

Title	移動ロボットの自律化を目的とする人間の行動知能の模倣
Author(s)	Shang, Tao
Citation	高知工科大学, 博士論文.
Date of issue	2006-03
URL	http://hdl.handle.net/10173/319
Rights	
Text version	author



Kochi, JAPAN

<http://kutarr.lib.kochi-tech.ac.jp/dspace/>

平成 18 年 3 月修了
博士（工学）学位論文

移動ロボットの自律化を目的とする人間の行動知能の模倣

Imitation Approach of Human Action Intelligence
for Autonomous Control of Mobile Robot

平成 17 年 12 月 16 日

高知工科大学大学院 工学研究科 基盤工学専攻
留学生特別コース

学籍番号 1076012

尚 濤

Tao Shang

目次

第1章 序論	6
第1節 研究の背景.....	6
第2節 本研究の目的	10
第3節 本論文の構成.....	11
参考文献.....	16
第2章 障害物回避行動の計測.....	18
第1節 はじめに.....	18
第2節 運転シミュレータの構築.....	19
第1項 運転シミュレータの仕様.....	20
第2項 運転制御のブロック線図.....	23
第3節 衝突の危険度.....	25
第1項 危険度の定義.....	25
第2項 ハンドルの修正周期と速度,距離,方向との関係.....	27
第3項 ハンドルの修正周期による危険度の定量化.....	33
第4節 行動決定の類似度.....	35
第5節 静的障害物を有する環境における行動戦略の特徴.....	38
第1項 行動戦略と障害物の個数,大きさとの関係.....	38
第2項 行動決定の類似度を利用した行動戦略の計測.....	43
第6節 動的障害物を有する環境における行動戦略の特徴.....	47
第1項 行動戦略と障害物の個数,大きさ,速度との関係.....	47
第2項 行動決定の類似度を利用した行動戦略の計測.....	52
第7節 まとめ	55
参考文献.....	56
第3章 模倣システムの構築.....	57
第1節 はじめに.....	57
第2節 模倣問題の設定.....	58
第3節 問題解決システム.....	59

第4節 模倣システムの展開.....	61
第5節 まとめ.....	62
参考文献.....	63
第4章 模倣に向ける知識の表現法と獲得法.....	64
第1節 はじめに.....	64
第2節 知識の表現法.....	65
第1項 知識獲得の入力空間.....	65
第2項 ルールベース.....	68
第3節 知識の評価法.....	72
第1項 GA 進化による知識の評価法.....	73
第2項 シミュレーションによる検討.....	76
第4節 データ学習による知識獲得法.....	79
第1項 距離型ファジィ推論法.....	80
第2項 学習アルゴリズムの提案.....	83
第3項 シミュレーションによる検討.....	88
第5節 まとめ	91
参考文献.....	92
第5章 模倣に向ける知識の使用と推論法.....	93
第1節 はじめに.....	93
第2節 知識使用を融合した推論法.....	94
第1項 知識半径の概念.....	94
第2項 知識半径を導入した距離型ファジィ推論法.....	96
第3項 知識半径の性質.....	101
第4項 知識半径の確定法.....	104
第3節 知識半径を用いた障害物回避行動の模倣.....	106
第1項 経路面積誤差の定式化.....	107
第2項 知識半径の確定.....	111
第3項 シミュレーションによる検討.....	112

第4節 動的な知識半径の導入による障害物回避行動の模倣.....	119
第1項 動的な知識半径の導入.....	119
第2項 予見機能を導入した評価基準.....	120
第3項 知識半径の確定.....	123
第4項 シミュレーションによる検討.....	125
第5節 まとめ	133
参考文献.....	134
第6章 実験による模倣システムの検証.....	136
第1節 はじめに.....	136
第2節 知能ロボットの遠隔操作システム.....	137
第1項 移動ロボットの構造.....	138
第2項 移動ロボットの移動機能.....	143
第3項 操縦インタフェースの構造.....	146
第3節 走行実験による検討.....	149
第4節 まとめ	155
参考文献.....	156
第7章 結論	157
謝辞.....	161
著者発表文献リスト.....	163
付録.....	165

第1章 序論

第1節 研究の背景

近年、機械工学と電気工学が融合したメカトロニクス技術を基盤として、ロボット工学技術が急速に進展した。高度に発達した社会が到来した今、作業の精度・速度・効率を追求するロボットのみならず、柔軟性に富み、複雑な実世界での作業環境や人間の要求に良く適応し、自由で自然な動きを見せるロボットに対する意味が必然的に高まってくる。このため、従来の産業用ロボットに比べて、自律・移動型ロボットで特に重要性が増している。日本ロボット工業会の調査[1]により、図1-1に示す非製造業やパーソナルロボットの国内需要は、自律・移動型ロボットの研究開発が進むと、2010年には製造業を超える1兆円規模にまで拡大すると予測されている。今後、自律・移動型ロボットは一般の人々の間で活動し、さらに急速な高齢化社会の進展を背景として、医療・福祉ロボットのように人々と直接接触する機会が増えると考えられる。

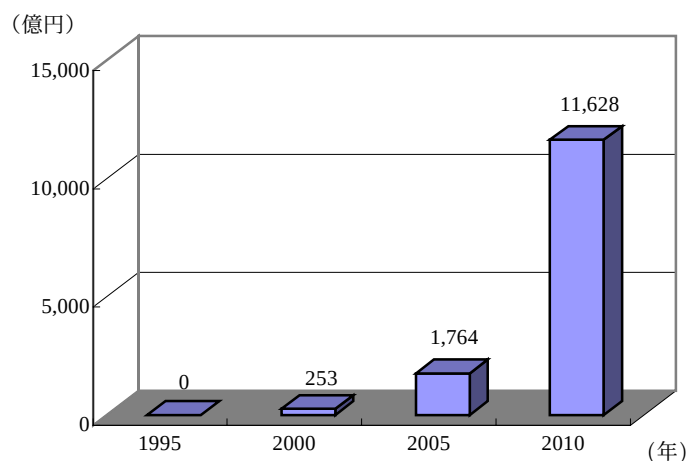


図1-1 パーソナルロボットの国内需要予測

ヨーロッパでもロボットの自律性を高める研究の重要性が認識されはじめている。これを推進すべく、フランスでは国立情報科学研究センターから、NEDOプロジェクトの一環として開発されているロボットのシミュレータ OpenHRP を

使用してロボットの自律性高度化の実現を目指す研究テーマが 2002 年に公募された。また、人間型ロボット研究に関する研究会がドイツ・フランス共同で立ち上げられたところである。

自律・移動型ロボットに対しては、移動のための本体構造や、人工知能等により制御を智能化する技術、移動型ロボットを自律的に分散制御したり遠隔地から制御したりするための技術、自律性を有するために必要となる外部情報の取得と意味解析に関する技術などが重要となる。NEDO 技術開発機構の平成 16 年度国際共同研究先導調査報告[2]により、自律性高度化の観点から、自律性に貢献する技術を抽象レベル別に分類すると、(1)行動の決定に高度な状況判断を必要とする高いレベルの行動計画技術、(2)これらの行動を達成するためにどのように運動を計画するかという中位レベルの運動計画・経路計画技術、(3)モデルと実機の誤差や外乱が存在する場合でも、柔軟に対応して所期の運動を達成する下位レベルの制御技術に分けられる。

この分類に基づいて、行われている研究を整理すると以下ようになる。

行動計画技術については、環境情報により自らの行動を決定する高度な行動を計画する手法に関する研究[3-5]が行われている。通常、行動を有限状態の遷移として表現し、行動の設計を支援する技術を開発するとともに、自律的に行動を獲得する技術についても研究が行われている。特に人間型ロボットなど多自由度ロボットの高度な行動計画は、問題の複雑性が高いため解決が難しく、東京大学などで研究[6-8]がはじめられたばかりで、シミュレーションによる研究が多く、今後この技術に関する研究は、実機による実証を含めてさらに重要になると考えられる。

運動計画・経路計画技術については、それぞれ異なった目的を持つ自律ロボットが多数存在するとき、それら個々の計画をどのように協調させるかの研究が行われてきた[9]。さらに、作業におけるシンボリックな制約を取り入れた運動計画にこれを発展させている。障害物回避をしながらロボットを目的位置に到達させるための手法が盛んに研究されているが、幾何制約を中心に計画を行うものが多く、ロボットの物理的な制約を考慮した計画はまだ少ない。

制御技術については、作業達成のための運動やバランス制御に優先度をつけて、動力学を考慮して多くの自由度をどのように多くの自由度に分配するかという研究が従来から行われている。最近では、作業に必要な運動を第1優先度に、また筋肉モデルを考慮してトルク限界に対する余裕を第2優先度に設定してこの理論を適用して人間の動きを解析した結果、人間は理論に基づいて求められる最適な運動に近い運動をしていることを示した研究も行われた。

一方、ロボットの自律化を実現するために、様々なアプローチ[7,10]が試されている。その中、人間の形態と機能を模したロボットの研究開発は盛んになりつつある[11, 12, 13]。ソニーの AIBO や GRIIO, トヨタのパートナーロボット, テムザックのロボリアや援竜, ホンダの ASIMO, そして、ロボカップや ROBO-ONE のようにロボコンに参加するための各種ロボットなど、数え上げればきりが無い。それらの研究は模倣というアプローチがロボットの自律化を実現する有効性を主に示した。代表的例として、2005 年 1 月 12 日に、東京大学・生産技術研究所と産業技術総合研究所は、人間の全身行動の模倣が可能なヒューマノイドロボット「HRP-2」を用いて「会津磐梯山踊り」を成功に踊らせることに発表した[14]。早稲田大学は、2003 年に、人間の演奏を模倣する人間形フルート演奏ロボット「WF-4」が開発された[15]。そのロボットはフルート演奏に必要な人体各器官の機能を再現した各部機構を備える。設計時に人間とロボットの寸法比較とフルート演奏時の姿勢の解析結果を設計に生かし、従来のロボットに比べて寸法バランスと演奏をする時の姿勢を大幅に人間形に近づけることに成功した。また、オムロン株式会社は 2001 年 10 月 16 日に、ネコ型コミュニケーションロボット「NeCoRo:EPA-R01」を発表した。オムロン独自の感情生成モデル MaC (Mind and Consciousness) により、外界からのインプットに対して、人間の持っている喜怒哀楽などの感情を模倣することができる[16]。これらにより、模倣アプローチを通して、人間の部分的な運動や行為を人間の形態や機能を模したロボットで再現していると見ることができる。模倣アプローチは、ロボットの自律化の実現に役に立つことがわかる。人間の生活環境で行動するロボットが、同じ環境に存在する人間の行動戦略を自分の行動

戦略として取り込み記憶し、適切な時と場所でその行動戦略を再現するといった、「行動知能の模倣」が実現できれば、ロボットが人間の行動戦略を運用することで人間の行動知能を持つことができるという点で、ロボットの自律走行の実現にとって大変重要になると考えられる。

以上から、移動ロボットの自律化への要点としては、人間の行動知能の模倣に着目することにする。従来の模倣と比べて、模倣の有効性を示すだけでなく、環境条件の変化を強調する一般的模倣方法を目指している。つまり、移動ロボットは自分で最適な行動を選び、動作の相異により模倣の程度を可変するなど高い行動表現力を持っているという機能を実現する模倣法を目指している。

第2節 本研究の目的

本研究の目的は、人間の障害物回避行動を真似ることで、移動ロボットの自律走行を実現することである。人間の障害物回避行動は脳における高度な戦略によって支配されるので、行動後の結果として残された何らかのデータにその戦略が隠されている。基本的な立場としては、人間の障害物回避行動後のデータから回避戦略を抽出することが可能であると主張している。そして、抽出した戦略を実装した移動ロボットは、人間の障害物回避行動を模倣することができ、結果的に自律走行が実現できる。従って、本論文では、データから障害物回避戦略の抽出法を開発し、シミュレーションおよび実験を通して、その有用性を示す。以後では、障害物回避問題を取り上げている意味では、「戦略」は「知能」と同意語として使用する。つまり、障害物回避の戦略は障害物回避の知能、行動戦略は行動知能と同じ意味とする。

まず、人間の障害物回避戦略の特徴を把握するためには、障害物回避行動の運転シミュレータを開発する。シミュレータを用いて人間の障害物回避行動を計測することにより、障害物回避戦略と環境との関係に関する解明に着目する。これらの結果は人間の障害物回避戦略を真似る際に重要な根拠になる。

つぎに、得られた知見に基づいて、障害物回避の結果として残されるデータから人間の回避戦略の抽出法を構築する。具体的に、人間の障害物回避行動の模倣を実現する模倣システムは問題解決システム的一种として、知識ベース、学習方法、推論方法から備えるべきだと考え、それぞれ構築している。

最後に、回避戦略の抽出法を移動ロボットに実装して、行動知能の模倣で実現された自律走行実験により、提案した戦略抽出法の有用性を確認する。

第3節 本論文の構成

以下では、論文の構成各章について説明する。

第1章では、目的と研究立場について述べている。

本研究の目的は、人間の障害物回避行動を真似ることで、移動ロボットの自律走行を実現することである。人間の行動はすべて脳の活動に支配されている。行動の知的レベルが高いほど、脳も高度な活動をする。例えば、自動車の運転過程では、環境を正しく認識、素早く次の動作を判断、そして確実に動作を実施するといった一連の知的行動はやはり脳における高度な戦略によって制御されていることが分かる。知的行動の結果として、場面に応じてハンドルの切り方やアクセルの踏み具合などがデータの形で残される。基本的な立場としては、これらのデータ中に脳の高度な戦略が隠されているはずであるので、人間の障害物回避行動後のデータから回避戦略を抽出することが可能であると主張している。抽出した戦略を実装した移動ロボットは、人間の障害物回避行動を模倣することができ、結果的に自律走行が実現できる。本論文では、データから障害物回避戦略の抽出法を開発し、シミュレーションおよび実験を通して、その有用性を示す。

第2章では、人間の障害物回避行動の計測について述べている。

人間の障害物回避行動の模倣を通して、移動ロボットの自律走行を実現するにあたって、人間の障害物回避戦略の特徴、特に衝突直前など極限状態における戦略の限界特性を解明することが必要である。しかし、実際の極限状態における人間の行動をデータの形で獲得することが困難である。例えば、自動車運転においては同時に避けられる相手の車の最大台数を測定しようとしても、安全などの問題があるため、実環境での実施が非常に難しい。そこで、本章では、人間の障害物回避行動に焦点を絞って、戦略とその限界を計測するために、運転シミュレータを開発した。次に、ハンドルの修正周期に着目した危険度と、人間の戦略変化を示す行動決定の類似度を定義した上、運転シミュレータを用いて、20代～30代で視力正常の10人を対象として測定実験を行った。仮想環

境における人間の障害物回避行動を計測することにより、環境と人間の戦略と限界特徴との関係を定性的および定量的に解明できた[17]。これらの結果は人間の障害物回避戦略を真似る際に重要な根拠になる。ここでは、環境は、静的障害物、動的障害物、障害物の個数、大きさ、速度を意味する。

第3章では、模倣システムの構築について述べる。

模倣とは、自分で創りだすのではなく、すでにあるものをまねなうことを意味しており、他者と類似あるいは同一行動をとることである。移動ロボットは、人間と同様な障害物回避行動、つまり人間の障害物回避戦略を模倣できれば、障害物回避問題に限りでは、人間と同レベルの行動知能を持っていると判断する。すなわち、ここでは「脳を創る」という立場に立っている。一方、同じ場面や環境に対して全く同じ行動を実現する場合では、人間の行動後に残されるデータを全部記憶していれば、簡単なサーチ機能によりそのまま再現すればよい。しかし、多少違った場面や環境においても人間のような回避行動を実現しようとする、どうしても推論機能を備えることが必要である。実際の場面や環境を入力として、回避行動を出力とすれば、人間の障害物回避戦略の模倣システムは問題解決システムの一種類である。問題解決システムでは、まず推論機能が要る。推論するには知識が要る。知識が学習により獲得される。したがって、人間の障害物回避戦略の模倣システムは、知識ベース、学習方法、推論方法から備えるべきだと考え構築している。

第4章では、模倣に向ける知識の表現法と獲得法について述べる。

問題解決システムにおける知識の表現手法の一つとして、if-then 型宣言的知識表現がよく使われており、これは、様々な状況に応じて物事を細かく表現することや、個別な状況に応じて異なる判断を下す、といったような人間が普段自然に行われている行動をよく表しているからである。前件部と後件部をファジィ集合とすれば、Yes か No で表現するはっきりした概念は勿論のことであるが、曖昧な概念の定量化表現も可能である。そのため、ファジィ集合を利用した if-then 型宣言的知識表現法は、高次脳機能の工学実現においては有力な表現法の一つであると考え。ここでは、if-then 型宣言的知識表現法を用い

て、人間の障害物回避戦略を表すことにする。if-then による知識表現を利用した、ファジィ推論法として、Mamdani の推論法、関数型ファジィ推論法、簡略型ファジィ推論法を始め、最近ファジィ集合間の距離情報に基づく距離型ファジィ推論法も提案されている。これらの推論モデルは、二値論理における知識だけではなく、あいまいな概念も取り扱えるので、実システムへの応用実績を数多く持っており、脳の推論機能のある側面を実現していると言える。本章では、知識の表現方法を通して人間の障害物回避戦略を表示する有効性を確認した。そして、知識は問題解決能力をチェックするために、遺伝アルゴリズム GA を用いた知識の評価法を提案した[18]。知識の不足と知識獲得の困難という問題点を確認したうえで、多変量で記述しにくい場合でも高速で知識獲得を実現するために、改善した距離型ファジィ推論法の学習ルゴリズムを提案した[18]。距離型ファジィ推論法のオリジナルな学習アルゴリズムに知識（ルール）の生起確率を導入することにより、知識の最適化を行った。最後に、シミュレーションにより提案する新学習ルゴリズムの有効性を示す。

