

相互排他性原理に基づくマルチモーダル共同注意

中 野 吏^{*1*2} 吉 川 雄一郎^{*2} 浅 田 稔^{*1*2} 石 黒 浩^{*1*2}

Multimodal joint attention based on mutual exclusivity principle

Tsukasa Nakano^{*1*2}, Yuichiro Yoshikawa^{*2}, Minoru Asada^{*1*2} and Hiroshi Ishiguro^{*1*2}

In this paper, we propose a method for simultaneous learning of multi-modules for joint attention: gaze-driven attention and word-driven attention. Inspired from child language acquisition, mutually exclusivity bias is utilized for mutual facilitative learning both in an intra- and inter-module manner by extending a modified Hebbian learning rule. Experiments on a human-robot interaction and on the computer simulations, we analyzed that the proposed method enabled mutually facilitative learning of a mapping for gaze-following and a label-to-object mapping by which the learner performs multimodal joint attention with its caregiver. Finally, through a computer simulation resembling mother-infant interaction, we argued a possibility of the proposed learning mechanism as a constructivist model for infant's cognitive development.

Key Words: Joint attention, Mutual exclusivity, Simultaneous learning of multi-functions

1. は じ め に

他者とコミュニケーションを行う際には、他者と共同注意をすることが重要である。共同注意とは他者と同一の対象物に注意を向ける行動であり、複数のモダリティの情報（視線、言葉、指差し等）を手がかりに達成することができる。発達心理学の研究では、人の幼児は共同注意に必要な視線追従能力 [1] や語彙 [2] を 12~18ヶ月頃の同時期に獲得しているとされている。これらの能力の発達過程は相互に促進するものであることが伺える [3] [4] が、このような複数の機能の相互促進を支える学習メカニズムについては明らかでない。

これに対し、発達するロボットを構築することを通じて幼児の発達メカニズムの新たな理解を目指す、構成論的アプローチが注目を集めている [5]。構成論的アプローチの従来研究として、共同注意の機能である視線追従 [6] [7] や語彙 [8]~[10] のマッピングをロボットやエージェントに学習させる研究がある。しかし、これらの研究ではいずれかのモダリティの学習に焦点が当てられており、マルチモーダルな共同注意の機能の相互促進的学習は考慮されていなかった。一方、Yu et al. [11] は複数の物体が存在する状況で、教示者の視線情報を利用することで効率的に語彙マッピングが学習できることを示したが、視線を追従す

る機能は設計者によりあらかじめ与えられていた。マルチモーダル共同注意の複数機能を同時学習させることを考えた場合、はじめから、それぞれの機能が他方の学習を促進することは期待できず、どのようにそれぞれの機能の学習を進め、統合させれば相互促進が可能であるかを考える必要がある。

そこで本研究ではマルチモーダルな共同注意による語彙および視線追従能力のマッピングの同時学習を実現する相互促進的学習が可能な学習メカニズムを提案し、幼児の共同注意発達過程の構成論的理解を図る。幼児の初期の語彙学習過程では、物体とラベルが相互排他的に対応づけられるバイアスがあることが知られ、相互排他性バイアスと呼ばれている [12]。ここで、二つのものが相互排他的に対応付けられている、とは双方が他方に対して一意的に対応すると捉えられていることを指し、幼児が初めて聞く物体の名称をすでに他の名称が与えられた物体へは結びつけないことはこのようなバイアスによるものであると考えられている。これは乳児の初期のマッピング学習において、モダリティを問わず、既知の対応関係に照らし合わせてどの経験が学習に用いられるべきかが取捨選択される結果、効率的な 1 対 1 マッピングの構築に利用されている可能性を示唆している。そこで本研究では、このような相互排他性バイアスに基づく学習メカニズムを語彙と視線追従の両方のマッピング学習に適用することで、各モダリティの 1 対 1 マッピングの学習を効率化すると同時に、モダリティを跨いだ相互促進の実現を図る。具体的には、自身と結合先のノードの排他性を考慮してノード間の対応を学習する交差投錨型ヘッブ則 [13] を語彙と視線追従の両方のマッピング学習に応用する。これによりノード間の結

原稿受付

^{*1}大阪大学大学院工学研究科

^{*2}JST ERATO 浅田共創知能プロジェクト

^{*1}Graduate School of Engineering, Osaka University

^{*2}Asada Synergistic Intelligence Project, ERATO, JST

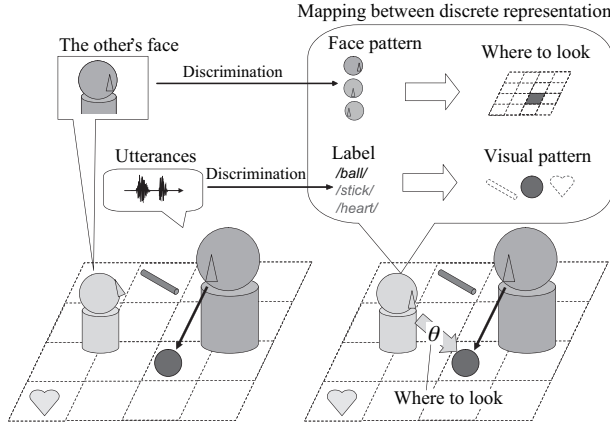


Fig. 1 Setting for learning multimodal joint attention

合強度として個々のマッピングの相互排他性が明示的に表現されることとなり、各モダリティの出力を相互排他性に応じて統合することができる結果として、相互促進学習を実現する。

本稿では2章でマルチモーダルな共同注意における視線追従および語彙の同時学習のための相互排他性原理の実装方法について説明する。3章では比較的簡単な環境で実施した実ロボットと人の共同注意学習実験において提案手法の有効性を検証した後、より複雑な環境を設定した計算機シミュレーションを行い、提案手法によりどのように各モジュールが相互促進可能であるかを検証する。最後に親子インタラクションを模した状況での計算機シミュレーション結果を示し、幼児の発達過程との相同性および今後の課題について議論する。

2. 相互排他性原理に基づく共同注意

本研究で想定する学習者と養育者のインタラクションの状況を Fig.1 に示す。養育者は周囲に配置された複数の物体の中から一つを選択し、これを見ながら、その名称（以下ラベル）を発話する。学習者のタスクはこのような養育者の振る舞いを手がかりに、試行錯誤を通じて養育者と同じものを見ること、すなわち、共同注意を獲得することである。Fig.2 に本研究で用いる学習者の共同注意学習システムを示す。共同注意の実現のため、学習者は以下で詳述する視線による注意モジュール (Fig.2 上部) と言葉による注意モジュール (Fig.2 下部) の出力を統合して注視方向を決定し、その結果である物体への注視経験をもとに、これらのモジュールのパラメータを更新する。この試行錯誤を通じて、学習者は養育者の視線とその視線の先の位置のマッピングと、ラベルとラベルが指す物体のマッピングを学習する。ただし、共同注意が成功したか否かの直接的な情報（教示や強化信号）は学習者に与えられていないことに注意されたい。

