

Title	感覚と報酬の予測誤差に基づく内部順モデルの適応 - 計算論的モデルと行動実験検証
Author(s)	佐藤, 仁是
Citation	
Issue Date	2016-03
Type	Thesis or Dissertation
Text version	author
URL	http://hdl.handle.net/10119/13623
Rights	
Description	Supervisor : 田中 宏和, 情報科学研究科, 修士

感覚と報酬の予測誤差に基づく内部順モデルの適応 - 計算論的モデルと行動実験検証

北陸先端科学技術大学院大学
情報科学研究科

佐藤 仁是

2016 年 3 月

修 士 論 文

感覚と報酬の予測誤差に基づく内部順モデルの適応 - 計算論的モデルと行動実験検証

1410025 佐藤 仁是

主指導教員	田中 宏和
審査委員主査	田中 宏和
審査委員	党 建武
	鵜木 祐史

北陸先端科学技術大学院大学

情報科学研究科

2016 年 2 月

概 要

本研究では、内部順モデルが何が要因となりどう学習するのかを行動実験と計算論モデルによって検証することを目的とする。

ヒトの内部順モデルは、運動指令を基に自己の運動を予測する機構である。この内部順モデルによる予測値は、ヒトが新奇の環境へ適応するときに重要な要素として用いられる。ヒトが今までできたことを新奇の環境でできるようになることを、運動適応という。先行研究では回転に対して適応する運動回転適応において、観測値と内部順モデルによる予測値の差が回転への適応に重要であると示している。その先行研究では、被験者は視覚情報を与える場合とそうでない場合の二条件において与えられる回転に適応した。被験者は二条件において同一の回転に適応したが、運動回転適応と異なる実験において各条件での学習プロセスは異なる結果を得た。この結果、内部順モデルの回転への適応は、感覚予測誤差という内部順モデルによる運動予測値と外界からの視覚情報による観測値の差が引き起こすと述べた。しかし、先行研究の結果はヒトの認知プロセスを経由するような実験の結果であり、内部順モデルの適応状況を直接説明できていないと考える。このことから、本研究では被験者の運動によって直接内部順モデルの適応状況を確認することが必要であると考え、目標点跳躍課題を計画した。この実験を行動実験と計算論モデルによる数値シミュレーションを行い、内部順モデルの適応に影響する要素を検証した。

新しく提案した行動実験結果と数値シミュレーション結果により、感覚予測誤差が与えられる場合でも与えられない場合でも内部順モデルが回転へ適応する結果を得た。視覚情報が与えられる場合での目標点跳躍課題では、先行研究同様、感覚予測誤差が内部順モデルの回転への適応において重要な役割を示す結果を得た。しかし、視覚情報が与えられていない条件の被験者では先行研究とは異なる結果を得た。これら結果より、本研究では先行研究における学習モデルでは説明できない結果を数値シミュレーションにより検討し、ヒトが視覚的な情報が与えられない場合においても内部順モデルの回転への適応を引き起こす可能性を示唆した。

目次

第1章	序論：ヒトの内部順モデルは何が影響し環境へ適応するのか	1
1.1	研究の背景	1
1.1.1	内部順モデルは感覚予測誤差により回転を学習するとの提案	1
1.1.2	先行研究は内部順モデルの適応状況を直接確認できていない可能性	1
1.2	本研究の目的：従来と異なる方法で内部順モデルの適応を直接的に確認	2
1.3	本論文の構成	2
第2章	先行研究の問題点から新実験の提案	4
2.1	運動適応：ヒトは新奇の環境へもすばやく適応可能	4
2.1.1	運動回転適応	4
2.1.2	運動回転適応と内部順モデル	6
2.2	先行研究の問題点：運動結果を直接観測できていない点	10
2.2.1	問題点1：到達運動と到達点同定課題で用いる腕の不一致	10
2.2.2	問題点2：知覚および認知処理系からの影響	11
2.2.3	問題点3：報告による運動結果の確認	12
2.3	新実験の提案：運動結果によって内部順モデルの適応状況を直接確認	12
2.3.1	提案行動実験：目標点跳躍課題	15
2.3.2	行動実験方法	17
第3章	数値シミュレーション：二つの学習プロセスで運動回転適応を再現	20
3.1	運動回転適応実験のシミュレーション方法	20
3.2	目標点到達課題のシミュレーション方法	24
第4章	数値シミュレーション結果：誤差条件は右へ報酬条件は新目標点へ到達	25
4.1	運動回転適応実験の数値シミュレーション結果	25
4.2	目標点到達課題の数値シミュレーション結果	27
第5章	被験者実験：報酬・誤差両条件とも右側へ運動補正を生成	30
5.1	運動回転適応実験の結果	32
5.2	目標点到達課題の結果	32
5.3	被験者実験結果の考察	35
5.3.1	考察：目標点の跳躍時には運動補正が生成されているかの検討	35

5.3.2	考察：なぜ報酬条件では手先が目標点の右側へ到達したのかの検討	41
5.3.3	考察：被験者は異なる距離の目標点跳躍を区別できるのかの検討	46
第6章	結論	53
6.1	本研究で明らかになったこと：誤差条件だけでなく報酬条件でも内部順モ デルの回転への適応する可能性を示唆	53
6.2	今後の展望	54
6.2.1	実験中でのインストラクションの変化	54
6.2.2	カウンターバランスの考慮	54
付録1	カルマンフィルター器の導出過程	56
6.3	記号の定義	56
6.4	カルマンフィルターの導出過程	57
付録2	確認実験結果	61
付録3	確認実験	69
6.5	運動回転適応実験の結果	69
6.6	目標点到達課題の結果	73
6.7	確認実験結果の考察	73
6.7.1	考察：目標点の跳躍時には運動補正が生成されているかの検討	73
6.7.2	考察：確認実験の結果は数値シミュレーションと異なる結果	78

第1章 序論：ヒトの内部順モデルは何が 影響し環境へ適応するのか

1.1 研究の背景

1.1.1 内部順モデルは感覚予測誤差により回転を学習するとの提案

ヒトは、時々刻々と変化する世界に住み、その世界の中でより快適に生きられるように日々学習し、環境へ適応している。このヒトのもつ優れた学習能力に関して、ヒトがどのように学習して行くのかを研究することは教育方法の検討やスポーツにおいてより効率的なトレーニングを行える等の様々な観点で非常に重要である。ヒトは、様々な環境へ柔軟な適応が可能であり、この優れた適応能力の一つに運動適応がある。運動適応は、新奇の環境でも通常的环境で行う運動を回復する能力である。これは、脳内の運動予測をする内部モデルが環境を学習する事で、新規の環境にも適応できるためである。例えば、宇宙飛行士が宇宙空間に適応する事は内部モデルが誤差を学習する事による。また、脳内で生成する運動指令を基に運動を予測するモデルを、内部順モデルと呼ぶ。内部順モデルと運動については、力場適応や運動回転適応、プリズム適応の実験により多く研究されている [1-7]。Izawa と Shadmehr ら [5] は、運動回転適応実験において内部順モデルが何に影響されて回転へ適応するのかを検証した。この運動回転適応実験は、視覚情報を与える誤差条件と視覚情報を与えず報酬情報のみをあたえる報酬条件の二つの班を作り実験に取り組んだ。彼らは、運動回転適応後に被験者からの聴取結果と提案した学習モデルの数値シミュレーション結果によって内部順モデルの適応がどう起こるかを考察した。その結果、彼らは視覚的に与える情報と脳内での運動予測の差である感覚予測誤差が内部順モデルが回転へ適応するのに重要な役割を示すと提案した。

1.1.2 先行研究は内部順モデルの適応状況を直接確認できていない可能性

一方で、被験者の聴取結果より上記の提案には、実験結果に対して誤差を含めてしまうとも考えられる。ヒトの認知結果は、実験時の条件や事前知識による思い込みによって実験結果に変化をもたらしてしまう可能性がある [8-12]。この現象は錯視等の特殊な環境だけでなく、経験により積み重ねてきた知識が原因となることもある。このため本研究では、行動実験の結果を確認するには認知処理を行わずに直接運動で確認する必要があると考える。

1.2 本研究の目的：従来と異なる方法で内部順モデルの適応を直接的に確認

本研究の目的は、実験における認知処理を排除できる行動実験を計画し、内部順モデルの回転への適応状況を運動により直接確認することで、内部順モデルがどのような情報に影響されて回転へ適応するのかを検討することである。内部順モデルの回転への適応状況を運動により直接確認することができれば、実験結果に誤差が生じることがなく結果を検証することができる。そこで、本研究では目標点跳躍課題を提案する事で、上記の問題を解消できると考える。目標点跳躍課題は、被験者へ視覚的な外乱による運動補正を確認することで内部順モデルの適応状況を確認できる実験として計画した。上記の条件の行動実験を計画し実施することで、内部順モデルが回転へ適応する際に重要となる要素を検証する。また、先行研究 [5] で提案されている学習モデルを基に数値シミュレーションを行い、その結果を本研究での行動実験と比較することで先行研究との結果の違いについても考察する。

1.3 本論文の構成

本論文は、全 6 章と付録にて構成される。第 1 章では、本研究で対象とする研究の背景や問題点を述べ、目的を述べる。本研究では目的を達成するために、目標点跳躍課題を提案した。まず、運動回転適応と内部順モデルの詳しい説明について次章で説明し、目標点跳躍課題についての説明はそれ以降の章で行う。

第 2 章では、まず本研究で取り扱う運動回転適応について詳細を述べる。ここではヒトの持つ学習・適応能力の 1 つである運動回転適応について、先行研究を踏まえてどのように研究され、ヒトの適応能力には一体どのような物が存在するのか述べる。その後、先行研究における問題点の提案および本研究で提案する新実験の目標点跳躍課題について説明する。ここでは先行研究では一体何が問題なのかを議論し、その問題点から実験において改善すべき点を考察し、それを達成するために目標点跳躍課題を提案する。また、目標点跳躍課題を含めた行動実験全体の方法について述べる。

第 3 章では、本研究で行う運動回転適応実験および目標点跳躍課題の数値シミュレーション方法について述べる。本研究で提案する目標点跳躍課題を数値シミュレーションで再現する。運動回転適応の数値シミュレーション結果は、先行研究 [5] で提案されている学習モデルを用いて示されている。この学習モデルは、カルマンフィルター器と強化学習器の二つのプロセスによって再現される。本研究ではまず、この学習モデルに基に運動回転適応の数値シミュレーション方法を述べる。さらに、先行研究の学習モデルを基に計画した目標点跳躍課題のシミュレーション方法を示した。

第 4 章では、先の章で述べた数値シミュレーションを実施し、先行研究の学習モデルではどのような結果を得られるのかを確認し、考察した。まず、運動回転適応実験の数値シミュレーションを行い、先行研究と同様の結果を得た。その後に目標点跳躍課題の数値シ

ミュレーションを行い，誤差条件では新目標点よりも右側へ報酬条件では新目標点へ到達する結果を得た．次に，実際にナイーブな被験者による行動実験ではどのような結果が出るのかを確認する．

第5章では，被験者実験の結果と考察を示す．被験者実験結果において，先行研究と同様に両条件共に回転へ適応できていた．目標点跳躍課題結果は，誤差条件では新目標点よりも左側へ報酬条件でも新目標点の右側へ到達する結果を得た．誤差条件においては，確認実験および数値シミュレーションの結果を説明できる結果を示した．しかし，報酬条件において数値シミュレーション結果と異なる結果を示している．この結果を踏まえ，報酬条件における学習モデルの働き方について考察し，内部順モデルがどう回転へ適応するかを考察する．

第6章では，結論として視覚的に与えられる感覚情報だけではなく，報酬情報のみを得る場合においてもが内部順モデルが回転へ適応する可能性が明らかとなった．感覚情報が内部順モデルに大きく影響する結果は，先行研究 [5] の結論と同様の結果である．しかし，報酬情報のみの場合においては，先行研究の学習モデルでは説明できない結果を得た．そこで，本研究では先行研究の学習モデルに補足し，何が内部順モデルの適応に影響するかを数値シミュレーションにより考察した．以上の結果を踏まえ，本研究の今後の課題と展望について述べる．

終わりに，付録をつける．ここでは，第3章において述べたカルマンフィルター器の導出方法に付いて詳細を述べると共に，被験者実験の結果（5章未掲載分）および本実験で行った確認実験の結果も述べる．

第2章 先行研究の問題点から新実験の提案

2.1 運動適応：ヒトは新奇の環境へもすばやく適応可能

我々の生活の中には、生きる上で学習や適応といった能力が必要となる場面が様々ある。ヒトが運動することによってその場や状態などを学習することを、運動学習 (Motor learning) と呼ぶ。例えば、自転車を乗れるようになることや水泳で平泳ぎができるようになること、または交通事故等により車のタイヤシャーシが歪んでしまっただけで走らない車を運転しているうちに難なく操作できるようになることなどがある (事故車は早急に修理か買い替えする事をお勧めする。) 前者二つは、新しい技術を獲得することに関係しており、今までできなかった事をできるようになる学習を行うので、運動学習の中でも運動技能学習やスキル学習 (Skill learning) などと言われている。対して後者は、もう既に獲得している技能に対して、新しい環境においても同様の技能を発揮出来るようにその環境へ適応していく能力であり、運動適応 (Motor adaptation) と言われている。しかし、このスキル学習と運動適応は、時折区別されずに使われる事がある。そこで、本研究ではこれらスキル学習と運動適応を以下のように定義する。

スキル学習：新しい技能 (運動制御) を獲得するための学習

運動適応：通常的环境でできることを新奇の環境でもできるようにする学習

スキル学習では、今までに経験した事のない運動や慣れていない運動をする事が多い [13]。また、運動適応では被験者の運動中に外力を与えて新たな環境へ適応させる実験と、視覚的な外乱による誤差 (自分の行う運動と観察される運動との誤差) に適応する実験の2種類に分かれる。本研究では、視覚変化における運動適応についてを取り扱うこととする。

2.1.1 運動回転適応

視覚的な外乱を与える事で生じる誤差へ適応する運動適応の一つに、運動回転適応 (Visuomotor adaptation) がある。先にも説明したが、これは視覚的な外乱により、自分が行っている運動と実際に観察される運動の間に生じる誤差から、脳内にある到達運動を生成するために必要な運動制御機構が回転に適応する事である。運動回転適応は、被験者に運動開始点から目標点までの運動をさせる到達運動と組み合わせて行われる。到達運動時に被

験者は、運動を行う実際の腕の位置を見えないように壁などで覆われ、代わりにディスプレイやスクリーンを用いて視覚情報を与えられた状態で、ロボットアームを握り目標点まで到達運動を行う。スクリーン上に与えられる視覚情報であるカーソルは被験者の手先位置と対応しており、腕を動かせばカーソルもそれに準じて動く。この到達運動で与える視覚情報に回転を付与し、被験者の意図している運動と実際に観測される運動の間に誤差を作り出すのが運動回転適応の一連の流れである。以下に、図 2.1 として運動回転適応についてまとめた図を示す。被験者は到達運動を繰り返していくうちに人為的に作られた誤差に適応していく。この運動回転適応は 1980 年代後半に提案されてから [6]、学習や記憶の研究によく用いられている。運動回転適応について、Mazzoni と Krakauer は被験者の回転への適応時においてここを狙えというような認知的な学習戦略を与えると、通常の学習能力に阻害された運動適応が起こる事を示し [7]、Braun らは運動回転適応において様々な回転を学習した方がそうではない場合に比べて、初期誤差も適応速度も効率の良い適応結果を示した [14]。また、西條と五味は付与される回転の大きさによる学習戦略の変化について考察し、突然大きな回転が加わる場合だと脳内の学習戦略が変化し、微小角度を加えて徐々に回転を付与する場合には到達運動時に用いられる運動制御機構が回転に適応している事を示し [15]、脳内で働く学習戦略の多様性についても研究されている。さらに、運動回転適応は回転付与ではなく、視覚情報の鏡面反転により新奇の学習戦略時にどのような反応を示しどう学習していくのか運動回転適応と比較されて研究されている [16, 17]。回転へ適応するような運動適応のみではなく、到達運動中に加えられる外力に適応する研究においても、適応能力について様々な議論が繰り返されている [18–20]。このように運動適応については非常に多くの知見が存在し、活発に研究されている。しかし、多くの研究では新たな学習戦略を見つけるといった研究が多く、なぜこのような運動適応が起こるのかといった運動回転適応の引き起こる根本に対しての研究は少ない。

Izawa と Shadmehr は、運動回転適応において被験者に与えられる視覚情報（文献内では感覚情報）が、運動制御機構の回転への適応に密接に関係すると述べた [5]。この研究では、被験者を、感覚情報と報酬情報を与える誤差条件と報酬情報のみを与える報酬条件の二班に分け、微小な回転を段階的に付与して運動回転適応させた。以下に図 2.2 として二班それぞれの模式図を示す。このとき、感覚情報は先に述べた様にスクリーン上に提供される視覚情報の事で、到達運動時の手先の軌跡が映像で常時与えられる。また報酬情報は音声情報であり、被験者が目標点への到達運動が成功した場合にのみ付与される。この二条件で運動回転適応実験をした場合に、両条件とも同一の回転へと適応可能であった。しかし、被験者へ到達点同定課題（Localization task）を行うと、二条件でそれぞれ違う結果が得られた。到達点同定課題では、最大回転（Izawa と Shadmehr らの実験では反時計回りに 8 [deg] の回転）に適応した後に、被験者に自分の手先の到達位置を到達運動していない方の腕（実験では左腕）で指し示すように指示された。以下に図 2.3 として到達点同定課題の模式図とその結果を示す。図中の結果より、両条件とも手先の描く軌跡は同一であるのにも関わらず、誤差条件では目標点付近に手先が到達したと示し、報酬条件では視認出来ていない実際の手先到達地点を示した。この結果から Izawa と Shadmehr

らは、同一の回転には適応するがその際にも用いる学習機構が異なるだろうと予測し、学習モデルを構築して数値シミュレーションを行った。以下に図 2.4 として先行研究 [5] の提案した学習モデルを示す。この学習モデルは、感覚情報によるカルマンフィルター器での誤差修正機構と、報酬情報による行動選択機構 (Action selection) の二つに分かれており、このモデルによる結果は運動回転適応実験の結果を再現する事ができた。その結果、脳内の運動制御機構が回転に適応するには、感覚情報と脳内での運動予測との差が必要であり、この誤差によって回転を学習すると提案された。この運動予測は、内部順モデル (Internal-Forward model) という脳内の機構が、実際にヒトが運動する前に運動指令を基にして自分がこれから行う運動を計算し予測することにより行っていると考えられている。Izawa と Shadmehr らの研究結果から、内部順モデルが運動回転適応に密接に関係しているという事が示された。

2.1.2 運動回転適応と内部順モデル

ヒトは、足の裏をくすぐられると非常にくすぐったくてむずむずするが、自分でくすぐると思ったよりくすぐったくない。この現象は内部順モデルにより説明がつく。内部順モデルは、運動指令によって次に自分がどう運動するかを予測する脳内の器官であり、脳機能イメージングによって小脳がその役割を担っているのではないかと研究がされている [3]。自分でくすぐるときには、自分の腕の運動というのは内部順モデルで予測している。よって、もとより自分のする運動がわかっているのに、何をされるかわからない他人からのくすぐりよりもくすぐったく感じない。この原理は論文 [21] でも説明されており、内部順モデルで生成された運動予測と、くすぐられるという感覚情報からの入力との差がくすぐったさになると説明している。この差は感覚予測誤差ともいい、この誤差により内部順モデルが回転へ適応すると提案されており [7]、Izawa と Shadmehr も感覚予測誤差が内部順モデルの適応に影響していると提案している。さらに、生理的な知見とモデルベースの結果を比較することで脳がどの部分がどう活動するのかという知見について議論されている [2, 3, 22]。図 2.4 で示したように、内部順モデルの運動予測結果と視覚的に与えられる感覚情報による感覚予測誤差 (Sensory Prediction Error) を利用して、カルマンフィルターが次の運動をどう変化させるかの学習を行う。これが提案された内部順モデルが回転に学習するメカニズムであり、実際にこのモデルによるシミュレーション結果は、行動実験結果を説明できる。先行研究 [5] は、表 2.1 の条件で運動回転適応実験を行った。この実験結果は誤差条件 (ERR Condition) と報酬条件 (RWD Condition) の二つの条件班において、横軸を試行回数にとり縦軸を到達角度とし、その試行で被験者が運動開始地点から見て何度かの場所に到達したのかを確認している。実験結果より、両条件共に回転へ適応しているが、報酬条件では視覚情報であるカーソルが与えられないため到達角度の分散が非常に大きい。数値シミュレーション結果でも両条件における実験結果を再現できている事がわかる。この結果から運動回転適応において、感覚予測誤差が内部順モデルの適応に大きく影響すると考えられる。

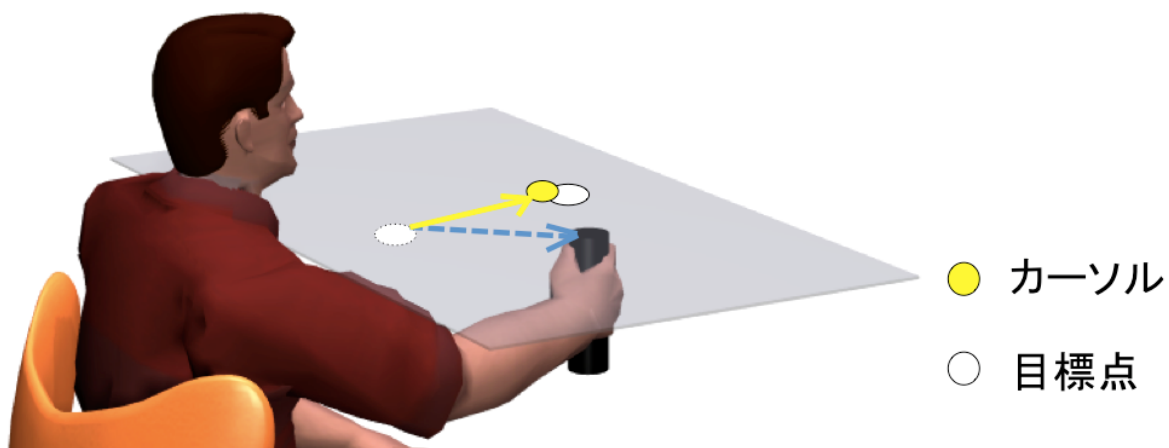
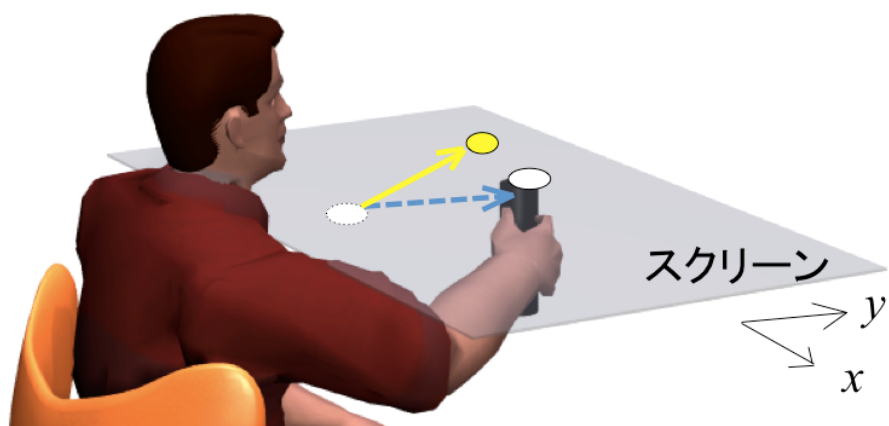
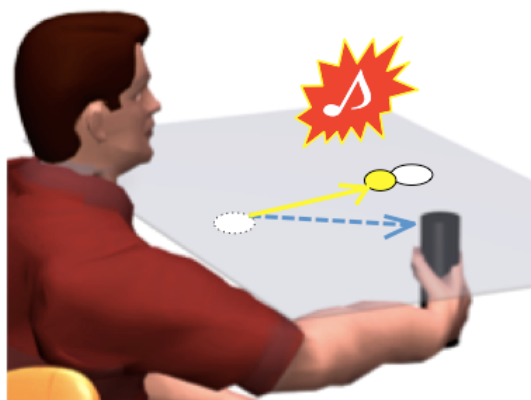


図 2.1 運動回転適応実験



視覚情報 + 音情報

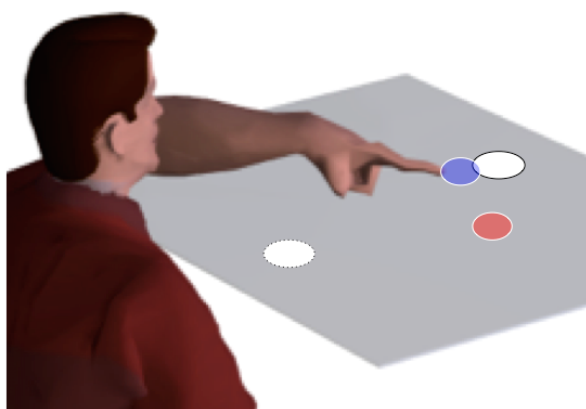
a. 誤差条件



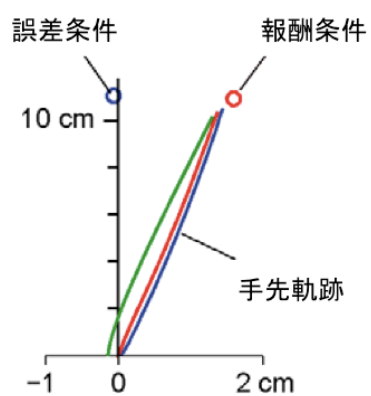
音情報のみ

b. 報酬条件

図 2.2 運動回転適応実験の条件 (Izawa et al., 2011.)



a. 実験方法



b. 実験結果

図 2.3 到達点同定課題 (Izawa et al., 2011.)

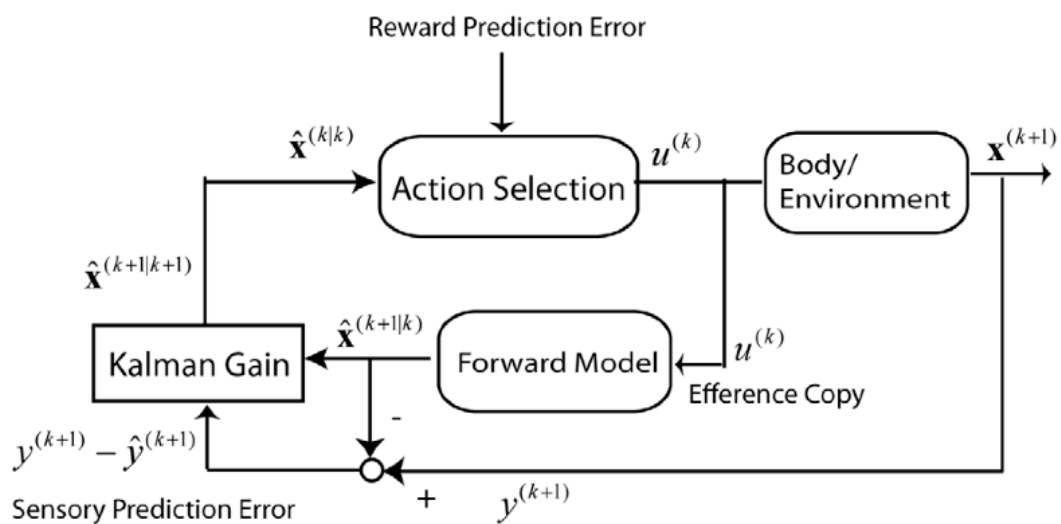


図 2.4 運動回転適応の学習モデル (Izawa et al., 2011.)

表 2.1 運動回転適応実験の条件 (Izawa et al., 2011.)