第5章では、模倣に向ける知識の使用と推論法について述べる。

推論は問題を解決するための高度な知的行為である。推論するには知識が要る。推論結果としての問題解決策の良し悪しは一般に推論アルゴリズムと知識と知識の使用に左右される。したがって、よりよい推論結果を得るために、適切な推論アルゴリズムと正確な知識と知識の用法とも必要である。

本章では、推論法として距離型ファジィ推論法を使用する。理由としては、ファジィ集合間の距離を利用して推論を行うので、推論用ルールの物理意味が明確であり、推論の方向性が大雑把に予測できる。または、前章に説明したよう、距離型ファジィ推論法のアルゴリズムを利用すれば、学習アルゴリズムとの結合も簡単にでき、かつ獲得した障害物回避戦略に関する知識の最適化も容易に行うことができる。

これまでの研究では、知識の用法については殆ど言及していない。しかし、人間が推論により何かの結論を下す時に、脳にある全ての知識を同時に利用す

るのではなく、状況に応じて選択的に知識を利用している。知識の選択的利用は、最も関連性のある知識を利用してすばやく結論を下すことが背景にあって、自然に最適化されている自然的な行動だと考えられる。本章では、第4章の学習法により得られた if-then 型宣言的知識が正確なものとして、脳内に行われる知識使用の選択策略を表現する一手法を与え、知識を選択的に利用する推論法を提案する。具体的に、まず事実にも最も関連性のある知識の範囲を意味する知識半径の概念を導入することにより、知識の選択的使用行為を表現する。次に知識半径を考慮した距離ファジィ推論アルゴリズムを提案する[19]。最後に障害物回避問題のシミュレータを用いて実験を行うことにより、提案手法の有効性を示す。

第6章では、実験による模倣システムの有効性について検討する。

提案の諸手法及びシミュレータの有用性について実証するために、全方向行動機能を持つ知能ロボットの遠隔操作システムを開発した[20]。全方向移動ロボットを用いる理由としては、人間のように素早く障害物を回避することが出来るからである。実験の流れとしては、まず試験者が無線 LAN を介して、全方向移動ロボットを、障害物回避をしながらゴールまで運転する。人間の障害物回避行動をした結果として、どの場面・環境に応じて、どのように操作されたのかがデータとして残されている。第4章の学習法を用いて、これらのデータから人間の回避戦略を知識として抽出する。得られた知識を利用して、第5章で述べた知識の使用法と推論法に基づいて、全方向移動ロボットの自律走行を実現する。実証実験を通じて、考案した概念、提案した諸手法、開発したシミュレータが、移動ロボットの自律走行に適用できることが分かった。

本論文の構成ブロック線図は図1-2に示す。

以上、人間の障害物回避行動後のデータから障害物回避戦略を抽出し、そして抽出した戦略を実装した移動ロボットは、人間の障害物回避行動を模倣することができ、結果的に自律走行が実現できるという知見を導くことが本研究の出発点である。運転シミュレータの構築に端を発し、知能ロボットの遠隔操作システムの構築に端を終わり、人間の障害物回避戦略の特徴とその戦略の限

界特性を解明し，障害物回避行動をより効果的に再現する模倣システムの構築とすることで，移動ロボットの自律化を図る人間の行動知能の模倣が可能であるという基本的な研究の枠組みの構築を行ったものである．

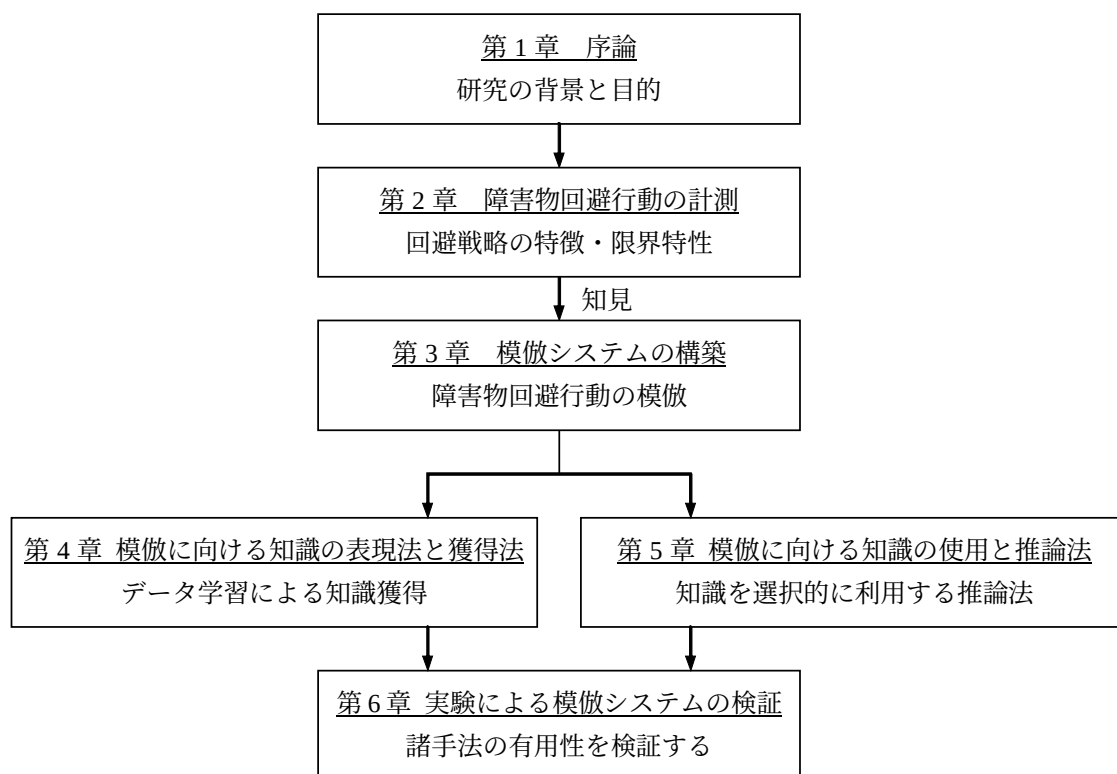


図1-2 論文の構成

参考文献

1. 日本ロボット工業会広報委員会：ロボットハンドブック，日本ロボット工業会，1997.
2. <http://www.nedo.go.jp/itd/fesendo/h16/>
3. 篠田芳明：地上用自律ロボット，日本ロボット学会誌，Vol.18，No.7，pp.26-30,2000.
4. 浦環：自律型海中ロボット，日本ロボット学会誌，Vol.18,No.7，pp.31-34,2000.
5. 嘉数侑昇：自律農業ロボットの課題と現状，日本ロボット学会誌，Vol.18，No.7,pp.49-52,2000.
6. M. Hirose, Y. Haikawa, T. Takenaka, and K. Hirai: Development of Humanoid Robot ASIMO, Proc. of Workshop on Explorations towards Humanoid Robot Applications of IEEE/RSJ International Conference on Intelligent Robots and Systems, 2001.
7. 比留川博久，加賀美聡：ロボットの知能とシステム統合：ヒューマノイドを例にとって，日本ロボット学会誌，Vol.20，No.5，pp.464-469，2002.
8. 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 Sytem, Vol.37, pp.185-193, 2001.
9. 倉林大輔，長川研太:幾何条件による自律移動ロボット群の編隊構造遷移，日本ロボット学会誌，Vol.23，No.3，pp.376-382，2005.
10. 稲葉雅幸:ロボット知能のアーキテクチャ,日本ロボット学会誌，Vol.20，No.5，pp.470-473，2002.
11. 國吉康夫：模倣ロボットは人間的知能を獲得するか？,学術月報,Vol.53，No.9,pp.11-18，2000.
12. Stefan Schall: Is Imitation Learning The Route to Humanoid Robots?, Trends in Cognitive Science, Vol. 3, pp.233-242 ,1999.

13. 岡田慧, 鈴木義久, 國吉康夫, 稲葉雅幸, 井上博允: 視覚による人間動作認識と全身行動表現に基づくヒューマノイドの行動獲得, 記憶, 再現, 第19回日本ロボット学会学術講演会, pp.431-432, 9月18-20日, 2001.
14. 東大生研・産総研共同プレス発表資料: ヒューマノイドロボットと踊り師範による会津磐梯山踊りの共演---ヒューマノイドロボットをメディアとした民俗芸能・技法の動的アーカイブに向けて, 1月12日, 2005.
15. <http://ascii24.com/news/i/hard/article/2001/10/16/630482-000.html?geta>
16. Shuzo Isoda, Manabu Maeda, Yuji Hiramatsu, Yu Ogura, Hideaki Takanobu, Atsuo Takanishi and Kunimitsu Wakamatsu: Realization of an Autonomous Search for Sound Blowing Parameters, Proceedings of the 2003 IEEE International Conference on Robotics & Automation, pp.3582-3587, May 12~17, Taipei, Taiwan, 2003.
17. 尚濤, 王碩玉: 複雑な環境における人間の障害物回避能力の計測, 第9回知能メカトロニクスワークショップ講演論文集, pp.121~126, 和歌山, 8月5~6, 2004.
18. Tao Shang and Shuoyu Wang : Knowledge Acquisition and Evolution Methods for Human Driving Intelligence, International Journal of Innovative Computing, Information and Control, Vol.2,No.1, pp. 1~16,2006.
19. 尚濤, 王碩玉: 知識半径を利用した距離型ファジィ推論法に基づく人間の操縦経路の模倣, 第21回ファジィシステムシンポジウム, pp.239~244, 東京, 9月7~9日, 2005.
20. 尚濤, 王碩玉: 知能模倣型ロボットのための遠隔操縦システムの開発, 第10回日本知能情報ファジィ学会中国・四国支部大会, pp.15-18, 広島, 12月3日, 2005.

第2章 障害物回避行動の計測

第1節 はじめに

人間の行動はすべて脳の活動に支配されている。行動の知的レベルが高いほど、脳も高度な活動をする。例えば、自動車の運転過程では、環境を正しく認識、素早く次の動作を判断、そして確実に動作を実施するといった一連の知的行動はやはり脳における高度な運転知能によって制御されていることが分かる。運転者は障害物を回避しながらできるだけゴールまで運転する。結果的に、車運転に関する人間の行動知能は移動ロボットの自律化に重要な参考になると考えられる。

人間の障害物回避戦略を抽出し、人間の障害物回避行動を模倣する前に、人間の障害物回避戦略の特徴、特に衝突直前など極限状態における戦略の限界特性を解明することが必要である。しかし、実際の極限状態における人間の行動をデータの形で獲得することが困難である。例えば、自動車運転においては同時に避けられる相手の車の最大台数を測定しようとしても、安全などの問題があるため、実環境での実施が非常に難しい。

従って、本章では、人間の障害物回避行動に焦点を絞って、障害物回避戦略とその限界特性を計測するために、まず日常の車に代わる運転シミュレータを開発した。次に、ハンドルの修正周期に着目した危険度と、人間の行動戦略変化傾向を示す行動決定の類似度を定義した上、運転シミュレータを用いて、仮想環境における人間の障害物回避行動を計測することにより、環境と人間の戦略と限界特性との関係を定性的および定量的に解明できた。これらの結果は人間の障害物回避戦略を真似る際に重要な根拠になる。ここでは、環境は、静的障害物、動的障害物、障害物の個数、大きさ、速度を意味する。

第2節 運転シミュレータの構築

人間の障害物回避戦略とその限界特性を計測するために、運転シミュレータを開発した。運転シミュレータの構成ブロック線図を図2-1に示す。実験者(ドライバ)は図2-2に示す運転シミュレータを用いて仮想的運転を行う。車(エージェント)と障害物とゴールはディスプレイ中で実現された仮想的環境にある。測定実験様子を図2-3に示す。実験者はハンドルとアクセル・ブレーキを利用してエージェントをゴールまでに障害物を回避しながら運転する。エージェントの走行方向はハンドルの回転角度、エージェントの推進力はアクセルとブレーキの踏み具合によって指示される。指示された方向信号と推進力信号をコンピュータに入力すると、エージェントの運動方程式を解くことにより得たエージェントの運動軌道を仮想的環境にプロットされる。

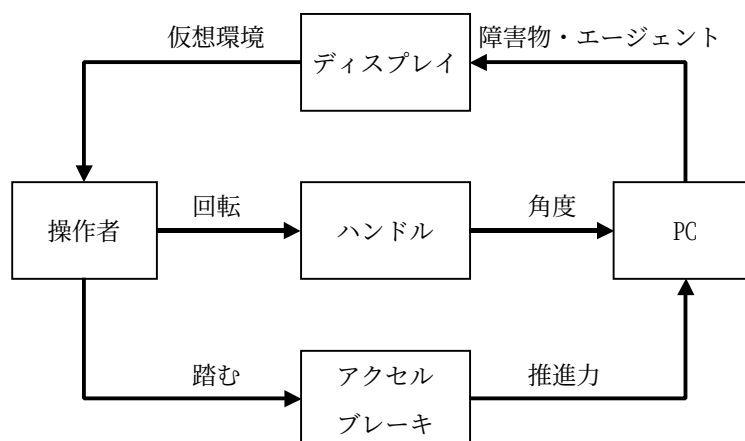


図2-1 運転システムのブロック線図



図 2-2 ハンドルとブレーキと環境



図 2-3 測定実験の様子

第1項 運転シミュレータの仕様

仮想的環境は、17インチの液晶ディスプレイ（I-O DAT社製LCD-AD17CESで、解像度は1280×1024pixel、視野角度は80°、応答速度は40ms）で実現する。ハンドルとブレーキとアクセルともマイクロソフト社製であり、ハンドルはSideWinder precision racing wheel(ver1.0)、ブレーキとアクセルはSideWinder pedals (ver1.0)を利用する。エージェントの運動軌道計算と環境生成を行うPCの仕様は、CPU: Pentium4-2.40GHz、メモリ:256MB、プログラムはWindows XPの下でVisual C++ 6.0 とDirectX 9.0で開発した。

よりリアルな運転を再現するために、エージェントに(2-1)式で表すダイナミックスを持たせることにした。

$$m\ddot{Y} + D\dot{Y} = F \quad (2-1)$$

ただし、

$m=10\text{Kg}$: エージェントの質量

$Y=[y_1(t), y_2(t)]^T$: エージェントの位置座標

T : マトリクスの転置

$\dot{Y}, \ddot{Y}$: エージェントの速度, 加速度

$D=\begin{bmatrix} 0.7 & 0 \\ 0 & 0.7 \end{bmatrix}$: 摩擦係数行列

$F=f(t)[\sin(\alpha(t)), \cos(\alpha(t))]^T$: エージェントが受ける合力

$f(t), \alpha(t)$: エージェントが受ける制御入力, つまり推進力と回転角度

図2-4に示す仮想的環境においては、幾何学的な構造と位相的な関係の定量化を容易にするため、エージェントは方向付きの円形、障害物は四角形、ゴールは円形で表す。エージェントの初期位置は環境内に自由に設定できるが、初期速度を0とする。エージェントの運動はハンドルとアクセル・ブレーキから入力される合力のベクトルを式(2-1)に代入することによりリアルタイムで得られる。障害物の出発位置と速度方向と速度の値は自由に設定できる。また、人間の運転はどのように障害物の個数と大きさに影響されるかを測定するた

めに，障害物の個数と大きさは自由に選べる．ゴールの目標点は環境内における任意の位置に設定できる．エージェントが目標点へ無衝突に到着する場合のみを成功，それ以外の場合を失敗とする．

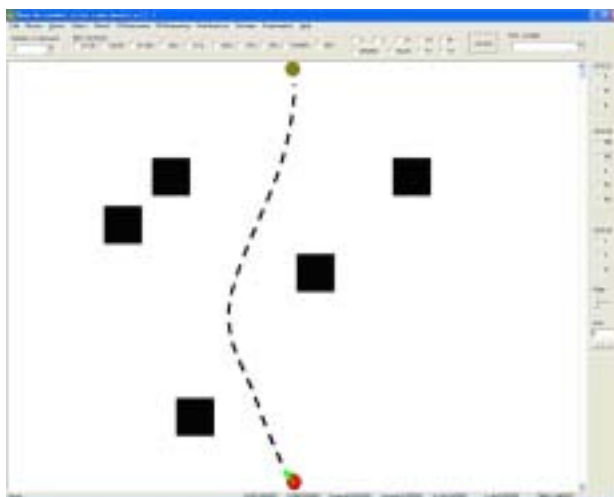


図 2-4 仮想環境

操作者は運転シミュレータを用いて仮想的運転を行う際に，ディスプレイ中における車（エージェント）と障害物とゴールを観察することで，運転行動の判断を行う．Cardらの認知情報処理モデル[1,2]により，人間の認知情報処理過程は知覚系，認知系，運動系の三つのシステムを含み，更に運転シミュレータによる作業反応時間が表2-1に予測されることができる．

表 2-1 Card らの認知情報処理モデルによる作業反応時間の予測

認知過程	知覚処理/仮現運動視	認知処理	運動処理	出力：行為
平均処理時間	$\tau_p = 100\text{ms}/50\text{ms}$	$\tau_c = 70\text{ms}$	$\tau_m = 70\text{ms}$	$\tau_o = 1\text{s}$
反応時間	←-----→ 240ms/190ms			
	←-----→ 1240ms/1190ms			

