誰にでも使えるインタフェースの実現のためには, 個人ごとの言葉の違いに適応できることが望ましい. さらに, 日常生活の中で使われる機器のインタフェースの場合には, 様々な働きの発話が混在した中で, それらの意味を獲得できる必要があると考える. 例えば, 掃除ロボットに対して人が掃除の仕方を教える場面では, 動作の指示や, 実行された動作への評価が入り混じった発話がなされるだろう.

言葉の意味学習の計算モデルに関するこれまでの研究の流れの一つは, 知覚や行動を単語と対応づけようとするものであり, 多くの研究が行われてきている(総説としてたとえば[3])が, 言葉の社会的な働きとしての意味の獲得はあまり扱われてこな

かった. これに対して, 鈴木ら[5]は, 行動の指示と行動の評価という異なる社会的働きを持つ言葉が混在する中での意味学習を目指した.

鈴木らは, 迷路を移動するロボット(ただし, 実機ではなくシミュレータを使用)を教示により誘導するタスクを設定し, ある状況でどのような行動が適切であるかの知識をロボットが利用できることを前提として, ロボットによる教示意味の学習モデルを提案した. 本研究においても, 鈴木らと同様に, 行動の指示と行動の評価のような異なる社会的働きを持つ言葉が混在する中での意味学習を目指す. また, 我々も, ある状況でどのような行動が適切であるかの知識を前提として意味学習を行う.

鈴木らは, 教示は計算機が与えることとし, a) 与えられる教

示の種類は、行動の指示（「上／下／左／右」の4種類）と行動の評価（「よし／だめ」の2種類）だけに限る。b) 教示は常に正しい。c) 教示は行動ごとに毎回必ず与えられる。d) 行動の指示と評価のどちらが与えられるかはランダムに選ばれる。—とした。これに対して本研究では、人が自由に発話した教示からの学習を目指す。したがって、以下の4つの特徴を持つ発話から意味学習を行う必要がある。

(1) 教示は、行動の指示と評価だけには限定されないし、また、行動指示や評価を伝える場合でも、様々な言い方でそれが行われる。たとえば、《下に行け》の意味で「下」と言ったり「下りて」と言ったりする。また、《直前の行動が不適切だった》の意味で「違う」と言ったり「あかん」と言ったりする。なお、本論文では、発話やフレーズは「」で囲んで示し、フレーズの意味は《》で囲んで示す。

(2) 誤った教示やタイミングがずれた教示も含まれる。

(3) 教示は行動ごとに毎回与えられるとは限らない。

(4) 各教示の生起頻度にも特に制限はない。計算機が一定の仕方で教示をした鈴木らの実験においても、各教示の出現頻度や各教示と対応する状況の出現頻度には、迷路の特徴やロボットの振る舞いの特徴に起因する偏りがあったが、人による自由教示では、この偏りがさらに増加すると考えられる。

このような極めて自然な教示が与えられる環境において、教示の意味の獲得を実現することが、本研究の目的である。自然な教示を受け入れ可能とすることは、たとえば車の運転中のような、余分な認知負荷を人に与えてはいけな状況や、緊急事態が発生するかもしれない状況における人と機械のインタラクションを実用レベルに高めるためには、避けて通れない技術課題である。

したがって、実際の応用場面で教示の意味学習を可能にするためには、ここで観察されたような性質を持つインタラクションから学習可能なアルゴリズムを構築することが必要となる。

本論文では、以降、まず、我々が提案する学習アルゴリズムの比較対象とする鈴木らの意味学習アルゴリズムを紹介し(2.節)、つづいて、学習対象データの収集実験について記す(3.節)。さらに、Fisherの直接確率計算法(付録1.参照)を用いることにより、上記の4つの性質を持つデータの中から、特異的に共起している言葉と状況を見つけ、それらに対応づけることにより、言葉の意味を獲得する方法を提案する(4.節)。さらに、提案方法の学習精度が鈴木らの方法による学習精度を上回ることを、実験的に示す(5.節)。

