

特集

ロボット対話

実世界情報を用いたコミュニケーションの学習

Robot Dialogue

Robots That Learn Physically-Situated Communication

杉浦 孔明
Komei Sugiura(独) 情報通信研究機構
National Institute of Information and Communications Technology
komei.sugiura@nict.go.jp, <http://komeisugiura.jp/>**keywords:** imitation learning, human-robot dialogue, robot language acquisition, multimodal dialogue systems

1. はじめに

ロボットが日常生活に浸透するにしたがって、人とロボットのコミュニケーションが課題となる [西田 03]。どれほど豊富な機能を持つロボットを構築したとしても、ユーザが手軽に機能を使えなければ普及しないであろう。

ロボットとの音声によるコミュニケーションは、ユーザにとって手軽であるというメリットがあるが、実現は簡単ではない。もちろん、雑音抑圧、発話検出、話者分離などは音声処理における重要な課題である。しかし、本稿で強調したい本質的な課題は、発話の解釈が実世界情報や経験により影響を受けることである。例えば「コップ取って」という命令を実行するには、どのコップを取るのか、「取る」とはどのような動作軌道を表すのか、を推論しなければならない。また、取る理由は食事の準備か片付けのためか、などによってもユーザに渡すべき対象は変わり得る。

一方、家族にコップを取ってもらう場合は省略した言い方でも通じることが多いうえ、わからなければ聞き返すだろうという期待がある。つまり、人間はそれまでの経験から、相手が自分の発言をどれくらい理解できるかに関して暗黙的な知識がある。しかし、ロボットと対話するユーザにとって、ロボットが状況をどれだけ理解しているかを推定するのは難しいだけでなく、ロボット側にもユーザの理解モデルがないため、両者の対話は支離滅裂になりがちである。

現状のロボット対話処理機構では、動作コマンドの伝達を目的とすることが多いにも関わらず、動作情報と音声認識は別々に処理されることが多い。それに対し、ロボットの対話機構にグラウンドしない知識を導入することが、シンボルグラウンディング問題を孕むことは古くから指摘されている [Harnad 90, Pfeifer 99]。したがって、ロボットに人間と自然にコミュニケーションさせるためには、ユーザや状況への適応手法が重要課題となる。

本稿では、記号創発ロボティクス [谷口 10] の中心課

題の一つである、実世界にグラウンドしたコミュニケーションの学習を扱う。まず、ロボットとの音声対話における研究を分類し課題を整理する。3 章以降では、著者らの研究グループがこれまでにこなってきたロボット対話を実現するための要素技術の開発として、動作の言語化技術（動作学習、動作生成）、ロボット対話技術（確認発話、指示発話の生成）を紹介する。

2. ロボットとの音声対話

実世界にグラウンドした対話を実現するためには、実世界のオブジェクトや動作を記号化・言語化することが極めて重要な課題である。ロボティクスの分野では、動作-記号の相互変換に関する試みが近年注目されている [Inamura 04, Ogata 07, 高野 09]。高野らは、運動の分節化を通じてヒューマノイドロボットが獲得した原始シンボルを用いて、運動認識・生成を行っている [高野 09]。Ogata らは、動作系列と記号列の間の多対多対応問題を扱うリカレントニューラルネットに基づく手法を提案している [Ogata 07]。Kollar らは、ロボットに与える移動指示に関して、ランドマークオブジェクトや動作にグラウンドした言語表現を学習する手法を提案した [Kollar 10]。学習されたモデルを用いることにより、指示が示す最も確からしい経路を推論する。

オブジェクト-記号の相互変換に関しては、人工知能や自然言語生成の分野で多くの研究が行われてきた [山肩 04, Roy 02, Dale 95]。これらの研究では、オブジェクトを指示する言語表現の生成が試みられている。山肩らはコップ類の名称とイメージモデルの参照関係における曖昧性に一貫した個人差があることを示している [山肩 04]。さらに、音声・言語・画像レベルの情報を統合して、複数のオブジェクトの中からユーザが意図したオブジェクトを選択する手法を提案している。Roy は、ディスプレイ上の複数の長方形のうちひとつを指示する言語表現をテキストベースで生成する手法を提案している [Roy 02]。[Roy

02]で提案された手法では、単語のカテゴリは教師なし学習の枠組みでクラスタリングされるため、設計者が属性を用意する必要がない。オブジェクト-記号の相互変換の詳細については、本特集の長井らによる解説 [長井 12]を参照されたい。

自然言語生成 (NLG) の分野では、コーパスに基づいて機械に言語表現を生成させる試みが行われてきた [Dale 95, Jordan 05, Funakoshi 09]。例えば、TUNA コーパスは、人間が生成する自然言語表現と近い表現を生成する手法を比較評価するために構築されたコーパスである [Sluis 06]。TUNA コーパスでは、システムへの入力 は家具の画像と属性値 (XML テキスト) である。例えば、サイズ属性は大・小のいずれか、タイプ属性は椅子・ソファ・机・扇風機のいずれか、など事前に定義された属性とその値がシステムに直接与えられる。