以下では、環境は N 個に分割されているとし、一つのスポットに置かれる物体は最大一つであるとする。また環境に配置される物体の候補は M 個あり、一度に置かれる物体の数は M_0 であるとする。どの物体がどこに配置されるかは、 T ステップごとにランダムに決定しなおされる。ここで1ステップとは、教示者が一つの物体を注視し、そのラベルを発話してから、ロボットがこれらの情報をもとに注視対象を決定し、視線追従と

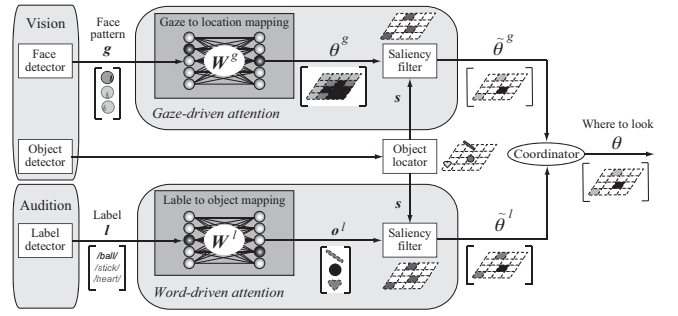


Fig. 2 Learning mechanism of joint attention with two modules: gaze-driven attentional one (top) and word-driven attentional one (bottom)

語彙のマッピングを更新するまでの間を指す。

2.1 視線による注意モジュール

視線による注意モジュールは、養育者が物体を注視している様子、すなわち顔についての観測情報を表す G 次元ベクトル g を入力として、共同注意のために注視されるべき位置を

$$\theta^g = W^g g \quad (1)$$

のように出力する。ここで、 θ^g は、養育者の視線を手がかりとしたときに環境中の各スポットが注視されるべき程度を要素とする N 次元ベクトルである。また W^g はマッピングの結合荷重であり、 $N \times G$ 行列である。本論文の実験では、 g を観測された養育者の顔画像と予め登録しておいた様々な向きの顔画像との輝度に関する一致度に基づいて算出する。ただし、最も近いと判定された顔の要素のみに1、残りの要素には0が割り当てられる。

ここで、学習者は環境中のいずれかの物体を必ず見ると仮定すると、 θ^g はその i 番目の要素が

$$\tilde{\theta}_i^g = \begin{cases} \theta_i^g & \text{if } s_i > 0 \\ 0 & \text{otherwise} \end{cases} \quad (2)$$

であるような N 次元ベクトル $\tilde{\theta}^g$ に修正される。ここで s_i は環境中の物体配置を表現する N 次元ベクトル s の i 番目の要素とし、 i 番目のスポットに物体が配置されていればその物体のIDである id 、そうでなければ0が i 番目の要素に設定されるとする。学習者は s を観測可能であると仮定する。

2.2 言葉による注意モジュール

言葉による注意モジュールは、養育者が発した物体のラベルについての観測情報を表す L 次元ベクトル l を入力として、共同注意のために注視されるべき物体を

$$\theta^l = W^l l \quad (3)$$

のように出力する。ここで、 θ^l は、養育者の発話を手がかりとしたときに各物体が注視されるべき程度を要素とする M 次元ベクトルである。また W^l はマッピングのパラメータである $M \times L$ 行列である。本論文の実験では、 l として、養育者の発話と予め登録された単語との一致度を要素とするベクトルを用いる。ただし、最も近いと判定された単語の要素のみに1、残りの要素

には0が割り当てられるとする。

ここで、養育者の発話を手がかりとしたときに環境中の各スポットが注視されるべき程度を要素とする N 次元ベクトル $\tilde{\theta}^l$ は、その i 番目の要素を

$$\tilde{\theta}_i^l = \begin{cases} o_{s_i}^l & \text{if } s_i > 0 \\ 0 & \text{otherwise} \end{cases} \quad (4)$$

のように決定することで求められる。

2.3 統合

各注意モジュールの出力 $\tilde{\theta}^g$ と $\tilde{\theta}^l$ を統合して、注視対象が決定される。決定には、統合されたベクトル $\tilde{\theta} = \tilde{\theta}^g + \tilde{\theta}^l$ を大きさが1になるように正規化したものを用い、これを各スポットの確率として、確率的に注視対象のスポットが選択される。

この確率的な選択により、 i 番目のスポットが選択されたとき、これに対応する i 番目の要素が1、その他の要素が0と設定される N 次元ベクトルを θ と呼ぶ。またその結果 j 番目のIDを持つ物体が注視されたとき、この物体に対応する j 番目の要素が1、その他の要素が0と設定される M 次元ベクトルを o と呼ぶ。ここで θ_d および o_d を、養育者が注視しているスポットおよびそのスポットにある物体のIDに対応する要素が1、その他の要素が0と設定される N および M 次元ベクトルであるとする。 $\theta = \theta_d$ および $o = o_d$ のとき、共同注意に成功したと見なされる。ただし、学習者には共同注意が成功したか否かの情報は与えられない。

2.4 相互排他性原理に基づく学習則

各試行において、学習者は観測情報 g および l 、そしてこれらをもとに実行された注視 θ とその結果観測された物体の情報 o をもとに、各モジュールの学習を行う。ここで養育者の、各スポットの注視の仕方と各物体のラベルのつけ方がそれぞれ相互排他的であると仮定する (Fig.3 参照)。すなわち、養育者の注視の様子 g と注視しているスポット θ_d 、および養育者の発話したラベル l と注視している物体のID o_d が、それぞれ1対1の対応関係を持っているとする。このとき、養育者と共同注意を達成するため、学習者はマッピングの結合荷重である W^g および W^l を変更することで、この対応関係を学習する必要がある。

学習者が観測できる情報 θ と o は、養育者の注視を表す θ_d と o_d とは必ずしも一致しない。しかし、周囲の物体あるいはその位置が変化する状況の中で経験を積み重ねていくと、 g と l が観測されたときに θ と o が θ_d と o_d に一致する傾向が高いという統計的性質を見出すことができる。これを利用して、ヘッブ学習のような統計的な対応を学習する手法により、共同注意のための視線追従 [6] [7] や語彙 [8] [9] のマッピングを学習可能であることが先行研究で示されている。しかしこれらの先行研究では、学習すべき複数の対応関係がそれぞれ相互排他的なものであること、また共同注意に利用可能な注意モジュールとしてそれらが相補的であることを利用できていなかった。

これに対し、対応関係がどれだけ相互排他的に見えるかに応じて、観測された対応関係をマッピングに反映させる度合いを調整する交差投錨型ヘッブ則が提案されている [13]。本研究ではこれを修正した学習則をマルチモーダルな共同注意学習の間