試行数	500 [trial]
到達距離	100 [mm]
回転方向	反時計回り
最大回転	8 [deg]
到達範囲	± 3 [deg]
回転付与	$+1$ [deg]/40 [trial]

2.2 先行研究の問題点：運動結果を直接観測できていない点

Izawa と Shadmehr の研究結果は、運動回転適応を起こす脳内の内部順モデルが、視覚的に与えられる感覚情報により影響を受けて回転を学習すると示した。これは、誤差条件と報酬条件の二条件で回転適応を行った後に、図 2.3 に示すような到達点同定課題により被験者の内部順モデルが回転に適応しているかを確認した結果によるものである。

もし、被験者が内部順モデルが回転へ適応しているならば、被験者は自分の手先が目標点へ到達していると感じるために図 2.3 の青点のように目標点位置を自分の手先到達位置だと指し示す。しかし、内部順モデルが回転へ適応していなければ、被験者は違和感を感じながらこちら辺を狙えば報酬が得られるという戦略を取るようになり、自分の手先位置を理解した上で到達運動を行っていると考えられる。よって自分の手先到達位置は図 2.3 での赤点のように実際の手先到達地点を示すはずである。実際に到達点同定課題の結果からは、感覚情報が与えられている誤差条件が目標点位置を指し示し、目標点へ到達できたか否かのみの報酬情報が与えられる報酬条件は実際に到達した手先位置を指し示した。この結果と提案された学習モデルによる数値シミュレーション結果によって、内部順モデルは視覚的な感覚情報が与えられることによって回転に適応することが示された。しかし、本研究では先行研究 [5] での内部順モデルの適応状況の確認方法に問題があると考えられる。

2.2.1 問題点 1：到達運動と到達点同定課題で用いる腕の不一致

本研究で先行研究の問題点の一つは、右腕での到達運動における回転への適応状況の確認に左手を用いたことであると考えられる。運動回転適応実験では、右手が利き手の被験者を集めて実験を行う。実験では右腕を運動開始点から目標点まで到達運動し、与えられた回転へ適応していく。しかし、この時の内部順モデルの適応状況を確認するために行われた到達点同定課題は、左手にて指し示すようを指示した。ここで、右腕で学習した結果を左手にて報告してよいのかという疑問が出る。各腕の学習プロセスが同じであって右腕で学習した事を左腕でも難なく使えれば良いが、片腕で学習した事がうまく他方の腕で使えるという事ではない。Nozaki らは、片腕運動と両腕運動では学習の過程が異なる事を示した [23]。この研究では、新奇の力場を左腕のみで学習する場合と両腕で学習する場合において、両手運動での学習は右腕と左腕の学習を単なる組み合わせでは無いことを示しており、それぞれ異なる学習プロセスが存在するだろうとしている。また、Yokoi らは右腕と左手での学習方法が異なり、右利き被験者の場合に左手の方が右手に比べて学習速度が速いことを述べている [24]。これらの研究と Izawa と Shadmehr らの先行研究とでは、実験方法（視覚的变化を与えるか新奇の力場を与えるか）の違いや、単に学習内容が右腕と左腕でどう変化するという関係を述べているわけではないこと等の相違点は存在するが、ここでは、各腕はそれぞれ異なる学習プロセスが存在し、片方の腕で学習した内容がもう片方の腕にそのまま受け継がれる訳ではない結果を示していることが重要である。

これらの結果より、異なる腕を経由して右腕の回転への適応状況を報告させることは、適応状況の結果に対して誤差や他の処理プロセスからの影響が存在すると考えられる。

2.2.2 問題点 2：知覚および認知処理系からの影響

2 つ目の問題点として，先行研究の確認方法では運動結果を確認したいのにも関わらず認知系の影響が存在する可能性がある．先行研究では，被験者は到達点同定課題において右手運動の回転へ適応後の結果を左手にて報告した．しかし，右腕の運動結果を左手で報告させるには，以下の手順を踏む必要がある．

1. 到達運動の軌道計画（軌道計画プロセス）
2. 運動開始（右腕運動）
3. 運動終了（右腕運動）
4. 運動終了地点確認および記憶（記憶プロセス）
5. 運動開始点へ腕が戻る（右腕運動：ロボットアームによる受動動作）
6. 到達点同定課題開始指示の確認（知覚・認知プロセス）
7. 到達地点の記憶想起（記憶プロセス）
8. 到達運動の軌道計画（軌道計画プロセス）
9. 運動開始（左腕運動）
10. 運動終了（左腕運動）

左手での報告は少なくとも，記憶・知覚および認知・想起・軌道計画の 4 つの脳内プロセスが関与することになる．自分の到達位置の記憶や想起だけではなく，到達点同定課題開始の指示により左手で報告しなくては行けないという認知処理が発生する．しかし，運動系の処理と認知系の処理での結果に違いが生じるという研究結果がある [8–10]．これらの研究では，エビングハウスの錯視と言う，図 2.5 に示すような有名な錯視を用いて実験を行い，被験者の静止時と動作時における認知系の働きが異なること示した [8]．被験者は通常，小さな円群に囲まれている円と大きな円群に囲まれている円では，小さな円群に囲まれている円の方がそうでない物に比べて大きく見える．しかし，実際は両者の円の大きさは同じ物であり，錯視が生じていることがわかる．ここで先行研究 [9] では，図 2.6 のように被験者には内円をつまむ様に指示をする．被験者の 2 本の指先にはセンサーが設置されており，つまんだ大きさによりその距離が計測出来るようになっている．このとき，不思議と小さな円群に囲まれている円と大きな円群に囲まれている円では，被験者がつまんだ距離が変化しなかった．つまり，知覚処理には周りの囲む円が大きく影響することで錯視が生じるが，運動処理にはそれが影響しないことがわかり，認知系プロセスと運動系プロセスではことなる視覚処理が行われていると示唆されている．また，Ganel らは上記の先行研究と同様の実験を行い（図 2.7 参照），錯視では運動系プロセスと認知系プロセスは乖離するような結果を得ることを示した [10]．

さらに、認知と運動の乖離は錯視だけにはとどまらず、物体の把持運動にも影響することが知られている [11,12]。Flanagan と Beltzner は、大きさの異なる重さが同一な二つの箱を用意し、被験者にそれらを持ち上げるように指示した [11]。そのとき、被験者は同一の重さの箱であっても、大きな箱よりも小さな箱の方が重く感じてしまう。これを、Size-Weight illusion という（実験内容は図 2.8 を参照）。この実験では、被験者が二つの物体を把持するときの把持力と負荷力を計測し、内部モデルにより物体の外見から重さを予測して物体を持ち上げていることがわかった。

これらの結果から、ヒトの行動は認知系の処理に大きく影響されることがあり、錯視のような特殊な状態でなくとも認知系と運動系の処理は、同一パラダイム内で行わない方がよいことがわかる。よって本研究では、先行研究 [5] で行われてた到達点同定課題には認知系の影響が介在する可能性があると考え、これを排除する必要があると考える。

2.2.3 問題点 3：報告による運動結果の確認

三つ目の問題点に、被験者からの報告によって運動の予測をする内部順モデルの適応状況を確認する点がある。先行研究 [5] は、到達運動後に自分の腕がどこへ到達したかを報告させた。この実験では、内部順モデルが適応しているか適応していないかで指し示す位置が異なるだろうという予測のもとに行われ、図 2.3 のように実際に実験結果では指し示す位置が異なった。これは、内部順モデルによる手先予測位置が異なることで起こる結果である。しかし、到達点同定課題では、到達運動により回転に学習している被験者に対して運動終了後に実験の主となる運動とは異なる方法で内部順モデルの回転への適応状況を確認している。内部順モデルは先にも述べているように、運動指令を用いて運動の予測を行う機構である。もし、回転適応時の内部順モデルの状況を確認するのなら、被験者の到達運動中に行うべきではないかと考えられる。そこで、本研究では、内部順モデルの予測結果を見るには到達運動によって直接的に確認すべきであると考えられる。

2.3 新実験の提案：運動結果によって内部順モデルの適応状況を直接確認

先の章で、先行研究内で内部順モデルの適応状況の確認方法に幾つか問題があると述べた。一つ目に、右手で到達運動して回転に適応しているのにも関わらず、左手にて適応状況を確認していること。二つ目に、被験者が左手にて報告することにより、認知系の処理が関与する可能性があること。三つ目に、内部順モデルの適応状況を、到達点同定課題では適応した運動ではない方法で確認していること。本研究では、先行研究 [5] で考えられるこの三つの問題点を解消する実験を提案する。この三つの問題点は、内部順モデルの適応状況の確認において運動回転適応と独立した到達点同定課題を行うことが問題であることを述べている（図 2.9 を参照）。よって、本研究では到達点同定課題のように到達運

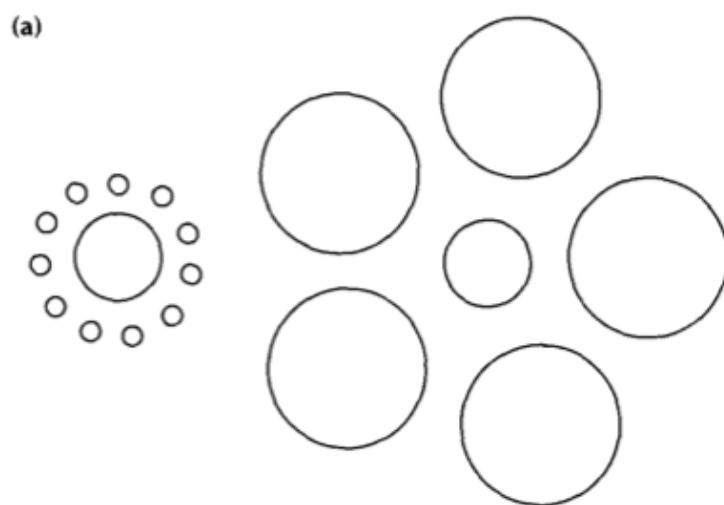


図 2.5 エビングハウス錯視 (Aglioti et al., 1995.)

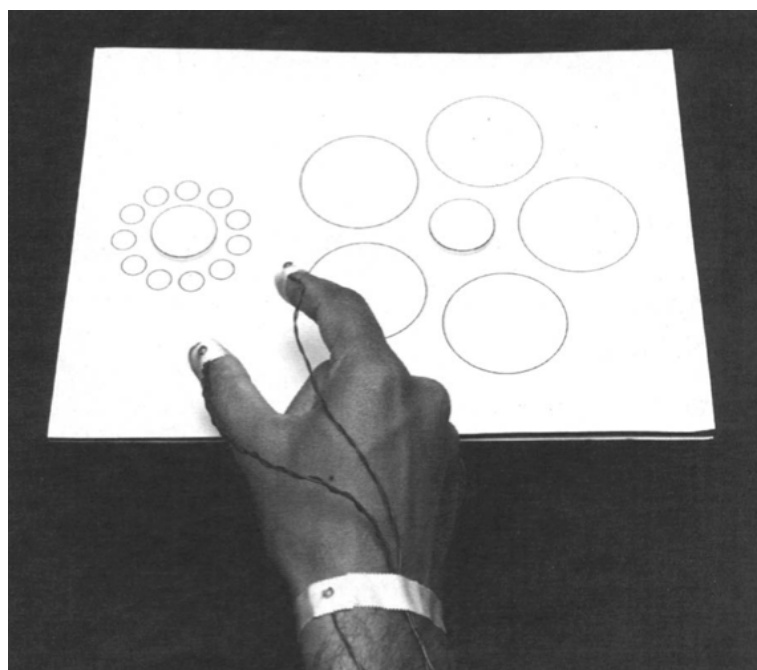


図 2.6 把握実験：エビングハウス錯視 (Aglioti et al., 1995.)

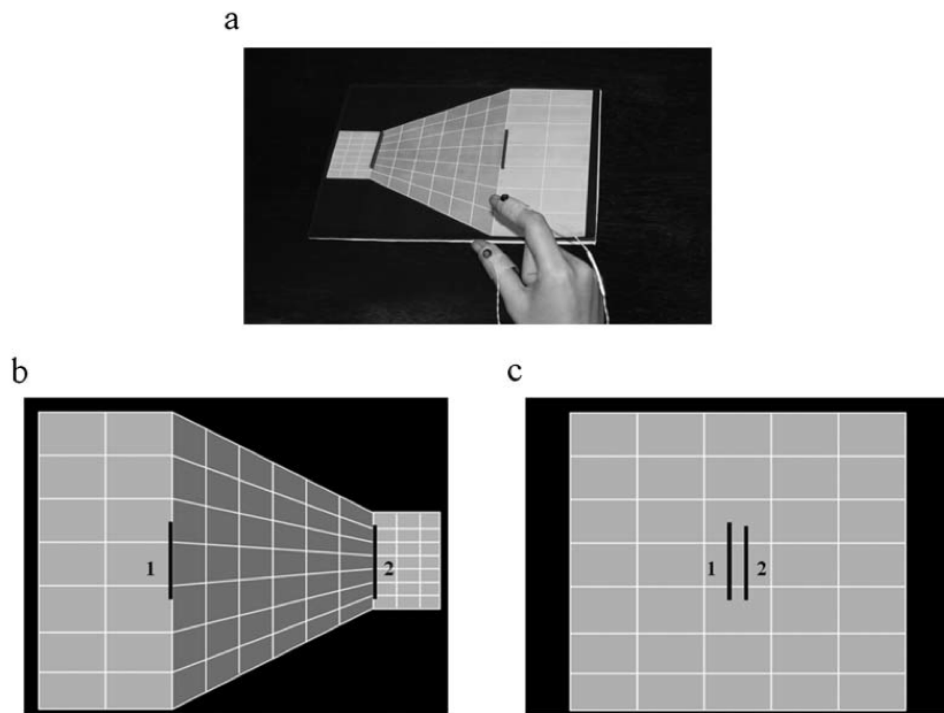


図 2.7 把握実験：奥行き錯視 (Ganel et al., 2008.)

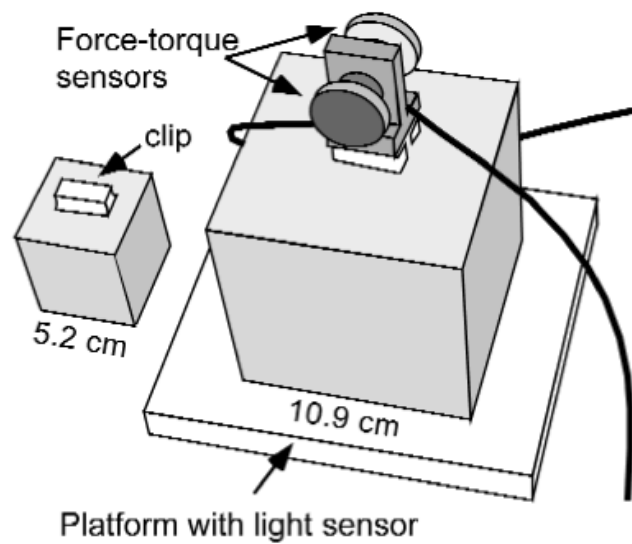


図 2.8 Size-Weight illusion (Flanagan et al., 2000.)

転と独立した課題を行わずに、適応後の到達運動内で内部順モデルの状況を確認する方法を提案する（図 2.10 を参照）。

2.3.1 提案行動実験：目標点跳躍課題

本研究は、先行研究の問題点を克服した上で、運動回転適応における内部順モデルの回転への適応状況を確認する方法として、目標点跳躍課題を提案する。この課題は、運動回転適応実験で与えられた回転に適応した後の到達運動中に目標点が突然飛ぶという課題である。以下に、図 2.11 として目標点跳躍課題の模式図を示す。被験者は、反時計回り方向（12 時の方向から左方向）に与えられた回転に適応しているので、誤差を修正するために時計回り（12 時の方向から右方向）に向けて到達運動するようになる。このとき、感覚情報が与えられる誤差条件と報酬情報が与えられる報酬条件では与えられる回転に適応できるが、内部順モデルが回転に適応しているのは誤差条件のみであることが提案されている [5]。この結果より、誤差条件と報酬条件では内部順モデルに適応しているかしていないかで分けられるので、内部順モデルの働きによる課題を到達運動中に再現することで、その課題終盤の運動結果に違いが現れると考える。

先行研究 [26,27] は、到達運動中に目標点が跳躍（任意の方向への突然移動）するという運動中で外乱を与えることで、今まで行っていた課題中に新たな課題を導入している例の一つである。また、Miall らは被験者の運動中に小脳へ経頭蓋磁気刺激（TMS）を打つことで到達運動を阻害し、小脳における腕の位置を予測している内部モデルの存在やリアルタイムでの運動補正は自らの腕の予測位置と目標点の位置の差を利用していることを述べている [25]。これらの先行研究の結果より、到達運動中に外乱を加えれば、内部順モデルによる予測手先位置と新しい目標点の差を認識して運動補正を行うため、内部順モデルの予測位置が異なる誤差条件と報酬条件では生成されるそれぞれの運動補正量は大きく異なるであろうと考える。そこで本研究では、到達運動中に内部順モデルの適応状況により異なる運動補正が得られる目標点の突然跳躍という外乱を利用し、誤差条件および報酬条件において被験者の到達運動中に内部順モデルの適応状況を確認しようと考えた。以下に、図 2.12 として、先行研究 [5] の結果を基にした誤差条件と報酬条件で生成すると予測される運動補正を示す。先行研究 [5] では、誤差条件は内部順モデルは回転へ適応し、報酬条件は内部順モデルは回転へ適応しないが戦略を構築することで回転が加わった状態で目標点までの到達運動を可能としていると述べている。運動途中での跳躍による運動補正は、内部順モデルの適応状況の違いにより大きく異なる（図 2.12 左端）。この結果より、目標点が右側に跳躍する場合における誤差条件では、目標点の跳躍により元々計画されていたの到達点より右へ到達すると考える（図 2.12 誤差条件を参照）。これは、誤差条件の到達運動中の被験者は内部順モデルが回転に適応しているため、自分の手先位置をカーソルの延長線上にある赤点のように予測し、運動補正は赤点位置と新目標点の差となり、その差を修正するように運動補正するためである。しかし、目標点が右側に跳躍する場合における報酬条件では、目標点の跳躍により新しい目標点へ運動補正を行うと考える（図

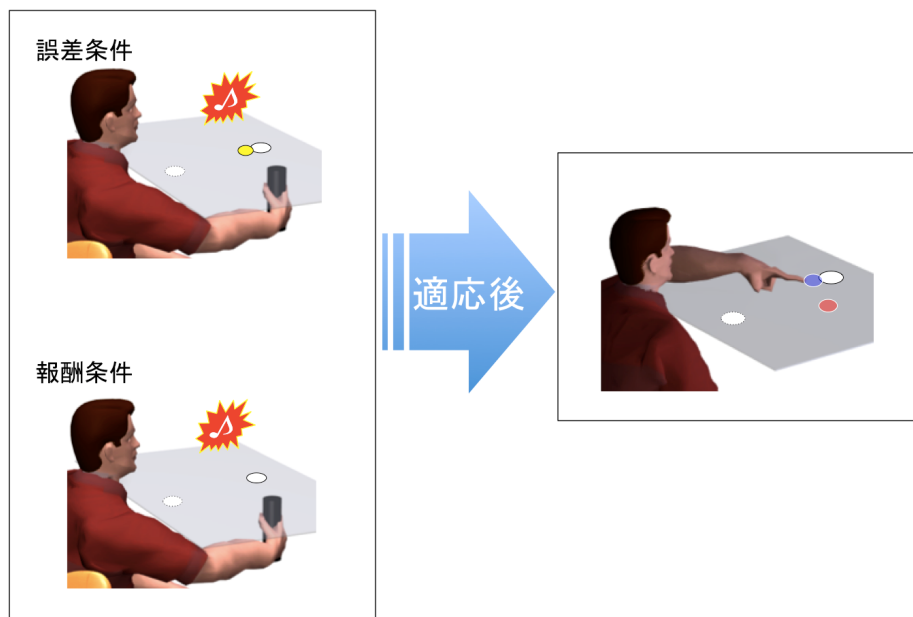


図 2.9 先行研究での問題点について

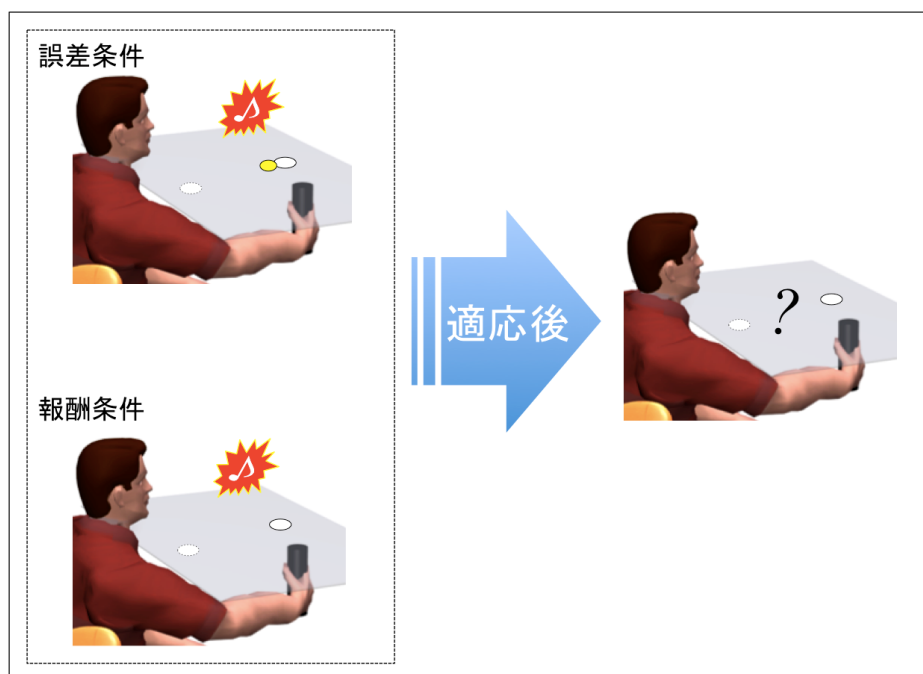


図 2.10 本研究における問題点の改善策について

2.12 報酬条件を参照)。これは、報酬条件での被験者は内部順モデルが回転に適応していないので自分の手先位置を理解しており、ここからの運動補正は本当の手先位置（到達運動中は視認不可）の位置と新目標点の差を修正すると考えられるためである。

この目標点跳躍課題の予測結果は先行研究で提案されている回転学習の学習モデルを基準にし予測しているため、本研究では行動実験を行うと共に、その学習モデルを基に数値シミュレーションを行う本研究では、この予測結果と行動実験における目標点跳躍課題の結果を比較する。

2.3.2 行動実験方法

本章では、先に述べた予測結果を被験者の行動実験によって検証する。本行動実験は、運動回転適応実験に加えて目標点跳躍課題を行う。図 2.12 に示す目標点跳躍課題は、最大回転に適応した到達運動中に目標点が跳躍し、内部順モデルの適応状況を到達運動の運動補正で確認する課題である。従って、この課題は被験者が回転へ適応した後の到達運動で行う必要がある。さらに、この目標点の跳躍に適応しては問題であるので、ランダムに出現する必要性がある。そこで、本研究ではこの課題を誤差条件と報酬条件の運動回転適応実験内で回転に適応していない状態と回転に適応した後の2つのセッションでランダムに発生させることとする。目標点跳躍課題の発生頻度は5回に1回の割合として、セッション中であればランダムに出現する。セッション中において目標点跳躍課題が発生しないときは、通常の運動回転適応実験の到達運動を行う。適応していない状況で目標点跳躍課題を行うのは、適応後の運動補正と比べることで適応による効果が現れているかを確認するためである。以下に、図 2.13 として本研究における行動実験の流れを示す。この流れで、運動回転適応実験を行い、目標点跳躍課題において内部順モデルの適応状況の確認を行う。また以下に、図 2.14 として運動回転適応実験における試行毎の到達運動課題の詳細を示す。到達運動課題時において被験者には、目標点が出現したときに到達運動を開始するように指示した。なので、被験者は試行毎に運動開始点にカーソルを置き、目標点の表示が出るまで待つ。このとき目標点が表示されるまでの時間は、50 ~ 200 [ms] の時間幅でランダムで設定している。目標点が表示された時、被験者には到達運動を始め目標点付近で運動を終了するように指示した。また、到達運動での速度は常に一定の結果を出すようにしたかったため、基準速度を上回るもしくは下回る場合において警告を出すように設定した。運動の奥行き方向のピーク速度を 500 [mm/s] を基準として設定し、 $\pm 10\%$ の誤差を許容するようにした。被験者には、運動終了とともに被験者自ら運動開始点まで腕を戻すことをさせず、運動終了時にその地点で静止するように指示した。被験者は自ら運動開始地点に戻れない代わりに、ロボットアームによって元の位置へ自動的に戻してもらえる。これは、運動開始点まで戻す運動によりその運動を学習させないようにするためである [19]。この流れで1試行が終了となり、次の試行が開始となる。なお、報酬条件の到達運動課題において図 2.14 のカーソルは常に与えられない。

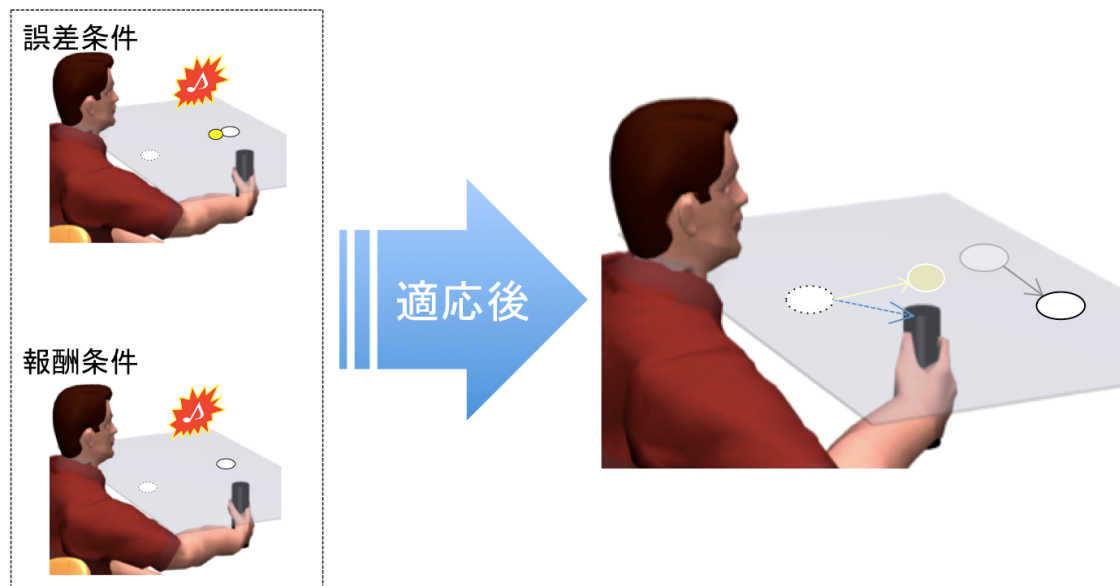


図 2.11 目標点跳躍課題

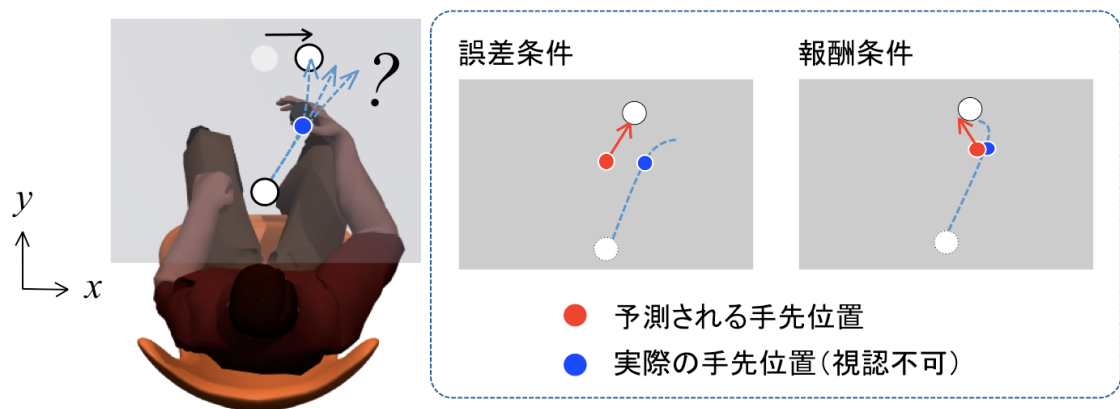


図 2.12 目標点跳躍課題の運動予測結果

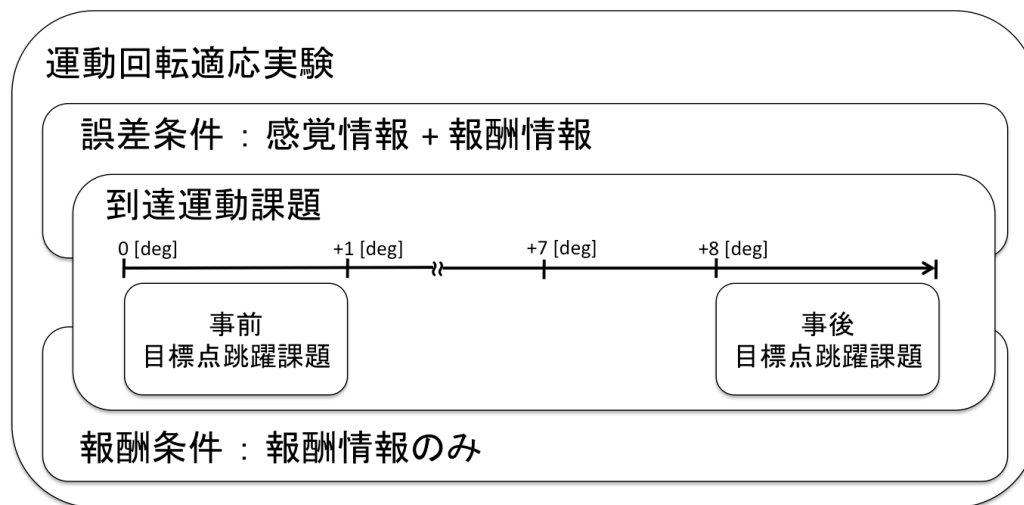


図 2.13 本研究で提案する行動実験の流れ

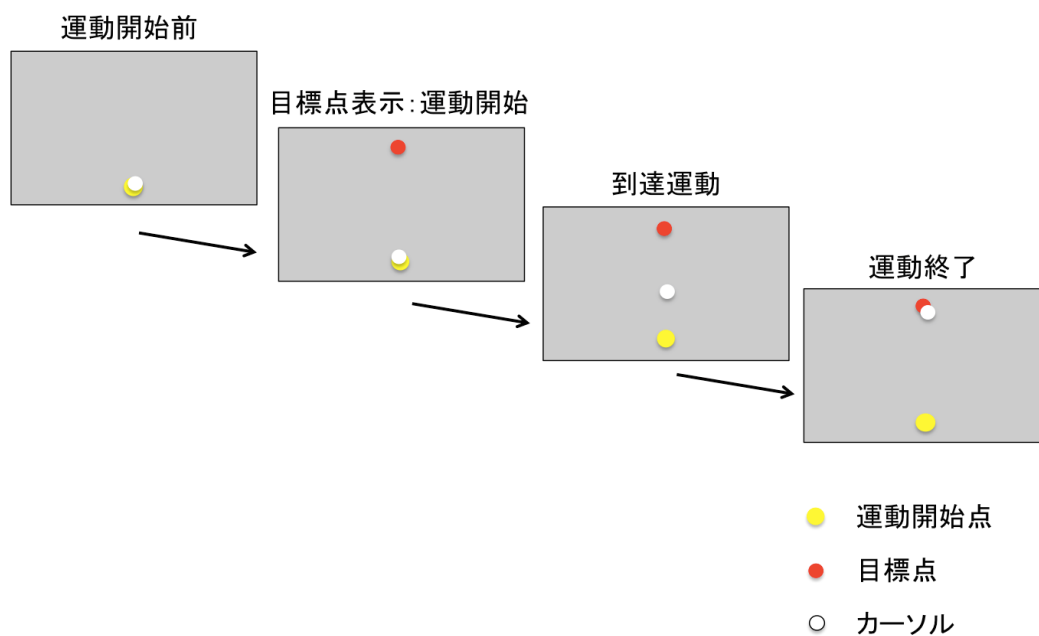


図 2.14 本研究での運動回転適応実験における到達運動課題

第3章 数値シミュレーション：二つの学習プロセスで運動回転適応を再現

本研究では，提案された学習モデルを基に数値シミュレーションを行い，予想した結果を得られるか確認する．また，図 2.4 において提案された学習モデルを示しているが，本研究ではこの学習モデルを再定義する．この理由は，後に述べる運動指令の計算式が先行研究の図では再現出来ないと考えたためである．以下に図 3.1 として，先行研究 [5] を基とした運動回転適応の学習モデルを示す．これは，先行研究 [5] ではカルマンフィルター器と強化学習による行動選択器（Action selection）の二つの器官によって回転を学習していると提案していたが，図 2.4 の学習モデルではカルマンフィルターで誤差を学習した後に行動選択器器が機能している．図 2.4 は各学習器は独立していない構成であるので，図 3.1 のように各学習器が独立に機能するよう再定義する．

3.1 運動回転適応実験のシミュレーション方法

まず，運動回転適応実験のシミュレーション方法を示す．ここで，試行ステップ数は k とし，式の上部に (k) のように記載している．例えば (k) は， k 試行目の結果を示している．この 1 試行では，1 回の到達運動の結果を表す．運動回転適応実験での結果は手先到達位置であるので，数値シミュレーションでは手先到達位置を計算する必要がある．本研究での運動回転適応のシミュレーションは離散系システムとして組み立てられており，1 試行毎に脳内で生成される運動指令を入力として手先の到達結果を出力する．以下に (3.1) 式として手先位置の計算式を示す．

$$h^{(k)} = u^{(k)} + n_h^{(k)} \quad (3.1)$$

ここで， h は手先位置， u は運動指令，そして n_h はノイズであり $n_h \sim N(0, \sigma_h^2)$ である．しかし，運動回転適応実験の被験者は自分の手先位置を視認することができず，代わりに手先位置と対応するカーソルが与えられる．そのカーソルには段階的に回転が加えられ，被験者はその回転に適応していく．以下に (3.2) 式としてカーソル位置 y の計算式を示す．

$$y^{(k)} = h^{(k)} + p^{(k)} + n_y^{(k)} \quad (3.2)$$

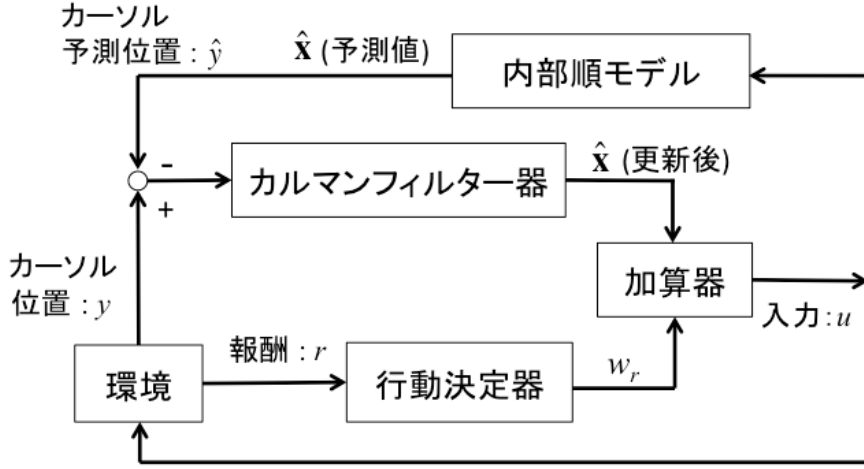


図 3.1 本研究での運動回転適応の学習モデル

ここで, y はカーソル位置, p はカーソルに加えられる摂動 (回転), そして n_y はノイズであり $n_y \sim N(0, \sigma_y^2)$ である. 被験者に与えられる感覚情報はこのカーソル位置であり, この感覚情報と被験者の内部順モデルが予想した予測値の差で回転に適応していく. 内部順モデルは, 運動指令の遠心性コピーを入力として自分の運動予測を行う. この運動予測というのは到達運動で言うと到達位置を示す. (3.2) 式で示したように到達位置を表現するには, 手先の位置と回転摂動が必要となる. つまり内部順モデルは, 運動指令を入力にして手先位置と回転摂動を予測するという出力をしていることとなり, 以下に (3.3) 式と (3.4) 式として内部順モデルの予測する手先位置 $\hat{h}$ と回転摂動 $\hat{p}$ の計算式を示す. ここで文字上部についている $\hat{\cdot}$ は予測値のことを示しており, ここでは脳内の内部順モデルが予測する値のことを指している. 例えば, p は現実にと与えられている回転摂動の値であるが, $\hat{p}$ と記入してあればこれは脳内で予測している回転摂動の値である.

$$\hat{h}^{(k+1)} = \hat{p}^{(k)} + u^{(k)} \quad (3.3)$$

$$\hat{p}^{(k+1)} = a\hat{p}^{(k)} + n_p^{(k)} \quad (3.4)$$

ここで, a は係数をそして n_p はノイズを示し, $n_p \sim N(0, \sigma_p^2)$ である. このノイズによって, 予測される回転摂動を更新していく. また, 試行数が $k+1$ となっているが, これは現試行の出力結果が前試行時の結果を用いて算出されていることを示している.

