

NAIST-IS-MT1551047

## 修士論文

# 仮想環境を用いた概念獲得のシミュレーションと転移 学習

笹野 仁

2018 年 3 月 14 日

奈良先端科学技術大学院大学  
情報科学研究科 情報科学専攻

本論文は奈良先端科学技術大学院大学情報科学研究科に  
修士（工学）授与の要件として提出した修士論文である。

笹野 仁

審査委員：

中村 哲 教授           （主指導教員）

松本 裕治 教授       （副指導教官）

吉野 幸一郎 助教     （副指導教官）

# 仮想環境を用いた概念獲得のシミュレーションと転移学習\*

笹野 仁

## 内容梗概

人間は視覚や聴覚, 触覚といった五感から得られるマルチモーダルな情報を用いて物体の概念を獲得していると言われている. こうした物体の概念獲得を, ロボットを用いて教師なしで学習する試みが行われている. このようにして獲得された概念は生活支援ロボットなど実世界でユーザと接触するシステムでの応用が考えられているが, 人間が普段の生活で扱う物体の数は膨大で, 学習対象が消耗品や高価である場合や, ロボット自体を長時間運用する必要があることなどが大きなコストになる. このような問題には転移学習が有効となるケースがある. ロボットによる物体概念獲得のシミュレーションを仮想環境上で行い, その学習結果を実機に転用できれば, こうしたコストを低減することができる. そこで, 本研究ではマルチモーダル情報を用いたロボットによる物体の概念獲得において転移学習が有効であることを示す. そのために, 転移学習で用いる具体的な概念学習アルゴリズムを既存研究で用いられたマルチモーダル LDA を拡張することで提案した. また, 実ロボットと仮想環境の両方で概念獲得の結果が一致することを示した. 最後に, 提案した転移学習アルゴリズムが有効であることを確認した.

## キーワード

マルチモーダル LDA, 概念獲得, 転移学習

---

\*奈良先端科学技術大学院大学 情報科学研究科 情報科学専攻 修士論文, NAIST-IS-MT1551047, 2018 年 3 月 14 日.

# Simulation and transfer learning of concept learning using virtual environment\*

Jin Sasano

## Abstract

Human beings acquire the concepts of objects by using multimodal information that comes from five senses, such as visual, auditory and tactile senses. Some existing works tried to learn the concept acquisition (concept learning) in unsupervised manner using robots. The concepts acquired in this manner will be used by assistant robots that communicate with people in daily life, such as living support robots. However the number of objects handled in a daily life is enormous, and it increases a cost of learning if the learning process requires expensive items. In addition, the operation cost of the robot itself for a long time is not insignificant. There are cases where transfer learning is effective for such problems. If we can run the concept learning on a robot model in a simulation environment for transferring the result to the real robot on actual environment, we can reduce these costs. In this research, we show that transfer learning is effective in concept learning of robot if we use multimodal information. We proposed a learning algorithm of transfer learning with extending the model of multi-modal latent Dirichlet allocation. We showed that the results of concept learning are transferable between both of real robot and the robot in virtual environment. We also investigated that the proposed transfer learning algorithm improved the accuracy of transfer learning from the simulation environment to the real environment.

## Keywords:

---

\*Master's Thesis, Department of Information Science, Graduate School of Information Science, Nara Institute of Science and Technology, NAIST-IS-MT1551047, March 14, 2018.

マルチモーダル LDA, 転移学習, 概念獲得

# 目次

|                               |     |
|-------------------------------|-----|
| 図目次                           | vi  |
| 表目次                           | vii |
| 第 1 章 緒言                      | 1   |
| 1.1 背景                        | 1   |
| 1.2 目的                        | 2   |
| 1.3 本論文の構成                    | 2   |
| 第 2 章 トピックモデルを用いた概念獲得         | 3   |
| 2.1 トピックモデル                   | 3   |
| 2.2 LDA                       | 4   |
| 2.3 マルチモーダル LDA               | 5   |
| 2.4 ギブスサンプリングによるパラメータ推定       | 7   |
| $z_{d,i}$ のサンプリング             | 7   |
| $\theta_d$ のサンプリング            | 8   |
| $\phi_k$ のサンプリング              | 10  |
| 第 3 章 ロボットを用いた概念獲得            | 12  |
| 3.1 実ロボットを用いた概念獲得             | 13  |
| 3.1.1 Bag-of-Features モデル     | 16  |
| 3.1.2 視覚情報の処理                 | 16  |
| 3.1.3 聴覚情報の処理                 | 16  |
| 3.1.4 触覚情報の処理                 | 17  |
| 3.1.5 LDA と MLDA による教師なし分類    | 17  |
| 3.2 シミュレーション環境における概念獲得        | 18  |
| 3.3 MLDA を用いた転移学習             | 21  |
| 第 4 章 評価実験                    | 22  |
| 4.1 実環境とシミュレーション環境における概念獲得の一致 | 22  |

|       |                                   |    |
|-------|-----------------------------------|----|
| 4.1.1 | 概念獲得の再現実験 . . . . .               | 22 |
| 4.2   | シミュレーション環境における概念獲得との比較 . . . . .  | 25 |
| 4.2.1 | 考察 . . . . .                      | 30 |
|       | トピックパラメータの正規化による分類制度の向上 . . . . . | 30 |
| 4.3   | 物体の概念獲得における転移学習 . . . . .         | 32 |
| 4.3.1 | 実験設定 . . . . .                    | 32 |
| 4.3.2 | 結果 . . . . .                      | 32 |
| 第 5 章 | 結言 . . . . .                      | 35 |
| 5.1   | 本論文のまとめ . . . . .                 | 35 |
| 5.2   | 今後の課題 . . . . .                   | 35 |
|       | 謝辞 . . . . .                      | 36 |
|       | 参考文献 . . . . .                    | 37 |
|       | 発表リスト . . . . .                   | 40 |

## 図目次

|     |                                               |    |
|-----|-----------------------------------------------|----|
| 2.1 | LDA のグラフィカルモデル . . . . .                      | 5  |
| 2.2 | MLDA のグラフィカルモデル . . . . .                     | 6  |
| 3.1 | ロボットのハンドに設置された 3 つの圧力センサー . . . . .           | 13 |
| 3.2 | ロボットが認識する 9 種類の物体 . . . . .                   | 14 |
| 3.3 | ロボットが物体に対して行う 3 種類の行動 . . . . .               | 15 |
| 3.4 | ロボットのハンドに設置された 3 つの圧力センサー . . . . .           | 17 |
| 3.5 | 仮想環境上でロボットが物体に対して行う 3 種類の行動 . . . . .         | 19 |
| 3.6 | 仮想環境上でロボットが認識する 9 種類の物体 . . . . .             | 20 |
| 4.1 | 実機での概念獲得において用いる情報の種類と分類精度の変化 . . .            | 24 |
| 4.2 | 実機で 9 種類の物体の識別を 100 回行った場合の精度分布 . . . . .     | 25 |
| 4.3 | 仮想環境上での概念獲得において用いる情報の種類と分類精度の<br>変化 . . . . . | 27 |
| 4.4 | 仮想環境上で 9 種類の物体の識別を 100 回行った場合の精度分布 .          | 27 |
| 4.5 | 実機と仮想環境上での実験結果の比較 . . . . .                   | 29 |
| 4.6 | 実機と仮想環境上で 9 種類の物体を 100 回識別した場合の精度分布           | 30 |
| 4.7 | 各物体情報におけるトピックパラメータの値 . . . . .                | 31 |
| 4.8 | 9 種類 3 カテゴリの転移学習の精度の変化 . . . . .              | 33 |
| 4.9 | 9 種類 9 カテゴリの転移学習の精度の変化 . . . . .              | 33 |

## 表目次

|      |                       |    |
|------|-----------------------|----|
| 2.1  | 3つの文書データセット . . . . . | 3  |
| 4.1  | 各物体の性質 . . . . .      | 22 |
| 4.2  | 視覚 . . . . .          | 23 |
| 4.3  | 聴覚 . . . . .          | 23 |
| 4.4  | 触覚 . . . . .          | 23 |
| 4.5  | 視覚・聴覚 . . . . .       | 23 |
| 4.6  | 聴覚・触覚 . . . . .       | 23 |
| 4.7  | 触覚・視覚 . . . . .       | 23 |
| 4.8  | 視覚・聴覚・触覚 . . . . .    | 24 |
| 4.9  | 視覚 . . . . .          | 25 |
| 4.10 | 聴覚 . . . . .          | 25 |
| 4.11 | 触覚 . . . . .          | 26 |
| 4.12 | 視覚・聴覚 . . . . .       | 26 |
| 4.13 | 聴覚・触覚 . . . . .       | 26 |
| 4.14 | 触覚・視覚 . . . . .       | 26 |
| 4.15 | 視覚・聴覚・触覚 . . . . .    | 26 |
| 4.16 | 視覚 . . . . .          | 28 |
| 4.17 | 聴覚 . . . . .          | 28 |
| 4.18 | 触覚 . . . . .          | 28 |
| 4.19 | 視覚・聴覚 . . . . .       | 28 |
| 4.20 | 聴覚・触覚 . . . . .       | 28 |
| 4.21 | 触覚・視覚 . . . . .       | 28 |
| 4.22 | 視覚・聴覚・触覚 . . . . .    | 29 |

# 第 1 章 緒言

## 1.1 背景

人間は、視覚や聴覚、触覚のような五感から得られる様々な情報を用いて物体の概念を獲得するといわれている。このような物体概念の獲得を、身体性を持つロボットが得る様々な情報から教師なしで学習する試みが行われている [1][2]。こうした学習の枠組みでは、人間が行うようにその物体に関する様々な情報を得るための試行を繰り返し、その結果の類似から試行に用いた物体のクラスタリングを行う。このクラス一つ一つを、ロボットが獲得した物体概念として扱う。生活支援システムや対話システムの研究 [3, 4, 5, 6, 7] ではロボットが人間の身体性に基づいた実世界の物体概念を扱う必要があるとされており、ロボットにより学習された物体概念はこのようなシステムへの応用が考えられている [8, 9]。

しかし、人間が普段の生活で扱う物体の数は膨大で、これらすべての物体概念を実際のロボットを用いて獲得することは難しい。これは、大量の学習対象の入手やロボットの運用などの物理的な制約が大きなコストとなるためである。このようにデータの入手や学習自体が大きなコストとなるタスクでは、転移学習が有効となるケースがある。転移学習とは、機械学習において、あるドメインでの学習結果を別のドメインに適用する技術である。ロボットを用いた機械学習においては、データを取得するコストが低いドメイン（シミュレーション環境）で学習を行い、その結果をデータを取得するコストが高いドメイン（実環境）に適応することで学習コストの低減や効率化を実現できることが知られている [10, 11]。転移学習を物体概念の獲得に応用し、ロボットによる物体概念の獲得のシミュレーションとして、仮想環境上で物体概念の獲得を行う。仮想環境上での学習結果を実環境上のロボット実機に転用することができれば、実機の運用などの学習に伴うコストを低減することができる。

## 1.2 目的

本研究では，マルチモーダル情報を用いたロボットによる物体概念の獲得にかかるコストを低減するため，転移学習の有効性を示す．そのためにまず，ロボットが実環境と仮想環境の両方で物体概念の獲得が可能であることを確認する．具体的には，関連研究 [1] で行われているように，ロボット実機での視覚と聴覚，触覚を利用した概念獲得を再現する．仮想環境上でロボットが物体概念の獲得が可能であることを確認するために，同様の物体クラスに対して，視覚と聴覚，触覚を再現した仮想環境上での概念獲得を行う（シミュレーション）．実機とシミュレーションでの実験の比較を行い，仮想環境上での概念獲得結果がどの程度実機での概念獲得結果と一致するかを確認する．

また，仮想環境上での概念獲得の結果をロボット実機に転用することでロボット実機が実環境上でも物体の概念を利用するための有効な転移学習アルゴリズムを提案する．また提案手法の有効性を確認するため，提案した転移学習アルゴリズムの仮想環境上での学習データ数を変化させ，実環境上での物体の概念獲得の精度の変化を確認する．

## 1.3 本論文の構成

本論文では，まず，本研究でロボットの学習に用いる統計モデルについて説明し，その応用例と具体的な学習法について述べる (2 章)．次に，ロボットを用いた実環境上と仮想環境上での物体の概念獲得について説明し (3 章)，実環境上と仮想環境上での概念獲得の一致実験と転移学習の実験の結果について述べ (4 章)，最後に，本論文のまとめを行い，今後の課題を示す (5 章)