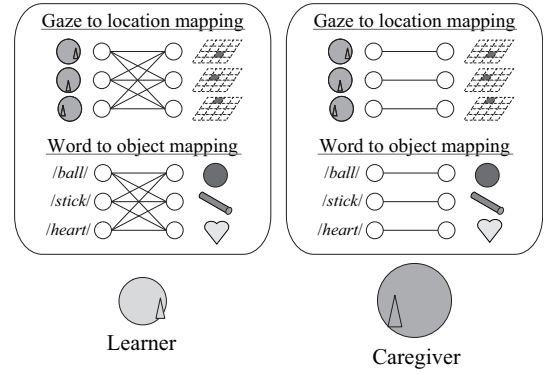


Fig. 3 Examples of initial mappings of the learner and mutually exclusive ones of the caregiver

題に適用することで、過去の経験を利用することによる各モダリティの学習の加速、および、モダリティ間の学習の相互促進を実現する。

注視行動を行った後、学習者は自身の注視経験に基づいて各マッピングの結合荷重を更新する。視線追従と語彙のマッピングには同様の更新則が適用されるため、ここでは視線追従のマッピングを例として説明する。視線追従マッピングの入力 g とシステム全体の出力 θ の最大値を持つ要素を、それぞれ勝者 i^* と j^* と呼ぶ。ここで t ステップ目に、 θ により、いずれかの物体が注視されたとき、視線マッピングの i^* 番目の入力要素と j^* 番目の出力要素の間の結合荷重 $w_{i^*j^*}^g(t)$ は

$$w_{i^*j^*}^g(t+1) = \mu_g(g_{i^*}\theta_{j^*} - w_{i^*j^*}^g(t)) \quad (5)$$

のように更新される。ここで μ_g は現在の入出力関係の視線追従マッピングにおける相互排他度を表す。すなわち、視線追従マッピングの結合荷重として、現在の入出力関係がどの程度相互排他的なものとして表されているかを表す。本論文では収束値近傍での学習の挙動がより安定になるように、係数 μ_g を $g_{i^*}\theta_{j^*}$ ではなく、 $(g_{i^*}\theta_{j^*} - w_{i^*j^*}^g(t))$ に掛ける形に交差投錨型ヘッブ則 [13] 修正している。 μ_g はマッピングの入力要素の排他度 η_{ij}^g と出力要素の排他度 η_{ij}^o の積 $\mu_g = \eta_{i^*j^*}^g \cdot \eta_{i^*j^*}^o$ として求められる。各入出力要素の排他度は

$$\eta_{ij}^g = \exp\left(-\frac{\sum_{k,k \neq j} w_{ik}^g(t)}{\alpha^2}\right) \quad (6)$$

$$\eta_{ij}^o = \exp\left(-\frac{\sum_{k,k \neq i} w_{kj}^g(t)}{\alpha^2}\right) \quad (7)$$

と計算される。ここで、 α は相互排他性の逆感度を表すパラメータである。従って、相互排他度 μ_g は入出力の勝者間が相互に排他的であるほど大きくなる。つまり、式 (5) の更新則では、より相互排他的な関係にある入出力要素間の結合関係については学習が進む一方、そうでない入出力要素間の結合関係についてはあまり学習が進まないことにより、相互排他的なマッピングが構築されることが期待される。

同時に、入力 of 勝者要素と出力の勝者以外の要素との結合荷重、および出力の勝者要素と入力の勝者以外の要素との結合荷重は側抑制によって

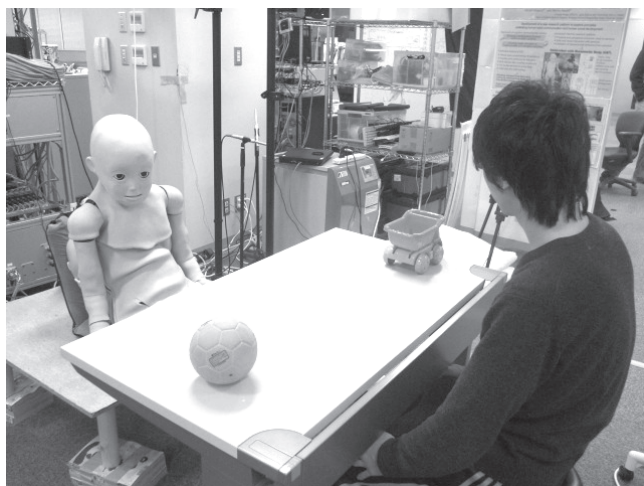


Fig. 4 A sample scene of human-robot interaction of joint attention

$$w_{ij*}^g(t+1) = w_{ij*}^g(t) - \beta w_{ij*}^g(t) \quad (8)$$

$$w_{i*j}^g(t+1) = w_{i*j}^g(t) - \beta w_{i*j}^g(t) \quad (9)$$

のように減らされる．ここで β は側抑制の強さを表すパラメータである．以上の学習則によって学習者は養育者のもつ相互排他的なマッピングを学習する．

2.4.1 提案手法によるモジュール内の相互促進作用

提案手法はモジュール内の相互促進的学習を実現するように設計されている．式 (5) の更新は入力・出力の勝者が他方の勝者以外の要素との間に相互排他的な関係が見出されていないとき，入力の勝者・出力の勝者との結合荷重が増やされやすい更新則になっている．また式 (8) および式 (9) の側抑制によって相互排他的な関係が見出された入力の勝者と出力の勝者は勝者以外の要素との結合荷重を次第に弱めるようになる．これにより，提案手法はモジュール内の対応学習を加速する相互排他性バイアスとして機能すると期待される．

2.4.2 提案手法によるモジュール間の相互促進作用

提案手法はモジュール間の相互促進的学習も実現するように設計されている．交差投鋳型ヘップ則によって学習された結合荷重は相互排他的な関係が発見されたときに大きな値となる．学習者は共同注意を行う際，各注意モジュールの出力を統合し，統合された出力に従って物体を注視する．従って，より相互排他的な対応関係が学習された注意モジュールの出力が共同注意の試行に利用されることとなり，他方の注意モジュールは，相互排他的な関係の学習が進む前から，相互排他的な関係をより頻繁に経験できることになる．これより，提案手法はモジュール間の対応学習を促進する相互排他性バイアスとして機能すると期待される．

3. 実 験

提案手法により，複数モジュールの学習がどのように相互促進可能であるかを検証する．まず事例として，比較的簡単な状況における，実ロボットと人とのインタラクションを通じた共同注意学習場面に提案手法を適用した結果について述べる．そ

Table 1 Success rates of face and object recognition

Gaze direction	Face recognition rate (%)	Object's label	Object recognition rate (%)
right	100	car	100
center	97.5	ball	98.1
left	88.2	bat	98.4