本研究では被験者の内部順モデルはこの二式を予測するように構成を定義し, 状態方程式 (State Space Model) によって定義すると,

$$\hat{\mathbf{x}}^{(k+1)} = A\hat{\mathbf{x}}^{(k)} + \mathbf{b}u^{(k)} \quad (3.5)$$

$$\begin{aligned}\hat{\mathbf{x}}^{(k)} &= [\hat{p}^{(k)} \quad \hat{h}^{(k)}]^T \\ A &= \begin{bmatrix} a & 0 \\ 1 & 0 \end{bmatrix} \\ \mathbf{b} &= \begin{bmatrix} 0 & 1 \end{bmatrix}^T\end{aligned}$$

となる．被験者は，カーソルが手先の運動と対応して動くことを聞かされてはいるが，カーソルに回転がかかることは知らない．なので，被験者にカーソルによって与えられる手先位置は被験者の予測した手先位置と同じということになり，

$$\hat{y}^{(k)} = C\hat{\mathbf{x}}^{(k)} + n_y^{(k)} \quad (3.6)$$

$$C = \begin{bmatrix} 0 & 1 \end{bmatrix}$$

と定義できる．試行毎の手先到達位置は，内部順モデルの予測が必要となる．この内部順モデルはカルマンフィルターによって，観測値と予測値の差から回転を学習する．以下に，(3.7) 式としてカルマンフィルターによる内部順モデルの学習を表現する計算式を示す．

$$\hat{\mathbf{x}}^{(k|k)} = \hat{\mathbf{x}}^{(k|k-1)} + K^{(k)}(y^{(k)} - C\hat{\mathbf{x}}^{(k|k-1)}) \quad (3.7)$$

$$\begin{aligned}K^{(k)} &= P^{(k|k-1)}C^T(CP^{(k|k-1)}C^T + \sigma_y^2)^{-1} \\ P^{(k|k)} &= (I - K^{(k)}C)P^{(k|k-1)}\end{aligned}$$

ここで， K はカルマンゲインで誤差がどれくらい学習へ影響するかを調整する係数であり， P は状態方程式の共分散で示す不確実性（Uncertainty）である．カルマンゲインは，感覚予測誤差： $(y^{(k)} - C\hat{\mathbf{x}}^{(k|k-1)})$ に影響しており，この誤差の大きさとカルマンゲインにより学習の度合いを調整する．感覚予測誤差は， $(y^{(k)} - \hat{y}^{(k)})$ で示されるが，(3.6) 式に示すように $(y^{(k)} - C\hat{\mathbf{x}}^{(k|k-1)})$ となる．不確実性はこのデータがどれくらい信頼出来るかという物であり，より信頼出来るデータを用いて学習する．上付の $(k|k)$ は，左側が現試行回数を表しており右側が使うデータの試行回数表現している．例えば $(k|k-1)$ の場合は (k) 試行目の到達位置を $(k-1)$ 試行目のデータを用いて計算していることを表している．このカルマンフィルターの導出については，付録にて詳細を記載する．

(3.7) 式の結果により回転摂動の予測値が計算でき，この予測値により運動指令が生成される．以下に，(3.8) 式として運動指令の詳細を示す．

$$u^{(k)} = -\hat{p}^{(k)} + w_r^{(k)} + n_u^{(k)} \quad (3.8)$$

$\hat{p}$ は学習する回転摂動であり，加わっている摂動の学習値なのでその回転を打ち消すように運動を行う必要があるので，負の方向に加わるようになっている． n_u は運動指令に加わるノイズであり， w_r は強化学習による結果を示している．報酬条件は強化学習によりこの時取るべき行動を選択する．その結果の値が， w_r である．強化学習とは，環境から得られる報酬を最大とするような評価関数を持つことで学習を行う．報酬を最大化するためにも，現在の状態がどのくらい良い状態なのかと計る必要があり，これを価値関数という．以下に，(3.9) 式として本研究での強化学習に用いる価値観数を示す．

$$V_k = E[r_k + \gamma r_{k+1} + \gamma^2 r_{k+2} + \gamma^3 r_{k+3} + \cdots + \gamma^{N-k} r_N] \quad (3.9)$$

r_k はその試行での報酬値であり， γ は遠い将来に得られる報酬ほど割り引く割引率という．この価値関数に従い，強化学習では報酬を最大化するように働く．この時得られる報酬と脳内で生成される予測報酬の差が報酬予測誤差となり，以下に示すような形となる．

$$\delta_k = r_k + \gamma \hat{V}_{k+1} - \hat{V}_k \quad (3.10)$$

先行研究 [5] では， $\hat{V}_k = w_v$ と提案しており，さらに割引率は $\gamma = 0$ として前回の報酬値を使わないようにしており，以下のような計算式を定義している．

$$\delta_k = r_k - w_v^{(k)} \quad (3.11)$$

この， w_v は強化学習における方策であり，以下に詳細を示す．

$$w_v^{(k+1)} = w_v^{(k)} + \alpha_v \delta_k \quad (3.12)$$

$$w_r^{(k+1)} = w_r^{(k)} + \alpha_v \delta_k n_u \quad (3.13)$$

(3.13) 式は報酬値を示しており，(3.9) 式の運動指令生成時に直接用いられる．(3.13) 式を計算するためには，(3.11) 式および (3.12) 式と外界からの報酬が必要であり，報酬予

測誤差が運動指令の生成に影響するようになっている．なお，外界からの報酬は以下の様に先行研究で定義されている．

$$\begin{aligned} r_k &= 1 - \beta u^2 & (c^{(k+1)} \in \text{goal area}) \\ r_k &= -\beta u^2 & (c^{(k+1)} \notin \text{goal area}) \end{aligned} \quad (3.14)$$

これらの式において運動回転適応実験の数値シミュレーションを行う．

3.2 目標点到達課題のシミュレーション方法

先に，運動回転適応の数値シミュレーション方法を提案した．目標点跳躍課題は運動回転適応実験の最中に行うので，このシミュレーション方法を用いて目標点跳躍課題の数値シミュレーションを行う．しかし，運動回転適応実験のシミュレーションは離散系で構築されており，試行毎の到達位置結果しか出力しない．目標点到達課題は到達運動途中に外乱を与えてその運動補正を確認しているので，連続的な計算過程が必要である．つまり，離散系の運動回転適応実験のシミュレーション方法を用いることでは連続系の目標点跳躍課題の計算はできないということである．そこで，連続的に検査しなくてはならない目標点跳躍課題を離散系に落とし込む方法を考える．

先行研究 [25] では，リアルタイムの運動補正は目標点位置と内部順モデルでの予測手先位置の差によって行われると提案している．よって運動補正に必要なのは，新しい目標点位置と現在の内部順モデルでは自分の手先位置をどう予測しているかという二点のみとなる．なので目標点跳躍課題の数値シミュレーションは運動回転適応実験の数値シミュレーション方法を用いて，新目標点の位置と現在の内部順モデルから計算される予測する手先位置を入力に，出力結果に運動補正後の手先到達位置を出力することで離散系での計算が可能である．以下に，(3.15) 式において運動補正を行う場合の運動指令 $u_{\text{correction}}$ の計算式を示す．

$$u_{\text{correction}}^{(k)} = T_{\text{new}}^{(k)} - \hat{h}^{(k)} \quad (3.15)$$

ここで， T_{new} は跳躍後の目標点位置であり， $\hat{h}$ は内部順モデルが予測する手先位置である．この計算式によって，目標点跳躍課題の結果をシミュレーションにて計算する．

第4章 数値シミュレーション結果：誤差条件は右へ報酬条件は新目標点へ到達

本章では先のシミュレーション方法から，図 2.12 の結果が再現できるか確認する．本シミュレーションは，MathWorks の MATLAB R2015a を利用してプログラムを製作した．

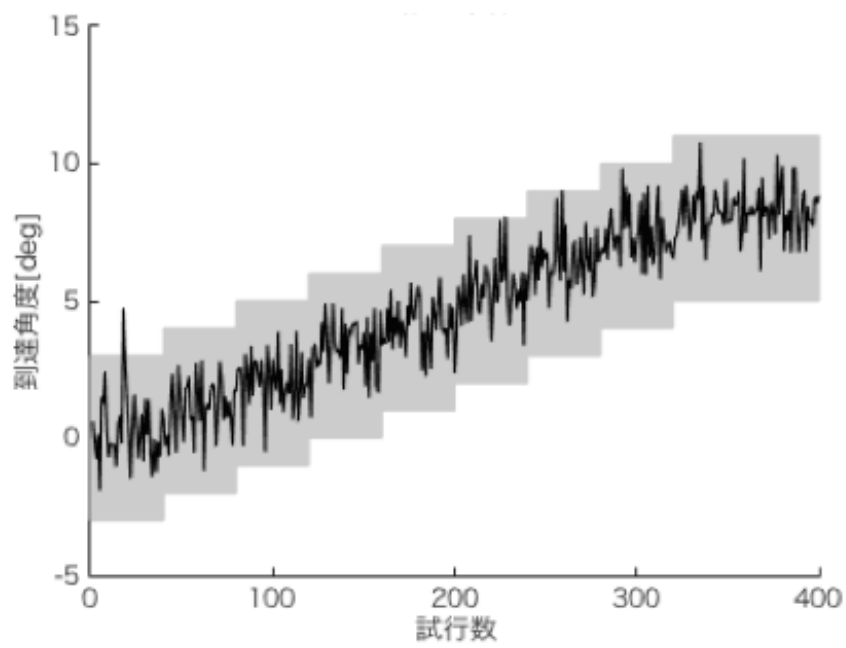
4.1 運動回転適応実験の数値シミュレーション結果

以下に運動回転適応実験の数値シミュレーション条件を述べる．本研究のシミュレーションの条件は，先行研究 [5] のシミュレーション条件に準ずる．また，以下に図 4.1 として表 4.1 の条件の基でのシミュレーション結果を述べる．シミュレーション結果で示すのは学習曲線であり，段階的に与えられる回転に対しての学習が確認できる．横軸には試行数，縦軸には運動終了時に到達していた位置を角度にて表した到達角度を示している．到達角度は， y 軸と運動開始地点から手先到達位置の直線がなす角度である．この到達角度が与えられた回転と同じ値へ遷移していけば，回転へ適応したことが確認できる．

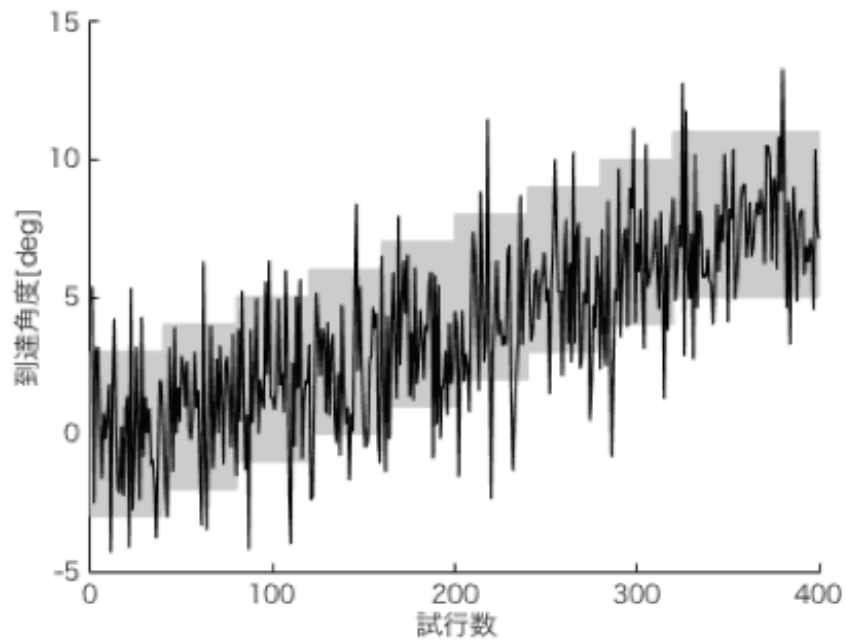
学習曲線は，a. に誤差条件，b. に報酬条件の結果を示している．誤差条件は非常に良い適応をしてしており，報酬条件は自分の手先位置の視覚情報が与えられない状況で回転に適応するので到達角度の分散が高い結果が得られている．以上の結果は，先行研究の実験結果およびシミュレーション結果と同様な結果を得られており，本研究で用いる学習モデルはヒトの運動回転適応学習を再現できていると考える．

表 4.1 運動回転適応実験の数値シミュレーション条件

試行数	400 [試行]
到達距離	100 [mm]
回転方向	反時計回り
最大回転	8 [deg]
到達範囲	± 3 [deg]
回転付与	+1 [deg]/40 [試行]



a. 誤差条件



b. 報酬条件

図 4.1 数値シミュレーションによる運動回転適応実験結果

表 4.2 数値シミュレーションによる目標点跳躍課題の条件

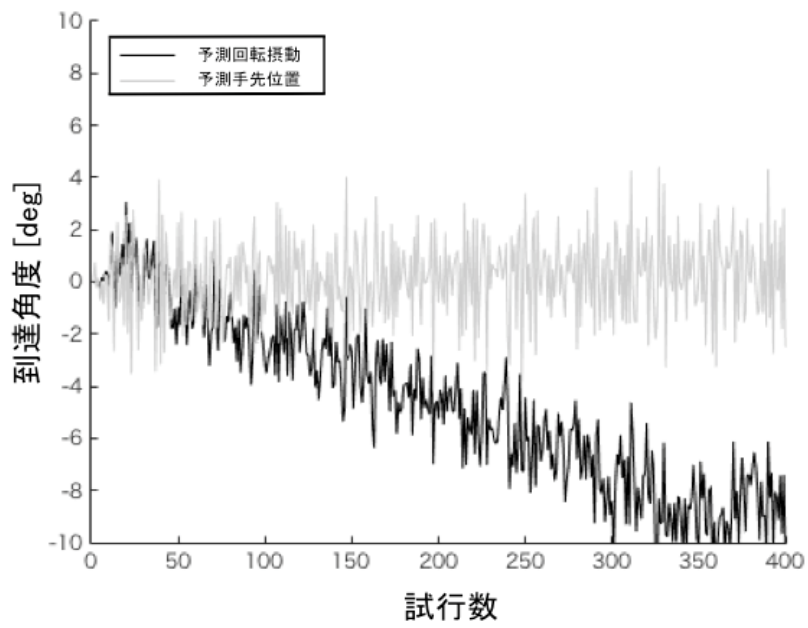
	条件 1	条件 2
跳躍方向	右方向	右方向
跳躍距離	7 [mm]	15 [mm]

以下に，図 4.2 として数値シミュレーションにおける内部順モデルの予測結果を示す．(3.5) 式でも示したように，内部順モデルの予測結果は手先位置を予測し，適応により回転摂動を変動させていく．ここで左図に誤差条件の，右図に報酬条件の結果を示し，縦軸は到達角度であり横軸は試行数である．図中に黒線で示す値が内部順モデルが適応する回転摂動の値であり，灰色の線がその試行のときに生成される運動指定から予測する手先の到達角度である．図 4.2 a. における誤差条件の結果は，予測する手先の位置は平均して 0 [deg] を示しており，常に腕の到達角度は直線的に運動していると予測する結果を示す．黒線に示す回転摂動は，段階的に与えられる回転を打ち消すように回転へ適応する結果を示している．また，報酬条件では感覚情報が与えられないので黒線で示す回転適応は起こらない結果を示し，灰色線は報酬値 w_r による行動選択器が回転を選択するために被験者自らの手先を右側へ到達させる実験時の内観を再現できている．これらの結果は先行研究 [5] と同様の結果を示しており，内部順モデルは感覚予測誤差で回転へ適応する事が示されている．

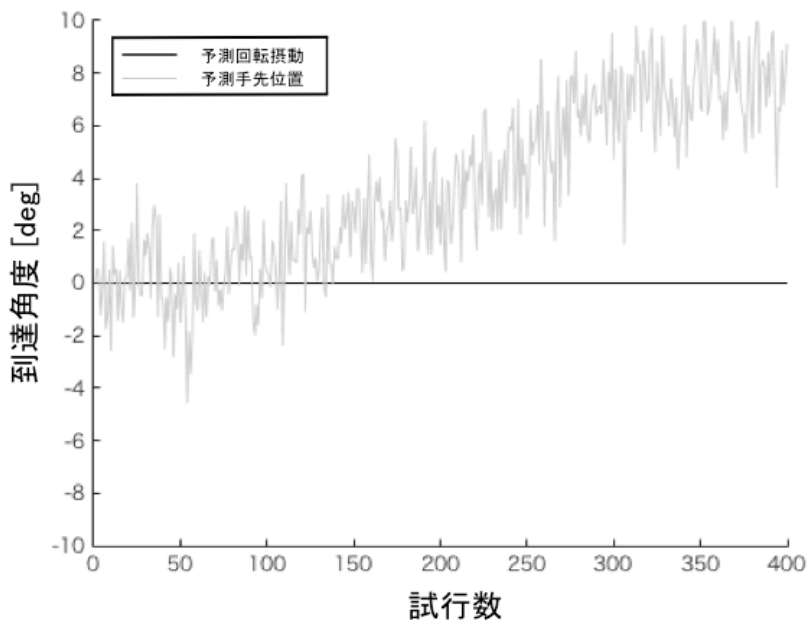
4.2 目標点到達課題の数値シミュレーション結果

先章に述べた目標点跳躍課題の方法で，数値シミュレーションを行う．本研究では目標点跳躍課題において，誤差条件は新目標点より右側へ到達して報酬条件では新目標点へ到達するという結果を予測している．この予測結果が先行研究を基とした学習モデルから出力されるのかを確認する．ここで，表 4.2 において目標点跳躍課題のシミュレーション条件を以下に示す．目標点の跳躍距離は右方向へ 7 [mm] と 15 [mm] の二種類とした．これは，シミュレーション条件によって異なる運動を生成できるかを確認するためである．以下に，図 4.3 として目標点跳躍課題の数値シミュレーション結果を示す．本研究の到達運動において横方向は x 軸，奥行き方向が y 軸としているので，図 4.3 の 縦軸 (y 軸) と横軸 (x 軸) はそれに対応している．

図 4.3 a. は 7 [mm] で，b. は 15 [mm] の跳躍した結果である．青線が誤差条件の赤線が報酬条件の手先軌跡であり，各色の丸で示すのが手先到達位置である．また，緑色の丸は跳躍後の目標点の位置を示す．以上の結果より，本シミュレーション結果は図 2.12 で予測していた目標点跳躍課題の結果を再現できしており，Izawa と Shadmehr ら [5] の学習モデルを基にし，(3.15) 式 によって目標点跳躍課題が再現できることを示す．

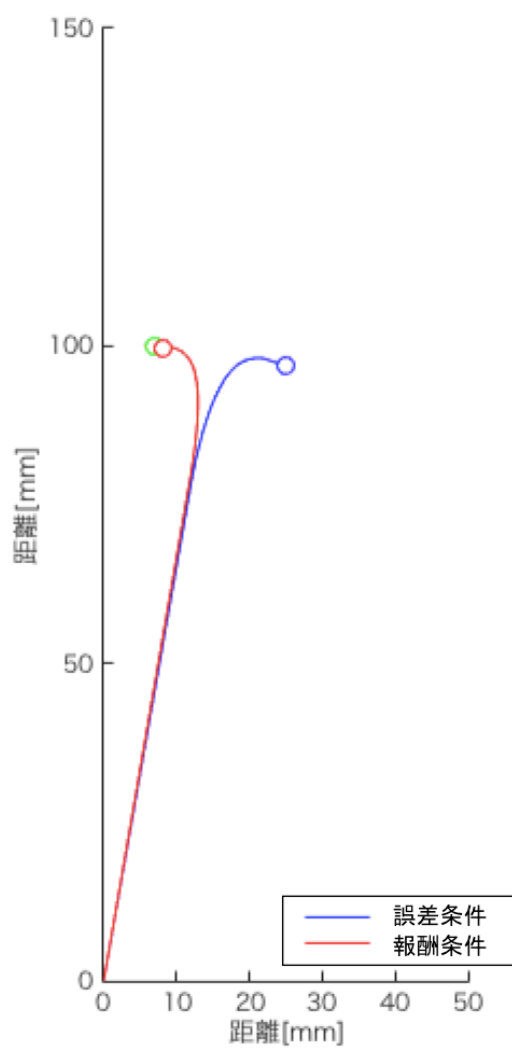


a. 誤差条件

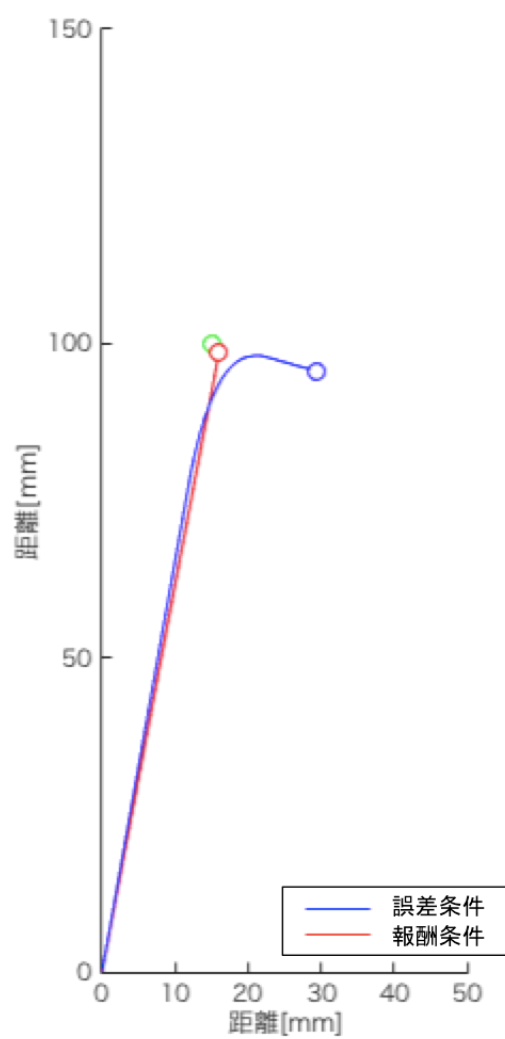


b. 報酬条件

図 4.2 数値シミュレーションによる内部順モデルの運動予測結果



a. 目標点 7 [mm] 跳躍



b. 目標点 15 [mm] 跳躍

図 4.3 数値シミュレーションによる目標点跳躍課題結果

第5章 被験者実験：報酬・誤差両条件とも右側へ運動補正を生成

本研究のすべての行動実験は、東京大学大学院教育学研究科身体教育学コースの野崎研究室の実験室をお借りして実施した。到達運動時の軌跡や速度を計測するためには、一般的にマニピュランダムと呼ばれるロボットアームのついた計測装置を使う。本研究では、以下の図 5.1 に示すマニピュランダム：Phantom Premium を用いて到達運動の計測を行った。また、本実験では被験者へ視覚情報としてカーソルを与えなくてはならないため、マニピュランダムで計測された運動はリアルタイムで可視化される必要がある。本研究では野崎研究室で開発された心理実験用 Phantom ライブラリを用いて、Phantom からの計測データをパソコンに保存しさらに視覚情報をプロジェクターを用いてリアルタイムにスクリーン上に投影した。以下の図 5.2 に、実際に実験を行った環境を示す。図 5.2 a. は被験者の背後から、b. は真上から見た状態である。被験者は、実験中は座席に付いたベルトによって姿勢を固定、さらに到達運動を行う右腕の手首にはサポータを装着させて手首が曲がらないように固定し、スクリーンは被験者の肩の下から全てを隠して実験中の運動を視認できなくした（図 5.2 a. を参照）。そのかわりにスクリーン上にプロジェクターによって手先位置と対応したカーソルが投影される（図 5.2 b. を参照）。また、カーソルは到達運動中は常に映し出されており、自分の運動中の手先の様子を確認することができる。図 5.2 a. において、下部にある黄色の点が運動開始点であり、上部の赤点が目標点、さらに他に比べてサイズが小さい白点はカーソルとなっており、図 2.14 に示す内容と同一の条件である。白点をロボットアームを操作することで制御し、黄点から赤点まで到達運動を行う。この試行を何度も繰り返し、到達運動課題を行うことで運動回転適応実験を実施する。

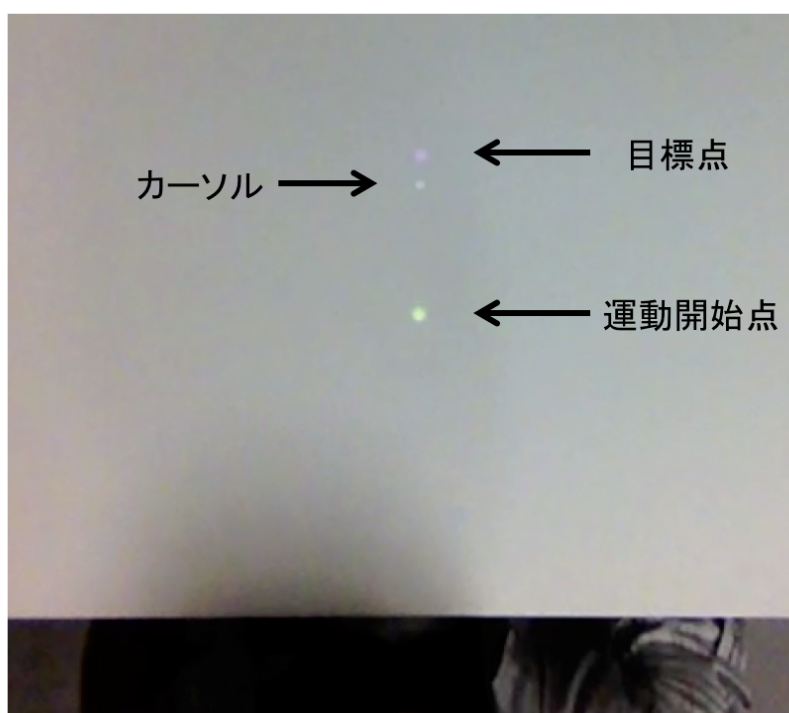
本研究では、被験者を募集し実験を行った。8名の被験者（男性、右利き、 20 ± 3 歳）には、誤差条件で実験を行った後に報酬条件での実験を行わせた。報酬条件を最初に行う実験を用意しなかったのは、報酬条件を最初に行う際には難易度が高くなかなか適応しない事例が存在したためである。この検討より本研究では、全ての被験者において誤差条件で運動回転適応実験を行った後に報酬条件の実験を行うように設定した。



図 5.1 マニピュラタム (Phantom Premium)



a. 実験の全体像



b. 被験者へ与える視覚情報

図 5.2 本研究の実験環境

5.1 運動回転適応実験の結果

以下に表 5.1 として実験条件を示す．この条件は先行研究 [5] に準ずる．表 2.1 にて示した先行研究 [5] の試行数は 500 試行あるが，本研究では 400 試行とした．これは，これは被験者の疲労の観点からも削除した方が良いと考えたためである．また，この 100 試行は最大回転に適応後に行っているので，本研究の実験では必要ないと考える．先行研究結果や事前に行った確認実験により 40 試行ほど繰り返せば被験者は回転へ適応することを確認しており，回転への適応において今回の試行の削除の影響はないと考える．