音声対話を行うロボットの先駆的事例としては、Jijo-2 が挙げられる [松井 00]。また、稲邑らは、移動ロボットの障害物回避において、ロボットが段階的に行動決定モデルを獲得する機構を提案した [稲邑 01]。センサ値と行動の関係をベイジアンネットを用いてモデル化し、推論結果の確信度を用いて応答決定を行う。

一方、対話システムの分野では、ホテル検索やバスの経路検索などが代表的なタスクである [Komatani 00, Bohus 06, Kawahara 98]。Singh らの研究 [Singh 00] は、強化学習ベースの対話戦略学習の最初の研究の一つであるが、その後、これを部分観測マルコフ決定過程 (POMDP) 上の強化学習に拡張した手法が注目されるようになった [Williams 07]。これらの手法は、グラウンドしない知識に基づく対話システムに適用される場合がほとんどであるが、将来的にはマルチモーダル対話システムに適用されることが期待される。

3. 動作の言語化

ロボットによるコミュニケーション能力の学習において、動作の言語化は重要な課題である。特に、生活支援ロボットにとっては、「食器棚からコップを取り出す」などの物体の操作は必要不可欠な機能であり、これらの動作を言語で指示できることが望ましい。

一方、各種の日用品や棚に対応する動作を事前にプログラムするコストは非常に大きい。事前にプログラムされた動作がユーザにとってイメージしにくいものである場合、安心して動作指示できないという問題がある。そこで我々は、模倣学習の枠組みにより物体操作を学習する手法の開発を行ってきた。以下では、我々が開発した物体操作の模倣学習手法を紹介する。

3.1 動作の学習

行動の模倣には、人間の多様な行動を観察し、実世界の事物とどう関連するかを推定することが不可欠である。

例えば、「X を Y に載せる」や「Z を回す」などの動作をユーザがロボットに教示することを考える (図 1)。認知言語学では、動かされるオブジェクトをトラジェクタ、トラジェクタの基準になるオブジェクトをランドマークと呼ぶ [Langacker 87]。上記の 2 種類の動作のように、ランドマークを必要とする動作と必要としない動作を同じ枠組みで学習させることは簡単ではない。このような情報は、通常訓練データ中に明示されていないためである。

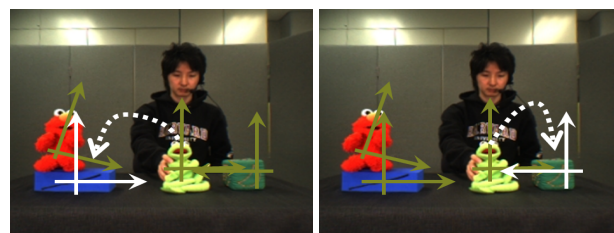


図 1 上図はともに「載せる」軌道であるが、世界座標系では汎化できない。一方、RPD-HMM では各軌道の汎化に最適な座標系を探索できる。上図において、明い実線は最適座標系、点線はユーザが実行した軌道を表す。

そこで我々は、適切な座標系を推定し軌道を汎化する手法 RPD-HMM (Reference-Point-Dependent HMM) を開発した。最適な座標系は、EM アルゴリズムにより探索される。つまり、図 1 に示すように、トラジェクタやランドマークを原点とするいくつかの座標系を考え、各座標系において軌道を学習する。これを繰り返すことで、誤った座標系における軌道は一貫性がなくなるため、最終的に正しい座標系が選択される。図 2 は、本研究で仮定した固有座標系タイプ (右図) と、各動作に対する固有座標系タイプのテストセット尤度 (左図) を示したものである。図より、各動作に対して合理的な固有座標系タイプが選択されていることがわかる。

提案手法では、以下の尤度最大化基準により、固有座標系タイプ k 、参照物体インデックス列 r 、HMM パラメー

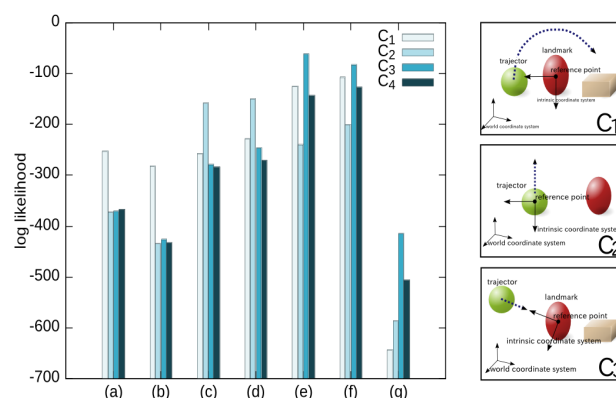


図 2 固有座標系タイプのテストセット尤度。(a) 載せる, (b) 飛び越えさせる, (c) ちかづける, (d) 離す, (e) 上げる, (f) 下げる, (g) 回す。座標系 C_4 は世界座標系である。

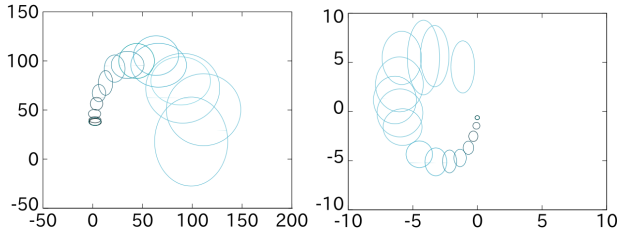


図3 「載せる」動作を学習後の HMM. 位置 (左図), 速度 (右図) の出力確率密度分布を示す. 状態遷移を淡色から濃色の变化で図示した. 各軸は単位は [pixel].