の後，計算機シミュレーションを用い，提案手法によるモダリティ内およびモダリティ間での相互促進効果について検証する．最後に，親子インタラクションを模した状況での計算機シミュレーション結果を示し，幼児の発達過程との相同性について議論する．

3.1 実ロボットを用いたインタラクション実験

本実験では子供型ヒューマノイドロボット CB² (*Child robot with Biomimetic Body*) [14] を用いた．ロボットと教示者はテーブルを挟んで互いに向かい合わせに座らせる (Fig.4 参照)．実験では，教示する物体の数を 3 とし，テーブルを右・中心・左に 3 等分し，各スポットに物体を配置する ($N = 3$, $M = 3$)．一度に置かれる物体の数は 2 とし，3 個の物体の中からランダムに選択したものを，テーブルのいずれかの位置に重複しないように配置する ($M_o = 2$)．配置する物体とその位置は 5 ステップ毎に入れ替えられる ($T = 5$)．ロボットのタスクは，共同注意ができるように，ロボットから見て右・中心・左に向いている教示者の顔画像を，それぞれ右・中心・左のスポットに対応付けること，およびそれぞれの物体とそれらに固有なラベルを対応付けることである．ここでロボットは共同注意が成功したかどうかを直接知ることができないため，その是非に関わらず，注視した物体をもとに各注意モジュールの結合荷重を更新することに注意されたい．

視線による注視モジュールの入力として g に眼球の方向ではなく顔の方向を用いる．高い精度で視線を捉えるためには，眼球の方向についても考慮することが望ましい．しかし，幼児の知見 [1] からも顔の方向は人の注視方向を知る有効な手がかりの一つであることが知られており，本研究では簡単のため顔の方向についての観測から g を求める．具体的には，観測された画像に対して予め登録された様々な向きの顔画像の輝度値のテンプレートマッチングを行い，各顔画像の最大の相関値を一致度として g を求めた．また同様に，観測された画像に対して予め登録された物体画像の色差ヒストグラムのテンプレートマッチングを行い，各画像の最大の相関値の最も高い値をもつものから物体の ID を求めた．顔認識成功率および物体認識成功率を Table 1 に示す．それぞれの平均値は 95.2%, 98.8%であった．言葉についての入力ベクトル l は汎用大語彙連続音声認識エンジン Julius を用いて求めた．教示者は物体のラベルのみを発話することとしたため，本実験での音声認識成功率は 100%であった．物体の配置の方法については，ランダム性を保持するようシステムに決定させた．この配置については画像処理で特定可能であるが，本実験では簡単のためロボットに直接教示した．更新則のパラメータは経験的に $\alpha = 1.2$, $\beta = 0.5$ とした．

Fig.5(a) は，50 ステップのインタラクション実験を 5 回行った際の，直近の過去 10 ステップの共同注意成功率の平均値の

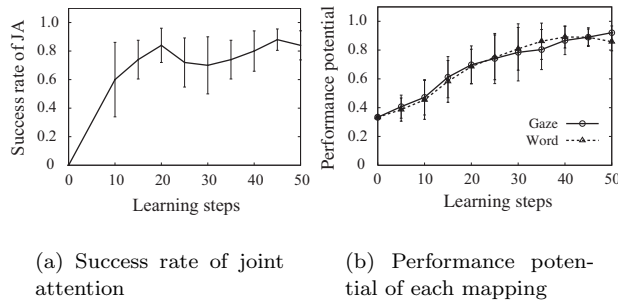


Fig. 5 Learning progresses in human-robot interaction

遷移を示している。また Fig.5(b) は、そのときの視線追従および語彙のマッピングの習熟率の遷移を示している。習熟率とは、そのマッピングを用いて、正しい注視位置が選択される確率の平均値のことであり、各時点でのマッピングに対して、起こりうる全ての場合の g と l を入力し、正しい選択がなされる確率を計算することで求めた。Fig.5(a) から、共同注意の成功率は 20 ステップあたりでおよそ 80% に達しており、また Fig.5(b) から各マッピングの習熟率は 30 ステップあたりでおよそ 80% に達していることがわかる。ここではパラメータ α , β を固定して実験を実施したが、 $\alpha = [0.3, 1.6]$, $\beta = [0.3, 0.6]$ の範囲で同様の学習が可能であることをこの実験設定を再現した計算機シミュレーションにより確認している。

本インタラクション実験において 1 ステップの所要時間は 20[sec] くらいであったため、10 分程度という早い時間で、人の右・中心・左の視線方向および 3 つの対象物の名前を理解して共同注意が可能であるといえる。また、入力ベクトルの認識率は 100% でなかったが、学習は可能であった。認識率の低さがどのように影響するかについては、次節でより詳しく検討する。

3.2 計算機シミュレーションによる相互促進効果の検証

前節での実例のように、実環境において養育者の言葉や視線を常に 100% 認識することは難しい。認識にエラーを含む場合、統計的な対応関係を発見し、学習することはさらに難しくなる。これに対し、提案手法はモジュール内相互促進とモジュール間相互促進によって、認識エラーを含む状況においてもより効率的な学習が可能であることが期待される。そこで人工的に入力の認識にエラーを導入した計算機シミュレーションを実施し、その効果を検証する。具体的には、入力の認識率 p_r と呼ぶ確率を定め、確率 p_r で正しいベクトルが、確率 $1 - p_r$ で誤ったベクトルが入力される、という形の認識エラーを導入する。ここで、もともとの入力ベクトルに近いものの誤認識が起きやすいとし、具体的には、入力ベクトルが正しく認識された際に選択されるべき要素の ID を中心とするガウス分布に従って、選択される入力の勝者要素を決定した。認識率は、ガウス分布の分散を変えることで調整した。

上記の認識エラーを導入した共同注意学習の計算機シミュレーションにおいて、次節ではまず単一のモジュールを用いた共同注意学習を想定し、視線による共同注意に対して相互排他性原理に基づく学習則を適用した場合のパフォーマンスを、関連学

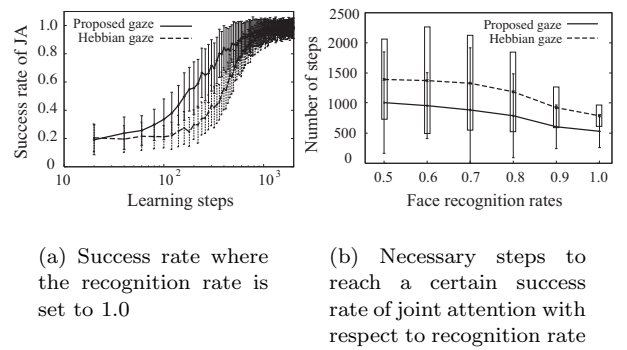


Fig. 6 Comparison of learning results with single module between the proposed method and Hebbian learning