2. 関連研究

本節では、鈴木ら[5]による意味学習アルゴリズムを示す。本論文では、この鈴木らのアルゴリズムと我々が提案するアルゴリズムの比較評価を5.節にて行う。

鈴木らのアルゴリズムは以下の通りである：

(1) 各教示がどの状況で発せられたかを記憶しておく。ここで状況とは、《直前の行動が良かった／悪かった》《直前の行

動は上／下／左／右に行くべきであった》^(注1)の合計6種類である。

(2) ある教示 I_j の、状況頻度基準（ある教示が与えられたときに、もっとも発生頻度の高い状況をその教示の意味とする）に基づく確信度 $Bel_s(\sigma_i, I_j)$ ($1 \leq k \leq n, n$ は状況の数) を次の式で計算する。

$$Bel_s(\sigma_i, I_j) = \frac{times(\sigma_i, I_j)}{\sum_k times(\sigma_k, I_j)} \quad (1)$$

ただし $times(\sigma_i, I_j)$ は、状況 σ_i で教示 I_j が与えられた累積回数である。

(3) $Bel_s(\sigma_i, I_j)$ が一番高い確信度の $1/2$ に満たない状況 σ_i は教示 I_j の意味候補から外す。(2倍基準)

(4) もし教示 I_j の意味候補となる状況が1つになれば、その状況を教示 I_j の意味として確定する。

(5) もし二つ以上の状況が I_j の意味候補として残ったら、教示頻度基準（ある状況でもっとも頻度の多い教示がその状況を意味する教示である）に基づく確信度を各 l ごとに次の式で計算する。

$$Bel_t(\sigma_l, I_j) = \frac{times(\sigma_l, I_j)}{\sum_k times(\sigma_l, I_k)} \quad (2)$$

(6) $Bel_t(\sigma_l, I_j)$ が一番高い確信度の $1/2$ に満たない場合、教示 I_j の意味候補から σ_l を外す。(2倍基準)

(7) 残った状況 σ_l が教示 I_j の最終的な意味候補となる。

3. 発話データ収集実験

我々は、計算機の画面上の迷路(図1)を用いて、仮想環境において人が学習エージェントを声でゴールまで誘導するタスクを設定し、インタラクションデータを収集した。

実験協力者に与えた教示は次の通りである。

- 画面上の迷路を移動するキャラクターを星の所(ゴール)まで自由に声を掛けて誘導して下さい。

- 好きなタイミングに好きな言葉で教示を与えて下さい。

- スタートからゴールまで教示を続けて下さい。

実験協力者は20代前半の工学系の大学生／大学院生で男性4名、女性1名であった。発話データの収集はビデオカメラを用いて行った。なお、エージェントが迷路探索を開始してからゴールに到達するまでを1試行とする。

エージェントは全ての地点における最適行動(ゴールまでの最短の道筋)を知っており、0, 0.33, 0.50, 0.66のいずれかの確率で最適行動を行い、それ以外の場合はランダム(最適行動を含む)に行動を行う。ここでの行動とは、上、下、左、右いずれかの移動が可能な方向への、1秒ごとの移動である。

実施した収集実験の内訳は表1の通りである。実験協力者DおよびEにおいてエージェントの行動の確率を変動させた理由

(注1)：鈴木らは、このように、行動の指示を《直前の行動は上／下／左／右に行くべきであった》の意味であるとしたが、我々が行った実験(3.節)の結果、人が教示を与える場合は、行動の指示は、《次は上／下／左／右に行くべきである》の意味でなされることが分かった。したがって、学習精度の比較実験(5.節)においては、鈴木らの方法も、《次は上／下／左／右に行くべきである》の学習を行うよう修正して実施した。



図 1 発話データ収集実験で用いた迷路の画面構成
Fig. 1 A maze used in the experiment.

表 1 発話データ収集実験の内訳
Table 1 The details of the data collection experiment.