タ λ を探索する.

$$(\hat{\lambda}, \hat{k}, \hat{r}) = \operatorname{argmax}_{\lambda, k, r} \sum_{l=1}^L \log P({}^C \Xi_l; {}^C \lambda) \quad (1)$$

ここに Ξ_l は l 番目の訓練サンプルにおける軌道を表し, 固有座標系タイプ k と参照点 r によって決定される固有座標系 C 上での軌道を ${}^C \Xi$ と表記することにする. (1) の解は [Sugiura 11b] を参照されたい.

HMM を用いて動作を学習する場合, 動作に応じて状態数を自動推定できることが望ましい. 生活環境ではどのような動作を学習するかあらかじめ決められないため, 仮に状態数が多すぎると適切にパラメータを学習できない可能性がある. 提案手法では, HMM の状態数を leave-one-out cross-validation により選択する. つまり, ユーザにより軌道が N 回提示されたとすると, $N-1$ の軌道サンプルを学習セットとしてモデルを学習し, 残りの 1 サンプルの検証集合 (validation set) の尤度が最も高い状態数を最適状態数とする. 実験の結果, 動作の種類およびサンプル数に応じて, 適切に状態数が変化することを確認した.

図3に, 「載せる」動作モデルとして学習された HMM を示す. 左図および右図はそれぞれ, 位置および速度の出力確率密度分布を楕円で図示したものである. 図3左図より, 状態が遷移していくに従い, 位置の分散が小さくなるのがわかる. このことは, 「載せる」動作が特定の点へ収束するような軌道であることを示唆している.

本節では, 主に 2 次元空間における模倣学習について述べた. 学習した RPD-HMM による動作認識の精度は, 3 人の被験者に対して約 90% であった. RPD-HMM の 3 次元動作への拡張については, [杉浦 10] を参照されたい.

3.2 動作の生成

次に, 模倣学習における動作の生成について考える. 動作の生成時には, 学習時とまったく同じように物体が配置されている訳ではない. そのため, たとえ同じ「載せる」動作であっても, 学習時の軌道をそのまま用いることは無意味である. 以降, 与えられた物体の配置 (および画像特徴量) を「状況」と呼ぶことにする.

前節で述べたように, 学習した RPD-HMM は固有座標系において汎化されたものであった. 状況に応じた生

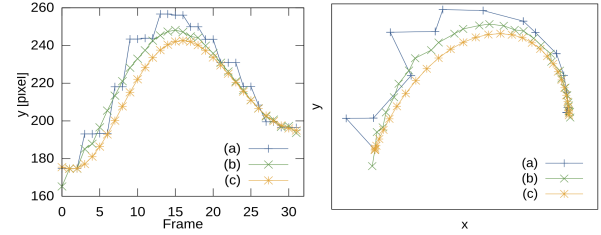


図5 デルタパラメータを用いた軌道生成の例. (a) 位置のみ, (b) 位置および 1 次のデルタパラメータ, (c) 位置および 1,2 次のデルタパラメータ, を用いた軌道生成の結果.

成を行うために, 我々は RPD-HMM を固有座標系 C から世界座標系 W へ変換する手法を提案している [Sugiura 11a]. 提案手法の概念図を図4に示す. RPD-HMM の位置・速度・加速度の平均ベクトルおよび共分散行列は, 同時変換行列により C から W 上のモデルに変換される.

HMM からの動作生成ではモンテカルロサンプリングが広く使われているが, 有限のサンプリングでは尤度最大化が保証されない. 一方, 我々の手法では, トラジェクトリ HMM [Tokuda 00] を用いて尤度最大化基準による連続的な軌道の生成を行なっている. トラジェクトリ HMM は, 位置だけでなく速度と加速度を用いて軌道生成を行う点が特徴である. 速度・加速度は音声処理分野では 1 次および 2 次のデルタパラメータと呼ばれる. デルタパラメータによる軌道生成の効果を図5に示す. 図5は, 動作「載せる」に関して, 軌道生成を行った結果である. 図において, 左図に t - y プロット, 右図に x - y プロットを示す. 図より, 位置のみを用いた場合には不連続な軌道が生成されているのに対し, 1, 2 次のデルタパラメータを用いることで滑らかな軌道が生成できることがわかる.

いま, HMM λ および状態系列 q が与えられたうえで, 生成軌道の尤度について考える. 出力確率密度関数はガウス分布であるので, 尤度を最大化する軌道 $\hat{\Xi}$ は次のようになる [Tokuda 00].

$$\hat{\Xi} = \operatorname{argmax}_{\Xi} \log P(\Xi | q, \lambda) \quad (2)$$

$$= \operatorname{argmax}_{\Xi} \left\{ -\frac{1}{2} \Xi^T \Sigma^{-1} \Xi + \Xi^T \Sigma^{-1} \mu + K \right\} \quad (3)$$

ここに, μ は q の各要素に対応する平均ベクトルを並べたベクトルであり, Σ は q の各要素に対応する共分散行列を対角に並べた行列である.