以下に，図 6.1 に行動実験における被験者一人分の運動回転適応実験結果の学習曲線を示す．他の 7 名の被験者の結果については付録にて述べる．これらの学習曲線は a. が誤差条件，b. が報酬条件の結果である．確認実験と同様の結果で，誤差条件では非常に良く回転に適応しており，報酬条件では分散が高いが回転へ適応している．また，報酬条件において，全被験者が誤差条件とは異なり自分の手先を明らかに右側へずらして到達させたという内観を得た．

5.2 目標点到達課題の結果

以下に，表 5.2 として実験条件を示す．このときの跳躍距離は，8 [deg] の回転に適応した場合の到達位置と同一位置の 15 [mm] とその半分の 7 [mm] となっている．なお，軌跡の平均には各試行毎にスプライン補間を用いた値を用いている．これは，生データでは運動開始時刻と運動終了時刻が異なることで 1 試行毎の軌跡を表現するデータ数が異なるため，各被験者共通のサンプリング数が必要となるためである．運動開始（Onset）と運動の終了判定（Offset）は，運動速度がピーク速度の 5% を超えるか否かを判定基準としており，以下の式で定義する．ここで， t は時間を表し，変数上部の点記号は微分を表す．

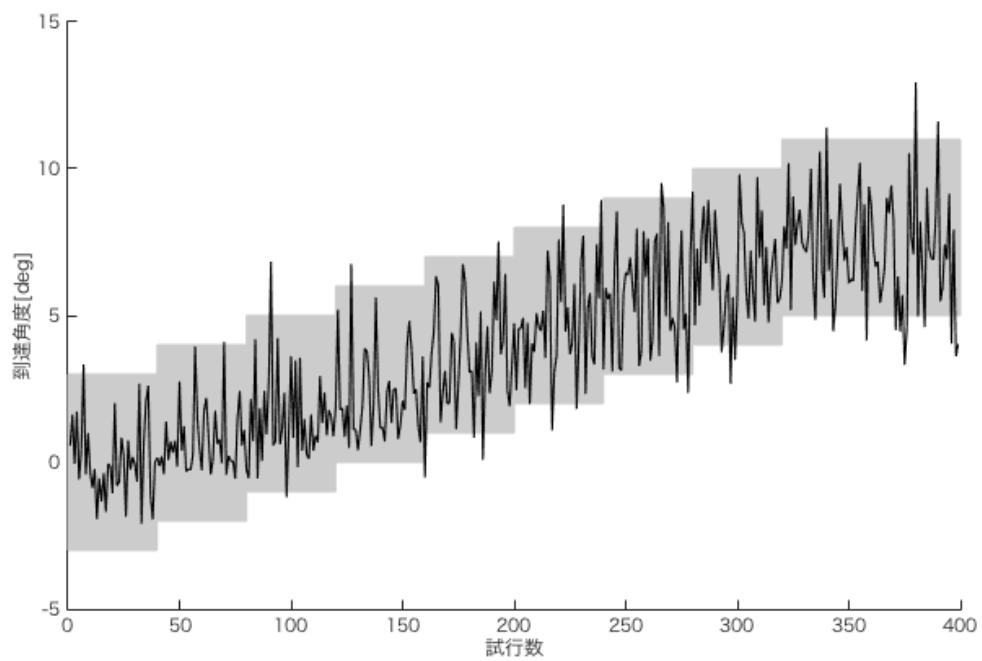
$$\text{Onset} : V(t) \geq V_m \times 0.05 \quad (5.1)$$

$$\text{Offset} : V(t) \leq V_m \times 0.05 \quad (5.2)$$

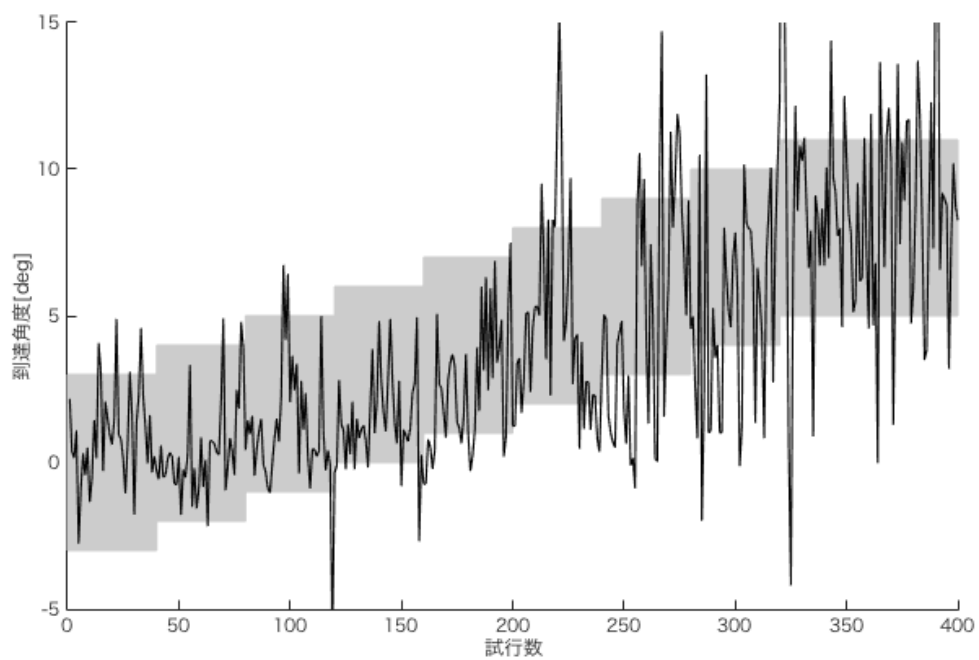
$$V(t) = \sqrt{\dot{x}(t)^2 + \dot{y}(t)^2}$$

$$V_m = \text{Peak Velocity}$$

以下に，図 5.4 として目標点到達課題の実験結果を示す．この結果は誤差条件と報酬条件において，全被験者の回転適応後の目標点跳躍課題の平均した軌跡である．平均に用いたデータ数は，目標点の跳躍距離が 7 [mm] の場合，誤差条件で 80 データで報酬条件は 58 データであり，跳躍距離が 15 [mm] の場合，誤差条件は合計 79 データで報酬条件は 57 データである．また，二種の跳躍距離において各条件の到達位置結果を比べると有意な差が示された（2 標本 t 検定, 7 [mm]: $p < 3.8284 \times 10^{-13}$, 15 [mm]: $p < 4.3003 \times 10^{-9}$ ）．

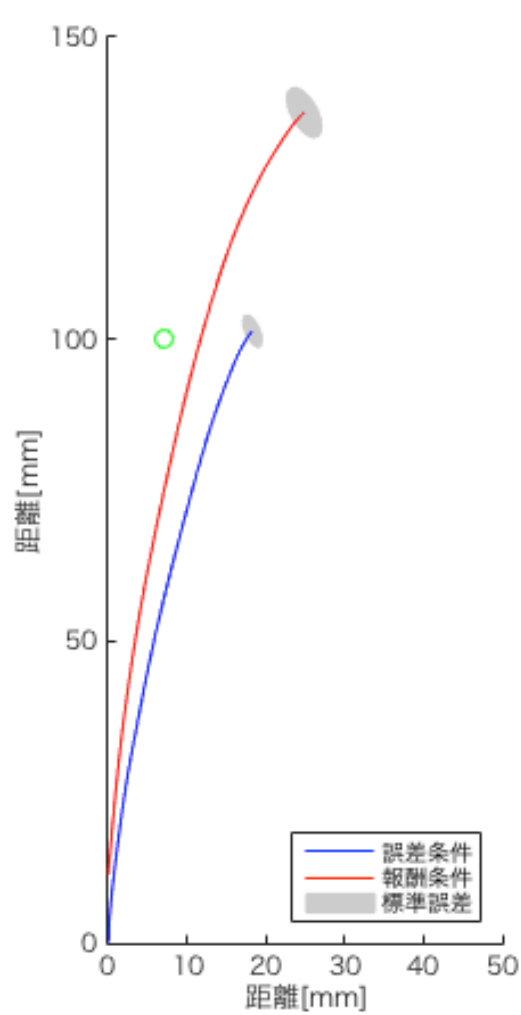


a. 誤差条件

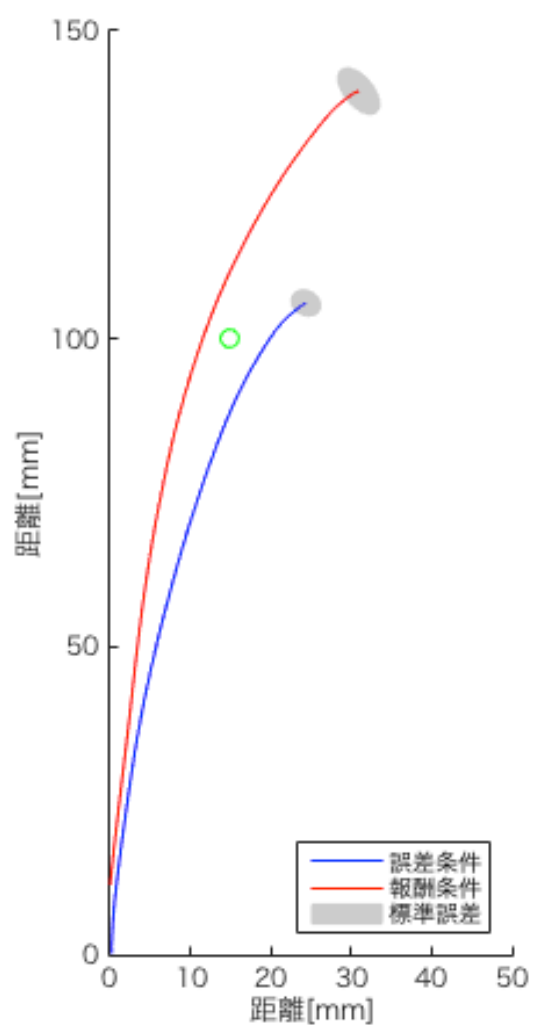


b. 報酬条件

図 5.3 被験者実験による運動回転適応実験結果：被験者 01



a. 目標点 7 [mm] 跳躍



b. 目標点 15 [mm] 跳躍

図 5.4 被験者実験による目標点跳躍課題の結果

表 5.1 運動回転適応実験の条件（被験者実験）

被験者数	8 名/各班
試行数	400 [試行]
到達距離	100 [mm]
回転方向	反時計回り
最大回転	8 [deg]
到達範囲	± 3 [deg]
回転付与	+1 [deg]/40 [試行]

表 5.2 目標点跳躍課題の条件（被験者実験）

	条件 1	条件 2
試行数	10 [試行]	10 [試行]
跳躍方向	右方向	右方向
跳躍距離	7 [mm]	15 [mm]

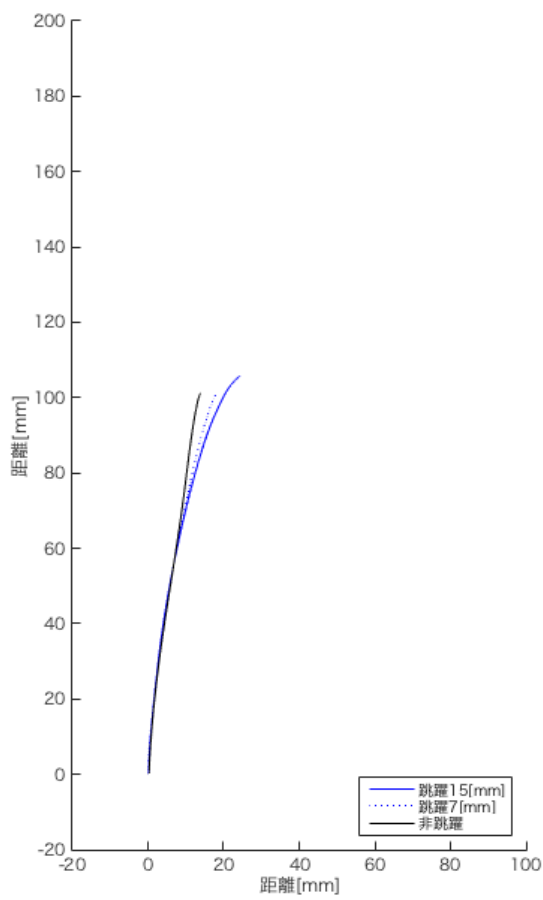
5.3 被験者実験結果の考察

図 5.4 より，両条件での手先到達位置には優位な差が存在する結果が得られた．この結果より，感覚情報を与えて回転へ適応させている誤差条件は内部順モデルが回転に適応していると考えられる．しかし，報酬条件の手先軌跡は右側へ到達する結果を得られ，感覚情報を与えずに目標へ到達できたかできていないかのみでの情報しか与えられない報酬条件については，数値シミュレーションと異なる結果を得た．本章の考察では，数値シミュレーション，被験者実験の結果を踏まえて誤差条件および報酬条件ではどのような学習プロセスによるのかを考察する．

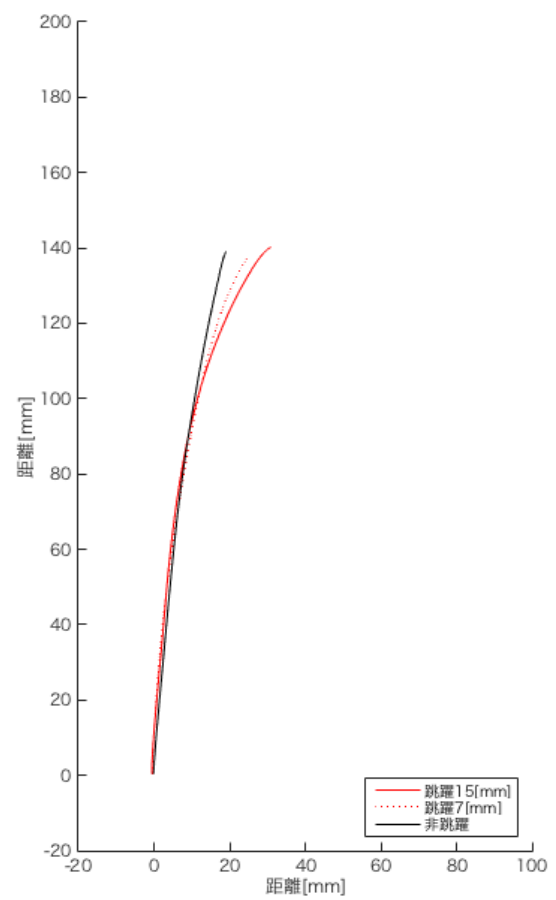
5.3.1 考察：目標点の跳躍時には運動補正が生成されているかの検討

まず，目標点跳躍により運動補正が行われているかを検証するために，最大回転に適応した後の目標点の跳躍時と非跳躍時における手先軌跡および手先速度の比較を行う．以下に，図 5.5 から 5.9 に 8 [deg] に最大回転に適応した後における跳躍時と非跳躍時の到達運動の手先軌跡と手先速度を示す．ここで，黒線が 8 [deg] の回転へ適応した後の非跳躍時の平均データである．また，青線が報酬条件での平均データで赤線が報酬条件での平均データを示しており，実線が 15 [mm] 跳躍時の結果であり点線が 7 [mm] 跳躍時の結果を示している．この色の差異は図 5.5 から 5.9 の全てに共通している．また，各色で半透明で示している範囲は標準偏差である．

図 5.5 からは，黒線で示す非跳躍時の軌跡と赤および青色の各種線の跳躍時の軌跡が異なっていることがわかる．また，図 5.6 から図 5.9 で示す速度では，補正がかかる x 軸方向の速度（両図左側）が，非跳躍時の速度はピークが 1 つなのに対して跳躍時の速度はピークが 2 つあることがわかる．これら結果から，誤差班および報酬班の目標点到達課題では，目標点の跳躍によって運動の補正が行われている結果が示されている．よって，目標点跳躍課題において被験者は到達運動中に運動の補正を行い，与えられる情報が異なる条件によって異なる運動を生成する結果を得た．

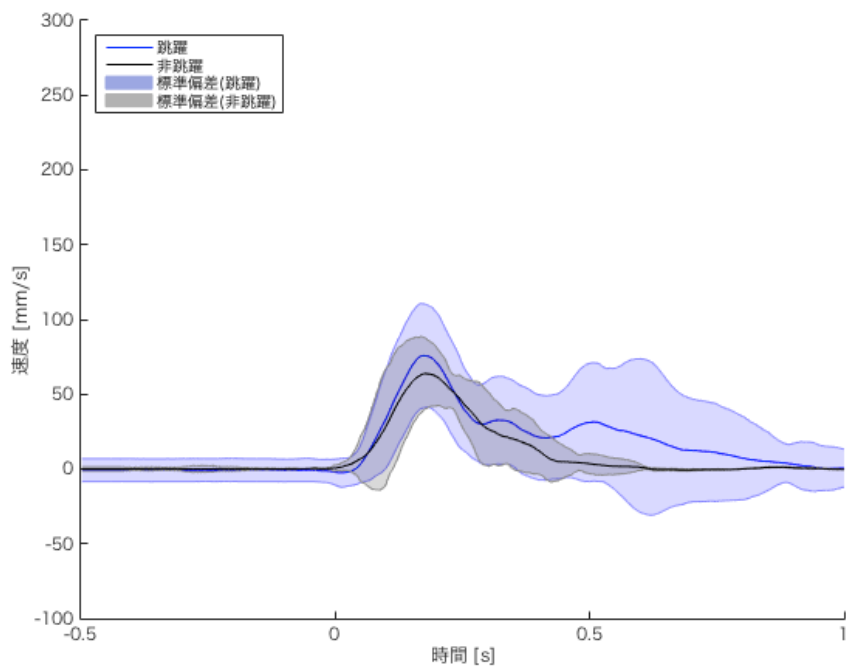


a. 誤差条件

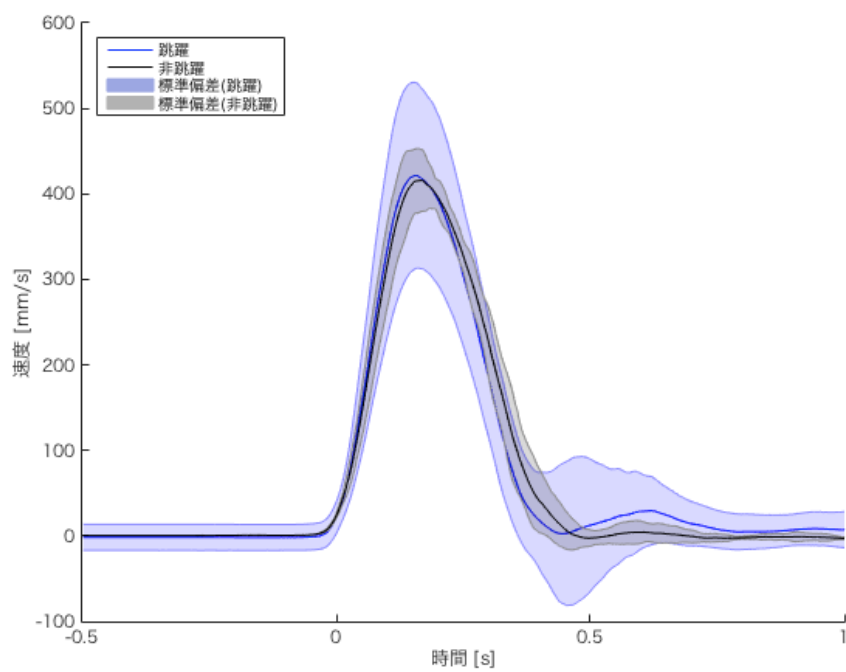


b. 報酬条件

図 5.5 跳躍時と非跳躍時における軌跡の比較

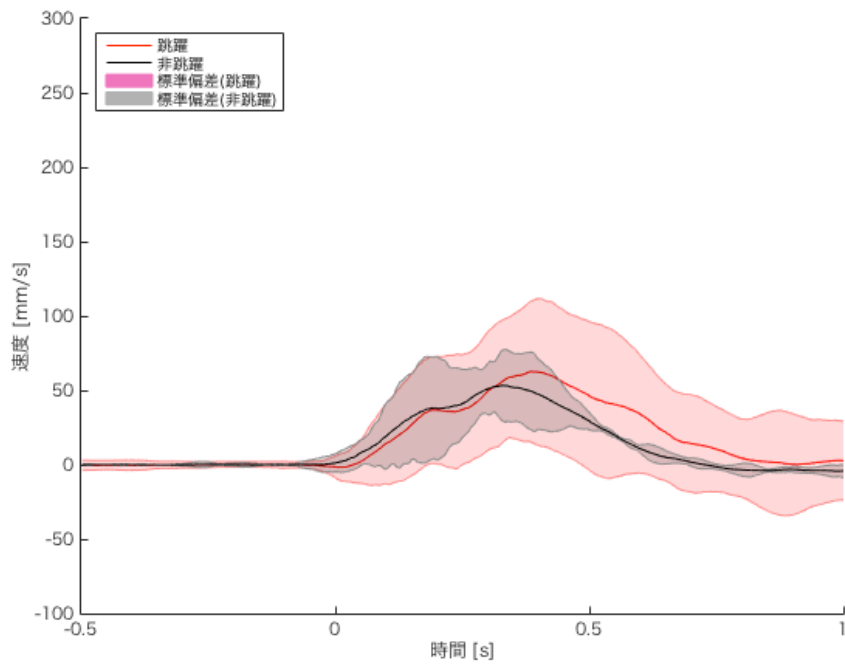


a. 誤差条件: x 方向速度

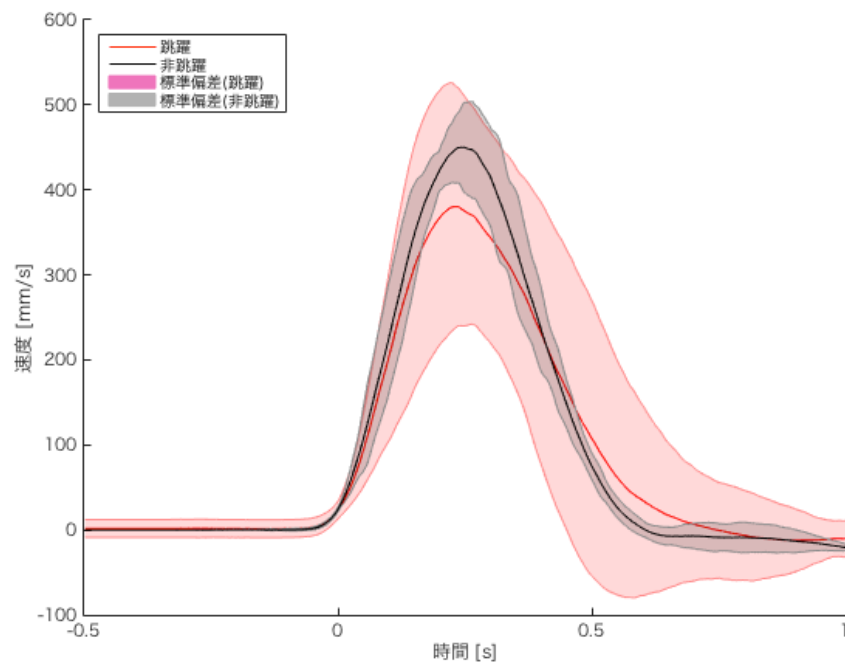


b. 誤差条件: y 方向速度

図 5.6 誤差条件における跳躍時と非跳躍時の速度比較: 7 [mm] 跳躍

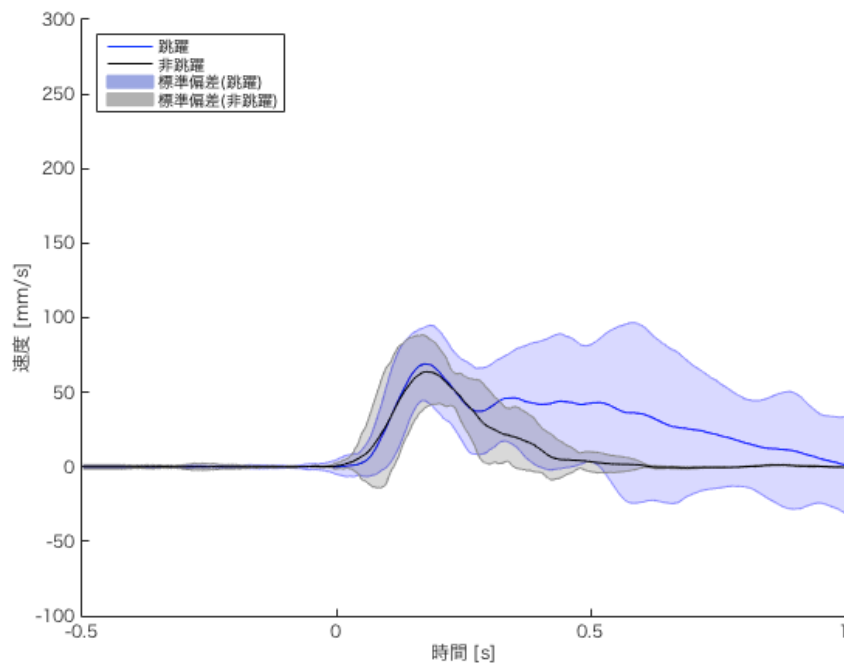


a. 報酬条件: x 方向速度

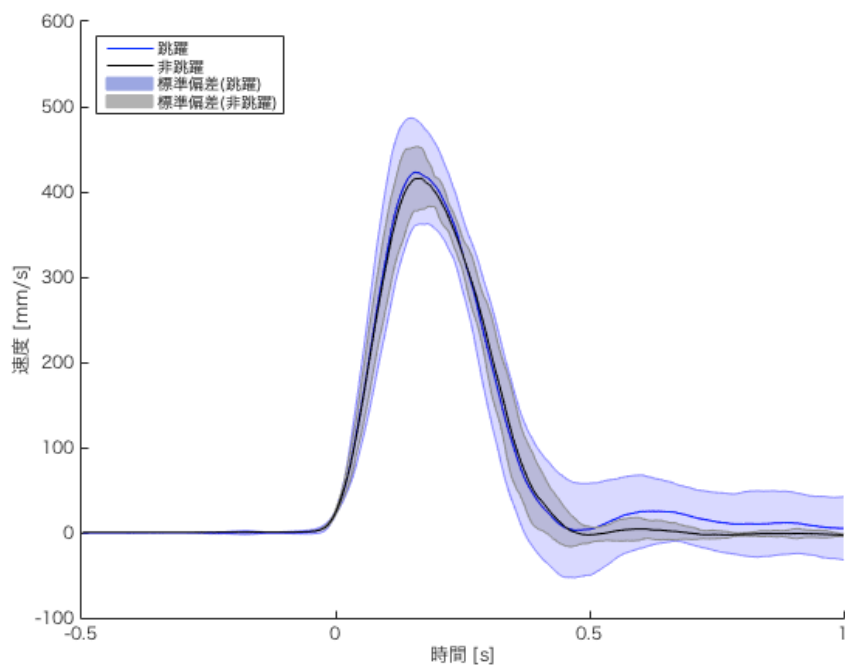


b. 報酬条件: y 方向速度

図 5.7 報酬条件における跳躍時と非跳躍時の速度比較 : 7 [mm] 跳躍

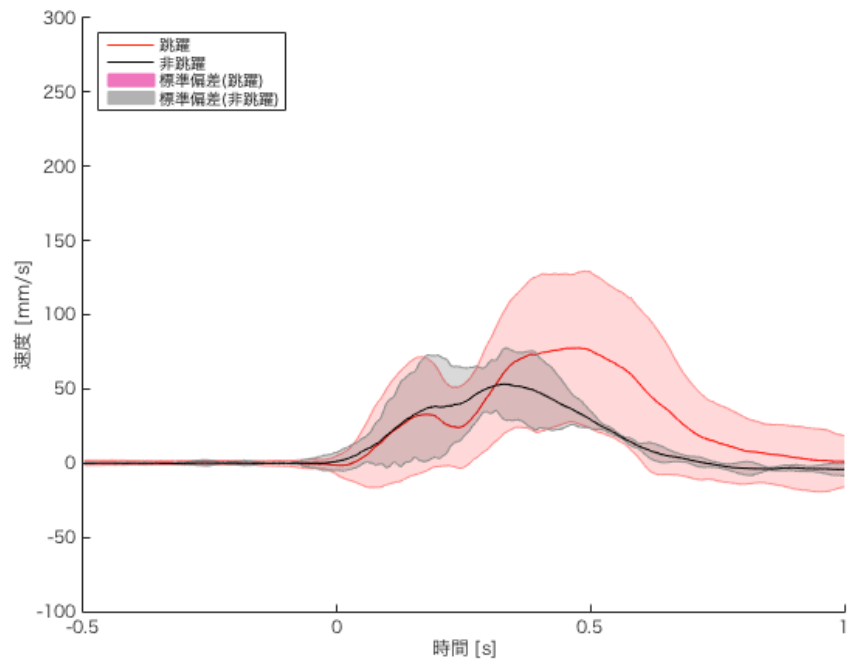


a. 誤差条件: x 方向速度 (15 [mm] 跳躍)

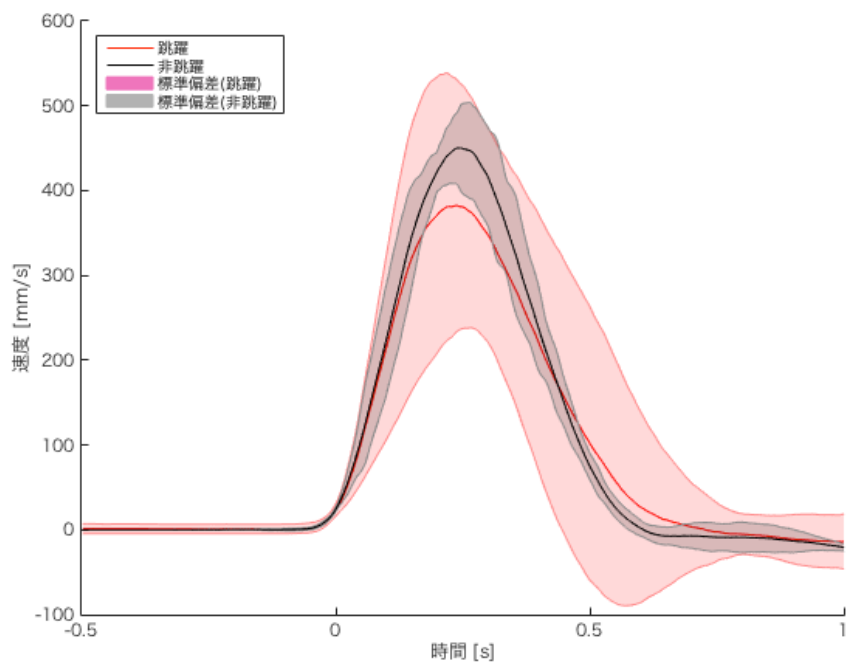


b. 誤差条件: y 方向速度

図 5.8 誤差条件における跳躍時と非跳躍時の速度比較: 15 [mm] 跳躍



a. 報酬条件: x 方向速度



b. 報酬条件: y 方向速度 (15 [mm] 跳躍)