## 第 2 章 トピックモデルを用いた概念獲得

### 2.1 トピックモデル

トピックモデルとはデータの潜在的な意味を数学的に記述する統計的学習モデルで，代表的な例として潜在的ディリクレ分配法 (Latent Dirichlet Allocation: LDA) [12] がある．自然言語処理の分野では文書データに潜む意味という概念的な情報を数学的に記述し，解析する手法として潜在的意味解析に用いられている．潜在的意味解析では意味の情報を扱う一つの見通しとして単語の共起性に着目する．例えば表 2.1 のような 3 つのデータセットを考える．

表 2.1: 3 つの文書データセット

| 文書番号 | 文書 (単語) データ          |
|------|----------------------|
| 1    | テニス，野球，サッカー，バスケットボール |
| 2    | 焼肉，うどん，パスタ，ビザ        |
| 3    | ピアノ，歌，演奏，ドラム，ギター     |

この例では，文書 1 はスポーツ，文書 2 は食べ物，文書 3 は音楽に関する単語の集まりだとわかる．3 つのデータセットにはそれぞれスポーツ，食べ物，音楽といった単語は出現しないが容易に想起することができる．このように複数の単語の共起によって創発される情報を潜在的意味と考える．潜在的意味解析では文書データには複数の潜在的意味が含まれていると考え，それぞれの潜在的意味のカテゴリを潜在トピック，または単にトピックと呼ぶ [13]．トピックモデルは文章データ中の単語と単語の出現頻度の組 (Bag-of-Words: BoW) がトピックの情報をもとに生成されていると仮定した生成モデルである．トピックモデルを用いることで文章データからトピックに関する情報を推定することで文章の潜在的意味を解析することができる．また，トピックモデルを自然言語以外のデータに適応することで様々なデータにおける意味の分析に応用することができる．次節では，本研究で用いる LDA のモデルと学習法について述べる．また，LDA の応用について述べ，LDA の拡張モデルであるマルチモーダル LDA [1] について述べる．

## 2.2 LDA

LDA は文書の確率的生成モデルとして提案された生成モデルである．LDA は文書中の単語の順序は無視し，文書の BoW 表現をモデル化するため，文書データ中の共起しやすい単語の集合を解析するためのモデルと言い代えることもできる．LDA では，文書中には潜在トピックがあると考え，統計的に共起しやすい単語の集合の出現を潜在トピックという観測できない潜在変数から生成されると仮定して定式化する．与えられた文書データに対し，この潜在トピックを計算することにより，文書のトピックや意味に関する様々な情報を分析することができる．

LDA では，トピックを単語の出現頻度で定義する．文書は複数のトピックで構成されており，文書の各単語はまず，各文書を構成するトピックの割合から多項分布に従いトピックが選択され，そのトピックから多項分布に従い単語が選択されると仮定する．また，各文書を構成するトピックの割合を示すパラメータとトピック (トピックパラメータ) がディリクレ分布から生成されると仮定する．

分析する文書数を  $M$ ，文書番号を  $d$ ，文書に出現する単語の種類数を  $V$ ，文書を構成するトピック数を  $K$  で表す．また，トピックは要素の和が 1 になるような  $V$  次元のベクトル  $\phi_k (\phi_k = (\phi_{k,1}, \dots, \phi_{k,V}), \sum_{v=1}^V \phi_{k,v} = 1)$  と表し， $d$  番目の文書のトピックパラメータを要素の総和が 1 になるような  $K$  次元のベクトル  $\theta_d (\theta_d = (\theta_{d,1}, \dots, \theta_{d,K}), \sum_{k=1}^K \theta_{d,k} = 1)$  で表す．さらに， $d$  番目の文書の単語数を  $n_d$ ， $d$  番目の文書の  $i$  番目の単語を  $w_{d,i}$  と表し，その単語が選出されたトピック番号の情報を潜在変数  $z_{d,i} (z_{d,i} \in \{1, \dots, K\})$  で表す． $z_{d,i}$  と  $w_{d,i}$  はそれぞれトピックパラメータとそれが示すトピックから多項分布に従って生成されている．多項分布を  $Multi(\cdot)$  とすると，

$$z_{d,i} \sim P(z_{d,i} | \theta_d) = Multi(\theta_d) \quad (2.2.1)$$

$$w_{d,i} \sim P(w_{d,i} | \phi_{z_{d,i}}) = Multi(\phi_{z_{d,i}}) \quad (2.2.2)$$

と表すことができる．また， $\phi_k$  と  $\theta_d$  はそれぞれディリクレ分布から生成されるものと仮定されている．これらの分布は  $\phi_k$  と  $\theta_d$  を生成するディリクレ分布の  $V$  次元と  $K$  次元のパイパーパラメータを  $\beta$  と  $\alpha$  とおき，ディリクレ分布を  $Dir(\cdot)$  とすると，

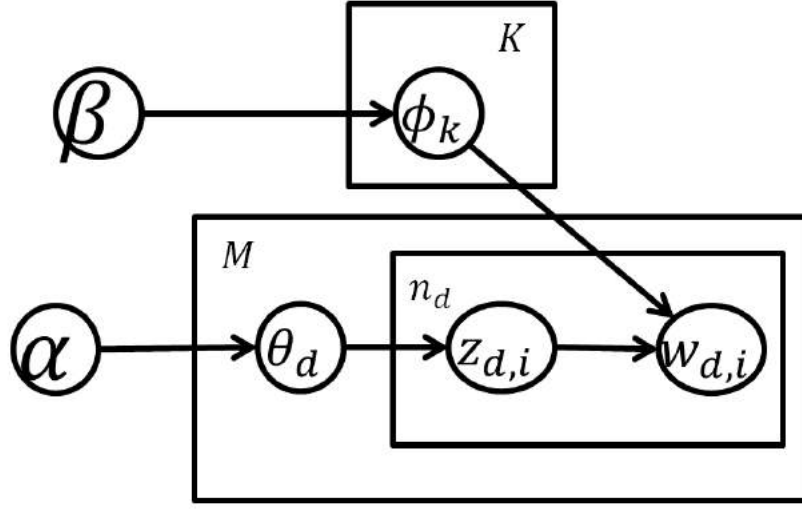


図 2.1: LDA のグラフィカルモデル

$$\theta_d \sim P(\theta_d|\alpha) = \text{Dir}(\alpha) \quad (2.2.3)$$

$$\phi_k \sim P(\phi_k|\beta) = \text{Dir}(\beta) \quad (2.2.4)$$

と表すことができる．LDA のグラフィカルモデルを図 2.1 に示す．

これまで様々な LDA の改良が提案され，総称して潜在トピックモデル (Latent topic model) と呼ばれている．また，LDA は自然言語処理に限らず，データ中の特徴とその特徴の出現頻度の組 (Bag-of-Features: BoF) を用いることで画像処理，音声処理，ソーシャルネットワーク解析，情報推薦など，様々な分野へ応用されている [14][15][16] ．

## 2.3 マルチモーダル LDA

マルチモーダル LDA (Multimodal LDA: MLDA) は中村らによって提案された LDA の拡張モデルである [1]．中村らは，人間のマルチモーダルな情報を用いた物体のカテゴリゼーション能力の重要性が指摘しており [17][18]，ロボットが様々な環境において柔軟に物体認識を行う際には同様にマルチモーダル情報を用いてカテ

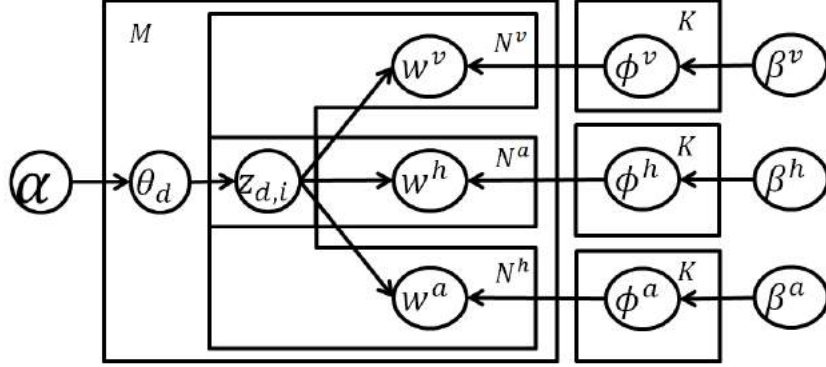


図 2.2: MLDA のグラフィカルモデル

ゴリー認識を教師なしで行うことが有効であることが示されている．マルチモーダル情報を用いた教師なしでの物体のカテゴリ分類にトピックモデルを用いることで，学習サンプルの分類だけでなく，未学習物体のカテゴリ認識や，未観測のモダリティ情報の推定が可能であることが示されている．モデルとしては，LDA をマルチモーダル情報を扱えるように拡張したモデルである MLDA が提案され，MLDA を用いることで，マルチモーダル情報を用いたロボットは人間の感覚に即した物体のカテゴリ（概念）を形成できることを示している [19][20]．MLDA のグラフィカルモデルを図 2.2 に示す．

MLDA の  $w^v$ ,  $w^a$ ,  $w^h$  は  $M$  個の物体情報における，視覚情報と聴覚情報，触覚情報の BoF インデックスを表し，LDA における文書の単語情報を表す  $w$  に相当する変数である．また， $\phi^v$ ,  $\phi^a$ ,  $\phi^h$  は LDA における文書中の単語の出現頻度情報である  $\phi$  に相当する変数で  $\beta^v$ ,  $\beta^a$ ,  $\beta^h$  はそれらを生成するディリクレ分布のハイパーパラメータである．MLDA の学習では LDA の文書集合の場合と同様に  $M$  個の各物体情報に一つずつ  $\theta_d$  が計算される． $\theta_d$  はトピックパラメータとも呼ばれ，あらかじめ定めたカテゴリ数  $K$  と等しい次元を持つベクトルで，MLDA においては  $d$  番目の物体情報（視覚情報，聴覚情報と触覚情報）の BoF インデックスを生成するカテゴリの比率を表している．関連研究 [1] ではこのトピックパラメータのカテゴリの比率をもとに物体のカテゴリの識別を行っている．

## 2.4 ギブスサンプリングによるパラメータ推定

本章では LDA と MLDA のギブスサンプリング [21] によるパラメータ推定について説明する．まず， $w_{d,i}$  全体の集合として変数  $w(= \{w_{d,i}\})$  と定義し， $w$  からある単語情報  $w_{d',i'}$  を除いた変数を  $w^{\setminus d',i'}$  と表す． $z_{d,i}$  の場合も同様に  $z$  と  $z^{\setminus d',i'}$  を定義する．また， $\theta_d$  全体の集合を  $\theta$  と定義し， $\theta$  からあるトピックパラメータ  $\theta_d'$  を除いた変数を  $\theta^{\setminus d'}$  と表す． $\phi_k$  の場合も同様に  $\phi$  と  $\phi^{\setminus k'}$  を定義する．これに加え，指示関数  $\delta()$  を以下のように定義する．

$$\delta(z_{d,i} = k) = \begin{cases} 1 & (z_{d,i} = k) \\ 0 & (otherwise) \end{cases} \quad (2.4.1)$$

$$\delta(z_{d,i} = k, w_{d,i} = v) = \begin{cases} 1 & (z_{d,i} = k, w_{d,i} = v) \\ 0 & (otherwise) \end{cases} \quad (2.4.2)$$

LDA において，ディリクレ分布のハイパーパラメータの  $\alpha$ ,  $\beta$  と単語情報  $w$  が与えられている場合に  $z$ ,  $\phi$ ,  $\theta$  をサンプリングするには事後分布  $P(z, \phi, \theta | w, \alpha, \beta)$  を構成すればよいが，このような同時分布は解析的に計算が困難で，計算コストも非常に高いため現実的ではない．そこで，ギブスサンプリングを用いることにより計算が簡単な条件つき確率を構成することでこのような同時分布からのサンプリングを可能にする．具体的には，ギブスサンプリングで値を計算する変数  $z_{d,i}$ ,  $\phi_k$ ,  $\theta_d$  の事後分布を構成し，それらを交互にサンプリングすることでパラメータを推定する．

$$z_{d,i} \sim P(z_{d,i} | z^{\setminus d,i}, \theta, \phi, w, \alpha, \beta) \quad (2.4.3)$$

$$\theta_d \sim P(\theta_d | z, \theta^{\setminus d}, \phi, w, \alpha, \beta) \quad (2.4.4)$$

$$\phi_k \sim P(\phi_k | z, \theta, \phi^{\setminus k}, w, \alpha, \beta) \quad (2.4.5)$$