実験協力者	最適行動の割合	試行数
A, B	0	1 試行
C	0	2 試行
D, E	0.33 → 0.50 → 0.66	左記の順で各 1 試行

は、エージェントが学習しているかのような印象を実験協力者に与えるためである。

4. 学習方法

学習手順は次の通りである：

(1) 手作業で音声データから意味学習の対象とする言葉の抜き出しを行い、その言葉が発せられた状況とともに記録する。この部分は、将来自動化する予定であるので、自動化を想定した処理を手作業で行う。(4.1 節)

(2) (1) で抜き出された言葉と状況に対して、共起頻度に基づいて関連の深さを計算することで、言葉と状況の対応付けを行い、この対応を言葉の意味とする。(4.2 節)

4.1 学習データの抽出

学習データ抽出の際に行った一連の操作について以下に示す。

(1) 録画された映像を基にして実験協力者の発話をエージェントの位置ごとに人手で記録する。

(2) 記録された発話からフレーズ（意味付与の単位）を抽出する（図 2）。抽出の方法は、エージェントの行動タイミングおよび無音区間で発話を区切り、1 フレーズとする。ただし、繰り返される音は同じフレーズの重複とみなし、そのフレーズの 1 回の出現と数える（たとえば、「右右」という発話は「右」というフレーズとみなす）。また、ある発話が、既出のフレーズ（音節数が 2 以上）を含む場合は、そのフレーズと他のフレーズが連続しているとみなす（たとえば、「右」というフレーズが既出の場合、「右行って」は「右」と「行って」の 2 つのフレーズの連続であるとみなす）。常に他のフレーズと連続して出現し、音節数が 1 のものはフレーズとはしない（たとえば、

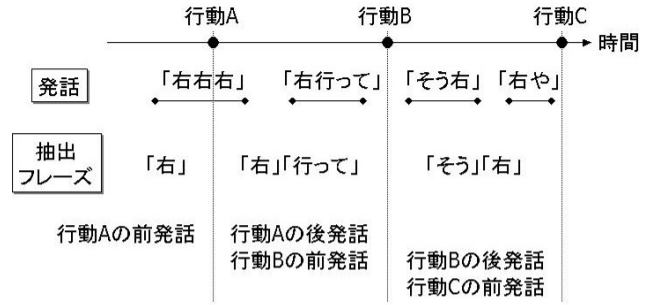


図 2 ある行動の前発話、後発話としてのフレーズの抽出
Fig. 2 Phrase extraction as pre- or post-action utterances.

「右や^(注2)」や「左や」などの常に「右」や「左」などのフレーズと共に出現する「や」はフレーズとはしない）。

(3) ある行動の直前の発話（1 つ前の行動からその行動までの間の発話）に含まれるフレーズをその行動の前発話、ある行動の直後の発話（その行動から次の行動までの間の発話）に含まれるフレーズをその行動の後発話として収録する（図 2）。今回の迷路誘導タスクでは、エージェントの行動間隔が 1 秒と短いため、ある発話が直前の行動の評価に関するものか、直後の行動を指示するものかを、発話タイミングだけに基いて区別することは困難である。したがって、両方の可能性を考慮して、1 つの発話を両方に分類した（図 2）。

4.2 意味学習アルゴリズム

4.2.1 学習アルゴリズムの概要

本アルゴリズムの外部仕様は次の通りである。本アルゴリズムは、エージェントが、行動学習により得られた知識（ある状況でどのように行動すべきか）をすでに持っている想定し、その知識に基づいて、行動を指示する言葉の意味、および、行動を評価する言葉の意味を獲得する。その際、行動を指示する言葉や行動を評価する言葉以外の言葉が混ざっていてもよい。また、言葉の種類（どれが行動を指示する言葉で、どれが評価する言葉か等）があらかじめ分かっている必要はない。また、誤った教示やタイミングがずれた教示も含まれていてよいし、各教示の生起頻度にも特に制限はないとする。なお、意味を付与する候補となる言葉の（連続音声入力からの）切り出しは、外部プログラムまたは人手によりすでに行われているとする。