実験では, 学習に必要なサンプル数を調べるために, 学習サンプル数に対する生成誤差を調べた. ユーザの提示した軌道データを学習データとテストデータに分け, 学習データから作成した HMM から軌道を生成させて, テストデータの軌道と比較した. 図6は, ユーザが実行した軌道 Ξ と生成軌道 $\hat{\Xi}$ の間の生成誤差 $D(\Xi, \hat{\Xi})$ をプロットしたものである. 生成誤差 $D(\Xi, \hat{\Xi})$ は, フレーム長 T

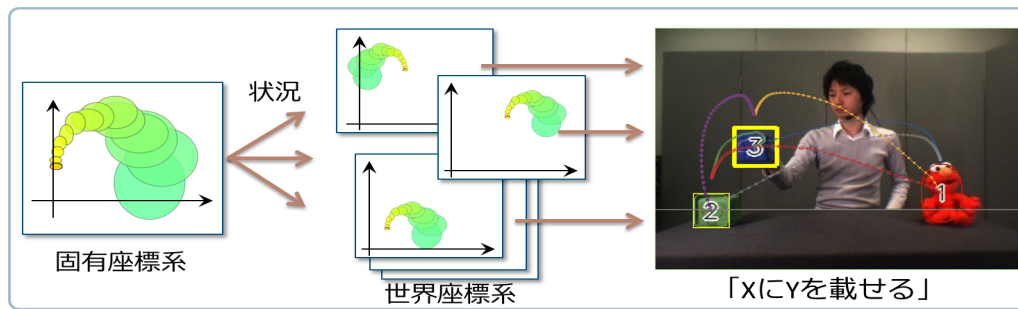


図4 軌道生成におけるHMMの座標変換

で正規化されたユークリッド距離で与えられる。

$$D(\Xi, \hat{\Xi}) = \frac{1}{T} \sum_{t=1}^T \sqrt{|x_t - \hat{x}_t|^2} \quad (4)$$

ここに、 x_t および $\hat{x}_t$ は、それぞれ Ξ と $\hat{\Xi}$ のフレーム t における位置ベクトルを表す。ただし、 $\hat{x}_t$ は、 x_t と同数のフレームを含むようにリサンプリングを行った。図において、縦軸は生成誤差、横軸は訓練サンプル数を表す。この実験より、7 個程度のサンプルで生成誤差が収束することがわかった。このように少ないサンプルで動作を学習できることは、インタラクティブに動作を学習するロボットにとって重要な要求仕様である。

我々は、RPD-HMM の実世界への応用にも取り組んでいる。実際に、RoboCup@Home 環境における家事動作の模倣学習に適用され、成功を収めている。実世界における模倣学習の応用では、障害物回避を考慮した軌道を生成することが重要である。山内らは、Rapidly-Exploring Random Tree (RRT) と RPD-HMM を組み合わせることにより、障害物を回避しつつ、人間の軌道に近い軌道を生成させる手法を提案している [山内 11]。

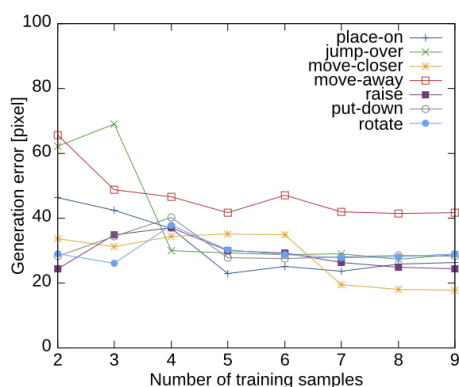


図6 訓練サンプル数と生成誤差 $D(\Xi, \hat{\Xi})$ の関係

4. コミュニケーションの学習

ロボットが生活環境に浸透するに従い、ロボットのコミュニケーション能力が課題となる。現状の技術では、「コップ（テーブルに）置いて」などの曖昧な発話の解釈

は非常に難しい。省略された語が表すオブジェクトを推定する必要があるうえ、環境中に存在する「コップ」の候補から正しいオブジェクトを選択しなければならない。

本節では、我々が開発してきたコミュニケーション学習基盤 LCore [Iwahashi 07, Iwahashi 09] を紹介する。LCore は、画像、動作、アフォーダンス、履歴などの実世界情報を学習し、発話・動作の生成が可能である。

4.1 コミュニケーション学習基盤 LCore

現状のロボットの対話処理機構では、動作コマンドの伝達を目的とすることが多いにも関わらず、動作情報と音声認識は別々に処理されていることが多い。ユーザの発話の意味はグラウンドされない知識に基づいて解釈されるため、動作が状況にふさわしいかどうかは音声認識時には考慮されない。しかしながらこのような手法では、ユーザの発話の意味が状況に応じて適切に理解されない、という問題がある。