### $z_{d,i}$ のサンプリング

LDA は図 2.1 のグラフィカルモデルからその結合確率を以下のように計算することができる．

$$P(z, \phi, \theta, w, \alpha, \beta) = P(w, z|\phi)P(z|\theta)P(\theta|\alpha)P(\alpha)P(\phi|\beta)P(\beta) \quad (2.4.6)$$

これを用い,

$$P(z_{d,i}|z^{\setminus d,i}, \theta, \phi, w, \alpha, \beta) = \frac{P(z, \phi, \theta, w, \alpha, \beta)}{\Sigma_{z_{d,i}} P(z, \phi, \theta, w, \alpha, \beta)} \quad (2.4.7)$$

$$= \frac{P(w|z, \phi)P(z|\theta)P(\theta|\alpha)P(\alpha)P(\phi|\beta)P(\beta)}{\Sigma_{z_{d,i}} P(w|z, \phi)P(z|\theta)P(\theta|\alpha)P(\alpha)P(\phi|\beta)P(\beta)} \quad (2.4.8)$$

$$= \frac{P(w_{d,i}|z_{d,i}, \phi)P(z_{d,i}|\theta)P(w^{\setminus d,i}|z^{\setminus d,i}, \phi)P(z^{\setminus d,i}|\theta)}{\Sigma_{z_{d,i}} P(w_{d,i}|z_{d,i}, \phi)P(z_{d,i}|\theta)P(w^{\setminus d,i}|z^{\setminus d,i}, \phi)P(z^{\setminus d,i}|\theta)} \quad (2.4.9)$$

$$= \frac{\theta_{d,z_{d,i}} \phi_{z_{d,i},w_{d,i}}}{\Sigma_{z_{d,i}=1}^K \theta_{d,z_{d,i}} \phi_{z_{d,i},w_{d,i}}} \quad (2.4.10)$$

の確率に基づいて  $z_{d,i}$  のサンプリングを行う.

$\theta_d$  のサンプリング

$$P(\theta_d|z, \theta^{\setminus d}, \phi, w, \alpha, \beta) = \frac{P(z, \phi, \theta, w, \alpha, \beta)}{\int P(z, \phi, \theta, w, \alpha, \beta) d\theta_d} \quad (2.4.11)$$

$$= \frac{P(w|z, \phi)P(z|\theta)P(\theta|\alpha)P(\alpha)P(\phi|\beta)P(\beta)}{\int P(w|z, \phi)P(z|\theta)P(\theta|\alpha)P(\alpha)P(\phi|\beta)P(\beta) d\theta_d} \quad (2.4.12)$$

$$= \frac{P(z|\theta)P(\theta|\alpha)}{\int P(z|\theta)P(\theta|\alpha) d\theta_d} \quad (2.4.13)$$

$$= \frac{P(z_d|\theta_d)P(z^{\setminus d}|\theta^{\setminus d})P(\theta_d|\alpha)P(\theta^{\setminus d}|\alpha)}{\int P(z_d|\theta_d)P(z^{\setminus d}|\theta^{\setminus d})P(\theta_d|\alpha)P(\theta^{\setminus d}|\alpha) d\theta_d} \quad (2.4.14)$$

$$= \frac{P(z_d|\theta_d)P(\theta_d|\alpha)}{\int P(z_d|\theta_d)P(\theta_d|\alpha) d\theta_d} \quad (2.4.15)$$

ここで式 2.4.15 の分母は定数であり, 分子の  $P(z_d|\theta_d)$  と  $P(\theta_d|\alpha)$  はそれぞれ, 多項分布とディリクレ分布であるから.

$$P(z_d|\theta_d) \propto \prod_{i=1}^{n_d} \theta_{d,z_{d,i}} \quad (2.4.16)$$

$$P(\theta_d|\alpha) \propto \prod_{k=1}^{K=1} \theta_{d,k}^{\alpha_k-1} \quad (2.4.17)$$

また，式 2.4.16 は  $d$  番目の文書の潜在変数  $z_{d,i}$  がトピック番号  $k$  と一致した回数を  $n_{d,k} = \sum_{i=1}^{n_d} \delta(z_{d,i} = k)$  とおくと

$$\prod_{i=1}^{n_d} \theta_{d,z_{d,i}} = \prod_{k=1}^K \theta_{d,k}^{n_{d,k}} \quad (2.4.18)$$

と表すことができるので

$$P(\theta_d|z, \theta^{\setminus d}, \phi, w, \alpha, \beta) \propto \prod_{k=1}^K \theta_{d,k}^{\alpha_k + n_{d,k} - 1} \quad (2.4.19)$$

であると分かる．ここで，多項分布とディリクレ分布の共役性から式 2.4.31 もまたディリクレ分布である．よって

$$P(\theta_d|z, \theta^{\setminus d}, \phi, w, \alpha, \beta) = \frac{\Gamma(\sum_{k=1}^K (\alpha_k + n_{d,k}))}{\prod_{k=1}^K \Gamma(\alpha_k + n_{d,k})} \prod_{k=1}^K \theta_{d,k}^{\alpha_k + n_{d,k} - 1} \quad (2.4.20)$$

の確率に基づいて  $\theta_d$  のサンプリングを行う．

$\phi_k$  のサンプリング

$$P(\phi_k|z, \theta, \phi^{\setminus k}, w, \alpha, \beta) = \frac{P(z, \phi, \theta, w, \alpha, \beta)}{\int P(z, \phi, \theta, w, \alpha, \beta) d\phi_k} \quad (2.4.21)$$

$$= \frac{P(w|z, \phi)P(z|\theta)P(\theta|\alpha)P(\alpha)P(\phi|\beta)P(\beta)}{\int P(w|z, \phi)P(z|\theta)P(\theta|\alpha)P(\alpha)P(\phi|\beta)P(\beta) d\phi_k} \quad (2.4.22)$$

$$= \frac{P(w|z, \phi)P(\phi_k|\beta)P(\phi^{\setminus k}|\beta)}{\int P(w|z, \phi)P(\phi|\beta)P(\phi^{\setminus k}|\beta) d\phi_k} \quad (2.4.23)$$

$$= \frac{P(w|z, \phi)P(\phi_k|\beta)}{\int P(w|z, \phi)P(\phi_k|\beta) d\phi_k} \quad (2.4.24)$$

ここで、式 2.4.24 の分子の  $P(w|z, \phi)$  と  $P(\phi_k|\beta)$  はそれぞれ多項分布とディリクレ分布なので、潜在変数  $z_{d,i}$  と単語  $w_{d,i}$  をそれぞれ  $k, v$  とおくと

$$P(w|z, \phi) \propto \prod_{d=1}^M \prod_{i=1}^{n_d} \phi_{z_{d,i}, w_{d,i}} \quad (2.4.25)$$

$$P(\phi_k|\beta) \propto \prod_{v=1}^V \phi_{k,v}^{\beta_v-1} \quad (2.4.26)$$

と表せる。全ての  $z_{d,i}$  と  $w_{d,i}$  についてトピック  $k'$  が選ばれたときに単語  $v$  が選ばれる回数を  $n_{k',v} = \sum_{d=1}^M \sum_{i=1}^{n_d} \delta(z_{d,i} = k', w_{d,i} = v)$  とおくと式 2.4.25 は

$$\prod_{d=1}^M \prod_{i=1}^{n_d} \phi_{z_{d,i}, w_{d,i}} = \prod_{k'=1}^K \prod_{v=1}^V \phi_{k',v}^{n_{k',v}} \quad (2.4.27)$$

と表すことができる。これから 2.4.24 は以下のように計算できる。

$$\frac{P(w|z, \phi)P(\phi_k|\beta)}{\int P(w|z, \phi)P(\phi_k|\beta) d\phi_k} = \frac{\prod_{k'=1}^K \prod_{v=1}^V \phi_{k',v}^{n_{k',v}}}{\int \prod_{k'=1}^K \prod_{v=1}^V \phi_{k',v}^{n_{k',v}} d\phi_k} \quad (2.4.28)$$

$$= \frac{\prod_{k' \neq k}^K \prod_{v=1}^V \phi_{k',v}^{n_{k',v}} \prod_{v=1}^V \phi_{k,v}^{n_{k,v}}}{\prod_{k' \neq k}^K \prod_{v=1}^V \phi_{k',v}^{n_{k',v}} \int \prod_{v=1}^V \phi_{k,v}^{n_{k,v}} d\phi_k} \quad (2.4.29)$$

$$\propto \prod_{v=1}^V \phi_{k,v}^{n_{k,v}} \quad (2.4.30)$$

であると分かる．ここで，多項分布とディリクレ分布の共役性から式 2.4.30 もまたディリクレ分布である．よって

$$P(\phi_k | z, \theta, \phi^{\setminus k}, w, \alpha, \beta) = \frac{\Gamma(\sum_{v=1}^V (\beta_v + n_{k,v}))}{\prod_{v=1}^V \Gamma(\beta_v + n_{k,v})} \prod_{v=1}^V \phi_{k,v}^{\beta_v + n_{k,v} - 1} \quad (2.4.31)$$

の確率に基づいて  $\phi_k$  のサンプリングを行う．

MLDA の場合も，LDA と同様に，結合分布から余因子を削除し，多項分布とディリクレ分布の共役性を利用した計算を行うことで，ディリクレ分布のハイパーパラメータ  $\alpha, \beta^m$  と単語情報  $w^m$  が与えられた場合の潜在変数  $z_{d,i}^m$  トピックパラメータ  $\theta^d$  トピック情報  $\phi_k^m$  を以下のように生成する条件つき確率分布を構成できる．ただし， $m$  は視覚情報と触覚情報，聴覚情報を表す添え字  $a, v, h$  を表すものとし， $*$  は結合分布の変数のうちサンプリングを行う変数以外の変数を表す．

$$z_{d,i}^m \sim P(z_{d,i}^m | *) = \frac{\theta_{d,z_{d,i}^m} \phi_{z_{d,i}^m}^m w_{d,i}^m}{\sum_{z_{d,i}^m=1}^K \theta_{d,z_{d,i}^m} \phi_{z_{d,i}^m}^m w_{d,i}^m} \quad (2.4.32)$$

$$\theta_d \sim P(\theta_d | *) = \frac{\Gamma(\sum_{k=1}^K (\alpha_k + n_{d,k}^a + n_{d,k}^v + n_{d,k}^h))}{\prod_{k=1}^K \Gamma(\alpha_k + n_{d,k}^a + n_{d,k}^v + n_{d,k}^h)} \prod_{k=1}^K \theta_{d,k}^{\alpha_k + n_{d,k}^a + n_{d,k}^v + n_{d,k}^h - 1} \quad (2.4.33)$$

$$\phi_k^m \sim P(\phi_k^m | *) = \frac{\Gamma(\sum_{v=1}^V (\beta_v^m + n_{k,v}^m))}{\prod_{v=1}^V \Gamma(\beta_v^m + n_{k,v}^m)} \prod_{v=1}^V \phi_{k,v}^{m \beta_v^m + n_{k,v}^m - 1} \quad (2.4.34)$$

### 第 3 章 ロボットを用いた概念獲得

人間は、五感から得られる情報を用いて物体を認知し、その繰り返しによって同じ種類のものが同じクラスに属するという概念を獲得すると考えられている [2] . 例えばコーヒーカップを認知する場合は、コーヒーカップを手で触ったときの“硬い”や“つめたい”といった触覚の情報や、目で見たときの“白い”や“滑らかな形”であるといった視覚の情報など、複数の感覚器から得られる情報を統合することで触覚と視覚から得られた情報とは独立したコーヒーカップという物体に関する情報を認知・学習すると言われている [22] . こうした情報は単一の種類の情報ではなく、多様な種類の情報（マルチモーダルな情報）が用いられている．マルチモーダルな情報により、同じ種類のものを同じクラスのものとして分類する、物体概念という物体の認知が形成される．