提案アルゴリズムの内部仕様は次の通りである。言葉と状況の共起頻度に基づき、言葉と状況を対応づける。ここで状況とは、《次にとるべき行動は上／下／左／右への移動である》《直前の行動は適切／不適切であった》の 6 種類を想定している。対応づけるために、周辺度数（ある言葉の度数やある状況の度数）から特定の組度数（ある言葉とある状況が共起する度数）が得られる確率を、Fisher の直接確率計算法（付録 1.）を用いて算出する。この確率が小さい程、その言葉の意味がその状況と対応する可能性が高いというのが、本アルゴリズムの基本的なアイデアである。

なお、Fisher の直接確率計算法は、通常は検定の一手法とし

(注2)：これは関西地方の方言であり、関東地方の方言では「右だ」に相当する。

表 2 前発話の度数分布表の一例（被験者 E のデータの一部）．表中の
 $\uparrow$, $\downarrow$, $\leftarrow$, $\rightarrow$ は上下左右に進むべきだという状況を表し，下，
 右，違う，などはフレーズを示す．

Table 2 A sample table of frequency distribution of pre-action utterances.

	下	右	違う	左	上	下りて	そのまま	よし	...	横計
$\uparrow$	0	2	2	0	7	0	0	2	...	13
$\downarrow$	41	0	8	1	0	4	2	0	...	58
$\leftarrow$	2	0	4	7	0	0	0	0	...	16
$\rightarrow$	0	18	3	0	0	0	0	0	...	22
縦計	43	20	17	8	7	4	2	2	...	109

表 3 後発話の度数分布表の一例（被験者 E のデータの一部）．表中
 の○，×は直前の行動が適切／不適切だったという状況を表し，
 下，右，違う，などはフレーズを示す．

Table 3 A sample table of frequency distribution of post-action utterances.

	下	右	違う	左	上	下りて	そのまま	よし	...	横計
○	29	13	0	2	4	2	2	2	...	58
×	14	7	17	6	3	2	0	0	...	51
縦計	43	20	17	8	7	4	2	2	...	109

て用いられるものであるが，ここでは，検定ではなく，共起度数の特異性の判断のために用いていることに注意されたい．共起度数が特異であるということは，ある言葉と，ある状況が特別の関係にある（その言葉は，その特定の状況で発せられやすい，あるいは，発せられにくい）ことを示す．したがって，その言葉の意味は，「その特定の状況であること（あるいは，ないこと）」であると判定するのである．本アルゴリズムでは，この状況を，行動学習により得られた知識に基づき，《右に行くべき状況》《適切な行動をとった（褒められるべき）状況》などと設定して，言葉との対応づけを試みる．

4.2.2 学習アルゴリズムの詳細

本研究では，実験協力者の発話の中で，下記の 6 種類の意味を持つフレーズの意味をエージェントが獲得することを目指す．1 つの意味を持つフレーズは複数種類あることを想定している．また，これら 6 種類以外の意味を持つフレーズの存在も想定しているが，それらに対しては，6 種類以外の意味を持つフレーズであることをシステムが学習することを目指す．

学習対象とするフレーズの意味は次の 6 種類である：エージェントが次に進むべき方向を示す行動指示 4 種《 $\uparrow$ 》《 $\downarrow$ 》《 $\leftarrow$ 》《 $\rightarrow$ 》，エージェントの直前の行動に対する評価を示す評価指示 2 種《○》《×》．ここで《 $\uparrow$ 》《 $\downarrow$ 》《 $\leftarrow$ 》《 $\rightarrow$ 》はそれぞれ《上／下／左／右に進め》の略号，《○》《×》はそれぞれ《適切だった》《不適切だった》の略号である．

意味学習アルゴリズムは以下の通りである：

(1) 《上／下／左／右に進むべきである》の各状況と各前発話との共起度数を記録し，度数分布表を作成する．度数分布表の一例を表 2 に示す．

(2) 《直前の行動が適切／不適切であった》の各状況と各後発話との共起度数を記録し，度数分布表を作成する．度数分布表の一例を表 3 に示す．

表 4 表 2 から作成した四分表の一例

Table 4 An example of four-fold table extracted from Table 2.