運転シミュレータにおいて、知覚処理における仮現運動視の平均処理時間 τ_p は計 50ms 必要であり、しかも運転環境の画面は 10ms 毎に更新するので、人間の反応は環境状況の変化より必ず遅いことがわかる。従って、その運転シミュレータを用いて人間の限界特性を計測することが可能である。表 2-2 には運転シミュレータの極限パラメータを詳細に説明する。

表 2-2 運転シミュレータの極限パラメータ

パラメータ		値
ディスプレイの面積		1280×1024 pixel ²
運転空間の面積		1177×768 pixel ²
ゴール	形	円形 (半径=15 pixel)
	最大数	1
エージェント	形	円形(半径=5 pixel)
	最大数	1
	回転角度	[-90 度, +90 度]
	推進力	[-18N, +18N]
	最大速度	257pixel/sec
障害物	形	正方形 (辺=80×K pixel) $K \geq 1$
	最大数	141 (重なり合わない)
	最大速度	300 pixel/sec
	最大移動距離	300 pixel
エージェントと 障害物の関係量	最大相対距離	1833 pixel
	最大相対速度	557 pixel/sec
	相対角度	[-180 度, +180 度]

第2項 運転制御のブロック線図

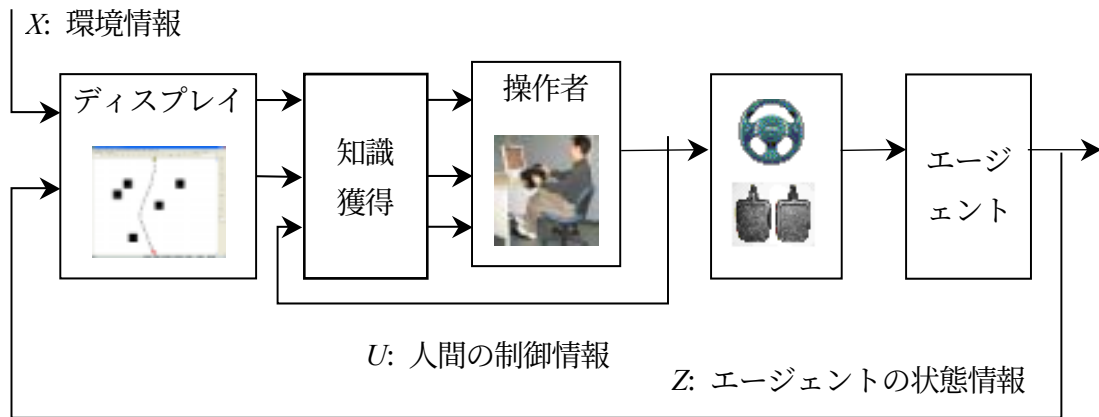


図2-5 運転制御のブロック線図

前述の運転シミュレータによる運転行動が図2-5の示すような運転制御のブロック線図に表されることが出来る。そのブロック線図の入力空間は環境情報 X とエージェントの状態情報 Z と人間の制御情報 U から組み立て、式(2-2)に表示される。

$$\text{Input space} = \{X, Z, U\} \quad (2-2)$$

$$X = [\text{Goal}^{n_g}, \text{Obstacle}^{n_o}]$$

$$\text{Goal} = [\text{position}, \text{speed}, \text{size}, \text{shape}]$$

$$\text{Obstacle} = [\text{position}, \text{speed}, \text{size}, \text{shape}]$$

$$Z = [\text{Agent}]$$

$$\text{Agent} = [\text{position}, \text{speed}, \text{size}, \text{shape}]$$

$$U = [\text{force}, \text{direction}]$$

ただし、

Goal : ゴール

Obstacle : 障害物

Agent : 操作者の制御されるエージェント

$position, speed, size, shape$: それぞれ位置と速度と大きさと形という対象の属性, 対象はゴールと障害物とエージェントを含む

$force, direction$: 操作者からエージェントの受ける制御入力, つまり推進力と回転角度

n_g, n_o : それぞれ環境において存在するゴールと障害物の個数

入力空間(2-2)は人間の運転モデルの解析に汎用な枠組みを提供する. 運転シミュレータにおけるエージェント・ゴール・障害物のパラメータの組み合わせ方を利用して, 異なる環境における運転戦略の区別を了解し, 能力の大きさを定量化することができる.

第3節 衝突の危険度

本研究では、人間の障害物回避行動に焦点を絞って、人間の障害物回避戦略と戦略の限界特性を解明するに着眼するので、人間は置かれた運転環境から一体何かに基づいて、障害物回避戦略を適切に運用するかという基本問題が出てくる。文献[3]では、人間が自動車を運転するとき、衝突のある危険性を感じると、操舵を修正する周期が急に短くなると指摘される。文献[4]では、人間の運転行動をモデルにして、危険度の概念に基づいて障害物回避のための動作計画法が提案されている。文献[5]では、交通事故を減らすための自動車運転作業時の認知応答特性解析と、文献[6]では、先行車両の減速に伴う運転者の減速動作のモデル化は衝突の危険性を言及した。それらにより、置かれた運転環境から感じる危険の度合いに基づいて、行動を適切に選択することが人間の障害物回避特徴の1つである。一括、危険性の関連研究においては、ファジィ推論に基づく定量化方法[4,7,8]と数式による定量化方法[9]とを分けて危険性を判断している。もちろん、もし物理的環境から感じる心理的な危険度を適合に測定できれば、危険性に関する理論は広い領域に明確で本質的な役割を果たして来ないかと思う。

自動車の運転においては、物理的環境から感じる心理的な危険度を測定するには、心理物理量と生理信号との二つのアプローチが考えられるが、ここでは定量化表現を容易に図るために、心理物理量に着目する。具体的に、まず、危険度の関数を定義した。次に、操作者によるハンドルの修正周期を危険度の心理物理量として、障害物との相対速度と距離と方向はどのように危険度に影響するかを測定し、そして測定結果に基づいて危険度の定量計算式を与えてみる。

第1項 危険度の定義

危険は、1つの対象を中心として、別の対象がその対象に危害または損失の生ずる恐れである。危険は少なくとも2つの対象の間に起きると言える。そこで、対象(O_1, O_2)間の危険関係を適切に表現する危険度関数 $\eta(O_1, O_2)$ を定義する

上に、更に1つの対象と複数の対象の間、つまり1つの対象と環境の間に危険関係を描くことができる。具体的に、もし人を O_1 、環境に存在している i 番目障害物 OB_i を O_2 とすれば、走行経路 P の各位置で障害物ごとに対する危険度の積分として式(2-3)の E で表すことができる。

$$M(x) = \sum_{i=1}^{n_o} \eta_x(O_1, OB_i) \quad (2-3)$$

$$E = \int_P (M(x)) dx$$

ただし、

$\eta_x(A, OB_i)$ ：位置 x における i 番目障害物の持つ危険度

n_o ：走行環境に存在している障害物の個数

$M(x)$ ：全部の障害物による危険度の和、つまり位置 x における環境危険度

E ：走行経路 P による環境危険度

人間が運転際に最大危険性をなるべく避ける限りで、各位置における $M(x)$ が E の積分量として一定値の範囲に変化する。それで、経路が有限に存在すれば、 E も一定値の範囲に変化する。従って、最も危ない障害物を優先的に回避して生じる経路が少なくとも安全な経路だと言える。もし走行経路の各位置で危険度の積分を経路の評価値とすれば、各危険度が生じる面積の合計値 E を最小とする経路が最も安全な経路である。

危険度関数の計算を通じて、人に最も危ない障害物を判定することができる。危険度関数の設計は今後のより高度な障害物回避戦略モジュールをロボットに導入するに役に立つと考えられる。

第2項 ハンドルの修正周期と速度,距離,方向との関係

前項の危険度関数 $\eta(O_1, O_2)$ を定義するうえで、更なる危険度の定量化は人間の障害物回避戦略特徴の解析にもっと役に立つ。人間が自動車を運転するとき、衝突のある危険性を感じると、操舵を修正する周期が急に短くなるので、文献[3]では、操舵の修正周期を危険性の心理物理量として車速一修正操舵周期特性を解析した結果に基づいて、正常運転時の操舵周期特性より長くなると、意識低下で危険性を感じなくなり、つまり居眠り運転と判断したら、居眠り運転に対して警報を鳴らす技術を開発された。実際では、速度のみならず障害物との距離と方向も重要なファクターであると考える。ここで、ドライバによるハンドルの修正周期を危険度の心理物理量として、障害物との相対速度と距離と方向はどのように危険度に影響するかを測定する。

【アプローチ】

第2節の運転シミュレータ

【対象】

被験者には、大学院生は1名採用された。

年齢：28歳

性別：男

【方法】

計測方法：図2-6の示すように、エージェントの出発点を最下行の midpoint に設定する。被験者の意識低下を避けるために、1つの動的障害物とゴールを一定確率でエージェントの前方に現れさせる。環境の動的な変動により被験者の運転行為を計測する。

具体的条件設定については、障害物は、1%の確率でエージェントの前方左右15度の範囲内にある200pixelsと離れるところで現れて来ており、その運動方向が0度から360度までの範囲に乱数で決まる。現れて来た障害物は200pixel/secの速度で300pixelsの移動距離で一回の往復運動をする。ゴールは0.5%の確率でエージェントから300pixels離れる円弧上でランダムに移動す

る。ゴールの確率と障害物の確率の比率は、人間がゴールを間もなく接近する前に、思いもよらないゴールの変化がぴったり発生し、趣味を激励し続けことを保証するパラメータである。

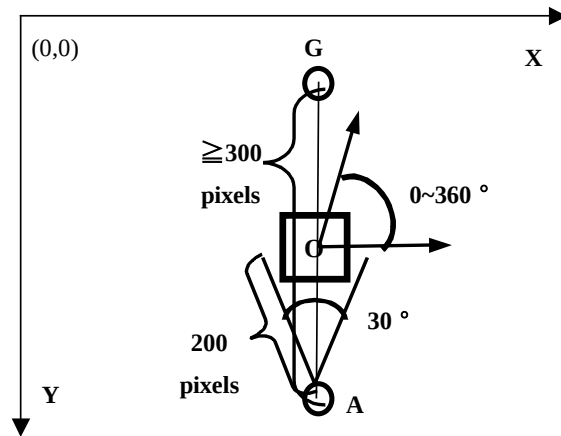


図 2-6 運転環境の配置

(G:ゴール, A:エージェント, O:障害物)

エージェントの進行方向を 0° つまり基準、左側を負の角度範囲、右側を正の角度範囲として、被験者のハンドル操作による操舵角度情報 α は $-90^\circ \sim +90^\circ$ の範囲で連続的に変化する。ここで、操舵情報を $F(-15^\circ \leq \alpha \leq +15^\circ)$ と $L(-90^\circ \leq \alpha < -15^\circ)$ と $R(+15^\circ < \alpha \leq +90^\circ)$ との三つの修正レベルに大まかに分類して、行動の修正パターンとしては $FL(F \rightarrow L)$ と $FR(F \rightarrow R)$ と $RF(R \rightarrow F)$ と $LF(L \rightarrow F)$ との4つの種類がある。運転途中では何らかの行動修正パターンが現れると、ハンドル操作による操舵角度が一回修正されたと見なして、修正の周期が短ければ短いほど、ドライバは障害物から感じる危険性が高い。修正周期は、ランダム成分を含める環境条件変更に伴い行為が変化した時に心理的危険程度に適応できるようにする意味を持つ。

計測時間：疲労による不随意動作を避けるために、1回の測定実験時間を5分間とし、特定の被験者にとっては1回の実験が終わってから十分な休憩時間を挟んだ後に次回の実験を再開する。成功または失敗すると、新たな障害物とゴールとエージェントが生成され、測定実験を続ける。

計測結果：エージェントがゴールへ無衝突に到着する場合のみを成功、それ以外の場合を失敗とする。

計測回数：被験者は10回の実験を行う。

障害物とエージェントとの面積比率：80:1(通常)

【結果】

相対速度と距離と方向角度が対象間の関係を表現する計測量として得られた。蓄積された測定結果から、かなり少ない個数（<15）を持ったデータレベルを除いて、ハンドルの修正周期と各計測量との関係を図2-7、2-8、2-9に示す。図2-7、2-8、2-9における縦軸は修正周期の平均値、つまり修正回数の逆数の平均値を表す。

図2-7における横軸は、障害物とエージェントとの相対速度のレベルを表すが、最大相対速度 557pixel/sec を10等分にして、レベルを付けた各々の相対速度範囲を表す。文献[3]と比べて、極限条件における修正周期と相対速度の関係も測定した。しかしながら、相対速度が小さすぎる、例えば車（エージェント）と障害物がほとんど動かない場合に、車と障害物の衝突の可能性が非常に低いので、随意操舵による修正周期は車と障害物の衝突の危険性を反映しない。相対速度が大きすぎる、例えば車と障害物との衝突が起きる場合に、人間の障害物回避能力の制限なので、修正周期は心理の衝突危険性も反映しない。そこで、衝突の危険性と相対速度の関係を解析するため、1レベルと8レベルが削除される。また、不足なデータを持つ9レベルと10レベルも削除される。

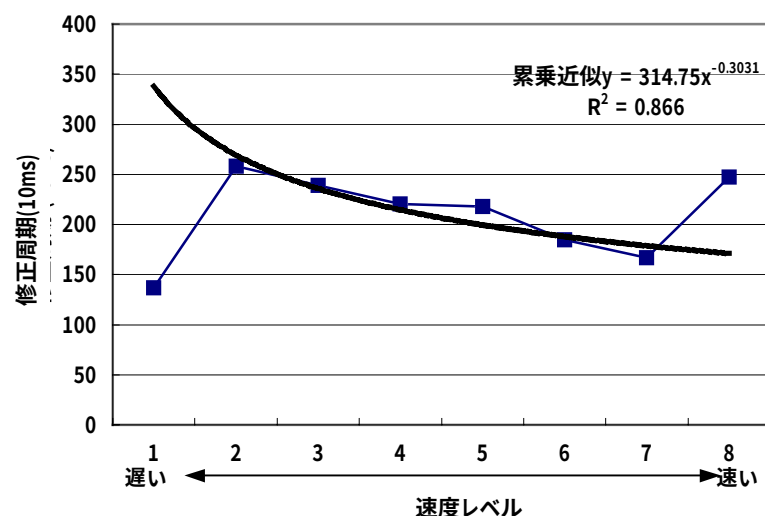


図2-7 修正周期と相対速度との関係

文献[3]と比べて、相対速度のみならず障害物との距離と方向角度も測定した。図2-8における横軸は、障害物とエージェントとの距離のレベルであり、最大距離 1833pixels を 10 等分にして、レベルを付けた各々の距離範囲を表す。不足なデータを持つ7レベルと8レベルと9レベルと10レベルが削除される。図2-9における横軸は、障害物とエージェントとの方向角度のレベルであり、最大角度 180 度を 10 等分にして、レベルを付けた各々の角度範囲を表す。

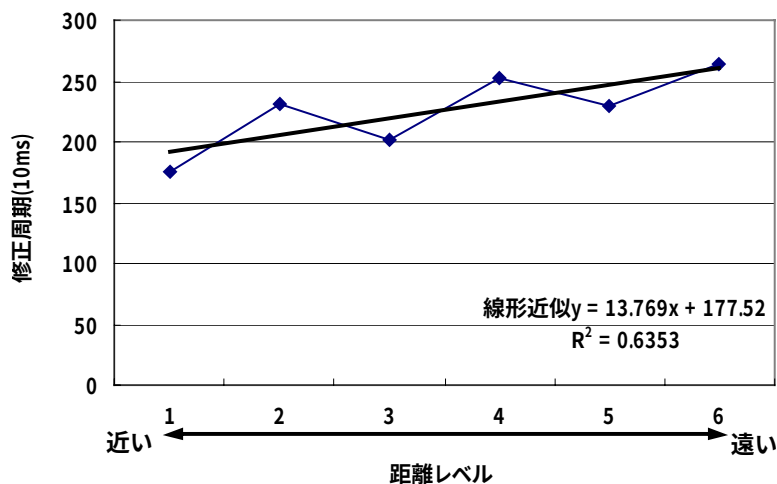


図2-8 修正周期と距離との関係

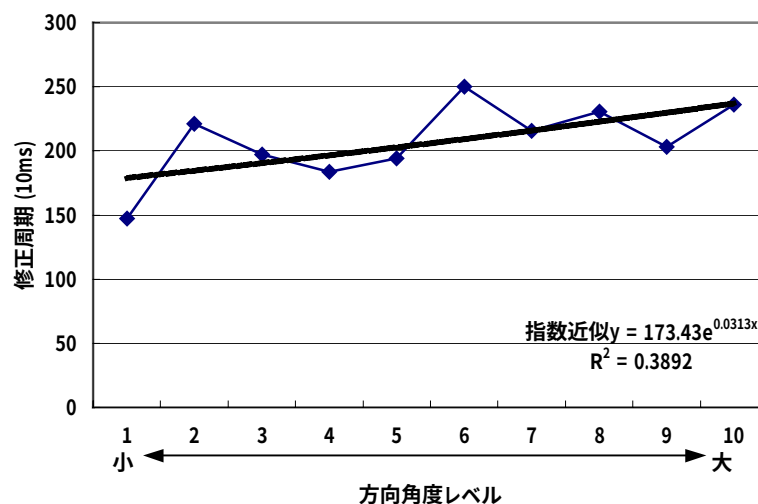


図2-9 修正周期と方向角度との関係

【考察】

図2-7における縦軸は修正周期の平均値、つまり修正回数の逆数の平均値を表す。全体実験においては、各相対速度区間内での修正周期の平均値をプロットし、折線で繋がったグラフは、各相対速度と危険性との関係を示している。同図中の点線は傾向の累乗近似である。したがって、図2-7に示すように、障害物とエージェントとの相対速度の増加につれて、ハンドルの修正回数が増え、すなわち、被験者の感じる危険性が高くなることが分かる。

図2-8における縦軸は各距離区間内での修正周期の平均値を表して、折線で繋がったグラフは距離と危険性との関係を示している。同図中の点線は傾向の線形近似である。したがって、図2-8に示すように、障害物とエージェントとの距離の増加につれて、ハンドルの修正回数が減り、被験者の感じる危険性が低くなる。

図2-9における縦軸は各方向角度区間内での修正周期の平均値を表して、折線で繋がったグラフは方向角度と危険性との関係を示している。同図中の点線は傾向の指数近似である。したがって、図2-9に示すように、障害物とエージェントとの方向角度が増えるにつれて、ハンドルの修正回数が相対的に減り、被験者の感じる危険性が低くなる。

その中に、ハンドルの修正周期との関係については、障害物との相対速度の変化は、距離と方向角度より大きく変化することがわかる。

【結論】

ドライバによるハンドルの修正周期を危険度の心理物理量として、障害物とエージェントとの相対速度の増加につれて、被験者の感じる危険性が高くなる。障害物とエージェントとの距離の増加につれて、被験者の感じる危険性が低くなる。障害物とエージェントとの方向角度が増えるにつれて、被験者の感じる危険性が低くなる。更に、相対速度は、距離と方向角度よりドライバの危険度に大きく影響することがわかる。相対速度が衝突の危険度の主要要素として確認されるが、距離と方向角度という要素を結びつければ、より適合な危険度関数の設計が可能となる。従って、相対速度と距離と方向角度の近似変化により、

以下の関係式を提案した.

$$\text{衝突の危険度} = \frac{|\text{相対速度}|}{\text{距離} \times |\text{方向角度}|} \quad (2-4)$$

ただし, $||$ は絶対値を表す.