このような物体概念の獲得を、マルチモーダル情報を用いてロボットにより獲得する試みがすでに研究されている [1] . ロボットにより獲得された物体概念は、様々な応用が考えられる．例えば、生活支援システムの研究 [3, 4] では、システムは人間の身体性に基づいて物体の概念を理解する必要があるとされている．また対話システムの研究においては言葉だけでなく表情や身振り手振り、語調といった複数のチャンネルを使用したマルチモーダルなコミュニケーションに関する議論が盛んに行われている [5, 6, 7] . これらの研究では、人間の身体的な感覚を理解するうえでマルチモーダルな概念獲得が有効であるとされている．しかし、物体概念の学習・獲得のために実際のロボットを用いることはコストが大きい．

例えば、学習対象が消耗品である場合や高価なものである場合などである．また、実ロボットを運用する費用もコストを増大させる要因である．加えて、実ロボットには物理的な制約があり、学習に長い時間を要する．例えば、Aoki らの研究 [23] では 81 カテゴリ 499 個の日用品の物体概念をマルチモーダル情報を用いるロボットにより獲得する試みが行われているが、この研究ではロボットが物体のマルチモーダル情報を取得するために 100 時間以上の手動での作業を要した．

これらに対して、仮想環境で物体概念の獲得を行うことができれば、こうしたコストの問題を解決できる可能性が考えられる．そこで本研究では、マルチモーダル情報を用いたロボットによる物体概念の獲得における転移学習の有効性の検証と

具体的な学習法を提案する．まず，仮想環境で概念獲得が正しく行われていることを確認するため，仮想環境と実環境の両方で物体の概念獲得の結果を比較した．また，仮想環境で学習に用いるデータ数を増やすことで概念獲得の精度の変化を確認した．次節以降では，本研究で行う実験で用いるロボットのマルチモーダル情報からの特徴量抽出について述べ，実ロボットと仮想環境上での物体概念獲得について説明する，また，MLDA を用いた物体情報の教師なし分類について述べ，最後に，MLDA を用いた物体の概念獲得の転移学習について述べる．

### 3.1 実ロボットを用いた概念獲得

まず先行研究同様，実機によって視覚・聴覚・触覚 を用いた概念獲得を行う．実験に用いるロボットとして Aldebaran Robotics の NAO [24] を使用した．ロボットは頭部にカメラ (CMOS 640 x 480 camera) 2 台とマイク 4 台を搭載している．これらに加え，我々はロボットのハンドに圧力センサー [25] を設置した (図 3.1)．

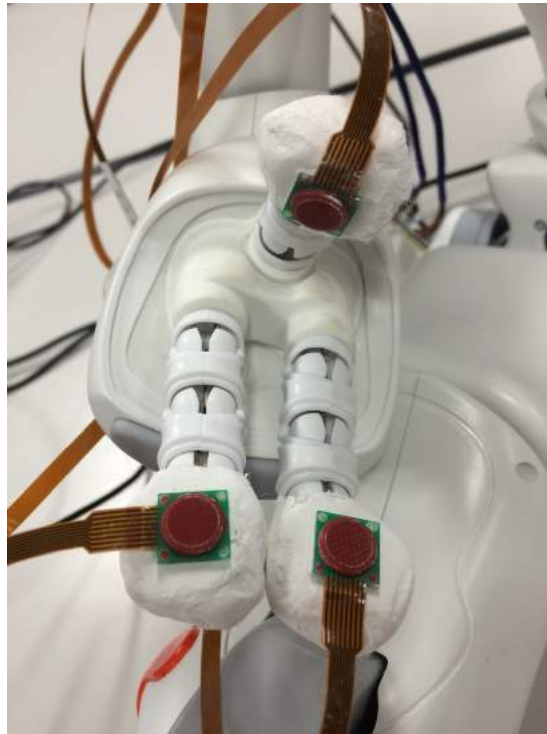
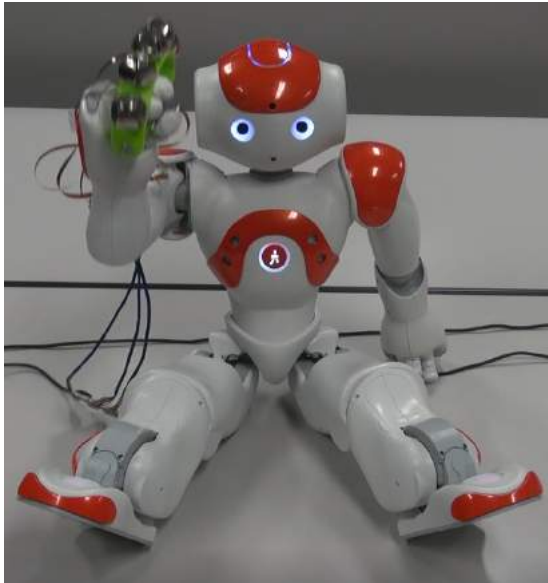


図 3.1: ロボットのハンドに設置された 3 つの圧力センサー

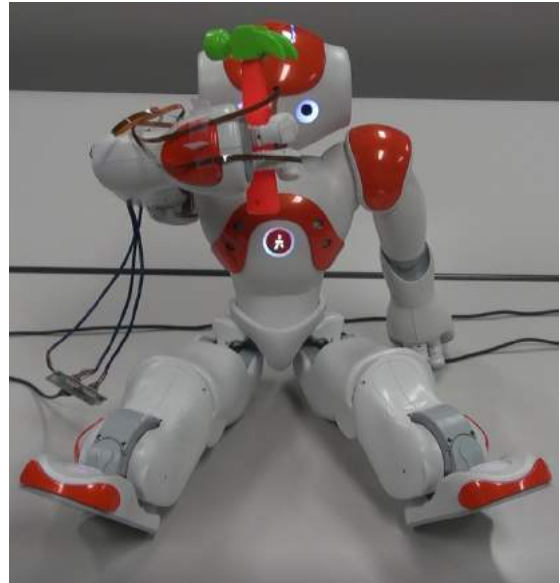
これらから得られる情報をそれぞれロボットから得られる視覚情報・聴覚情報・触覚情報として利用することができる。ロボットは9種類の物体(タンバリンや指人形, プラスティック製の工具などの幼児向けのおもちゃ)(図3.2)をカメラに近づけ撮影する行動, 物体を振りそれにより生じる音を録音する行動, 物体を握り, 圧力センサーの変化を計測する行動を各物体に対して行う(図3.3)。



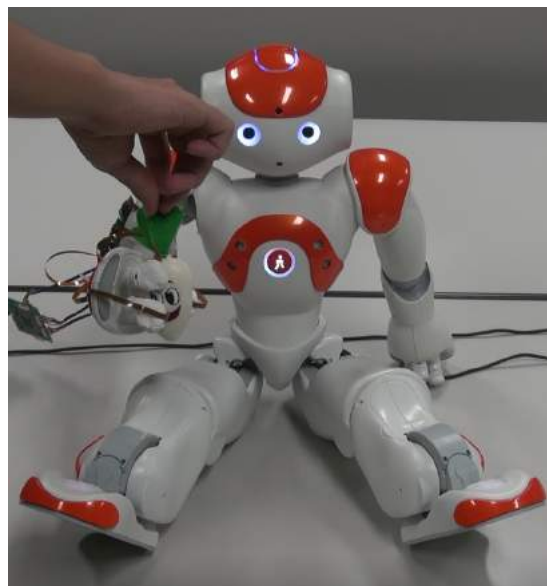
図 3.2: ロボットが認識する 9 種類の物体



(a) 振って音を聞く (聴覚)



(b) 近づけて見る (視覚)



(c) 物体を握る (触覚)

図 3.3: ロボットが物体に対して行う 3 種類の行動

これらの行動によって、得られた視覚情報・聴覚情報・触覚情報を，LDA と MLDA によって教師なしで分類する．MLDA で扱う特徴量には，Bag-of-Features

モデルを用い、次節以降に述べる方法で特徴量を計算する。

### 3.1.1 Bag-of-Features モデル

Bag-of-Features は各離散特徴量の出現頻度ヒストグラムをベクトルにしたもので、ベクトルの成分が各特徴量の出現回数となる [26]。この特徴量を作成するため、あらかじめ与えられたサンプル画像の特徴量ベクトルを k-means を用いてクラスタリングし、サンプル画像を代表する特徴量（コードブック）を計算する。分析対象の画像ベクトルをこのコードブックを用いてベクトル量子化することで、各コードブックにおける特徴量の出現頻度ヒストグラムを計算することができる。関連研究 [1] ではこの Bag-of-Features モデルを視覚情報だけでなく、聴覚情報と触覚情報にも拡張したモデルであるマルチモーダル Bag-of-Features モデルを提案している。本研究ではこれにならい、視覚情報、聴覚情報、触覚情報を以下のように処理する。

### 3.1.2 視覚情報の処理

ロボットのカメラで撮影した画像は、SIFT 特徴量 [27] を用いて 1 枚ごとに 128 次元の特徴ベクトル 300~400 個に変換する。変換された特徴ベクトルは、学習とは関係のない画像 10 枚から計算された 300 次元の代表ベクトルを用いてベクトル量子化する。つまり視覚情報は、300 次元のヒストグラムに変換される。このヒストグラムのインデックスが LDA と MLDA で扱う特徴量となる。

### 3.1.3 聴覚情報の処理

ロボットがマイクでとらえた音声は、MFCC (Melfrequency cepstrum) [28] を用いて 13 次元の特徴量ベクトルに変換する。この特徴ベクトルは白色雑音や複数の音楽、人間の音声を用いて計算した 100 の代表ベクトルによりベクトル量子化する。視覚特徴量は 100 次元のヒストグラムに変換され、このヒストグラムのインデックスが LDA と MLDA で扱う特徴量となる。

### 3.1.4 触覚情報の処理

今回の実験ではロボットのハンドの先端に 3 つの 3 軸圧力センサー [25] を設置した．3 軸圧力センサーはロボットのハンドに紙粘土で固定されており 3.1, 接地面の圧力方向だけでなく，その圧力方向と垂直な 2 つのせん断応力方向の力を計測することができる．この圧力センサーを設置したハンドで指人形を把持したときに得られた圧力データを (図 3.4) に示す．

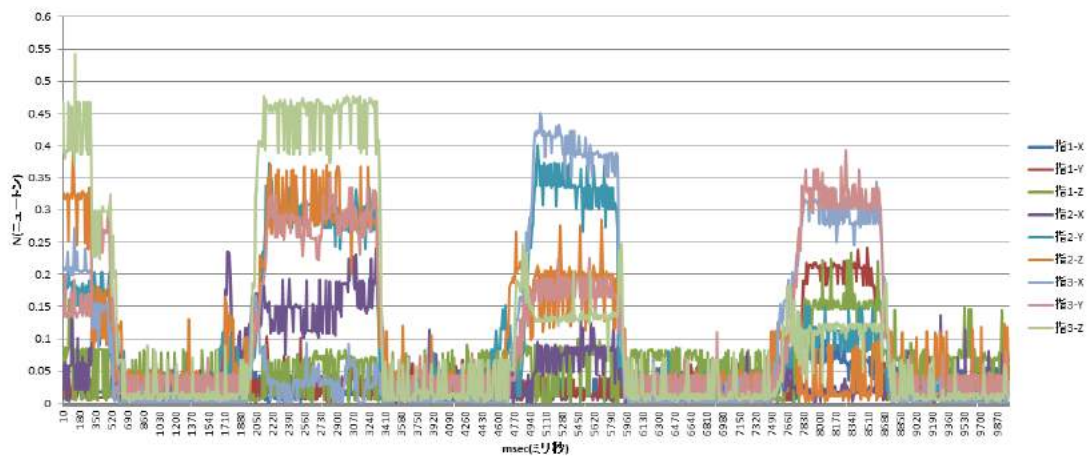


図 3.4: 指人形を把持したハンドから得られた圧力データ

我々は，この圧力センサーから得られた数値も他と同様にベクトル量子化するため，圧力センサーの変化量から得られる特徴量として valiation unigram 特徴量を提案した．valiation unigram では圧力センサーに働く力を 10msec ごとにサンプリングし，隣接するサンプリング間の力の差を変化量とする．この変化量を 0.01N 単位で量子化し，その頻度を計算したものを特徴量ベクトルとする．今回の実験では変化量の範囲を-0.59～0.60 の 120 段階とすることで 120 次元の特徴量ベクトルを計算した．valiation unigram により，ロボットが物体を 1 回握るごとに 3 つの 3 軸センサーから得られる 9 個の圧力データから 9 個の特徴量ベクトルを計算できる．この特徴量ベクトルは実験に使用した物体とは異なる物体から得られる圧力データから得られた 100 の代表ベクトルにより 100 次元のヒストグラムに変換される．このヒストグラムのインデックスを LDA と MLDA で扱う特徴量とする．