	「下」	「下」以外	合計
$\downarrow$	41	17	58
$\downarrow$ 以外	2	49	51
合計	43	66	109

(3) (1) (2) で得られた度数分布表から，ある 1 つの状況 x とある 1 つのフレーズ y に着目し，四分表を作成する．(状況数 $\times$ フレーズ種類数) の枚数の四分表が作成される．四分表の一例を表 4 に示した．

(4) これらの四分表に対して，周辺度数を固定した場合に各セルの度数パターンが観測される確率と，観測されたパターン以上に偏ったパターンが観測される確率の総和 P_t を，Fisher の直接法（付録 1.）を用いて算出する．この確率の総和 P_t は，状況 x とフレーズ y の共起度合いの特異性を示すものである．すなわち， P_t が小さいことは，周辺分布を基準とした場合，状況 x とフレーズ y の共起頻度が特異的に高い，または，特異的に低いことを意味する．特異的に高いのか低いのかは，付録の表 A.1 の記号を使用すると， $a - \frac{p_t}{t}$ の値の正負により区別できる．この値が負であることは，状況 x とフレーズ y の共起度数が周辺度数を基準とした標準的な度数より小さいことを示す．ここでは共起頻度が特異的に高い組を見つけないので， $a - \frac{p_t}{t}$ が負の場合は状況 x はフレーズ y の意味候補とはしない．

(5) i) (1) から作成された四分表において共起頻度が特異的に高い状況 x とフレーズ y の組が見つかった（すなわち， $a - \frac{p_t}{t} > 0$ ，かつ， $P_t < \theta$ (θ は，あらかじめ定めておく閾値) となる，状況 x とフレーズ y の組が見つかった) とすると，フレーズ y の意味は，次に進むべき方向を指示する行動指示（状況 x に対応する行動指示）であると推定できる．表 5 の例では， $a - \frac{p_t}{t} > 0$ ， $P_t < \theta$ の各条件を満たすデータには波線をつけて示したが，この両条件を満たす状況は《 $\uparrow$ 》《 $\downarrow$ 》《 $\leftarrow$ 》《 $\rightarrow$ 》の中には無く，「違う」の意味は行動の指示ではないと推定される．ii) (2) から作成された四分表において共起頻度が特異的に高い状況 x とフレーズ y の組が見つかった（すなわち， $a - \frac{p_t}{t} > 0$ ，かつ， $P_t < \theta$ ）とすると，フレーズ y の意味は，直前の行動を評価する指示（状況 x に対応する評価指示）であると推定できる．表 5 の例では，「違う」の意味は《×》であると推定される．iii) いずれの四分表においても， $a - \frac{p_t}{t} > 0$ ，かつ， $P_t < \theta$ となることがなかったフレーズ y については，上記 6 種類以外の意味，すなわち，《その他の意味》を示すものであると推定できる．

(6) ステップ 5 において，あるフレーズ y に対して，共起頻度が特異的に高い状況 x が 1 つだけ見つかったとすると，そのフレーズ y の意味は状況 x であるとし，その確信度は 1 とする．表 5 の例は，1 つだけ見つかった場合を示している．複数の状況 x に対して共起頻度が特異的に高いことが分かったとすると，それらの各状況に対して，確率 P_t の逆数をそれぞれ求め，総和が 1 となるように正規化し，その値をそのフレーズ y の意味がそれぞれの状況 x である確信度とする．なお，《その他

表 5 教示の意味の推定のための計算過程の一例. 実験協力者 E の「違う」の意味の推定過程. $\theta = 0.2$ とした場合.