LCore [Iwahashi 07] では、マルチモーダル入力から学習されたモデルを用いてユーザの発話を理解する。本論文では、音声・画像・動作などの各モダリティに対応するモデルを**信念モジュール**と呼ぶ。また、(1) 音声、(2) 動作、(3) 視覚、(4) 動作-オブジェクト関係、(5) 行動コンテキスト、の 5 つの信念モジュールを統合したモデルを**共有信念 Ψ** と呼ぶ。

各モジュールは以下のように定義される。

- 音声信念 B_S

単語の n -gram、および節の接続確率として学習される。 B_S は、発話 s に対する概念構造 z の条件付き確率の対数として表す。

- 視覚信念 B_I

ガウス分布により学習される。 B_I は、オブジェクトの視覚特徴量が与えられたときの対数尤度である。

- 動作信念 B_M

3 章で述べた RPD-HMM により学習される。 B_M は、可能な行動に対して軌道を仮想的に生成したうえで、その軌道の尤度として得られる。

- 動作-オブジェクト関係信念 B_R

「平らなものはオブジェクトを載せられやすい」など、動作と視覚的特徴の関係を表す。 B_R は、2 個の

オブジェクトの視覚特徴量に対するガウス分布の対数尤度である。

● 行動コンテキスト信念 B_H

B_H は、あるコンテキスト（「把持されている」、「直前に操作された」など）のもとでの指示対象としてのオブジェクトの適切さ（スコア）を表す。

以上より、共有信念関数 Ψ を、各信念モジュールの重み付き和として定義する。

$$\Psi(s, a_k, O, q^{(i)}) = \max_z (\gamma_1 B_S + \gamma_2 B_I + \gamma_3 B_M + \gamma_4 B_R + \gamma_5 B_H) \quad (5)$$

ここに、 z は各単語の概念構造への分割であり、 $(s, a_k, O, q^{(i)})$ はそれぞれ、発話、行動、状況、トラジェクタに関するコンテキストを表す。また、 $\gamma = (\gamma_1, \dots, \gamma_5)$ は各信念に対する重みであり、MCE 学習 [Katagiri 98] を用いて学習される。

コンテキスト q 、状況 O 、発話 s が与えられたときの最適行動 $\hat{a}$ は以下で得られる。

$$\hat{a} = \operatorname{argmax}_k \Psi(s, a_k, O, q) \quad (6)$$

4.2 確信度に基づく動作と発話の生成

ロボットとの音声によるコミュニケーションは、ユーザにとって手軽であるというメリットがあるものの、自然発話に含まれる曖昧性の処理は簡単ではない。例えば、ユーザが「コップ（テーブルに）置いて」などの曖昧な発話を行ったとする。曖昧性を解消するためには、「赤いコップですか、青いコップですか」のような確認発話を毎回行ってもよいが、曖昧なときのみ確認発話を行う方が望ましい。

そこで我々は、発話理解確率（発話を正しく解釈できる確率）を推定し、効用最大化により応答を生成する手法を提案した [Sugiura 11c]。発話が曖昧であるとき、共有信念関数の第 1 候補と第 2 候補のスコアの差（マージン）が小さいことから、これを曖昧性の尺度として用いている。提案手法では、統合確信度を発話理解確率の推定値としてモデル化した。統合確信度は、ベイズロジスティック回帰により学習される。動作応答 b_1 と確認発話応答 b_2 は、対応する期待効用 $\mathbb{E}[R_i] (i = 1, 2)$ の最大化により選択される。

$$\mathbb{E}[R_i] = r_{i1} f(d; w) + r_{i2} (1 - f(d; w)) \quad (7)$$

$$d = \Psi(s, \hat{a}, O, q) - \max_{j \neq k} \Psi(s, a_j, O, q) \quad (8)$$

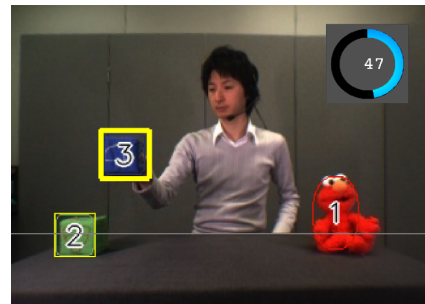
ここに、 $f(d; w)$ は、 $w = (w_0, w_1)$ をパラメータとするロジスティックシグモイド関数である。また、 r_{i1}, r_{i2} はそれぞれ、行動 $\hat{a}$ がそれぞれ正解、不正解のときの応答 b_i に対する効用である。

実験では、オブジェクトを操作するよう、ユーザはロボットに指示を与える。両者は音声対話により曖昧性を

解消し、ロボットがユーザの意図した行動をとればタスク成功とした。ロボットが失敗行動（正解でない行動）を行ったときは、アームに取り付けたセンサを叩くことで教師信号を与えることができる。このようにして得たマージンと教師信号の組を用いて、ベイズロジスティック回帰により $f(d)$ を学習させた。