### 3.1.5 LDA と MLDA による教師なし分類

本研究では、LDA と MLDA を用いて物体の情報の教師なしの分類を行う。LDA を用いる場合は、ロボットから得られた視覚情報と聴覚情報、触覚情報の 300 次元と 100 次元のヒストグラムをそれぞれ、1～300 と 301～400, 401～500 のインデックスとみなし、一つのベクトルとして扱うことで分類を行う。MLDA は視覚情報と聴覚情報、触覚情報をそれぞれ対応する単語ベクトルとして扱い、分類を行う。クラス数は実験の種類に応じてあらかじめ与え、LDA と MLDA のパラメータ推定にはギブスサンプリングを用いた。クラス分類は LDA と MLDA の両方で物体情報のトピックパラメータである  $\theta_d$  を用いた。 $\theta_d$  はクラス数に等しい次元を持つベクトルであり、学習を行った  $d$  個の各物体情報に対して一つの  $\theta_d$  が計算される。関連研究 [1] では、 $d$  番目の物体情報のラベル  $R_d$  を  $\theta_d$  の最大値の要素番号、すなわち

$$R_d = \arg \max_i \theta_{d,i} \quad (3.1.1)$$

で与える。本研究では、これに併せて  $\theta_{d,i}$  を  $d = 1, \dots, M$  について正規化した  $\theta'_d$  の最大値の要素番号、すなわち

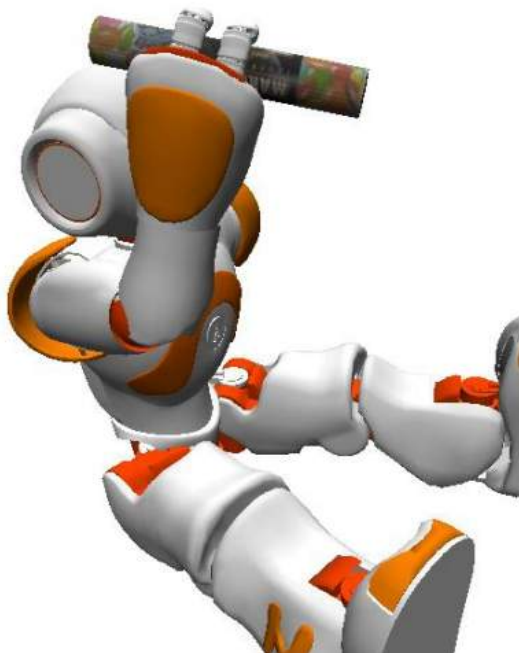
$$\theta'_{d,i} = \frac{\theta_{d,i}}{\sum_{d=1}^M \theta_{d,i}} \quad (3.1.2)$$

$$R_d = \arg \max_i \theta'_{d,i} \quad (3.1.3)$$

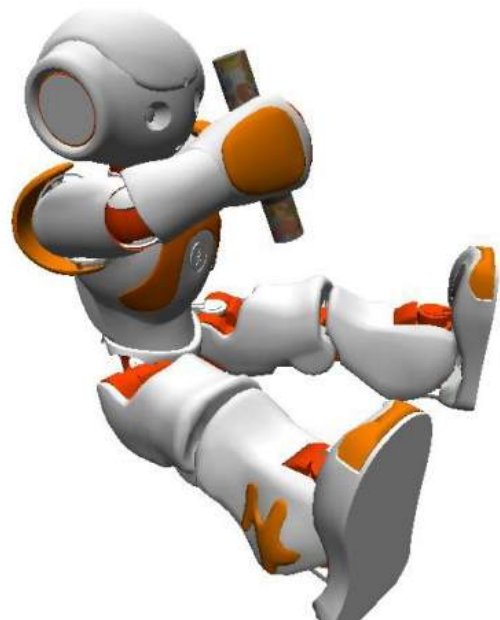
で与えた。

## 3.2 シミュレーション環境における概念獲得

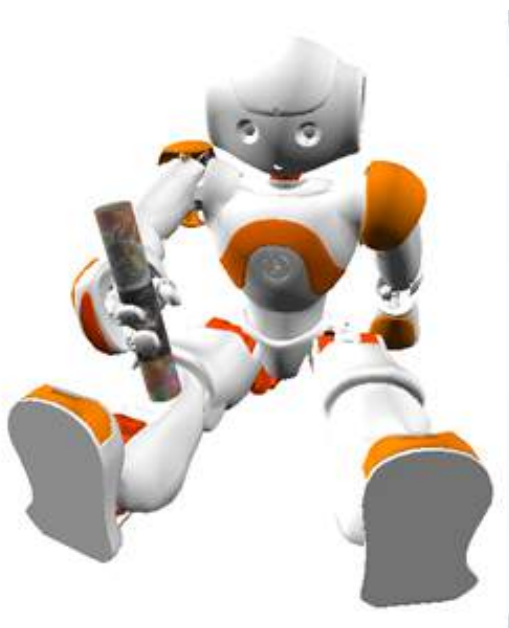
次に、実機と同様の状況を仮想環境で再現し、概念獲得のシミュレーションを行う。仮想環境には Unity [29] を使用した。Unity 上で Aldebaran Robotics が提供する 3DCG [30] モデルを用いて実機での実験を再現した。図 3.5 に実験の様子を示す。



(a) 振って音を聞く (聴覚)



(b) 近づけて見る (視覚)



(c) 物体を握る (触覚)

図 3.5: 仮想環境上でロボットが物体に対して行う 3 種類の行動

シミュレーションに使用する物体は AUTODESK REMAKE [31] を用い，実機での実験で使った物体の写真から 3DCG モデルを自動生成した（図 3.6）.



図 3.6: 仮想環境上でロボットが認識する 9 種類の物体

視覚のシミュレーションとしては，Unity 上でロボットのカメラの位置から捉えられる画像情報を用いた．また聴覚のシミュレーションとしては，物体の中の鈴のシミュレーションを行い，その衝突に合わせてあらかじめ録音した鈴の音を鳴らすことで，実機のように音声データを得ることができるようにした．触覚のシミュレーションとしては，ロボットのハンドのモデルにばねオブジェクト（ばねの働きをする Unity スクリプトを持つオブジェクト）を設置し，ロボットのハンドが物体を握ったときのばねオブジェクトの変位から圧力データを計算した．また，軟体は Unity の Interactive Cloth オブジェクトを用いることで再現した．

### 3.3 MLDA を用いた転移学習

本研究では，MLDA を転移学習に応用する手法を提案した．提案手法ではまず，仮想環境上で得られた物体情報を用いて MLDA のパラメータを推定する．次に，語彙情報を表す変数である  $\phi^a, \phi^v, \phi^h$  以外のパラメータを初期化する．最後に， $\phi^a, \phi^v, \phi^h$  を固定して，実環境上で得られた物体情報を用いて MLDA のパラメータを推定することで転移学習を行う．学習にはギブスサンプリングを用い， $\phi^a, \phi^v, \phi^h$  を固定し，学習は  $\phi^a, \phi^v, \phi^h$  のみサンプリングによる更新を省くことで行った．このアルゴリズムは  $\phi^a, \phi^v, \phi^h$  を固定した学習を行うことで仮想環境上で得られた物体情報に基いた実環境上の物体情報の認識を行うことができる可能性が考えられる．仮想環境上で得られる物体情報の数を変化させた場合に実環境上で得られた物体情報の識別率に変化があれば，このアルゴリズムは仮想環境上での学習に基づいた識別を行う転移学習アルゴリズムであると言える．4 章ではこのアルゴリズムの有効性を確認するため，仮想環境上で得られるデータセット数を増加させた場合に実環境上での物体情報の識別率の変化を調査する．

## 第 4 章 評価実験

本章では，3 章で述べた概念獲得手法の有効性を確認するために評価実験を行う．まず，関連研究 [1] で行われた概念獲得が我々が用意したロボットで再現できることを確認する．次に，仮想環境上で物体の概念獲得が正しく行われていることを確認するために，実ロボットを用いた物体の概念獲得と仮想環境上での物体の概念獲得の結果を比較し，その一致率の評価を行う．最後に，MLDA を用いた転移学習を行い，仮想環境上で学習に用いるデータを増加させた場合に，実環境上での物体の識別精度の変化を調査する．

### 4.1 実環境とシミュレーション環境における概念獲得の一致

#### 4.1.1 概念獲得の再現実験

本節では，まず関連研究 [1] の物体概念の獲得を再現する．我々は，3 章で述べたロボットを用い，9 種類の物体を用いた概念獲得を行った．9 種類の物体は“音が鳴る”，“かたい”，“やわらかい”といった性質があり，それに応じてあらかじめ 3 つのカテゴリに人手で分類されている．カテゴリとそれに対応する物体の性質を表 4.1 に示す．

表 4.1: 各物体の性質

| クラス   | 1                                                                                   |                                                                                     | 2                                                                                   |                                                                                     |                                                                                     |                                                                                       | 3                                                                                     |                                                                                       |                                                                                       |
|-------|-------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------|
| 物体番号  | 1                                                                                   | 2                                                                                   | 3                                                                                   | 4                                                                                   | 5                                                                                   | 6                                                                                     | 7                                                                                     | 8                                                                                     | 9                                                                                     |
| 物体    |  |  |  |  |  |  |  |  |  |
| 音がなる  | ○                                                                                   | ○                                                                                   | ×                                                                                   | ×                                                                                   | ×                                                                                   | ×                                                                                     | ×                                                                                     | ×                                                                                     | ×                                                                                     |
| かたい   | ○                                                                                   | ○                                                                                   | ×                                                                                   | ×                                                                                   | ×                                                                                   | ×                                                                                     | ○                                                                                     | ○                                                                                     | ○                                                                                     |
| やわらかい | ×                                                                                   | ×                                                                                   | ○                                                                                   | ○                                                                                   | ○                                                                                   | ○                                                                                     | ×                                                                                     | ×                                                                                     | ×                                                                                     |

この実験では，9 種類の各物体に対しそれぞれ 20 回ずつ計 180 回物体の視覚情

報と聴覚情報，触覚情報を取得し，それらの情報を入力として LDA を用い，3 章で述べた方法で物体情報にラベルを割り当てた．この LDA によるラベルの割り当て結果と，実際の物体番号の一致率を表 4.2～4.8 示す．

表 4.2: 視覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 0.46      | 0.90   | 0.61 | 40      |
| 2     | 0.79      | 0.53   | 0.63 | 80      |
| 3     | 0.35      | 0.28   | 0.31 | 60      |

表 4.3: 聴覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 0.95      | 0.97   | 0.96 | 40      |
| 2     | 0.59      | 0.60   | 0.59 | 80      |
| 3     | 0.44      | 0.42   | 0.43 | 60      |

表 4.4: 触覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 0.42      | 0.57   | 0.48 | 40      |
| 2     | 0.85      | 0.93   | 0.89 | 80      |
| 3     | 0.68      | 0.43   | 0.53 | 60      |

表 4.5: 視覚・聴覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 0.93      | 0.97   | 0.95 | 40      |
| 2     | 0.67      | 0.56   | 0.61 | 80      |
| 3     | 0.52      | 0.62   | 0.56 | 60      |

表 4.6: 聴覚・触覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 0.97      | 0.95   | 0.96 | 40      |
| 2     | 0.97      | 0.49   | 0.65 | 80      |
| 3     | 0.58      | 0.98   | 0.73 | 60      |

表 4.7: 触覚・視覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 0.69      | 0.45   | 0.55 | 40      |
| 2     | 0.78      | 1.00   | 0.88 | 80      |
| 3     | 0.81      | 0.70   | 0.75 | 60      |

表 4.8: 視覚・聴覚・触覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 0.97      | 0.97   | 0.97 | 40      |
| 2     | 0.99      | 0.94   | 0.96 | 80      |
| 3     | 0.92      | 0.98   | 0.95 | 60      |