Table 5 An example of the meaning estimation process of a phrase.

「違う」				
状況	度数	確率 P_t	$a - \frac{pr}{t}$	確信度
↑	2	0.67	-0.03	1.00
↓	8	0.42	-1.05	
←	4	0.23	<u>1.50</u>	
→	3	0.54	-0.43	
○	0	<u>0.00</u>	-9.05	
×	17	<u>0.00</u>	<u>9.05</u>	

表 6 本研究で提案するアルゴリズム ($\theta = 0.2$) と鈴木らのアルゴリズムの正答率の比較

Table 6 The comparison of the correctness between the proposed algorithm ($\theta = 0.2$) and Suzuki et al.'s algorithm.

実験協力者	延べ語数	提案アルゴリズム		鈴木らのアルゴリズム	
		延べ語数正答率 (%)	異なり語数正答率 (%)	延べ語数正答率 (%)	異なり語数正答率 (%)
A	33	62.4	47.4	57.6	45.0
B	76	58.5	62.3	49.0	54.4
C	61	65.4	60.6	53.3	54.8
D	100	100.0	100.0	78.3	76.6
E	109	93.6	64.3	72.9	75.0

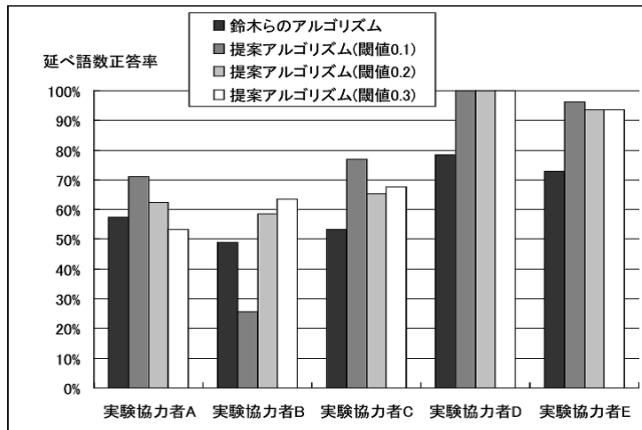


図 3 閾値 θ の大きさが正答率に与える影響

Fig. 3 The effect of the threshold θ on the percentage of correct estimations.

の意味》に判定された場合の確信度は、1 とする。

5. 学習性能の評価と考察

本研究において提案する学習アルゴリズムと、鈴木ら [5] が用いた学習アルゴリズムとを、上記収集データ (3. 節) から抽出した学習データ (4.1 節) に適用し、その学習性能を比較検討した。 $\theta = 0.2$ としたときの結果を表 6 に示す。今回提案したアルゴリズムは、鈴木らのアルゴリズムをほぼ一貫して上回る学習精度を示していることが分かる。なお、閾値 θ の大きさが学習精度に与える影響については、現在検討中であるが、実験結果の一部を図 3 に示す。

ここで、“延べ語数正答率”、“異なり語数正答率” は、それぞれ、次の通り定義される：

$$\text{延べ語数正答率} = \frac{\text{正答フレーズの発生回数}}{\text{全フレーズの発生回数}}$$

$$\text{異なり語数正答率} = \frac{\text{正答フレーズの種類の数}}{\text{全フレーズの種類の数}}$$

意味を推定する上で、より多く話された言葉を正しく推定できた方が有益であると我々は考える。したがって、異なり語数正答率の結果よりも、延べ語数正答率の結果 (表 6 中では太字で示した) が重要であると考えられる。

なお、両アルゴリズムとも、複数の意味候補が挙げられる場合があるが、その場合の正答率は、提案アルゴリズムでは確信度の重みを付けて算出し、鈴木らのアルゴリズムでは (1/候補数) の重みを付けて算出した。また、想定する 6 種類以外の意味のフレーズに対しては、《その他の意味》と判定した場合を正答とした。