衝突の危険度: ある対象が人に生ずる危険度

相対速度: ある対象に対する人の相対速度

距離: ある対象と人の距離

方向角度: ある対象に対する人の相対速度と位置ベクトルとのなす角

第3項 ハンドルの修正周期による危険度の定量化

第2項の危険度の関係式により、衝突の危険度の定量化として危険度関数は、
相対速度 V_r と距離 D_r と方向角度 α_r の近似変化に応じて、

$$\begin{aligned} \text{Danger} &\sim \frac{1}{T}, \quad T \text{ 修正周期} \\ T &\sim a_1(|V_r|)^{-b_1} \quad a_1 > 0, b_1 > 0 \\ T &\sim a_2 D_r + b_2 \quad a_2 > 0, b_2 > 0 \\ T &\sim a_3 \exp(b_3 |\alpha_r|) \quad a_3 > 0, b_3 > 0 \end{aligned} \quad (2-5)$$

その3つの要素を基にして簡略化のパラメータを利用すると、対象(O_1, O_2)間の危険関係を評価する危険度関数は、次式に定義される：

$$\begin{aligned} \eta(O_1, O_2) &= f(D_r, V_r, \alpha_r) \\ &= e^{-(|\alpha_r|/180)} \times |V_r| / (D_r + \delta) \end{aligned} \quad (2-6)$$

ただし、

O_1, O_2 ：環境における2つの対象、 O_1 を中心とする

V_r ：最大相対速度による正規化した対象 O_2 に対する対象 O_1 の相対速度

D_r ：最大距離による正規化した対象 O_1 と O_2 の距離

α_r ：対象 O_1 の相対速度 V_r と2つ対象の位置ベクトルとのなす角であり、特に、 V_r が0になると、 α_r を0とする。

δ ：定数、0.001とする

(2-6) 式は静的対象や動的对象に対して、2つの対象間の危険関係をうまく説明できる。一般に、 D_r は小さいほど、危険度が大きくなる。 $|V_r|$ は大きいほど、危険度が大きくなる。 $|\alpha_r|$ は小さいほど、危険度が大きくなる。ここでは、質点間のみの関係を計算するので、三次元以上の特徴を考慮しない。

図2-10, 2-11, 2-12はそれぞれ $\alpha_r = 0, 90, 180$ の場合に危険度関数の変化状況を示す。

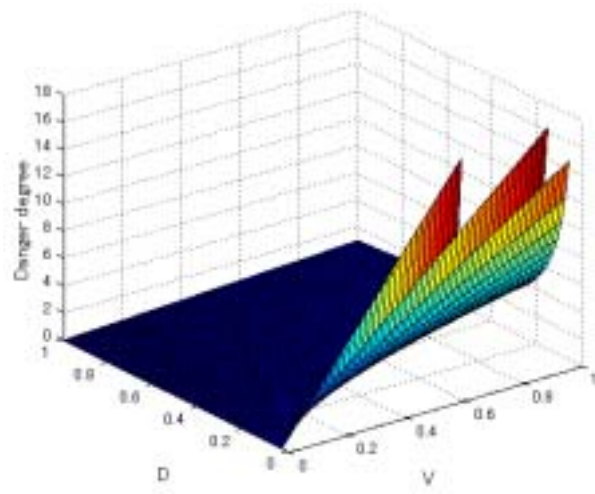


図 2-10 危険度関数($\alpha_r = 0$)

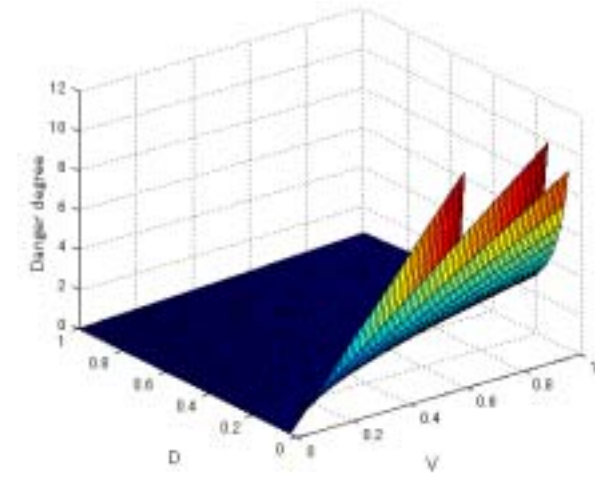


図 2-11 危険度関数($\alpha_r = 90$)

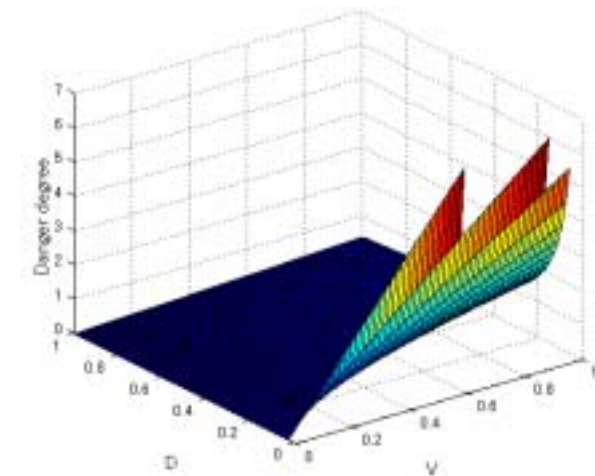


図 2-12 危険度関数($\alpha_r = 180$)

第4節 行動決定の類似度

運転シミュレータにおいて、相同や相似な障害物回避行為がよく発生した。この潜在的行為再現性はある障害物回避戦略の特徴を意味する可能性がある。その戦略特徴を捉えるために、行動決定の類似度を定義した。行動決定の類似度（ADSD：Action Decision Similarity Degree）とは、決定符号照合により計算される経路間の相似程度である。そこで、行動決定の類似度の計算は、行動決定の符号化と符号による照合という2つ段階を含む。

行動決定の符号化について、機械の揺らぎに対してのロバスト性を向上すると考えられると、各時刻の行動決定を符号に対応して、時系列の決定符号をベクトル化することで、決定変化パターンの多様さを追いかけていく。

具体的に、 θ_t は t 時刻における回転角度を行動決定として表す。ただし、 $t \in \{1, 2, \dots, T\}$ 、 T は時間の長さである。結果的に、1 つの経路は回転角度系列 $\{\theta_1, \dots, \theta_t, \dots, \theta_T\}$ に対応する。便利な分析のために、図 2-13 のように回転角度範囲を小さい角度区間に分けることで、時刻あたりの回転角度を数式(2-7)により離散的な決定符号に変換させる。

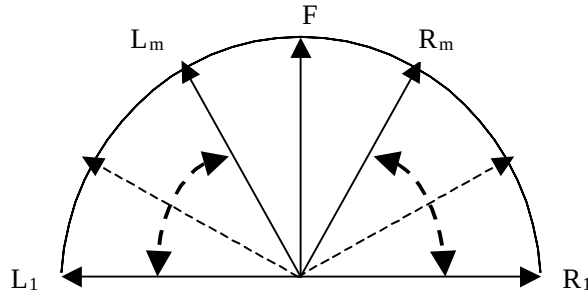


図 2-13 回転角度の対応決定

$$v(\theta) = \begin{cases} F, \theta \in [-\delta, +\delta] \\ R_i, \theta \in (\delta + (i-1)(90-\delta)/m, \delta + i(90-\delta)/m] \\ L_i, \theta \in [-\delta - i(90-\delta)/m, -\delta - (i-1)(90-\delta)/m) \end{cases} \quad (2-7)$$

$$i \in \{1, \dots, m\}$$

ただし、

回転角度範囲： $[-90, +90]$

δ : 指定閾値

m : 左右回転角度範囲の割り数, ≥ 1

図 2-13 の示すように, 離散的な決定符号は $\{F, L_1, L_2, \dots, L_m, R_1, R_2, \dots, R_m\}$ を含む. m が大きいほど, 離散的な決定符号が多い. $2m+1$ の種類の離散的決定符号に応じて, 決定変化パターンの総数は $N = (2m-1)*2+2 = 4m$ である. 特に, $m=1$ の場合に, $N=4$ になり, 離散的な決定符号は $\{F, L, R\}$ で表す.

結果的に, 回転角度系列による経路は決定符号系列に表示される.

$$path = \{v_1^{n_1}, v_2^{n_2}, \dots, v_L^{n_L}\} \quad (2-8)$$

ただし, L は符号化された経路の長さである. 同じ決定符号が連続的に並べることが可能なので, n_i で連続な同一の決定符号の数を表す. かつ,

$$\sum_{i=1}^L n_i = T, L \leq T, v_i \neq v_{i+1}, i \in \{1, 2, \dots, L-1\} \text{ が成立する.}$$

更に, 行動決定の変化傾向として $V = \{v_1, v_2, \dots, v_L\}$ を保留して行動決定の類似度への計算を行う.

符号による照合については, 行動決定の符号化の結果として V を用いて行動決定の類似度を計算する. 2 つの符号化された経路 V^1, V^2 に対して, 行動決定の類似度は(2-9)式によって計算される:

$$PDSD(V^1, V^2) = \left(\sum_{k=1}^L p_k \right) / L, \quad p_k = \begin{cases} 1, & v_k^1 = v_k^2 \\ 0, & v_k^1 \neq v_k^2 \end{cases} \quad (2-9)$$

$$L = \begin{cases} \min(L_1, L_2), & \min(L_1, L_2) > 2 \\ \max(L_1, L_2), & \text{others} \end{cases}$$

ただし,

L_1, L_2 : それぞれ経路 V^1, V^2 の長さ

L : 比較の長さ

比較の長さ L については, 通常, $\min(L_1, L_2)$ を用いて余分な決定符合を除き, 主な決定変化パターンを追いかけることができる. 経路が短すぎる場合に, $\max(L_1, L_2)$ を用いて行動決定の類似度を計算することができる.

$ADSD(V^1, V^2)$ の性質について説明する:

- 1) 符号化された経路は時間に関係がなし，経路における決定変化パターンを描くので， $ADSD(V^1, V^2)$ は経路における決定変化パターンの相似性を描く．
- 2) $ADSD(V^1, V^2) = ADSD(V^2, V^1)$.
- 3) もし V^1 ， V^2 は完全に同じ，あるいは，一定の長さの V^1, V^2 は $V^1 \subset V^2, or V^2 \subset V^1$ を満たせば， $ADSD(V^1, V^2) = 1$.
- 4) 離散的な決定符号の種類が増えると， $ADSD(V^1, V^2)$ が減る．

第5節 静的障害物を有する環境に関する行動戦略の特徴

エージェント・ゴール・障害物のパラメータ組み合わせ方により配置される環境において、人間の運転行為が適応的に変化する、つまり人間の行動戦略は環境の変化につれて変化する。人間の障害物回避戦略と戦略の限界特性を解明するために、運転シミュレータを利用して、人間の障害物回避戦略特徴の計測を行う。ここでは、運転成功率を人間の障害物回避能力の評価基準として、環境を次第に複雑とするならば、障害物回避戦略特徴の計測を行う。環境の区別を考慮しており、それぞれ静的・動的障害物を配置する環境を通じて、位置や個数や大きさや速度など障害物の属性を調整して環境複雑性を変えることで、運転成功率と障害物の属性との関係、つまり、障害物回避戦略と環境複雑性との関係を解析してみる。

第1項 行動戦略と障害物の個数,大きさとの関係

【対象】

被験者には、大学院生は2名採用された。ただし、

年齢：28歳 29歳

性別：男 女

【方法】

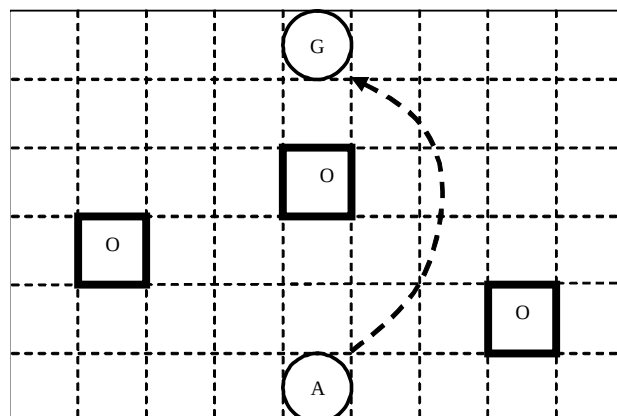


図 2-14 運転環境の配置

(G:ゴール,A:エージェント,O:障害物,矢印付きの点線:可能な被験者の経路)

計測方法：図 2-14 の示すように，エージェントの出発点を最下行の中心に，ゴールを最上行の中心に固定し，障害物の位置を乱数で生成する．環境を次第に複雑にするため，障害物の大きさを固定して，障害物の個数を単調増加に調整して計測を行う．そして，障害物の大きさを増加して，重ねて障害物の個数を単調増加に調整して計測を行う．

計測時間：各場合に対して，疲労による不随意動作を避けるために，1回の測定実験時間を3分間とし，特定の被験者にとっては1回の実験が終わってから十分な休憩時間を挟んだ後に次回の実験を再開する．

計測結果：エージェントが目標点へ無衝突に到着する場合のみを成功，それ以外の場合を失敗とする．

各場合の計測回数：障害物の個数と大きさが固定された場合に，被験者は20回の実験を行う．

障害物の個数：1～100

障害物とエージェントとの面積比率：80:1(通常)，180:1(やや大)，320:1(非常大)

【結果】

各場合の成功率を平均した．異なる面積比率に対する運転成功率と障害物の個数との関係は図 2-15～2-17 に示す．

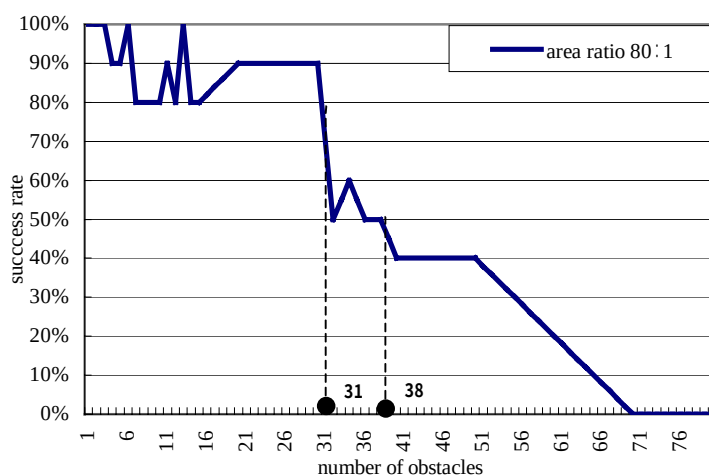


図 2-15 運転成功率と障害物の個数との関係(面積比率 80:1)