表は視覚情報・聴覚情報・触覚情報の3種類すべてを用いて分類したもの、2種類の組み合わせ、1種類のを合わせた計7通りでの分類結果を示している。評価には正解ラベルに対する適合率、再現率およびその調和平均（F 値）を用いる。support は各正解ラベルの個数である。分類に用いる情報の種類を増やすことで精度が向上する傾向があり、3 つすべてを用いた場合全体の精度が 96% に達した (図 4.1)。また、3 種類のクラス分類とは別に LDA を用いて 9 種類の各物体ごとの分類も行った。この場合全体の分類の精度は 46% であった。

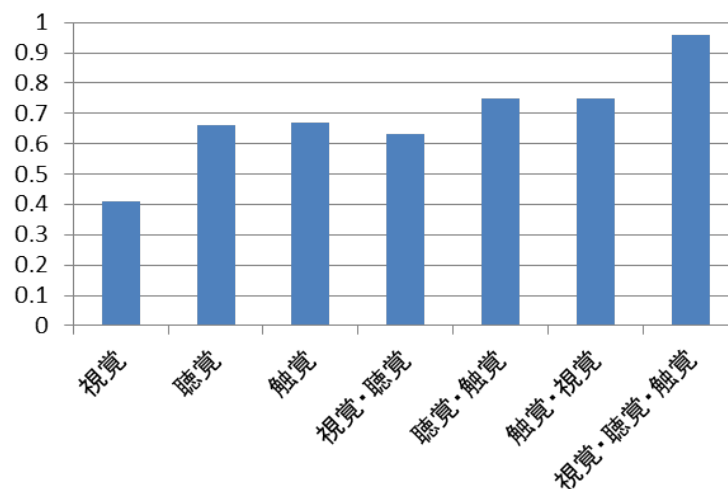


図 4.1: 実機での概念獲得において用いる情報の種類と分類精度の変化

MLDA を用い 9 種類の各物体ごとの分類を 100 回行った場合の精度の分布を図 4.2 に示す。トピックパラメータの情報から関連研究と同様の方法で分類を行うと 40% 程度の精度であるが、トピックパラメータの各要素を各物体に関して正規化を

行うことで分類の精度は 60% 程度に向上した.

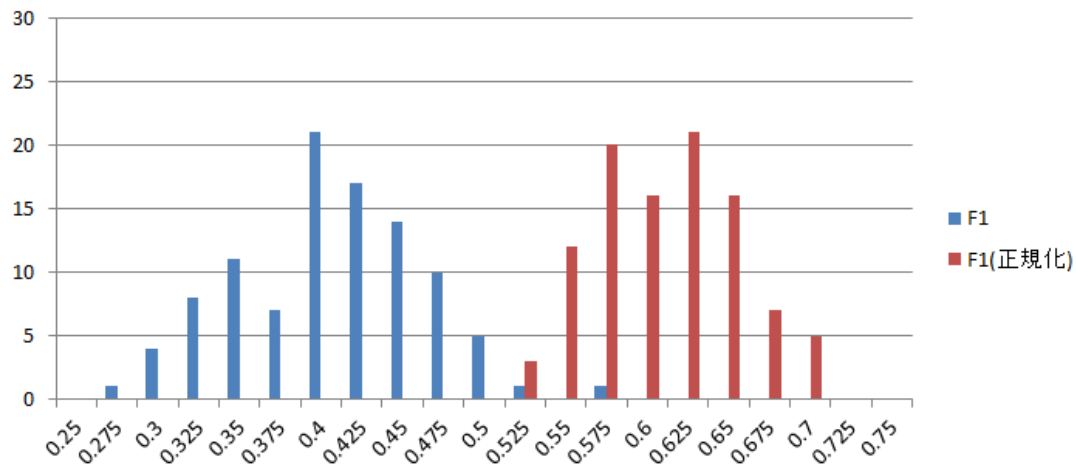


図 4.2: 実機で 9 種類の物体の識別を 100 回行った場合の精度分布

## 4.2 シミュレーション環境における概念獲得との比較

シミュレーション環境で, 実機と同様の試行 (9 種類の物体を それぞれ 20 試行) を行い, 獲得された特徴量に対する LDA を用いたクラスタリングを行った. この結果を表 4.9~4.15 に示す.

表 4.9: 視覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 0.79      | 0.75   | 0.77 | 40      |
| 2     | 0.84      | 0.60   | 0.70 | 80      |
| 3     | 0.59      | 0.83   | 0.69 | 60      |

表 4.10: 聴覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 1.00      | 0.95   | 0.97 | 40      |
| 2     | 0.56      | 0.91   | 0.70 | 80      |
| 3     | 0.33      | 0.07   | 0.11 | 60      |

表 4.11: 触覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 0.00      | 0.00   | 0.00 | 40      |
| 2     | 1.00      | 0.88   | 0.93 | 80      |
| 3     | 0.60      | 1.00   | 0.75 | 60      |

表 4.12: 視覚・聴覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 0.95      | 0.97   | 0.96 | 40      |
| 2     | 0.58      | 0.57   | 0.58 | 80      |
| 3     | 0.45      | 0.45   | 0.45 | 60      |

表 4.13: 聴覚・触覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 0.98      | 1.00   | 0.99 | 40      |
| 2     | 1.00      | 0.99   | 0.99 | 80      |
| 3     | 1.00      | 1.00   | 1.00 | 60      |

表 4.14: 触覚・視覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 0.95      | 0.93   | 0.94 | 40      |
| 2     | 1.00      | 1.00   | 1.00 | 80      |
| 3     | 0.95      | 0.97   | 0.96 | 60      |

表 4.15: 視覚・聴覚・触覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 1.00      | 0.97   | 0.99 | 40      |
| 2     | 0.98      | 1.00   | 0.99 | 80      |
| 3     | 1.00      | 0.98   | 0.99 | 60      |

実機と同様，分類に用いる情報の種類を増やすごとに精度が向上する傾向が見られ，最終的な 3 クラスへの分精度は 99% であった (図 4.3). また，実機の実験と同様に各物体ごとの分類も行った．こちらの実験の分類の精度は実機と同程度の 43% であった．

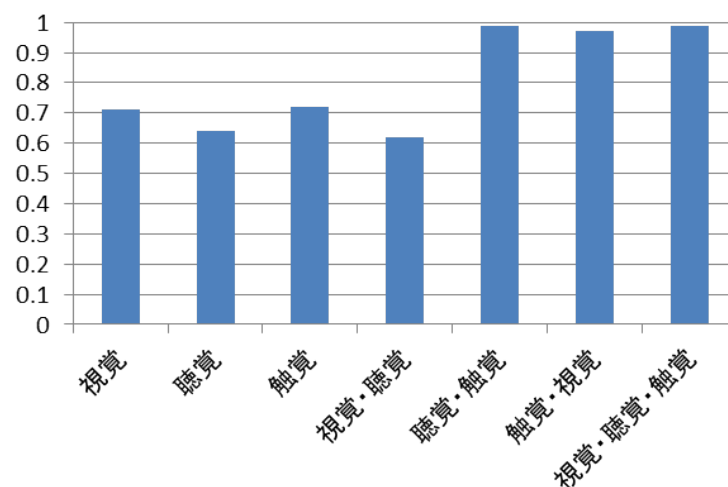


図 4.3: 仮想環境上での概念獲得において用いる情報の種類と分類精度の変化

実機の場合と同様に、仮想環境で MLDA を用いて 9 種類の各物体ごとの分類を 100 回行った結果を図 4.4 に示す。分類の精度はトピックパラメータの正規化なしの場合に 50% 程度で、正規化を行った場合は 60% 程度であった。

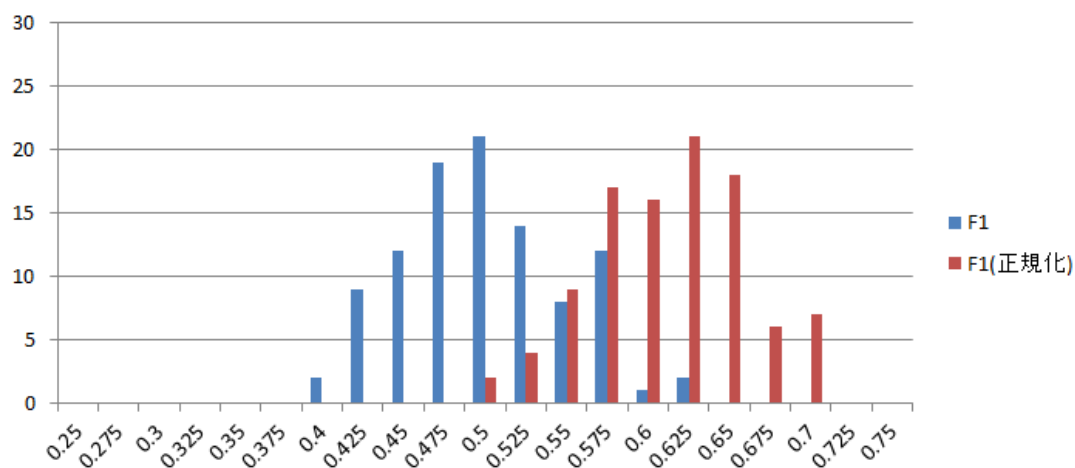


図 4.4: 仮想環境上で 9 種類の物体の識別を 100 回行った場合の精度分布

これらの結果から 3 クラス分類に関してはシミュレーションにおいても実機で行う場合と同程度の精度で物体に対する概念獲得が可能であることが示された。しか

し，この学習結果を実機で利用する概念として転用する場合，この結果がどの程度実機での結果と一致しているかが重要となる．そこで，実機とシミュレーションで得られた 180 試行のラベルの割り当ての一致率を確認した．これを表 4.16～4.22 に示す．

表 4.16: 視覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 0.40      | 0.47   | 0.43 | 49      |
| 2     | 0.68      | 0.33   | 0.45 | 78      |
| 3     | 0.29      | 0.47   | 0.36 | 53      |

表 4.17: 聴覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 1.00      | 0.93   | 0.96 | 41      |
| 2     | 0.58      | 0.93   | 0.72 | 82      |
| 3     | 0.33      | 0.07   | 0.12 | 57      |

表 4.18: 触覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 0.10      | 0.03   | 0.04 | 38      |
| 2     | 0.93      | 0.75   | 0.83 | 87      |
| 3     | 0.55      | 1.00   | 0.71 | 55      |

表 4.19: 視覚・聴覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 0.98      | 0.95   | 0.96 | 42      |
| 2     | 0.53      | 0.45   | 0.49 | 71      |
| 3     | 0.52      | 0.61   | 0.56 | 67      |

表 4.20: 聴覚・触覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 0.93      | 0.97   | 0.95 | 39      |
| 2     | 0.48      | 0.95   | 0.64 | 101     |
| 3     | 0.48      | 0.95   | 0.64 | 40      |

表 4.21: 触覚・視覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 0.38      | 0.58   | 0.46 | 26      |
| 2     | 1.00      | 0.78   | 0.88 | 102     |
| 3     | 0.66      | 0.77   | 0.71 | 52      |

表 4.22: 視覚・聴覚・触覚

| class | presicion | recall | f1   | support |
|-------|-----------|--------|------|---------|
| 1     | 1.00      | 0.97   | 0.99 | 40      |
| 2     | 0.91      | 0.99   | 0.95 | 76      |
| 3     | 0.98      | 0.91   | 0.94 | 64      |

ここでは実機で獲得されたラベルを正解とし，シミュレーションで獲得されたラベルの精度を調べたものを，これまでの実験同様適合率，再現率とその調和平均で表す．これまでの結果と同様に用いる情報の種類が増えるごとに精度が向上する傾向が見られた (図 4.5)．3 クラス分類における最終的な一致率は 96% に達し，実機とシミュレーションでの概念獲得結果がほぼ同じ結果となっていることがわかる．また，実機と仮想環境上での概念獲得を各物体について比較した．分類の精度は 43% であった．

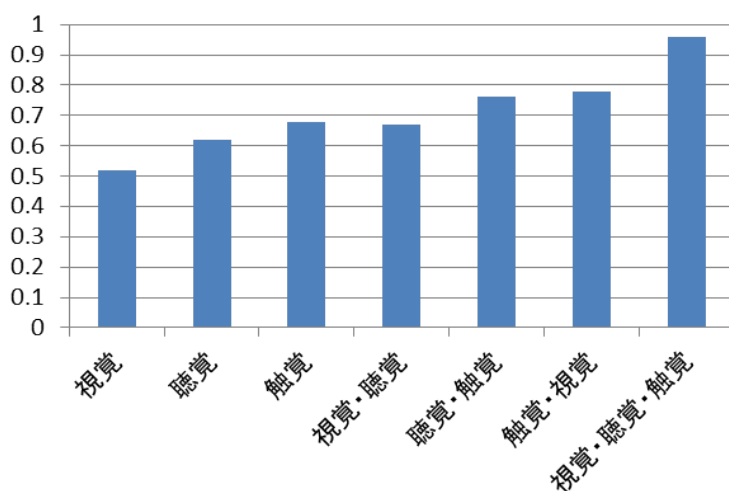


図 4.5: 実機と仮想環境上での実験結果の比較

実機と仮想環境上での物体概念獲得を 9 種類の各物体ごとに 100 回比較した結果を図 4.6 に示す．トピックパラメータの情報から関連研究と同様の方法で分類を行うと 40% 程度の精度であるが，トピックパラメータを正規化することで分類の精

度は 50% 程度に向上した.