両アルゴリズムの正答率の差の原因として、次の 3 点を挙げることができる：

(1) 鈴木らのアルゴリズムでは単純に共起頻度が高いという基準で意味を付与している。これに対して、我々のアルゴリズムでは周辺度数を基準として計算した組度数の実現確率が小さいことを判断基準としている。この違いにより、提案アルゴリズムは、各状況の生起頻度や各教示の生起頻度が均一でない場合にも高い精度で学習可能となっている。

この具体例として、表 2 を参照して説明する。表 2 の「違う」の列に注目する。鈴木らの状況頻度基準に従うと、「違う」は《下に行け》を意味するという誤った判断をしてしまう。他の状況と比べて、下に行くべき状況において、「違う」が多く発せられているのは、単に、今回の実験環境においては、下に行くべき状況が他の状況より生じやすかった (ゴールがエージェントの初期位置よりも下方に設定されていた) からに過ぎないのであるが、状況頻度基準ではこの情報を考慮できない。これに対して、我々のアルゴリズムでは周辺度数を基準として組度数の実現確率を計算するため、「違う」が特定の方向教示とは対応しないことを正しく判断できる (表 5 を参照)。

(2) 提案アルゴリズムでは、Fisher の直接確率計算法を用いたため、データが少ないときにも、正確な確率 (共起度数の特異性) が計算でき、データが少ないなりの適切な推測が可能となった。これに対して、鈴木らのアルゴリズムではデータ数の多少による推測の確からしさの変動を考慮する手段を持たないため、データが少ないときに判断を誤ることがある。

人とのインタラクションを通じた学習では、データは人とのインタラクションを通じてしか得られないため、大量のデータを利用できることは期待できず、また、学習の早い段階で学習効果を見せないと、使い続けてもらうことが難しいため、少数データからの学習性能は極めて重要であると考えられる。

具体例として、再び表 2 を参照して説明する。表 2 の「そのまま」の列に注目する。鈴木らの状況頻度基準に従うと、「そのまま」は《下に行け》を意味するという誤った判断をしてしまう。これに対して、Fisher の直接確率計算法を用いると、こ

の2例がたまたま下にいくべき状況と共起する確率を正確に計算することができるため、鈴木らの方法のように少数例の偶然による共起を過大評価することはない。鈴木らは2倍基準（2.節）というヒューリスティックを導入して判定の確からしさの問題に対処しようとしているが、これは少数データに対してはうまく働くとは限らない。

（3）鈴木らのアルゴリズムは2倍基準以外にも、「まず状況頻度基準で判定し次に教示頻度基準を用いる」というヒューリスティックな順序を導入している。これに対し、提案手法では、状況頻度基準・教示頻度基準のような、度数分布表上の非対称な処理は排除し、正確な確率計算に基づくことにより、可能な限りヒューリスティックな判断を排除した。ただし、我々のアルゴリズムにおいても、《その他の意味》に対応するための閾値 θ の設定や、Fisherの直接法により求めた確率から確信度に変換する部分にヒューリスティックな判断が残っており、改善の余地がある。

6. おわりに

本論文では、言葉とその発せられた状況（次に進むべき方向がどちらであるか、および、直前の行動が適切であったかどうか）の共起頻度に基づいて、言葉と状況を対応づけることにより、言葉の意味を学習するアルゴリズムを提案した。

Fisherの直接確率計算法を利用することにより、先行研究である鈴木らのアルゴリズムと比較して、学習精度が向上することを実験的に確認し、向上の原因について分析した。

今後、ネットワーク機器やロボットなどの複雑で高度な装置が家庭にますます普及することを考えると、人とのインタラクションデータに基づく学習は、大変重要な技術であると考えられる。そこでは、誤りを含み、頻度に偏りがある、少数のデータからの学習が求められるため、本論文で提案したような正確な確率計算に基づく推定技術は地味ではあるが必須のものとなると我々は考えている。