図 7 は、ユーザ (U) とロボット (R) の対話例を示したものである。図において、右上の数値は統合確信度 $f(d)$ を表す。図 7 では、最適行動の確信度は $f(d) = 0.478$ であり、確認発話「アオイ ハコをちかづけて」が最適応答であった。この言語表現は、オブジェクト 2 と 3 の視覚的特徴のなかで最も異なる属性^{*1}について述べており、ユーザにとって理解しやすい。ランドマークについては確認発話を行わなくても確信度に影響はないため、確認を省略していると考えられる。

実験の結果、動作や視覚などの情報を用いず、音声のみで発話理解を行った場合の行動失敗率は 83.4% であった。ベースライン手法（発話を行わず動作のみで応答する）における行動失敗率は 12.0% である一方、提案手法による行動失敗率は 2.6% であった。このことから、提案手法はベースライン手法に比べて行動失敗率を大幅に低減できたといえる。



【状況】オブジェクト 2 が直前に操作された

U: ハコ エルモ ちかづけて。

R: ミドリハコをちかづけて？

U: いいえ。

R: アオイハコをちかづけて？

U: はい。

R: (動作実行：オブジェクト 3 をオブジェクト 1 に近づける)

図 7 グラウンドした語彙による確認発話の生成

4.3 能動学習に基づく対話戦略最適化

一般の音声対話システムと、音声対話を行うロボットとの本質的な違いは、ロボットが（ユーザと類似した）身体性を有していることである。そのため、物体を物理的に操作したり、物理世界の状態を変化させたりするこ

*1 カラー画像ではオブジェクト 2 は緑、オブジェクト 3 は青色である。

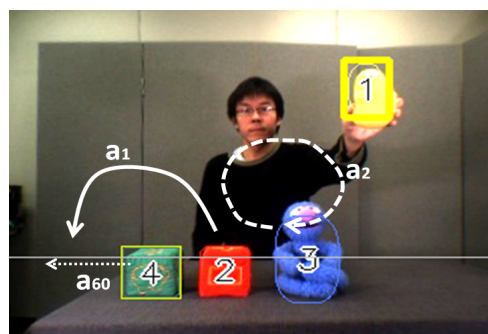


図8 能動学習に基づく発話の生成例

Target action	Robot utterance	Margin	ICM	ELL
a ₁	"Jump-over red box"	15.0	0.913	3.41
a ₁	"Jump-over box red box green"	17.9	0.945	3.43
a ₁	"Jump-over"	0.077	0.473	4.33
a ₁	:	:	:	:
a ₂	"Rotate Glover"	574	1.000	3.52
a ₂	:	:	:	:
:	:	:	:	:
a ₆₀	"Move-away box green box red"	26.8	0.987	3.48

とが可能である。したがってロボット対話では、ユーザ・ロボット共に実行可能な語彙（「載せる」など）を用いた双方向的なコミュニケーションを考えることができる。

前節までは、ユーザが指示を行い、ロボットが確認発話や物体操作を行う枠組みについて述べてきた。このようなインタラクションでは、学習フェーズと実行フェーズを明示的に分けるか否かに関わらず、学習の過程で過学習を防ぐための動作失敗が必要である。一方、ロボットが発話し、ユーザに物体操作させることによって確信度のパラメータを学習できれば、学習サンプル数を少なくすることができると考えられる。

それでは、どのような発話を行えばよいのだろうか？ここで難しい点は、ある状況において可能な動作は少なくないうえ、ある動作を表す言語表現も唯一ではないことである。

この課題に対し、我々はモデルの学習に有効な発話を生成してユーザに動作を実行させるアプローチを提案した [Sugiura 11c]。発話の生成には、能動学習の一種である Expected Log Loss Reduction (ELLR) [Roy 01] を用いている。ELLR は、テストセットにおける誤分類の期待値を最も低くするような入力を選択する。ELLR ベースの発話生成では、マージンを入力として確信度関数のパラメータを効率的に学習するように発話を選択される。実験では、発話指示が曖昧なため行動失敗を続けるユーザに対しては、発話がより明確になることが確認された。

提案手法による発話生成の例を図8に示す。図8左図の状況においては、182個の発話候補が生成された。図8右表は、各発話のマージン、確信度 (ICM)、対数損失の期待値 (ELL) を示したものである。この状況では、発話「赤い箱 飛び越えさせて (Jump-over red box)」が最小の ELL を持つため、この発話が音声合成され、それを受けてユーザが動作を実行する。実験の結果、能動学習による対話戦略により、学習フェーズにおけるサンプル数および動作失敗数を低減できることがわかった。

4.4 コミュニケーションに意味付けされる単語の学習

ここまでは、実世界の事象にグラウンドした発話の意味の学習と生成について述べてきた。一方、人の発話には「なに」「どれ」などの疑問詞や、「おはよう」や「バイ

バイ」といった挨拶など、実世界の事象ではなくコミュニケーション機能そのものにグラウンドした単語も含まれる。これら2種類の単語と状況の対応を学習できれば、より自然で複雑な発話を理解できるであろう。

田口らは、発話応答を表現するダイナミックグラフィカルモデルを学習する手法を提案し、人の発話意図を推定して適切な応答を生成させている [Taguchi 09]。ロボットは学習の過程で、「なに？」と聞かれたらオブジェクトの名称を答え、「どれ？」と聞かれたらオブジェクトを指し示すことができるようになる。モデルの詳細については、[Taguchi 09] を参照されたい。