図 5.9 報酬条件における跳躍時と非跳躍時の速度比較: 15 [mm] 跳躍

しかし、ここで被験者実験の目標点跳躍課題において2つの疑問点が存在する。1つ目に、報酬条件における手先の到達位置が、新目標点へではなく目標点より右側に到達したことである。2つ目に、目標点の跳躍時と非跳躍時の手先軌跡の違いは見えても、両条件において図 5.5 に示すように跳躍距離が 7 [mm] と 15 [mm] 両方の手先軌跡にはほぼ違いが見られない結果を得たことである。これらを踏まえて、目標点跳躍課題において以下のことが考えられる。

- 1つ目の問題点より、被験者は目標点の跳躍距離 7 [mm] と 15 [mm] の2つの距離の差を区別できていないのではないか
- 2つ目の問題点より、誤差条件での経験が影響して報酬条件では見えないカーソルを新目標点へ到達させるように運動したのではないか

以下は、被験者実験における目標点跳躍課題の結果の考察について検討することとする。

5.3.2 考察：なぜ報酬条件では手先が目標点の右側へ到達したのかの検討

行動実験において誤差条件の到達位置結果は、跳躍後の新目標点の右側へ到達する予測結果と同様の結果を得た。しかし報酬条件の到達位置結果では、跳躍後の新目標点に向かって到達するとの予測結果と異なり、誤差条件と同様の右側に到達する結果を得た。この結果は、目標点の跳躍距離が 7 [mm] と 15 [mm] のどちらの条件でも右側へ到達する結果を得た。ここでは、なぜ報酬条件において予測した結果と異なる運動結果を得たのかを問題点とし考察する。本研究では、このような結果を得た原因を、

- (1) 前に行った誤差条件での経験が報酬条件の実験へ影響された
- (2) 報酬条件では見えないカーソルを想像することでカーソルと目標点間の誤差により内部順モデルが回転に適応した
- (3) 本研究の学習モデルで明らかにできていない機構の働きが存在する

と考える。

(1) は、本研究での行動実験が被験者へ視覚的な感覚情報が与えられる誤差条件を行いその後に報酬条件での運動回転適応実験を行ったことが問題点の原因と考える。最初に行った誤差条件の経験が次に行う報酬条件へ影響してしまい、不意に右側へ到達運動を行う運動補正を生成した可能性があると考えられる。この考察を検証するには、誤差条件の班と報酬条件の班に分けて行動実験を行う必要がある。本研究では時間の関係上そこまでの検証ができなかった。

(2) では、先行研究では報酬条件で引き起こされない内部順モデルの回転適応が、カーソルを想像する事により引き起こされたのではないかと考える。(1) の考察時にも述べたが、本研究において被験者はすべて誤差条件を行った後に報酬条件で運動回転適応実験を

行った．誤差条件では常にカーソルが見せられていた．その結果，報酬条件ではその見えないカーソルを想像し，感覚予測誤差を生成する事で内部順モデルが回転へ適応する事が考えられる．つまり，(3.7) 式が以下のように置き換わる可能性がある事を考える．

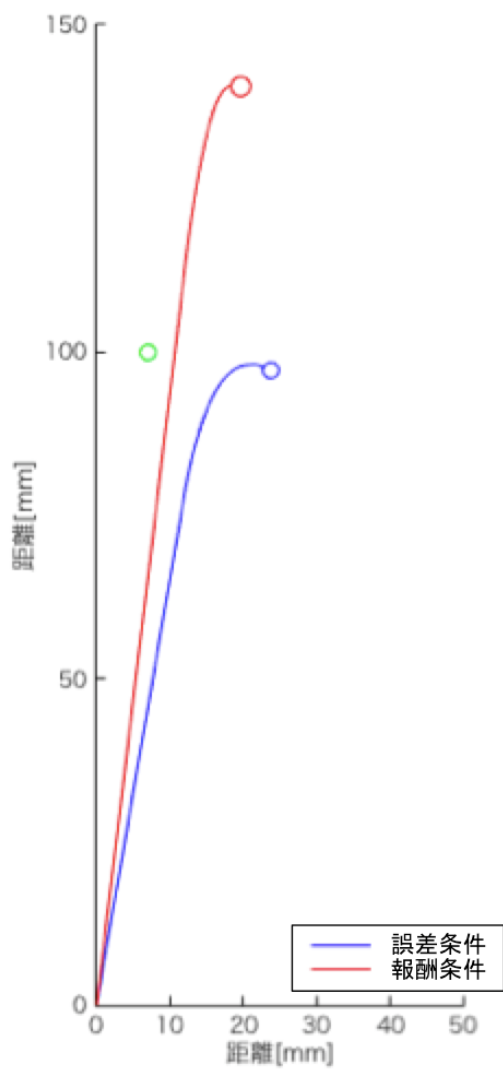
$$\hat{\mathbf{x}}^{(k|k)} = \hat{\mathbf{x}}^{(k|k-1)} + K^{(k)}(\hat{c}^{(k)} - C\hat{\mathbf{x}}^{(k|k-1)}) \quad (5.3)$$

この結果，先行研究では報酬条件では機能しないと提案されていたカルマンフィルターが(5.3) 式により回転を学習すると考えた．実際に(5.3) 式を基に，本研究における学習モデルの数値シミュレーションを行い，内部順モデルの回転への適応が起こるかを検討した．以下に，(5.3) 式を基とした目標点跳躍課題の数値シミュレーション結果を示す．この結果は，シミュレーションにより生成される手先軌跡がどの方向に運動するかとともに，この条件においても学習曲線に異常がなく正常に回転へ適応で来ているかを確認する．

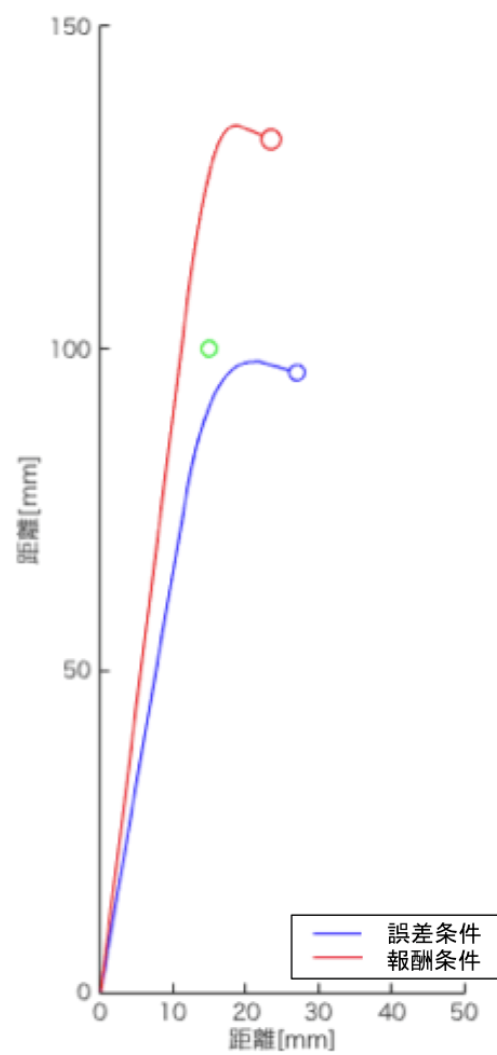
図 5.10 において左側が目標点の跳躍距離が 7 [mm] であり右側が 15 [mm] の運動結果を示しており，青線が誤差条件および赤線が報酬条件での手先軌跡結果である．図 5.11 に各条件における学習曲線を示す．a. が誤差条件，b. が報酬条件における運動回転適応での学習曲線である．両図の結果として，本研究における被験者実験の結果を良く説明できる結果となった．また，図 5.10 での手先軌跡はオーバーシュートしているが，これは誤差条件での結果を基準として報酬条件の結果にどれだけ誤差があるかの誤差率を計算し，その誤差率分だけ到達距離がずれると仮定して手先軌跡の到達距離を調整する事により再現した．さらに以下に図 5.12 として内部順モデルの回転適応状況を確認するために，内部順モデルで予測される回転摂動と手先位置を示す．ここで a. に誤差条件の，b. に報酬条件の結果を示し，縦軸は到達角度であり横軸は試行数である．図中に黒線で示す値が内部順モデルが適応する回転摂動の値であり，灰色の線がその試行のときに生成される運動指定から予測する自分の腕の到達位置である．

図 5.12 a. の誤差条件の結果は図 4.2 a. に示す結果と同様であるが，報酬条件の結果は図 4.2 b. の結果とは異なり，誤差条件特有の内部順モデルの回転への適応を示している．また図 5.12 b. より，内部順モデルの適応が誤差条件の結果ほど強くないことがわかる．この結果より図 5.12 の条件でのシミュレーション結果は，実験での誤差条件時のように内部順モデルが多少回転に適応するが，被験者実験時の報酬条件における結果のように右側にへ意識して手先を到達させたという内観も再現できた状態で回転へ学習できた結果を再現した事がわかる．よって報酬条件での目標点跳躍課題で手先が右側へ到達する原因の 1 つとして，誤差条件での経験が影響することで被験者が見えないカーソルを想像して回転へ適応してしまった事が考えられる．

最後の(3) では，先行研究における学習モデルに問題点があることが原因と考える．(2) の考察は，誤差条件を先に経験していないと得られない結果だと考えられる．ここでは，先行研究の学習モデルに問題点があるとして仮定して考察する事とする．運動指令は(3.8) 式で示すように，感覚予測誤差からの回転摂動 $\hat{p}$ と報酬予測誤差からの報酬値 w_r および

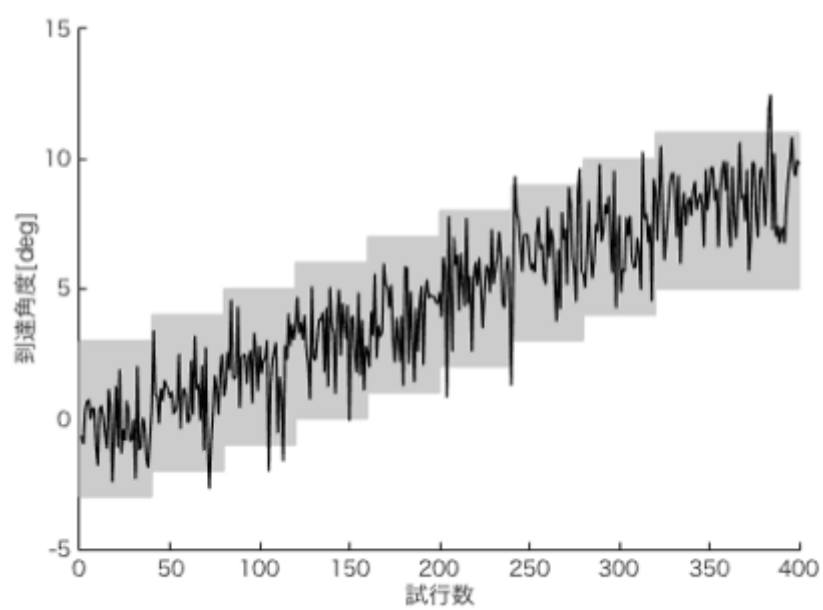


a. 目標点 7 [mm] 跳躍

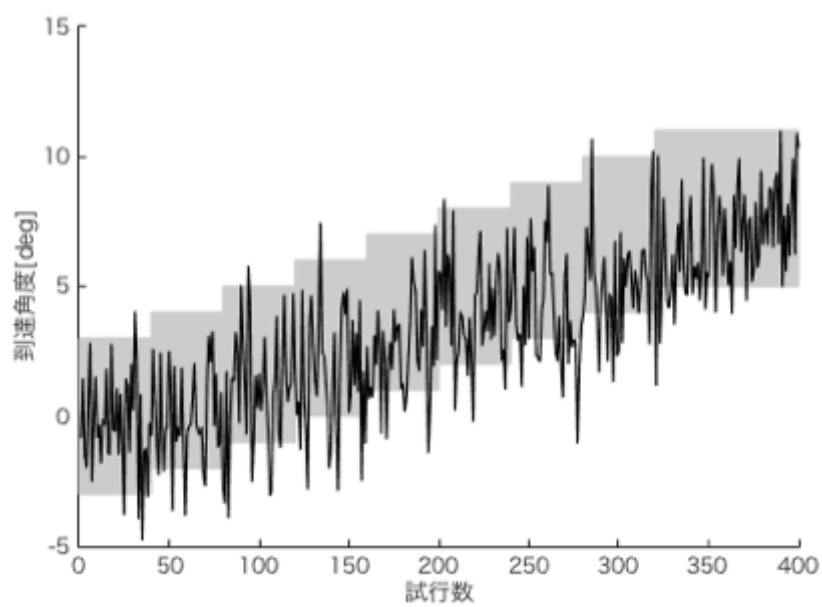


b. 目標点 15 [mm] 跳躍

図 5.10 数値シミュレーションによる目標点跳躍課題結果 ((5.3) 式ベース)

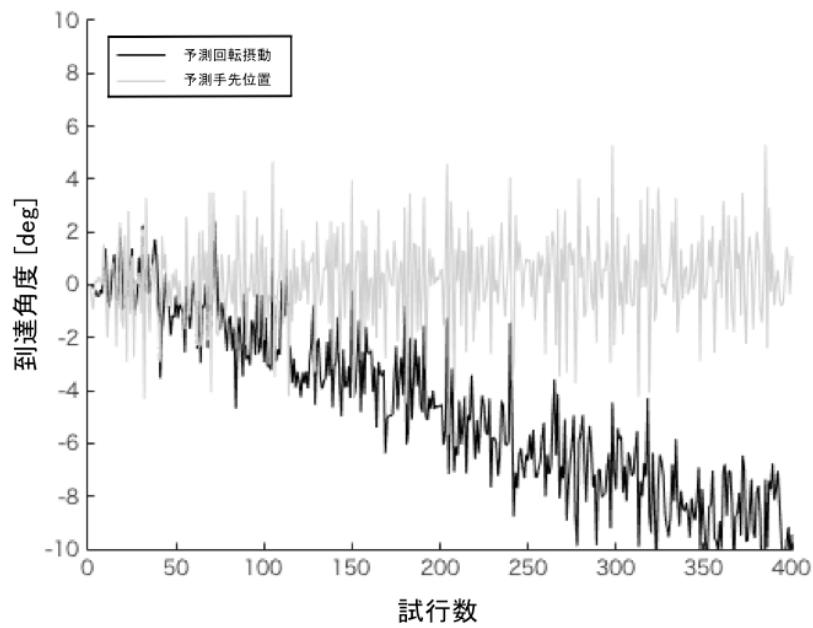


a. 誤差条件

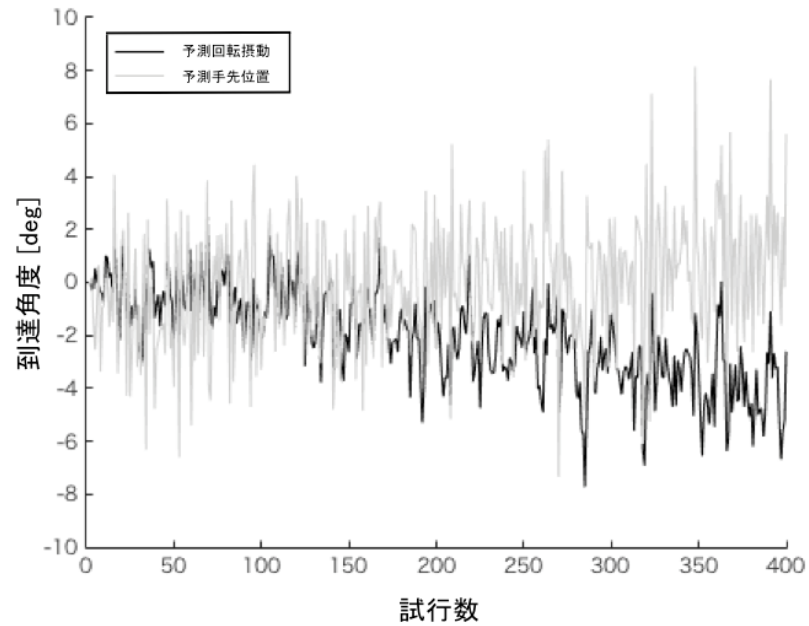


b. 報酬条件

図 5.11 数値シミュレーションによる運動回転適応実験結果 ((5.3) 式ベース)



a. 誤差条件



b. 報酬条件

図 5.12 数値シミュレーションによる内部順モデルの運動予測結果 ((5.3) 式ベース)

ノイズ n_u から構成される．報酬条件では，被験者は視覚情報が与えられないので感覚予測誤差による学習が起こらないと考える．また，報酬値 w_r は (3.13) 式に示すように運動指令ノイズが引き金となり計算されると考えられる．これは，被験者が視覚情報を与えられない状況で回転に適応するためにはあるきっかけが必要となり，先行研究ではそのきっかけを n_u に設定しているのである．被験者は報酬条件において到達運動を行う時に，目標点に対して常に直線的な運動を心がける．しかし，それでは被験者は回転摂動を段階的に与えられると突然報酬情報を得られなく時が訪れる．被験者はなぜ当たらなくなったのかはわからないが直線的な到達運動を繰り返すが，運動指令ノイズによりたまたま与えられた回転を消去できる方向へ到達運動する事がある．このときの報酬情報により，(3.13) 式に示すようにそのノイズ分だけの回転を学習するのが強化学習による行動選択器の動作である．しかし， n_u は脳内の内部変数であり，被験者はもちろん実験者も知る事ができない．であるにもかかわらず，ここでは被験者が運動指令ノイズ n_u を理解しているような運動を示している．これより，本研究では報酬条件においてこの運動指令のノイズである n_u が被験者の予測値に影響しないように，状態方程式を以下のように再定義する事を考える．また，報酬条件ではカルマンフィルタが働かない事を仮定するのでカルマンゲイン K はつねに 0 となるようにする．

$$\hat{\mathbf{x}}^{(k|k)} = \hat{\mathbf{x}}^{(k|k-1)} + K^{(k)}(y^{(k)} - \hat{y}^{(k)}) - \alpha_v \delta_k n_u^{(k)} \quad (5.4)$$

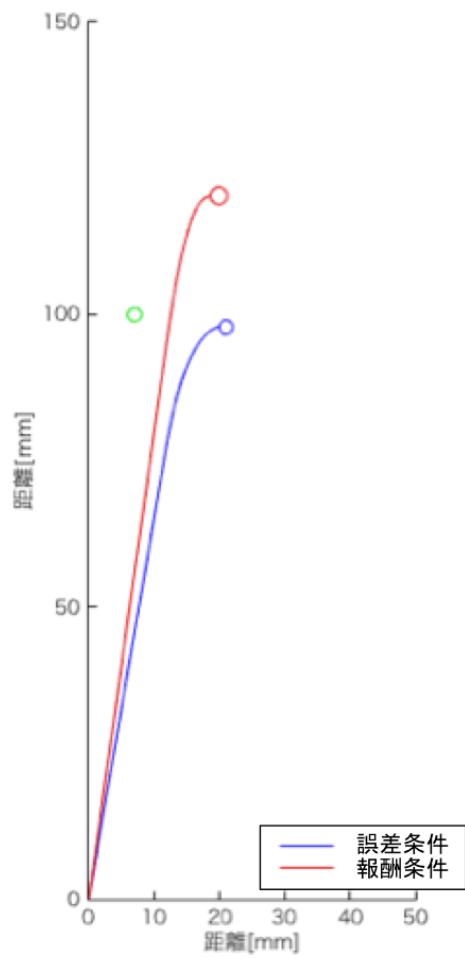
以下に，(5.4) 式に基づく目標点跳躍課題のシミュレーション結果と学習曲線を示す．この結果から，シミュレーションにより生成される手先軌跡がどの方向に運動するかとともに，この条件においても学習曲線は正常に回転へ適応できているかを確認する．

これらの結果より，(2) と同様に与えられた回転に適応し，実験結果を説明できる手先軌跡を得た．ここで，図 5.15 として (5.4) 式に基づく内部順モデルの予測結果を示す．この結果からも，(2) 同様に内部順モデルの回転への適応と共に被験者の内観を再現できる結果を得た．よって報酬条件での目標点跳躍課題で手先が右側へ到達する原因の 1 つとして，先行研究の学習モデルが実際のヒトの学習モデルとは異なる事も考えられる．

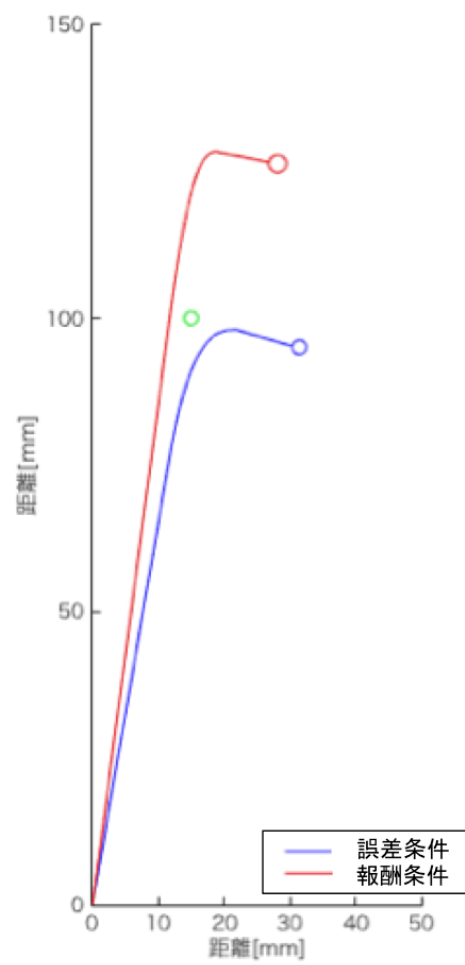
以上の考察より，本研究では先行研究とは異なる学習モデルが働く可能性がある事が明らかとなった．内部順モデルは感覚予測誤差のみで適応するのではなく，脳内で生成される誤差によって内部順モデルが環境に適応する場合が存在することを示唆する．

5.3.3 考察：被験者は異なる距離の目標点跳躍を区別できるのかの検討

被験者実験では，誤差条件と報酬条件の両班において二種類の目標点跳躍距離で目標点跳躍課題を行った．しかし，跳躍距離を変化させてもそれぞれの手先軌跡はほぼ変わらなかった．ここでは，この跳躍距離の異なる場合において変化しなかった原因を検討する．まず原因として考えられるのは，目標点跳躍課題の到達運動中には誤差条件および報

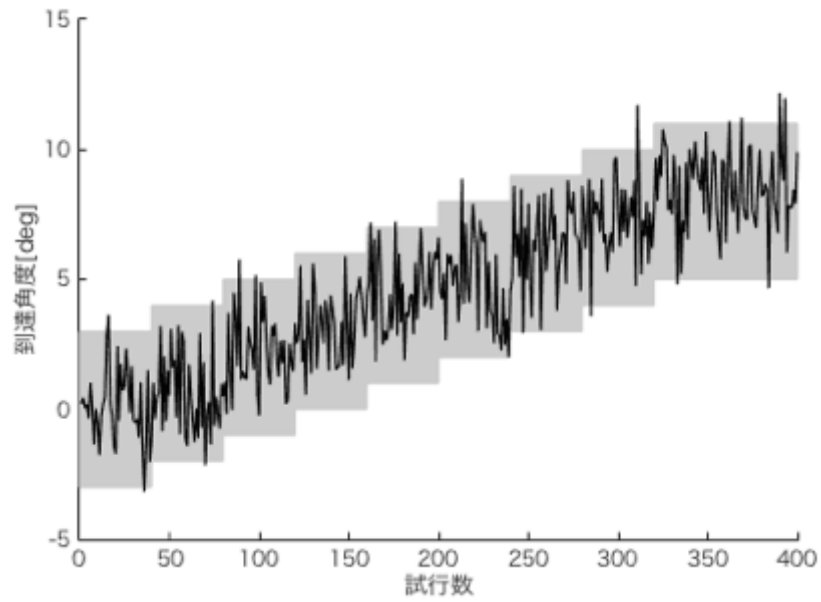


a. 目標点 7[mm] 跳躍

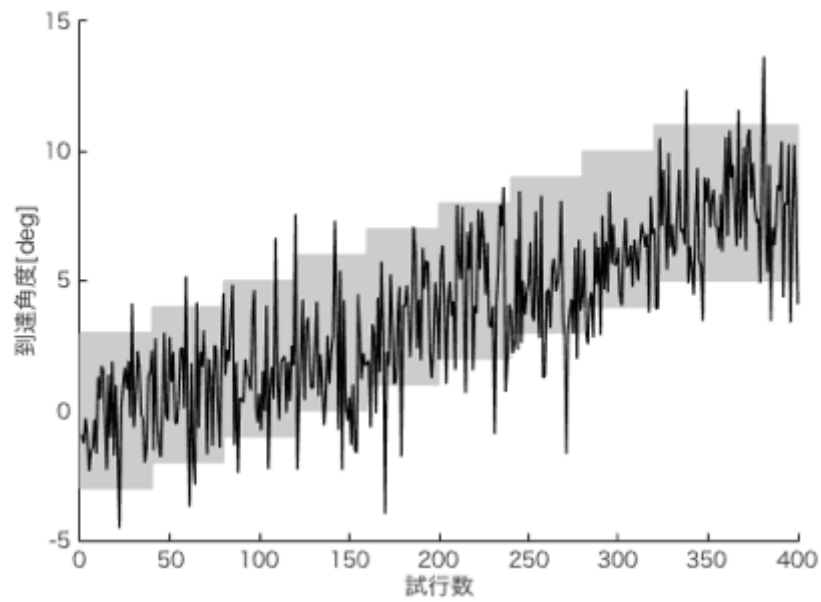


b. 目標点 15[mm] 跳躍

図 5.13 数値シミュレーションによる目標点跳躍課題結果 ((5.4) 式ベース)

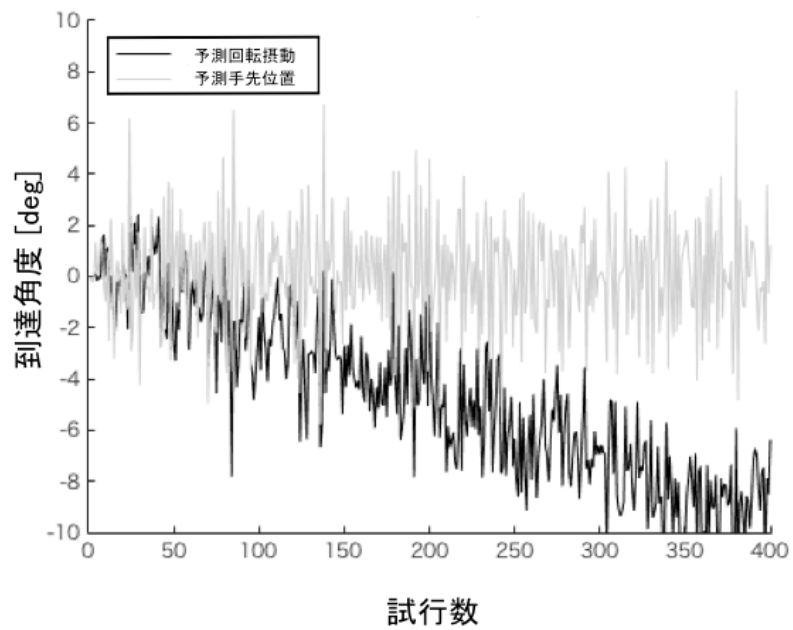


a. 誤差条件

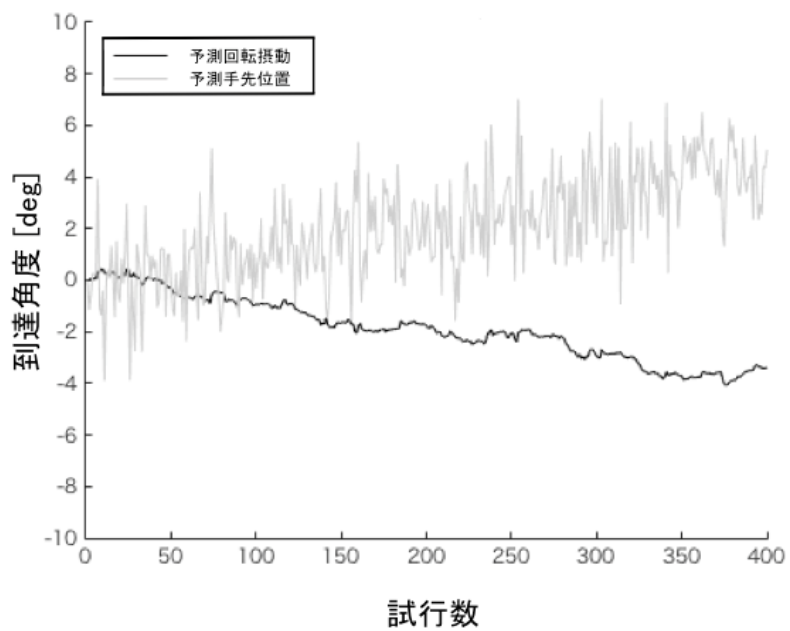


b. 報酬条件

図 5.14 数値シミュレーションによる運動回転適応実験結果 ((5.4) 式ベース)



a. 誤差条件



b. 報酬条件

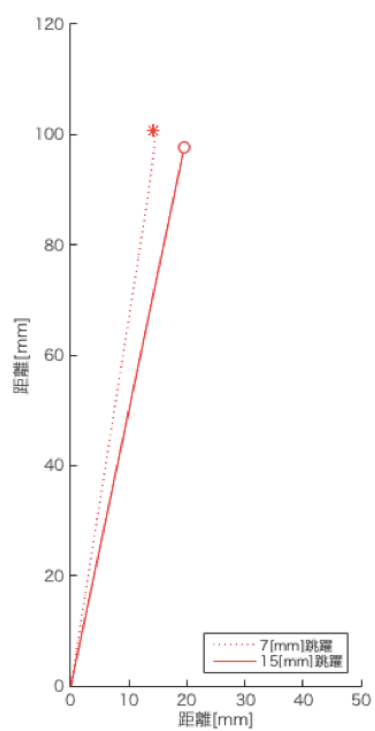
図 5.15 数値シミュレーションによる内部順モデルの運動予測結果 ((5.4) 式ベース)