習の方法の一つであるヘブ学習のそれと比較することで、モジュール内の相互促進効果を検証する。次に視線および言葉による注意モジュールを用いたマルチモーダルな共同注意を想定し、提案手法による学習とヘブ学習の比較をすることでモジュール間の相互促進効果を検証する。比較対象のヘブ学習では、視線追従マッピングの結合を

$$w_{i,j}^g(t+1) = w_{i,j}^g(t) + \gamma(g_i \cdot \theta_j - w_{i,j}^g(t)) \quad (10)$$

のように更新させた。語彙マッピングの結合荷重更新も同様とした。ここで γ は学習率である。

3.2.1 モジュール内での相互促進効果

認識率 p_r を 0.5 から 1.0 まで 0.1 ずつ変化させたそれぞれの場合について、2000 ステップのインタラクションを通じて視線による共同注意モジュールを学習させるシミュレーションを 50 回実施した。ただし、インタラクションのパラメータを $N = 10, M = 10, M_o = 5$ とした。提案手法とヘブ学習の各パラメータは経験的に $\alpha = 0.8, \beta = 0.05, \gamma = 0.01$ とした。

Fig.6(a) に、認識率を $p_r = 1.0$ に固定した際の、直近の過去 20 ステップの共同注意成功率の遷移を示す。ヘブ学習 (Hebbian gaze: 破線) に比べ、提案手法 (Proposed gaze: 実線) の学習が早いことがわかる。Fig.6(b) は、直近の過去 100 ステップの共同注意成功率が、認識率が変えられた各状況においてある程度よいパフォーマンスであると考えられる成功率を 1% 水準で有意に上回るのに要したステップ数の平均値を示したものである。ここで、ある程度よいパフォーマンスの基準となる成功率として、正しいマッピングをシステムに与えて共同注意を試みさせたときの成功率の 80% の確率を用いた。分散は大きいですが、ヘブ学習 (Hebbian gaze: 破線) に比べ、提案手法 (Proposed gaze: 実線) の方が認識率が低い場合においても、早くなっていることがわかる。従って、モダリティ内の相互促進が実現されているといえる。ここではパラメータ α , β , γ を固定して実験を実施したが、 $p_r = 1.0$ のとき、 $\alpha = [0.6, 1.2]$, $\beta = [0.01, 0.1]$, $\gamma = [0.01, 0.1]$ の範囲で同様の学習が可能であることを確認している。

3.2.2 モジュール間での相互促進効果

視線と言葉による共同注意モジュールの両方を用い、前節と同様に認識率 p_r を変えて、2000 ステップのインタラクション

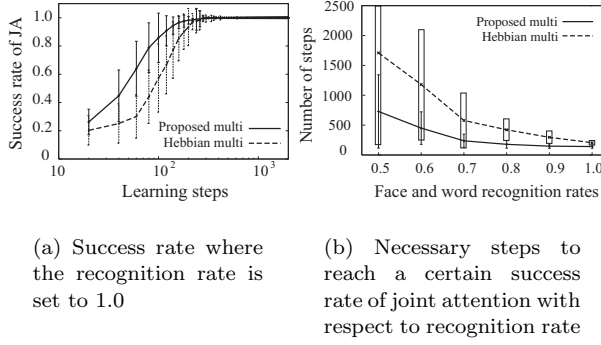


Fig. 7 Comparison of success rate of joint attention with gaze-driven and word-driven modules between the proposed method and a Hebbian learning where the recognition rate is set to 1.0

を通じて、これらの共同注意モジュールを学習させるシミュレーションを 50 回実施した。ただし、インタラクションのパラメータを $N = 10, M = 10, M_o = 5$ とした。また、学習のパラメータは経験的に $\alpha = 1.2, \beta = 0.5, \gamma = 0.05$ とした。

Fig.7(a) に視線およびラベルの認識率を 1.0 に固定して、直近の過去 20 ステップの共同注意成功率の遷移を示す。ヘッブ学習 (Hebbian multi: 破線) に比べ、提案手法 (Proposed multi: 実線) の学習が早いことがわかる。Fig.7(b) は、Fig.6(b) と同様に、直近の過去 100 ステップの共同注意成功率が有意にある程度よくなるのに要したステップ数の平均値を示したものである。ここで比較のため、ある程度よいパフォーマンスの基準となる成功率として、Fig.6 の描画の際に用いた値を用いた。ヘッブ学習 (Hebbian multi: 破線) に比べ、提案手法 (Proposed multi: 実線) の学習が速いことがわかる。Fig.6(b) と Fig.7(b) を比較すると、視線のみを利用した単一モジュールの共同注意システムより、マルチモジュールの共同注意システムの方の学習が速く、分散も小さいことがわかる。これは早期に学習できた側のマッピングを共同注意の試行に利用することによる、促進効果であるといえる。

そこで、それぞれのモダリティの認識率が異なる状況を設定し、提案手法の促進効果をより詳細に検証した。Fig.8 に視線とラベルの認識率をそれぞれ独立に変更したシミュレーションの結果を示す。ただし、提案手法については Fig.8(a),(c),(e) に、ヘッブ学習については Fig.8(b),(d),(f) に示した。Fig.8(a) と (b) は、直近の過去 100 ステップの共同注意成功率が有意にある程度よくなるのに要したステップ数の平均値を色の薄さで示している。ここで、ある程度よいパフォーマンスの成功率として、正しい視線およびラベルのマッピングをシステムに与えて共同注意を試みさせたときの成功率の 80% の確率を用いた。視線とラベルの認識率が比較的低い場合に特に顕著に、提案手法の方がある程度よいパフォーマンスに達するのが早いといえる。また Fig.8(c) と (d) は視線のマッピングの、Fig.8(e) と (f) はラベルのマッピングの習熟率が、同じ基準の確率を有意に上回るのに要するステップ数を色の薄さで示している。視線追従のマッピングについて比較すると、ヘッブ学習 (Fig.8(d)) では、ラベ