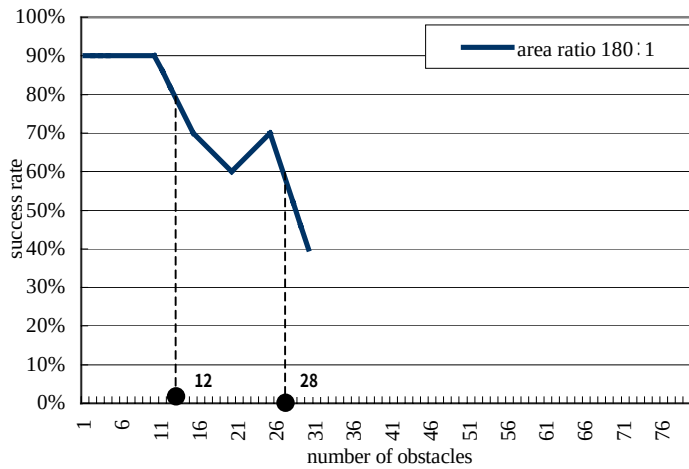


図 2-16 運転成功率と障害物の個数との関係(面積比率 180:1)

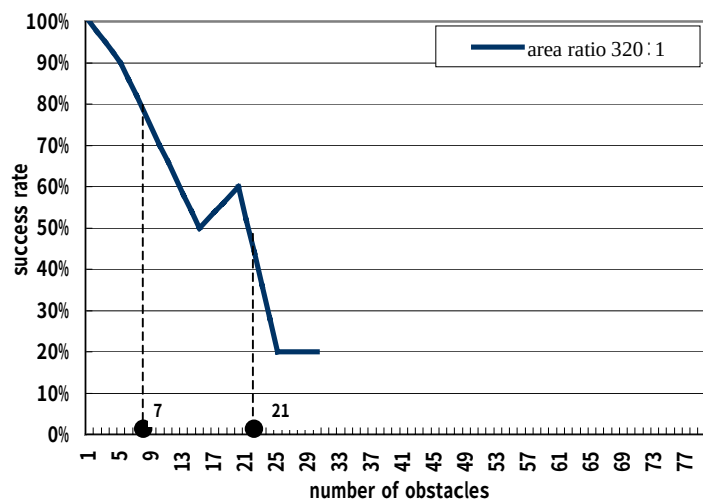


図 2-17 運転成功率と障害物の個数との関係(面積比率 320:1)

【解析方法】

障害物回避戦略の解析は、以下の手順で行った。

- ① 各場合の運転成功率を障害物回避能力の評価基準とする。
- ② 個数と大きさなど障害物の属性を環境複雑性の評価基準とする。
- ③ 運転成功率と障害物の属性との関係により、障害物回避能力と環境複雑性との関係を解析し、更に障害物回避戦略と環境複雑性との関係を解析する。

【考察】

図 2-15～2-17 の示すように、静的障害物を有する環境において、異なる面

積比率でも、障害物の個数が増えるとともに、運転成功率は最終的に必ずゼロになる。回避能力の限界が存在した、かつその限界は障害物の個数に関係があったと確認した。しかし、障害物の個数は1から100まで増えるにつれて、緩慢な運転成功率の変化は、限界と障害物の個数の関係があまり強くないことを示す。

面積比率 80:1 の通常な場合に、運転成功率の変化に対応する特別な区間が明白に存在したことがわかる。障害物の個数として 38 以内に、運転成功率は 50%以上であり、特に 31 以内に、成功率は 80%を超える。31 個の静的障害物はほとんど 20%の運転空間の面積を占めるので、これは複数の静的障害物が存在する環境におけるある人間の障害物回避戦略という意味を含むはずである。障害物の個数として 31 以内に、安定な成功運転率は障害物回避能力がほとんど障害物の個数に関係がない、即ち、障害物の個数はほとんど障害物回避能力に影響を与えないといえる。勿論、人間にとっては、行動戦略の特徴として、1 個ぐらいの最も重要な障害物を考慮して行動を決めれば充分である。そして、障害物の状況を換え切ることで、障害物回避を行うと考えられる。

又、高い運転成功率 ($\geq 80\%$) を保つ際に応じる障害物の個数を簡単に説明するため、安定限界という概念を導入して障害物回避戦略の特徴を示す。普通な運転成功率 ($\geq 50\%$) に対して、処理可能な限界という概念を導入して障害物回避能力の限界を示す。図 2-15~2-17 により、面積比 80:1 の場合には安定限界は 31 であり、処理可能な限界は 38 である。面積比 180:1 (やや大) の場合には安定限界は 12 であり、処理可能な限界は 28 である。面積比 320:1 (非常大) の場合には安定限界は 7 であり、処理可能な限界は 21 である。障害物とエージェントの面積比率が増えるとともに、処理可能な限界がやや下がる。同時に、障害物の占める空間面積 (障害物の大きさ $\times$ 限界値) が増える、あるいはやや減る。つまり、被験者はもっと狭い空間でも障害物回避を確実に行える。それにより、障害物の大きさは運転成功率に影響を小さく与える。

【結論】

静的障害物を有する環境において、人間の障害物回避能力の限界が存在した。しかし、障害物の個数と大きさは障害物回避能力に影響を小さく与えると結論した。また、行動戦略の特徴として、1 個ぐらいの最も重要な静的な障害物により行動を決める。そして、環境の変化につれて障害物の状況を換え切ると考えられる。更に、第 3 節に定義された危険度を利用して、もしエージェント A を O_1 、環境に存在している i 番目障害物 OB_i を O_2 とすれば、人間の障害物回避戦略に向ける障害物とは、条件 $\{i \mid \max(\eta(A, OB_i)), i \in [1, \dots, n_o]\}$ ，即ち最大危険度を満たす障害物と少なくとも定義される。ただし、 $\eta(A, OB_i)$ は走行環境における i 番目障害物の持つ危険度であり、 n_o は走行環境に存在している障害物の個数である。

第2項 行動決定の類似度を利用した行動戦略の計測

【対象】

被験者には，大学院生は 10 名採用された．ただし，普通運転免許を有する者は 3 名である．

年齢：22 歳から 29 歳まで，平均年齢 25 歳

性別：男 8 名 女 2 名

【方法】

計測方法：図 2-18 の示すように，エージェントの出発点を最下行の midpoint に，ゴールを最上行の midpoint に固定し，それぞれ三つの障害物を固定の座標位置に設定する上に，被験者の回避行為を計測する．

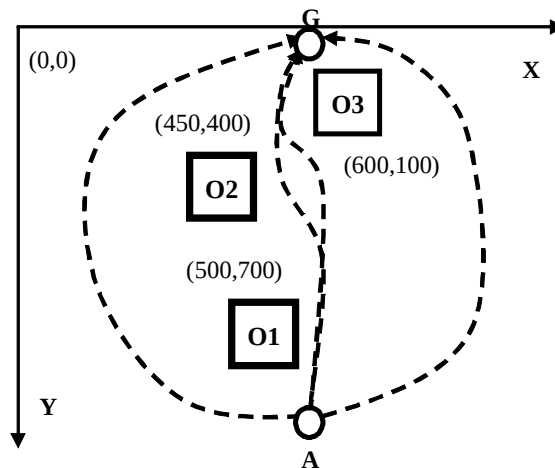


図 2-18 運転環境の配置

(G:ゴール, A:エージェント, O1,O2,O3:障害物, 矢印付きの点線:可能な経路)

計測時間：各場合に対して，疲労による不随意動作を避けるために，1 回の測定実験時間を 3 分間とし，特定の被験者にとっては 1 回の実験が終わってから十分な休憩時間を挟んだ後に次回の実験を再開する．

計測結果：エージェントが目標点へ無衝突に到着する場合のみを成功，それ以外の場合を失敗とする．

1 人あたりの計測回数：被験者は 11 回の実験を行う．

障害物：O1(500,700) O2(450,400) O3(600,100)

障害物とエージェントとの面積比率：80:1(通常)

【解析方法】

障害物回避戦略の解析は、以下の手順で行った。

- ① 行動決定の類似度の計算が1人の計測データの間に行われる。
- ② 連続的な行為の相似性を避けるために、行動決定の類似度の計算は非連続的な計測データの間に行われる。
- ③ 人あたりの1回目の計測データを類似度の計算の比較基準とする。

【結果】

表2-3は1番目の被験者の計測データによる行動決定の類似度の計算過程を示す。表2-4は、被験者あたりの行動決定の類似度を平均した通りで、10人の間に行動決定の類似度の比較を示す。

表2-3 1番目の被験者の行動決定の類似度 ADSD ($\delta=10$)

順番	ADSD	符号化された経路									
基準		F	R	F	L	F	L	F	R	F	R
長さ		134	64	10	126	55	61	38	146	34	32
1	83.33%	102	70	14	25	1	R228				
2	90.00%	72	54	9	122	48	179	29	121	68	L22
3	83.33%	37	53	22	118	215	R90				
4	100%	39	50	9	106	55	110	61	173	16	
5	87.50%	60	55	10	180	82	R125	32	64		
6	87.50%	55	64	17	119	181	112	15	L73		
7	83.33%	43	54	5	195	127	R171				
8	87.50%	67	51	8	261	15	R114	26	73		
9	80.00%	31	52	7	305	34	R59	13	164	4	L181
10	85.71%	92	45	9	124	149	R123	5			
平均 ADSD		86.8%									

(F/L/R：異なる決定符号，決定符号付きの数字：同じ決定符号の数)

表 2-4 10 名の被験者の平均 ADSD ($\delta=10$)

被験者の番号	平均 ADSD (静的障害物)
1	86.8%
2*	76.3%
3*	72.1%
4	80.0%
5	86.2%
6*	86.0%
7	71.7%
8	89.2%
9	76.0%
10	92.4%
平均値	81.7%

(*は普通運転免許を持つ被験者である)

【考察】

表 2-4 には、10 名の被験者は 71.7%から 92.4%まで、かつ平均値 $\geq 80\%$ の行動決定の類似度を持つことを示す。行動決定の類似度は経路間の相似性を決定符号照合で計算するので、平均値 $\geq 80\%$ の行動決定の類似度はその静的な障害物を有する環境において高い行為再現性を反映する。また、被験者が運転経験の有無には関係なく、運転経験がない方も高い行動決定の類似度を持ち、うまく個人戦略の特徴を表現することが分かる。この高い行為再現性は偶然な行為ではなく、行動前に、全般的な環境情報により予め行動経路を決めるのではないかと判断できる。

【結論】

第 1 項での障害物回避戦略に比べて、行動決定の類似度による行為再現性は、行動前に全般的な環境情報により予め行動を決めるという障害物回避戦略、つまり、大局的特徴を持つ障害物回避戦略を体現する。相対的に、第 1 項は局所

的特徴を持つ障害物回避戦略を体現する。以上の静的障害物を有する環境において、大局的再現性は人間の障害物回避戦略の特徴の一つであると結論した。

以下各章では、障害物回避問題を取り上げている意味では、障害物回避戦略の特徴を「障害物回避戦略の大局的特徴」と「障害物回避戦略の局所的特徴」に分けて使用する。

第6節 動的障害物を有する環境に関する行動戦略の特徴

静的な障害物を有する環境における計測に比べて、動的障害物を有する環境において障害物回避戦略がもっと複雑になるべきである。障害物の速度を結びつけて環境複雑性を変えることで、人間の障害物回避戦略と環境複雑性との関係を解析してみる。

第1項 行動戦略と障害物の個数,大きさ,速度との関係

【対象】

被験者には、大学院生は2名採用された、第5.1節と同じ被験者。ただし、
 年齢：28歳 29歳
 性別：男 女

【方法】

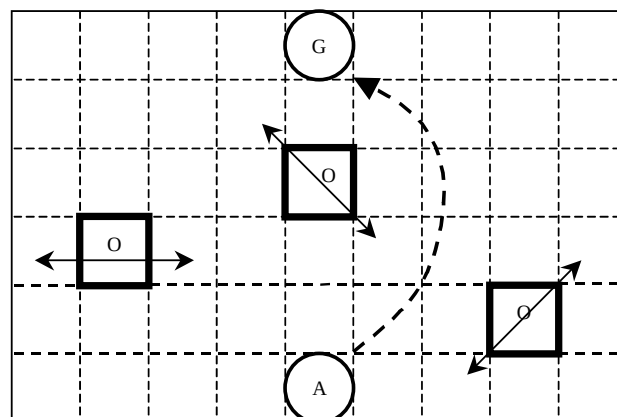


図 2-19 運転環境の配置

(G:ゴール, A: エージェント, O:障害物, →: 障害物の移動方向, 矢印付きの点線:可能な被験者の経路)

計測方法：図 2-19 の示すように、エージェントの出発点を最下行の midpoint に、ゴールを最上行の midpoint に固定し、障害物の初期位置と初期速度方向を乱数で生成する、かつ障害物は定常的な速度と移動距離で往復移動を行う。環境を次第に複雑にするために、まず障害物の大きさと速度を固定して、障害物の個数を単調増加に調整して計測を行う。として、障害物の大きさを増加して、重ね

て障害物の個数を単調増加に調整して計測を行う。一方、障害物の速度を増加して、重ねて障害物の個数を単調増加に調整して計測を行う。

計測時間：各場合に対して、疲労による不随意動作を避けるために、1回の測定実験時間を3分間とし、特定の被験者にとっては1回の実験が終わってから十分な休憩時間を挟んだ後に次回の実験を再開する。

計測結果：エージェントが目標点へ無衝突に到着する場合のみを成功、それ以外の場合を失敗とする。

各場合の計測回数：障害物の個数と大きさと速度が固定された場合に、被験者は20回の実験を行う。

障害物の個数：1～40

障害物とエージェントとの面積比率：80:1(通常), 180:1(やや大), 320:1(非常大)

障害物の速度：0～300 pixel/sec

障害物の移動距離：300 pixel

【結果】

各場合の成功率を平均した。障害物の速度固定の場合に、異なる面積比率に対する運転成功率と障害物の個数との関係は図 2-20～2-22 に示す。

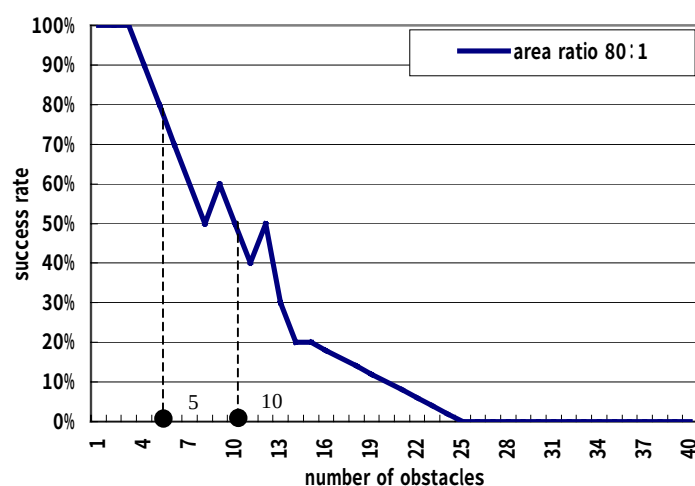


図 2-20 運転成功率と障害物の個数との関係（面積比率 80:1, 障害物の速度 $V=200\text{pixel/sec}$ ）

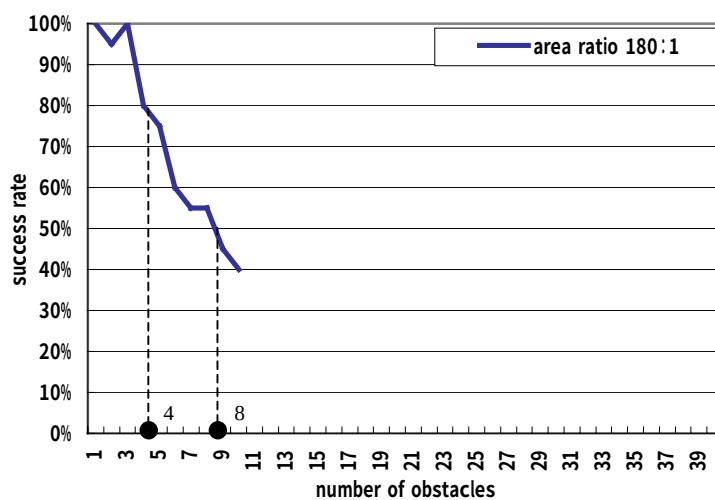


図 2-21 運転成功率と障害物の個数との関係（面積比率 180:1, 障害物の速度 $V=200\text{pixel/sec}$ ）

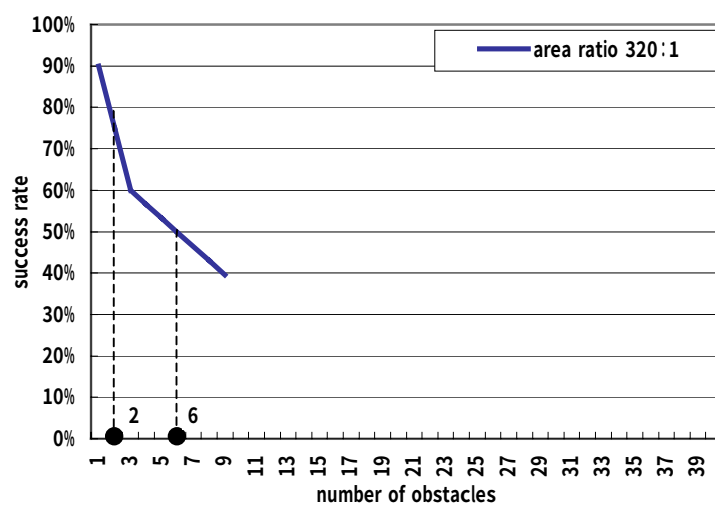


図 2-22 運転成功率と障害物の個数との関係（面積比率 320:1, 障害物の速度 $V=200\text{pixel/sec}$ ）

一方、障害物の大きさ固定の場合に、異なる速度に対する運転成功率と障害物の個数との関係は図 2-23 に示す。

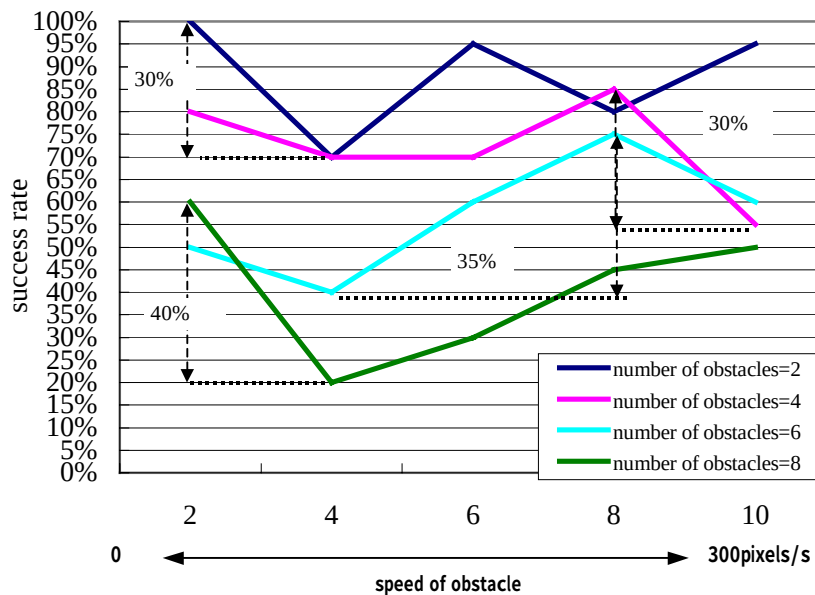


図 2-23 運転成功率と障害物の速度との関係(面積比率 80:1)