報酬条件の両条件で視覚情報であるカーソルが与えられない事により，自分の手先位置の不確実性が増加したため，予測する手先位置の分散が高くなり跳躍距離の差が明確に区別できないということが考えられる，この考察を検討するために，異なる目標点跳躍距離での手先軌跡の定量的な差を検討する．誤差条件の 7 [mm] と 15 [mm] および報酬条件の 7 [mm] と 15 [mm] の到達位置において 2 標本 t 検定を行うと，前者は $p = 0.270149$ となり後者では $p = 0.491681$ という結果を得た．この結果より，各条件における跳躍距離の異なる場合の到達位置に有意な差は見られない．よって自分の手先位置が確認できない様な不確実性の大きい実験計画だと，約 10 [mm] 程度の差異を区別できない可能性がある．

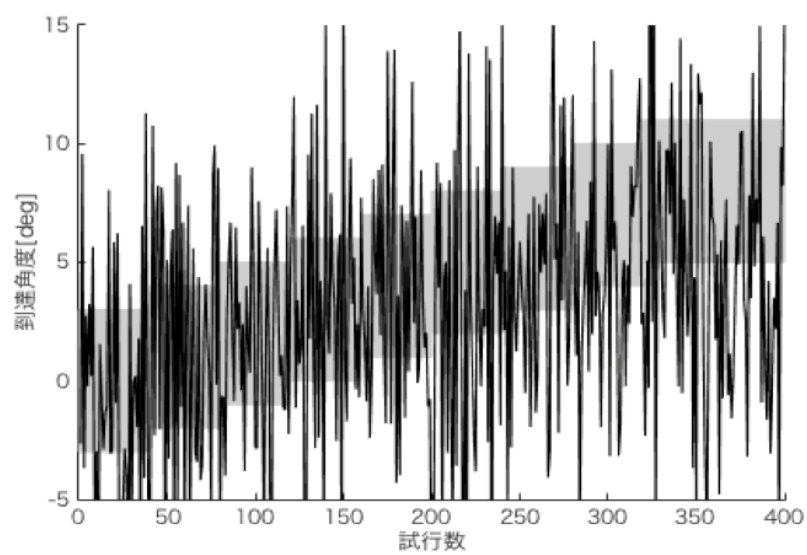
また，不確実性が大きい程に示された目標点の違いを区別できないのであれば，常にカーソルを消している報酬条件結果の方が誤差条件での結果より 7 [mm] と 15 [mm] の跳躍距離を区別できていないはずである．ここでは，この報酬条件の結果を誤差条件での結果と比べ，不確実性と目標点跳躍距離の関連性について検討する．先に述べたように，誤差条件および報酬条件での 7 [mm] と 15 [mm] 跳躍距離の到達位置結果は，誤差報酬両条件において有意差がないという検定結果を示した．ここで各条件において，異なる跳躍距離間での到達位置にどれぐらいの相関があるか確認した．各条件での相関係数を計算すると，誤差条件では $r = 0.2178$ となり報酬条件では $r = 0.6067$ という計算結果を得た．この結果より誤差条件より報酬条件の結果間ではやや強い相関がある結果を得た．よって被験者は，誤差条件より報酬条件において 7 [mm] と 15 [mm] の目標点の跳躍距離の違いを区別できずに到達運動を行ったと考える．また，不確実性が大きくなる事により目標点跳躍距離を区別できなくなるのか，数値シミュレーションにより検討する．本研究では，本学習モデルにおける数値シミュレーションで本実験における報酬条件の目標点跳躍課題結果を再現できるかを検証する．

報酬条件ではカーソルが視認不可なので自己の手先位置がわからなくなっていると考察したので，手先位置の分散を表す σ_h^2 の値を変更することとする．この値の変更は，被験者実験の到達位置の結果から分散値を計算することで行った．8 [deg] に適応した後の到達運動結果 20 試行分の結果から分散値を計算し，結果的に報酬条件の分散値は誤差条件に比べて約 1.9 倍高い値と計算された．この結果を利用し，数値シミュレーションの条件を変更して再計算した．上記の条件で再検証したところ，図 4.3 に示したシミュレーション結果と比べても差がない結果であり，分散値の増加による効果は見られなかった．しかし，図 5.16 に示すように， σ_h^2 の値を極端に高くして自分の手先位置が全くわからない状態を再現して目標点跳躍課題をシミュレーションしたところ，目標点跳躍距離が 7 [mm] と 15 [mm] の間の差は半分以上になり統計的に差がある結果を得た．以下に，数値シミュレーションによるその結果を示す．このとき，二つの跳躍距離条件における x 軸方向の差は 3 [mm] 程であり，従来の 7 [mm] のさに比べて半分以下の小さな差となっている．2 標本 t 検定において， $p = 0.0365$ となり，この二つの跳躍距離における到達位置結果は有意差が存在しない結果を示している．この結果より，本研究では不確実性が増加することは目標点跳躍距離の違いを区別できなくなる 1 つの原因だと考察する．

しかし，図 5.16 b. の学習曲線を確認すると不確実性が増加すると手先到達位置の分散



a. 手先軌跡



b. 学習曲線

図 5.16 報酬条件の数値シミュレーションによる手先軌跡と学習曲線結果

が非常に大きくなり、シミュレーション結果は被験者実験時の学習曲線を再現できると言える結果にはならない事が明らかとなった。たしかに、自己の手先位置がわからなくなる事により目標点の跳躍距離の区別がつかなくなる事も1つの考察と考えられるが、この結果より本研究では手先軌跡が目標点の跳躍距離に関わらずほぼ同一なのは、跳躍距離間の差は7 [mm] という非常に小さな値である事が一番大きな要因として考えられる。

これらの結果から、先行研究 [5] の内部順モデルは感覚予測誤差のみで回転に適応するという提案とは異なり、本研究は視覚情報が与えられる場合のみでなく視覚情報が与えられていない場合においても内部順モデルが回転へ適応する可能性が存在すると提案する。

第6章 結論

本研究では、内部順モデルの回転への適応は、視覚的に与えられる情報からの感覚時予測誤差が影響することの他にも、視覚情報が与えられない場合においても起こる結果を得た。これは、先行研究 [5] と同様の結果であり、ヒトの学習において感覚予測誤差を与える事が重要である事を示した。しかし、本研究では報酬条件においても内部順モデルが適応する被験者実験の結果を得た。これは先行研究とは異なる結果であり、数値シミュレーションから報酬条件においても内部順モデルが適応するメカニズムを考察した。

運動回転適応実験の先行研究では、回転適応に重要な役割を担う内部順モデルの適応状況を直接的に確認できていないと考える。そこで、本研究では内部順モデルの適応には何が重要な要素となるのかを検証するために、目標点跳躍課題という実験を計画した。まず先行研究の学習モデルを基にした数値シミュレーションにより目標点跳躍課題での結果を予測し、行動実験によりヒトが提案されている学習モデルによって運動回転学習を行っているのかを検証した。

今後の展望として行動実験におけるカウンターバランスの考慮、一人の被験者へ他条件の影響を最小限にすることを考える。これらの提案から実験の頑健性を高め、行動実験の報酬条件において内部順モデルがどのように変化しているかを検討する必要がある。

6.1 本研究で明らかになったこと：誤差条件だけでなく報酬条件でも内部順モデルの回転への適応する可能性を示唆

被験者には運動回転適応実験において、視覚情報が与えられる誤差条件と与えられない報酬条件に分けて、行い両条件とも付与された回転に適応した。最大回転に適応した後、本研究が提案する目標点跳躍課題を行い、各条件における内部順モデルの適応状況を比較しようと試みた。被験者実験の誤差条件での結果は、先行研究 [5] を基にした数値シミュレーションで予測した結果と同様の結果を得た。しかし、報酬条件における結果は事前に予測していた学習モデルからの数値シミュレーション結果とは異なる結果を得た。

これらの結果より、本研究では視覚的な感覚情報が与えられる場合は被験者の内部順モデルに影響することで付与された回転へ適応すると考える。これは、先行研究 [5] の結果と同様の提案であり、カルマンフィルター器が観測値と予測値の差である感覚予測誤差によって回転を学習していくというプロセスが内部順モデルの適応を引き起こしている結果となった。しかし、報酬条件においては運動指令ノイズにより偶発的に到達運動に成功し

て報酬を得た場合に，内部順モデルがその誤差によって回転へ適応すると考察した．

以上より，本研究ではこのシミュレーション結果と異なる実験結果となった原因を，誤差条件での経験が影響するため，もしくは報酬条件においても内部順モデルが適応すると考察した．この結果，内部順モデルは視覚的に与えられる情報のみではなく視覚情報が与えられない場合においても，回転への適応を起こす結果を示した．

6.2 今後の展望

6.2.1 実験中でのインストラクションの変化

本研究の行動実験結果では，誤差条件と報酬条件の両条件で手先が右方向へ到達する結果を得た．目標点跳躍課題に関しては，被験者に対して「手先を新しい目標点へ到達させてください．」とインストラクションした．しかし到達運動を 400 試行程繰り返した後に行われる目標点跳躍課題において，被験者が与えられたインストラクションを忘れてしまった可能性がある．被験者は，報酬条件において誤差条件のみに与えられるカーソルを想像してしまうことでインストラクションと異なった運動を起こしてしまった可能性がある．この結果より，カーソルを常に想像して誤差条件のような運動を引き起こしてしまったのではないかと考えられる．これに関しては，被験者に対して「なぜ目標点が跳躍したときに右へ動かしたのですか？」等と内観に対する調査を行わなかったために考察ができない状況である．今後の展望として，このようなインストラクションの変異が起こるような実験パラダイムの計画はするべきではないと述べる．改善策として例えば，本研究のように一人の被験者において誤差条件と報酬条件の両方を実施する事は，前条件の経験を引きずる可能性がある [19]，あるいは実験時間が長く被験者が実験に集中できないと考えられるので，時間的余裕がある実験パラダイムを設定する事があると考えられる．

6.2.2 カウンターバランスの考慮

本研究で行った行動実験は，先にも述べたが図 2.13 に示す様に被験者一人で誤差条件と報酬条件の二つの実験条件で行っている．さらに，全被験者で誤差条件で実験を行った後に報酬条件で実験を行った．多くの場合，行動実験ではカウンターバランスを取るが本研究の実験ではそれが取れていない．カウンターバランスとは釣り合いを取ることを意味し，本研究の場合では以下の点に付いてカウンターバランスを取るべきであると考えられる．

- 回転の付加する方向を反時計回りだけではなく，時計回りでも行うべき
- 目標点の跳躍方向を右方向だけではなく左方向でも行うべき
- 誤差条件から報酬条件の順番で実験を行うのであれば，報酬条件から誤差上条件の順序でも行うべき

しかし、これらのカウンターバランスを取るには問題がある。二つ目の左方向にも跳躍する場合には、跳躍後に生成される運動補正された手先軌跡の結果が各条件で区別がつきにくいという点がある。三つ目については先にも述べたが、報酬条件での運動回転適応は非常に難しいと言う点がある。これらの問題は、被験者一人当たりの実験にかける時間を延ばすことや被験者数を増やす事で解決できる問題が多い。報酬条件から始める場合には、課題の練習を通常より増やして行うことで解決できると考える。

これらの問題を解決するためには、一人の被験者に対して誤差条件か報酬条件かのどちらかのみの実験を行うことや、実施する行動実験のカウンターバランスを被験者一人分でも良いので取ることとするが重要であると考ええる。また、内観の取り方も非常に重要であると考えられる。事前に予想していた結果と同様の結果が得られた場合でも、「なぜそのような運動にしたのですか？」等の感想程度で良いので聴取すべきであると考ええる。これらの問題点を踏まえ、将来的には上記のように改善すべきと本研究では提案する。

付録1：カルマンフィルター器の導出過程

カルマンフィルターは，あるシステムにおいて現在の試行で得た結果と前回の試行で得た結果を入力値に利用することによって，そのシステムを推定する最適手法の1つである．はじめに，カルマンフィルターで用いる記号を定義し，その後にカルマンフィルターの導出について説明する．

6.3 記号の定義

以下に，表 6.1 にてカルマンフィルターで用いる記号を定義する．ここで，試行数は k で示しており，何試行目の値かは，式に上付で (k) と記載しているので確認することができる．また，予測値は $\hat{\cdot}$ で示しており， $\hat{x}$ は予測される状態変数を意味する．

表 6.1 カルマンフィルターで用いる記号

記号	定義
u	運動指令 (Motor command)
p	回転摂動 (Perturbation)
h	手先位置 (Hand position)
c	カーソル位置 (Cursor position)
y	提供されるカーソル位置 (Provided cursor position)
x	状態変数 (ここでは，回転摂動と手先位置)
P	状態変数の共分散
K	カルマンゲイン
n_u	運動指令へのノイズ
n_p	回転摂動へのノイズ
n_h	手先位置へのノイズ
n_y	提供されるカーソル位置へのノイズ

6.4 カルマンフィルターの導出過程

本章では，先行研究 [5] で提案されている学習モデルのカルマンフィルター器の働きを示す計算式群がどう導出されるのかを述べる．本導出は，運動回転適応実験に基づいている．運動回転適応実験時に被験者は，自分の手先位置を視認することができずに代わりに手先位置と対応しているカーソルが与えられる．そのカーソルは自分の手先の運動と対応して運動することが説明され，目標点まで到達運動するように教示される．ここで，視認できない本当の手先位置 h は，

$$h^{(k)} = u^{(k)} + n_h^{(k)} \quad (6.1)$$

で示され，カーソル位置 c は，

$$c^{(k)} = h^{(k)} + p^{(k)} \quad (6.2)$$

となり，被験者に映像として与えられる手先位置 y は，

$$\begin{aligned} y^{(k)} &= c^{(k)} + n_y^{(k)} \\ &= h^{(k)} + p^{(k)} + n_y^{(k)} \\ &= u^{(k)} + p^{(k)} + n_h^{(k)} + n_y^{(k)} \end{aligned} \quad (6.3)$$

と示すことができる．ここで， u は運動指令， p はカーソルに加えられる摂動（回転），そして n はノイズを表す． n_h は， $n_h \sim N(0, \sigma_h^2)$ であり，また n_y は， $n_y \sim N(0, \sigma_y^2)$ である．

これらの式より，(6.4) 式にヒトの内部順モデルが予測する手先位置 $\hat{h}$ と回転摂動 $\hat{p}$ の計算式を示す．ここで文字上部についている $\hat{}$ は予測値のことを示しており，ここでは脳内の内部順モデルが予測する値のことを指している．

$$\hat{h}^{(k+1)} = \hat{p}^{(k)} + u^{(k)} \quad (6.4)$$

$$\hat{p}^{(k+1)} = a\hat{p}^{(k)} + n_p^{(k)} \quad (6.5)$$

ここで， a は係数を示しており，そして n_p はノイズであり $n_p \sim N(0, \sigma_p^2)$ である．このノイズによって，予測される回転摂動を更新していく．また，試行数が $k+1$ となっているがこれは現試行の出力結果が前試行時の結果を用いて算出されていることを示しており，

予測を表している．ここで，(6.4) 式のノイズが無くなっているのはこれらの値は脳内で予測している値であるのでノイズは付加されないためである．

回転摂動と手先位置を，状態方程式 (State Space Model) によって定義すると，

$$\mathbf{x}^{(k+1)} = A\mathbf{x}^{(k)} + \mathbf{b}u^{(k)} + \mathbf{n}_x^{(k)} \quad (6.6)$$

$$\begin{aligned} \mathbf{x}^{(k)} &= [p^{(k)} \quad h^{(k)}]^T \\ A &= \begin{bmatrix} a & 0 \\ 1 & 0 \end{bmatrix} \\ \mathbf{b} &= \begin{bmatrix} 0 & 1 \end{bmatrix}^T \\ \mathbf{n}_x &= \begin{bmatrix} n_p^{(k)} & n_h^{(k)} \end{bmatrix}^T \end{aligned}$$

となる．また，被験者に映像として与えられる手先位置 y は，

$$y^{(k)} = C\mathbf{x}^{(k)} + n_y^{(k)} \quad (6.7)$$

$$C = \begin{bmatrix} 0 & 1 \end{bmatrix}$$

と再定義される．

この(6.6) 式と(6.7) 式によりカルマンフィルターは推定(*Prediction*)と補正(*Filtering*)の二つのプロセスを用いて，システムを最適な状況へ推移させる．ここでの推定式は，

$$\hat{\mathbf{x}}^{(k+1|k)} = A\hat{\mathbf{x}}^{(k|k)} + \mathbf{b}u^{(k)} \quad (6.8)$$

となる． $(k+1|k)$ では，左側に現在の試行を示して右側は利用された試行のデータを示す．この状態方程式の共分散は，

$$\begin{aligned} P^{(k+1|k)} &= \text{Var}[\hat{\mathbf{x}}^{(k+1|k)}] \\ &= E[(\mathbf{x}^{(k+1)} - \hat{\mathbf{x}}^{(k+1|k)})(\mathbf{x}^{(k+1)} - \hat{\mathbf{x}}^{(k+1|k)})^T] \\ &= AE[(\mathbf{x}^{(k)} - \hat{\mathbf{x}}^{(k|k)})(\mathbf{x}^{(k)} - \hat{\mathbf{x}}^{(k|k)})^T]A^T + \Omega_x \\ &= AP^{(k|k)}A^T + \Omega_x \end{aligned} \quad (6.9)$$

$$\Omega_x = \begin{bmatrix} \sigma_p^2 & 0 \\ 0 & \sigma_h^2 \end{bmatrix}$$

となる．また，補正式は観測値と予測値の差を利用して補正して次の試行の値を生成するので，カルマンゲインを K とすると，

$$\begin{aligned}
 \hat{\mathbf{x}}^{(k|k)} &= \hat{\mathbf{x}}^{(k|k-1)} + K^{(k)}(y^{(k)} - \hat{y}^{(k)}) \\
 &= \hat{\mathbf{x}}^{(k|k-1)} + K^{(k)}(C\mathbf{x}^{(k)} - C\hat{\mathbf{x}}^{(k|k-1)} + n_y^{(k)}) \\
 &= (I - K^{(k)}C)\hat{\mathbf{x}}^{(k|k)} + K^{(k)}(C\mathbf{x}^{(k)} + n_y^{(k)})
 \end{aligned} \tag{6.10}$$

となり，この補正式の共分散は，

$$\begin{aligned}
 P^{(k|k)} &= E[(\mathbf{x}^{(k+1)} - \hat{\mathbf{x}}^{(k+1|k)})(\mathbf{x}^{(k+1)} - \hat{\mathbf{x}}^{(k+1|k)})^T] \\
 &= (I - K^{(k)}C)E[(\mathbf{x}^{(k)} - \hat{\mathbf{x}}^{(k|k-1)})(\mathbf{x}^{(k)} - \hat{\mathbf{x}}^{(k|k-1)})^T](I - K^{(k)}C)^T + K^{(k)}\sigma_y^2 K^{(k)T} \\
 &= (I - K^{(k)}C)P^{(k|k-1)}(I - K^{(k)}C)^T + K^{(k)}\sigma_y^2 K^{(k)T}
 \end{aligned} \tag{6.11}$$

となる．補正式には，分散を最小にすることでシステムを安定状態へ遷移させる必要がある．そこで，分散を最小にする条件でのカルマンゲイン K を求める．分散の細小にするために，対角和を計算してその値を最小にする．分散 $P^{(k|k)}$ の対角和は，

$$\begin{aligned}
 Tr[P^{(k|k)}] &= Tr[(I - K^{(k)}C)P^{(k|k-1)}(I - K^{(k)}C)^T + K^{(k)}\sigma_y^2 K^{(k)T}] \\
 &= Tr[(I - K^{(k)}C)(I - K^{(k)}C)^T P^{(k|k-1)}] + Tr[K^{(k)T} K^{(k)}]\sigma_y^2 \\
 &= Tr[P^{(k|k)}] - 2Tr[K^{(k)T} P^{(k|k)} C^T] + Tr[K^{(k)T} C^T P^{(k|k)} C K^{(k)}] + K^{(k)T} K^{(k)}\sigma_y^2 \\
 &= Tr[P^{(k|k)}] - 2K^{(k)T} P^{(k|k)} C^T + K^{(k)T} K^{(k)} C^T P^{(k|k)} C + K^{(k)T} K^{(k)}\sigma_y^2
 \end{aligned} \tag{6.12}$$

となる．この共分散の対角和値を最小にするために，(6.12) 式 をカルマンゲインで微分する．

$$\begin{aligned}
 \frac{d}{dK^{(k)}}(Tr[P^{(k|k)}]) &= -2P^{(k|k-1)}C^T + 2K^{(k)}CP^{(k|k-1)}C^T + 2K^{(k)}\sigma_y^2 \\
 &= 0
 \end{aligned} \tag{6.13}$$

$$\begin{aligned}
 K^{(k)} &= P^{(k|k)}C^T(CP^{(k|k)}C^T + \sigma_y^2)^{-1} \\
 &= \frac{P^{(k|k)}C^T}{(CP^{(k|k)}C^T + \sigma_y^2)}
 \end{aligned} \tag{6.14}$$

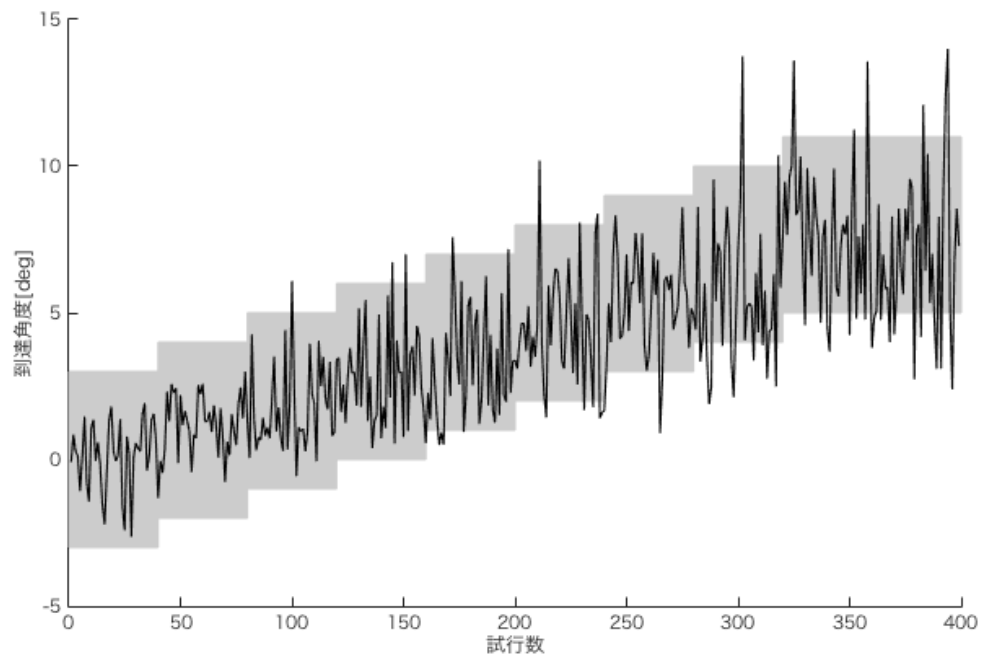
(6.14) 式の結果を (6.11) 式に代入すると，補正式の共分散は，

$$\begin{aligned} P^{(k|k)} &= (I - K^{(k)}C)P^{(k|k-1)}(I - K^{(k)}C)^T + K^{(k)}\sigma_y^2 K^{(k)T} \\ &= (I - K^{(k)}C)P^{(k|k-1)} \end{aligned} \quad (6.15)$$

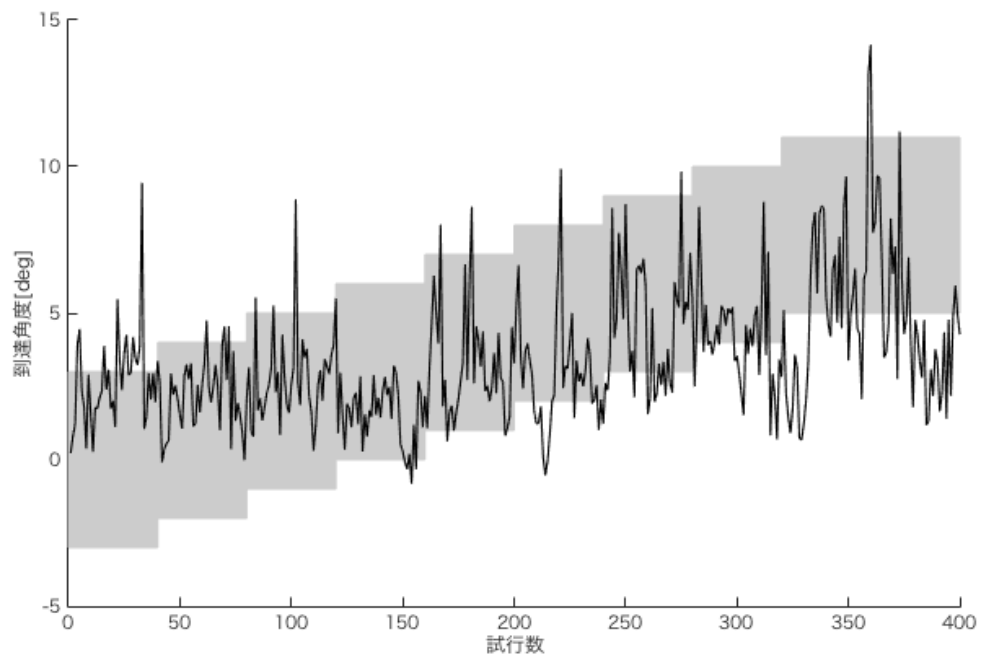
と再定義される．この結果は，第 5 章に述べた (3.7) 式のカルマンゲインと共分散の関係を示している．

付録2：被験者実験結果

ここでは、5章にて述べた行動実験の運動回転適応実験結果において載せていない結果を示す。また、実験条件は表 5.1 に準ずる。ここで、被験者 06 と被験者 08 の報酬条件の学習曲線が途中で切れているが、これは都合により途中で実験を中止したためである。