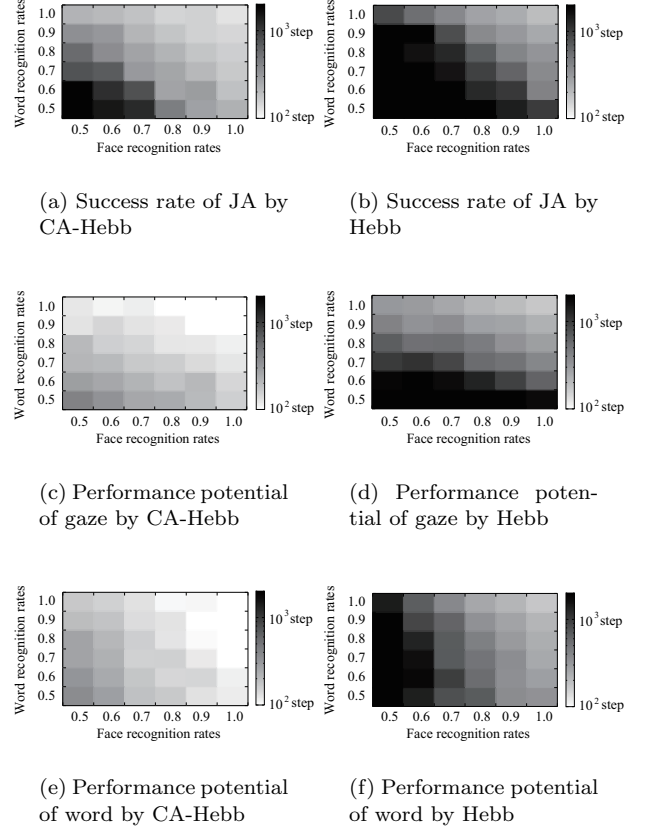


Fig. 8 Comparison of necessary steps to reach a certain level of learning performance between the proposed method (CA-Hebb: (a), (c), (e)) and a Hebbian learning (Hebb: (b), (d), (f)) under different pairs of recognition rates: success rates of joint attention (a) and (b), performance potential (c) and (d) of gaze-driven module, and performance potential of word-driven module.

ルの認識率が低い場合に、習熟率の向上が遅いのにに対し、提案手法 (Fig.8(c)) では、比較的に早くなっていることがわかる。一方ラベルのマッピングを比較すると、ヘッブ学習 (Fig.8(f)) では視線の認識率が低い場合に習熟率の向上が遅いのにに対し、提案手法 (Fig.8(e)) では比較的に早くなっていることがわかる。これらの結果は、より早期に相互排他性が向上した側のマッピングが共同注意の試行に利用されるため、他方が認識率が低くても、相互排他性を向上させることができていることを示しており、提案手法によるモジュール間の相互促進が実現できているといえる。ここではパラメータ α, β, γ を固定して実験を実施したが、 $p_r = 1.0$ のとき、 α がそれぞれ 1.0 以上、 $\gamma = [0.01, 0.3]$ の範囲で同様の学習が可能であることを確認している。

3.3 検証 2：養育者を考慮した学習過程

最後に親子インタラクションを模した状況での計算機シミュレーションを実施し、幼児の発達過程をどの程度再現できるかについて検討した。幼児は 18 ヶ月あたりでは、最大で 200 語、平均 100 語のラベルを発話できる [2] とされている。そこで本節では、提示する物体の数をこれと同程度の数に合わせ、 $M = 200$ とした計算機シミュレーションを実施する。環境の分割数は教

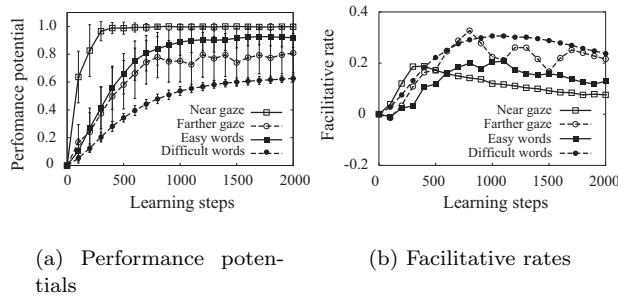


Fig. 9 Learning progress under the resembled mother-infant interaction: (a) performance potentials of near- and further-gaze as well as easy- and difficult-word mappings and (b) facilitative rates of them

示される物体の数に比べ、かなり小さいと考え、 $N = 10$ とした。ここで幼児が獲得するラベルの中には聞き取りやすいものとそうでないものが含まれており、また遠くの領域の異なる位置を見ている養育者の顔の違いは小さいのに対し、近くのそれは比較的大きいと考えられる。そこでこれらの違いを認識率の違いと考え、計算機シミュレーションでは各モダリティの入力を二つのグループに分けた。具体的には認識の容易な要素（近くの視線および簡単な言葉）の認識率を 1.0、認識の難しい要素（遠くの視線および難しい言葉）の認識率を 0.7 とし、各グループの要素数は近くの視線を 5、遠くの視線を 5、簡単な言葉を 10 および難しい言葉を 190 とした。

実際の親子インタラクションでは、養育者は“Learning from Easy Mission” [15] と呼ばれるような学習者の学習を促進させる、学習者に配慮した行動を行うと期待される。そこで本節のシミュレーションでは、養育者は学習者の語彙理解の能力に配慮して振舞うとする。具体的には、学習者がそのラベルを知っているように見える物体を優先的に選択させるように、幼児が 5 回連続で共同注意を成功させた物体があれば、これを選択させることとした。また学習者がラベルを知っているように見える物体がない場合には、近い場所にありかつ簡単なラベルの物体を優先して選択させるとする。しかし、学習が進むと遠くにも学習者が知っている物体が存在する場面が増えていくため、養育者は遠い領域でのインタラクションを増やしていくことになる。

上記の設定で 2000 ステップの共同注意学習シミュレーションを 50 回実施した。学習のパラメータは 3.2.2 節と同様に $\alpha = 1.2, \beta = 0.5$ とした。Fig.9(a) に、視線追従および語彙のマッピングの習熟率の平均値の遷移を 100 ステップごとに示す。およそ 300 ステップあたりで近くの視線追従 (Near gaze: □の実線) ができるようになり、遅れておよそ 800 ステップあたりで遠くの視線追従 (Farther gaze: ○の破線) ができるようになるという学習過程が現れた。これは学習者が言葉の注意モジュールを利用し、遠くのある場所にある知っているラベルの物体を注視することで、視線追従の学習が促進できたためであるといえる。一方幼児の共同注意は、12ヶ月頃には、幼児自身の視野内にある近くの物体に対するものに限られるのに対し、18ヶ月

頃では目の前に物体がない場合にも視野外の遠くの物体を探しあて、共同注意ができるようになること [1] が知られており、本節のシミュレーションの結果に一定の相同性が認められる。