【考察】

図 2-20～2-22 の示すように、動的障害物を有する環境において、障害物の速度固定の場合に、異なる面積比率でも、障害物の個数は 1 から 40 まで増えるとともに、運転成功率は 100% から 0% まで著しく下がる。障害物回避能力の限界は障害物の個数に強い関係があることがわかる。従って、動的障害物を有する環境において、同時にほぼ全ての動的な障害物に着目して行動を決める。

面積比 80:1 の場合には安定限界は 5 であり、処理可能な限界は 10 である。面積比 180:1 (やや大) の場合には安定限界は 4 であり、処理可能な限界は 8 である。面積比 320:1 (非常大) の場合には安定限界は 2 であり、処理可能な限界は 6 である。障害物とエージェントの面積比は増えるとともに、安定限界と処理可能な限界がやや下がるが、障害物の占める空間面積(障害物の大きさ×障害物の限界値)が増える。つまり、被験者はもっと狭い空間でも障害物回避を確実にできる。従って、障害物の大きさは障害物回避能力に影響を小さく与える。

また、図 2-23 の示すように、障害物の速度が 0 から 300pixels/sec まで変

わるとともに、運転成功率は40%以内の範囲に変動する。しかし、障害物の個数が2から8まで変わるとともに、運転成功率が明らかに下がる傾向がある。それから、障害物の速度より、障害物の個数の方が障害物回避能力にもっとも大きく影響を与えると判定する。

【結論】

動的障害物を有する環境において、静的障害物を有する環境より人間の障害物回避能力の限界が明らかに存在した。また、障害物の速度と大きさより、障害物の個数は障害物回避能力に影響を大きく与えると結論した。

障害物回避戦略の局所的特徴としては、1つの障害物だけではなく、必ず同時にほぼ全ての動的な障害物により行動を決める、そして、環境の変化につれて障害物の状況を換え切る。さらに、第3節で定義された危険度を利用して、対象($O_1 O_2$)間の危険関係を適切に表現できる危険度関数 $\eta(O_1, O_2)$ を通じて、もしエージェント A を O_1 、環境に存在している i 番目障害物 OB_i を O_2 とすれば、障害物回避戦略に向ける障害物とは、条件 $\{i | \eta(A, OB_i) \geq \delta_L, i \in [1, \dots, n_o]\}$ 、即ち、安定限界以内の数を持つ最大危険度を満たす障害物と少なくとも定義される。ただし、 $\eta(A, OB_i)$ は走行環境における i 番目障害物の持つ危険度であり、 L は安定限界の値であり、 δ_L は安定限界の応じる危険度の閾値である。

第2項 行動決定の類似度を利用した行動戦略の計測

【対象】

被験者には、大学院生は10名採用された。第5.2節と同じ被験者、ただし、

年齢：22歳から29歳まで、平均年齢25歳

性別：男8名 女2名

【方法】

計測方法：図2-24の示すように、エージェントの出発点を最下行の midpoint に、ゴールを最上行の midpoint に固定し、それぞれ三つの障害物を固定の座標位置に設定する上に、被験者の運転行為を計測する。

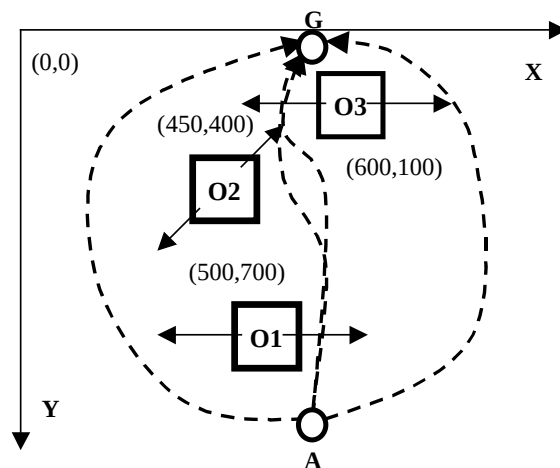


図2-24 運転環境の配置

（G:ゴール, A:エージェント, O1, O2, O3:障害物, 矢印付きの点線：可能な被験者の経路）

計測時間：各場合に対して、疲労による不随意動作を避けるために、1回の測定実験時間を3分間とし、特定の被験者にとっては1回の実験が終わってから十分な休憩時間を挟んだ後に次回の実験を再開する。

計測結果：エージェントが目標点へ無衝突に到着する場合のみを成功、それ以外の場合を失敗とする。

人あたりの計測回数：被験者は11回の実験を行う。

障害物：O1(500, 700) O2(450, 400) O3(600, 100)

障害物とエージェントとの面積比率：80:1(通常)

障害物の速度：200 pixel/sec

障害物の移動距離：300 pixel

【結果】

表 2-5 には、各被験者の行動決定の類似度により、行動決定の類似度の変化を比較する。

表 2-5 10 名の被験者の平均 ADSD ($\delta=10$)

被験者の番号	平均 ADSD (動的障害物)	ADSD の変化 (静的場合と比べる)
1	57.5%	-29.3%
2*	63.3%	-13.0%
3*	68.7%	-3.4%
4	72.5%	-7.5%
5	64.3%	-21.9%
6*	60.2%	-25.8%
7	67.4%	-4.3%
8	20.0%	-69.2%
9	57.5%	-18.5%
10	72.3%	-20.1%
平均値	60.4%	-21.3%

(*は普通運転免許を持つ被験者である)

【考察】

表 2-5 には、10 名の被験者は 20.0%から 72.5%まで、かつ平均値 60.4%の行動決定の類似度を持つことを示す。静的障害物を有する場合と比べて、すべての被験者による行動決定の類似度が下がり、かつ $\leq -20\%$ ぐらいの変化がある。行動前に全般的な環境情報により予め行動経路を決めても、動的障害物を有する環境が複雑になるので、運転の大局計画は障害物の速度によって確実に制限

されることが分かる．

【結論】

静的な障害物を有する環境に関する障害物回避戦略に比べて，動的障害物を有する環境において，障害物回避戦略の大局的再現性が無くなると結論した．

第7節 まとめ

本章では、人間の障害物回避行動に焦点を絞って、障害物回避戦略とその限界特性について報告した。まず日常の車に代わる運転シミュレータを構築した。次に運転シミュレータにおけるエージェント・ゴール・障害物のパラメータ組み合わせ方により配置される環境において、人間の障害物回避行動を計測することにより、障害物回避戦略と環境との関係を定性的および定量的に解明した。複雑な環境における人間の障害物回避戦略は以下に纏められる：静的障害物を有する環境では、人間は障害物回避戦略を予め決めておくという傾向があり、しかも障害物回避戦略の大局的再現性が明らかに存在する。人間の障害物回避戦略の局所的特徴については、障害物の個数にほとんど関係がない、人間は1つの最も危ない障害物に着目して障害物を回避すると少なくとも言える。動的障害物を有する環境において、成功率が低くなったために、障害物回避戦略の大局的再現性は見られなく、障害物回避能力に限界が明らかに存在することがわかる。人間の障害物回避戦略の局所的特徴については、障害物の個数に関係がある、人間は同時にいくつの危ない障害物、つまり、安定限界の値を持つ障害物に着目して障害物を回避すると少なくとも言える。障害物回避能力の限界は障害物の個数と速度と大きさというパラメータで定量化でき、障害物の個数が最も重要な要因であると考えられる。

障害物回避戦略の抽出については、その能力限界、つまり処理可能な限界数を用いてなるべく多い有用な障害物回避戦略の獲得に役に立ち、ガイドラインとして次のように纏められる：静的障害物を有する環境において、少なくとも1つの障害物の場合に対して戦略の抽出を行う。動的障害物を有する環境において、少なくとも処理可能な限界数以内の障害物の場合に対して戦略の抽出を行う。

これらの研究成果は、人間の障害物回避行動の限界を解明することができ、より高度な知的戦略モジュールをロボットに導入できれば、人間の行動を模倣するロボット行動の知能化にも有益であると考えられる。

参考文献

1. 村田厚生:認知科学—心の働きをさぐる,朝倉書店,1997.
2. John M. Carroll :HCI Models,Theories and Frameworks toward a Multidisciplinary Science, Elsevier Science (USA), 2003.
3. 阿久津英作,渥美文治:居眠り運転警報技術,計測と制御, Vol.36,No.3,pp.168-170, 1997.
4. 前田陽一郎,竹垣守一:ファジイ推論を用いた移動ロボットの動的障害物回避制御,日本ロボット学会誌, Vol.6, No.6, pp.50-54, 1988.
5. 中山仁寛,中易秀敏,前田多章:視覚作業時の認知応答特性解明のための心理物理量と生理信号解析,日本機械学会論文集(C編), Vol.70, No.696,pp.2443-2451, 2004.
6. 広瀬俊也,澤田東一,小口泰平:人間—自動車システムにおける減速動作のモデル化,日本機械学会論文集(C編), Vol.70, No.692, pp.1113-1140, 2004.
7. 久保田直行,盛岡利仁,小島史男,福田敏男:動的環境下における移動ロボットのファジイ制御,日本ファジイ学会誌, Vol.12, No.1, pp.55-63,2000.
8. 伊東敏夫,荒木秀夫:追突防止技術,計測と制御,Vol.36,No.3, pp.187-189,1997.
9. 永谷圭司,油田信一:衝突の危険性を評価関数とする移動ロボットの経路とセンシング点の計画,日本ロボット学会誌, Vol.15, No.2, pp.197-206,1997.

第3章 模倣システムの構築

第1節 はじめに

第2章では、開発した運転シミュレータを利用して障害物回避戦略とその限界特性を計測したことで、運転環境における人間の行動知能は移動ロボットの自律行動を多様化に実現するに役に立つと確認した。障害物回避行動の模倣は充分期待できるものであるが、模倣方法で再現した時に得られる効果が如何なるものであるかは未知のものである。従って、行動後の結果として残されたデータからより効果的に障害物回避行動を再現するに着目する。つまり、データから人間の回避行動戦略を抽出し運用する模倣システムを構築するに着目する。

本章では、問題解決システムの見方から模倣システムの構築について述べる。まず、人間の障害物回避戦略を模倣するという問題を設定する。つぎに、人間の障害物回避戦略の模倣システムは問題解決システムの一種類であると確認する。問題解決システムでは、推論機能が要る。推論するには知識が要る。知識が学習により獲得される。この考えに従い、人間の障害物回避戦略の模倣システムは、知識ベース、学習システム、推論システムから備えるべきだと考え構築している。

第2節 模倣問題の設定

近年のロボティクスや発達科学研究では、模倣が大変重要視されている[1-5]. 多少オーバーに表現すると、模倣こそが人間の学習の根源であると言われる. 国語辞書によって、【模倣/模倣】とはまねること、にせることであり、つまり行動・様子などが他の人や物と同じになるようにすることであることがわかる. 結果的に、偶発的に起こったアクションは必ずしも真似しない. われわれの日常生活で子どもが確かに模倣によって発音・発声しなければならないが、他人が行う行動を真似して学習を行うことがよく見られる.

また、ドイツ・マックスプランク研究所のトマセロ[6]によると、真の模倣こそ人間の文化の源であると指摘される. より良い模倣は、確実にそのグループに伝わる方法である. 他者が発見した行為をそのまましっかり真似しないと、正しいやり方は伝わらない.

一括、模倣とは、自分で創りだすのではなく、すでにあるものをまねなうことを意味しており、他者と類似あるいは同一行動をとることである. 移動ロボットは、人間と同様な障害物回避行動、つまり人間の障害物回避戦略を模倣できれば、障害物回避問題に限りでは、人間と同レベルの行動知能を持っていると判断する. すなわち、ここでは「脳を創る」という立場[7]に立っている. 同じ環境に存在する人間の行動戦略を取り込み記憶し、適切な時と場所でその行動戦略を再現するといった人間の行動知能の模倣を行う. 一方、同じ場面や環境に対して全く同じ行動を実現する場合では、人間の行動後に残されるデータを全部記憶していれば、簡単なサーチ機能によりそのまま再現すればよい. しかし、多少違った場面や環境においても人間のような回避行動を実現しようとする、どうしても推論機能を備えることが必要である. 以上により、人間の障害物回避戦略を模倣することは、行動後のデータから人間の回避行動戦略を抽出し運用する模倣システムの構築ということになる. このような模倣システムは問題解決システムの一種類であり、問題解決者を助けて解決策を迅速に達成し、より高水準の解決策を生み出し、そしてもっと容易に適用できる.

第3節 問題解決システム

問題解決は人工知能の最も基礎となる分野である。問題解決システム (Problem Solving System) は、複雑な問題の解決の方法をより簡単な問題の解決法の組み合わせにより構築できるという考え方で作られることである。図 3-1 には問題解決システムの一般的な構成を示す[8]。問題解決システム=推論機構+TMSということがわかる。ただし、推論機構(Inference Engine, IE)は問題の解決を、真偽維持システム(Truth Maintenance System, TMS)は推論管理を行う。

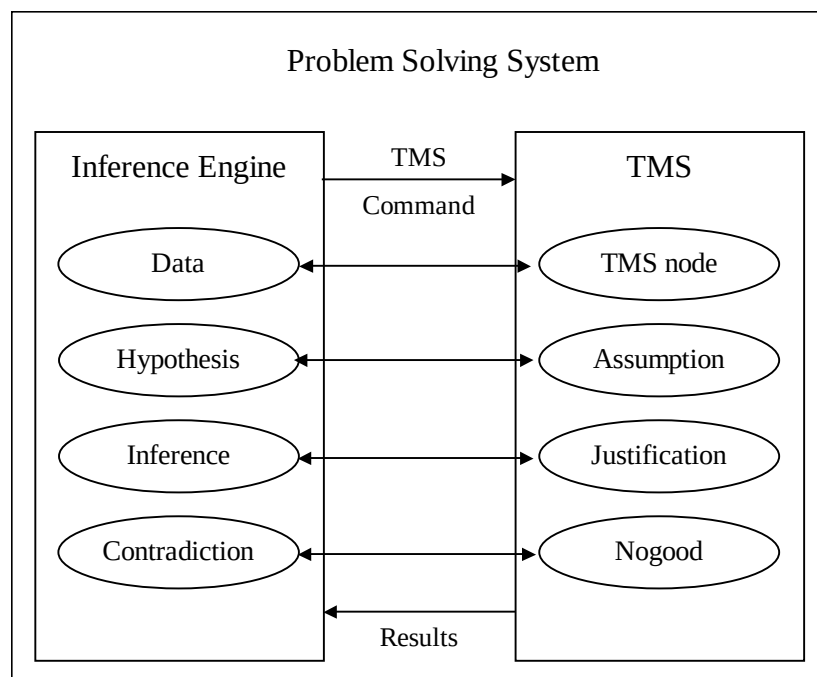


図 3-1 問題解決システムの一般的な構成

問題解決システムの高機能化と高速化という 2 つの相反要請に対して、従来から多くの研究が行われ、さまざまな手法が提案されてきたが、未だ実世界の問題が取り扱えるほどには技術が成熟していない。実世界の問題を扱うには、複数のことを同時に考え、それらの違いを検討することによって判断を下すことができるような機能が不可欠である。このような機能については、人工知能研究ではエキスパートシステム、非単調推論などの要素技術として重要性が認識されている。

通常、問題解決システムは、基本構成要素として知識ベースと推論機構をもつというプログラム構造上の特徴をもっている。知識ベースは、すべての知識を一定の形式で蓄積したものであり、推論機構は、知識ベース内の知識を使って推論を実行するための制御機構である。知識ベース内の知識を利用して、推論を実行し、結論を得る。適切な知識を知識ベースから取り出して適用し、推論を次に進めるということが行われる。知識ベースと推論機構を分離することにより、推論制御メカニズムを標準化でき、これとは独立に、知識ベース内の知識のみを、定義、確認、変更、削除、あるいは追加することができるというメリットがある。