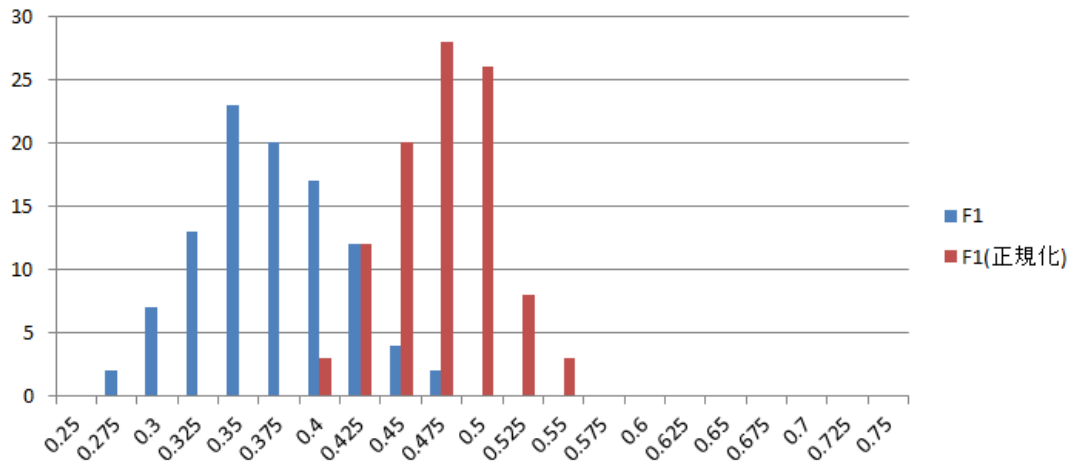


図 4.6: 実機と仮想環境上で 9 種類の物体を 100 回識別した場合の精度分布

#### 4.2.1 考察

実機と仮想環境で獲得されたラベルは、視覚、聴覚、触覚の 3 種の情報をすべて用いた場合についてはほぼ完全に一致した。このことから視覚・聴覚・触覚を用いた仮想環境上での概念獲得の結果を実機に転写し、高い精度で利用できる可能性があることがわかる。また、各物体ごとの分類の比較に関しては LDA を用いた場合は実機とシミュレーションとその比較についてほぼ同水準の精度であった。このことから低精度ではあるがある程度の一致があることがわかる。ただ、各物体の認識精度そのものに課題があり、特徴量の作成方法に改善の余地がある。一方、MLDA を用いた場合は、LDA の場合に比べ 20% 程度精度が向上することが分かった。また、トピックパラメータを正規化することで精度がさらに 10% 程向上することも確認された。

##### トピックパラメータの正規化による分類制度の向上

トピックパラメータはロボットが認識した物体情報に一つずつ計算されるベクトルであり、次元数はあらかじめ設定した物体のカテゴリ数である。今回の実験

では 9 種類の物体にそれぞれ 20 回づつ物体情報を取得することで計 180 個の物体情報を取得したので、9 次元のトピックパラメータが 180 個計算される．実機と仮想環境上での概念獲得を各物体について比較した場合のトピックパラメータ  $\theta_d = (\theta_{d,1} \cdots \theta_{d,9})(d = 1, \cdots, 180)$  の  $\theta_{d,6}$  と  $\theta_{d,7}$  の結果を 4.7 に示す．

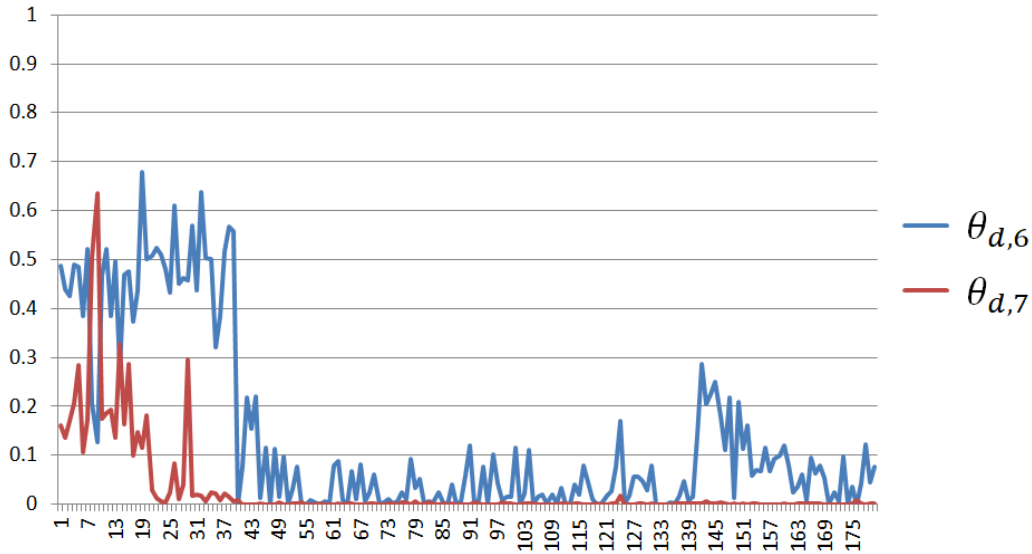


図 4.7: 各物体情報におけるトピックパラメータの値

グラフは横軸が物体情報の番号  $d$  を，縦軸が  $\theta_{d,6}$  (青線) と  $\theta_{d,7}$  (赤線) の値を表している．物体の情報番号は 1～20 番，21～40 番がそれぞれ物体番号 1 と物体番号 2 の物体を示している． $\theta_{d,7}$  では要素番号  $d=1\sim 20$  で値が上昇している．これは 7 番のトピックが物体番号 1 番の青い鈴を識別したものと考えられる．また， $\theta_{d,6}$  はでは要素番号  $d=1\sim 40$  で値が上昇している．これは物体番号 1 番と 2 番の青い鈴と緑の鈴を音の鳴る物体のトピックとして識別したものと考えることができる．このとき， $\theta_{d,6}$  は  $d=1\sim 40$  でほぼ  $\theta_{d,7}$  を上回る値を示しており， $\theta_{d,7}$  のデータは埋もれてしまっている．関連研究にならない，この中から最大値を選ぶことにより，物体のクラスラベルを計算すると物体 1 と物体 2 は同じ物体であると分類されてしまう可能性がある．トピックパラメータの正規化による分類精度の向上の理由として，このような埋もれた要素に注目できるようになることが考えられる．

### 4.3 物体の概念獲得における転移学習

4.1 節では実ロボットを用いた概念獲得と仮想環境上での概念獲得が 3 クラスの場合ではほぼ完全に一致し、9 クラス分類の場合でもある程度一致することが示された。これらの結果から仮想環境上での物体情報の学習結果を実ロボットに転写し、高い精度で利用できる可能性がある。本節ではこれらを踏まえ、仮想環境上で物体概念の獲得を行い、それを 3.3 節で述べた方法を用いて実機に転写し、実データの分類を行い、その精度を調査する。

#### 4.3.1 実験設定

転移学習実験では仮想環境上で MLDA を学習し、3.3 章で述べた方法で実ロボットから得られた物体情報を分類する。実環境上での物体情報は 4.1 章と同様に 9 種類の各物体に対し、実ロボットを用いてそれぞれ 20 回ずつ計 180 回物体情報を取得した。仮想環境上での物体情報は 9 種類の各物体に対し、それぞれ 1 回～200 回ずつ計 9 回～1800 回物体情報を取得した。物体 1 つから 1 回ずつ物体情報を取得したデータを 1 セットとすると、仮想環境上で MLDA の学習に用いるデータセット数を 1, 5, 10, 15, 20, 50, 100, 150, 200 と増加させたときの実環境上での分類精度を調査した。分類精度の調査は 3 クラス分類と 9 種類の各物体の分類について行い、3 クラス分類では 1～20, 9 クラス分類では 1～200 の間でデータセット数を変化させ、各データセット数に対して 10 回ずつ精度の計算し、その平均値を算出した。