Fig.9(b) に視線追従および語彙のマッピングの習熟率の促進度の遷移を 100 ステップ毎に示す。ここで促進度とはマルチモーダルな共同注意学習システムのマッピングの習熟率と単一モダリティのみを利用した共同注意学習システムのマッピングの習熟率との差とした。Fig.9(b) から難しい言葉の語彙マッピングの学習の促進度 (Difficult words: ●の破線) が比較的大きくなり始めるのは 300 ステップを越えたあたりからであることがわかる。ここで Fig.9(a) を見ると、その 300 ステップあたりはちょうど近くの視線追従 (Near gaze: □の実線) が獲得できた頃であり、言葉としては簡単な言葉 (Easy words: ■の実線) を 4 割、難しい言葉 (Difficult words: ●の破線) を 2 割覚え、合わせて 40 語程度をようやく覚えた頃である。このことから、遠くの視線追従 (Farther gaze: ○の破線) が未熟な段階からも近くの視線追従 (Near gaze: □の実線) を語彙の学習に利用でき始めていることがわかる。また Fig.9(b) から、1000 ステップあたりで両方の語彙マッピング学習の促進度 (Easy words: ■の実線、および Difficult words: ●の破線) は最大となることわかる。これは Fig.9(a) を見ると、学習できた遠くの視線追従 (Farther gaze: ○の破線) を十分に利用できるようになった頃に対応しているようで、言葉としては簡単な言葉を 9 割、難しい言葉を 5 割覚え、合わせて 100 語程度覚えた頃に対応するようである。一方幼児は、平均的に 100 語程度を獲得する 18ヶ月ごろ [2] に、物体のラベルの学習に親の視線情報を利用できるようになることが知られている [3]。このような乳児の振る舞いと、本シミュレーションにおいてある程度の語数を覚える頃と視線によるラベルマッピングの学習の促進度が最大になる頃が同時期であったことに一定の相同性が認められ、興味深い。しかし本シミュレーションは特定のパラメータのみを扱う限定的なものであるため、定量的な一致までを議論することは困難であり、パラメータの妥当性などについてのより詳細な検証が今後の課題となる。

4. おわりに

本研究では視線追従および言葉によるマルチモーダル共同注意の同時学習において、相互排他性バイアスを利用する学習システムを提案した。実ロボットを用いた人とのインタラクション実験および計算機シミュレーションに提案手法を適用し、各モダリティの共同注意の機能の学習において、モダリティ内およびモダリティ間での相互促進効果を確認した。

これまでにも複数モジュールを同時学習させることにより複雑な対象を効率的に学習させる研究が行われてきた [16]~[20]。MOSAIC [16] や階層混合モデル [17]、RNN エキスパート混合モデル [18] 等の分割統治型の複数モジュール学習の手法では、状態空間を分割し、各モジュールにそれぞれを分担させることで、効率的な学習を実現している。これに対し提案手法では、このような分担については議論せず、視線やラベルといった異なるモダリティの情報に基づく複数のモジュールの競合的な統合を通じて、いかに学習を相互促進的に進めさせられるかに注目

している点で異なる。一方 Uchibe et al. [19] は構造の異なる複数の強化学習モジュールの競合的な同時学習において、性能は高くないが学習の速いモジュールを用いたときの経験も学習に利用することで、性能は高いが学習の遅いモジュールを実時間で学習させられることを示した。提案手法も、視線および言葉を利用したマルチモーダルな共同注意の学習において、早期に性能が良くなった部分を双方に利用しあうことで効率的な学習を実現できている点では同じである。ただし、提案手法は1対1 マッピングの構築に特化したものであるという点と、これに有効な相互排他性バイアスを積極的に利用することで、モダリティ間のみならずモダリティ内においても相互促進的学習を実現したものであるという点で異なると考えられる。また Sugita et al. [20] は、動作の系列と言葉の系列を学習するリカレントニューラルネットワークをそれぞれのPB値と呼ばれる変数を共有させることで結合したネットワークを提案し、特定の行動についての教師あり学習により、これらの系列の汎化が可能であることを示した。この研究は異なるモジュールの学習を別のモジュールの学習に利用するという点で本研究と同じであるが、Sugita et al. [20] はこれを汎化に利用したのに対し、提案手法は双方の学習の相互促進に利用しているという点、またその結果、厳密に教師信号が与えられなくても学習を可能にしている点で異なると考えられる。しかし、提案手法はこれらの従来研究と排他的ではなく、これらと融合させることが今後の一つの課題である。

また親子インタラクションを模した状況において提案手法が、Butterworth et al. [1] の示す視線追従の段階的発達過程および Baldwin [3] の示す 18ヶ月児の視線情報を利用した語彙学習と対応する発達過程を再現可能であることを示した。本稿では一組のパラメータによる計算機シミュレーションの結果についてのみ議論したが、これらのパラメータの妥当性を検証するとともに、これらを変動させた場合に提案手法が幼児の発達のバリエーションをどの程度再現できるかを検討する必要がある。

人の幼児は、共同注意学習を始める以前に、視線追従マッピングの学習に必要な養育者の顔の認識 [21] や注視方向を向くための首振り運動 [22]、および語彙マッピングの学習に必要な養育者の発話するラベルの認識 [23] や物体の認識 [24] 等の能力をある程度獲得しているということが発達心理学の研究で知られている。ただし、我々はこれらの能力は未熟であると考え、本論文の実験では視線およびラベルといった入力への認識にエラーを導入し、このようなエラーに対する提案手法の有効性を確認した。実験では入力についてのエラーのみしか考慮しなかったが、各モジュールの出力値の処理に必要な物体の認識および首振り運動の能力についてのエラーも今後考慮していく必要がある。

しかし、本論文の実験では入力への認識にエラーがあることを想定していたものの、その程度については一定であるとみなしていた。一方、人の幼児を見ていると、言葉を覚え始める 14ヶ月ごろから音節認識が語彙学習に特化したものに変わっていくこと [25] や言葉を 100 語覚える時期の前後で物体のカテゴリ化の能力が変化すること [26] が示唆されており、幼児は共同注意を発達させていく中で共同注意に必要なこれらの離散化にかか

る処理の能力を発達させると考えられる。しかし、入力をどの程度性格に離散化できるか、また離散的なマッピングの出力をどの程度正確に行動に反映させられるかは、どのようなマッピングを学習すべきか、に相互に依存するため、マッピングと離散化は同時並行的に学習されることが望ましい。本論文では入力の認識エラーに対し、過去の経験や他のモダリティを積極的に利用することにより、認識率の低さがある程度許容されることを実験的に見た。従って、提案手法に基づく方法では、離散化の学習と並行したマッピングの学習が可能であることが期待され、これを実現することが今後の課題となる。このように、モデルを段階的に洗練していくことが、人の発達を支える学習メカニズムを解き明かすための有効なアプローチであると考えられる。