以上により、実際の場面や環境を入力として、回避行動を出力とすれば、人間の行動知能の模倣を実現する手法は図3-2の示すようなブロック線図で表す。ただし、上方の部分は人を、下方の部分は狙うべき人間の障害物回避戦略の模倣システムを表す。模倣システムは知識獲得と知識ベースと推論方法という三つの部分を含む。よい模倣効果を達成するために、評価基準による模倣評価も行われる。

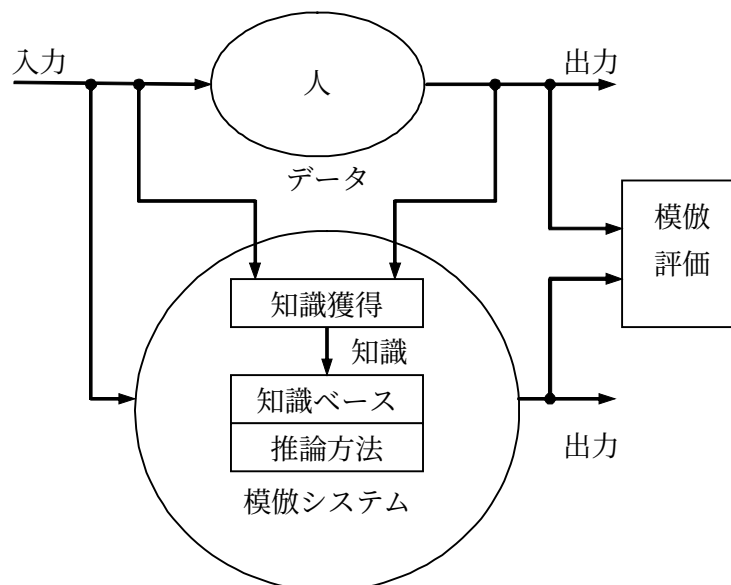


図3-2 人間の行動知能の模倣を実現する手法のブロック線図

第4節 模倣システムの展開

問題解決システムとして人間の障害物回避戦略の模倣システムでは、まず推論機能が要る。推論するには知識が要る。知識が学習により獲得される。したがって、人間の障害物回避戦略の模倣システムは、知識ベース、学習方法、推論方法から備えるべきだと考え構築している。具体的に、以下の点から模倣システムの構築を展開していく。

1. 推論方法（適切な時と場所で人間の行動戦略を再現する）

推論方法は、数学的論理学などを基礎にして発展してきたものであり、人工知能における問題解決のための最も基礎的なメカニズムである。模倣システムでは、ある推論方法を障害物回避戦略の運用メカニズムとして採用して有効な障害物回避行動の模倣を可能にする。

2. 知識表現（人間の行動戦略を表現する）

障害物回避戦略としての知識を表現するためには、推論方法のみでは不十分であり、知識をコンピュータ処理可能なデータ形式で表現することが重要である。コンピュータ処理に向くようにある表現形式は、主に宣言型知識と手続き型知識の2つに分類される。宣言型知識は、「～は～である」という事実の表現と「IF-THEN」のような特徴がある。手続き型知識は、「～の手続きは～」というようなものであり、記述の順序と解釈の順序が一致する。

3. 知識ベース（人間の行動戦略を記憶する）

人間の障害物回避戦略としての知識を運用するために、模倣システムは、記述された大量の知識の集合を蓄積・管理し利用する。知識ベースへの問い合わせに対して、直接の解が存在しない場合、推論機構を用いて新しい解を生成することができる。

4. 学習・知識獲得（人間の行動戦略を取り込む）

知識の特徴は外界の状況変化に適応して自分の情報処理機構を変更したり、さまざまな事例から新しい知識を獲得できる能力にある。模倣システムでは、ある学習・知識獲得方法を用いて、人間の障害物回避行動後データから障害物回避の戦略を定量的に抽出することを目指す。

第5節 まとめ

本章では、人間の障害物回避行動を模倣することにより、移動ロボットの自律走行の実現を目的として、人間の障害物回避戦略の模倣問題を設定して問題解決手法として模倣システムの構築を目指した。具体的には、模倣システムの構築が推論方法や知識表現や知識ベースや学習・知識獲得に深くかかわる知識ベース型システム構成の問題となる。この考えに従い、知識ベース、学習方法、推論方法の構築という角度から人間の障害物回避戦略の模倣システムの構築を行っていく。

参考文献

1. 國吉康夫：模倣ロボットは人間的知能を獲得するか？, 学術月報, Vol.53, No.9, pp.11-18, 2000.
2. Stefan Schall: Is Imitation Learning The Route to Humanoid Robots?, Trends in Cognitive Science, Vol. 3, pp.233-242, 1999.
3. 岡田慧, 鈴木義久, 國吉康夫, 稲葉雅幸, 井上博允：視覚による人間動作認識と全身行動表現に基づくヒューマノイドの行動獲得, 記憶, 再現, 第19回日本ロボット学会学術講演会, pp.431-432, 9月18-20日, 2001.
4. 吉川雄一郎, 浅田稔, 細田耕：エピソード幾何を利用した呈示者視野復元に基づく模倣の実現, 日本ロボット学会誌, Vol.22, No.1, pp.68-74, 2004.
5. Dautenhahn, K. and Nehaniv, C.L. : Imitation in animals and artifacts, Cambridge, Massachusetts: the MIT press, 2002.
6. 板倉昭二：模倣は意図を読むことから始まる, ATR 研究所.
7. 銅谷賢治, 五味裕章, 阪口豊, 川人光男: 脳の計算機構ボトムアップ・トップダウンのダイナミクス, 朝倉書店, 2005.
8. K. Forbus and J. de Kleer: Building Problem Solvers, MIT Press, 1993.

第4章 模倣に向ける知識の表現法と獲得法

第1節 はじめに

障害物回避行動の模倣は充分期待できるものであるが、模倣システムで再現した時に得られる効果が如何なるものであるかは未知のものである。従って、行動後の結果として残されたデータからより効果的に障害物回避戦略を抽出することから研究に着手した。第3章での考えに従い、問題解決手法として模倣システムの構築が知識表現や知識ベースや学習・知識獲得や推論方法に深くかかわる知識ベース型システム構成の問題となる。

本章では、知識の表現法・獲得法の角度から模倣システムの構築を試みた。ここで、if-then 型宣言的知識の表現法を用いて、人間の障害物回避戦略を表すことにする。知識の表現法を通して人間の障害物回避戦略を表す有効性・適応性を確認した。そして、知識の問題解決能力をチェックするために、遺伝アルゴリズム GA を用いた知識の評価法を提案した。現有知識の不足と知識獲得の困難という問題点を確認したうえ、主観的な経験に依頼しない知識獲得の重要性を確認した。多変量で記述しにくい場合でも高速で知識獲得を実現するために、改善した距離型ファジィ推論法の学習ルゴリズムを提案した。距離型ファジィ推論法のオリジナルな学習アルゴリズムに知識（ルール）の生起確率を導入することにより、知識の最適化を行った。最後に、シミュレーションにより提案する新学習ルゴリズムの有効性を示す。

第2節 知識の表現法

問題解決システムにおける知識の表現手法の一つとして、if-then 型宣言的知識表現がよく使われており[1]、これは、様々な状況に応じて物事を細かく表現することや、個別な状況に応じて異なる判断を下す、といったような人間が普段自然に行われている行動をよく表しているからである。前件部と後件部をファジィ集合とすれば、Yes か No で表現するはっきりした概念は勿論のことであるが、曖昧な概念の定量化表現も可能である。そのため、ファジィ集合を利用した if-then 型宣言的知識表現法は、高次脳機能の工学実現においては有力な表現法の一つであると考えられる。if-then による知識表現を利用した、ファジィ推論法として、Mamdani の推論法[2]、関数型ファジィ推論法[3]、簡略型ファジィ推論法[4]を始め、最近ファジィ集合間の距離情報に基づく距離型ファジィ推論法も提案されている[5]。これらの推論モデルは、二値論理における知識だけではなく、あいまいな概念も取り扱えるので、実システムへの応用実績を数多く持っており、脳の推論機能のある側面を実現していると言える。したがって、ファジィ集合を利用した if-then 型宣言的知識表現を採用して人間の障害物回避戦略を表示すると考えられる。また、運転シミュレータにおいて、操作者はハンドルとアクセル・ブレーキを利用してエージェントをゴールまでに障害物を回避しながら運転する。もし表現された知識を運転シミュレータにおけるエージェントに付加できれば、エージェントの行動効果により、知識の正確性をチェックし、更なる知識獲得も役に立つことができる。

人間の障害物回避戦略を表示するために、知識獲得の入力空間と知識の表現は以下に詳しく紹介される。

第1項 知識獲得の入力空間

第2章で図2-5の運転制御のブロック線図に示すように、環境情報 X とエージェントの状態情報 Z と人間の制御情報 U により、人間の障害物回避戦略の抽出、つまり知識獲得を行うことが可能となる。しかし、知識表示に対して、エー

エージェントと環境との関係を定量的に表す必要がある。そこで、入力空間(2-2)の環境設置とエージェントの運動特性(2-1)式により、図4-1の示すようにエージェントと環境との関係を代表する物理量が知識獲得の入力空間として計算される、つまり(4-1) 式で表す。

$$\text{Input space} = \{ [dg, vg]^{n_g}, [do, vo, dc, vc, vrx, vry]^{n_o}, U \} \quad (4-1)$$

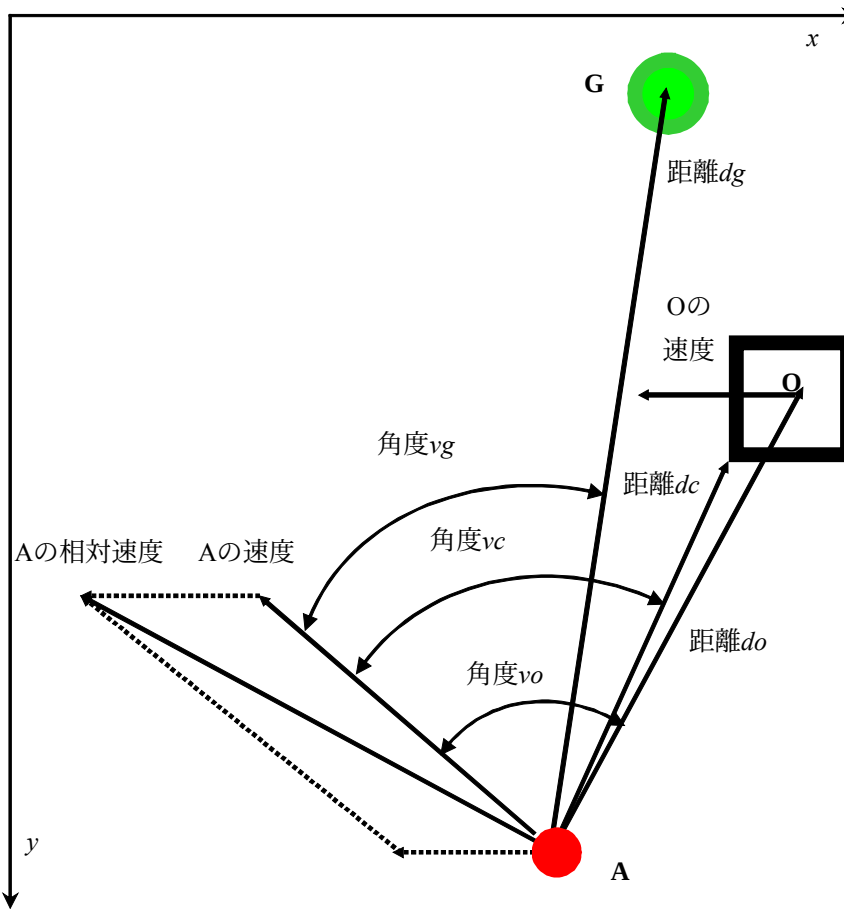


図 4-1 物理量の説明図(G:ゴール, A: エージェント, O:障害物)

ただし、

$do / dc / dg$: エージェントと対象(障害物/エージェントと最寄りの障害物の点/ゴール)との距離

$vo / vc / vg$: エージェントの速度方向とエージェントから見える2つ対象(エージェントと障害物/エージェントと最寄りの障害物の点/ゴール)の位置ベクトルとのなす角

vr_x/vr_y : エージェントと障害物との水平相対速度/垂直相対速度

知識獲得の入力空間の長さは $2n_g + 6n_o + 2$ であるが、一般に、1つのゴールと幾つの障害物、つまり $n_g = 1$ と $n_o \geq 1$ の場合に対して、入力空間の長さは $6n_o + 4$ になる。

第2項 ルールベース

通常は、知能エージェントへの Rasmussen モデルは、人間行動の制御に関する 3 階層モデルを提唱しており、知識ベースの行動、ルールベースの行動、技能ベースの行動を説明している [6,7]。議論対象としているエージェントの階層知識を論じるために、果たして、Rasmussen のモデルを引用するのが最適かどうかは自明ではないが、何らかの指針を示してくれるものと思われる。ここでは、エージェントの階層知識がマクロ処理レベルとマイクロ処理レベルに分けられる。

マクロ処理レベルに対しては、障害物回避とゴール到着という基本の行動目標を考慮するうえに、障害物回避に高い優先級をゴール到着より持たせる。具体的に、もしエージェントの前方に障害物が存在しない、つまりエージェントの速度方向とエージェントから見えるエージェントと障害物における凸点との位置ベクトルとのなす角はすべて相同の符号であれば、エージェントがゴール到着への前進を行う。またはそうなく、障害物が左前方にある場合に障害物回避への右転を行う、または障害物が右前方にある場合に障害物回避への左転を行う。

マイクロ処理レベルに対しては、従来のプロダクションシステムのとおり、知能エージェントの推論エンジンとして推論法に基づいてエージェントの行動を決める。ここでは、エージェントの推進力と回転角度とを決める。そして、エージェントの運動方程式(2-1)により、エージェントの運動軌道を計算する。必要な組み立てとして、知能エージェントの知識は予め経験的に引き出されることができる。そのような知識を先験的知識と呼ぶ。具体的に、エージェントの先験的知識は表 4-1, 4-2, 4-3, 4-4 に記述される。推進力のルールは表 4-1 に記述されるが、回転角度のルールは 3 つの種類のサブベースに分類される：ゴール到着への前進ルール(表 4-2)、障害物回避への左転ルール(表 4-3)と右転ルール(表 4-4)。ただし、ルールの前件部の詳細は以下に説明する。

force, direction : それぞれエージェントが受ける推進力と回転角度

$do / dc / dg$: エージェントと障害物/障害物における最寄りの凸点/ゴールとの距離

$vo / vc / vg$: エージェントの速度方向とエージェントから見えるエージェントと障害物/障害物における最寄りの凸点/ゴールとの位置ベクトルとのなす角

表 4-1 推進力ルール(*if do then force*)

do	S	M	B
$force$	Z0	PS	PB

表 4-2 前進ルール(*if vg then direction*)

vg	NB	NS	Z0	PS	PB
$direction$	NB	NS	Z0	PS	PB

表 4-3 左転ルール(*if vc and dc then direction*)

$dc \backslash vc$	NB	NS	Z0	PS	PB
S	PB	NB	NS	Z0	PB
M	PB	NB	NS	Z0	PB
B	NB	NS	Z0	PS	PB

表 4-4 右転ルール(*if vc and dc then direction*)

$dc \backslash vc$	NB	NS	Z0	PS	PB
S	NB	Z0	PS	PB	NB
M	NB	Z0	PS	PB	PB
B	NB	NS	Z0	PS	PB

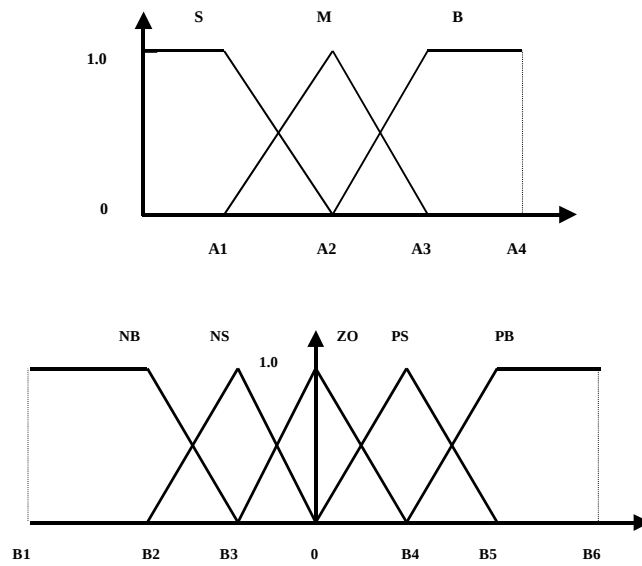


図 4-2 ファジィラベルとメンバーシップ関数

ファジィラベルは図 4-2 に示すようなメンバーシップ関数に応じ，関連パラメータは表 4-5，4-6 に経験的に確定される．

表 4-5 メンバーシップ関数のパラメータ

パラメータ 種類	A1	A2	A3	A4	単位
do, dc, dg	100	350	600	1833	pixel

(最大距離=1833 pixel)

表 4-6 メンバーシップ関数のパラメータ

パラメータ 種類	B1=-B6	B2=-B5	B3=-B4	単位
vo, vc, vg	-180	-90	-45	度
$direction$	-90	-60	-20	
$force$	-18	-12	-6	N

(方向の成す角: $[-180, +180]$, 回転角度: $[-90, +90]$, 推進力: $[-18, +18]$, 左-, 右+)