5. おわりに

人口構造の変化や共働き世帯の増加とともに、生活環境で人間と共存・支援するロボットへの期待が高まっている。実際に、掃除ロボットの国内市場規模は2006年に約2万台、2008年に約5万台であったものが、2010年には約26万台になっている。今後、人間と共存するロボットがさらに普及するにつれ、ロボットのコミュニケーション機能が重要になるであろう。

国際会議 HRI (Human-Robot Interaction) や AAAI の "Dialog with Robots" シンポジウムなど、ロボットのコミュニケーション機能を扱う研究は存在感を増している。さらに2011年から始まった DARPA の BOLT (Broad Operational Language Translation) プロジェクトでは、音声翻訳や対話と並んでグラウンドした言語獲得が主要なテーマになっている。ただし現状でも、ロボットの音声対話機構に必要な要素技術の開発は、音声言語処理分野とロボティクス分野で個々に進められることが多く、研究者の交流は十分とはいえない。今後、記号創発ロボティクスでは、コミュニケーションの学習論の深化を目指し、音声・画像・動作・言語の学習を統合した研究領域を開いていく。

本稿では、実世界にグラウンドした対話を行うロボットの開発に関する我々の取り組みについて紹介した。本稿で紹介した我々のロボットの動画は、http://komeisugiura.jp/video_gallery/video_gallery_jp.html で閲覧可能である。