参考文献

- [1] G. Butterworth and N. Jarrett: "What minds have in common is space: Spatial mechanisms serving joint visual attention in infancy", *British Journal of Developmental Psychology*, vol.9, no.1, pp.55-72, 1991.
- [2] E. Bates, P.S. Dale and D. Thal: "Individual differences and their implications for theories of language development", *The handbook of child language*, pp.96-151, 1995.
- [3] D.A. Baldwin: "Infants' contribution to the achievement of joint reference", *Child Development*, vol.62, no.5, pp.875-890, 1991.
- [4] P.E. Spencer: "Looking Without Listening: Is Audition a Pre-requisite for Normal Development of Visual Attention During Infancy?", *Journal of Deaf Studies and Deaf Education*, vol.5, no.4, pp.291-302, 2000.
- [5] M. Asada, K.F. MacDorman, H. Ishiguro and Y. Kuniyoshi: "Cognitive developmental robotics as a new paradigm for the design of humanoid robots", *Robotics and Autonomous Systems*, vol.37, pp.185-193, 2001.
- [6] Y. Nagai, K. Hosoda, A. Morita and M. Asada: "A constructive model for the development of joint attention", *Connection Science*, vol.15, no.4, pp.211-229, 2003.
- [7] C. Teuscher and J. Triesch: "To each his own: The caregiver's role in a computational model of gaze following", *Neurocomputing*, vol.70, no.13-15, pp.2166-2180, 2007.
- [8] D.K. Roy and A.P. Pentland: "Learning words from sights and sounds: a computational model", *Cognitive Science*, vol.26, no.1, pp.113-146, 2002.
- [9] C. Yu: "The emergence of links between lexical acquisition and object categorization: a computational study", *Connection Science*, vol.17, no.3, pp.381-397, 2005.
- [10] 菊池, 荻野, 浅田: "顕著性に基づくロボットの能動的語彙獲得", *日本ロボット学会誌*, vol.26, no.3, pp.226-270, 2008.
- [11] C. Yu, D.H. Ballard and R.N. Aslin: "The Role of Embodied Intention in Early Lexical Acquisition", *Cognitive Science*, vol.29, no.6, pp.961-1005, 2005.
- [12] 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.
- [13] Y. Yoshikawa, K. Hosoda and M. Asada: "Unique association between self-occlusion and double-touching towards binding vision and touch", *Neurocomputing*, vol.70, no.13-15, pp.2234-2244, 2007.
- [14] T. Minato, Y. Yoshikawa, T. Noda, S. Ikemoto, H. Ishiguro and M. Asada: "CB2: Child robot with Biomimetic Body for Cognitive Developmental robotics", *Proceeding of IEEE-RAS International Conference on Humanoid Robots*, 2007.
- [15] M. Asada, S. Noda, S. Tawaratsumida and K. Hosoda: "Vision-based reinforcement learning for purposive behavior acquisition", 2008.

- sition”, Proceeding of IEEE International Conference on Robotics and Automation, pp.146–153, 1995.
- [16] D.M. Wolpert and M. Kawato: “Multiple paired forward and inverse models for motor control”, Neural Networks, vol.11, no.7-8, pp.1317–1329, 1998.
- [17] M.I. Jordan and R.A. Jacobs: “Hierarchical Mixtures of Experts and the EM Algorithm”, Neural Computation, vol.6, no.2, pp.181–214, 1994.
- [18] J. Tani and S. Nolfi: “Learning to perceive the world as articulated: an approach for hierarchical learning in sensory-motor systems”, Neural Networks, vol.12, no.7-8, pp.1131–1141, 1999.
- [19] E. Uchibe and K. Doya: “Competitive-Cooperative-Concurrent Reinforcement Learning with Importance Sampling”, Proceeding of International Conference on Simulation of Adaptive Behavior: From Animals and Animates, pp.287–296, 2004.
- [20] Y. Sugita and J. Tani: “Learning semantic combinatoriality from the interaction between linguistic and behavioral processes.”, Adaptive Behavior, vol.13, no.1, pp.33–52, 2005.
- [21] T. Farroni, S. Massaccesi, E. Menon and M.H. Johnson: “Direct gaze modulates face recognition in young infants”, Cognition, vol.102, no.3, pp.396–404, 2007.
- [22] C.V. Hofsten: “An action perspective on motor development”, TRENDS in Cognitive Sciences, vol.8, no.6, pp.266–272, 2004.
- [23] J.R. Saffran, R.N. Aslin and E.L. Newport: “Statistical learning by 8-month-old infants”, Science, vol.274, no.5294, pp.1926, 1996.
- [24] B. Younger and L. Cohen: “Infant Perception of Correlations among Attributes”, Child Development, vol.54, pp.858–867, 1983.
- [25] C.L. Steger and J.F. Werker: “Infants listen for more phonetic detail in speech perception than in word-learning tasks”, Nature, vol.388, no.6640, pp.381–382, 1997.
- [26] L. Smith: “Learnig to Recognize Objects”, Psychological Science, vol.14, no.3, pp.244–250, 2003.

SMC societiesなどの会員。

(日本ロボット学会正会員)

石黒 浩 (Hiroshi Ishiguro)

1991 年大阪大学大学院基礎工学研究科物理系専攻修了。工学博士。同年山梨大学工学部情報工学科助手。1992 年大阪大学基礎工学部システム工学科助手。1994 年京都大学大学院工学研究科情報工学専攻助教授, 1998 年同大学大学院情報学研究科社会情報学専攻助教授。この間, 1998 年より 1 年間, カリフォルニア大学サンディエゴ校客員研究員。2000 年和歌山大学システム工学部情報通信システム学科助教授。2001 年より同大学教授。1999 年より, ATR 知能映像研究所客員研究員。現在, 大阪大学大学院工学研究科知能・機能創成工学専攻教授および ATR 知能ロボティクス研究所客員室長。知能ロボット, アンドロイドロボット, 知覚情報基盤の研究に興味をもつ。日本ロボット学会, 人工知能学会, 電子情報通信学会, 情報処理学会, IEEE, AAAI, ACM などの会員。
(日本ロボット学会正会員)

中野 吏 (Tsukasa Nakano)

2009 年大阪大学大学院工学研究科知能・機能創成工学専攻博士前期課程修了。同年同博士後期課程入学し現在に至る。認知発達ロボティクスの研究に従事。
(日本ロボット学会学生会員)

吉川 雄一郎 (Yuichiro Yoshikawa)

2005 年大阪大学大学院工学研究科知能・機能創成工学専攻修了。工学博士。同年 ATR 知能ロボティクス研究所研究員。2006 年より JST ERATO 浅田共創知能プロジェクト研究員となり現在に至る。認知発達ロボティクスの研究に従事。
(日本ロボット学会正会員)

浅田 稔 (Minoru Asada)

1982 年大阪大学大学院基礎工学研究科物理系専攻修了。工学博士。同年大阪大学基礎工学部制御工学部助手。1989 年同大学工学部電子制御機械工学科助教授。1995 年同教授。1997 年同大学大学院工学研究科知能・機能創成工学専攻教授となり現在に至る。この間, 1986 年から 1 年間米国メリーランド

大学客員研究員。知能ロボットの研究に従事。1989 年情報処理学会研究賞, 1992 年 IEEE/RSJ IROS'92 Best Paper Award 受賞。1996 年日本ロボット学会論文賞受賞。2001 年文部科学大臣賞 科学技術普及啓発功績者の表彰。ロボカップ国際委員会会長。阪大 FRC ロボカップヒューマノイド研究プロジェクトリーダー。日本ロボット学会, 電子情報通信学会, 情報処理学会, 人工知能学会, 日本機械学会, 計測自動制御学会, システム制御情報学会, IEEE R&A, CS,