今回は、学習データの抽出部分は手作業で実施したが、我々は、音声信号列からの繰り返しパターンに注目したフレーズ候補抽出技術を並行して開発中であり、今後は、取得データのオフライン解析でなく、オンライン環境における、提案した学習アルゴリズムの性能を検証することを計画している。また、我々は、行動学習を前提としない教示の意味学習の研究[1]や、非明示的な報酬からの意味学習の研究[4]も進めており、これらの技術との統合も考えていきたい。

謝辞 本研究の一部は科学研究費補助金基盤研究(c)「3項間インタラクションを通じた認知発達メカニズムの構成的解明」の支援を受けた。

文 献

- [1] 伊丹英樹, 川上茂雄, 野川博司, 岡 夏樹: 最終行動ヒューリスティクスに基づく教示意味の学習, 平成19年度情報処理学会関西支部大会講演論文集, 83-86, 2007.
- [2] 森敏昭, 吉田寿夫 編著: 心理学のためのデータ解析テクニカルブック, 北大路書房, 187-189, 1990.
- [3] Roy, D.: Grounding words in perception and action: computational insights, TRENDS in Cognitive Sciences, 9(8), 389-396, 2005.

表 A・1 四分表 (2 × 2 分割表)

Table A・1 Four-fold table (Two-by-two frequency table).

	カテゴリー 1	カテゴリー 2	合計
条件 1	a	b	p
条件 2	c	d	q
合計	r	s	t

- [4] 左 祥, 田中一晶, 嵯峨野泰明, 岡 夏樹: “No news is good news” 規準を利用した行動教示の学習法, 情報科学技術レターズ, 319-322, 2007.
- [5] 鈴木健太郎, 植田一博, 開一夫: 自律的な行動学習を利用した評価教示の計算論的意味学習モデル, 認知科学 9(2), 200-212, 2002.

付 録

1. Fisher の直接確率計算法

Fisherの直接法[2]は、対応がない2条件間の比率の比較を行う方法であり、得られたデータが少ない場合にも用いることができる方法である。

まず、比較の対象となるデータを表 A・1 の四分表の形に整理する。

Fisherの直接法では、次に示すように、周辺度数(p, q, r, s)を観測された値に全て固定した場合の各セルの度数(a, b, c, d)についてのある特定のパターンが得られる確率を求めることにより、帰無仮説「カテゴリー1が条件1のときに観測される比率は、カテゴリー1が条件2のときに観測される比率と差がない」に対する検定が行われる。

表 A・1 に示されるデータについて、全データ t 個のうちの r 個がカテゴリー1に属し、残りの s 個がカテゴリー2に属するという組み合わせは、 ${}_tC_r$ 通りである。このなかで、 a 個が条件1を満たし、 c 個が条件2を満たすという組み合わせは、 ${}_pC_a \times {}_qC_c$ 通りである。従って、表 A・1 のようなデータが得られる確率を P とすると、

$$P = \frac{{}_pC_a \times {}_qC_c}{{}_tC_r} = \frac{p!q!r!s!}{a!b!c!d!t!} \quad (\text{A} \cdot 1)$$

となる。フィッシャーの直接法では、式 A・1 を用いて観測された各セルの度数パターンが得られる確率と、観測されたパターン以上に対立仮説を支持するパターンになる確率を全て求める。そして、これらの確率の総和 P_t があらかじめ設定された有意水準以下であれば帰無仮説を棄却する。

観測されたパターン以上に対立仮説を支持するパターンとは、観測されたパターン（表 A・1）以上にデータが偏っている場合を指す。データが表 A・1 で示される場合、式 A・2 の値を算出する。

$$a - \frac{pr}{t} \quad (\text{A} \cdot 2)$$

式 A・2 の値が正ならば、 a の値を+1, +2, +3, ... したパターン (b または c の値が0になるまで) が、より対立仮説を支持するパターンとなる。負ならば、データ a を-1, -2, -3, ... したパターン (a の値が0となるまで) がより対立仮説を支持するパターンとなる。