謝 辞

本研究の一部は、科研費（若手 (B)24700188, および新学術領域「人口ロボット共生学」公募課題 24118710）の助成を受けて実施されたものである。

◇ 参 考 文 献 ◇

- [Bohus 06] Bohus, D., Langner, B., Raux, A., Black, A., Eskenazi, M., and Rudnick, A.: Online supervised learning of non-understanding recovery policies, in *Proceedings of the IEEE/ACL Workshop on Spoken Language Technology*, pp. 170–173 (2006)
- [Dale 95] Dale, R. and Reiter, E.: Computational interpretations of the Gricean maxims in the generation of referring expressions, *Cognitive Science*, Vol. 19, No. 2, pp. 233–263 (1995)
- [Funakoshi 09] Funakoshi, K., Spanger, P., Nakano, M., and Tokunaga, T.: A probabilistic model of referring expressions for complex objects, in *Proceedings of the 12th European Workshop on Natural Language Generation*, pp. 191–194 (2009)
- [Harnad 90] Harnad, S.: The Symbol Grounding Problem, *Physica D*, Vol. 42, pp. 335–346 (1990)
- [Inamura 04] Inamura, T., Tushima, I., Tanie, H., and Nakamura, Y.: Embodied symbol emergence based on mimesis theory, *International Journal of Robotics Research*, Vol. 23, No. 4, pp. 363–377 (2004)
- [Iwahashi 07] Iwahashi, N.: Robots That Learn Language: Developmental Approach to Human-Machine Conversations, in Sanker, N. ed., *Human-Robot Interaction*, pp. 95–118, I-Tech Education and Publishing (2007)
- [Iwahashi 09] Iwahashi, N., Taguchi, R., Sugiura, K., Funakoshi, K., and Nakano, M.: Robots that Learn to Converse: Developmental Approach to Situated Language Processing, in *Proceedings of International Symposium on Speech and Language Processing*, pp. 532–537 (2009)
- [Jordan 05] Jordan, P. and Walker, M.: Learning content selection rules for generating object descriptions in dialogue, *Journal of Artificial Intelligence Research*, Vol. 24, No. 1, pp. 157–194 (2005)
- [Katagiri 98] Katagiri, S., Juang, B., and Lee, C.: Pattern recognition using a family of design algorithms based upon the generalized probabilistic descent method, *Proceedings of the IEEE*, Vol. 86, No. 11, pp. 2345–2373 (1998)
- [Kawahara 98] Kawahara, T., Lee, C., and Juang, B.: Flexible speech understanding based on combined key-phrase detection and verification, *IEEE Transactions on Speech and Audio Processing*, Vol. 6, No. 6, pp. 558–568 (1998)
- [Kollar 10] Kollar, T., Tellex, S., Roy, D., and Roy, N.: Toward understanding natural language directions, in *Proceeding of the 5th ACM/IEEE international conference on Human-robot interaction*, pp. 259–266 (2010)
- [Komatani 00] Komatani, K. and Kawahara, T.: Flexible mixed-initiative dialogue management using concept-level confidence measures of speech recognizer output, in *Proceedings of the 18th conference on Computational Linguistics*, pp. 467–473 (2000)
- [Langacker 87] Langacker, R. W.: *Foundations of Cognitive Grammar: Theoretical Prerequisites*, Stanford Univ Pr (1987)
- [Ogata 07] Ogata, T., Murase, M., Tani, J., Komatani, K., and Okuno, H. G.: Two-way translation of compound sentences and arm motions by recurrent neural networks, in *Proceedings of the 2007 IEEE/RSJ International Conference on Intelligent Robots and Systems*, pp. 1858–1863 (2007)
- [Pfeifer 99] Pfeifer, R. and Scheier, C.: *Understanding Intelligence*, MIT Press, Cambridge, MA. (1999)
- [Roy 01] Roy, N. and McCallum, A.: Toward optimal active learning through sampling estimation of error reduction, in *Proceedings of 18th International Conference on Machine Learning*, pp. 441–448 (2001)
- [Roy 02] Roy, D.: Learning visually grounded words and syntax for a scene description task, *Computer Speech and Language*, Vol. 16, No. 3, pp. 353–385 (2002)
- [Singh 00] Singh, S., Kearns, M., Litman, D., and Walker, M.: Reinforcement learning for spoken dialogue systems, in *Advances in Neural Information Processing Systems 12 (NIPS)* (2000)
- [Sluis 06] Sluis, van der I., Gatt, A., and Deemter, van K.: Manual for the TUNA corpus: Referring expressions in two domains, Technical report, Technical Report AUCS/TR0705, University of Aberdeen (2006)
- [Sugiura 11a] Sugiura, K., Iwahashi, N., and Kashioka, H.: Motion generation by reference-point-dependent trajectory HMMs, in *Proc. IROS*, pp. 350–356 (2011)
- [Sugiura 11b] Sugiura, K., Iwahashi, N., Kashioka, H., and Nakamura, S.: Learning, Generation, and Recognition of Motions by Reference-Point-Dependent Probabilistic Models, *Advanced Robotics*, Vol. 25, No. 6-7, pp. 825–848 (2011)
- [Sugiura 11c] Sugiura, K., Iwahashi, N., Kawai, H., and Nakamura, S.: Situated Spoken Dialogue with Robots Using Active Learning, *Advanced Robotics*, Vol. 25, No. 17, pp. 2207–2232 (2011)
- [Taguchi 09] Taguchi, R., Iwahashi, N., and Nitta, T.: *Learning Communicative Meanings of Utterances by Robots*, pp. 62–72, Springer (2009)
- [Tokuda 00] Tokuda, K., Yoshimura, T., Masuko, T., Kobayashi, T., and Kitamura, T.: Speech parameter generation algorithms for HMM-based speech synthesis, in *Proceedings of ICASSP*, pp. 1315–1318 (2000)
- [Williams 07] Williams, J. and Young, S.: Scaling POMDPs for Spoken Dialog Management, *IEEE Transactions on Audio Speech and Language Processing*, Vol. 15, No. 7, pp. 2116–2129 (2007)
- [稲邑 01] 稲邑 哲也, 稲葉 雅幸, 井上 博允: ユーザとの対話に基づく段階的な行動決定モデルの獲得, 日本ロボット学会誌, Vol. 19, No. 8, pp. 983–990 (2001)
- [杉浦 10] 杉浦 孔明, 岩橋 直人, 柏岡 秀紀, 中村 哲: 日用品を扱う物体操作の模倣学習による動作の言語化, 第 24 回人工知能学会全国大会 (2010)
- [高野 09] 高野 渉, 中村 仁彦: 統計的相関に基づく動作パターンのリアルタイム教師なし分節化と原始シンボルの自律的獲得, 日本ロボット学会誌, Vol. 27, No. 9, pp. 1046–1057 (2009)
- [谷口 10] 谷口 忠大: コミュニケーションするロボットは創れるか - 記号創発システムへの構成論的アプローチ, NTT 出版 (2010)
- [長井 12] 長井 隆行, 中村 友昭: マルチモーダルカテゴリゼーション, 人工知能学会誌, Vol. 27, p. 6 (2012)
- [西田 03] 西田 豊明: 人とロボットの意思疎通, 情報処理, Vol. 44, No. 12 (2003)
- [松井 00] 松井 俊浩, 麻生 英樹, John, F., 浅野 太, 本村 陽一, 原 功, 栗田 多喜夫, 速水 悟, 山崎 信行: オフィス移動ロボット Jijo-2 の音声対話システム, 日本ロボット学会誌, Vol. 18, No. 2, pp. 142–149 (2000)
- [山内 11] 山内 祐輝, 杉浦 孔明, Sakriani, S., 戸田 智基, 中村 哲: 物体操作タスクにおける HMM を用いた障害物回避動作の生成, 第 12 回計測自動制御学会 システムインテグレーション部門講演会, pp. 1614–1617 (2011)
- [山肩 04] 山肩 洋子, 河原 達也, 奥乃 博, 美濃 導彦: 音声対話システムにおける物体指示のための信念ネットワークを用いた曖昧性の解消, 人工知能学会論文誌, Vol. 19, No. 1, pp. 47–56 (2004)

[担当委員: ××○○]

19YY 年 MM 月 DD 日 受理

—— 著 者 紹 介 ——

杉浦 孔明(正会員)

2007 年京都大学大学院情報学研究所博士後期課程修了。博士 (情報学)。日本学術振興会特別研究員, ATR 音声言語コミュニケーション研究所研究員を経て, 2009 年より情報通信研究機構専攻研究員。知能ロボティクス, 音声対話システム, 機械学習の研究に従事。日本ロボット学会, 計測自動制御学会, IEEE などの会員。