#### 4.3.2 結果

3 クラス分類の結果を図 4.8 に、9 種類の各物体の分類の結果を図 4.9 に示す。

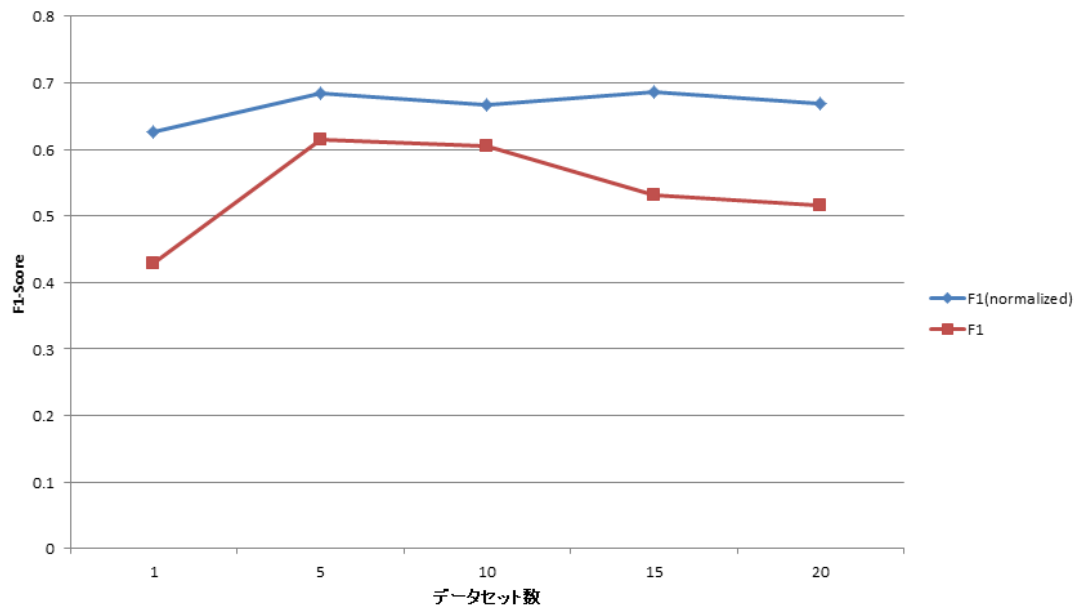


図 4.8: 9 種類 3 カテゴリの転移学習の精度の変化

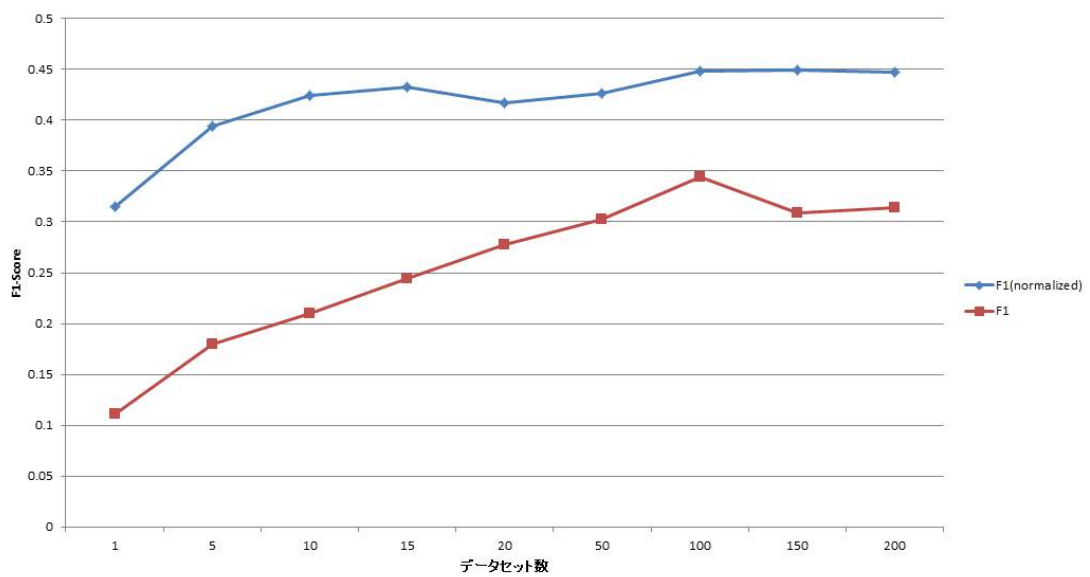


図 4.9: 9 種類 9 カテゴリの転移学習の精度の変化

3 クラス分類ではデータセット数を 5 にした場合に分類精度がトピックパラメー

タの正規化を行う場合 70%, 行わない場合で 60% とそれぞれ精度が最大になった。それ以降はデータセット数が増加するにつれ、精度は減少する傾向が見られる。データセット数が 10 以上の識別では仮想環境上での学習に MLDA が過学習を起こした可能性が考えられる。一方、9 クラス分類ではデータセット数を 1~200 に変化させることで全体を通して精度が上昇する傾向が見られた。分類精度はトピックパラメータの正規化を行った場合は最大 45%, 正規化を行わなかった場合で最大 35% であった。また、精度の上昇率はデータセット数を増加させる毎に小さくなり、トピックパラメータの正規化を行う場合ではデータセット数が 100 以上の場合は変化が見られなかった。いずれもデータセット数が増加するにしたがって実ロボット上での識別精度が上昇することがわかる。また、トピックパラメータの正規化を行うことでデータセット数や分類カテゴリ数にかかわらず分類精度が向上することも確認された。

## 第 5 章 結言

### 5.1 本論文のまとめ

本研究では，物体の概念獲得を仮想環境上で行い，実ロボットに転写する転移学習を行った．そのために，MLDA を用いた転移学習法を提案した．また，転移学習が有効であることを示すために，関連研究の実験を実ロボットを用いて再現し，仮想環境上でも同様の実験が成立することを示した．そして実ロボットでの概念獲得結果と仮想環境上での概念獲得結果が高い精度で一致することを示した．これに伴い，関連研究で示されたロボットによる物体概念の獲得を関連研究とは異なる実験環境で再現することに成功した．これに加えて，MLDA を用いた教師なしの物体概念獲得においてトピックパラメータの正規化が分類精度の向上に有効であることを示した．転移学習の精度は低く，改善の余地があるが，仮想環境上でデータセット数を増加させることで精度が向上することを示した．

### 5.2 今後の課題

まず，本研究で取り組んだ実験において 9 種類の物体の分類の精度には改善の余地がある．次に，転移学習の精度も低く，仮想環境で学習に用いるデータセット数を増加させるたびに精度が高くなる傾向が見られたが，データセット数が 100 個以上になるとデータセット数を増加させても，トピックパラメータを正規化した場合の精度が向上しなくなった．これに加えて，物体の概念獲得を生活支援システムなどに実際に応用するためには，物体の種類とカテゴリ数を増加させた場合について学習を行う必要がある．また，実生活で実用的な仕事をロボットが担うことを想定した場合，物体の概念を用いることで，ロボットが蛇口やコップ，スイッチといった物体の存在を認識するだけでなく，それらの物体に対し，コップを運ぶ，水を流す，スイッチを入れる，といった物体に対して行うロボットの動作の概念を獲得する必要があると考えられる．

## 謝辞

本学情報科学研究科の中村哲教授には、主指導教官として重要なご助言やご教示をいただきました。また、研究を進めるにあたり、素晴らしい研究環境を与えていただきました。心より感謝いたします。

本学情報科学研究科の松本裕治教授には副指導教官として貴重なご助言を賜りました。心より感謝いたします。

本学情報科学研究科の吉野幸一郎助教には副指導教官として研究の進め方や日常生活について数多くの貴重なご助言やご教示を賜りました。また、論文執筆や学会参加時に数々のご支援をいただきました。心より感謝を申し上げます。

本学情報科学研究科の須藤克仁准教授には、副指導教官として本研究の実験の考察に関して様々なご助言を賜りました。心より感謝申し上げます。

本学情報科学研究科の田中宏季助教には、副指導教官として研究内容に関して、ご助言を賜りました。心より感謝申し上げます。

本学知能コミュニケーション研究室の本研究室の皆様には、研究や日常生活に関する数多くのご助言をいただきました。心より感謝を申し上げます。

最後に、私をここまで育ててくれた家族に感謝を申し上げます。

## 参考文献

- [1] Tomoaki Nakamura, Takaya Araki, Takayuki Nagai and Naoto Iwahashi: “Grounding of Word Meanings in LDA-Based Multimodal Concepts”, *Journal of Intelligent and Robotic Systems*, pp.1-18, Jul.2011
- [2] Tadahiro Taniguchi, Takayuki Nagai, Tomoaki Nakamura, Naoto Iwahashi, Tetsuya Ogata, and Hideki Asoh, "Symbol Emergence in Robotics: A Survey", *Advanced Robotics*, Vol. 30, Issue 11-12, pp. 706-728, Jun. 2016
- [3] Kensho Miyoshi, Ryo Konomura, Koichi Hori, gEntertainment Multi-Rotor Robot that realises Direct and Multimodal Interactionh, *BCS The 28th British HCI Conference (BHCI)*, September 9-12, 2014.
- [4] Masahiro Araki, Akiko Kouzawa and Kenji Tachibana:“Proposal of a multimodal interaction description language for various interactive agents.”, *IE-ICE transactions on information and systems*, E88-D(11), pp.2469-2476.
- [5] Yasuyuki Sumi, Masaharu Yano and Toyoaki Nishida: Analysis environment of conversational structure with nonverbal multimodal data, *12th International Conference on Multimodal Interfaces and 7th Workshop on Machine Learning for Multimodal Interaction (ICMI-MLMI 2010)*, Beijing, China, November 2010.
- [6] X. Zuo, N. Iwahashi, R. Taguchi, S. Matsuda, K. Sugiura, K. Funakoshi, M. Nakano and Natsuki Oka. ?gRobot-Directed Speech Detection Using Multimodal Semantic Confidence Based on Speech, Image, and Motion,?h in *Proc. IEEE Int. Conf. Acoustics, Speech, and Signal Processing*, pp. 2458-2461, 2010.
- [7] T. Tagniguchi, K. Hamahata, and N. Iwahashi, “Unsupervised segmentation of human motion data using sticky hdp-hmm and mdl-based chunking method for imitation learning”, *Advanced Robotics*, vol. 25, no. 17, pp. 2143-2172, 2011.
- [8] Takaya Araki, Tomoaki Nakamura, Takayuki Nagai, Kotaro Funakoshi, Mikio Nakano, and Naoto Iwahashi : Online Object Categorization Using

- Multimodal Information Autonomously Acquired by a Mobile Robot Advanced Robotics, Vol. 26, Issue 17, pp. 1995-2020, 2012.
- [9] Takaya Araki, Tomoaki Nakamura, Takayuki Nagai, Shogo Nagasaka, Tadahiro Taniguchi, and Naoto Iwahashi : Online Learning of Concepts and Words Using Multimodal LDA and Hierarchical Pitman-Yor Language Model IEEE/RSJ International Conference on Intelligent Robots and Systems, pp. 1623-1630, Portugal, Oct.. 2012.
  - [10] MIT Technology Review “To Get Truly Smart, AI Might Need to Play More Video Games”, <https://www.technologyreview.com/s/601009/to-get-truly-smart-ai-might-need-to-play-more-video-games/>
  - [11] Andrei A Rusu, Matej Vecerik, Thomas Rothorl, Nicolas Heess, Razvan “ Pascanu, and Raia Hadsell. Sim-to-Real Robot Learning from Pixels with Progressive Nets. In CoRL, 2017.
  - [12] Blei, D. M.?CNg, A.Y. and Jordan, M.I.: “Latent Dirichlet Allocation”, Journal of Machine Learning Research 3, pp.993-1022, 2003.
  - [13] 奥村学・佐藤一誠 『自然言語処理シリーズ 8 トピックモデルによる潜在的意味解析』 コロナ社 (2015).
  - [14] R. Fergus, P. Perona, and A. Zisserman, “Object class recognition by unsupervised scale-invariant learning,” Proc. CVPR2003, vol.2, pp.264–271, June 2003.
  - [15] R. Fergus, P. Perona, and A. Zisserman, “Using multiple segmentations to discover objects and their extent in image collections,” Proc. CVPR2006, vol.2, pp.1605–1614, June 2006.
  - [16] L. Fei-Fei and P. Perona, “A Bayesian hierarchical model for learning natural scene categories,” Proc. CVPR2005, vol.2, pp.524–531, June 2005.
  - [17] F.G. Ashby and W.T. Maddox, “Human category learning,” Annu. Rev. Psychology, vol.56, pp.149– 178, Feb. 2005.
  - [18] D.J. Freedman and J.A. Assad, “Experiencedependent representation of visual categories in pariental cortex,” Nature, vol.443, no.7, pp.85–88, 2006.
  - [19] 安藤 義記, 中村 友昭, 荒木 孝弥, 長井 隆行, ” 階層マルチモーダルカテゴ

- リゼーションによる多様な概念と語意の学習,” 人工知能学会全国大会論文集 2013.
- [20] 中村 友昭, 長井 隆行, 岩橋 直人, 複数のマルチモーダル LDA を用いた抽象的概念の形成 The 24th Annual Conference of the Japanes Society for Artificial Intelligence, 2010
- [21] Griffiths, T.L. and Steyvers, M. (2004). Finding scientific topics. *Proceedings of the National Academy of Sciences*, 101(Supp.1), 5229-5235.
- [22] Erdogan, G., Yildirim, I., and Jacobs, R. A. “From sensory signals to modality-independent conceptual representations: A probabilistic language of thought approach.”, *PLoS Computational Biology*, 11, e1004610., 2015.
- [23] Aoki, T., Nishihara, J., Nakamura, T., and Nagai, T. (2016). Online Joint Learning of Object Concepts and Language Model using Multimodal Hierarchical Dirichlet Process. *IEEE/RSJ International Conference on Intelligent Robots and Systems*, 2636–2642.
- [24] “Aldebaran Robotics NAO”, <https://www.ald.softbankrobotics.com/en/cool-robots/nao>
- [25] “Touchence ShokacChip p/nFTSSI OD10 C10”, <http://www.touchence.jp/chip/index.html>
- [26] Csurka, G., Dance, C.R., Fan, L., Willamowski, J. and Bray, C: “Visual categorization with bags of keypoints”, *ECCV International Workshop on Statistical Learning in Computer Vision* (2004).
- [27] D. G. Lowe.: “Distinctive Image Features from Scale-Invariant Keypoints”, *International Journal of Computer Vision*, 60(2):91-110, 2004.
- [28] P. Mermelstein (1976), "Distance measures for speech recognition, psychological and instrumental," in *Pattern Recognition and Artificial Intelligence*, C. H. Chen, Ed., pp. 374-388. Academic, New York.
- [29] “Unity4.0”, <https://unity3d.com/jp/unity/whats-new/unity-4.0/>
- [30] <https://community.ald.softbankrobotics.com/en/resources/software/language/en-gb>
- [31] “AUTODESK REMAKE”, <https://remake.autodesk.com/about>

## 発表リスト

**国内学会** 笹野 仁, 吉野 幸一郎, 中村 哲. マルチモーダル情報を用いたロボットによる物体概念獲得のシミュレーション, 日本ロボット学会 2016

**大会講演** 笹野 仁, 吉野 幸一郎, 中村 哲. マルチモーダル情報を用いたロボットによる物体概念獲得のシミュレーション 情報処理学会関西支部大会講演論文集, 3p, 2016