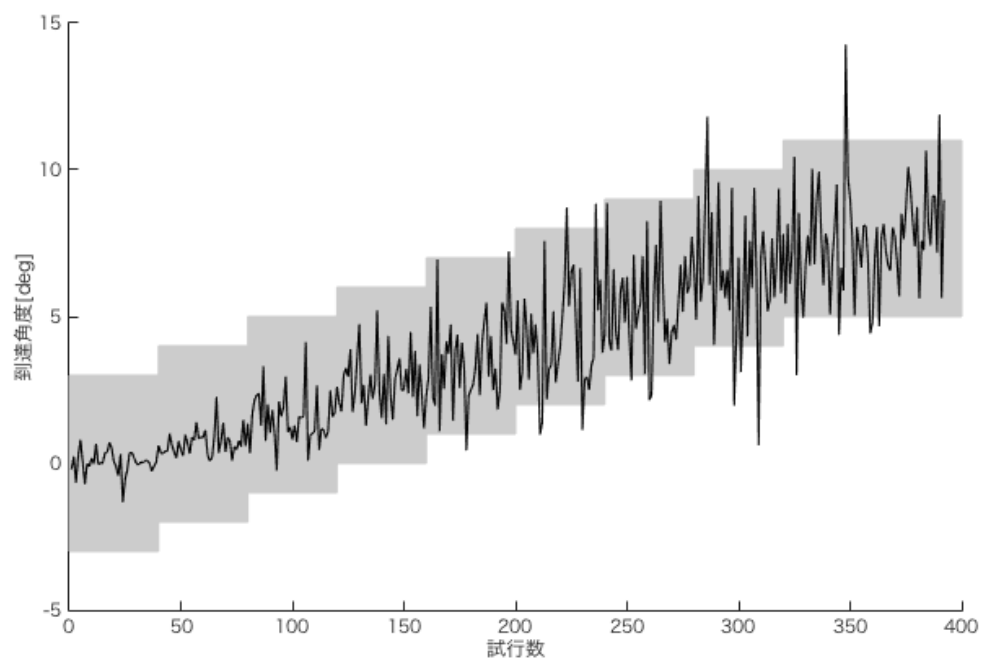


a. 誤差条件

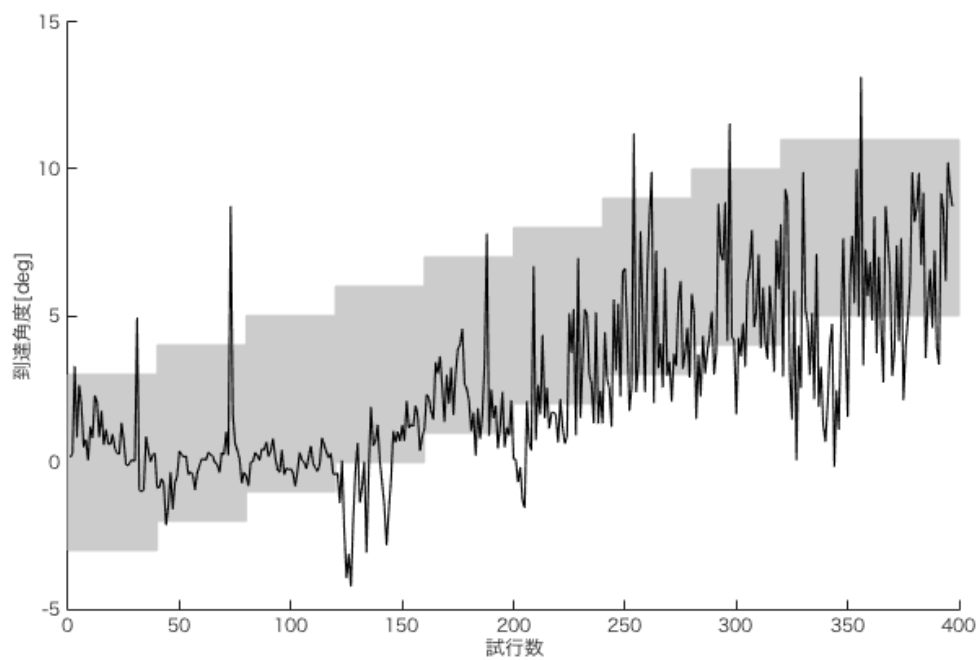


b. 報酬条件

図 6.1 被験者実験による運動回転適応実験結果：被験者 02

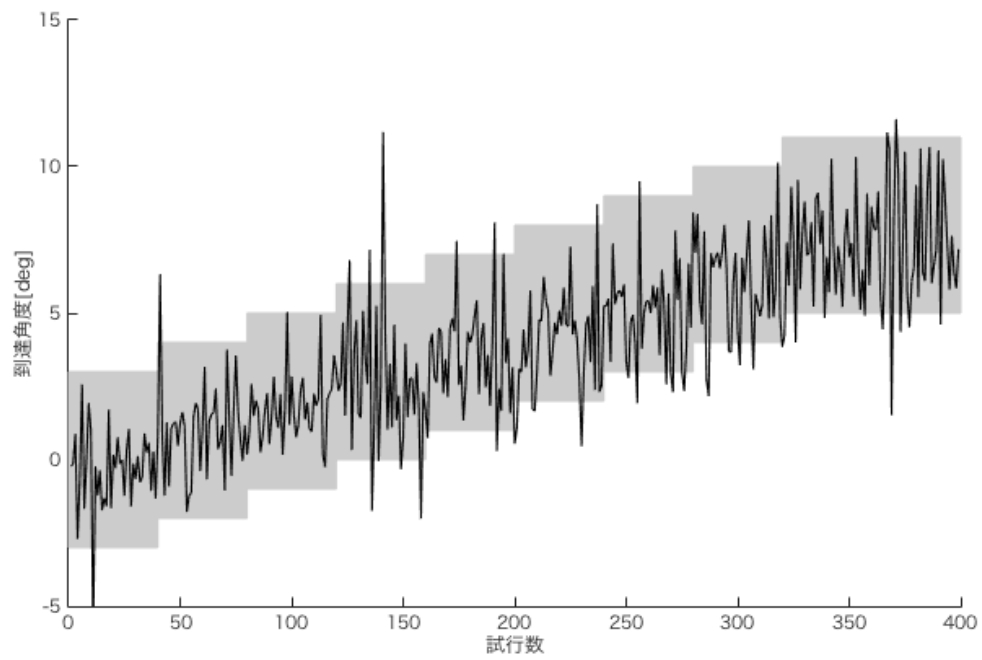


a. 誤差条件

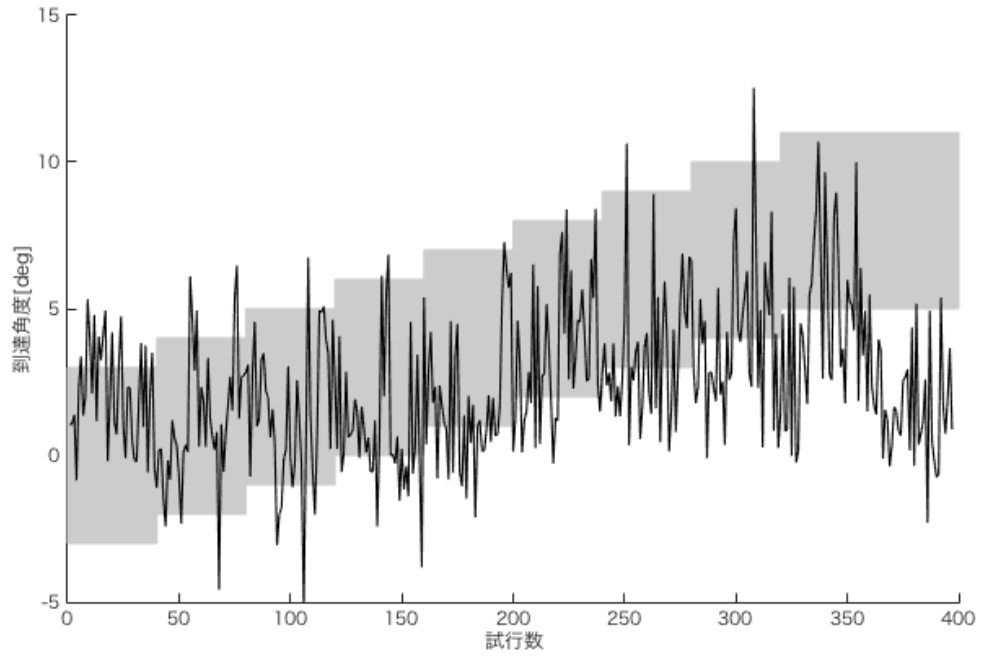


b. 報酬条件

図 6.2 被験者実験による運動回転適応実験結果：被験者 03

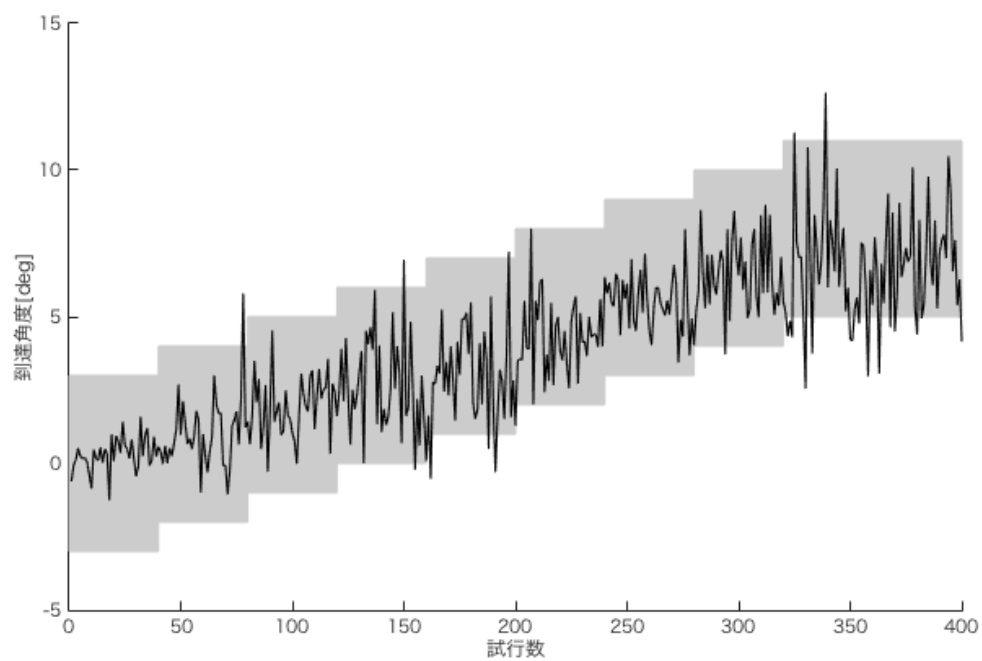


a. 誤差条件

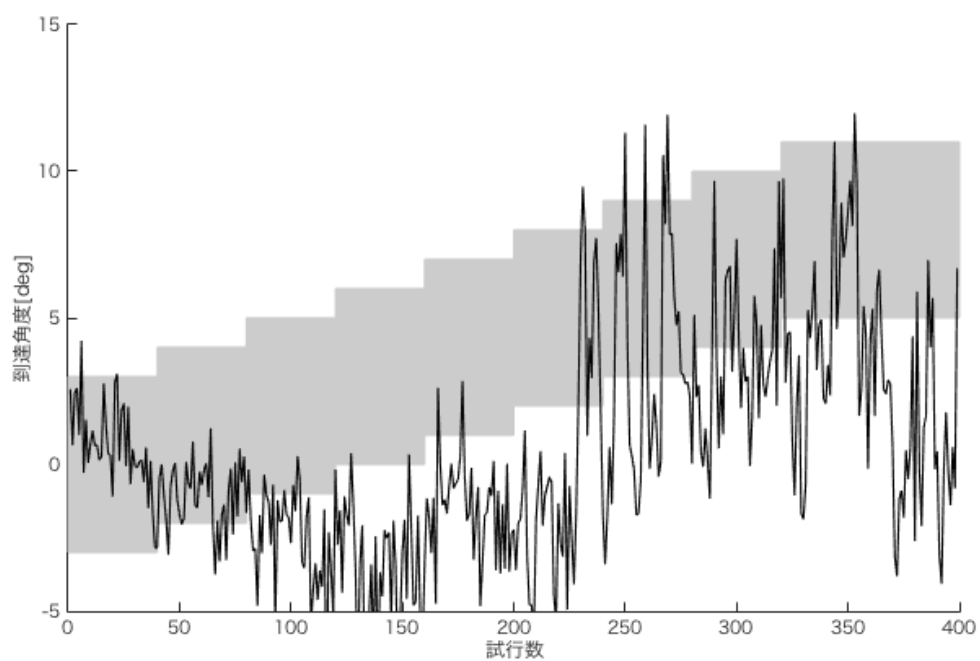


b. 報酬条件

図 6.3 被験者実験による運動回転適応実験結果：被験者 04

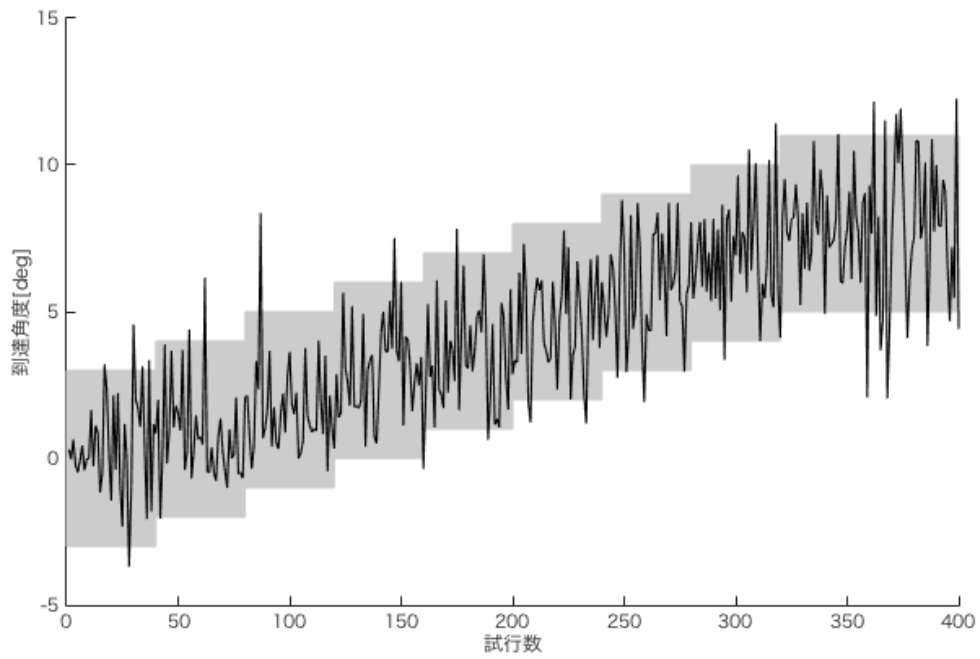


a. 誤差条件

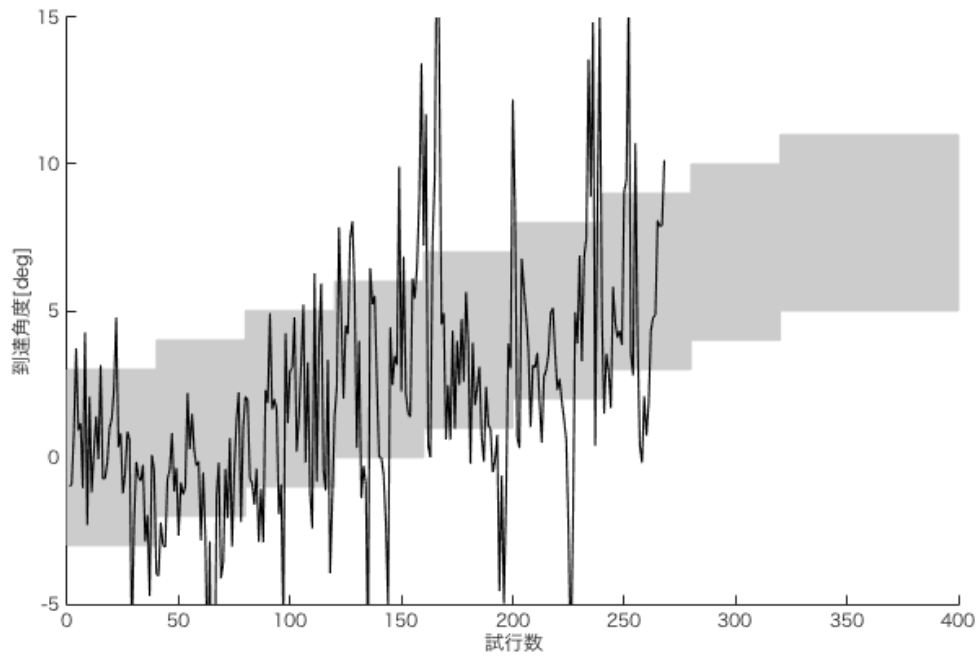


b. 報酬条件

図 6.4 被験者実験による運動回転適応実験結果：被験者 05

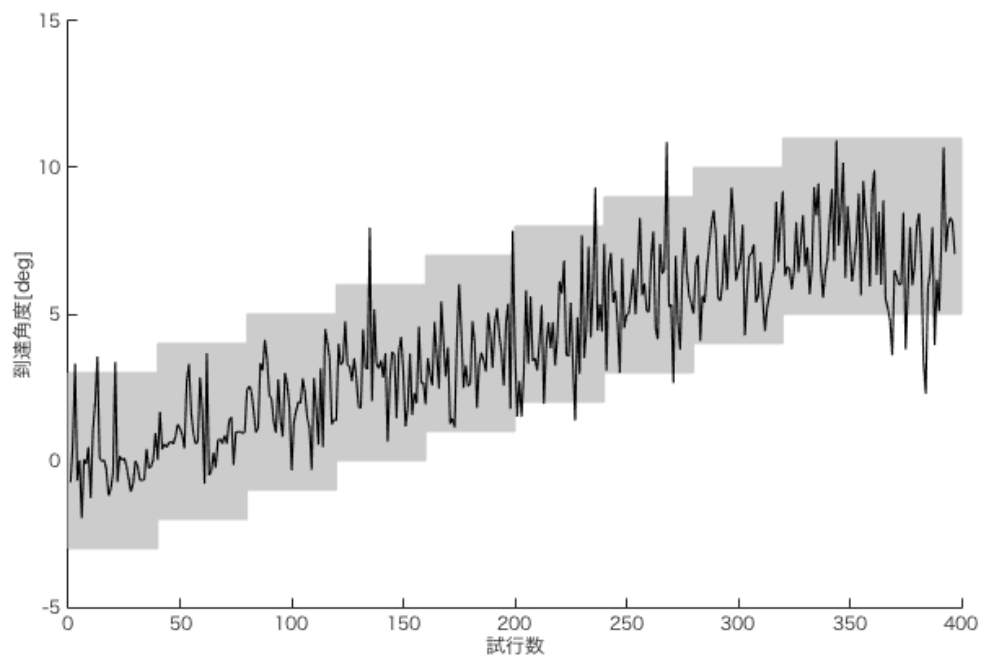


a. 誤差条件

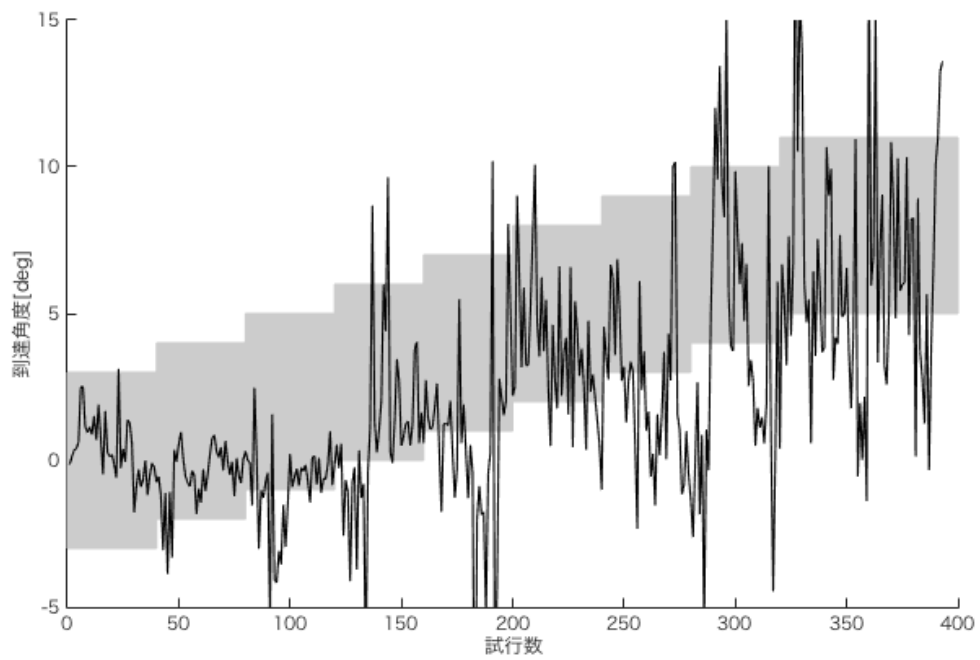


b. 報酬条件

図 6.5 被験者実験による運動回転適応実験結果：被験者 06

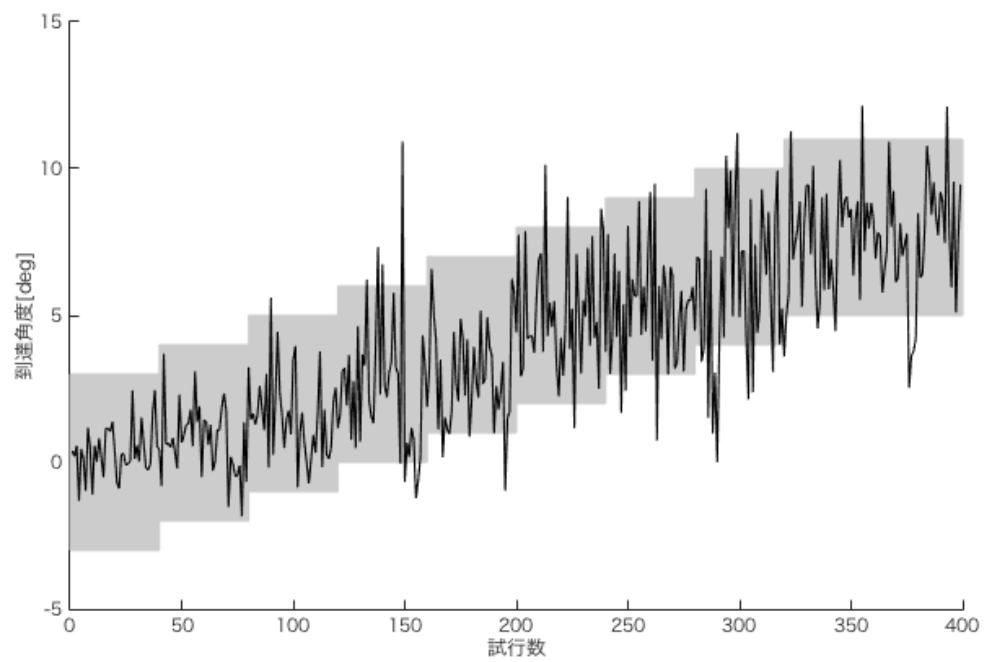


a. 誤差条件

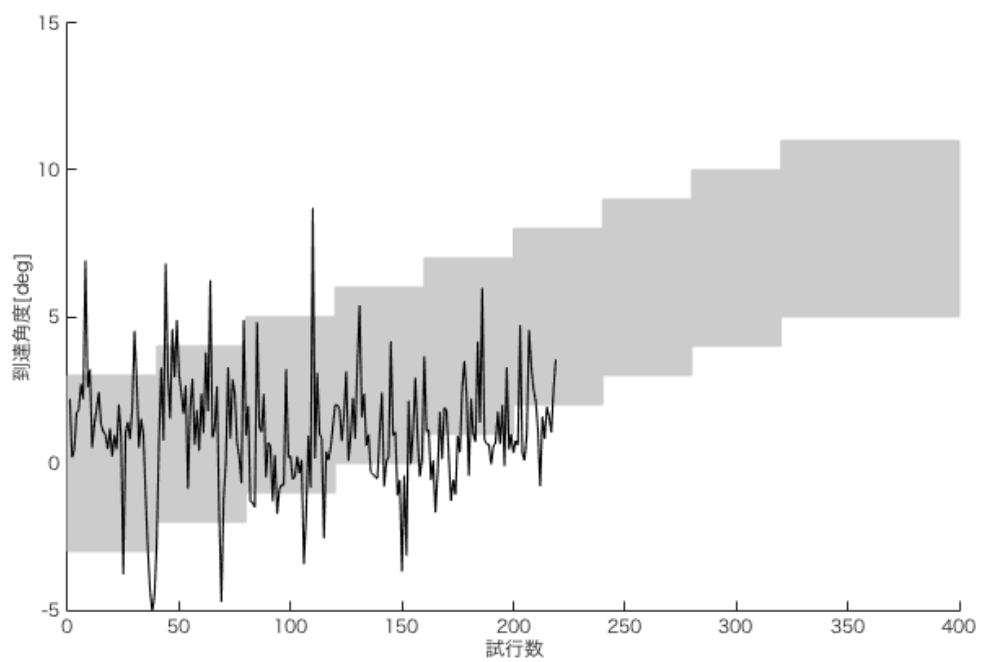


b. 報酬条件

図 6.6 被験者実験による運動回転適応実験結果：被験者 07



a. 誤差条件



b. 報酬条件

図 6.7 被験者実験による運動回転適応実験結果：被験者 08

付録3：確認実験

本研究では，被験者実験を行う前に提案した実験が行う事ができるのか，どのような結果なのかを確認するために予備的に確認実験を行った．この確認実験は，先に予測していた目標点跳躍課題の結果を確認し，被験者を募集し実験する前に実験計画に不足や問題点が存在しないか確認するためである．以下に，確認実験の条件および実験結果を示す．

6.5 運動回転適応実験の結果

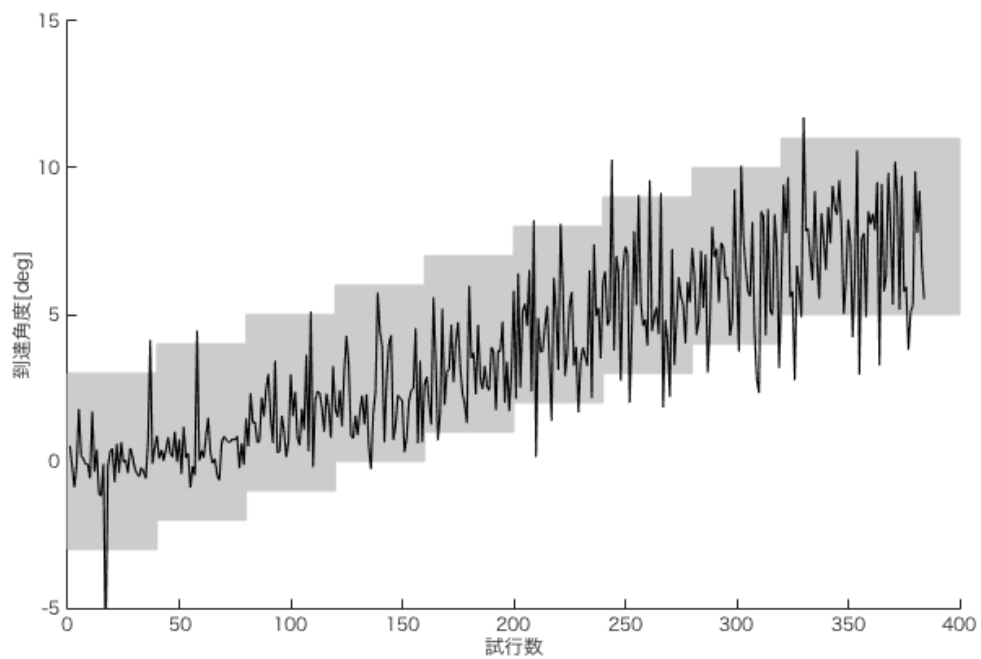
以下に表 6.2 として実験条件を示す．この条件は先行研究 [5] に準ずる．以下の図 6.8 から図 6.10 に，被験者 3 名の確認実験における運動回転適応実験の結果として学習曲線を示す．図において学習曲線は a. が誤差条件，b. が報酬条件での結果となっている．横軸に試行数，縦軸に到達位置と運動開始点とがなす角度として定義している到達角度を示している．また，灰色の範囲は到達範囲を示しておりその範囲に到達運動すれば，目標点へ到達できたとした．結果より，誤差条件においては全被験者とも回転への適応を示しており，報酬条件においては分散が非常に高いが段階的に付与される回転への適応が示された．報酬条件においては，ほぼ全ての被験者が誤差条件とは異なり自分の手先を明らかに右側へずらして到達させたという内観を得た．

表 6.2 運動回転適応実験の条件（確認実験）

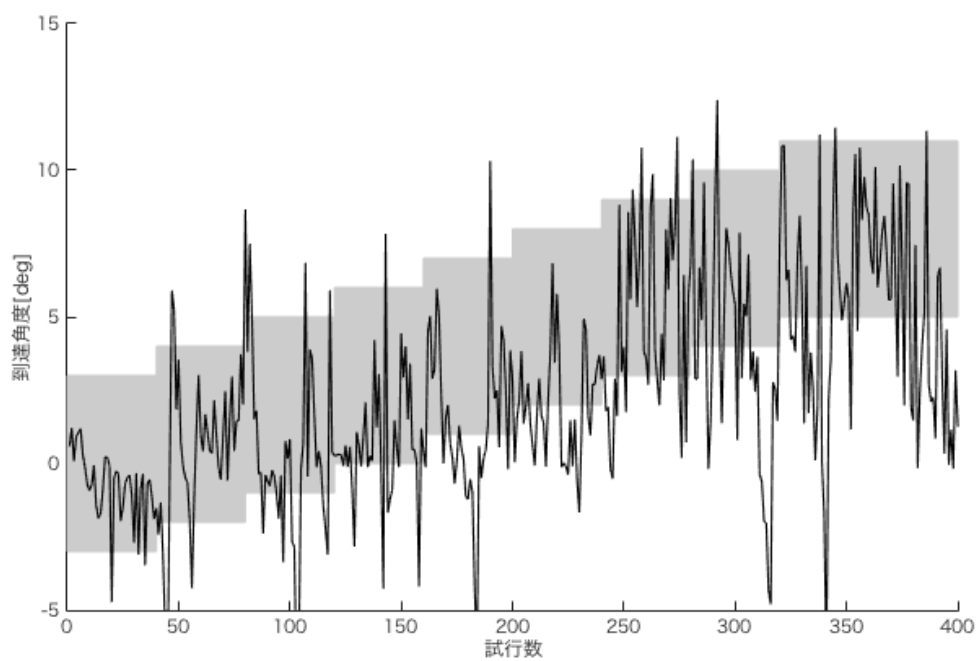
被験者数	3 名/各班
試行数	400 [試行]
到達距離	100 [mm]
回転方向	反時計回り
最大回転	8 [deg]
到達範囲	± 3 [deg]
回転付与	+1 [deg]/40 [試行]

表 6.3 目標点跳躍課題の条件（確認実験）

試行数	10 [試行]
跳躍方向	右方向
跳躍距離	15 [mm]

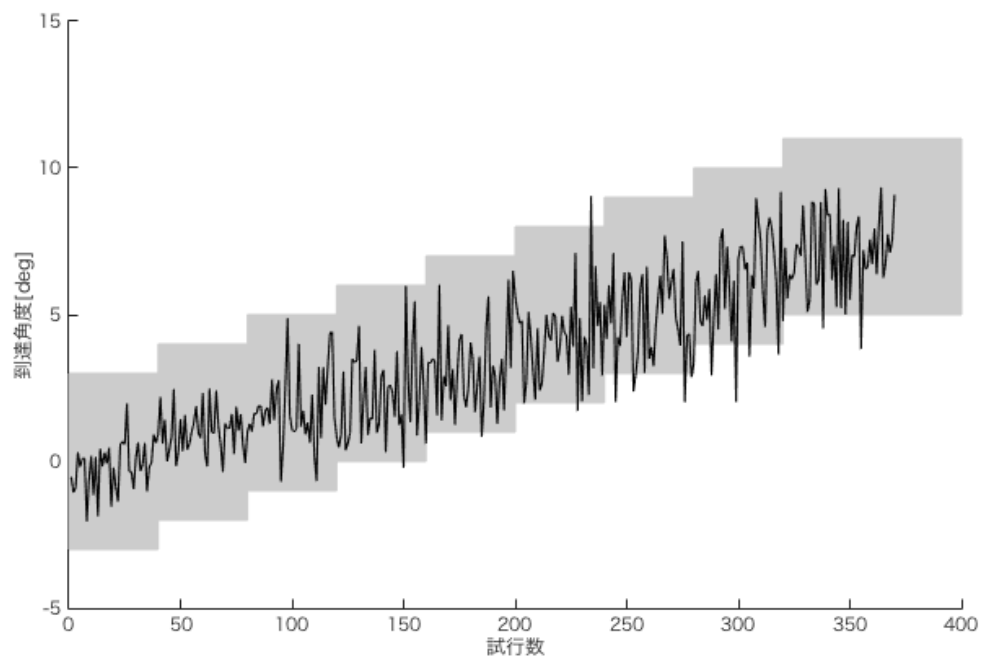


a. 誤差条件

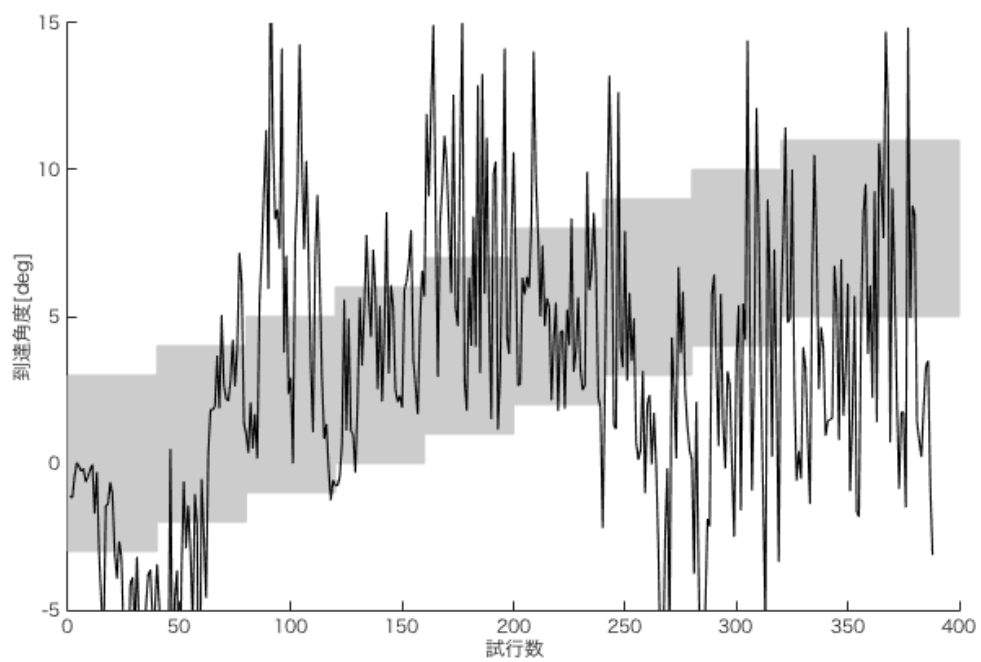


b. 報酬条件

図 6.8 運動回転適応実験結果（確認実験）：被験者 a

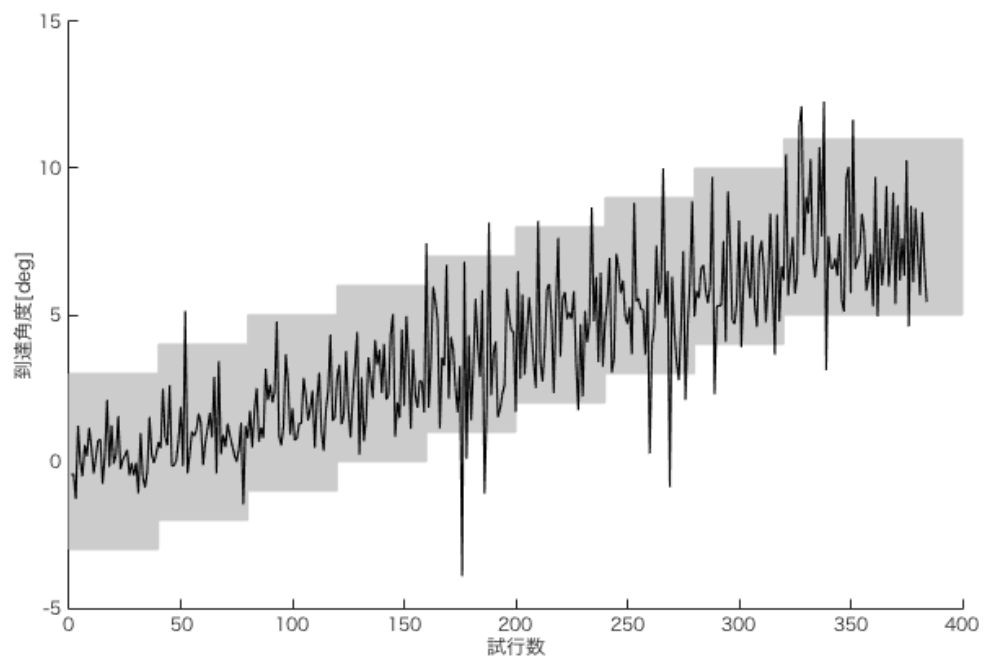


a. 誤差条件

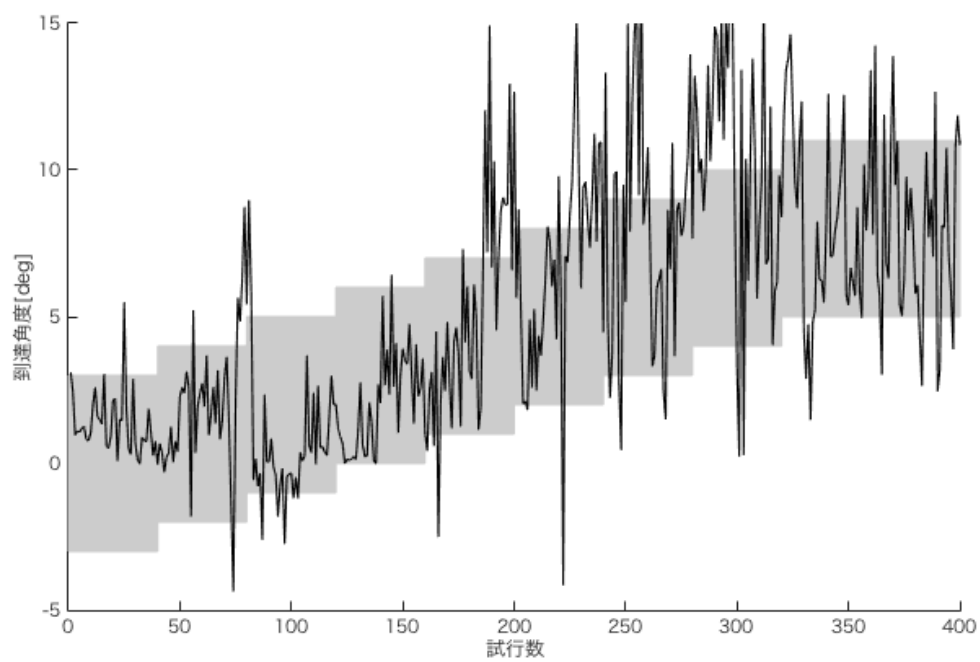


b. 報酬条件

図 6.9 運動回転適応実験結果 (確認実験): 被験者 b



a. 誤差条件



b. 報酬条件

図 6.10 運動回転適応実験結果（確認実験）：被験者 c

6.6 目標点到達課題の結果

以下に、表 6.3 として実験条件を示す。図 2.13 に示すように、目標点到達課題は回転が付与されていない時と最大回転に適応した後の 2 セッション行う。このときの跳躍距離は、8 [deg] の回転に適応した場合の到達位置と同一位置の右方向へ 15 [mm] の跳躍となっている。以下に、被験者 3 名分の確認実験における目標点跳躍試行結果を示す。この結果は誤差条件と報酬条件の両条件において、3 名分の被験者の回転適応後の目標点跳躍課題の平均した軌跡である。平均に用いたデータ数は誤差条件で 24 データ、報酬条件で 27 データである。これは、計測開始時に既に運動しているまたは運動距離が 80 [mm] 以下である場合は何かしらの問題により課題がこなせていないと判断して破棄しているためである。手先軌跡の青線が誤差条件、赤線が報酬条件の平均軌跡であり、緑丸が跳躍した目標点の位置である。灰色の領域は標準誤差であり、各被験者のすべての到達位置から求めた。また、誤差条件と報酬条件それぞれの到達位置には有意な差が示された（2 標本 t 検定, $p < 1.347 \times 10^{-10}$ ）。図 2.1 や図 2.11 で示しているように、本研究の到達運動においては横方向は x 軸、奥行き方向が y 軸としている。図 6.11 の縦軸（ y 軸）と横軸（ x 軸）はそれに対応している。

6.7 確認実験結果の考察

まず、確認実験の目標点跳躍課題において運動補正が行われているかを確認する必要がある。そこで、目標点跳躍課題において運動補正が起こっているのかを、非跳躍時と比較して検証する。

6.7.1 考察：目標点の跳躍時には運動補正が生成されているかの検討

以下に、図 6.12 から 6.14 に 8 [deg] の最大回転に適応した後における目標点の跳躍時と非跳躍時の到達運動の手先軌跡と手先速度を示す。ここで、黒線が非跳躍時の平均データであり、青線が誤差条件での平均データ、赤線が報酬条件での平均データを示しており、この色は図 6.12 から 6.14 の三図に共通している。また、各色で半透明で示している範囲は標準偏差である。図 6.12 からは、黒線で示す非跳躍時の軌跡と各色の線の軌跡の到達位置が異なっていることがわかる。また、図 6.13 および図 6.14 で示す速度では、補正がかかる x 軸方向の速度（両図左側）が、非跳躍時の速度はピークが 1 つなのに対して跳躍時の速度はピークが 2 つ存在することがわかる。これら結果から、誤差条件および報酬条件での目標点到達課題では、目標点の跳躍によって運動補正が行われており、その補正は異なる結果が示した。

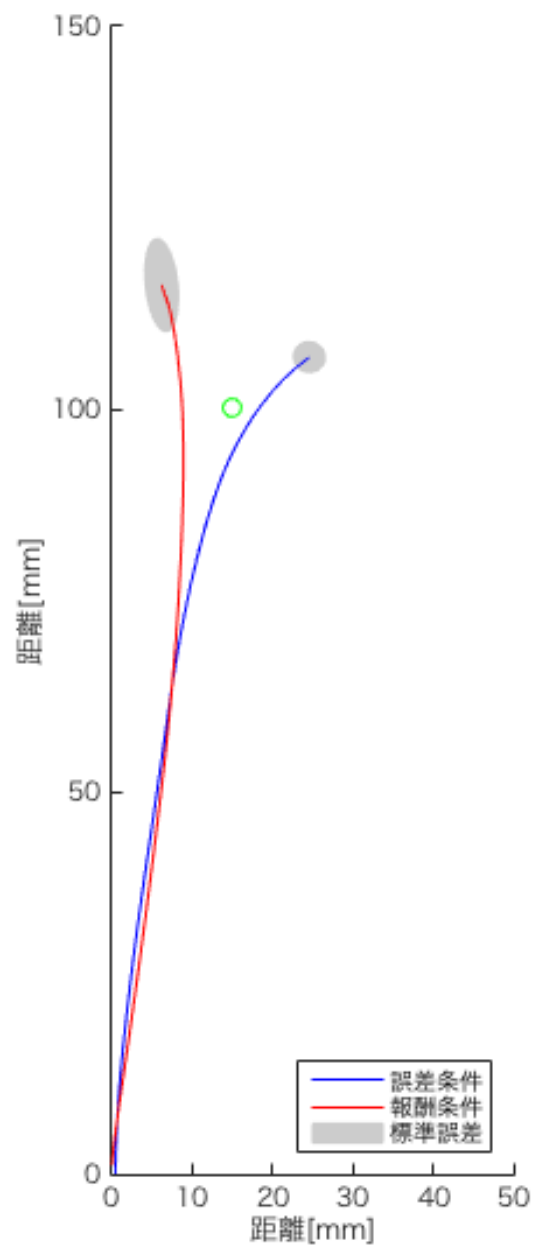
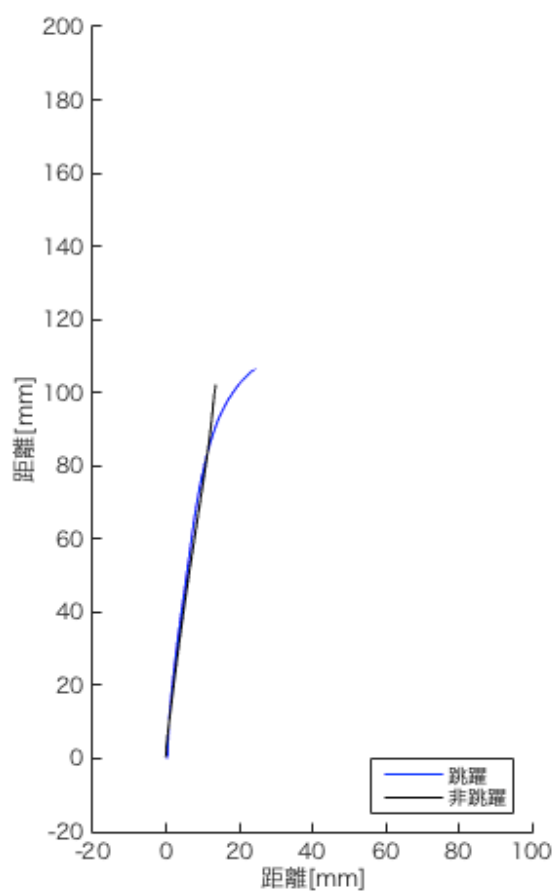
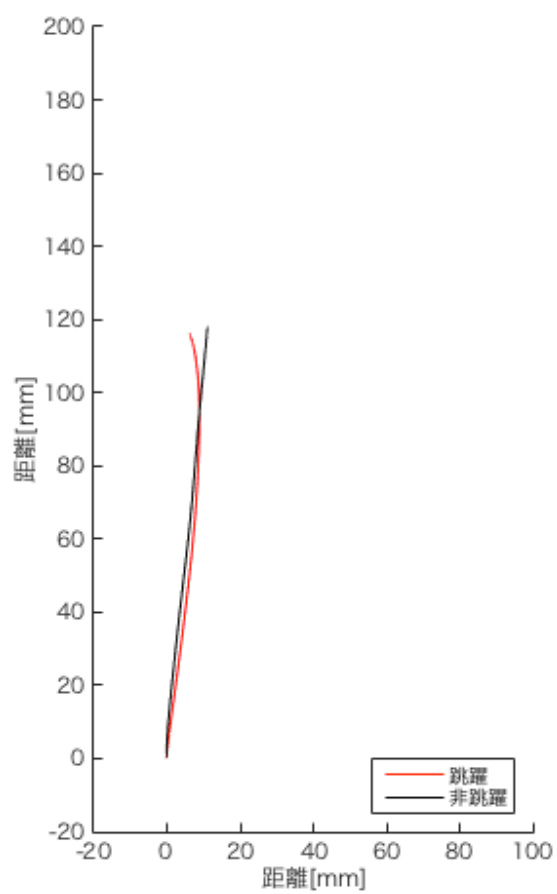


図 6.11 確認実験における目標点跳躍課題結果

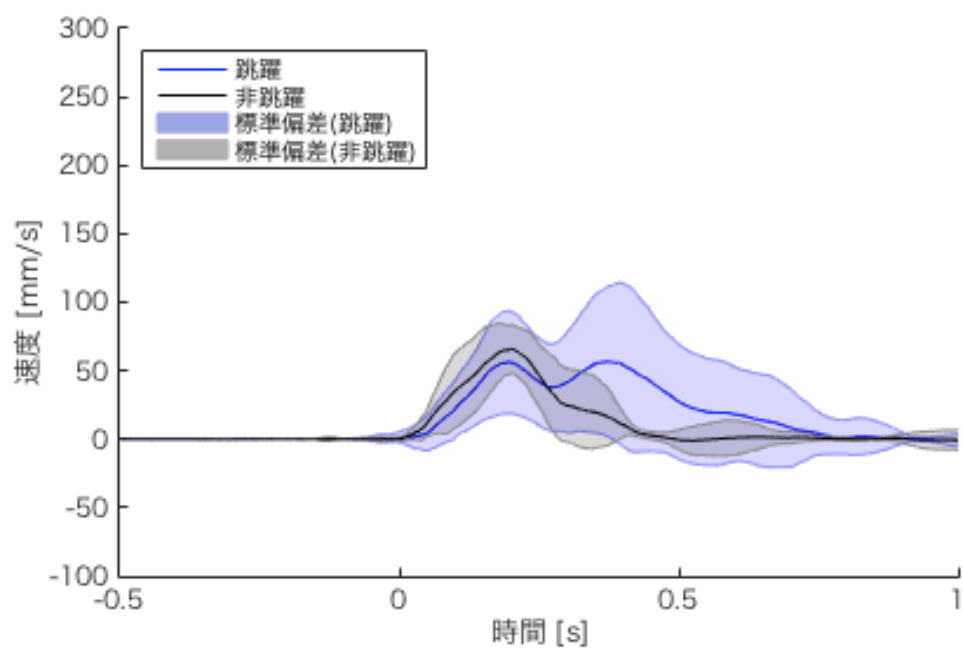


a. 誤差条件

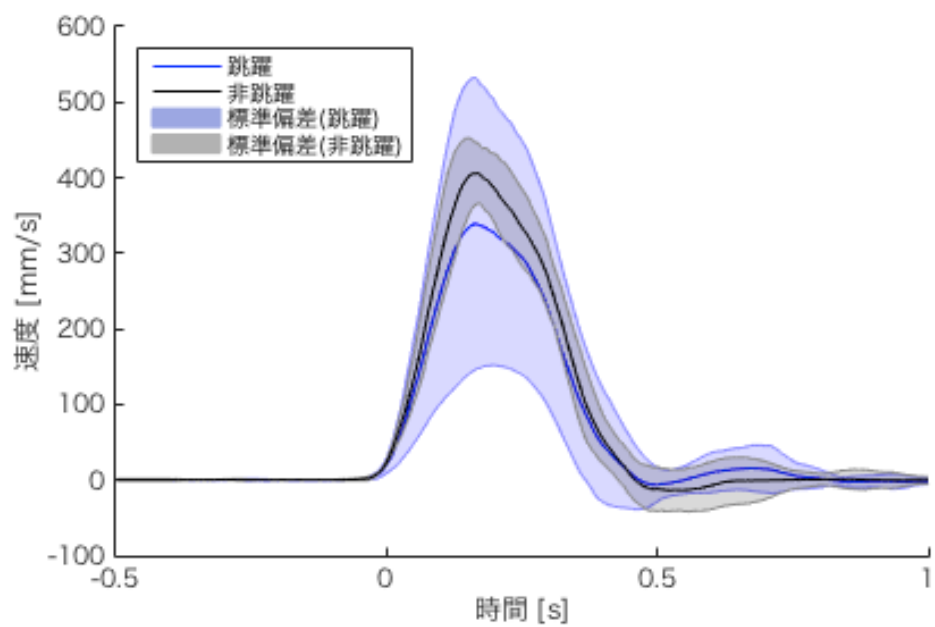


b. 報酬条件

図 6.12 跳躍時と非跳躍時における軌跡の比較

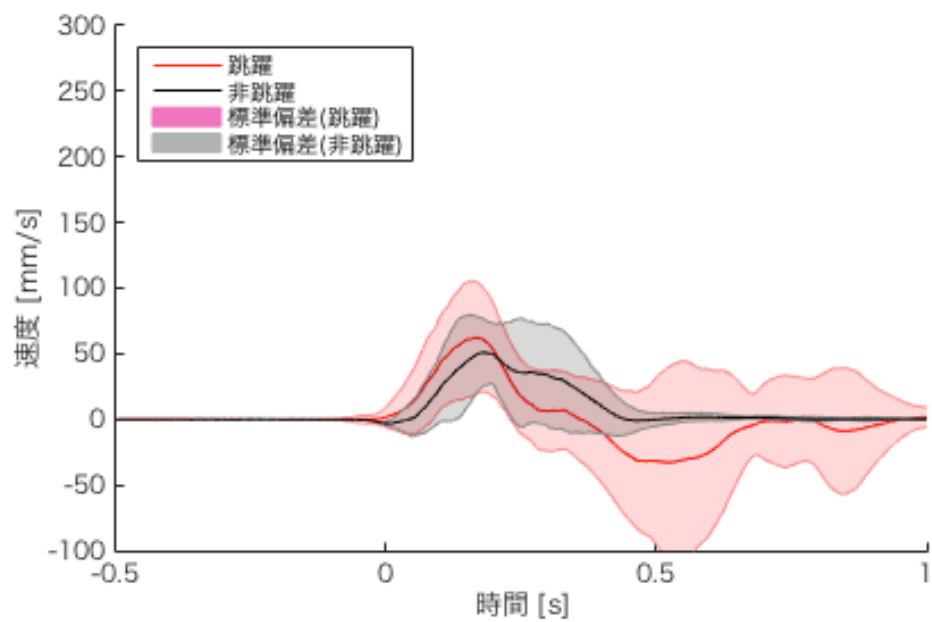


a. 誤差条件: x 方向速度

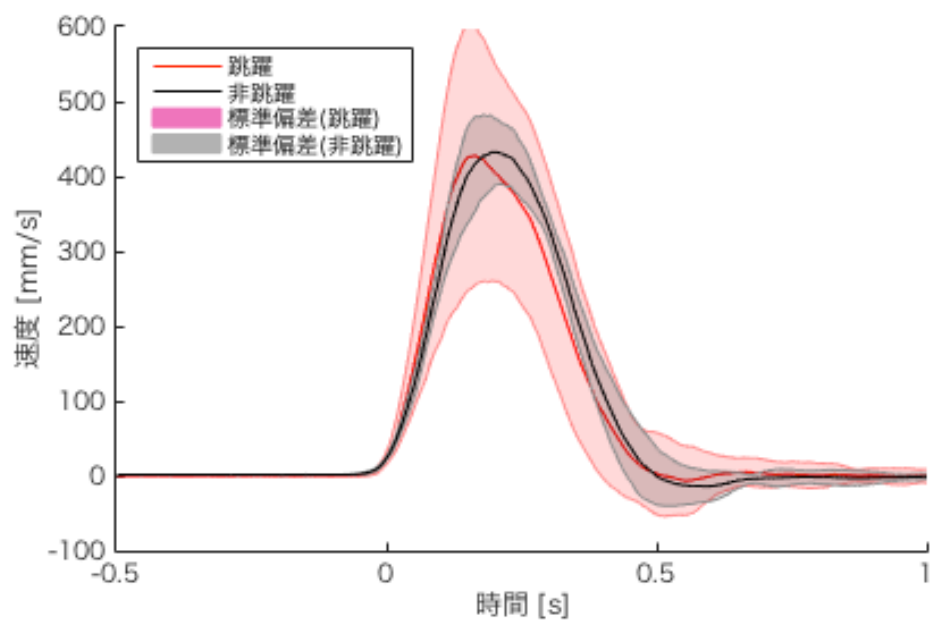


b. 誤差条件: y 方向速度

図 6.13 誤差条件における跳躍時と非跳躍時の速度比較: 15 [mm] 跳躍



a. 報酬条件: x 方向速度



b. 報酬条件: y 方向速度

図 6.14 報酬条件における跳躍時と非跳躍時の速度比較：15 [mm] 跳躍

6.7.2 考察：確認実験の結果は数値シミュレーションと異なる結果

本確認実験の結果と予測していた予測結果とは異なる結果を得た．報酬条件において目標点が右側に 15 [mm] に跳躍した行動実験結果は図 2.12 に示した予測結果とは異なり，新しい目標点へは到達しなかった．この行動実験の結果は，15 [mm] ではなく 7 [mm] の跳躍時の予測結果とほぼ同一の結果であった．この比較は，シミュレーション結果からも確認できる．この比較結果より，なぜ予測結果と異なる確認実験結果となったかを考察する．この異なった原因としては，被験者 3 人がすべて予測される結果を知っていた事が大きく影響しているのではないかと影響する．全被験者は事前に，目標点が右側に 7 [mm] の跳躍を行う事を知らされていた．しかし，実際の実験は 15 [mm] に跳躍を実施した．この結果，7 [mm] 跳躍時の予測結果を知っていることにより 15 [mm] 跳躍の実験に影響してしまったのではないかと考える．この結果を踏まえ，被験者実験では 7 [mm] と 15 [mm] の跳躍距離において目標点跳躍課題を実施した．その結果は，図 5.4 に述べているが，確認実験結果とは一致しなかった．よって，本確認実験では与えられた事前情報が運動結果に影響を与える可能性があると考えられる．この結果は，認知系の処理が運動結果に影響する事を述べている結果でもある．よって，内部順モデルの適応状況の確認は直接的に確認する方法をとるべきであると言え，ナイーブな被験者を取る必要性が確認できた．

参考文献

- [1] Wolpert D. M., Ghahramani Z., Jordan M. I., “ An Internal Model for Sensorimotor Integration. ”, Science, Vol.269, No. 5232, 1880 - 1882, 1995.
- [2] Kawato M., “ Internal models for motor control and trajectory planning. ”, Neuroreport 9, 718 - 727, 1999.
- [3] Imamizu H., Miyauchi S., Tamada T., Sasaki Y., Takino R., Puetz B., Yoshioka T., Kawato M., “ Human cerebellar activity reflecting an acquired internal model of a new tool. ”, Nature, Vol 403, 192 - 195, 2000.
- [4] Krakauer J. W., Pine Z. M., Ghilardi M. F., Ghez C., “ Learning of Visuomotor Transformation for Vectorial Planning of Reaching Trajectories. ”, The Journal of Neuroscience, Vol.20, Issue 23, 8916 - 8924, 2000.
- [5] Izawa J., Shadmehr R., “ Learning from Sensory and Reward Prediction Errors during Motor Adaptation. ”, Plos Computational Biology, Vol. 7, Issue 3, e1003377, 2011.
- [6] Cunningham H. A., “ Aiming error under transformed spatial mappings suggests a structure for visual-motor maps. ”, Journal of Experimental Psychology, Vol. 15, No. 3, 493 - 506, 1989.
- [7] Mazzoni P., Krakauer J. W., “ An implicit plan overrides an explicit strategy during visuomotor adaptation. ”, The Journal of Neuroscience, Vol.26, Issue 14, 3642 - 3645, 2006.
- [8] Goodale M. A., Milner A. D., “ Separate visual pathways for perception and action. ”, Trends in Neurosciences, Vol. 15, No. 1, 20 - 25, 1992.
- [9] Aglioti S., DeSouza J. F., Goodale M. A., “ Size-contrast illusions deceive the eye but not the hand. ”, Current Biology, Vol. 5, No. 6, 679 - 685, 1995.
- [10] Ganel T., Tanzer M., Goodale M. A., “ A double dissociation between action and perception in the context of visual illusions: opposite effects of real and illusory size. ”, Psychological science, Vol. 19, 221 - 225, 2008.

- [11] Flanagan J. R., Beltzner M. A., “ Independence of perceptual and sensorimotor predictions in the size-weight illusion. ”, *Nature Neuroscience*, Vol. 3, No. 7, 737 - 741, 2000.
- [12] Brayanov J. B., Smith M. A., “ Bayesian and “ Anti-Bayesian ” Biases in Sensory Integration for Action and Perception in the Size-Weight Illusion ”, *Journal Neurophysiology*, Vol. 103: 1518 - 1531, 2010.
- [13] 橋本 佳子, “ 運動スキル学習に関する考察 ―脳内経路の変化と記憶の固定をめぐって― ”, *新潟工科大学研究紀要*, 第 12 号, 133 - 147, 2007.
- [14] Braun D. A., Aertsen A., Wolpert D. M., “ Motor task variation induces structural learning. ”, *Current Biology*, Vol. 19, No. 4, 352 - 357, 2009.
- [15] 西條 直樹, 五味 裕章, “ 手先回転移動変換の変化に応じた腕到達運動の学習戦略 ”, *電気情報通信学会論文誌*, Vol. J92-D, No. 4, 552 - 561, 2009.
- [16] Telgen S., Parvin D., Diedrichsen J., “ Mirror reversal and visual rotation are learned and consolidated via separate mechanisms: recalibrating or learning de novo?. ”, *The Journal of Neuroscience*, Vol.20, Issue 23, 8916 - 8924, 2000.
- [17] Lillicrap T. P., Brisenno P. M., Diaz R., Tweed D. B., Troje N. F., Ruiz J. F., “ Adapting to inversion of the visual field: a new twist on an old problem. ”, *Experimental Brain Research*, Vol.228, 327 - 339, 2013.
- [18] Krakauer, J. W., Ghilardi, M. F., and Ghez, C, “ Independent learning of internal models for kinematic and dynamic control of reaching. ”, *Nature Neuroscience*, Vol.2, No. 11, 1999.
- [19] Herzfeld D. J., Vaswani P. A., Marko M. K., Shadmehr R., “ A memory of errors in sensorimotor learning. ”, *Science*, Vol.34 5, Issue 6202, 2014.
- [20] Izawa J., Rane T., Donchin O., Shadmehr R., “ Motor Adaptation as a Process of Reoptimization. ”, *The Journal of Neuroscience*, Vol.28, Issue 11, 2883 - 2891, 2008.
- [21] Blakemore S. J., Wolpert D. M., Frith C., “ Why can't you tickle yourself? ”, *Neuroreport*, Vol. 11, No. 11, 2000.
- [22] (財) ブレインサイエンス振興財団, 伊藤 正男, 川合 述史, “ ブレインサイエンス・レビュー 2001 ”, *医学書院*, 216 - 243, 2001.
- [23] Nozaki D., Kurtzer I., Scott H. S., “ Limited transfer of learning between unimanual and bimanual skills within the same limb. ”, *Nature Neuroscience*, Vol.9, No. 11, 1364 - 1366, 2006.

- [24] Yokoi A., Hirashima M., Nozaki D., “ Lateralized sensitivity of motor memorize to the kinematics of the opposite arm reveals functional specialization during bimanual actions. ”, *The Journal of Neuroscience*, Vol.34, Issue 27, 9141 - 9151, 2014.
- [25] Miall, R. C., Christensen, L. O., Cain, O., Stanley, J., “ Disruption os state estimate in the human lateral cerebellum. ”, *Plos Biology*, Vol. 5, Issue 11, e316, 2007.
- [26] Desmurget M., Epstein C. M., Turner R. S., Prablanc C., Alexander G. E., Grafton S. T., “ Role of the posterior parietal cortex in updating reaching movements to a visual target. ”, *Nature Neuroscience*, Vol. 2, No. 6, 563 - 567, 1999.
- [27] Desmurget M., Grafton S. T., “ Forward modeling allows feedback control for fast reaching movements. ”, *Trends in Cognitive Science*, Vol. 4, No. 11, 423 - 431, 2000.
- [28] Jordan A. Taylor, Richard B. Ivry, “ Cerebellar and prefrontal cortex contributions to adaptation, strategies, reinforcement learning. ”, *Prog Brain Res*2014, Vol. 210, 217 - 253, 2014.

謝辞

本研究を進めるにあたり、本学 田中宏和准教授には一年次の頃から熱意あるご指導をいただき、心から感謝申し上げます。田中先生には研究分野に関する勉強のご助言だけにとどまらず、様々な面でご教授いただきました。多彩かつ活発な勉強会に参加させていただき、脳科学に関する多くの知識と同世代のみでとどまらない研究仲間を得られたのは先生のお力添えがなければありえませんでした。この経験があったからこそ、ここまで研究を進める事ができました。本当にありがとうございます。副指導として多くのご助言ご指導頂きました党建武教授に心から感謝して申し上げます。党先生には、多方面からの観点での質問を頂き、より深く物事を考える視点をご教授いただいたことに非常に感謝しております。審査委員として多くのご助言ご指導頂きました鵜木祐史准教授に心から感謝して申し上げます。鵜木先生には、仮配属時から大変お世話になり、発表について何を述べるのが大切なのかと言う点やより良く人に物事を伝えるにはどうすればよいかをご指導いただきました。研究室で日々研究を進めるにあたり、多くのご助言ならびにご指導頂きました末光厚夫助教に心より感謝申し上げます。末光先生から教わりました基礎的な勉強内容から、週報等での多くの議論により今の研究があります。誠にありがとうございます。毎日の研究生活で共に議論を重ねてきた田中研の Li Longchuan さん、党研の近藤宗真さん、Guo Xiangjun さん、Do Khac Phong さん、また赤木研および鵜木研の皆様には非常にお世話になりました。心より感謝いたします。留学生の皆様にはつたない英語を理解していただき、様々な点において議論できたことは勉学面だけでなくこれからの人生においても大きな糧になると思います。

本研究の実験は、東京大学大学院教育学研究科身体教育学コース野崎研究室の皆様のおかげで行うことができました。野崎大地教授をはじめ、野崎研究室所属の修士課程 佐々木彰一さんには実験機器のセットアップから実験計画に至るまで大変お世話になりました。また、博士課程 林拓志さんには行動実験に対するアドバイスや他方向からのアプローチ方法などの様々なディスカッションをさせていただきました。誠にありがとうございます。ご多忙の中、多くのご尽力およびお力添えを賜り本当に感謝申し上げます。

さらに、モーターコントロール研究会 2014, 2015 や ASCONE2014, 2015, 第4回 多次元トレーニング&レクチャーにて共に議論させていただきました皆様にも、感謝申し上げます。この有意義な時間がなければ私の研究はここまで進んではいなかったと思います。

最後に繰り返しとはなりますが改めまして、この場を借りて皆様にお礼申し上げます。

学会発表リスト

- Hitoshi Sato, Akikazu Sasaki, Daichi Nozaki, Hirokazu Tanaka, "Formation of internal forward model with sensory and reward prediction errors: A behavioral confirmation", International Conference on Advanced Mechatronics 2015 (ICAM2015), 2015.