段階的な知識獲得を行うために、前述した先験的知識と対比してみると、知識獲得・学習法により派生される知識は、新しいルールベースに保存される。また、主観的に描きにくい属性として、エージェントと障害物との水平相対速度 vr_x と垂直相対速度 vr_y が新しいルールの表示に追加される：

$$\text{if } [do, vo, dc, vc, vr_x, vr_y]^1 \cdots [do, vo, dc, vc, vr_x, vr_y]^{n_o} [dg, vg] \text{ then } force / direction \quad (4-2)$$

ただし、 n_o は環境における障害物の個数であり、他の符号は 2.2 節と同じ意味を持つ。

関連パラメータは表 4-7 に経験的に確定される。

表 4-7 メンバシップ関数のパラメータ

パラメータ 種類	B1=-B6	B2=-B5	B3=-B4	単位
vr_x, vr_y	-557	-200	-100	pixel/sec

(最大相対速度=557 pixel/sec)

上述の結果から、表現された知識をエージェントに付加すると、エージェントの行動効果により、知識の正確性をチェックできるが、どのような効果を問題解決に達成するかということがまだ明確ではなく、更に如何に知識獲得を行うことも曖昧になる。

第3節 知識の評価法

以上に知識の表現法を用いて、障害物回避戦略としてファジィルールの収集を利用して模倣システムを構築することができる。そのようなルールの収集が模倣システムの構築に強い融通性を持つと確認した。しかしながら、問題解決を達成する効果はまだ明確ではない。また、新しいルールの追加とともに、結果的システムが必ず膨れ上がって能率が悪くなる。関連な理論的研究として、Wang [8]はファジィルールに基づくファジィシステムを普遍的な近似関数に対応され、しかも理論的に証明された。[9]では、不合理な複雑性を持つ構造は「over-fitting」という偏り効果を引き出しており、システムの近似性能を損なう可能性があるため、単純な構造はロバスト特徴を持つシステムの構造に役に立つと指摘された。従って、最適な知識組み合わせを求め、模倣システムを構築するために、獲得されたルールを評価して余計な部分を抜ける必要があると考えられる。

通常は、余計なルールの存在は偶然な問題ではなく、以下に示すような原因で起きる。

- 1) 知識を積み重ねるとともに、相同や相似や矛盾な知識が必然的に成り果てる。同時に、ルールの数も次第に増える。
- 2) 異なる学習対象としての人は異なる障害物回避戦略を持つ。全ての人は優れた運転手ではなく、結果的には、全てのルールは最適なルールではない。
- 3) 障害物回避に対して、優れた運転手さえとしては障害物回避能力の限界を持つことがわかる。結果的に、学習対象から獲得した全部のルールは成功的障害物回避に役に立つことが保証されない。

以上の原因に鑑みて、ファジィシステムの構造最適化の見方から、遺伝アルゴリズム GA 進化による知識の評価法を提案した。この方法で、最適なルール組み合わせを求め、ファジィシステムに定量的評価基準を提供し、更なる知識獲得に指導的役割を果たすことができる。

第1項 GA 進化による知識の評価法

遺伝アルゴリズム (Genetic Algorithm, GA) は、複数ある解の中から最適なものを選び出すという問題の解決方法として適していると考えられる。現有ルール集合による最適なルール組み合わせを考えると 100 万通り以上になる。その中から、最適なルール組み合わせを決定しなくてはならない。そのような場合に、GA は有効であると考えられる。GA の基本的な流れは個体の評価をもとにして、優秀なものを選んでいくというものであるが、アプローチの仕方の違いから、Pittsburg アプローチと Michigan アプローチに分けることができる [10]。

ここでは、既知のルール集合から代表的なルール集合を探すために、進化過程において選定されたルール集合の評価を行うという Pittsburgh アプローチを利用してファジィシステムを最適化する。GA の手法を自然に利用した方法で IF-THEN ルールの集合を個体として、さまざまなタイプの IF-THEN ルールの組み合わせを探す。交差や突然変異の方法としては、IF-THEN ルールの集合をルール単位で遺伝子座として扱い、GA オペレータを適用する。個体数が増えると遺伝子座に含まれるルール数は非常に大きなものになってしまう。しかし、GA の形態としては非常に自然で扱いやすいものとなる。特に、個体の評価法としては実際にそれぞれのルール集合を適用した場合の評価を、各個体となるルール集合ごとに求めるので、システム全体に関する評価値かついくつかの評価基準を受け取ることができる。式(4-3)のように、3つの評価基準を持つ適合度関数が定義される。結果的なルール集合に基づく推論性能を考察しながら、ファジィシステムの構造の複雑性も考察する。

$$E = \eta E_{\text{success}} \cdot E_{\text{averrule}} \cdot E_{\text{neverrule}} \quad (4-3)$$

ただし、 η は指定係数である。

第一項として E_{success} はあるルール集合を知識として付加するエージェントが評価方案に基づいて障害物回避を行う成功率を表す。

$$E_{success} = \sum_{k=1}^L P_k / L, P_k = \begin{cases} 1, & \text{success} \\ 0, & \text{fail} \end{cases} \quad (4-4)$$

ただし,

L : 評価方案における運転タスクの総数

P_k : エージェントが評価方案における k 番目の運転タスク環境に障害物回避を行う結果

あるルール集合を知識としてエージェントに付加したら, そのルール集合に基づく推論性能を考察するために, 通常, 幅広い運転タスクを含む評価方案のみで, 各ルールの有効性を評価することが可能となる. 図 4-3 の示すように, シミュレータ環境において, 環境対象の位置関係を縦横 $m \times n$ 通り変化させる. 評価方案にける 1 つの運転タスクはゴールや障害物やエージェントという環境対象の位置に対応する. もちろん, 通りが細かければ, 細かいほど, 運転タスクの個数が多くなり, 各ルールの有効性がもっと評価されることができる. しかし, m や n の増大につれて, 評価方案にける運転タスクの個数も指数的に増え, 最適なルール集合を探す進化過程において大変な時間がかかる.

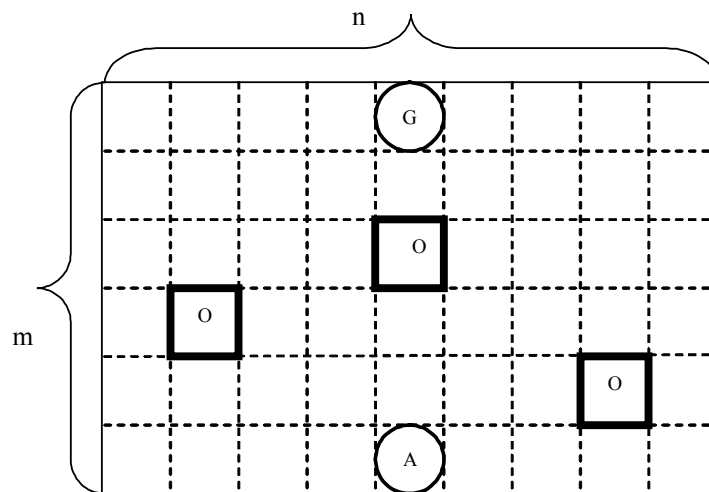


図 4-3 運転タスク (G:ゴール, A: エージェント, O:障害物)

エージェントが障害物をなるべく回避してゴールに到着するということについて考察するため, 評価方案が簡略化される: 図 4-3 に示すようにエージェ

ントの出発点を最低の通りに、ゴールを最高の通りに、障害物を両方の間に設定する．結果的に、評価方案におけるタスクの数 L は $(m * n)P(2 + n_o)$ から $n * n * ((m - 2) * n)Cn_o$ まで減少する．ただし、 n_o は環境に存在する障害物の個数を表す．

第二項として $E_{averrule}$ はあるルール集合を知識として付加するエージェントが評価方案に基づいて障害物回避を行う過程に、活発されたルールの平均数の逆数を表す．

$$E_{averrule} = L / \left(\sum_{t=1}^L \left(\sum_{j=1}^{N_e^t} R_j^t / N_e^t \right) \right) \quad (4-5)$$

ただし、

R_j^t : 評価方案の t 番目のタスクに対して、 j 番目の推論エピソードにおいて活発されたルールの数, ≥ 1

N_e^t : 評価方案の t 番目のタスクに対して、推論エピソードの総数, ≥ 1

L : 同上、評価方案のタスクの総数

もしあるルール集合に基づく推論エピソードごとに活発されたルールの数は1であれば、そのファジィシステムが高い近似性をもち、最も分かりやすい．

第三項として $E_{neverrule}$ はあるルール集合を知識として付加するエージェントが評価方案に基づいて障害物回避を行う過程に、いつでも活発されなかったルールを処理する．いつでも活発されなかったルールがファジィシステムの推論性能に悪い影響を与えるので、できるだけ少なくさせる必要がある．

$$E_{neverrule} = N / (R_z + 1) \quad (4-6)$$

ただし、

R_z : 評価方案におけるいつでも活発されなかったルールの数, ≥ 0

N : ルールの総数, ≥ 1

もしあるルール集合に基づくいつでも活発されなかったルールの数は0であれば、そのファジィシステムが余計なルールを含まない．

第2項 シミュレーションによる検討

表 4-8 には経験的に獲得されたルール集合を示す．そのルール集合を知識としてエージェントに付加したら，そのルール集合に基づく推論性能を考察するために，1つの静的な障害物を有する環境において，図 4-4 の示すように環境対象の位置関係を縦横5×3通り変化させると，評価方案に基づいて障害物回避を行った成功率は成功率=43.21%である．

表 4-8 ルール集合

制御の種類 ルールの種類	回転角度	推進力
基本ルール	35 個	3 個
性能	成功率=43.21% (35/81)	

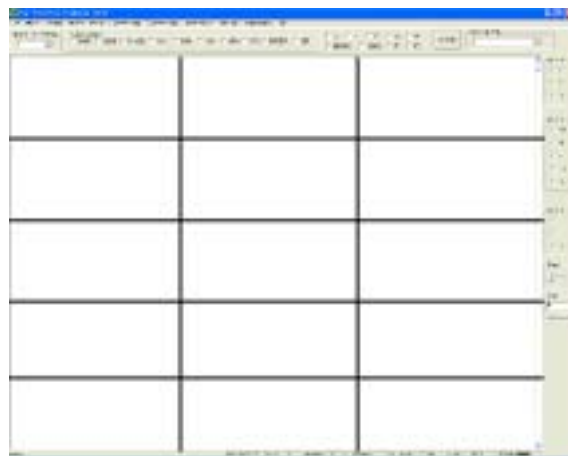


図 4-4 評価方案にける運転タスク

まず，知識進化については，GAのパラメータは以下に指定される：交叉の起こる確率を50%，交叉の種類を一様交叉，突然変異の起こる確率を10%とする．計算に影響する主要要素として，個体群のサイズと評価方案における運転タスクの総数が非常に重要であると思われる．個体群のサイズの増大につれて，世代あたりの時間が増えており，最適なルール集合を必ずしもより速く探さない．一方，評価方案における運転タスクの総数の増大につれて，ルールの有効性をもっと完全に考察されることが出来る．そこで，個体群のサイズと評価方案に

おける運転タスクの総数の間に適合な割合を確定できれば、ルールの有効性も考察し、計算時間も減少することができる。

個体群のサイズと評価方案における運転タスクの総数をそれぞれ調整することで、図4-5には異なる割合によるルールの進化過程を示す。評価方案における運転タスクの総数の増大につれて、推論に基づく運転成功率が下がることがわかる。しかしながら、個体群のサイズとタスクの総数の反比例関係の場合に、例えば、size=20とsize=60の場合に、size=20の成功率がsize=60の変化にほとんど近くなる。この場合に、もし個体群のサイズが減少されれば、計算時間が必ず大きく減少される。従って、快速的な計算を保証するために、全体に個体群のサイズを40個、評価方案当たりの81タスク($m=5, n=3$)として進化計算を行う。

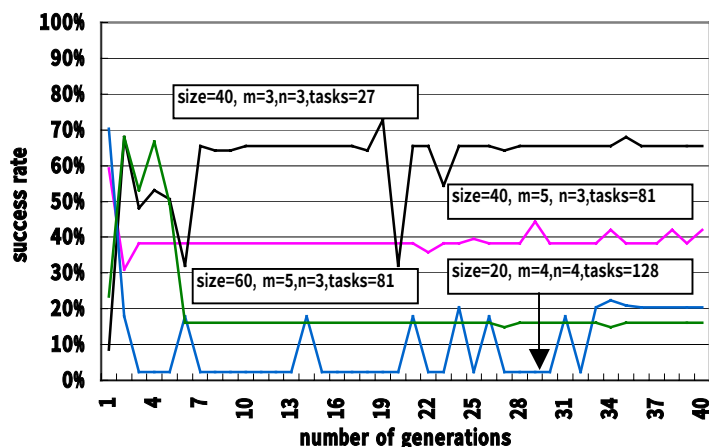


図 4-5 個体群のサイズと評価方案における運転タスクの総数

次に、GA 進化により表 4-9 における知識の評価を行う。結果的なルールは表 4-15 に纏められる。図 4-6 と 4-7 には、ファジィシステムの構造の複雑性を評価する 2 つ基準として、活発されたルールの平均数といつでも活発されなかったルールの数との進化過程を示す。図 4-6 の示すように、活発されたルールの平均数は 1.60 であり、良い近似性および理解性を持つファジィシステムを示す。図 4-7 の示すように、いつでも活発されなかったルールの数は 1 であり、ファジィシステムにおけるほとんどゼロの余計を示す。結果的ルールの性能がまだ低いものの、最初ルール集合に比べて、もっと少ないルールを用いて相似

の性能を達成したことが分かる．もし障害物回避失敗場合の記録により知識獲得を行っていけば，有効なルールを更に抽出することができる．

表 4-9 ルール集合の進化結果

制御の種類 ルールの種類	回転角度	推進力
基本ルール	24 個	3 個
性能	成功率=41.98% (34/81) 活発されたルールの平均数=1.60 いつでも活発されなかったルールの数=1	

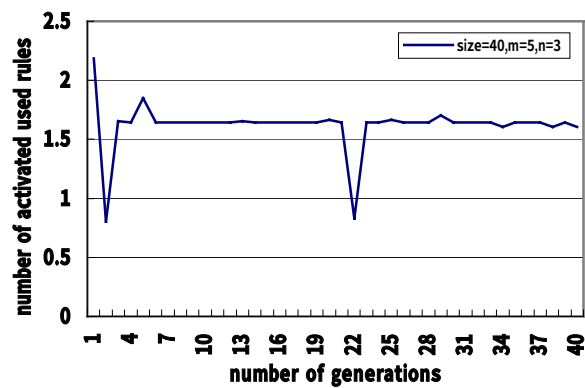


図 4-6 エピソードあたりの活発されたルールの平均数

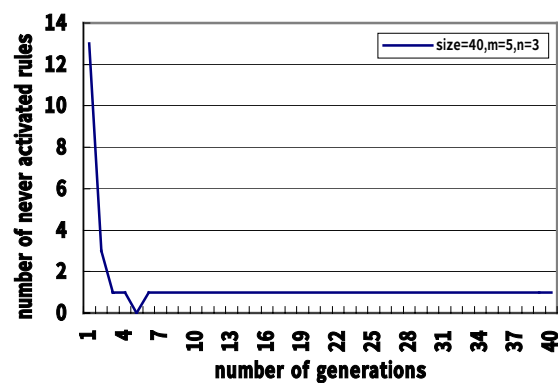


図 4-7 いつでも活発されなかったルールの数

第4節 データ学習による知識獲得法

第3節の知識の評価法により、現有知識の有効性を複数の評価基準で定量的に評価することができる。同時に、経験的に抽出されたルールの性能がまだ低く、複雑な環境において先験的知識を記述しにくい、結果的に知識獲得が能率的ではないことが分かる。そこで、多変量、つまり言語で主観的に記述しにくい場合でも豊富な知識を獲得するために、高速・有効な知識獲得への学習法の提案を目指している。行動後の結果として残されたデータから如何より効果的・効率的に人間の障害物回避戦略を獲得することも人間の障害物回避行動の模倣を実現する基盤である。従って、適切な学習方法の選択は問題解決に非常に重要である。ここでは、距離型ファジィ推論法の学習アルゴリズムを採用する。この方法は距離型ファジィ推論法を元にして発展して推論に直接に関係がある[11]。また、この方法は3つの独特な特徴を持つ：1)結果的に学習誤差が任意に指定でき、0を含む。2)GAかNNなど他の学習方法に比べて、有限個のデータによる学習時間がほとんど要らなく極めて速い。3)数値型データ学習に適合する。離散なデータからルールへの変換過程が避けられるから、通常の方法のようにルール生成がメンバーシップ関数の影響を大きく受けることがない。

しかし、距離型ファジィ推論法の学習アルゴリズムも欠点を持つ。分離規則を厳密に満たすため、学習過程に更新の見方から矛盾ルールの後件部を更新する。ただし、矛盾ルールとは同様な前件部かつ幾つの異なる後件部補選値を持つルールだと定義される。結果的に、獲得されたルールは、学習データの順番、つまり、時間に関係があり、最新な補選値が矛盾ルールの後件部に選択される。更に、学習後のルールに基づくファジィ推論を関数として評価する際に、初期教師データの出力誤差が末期教師データの出力誤差より相対に大きい現象が存在する。実際に、データ学習は時間に関係がなく、学習対象の状態に関係があるべきである。従って、学習過程に得る矛盾ルールの後件部補選値の分布情報により、最適な補選値の選択を着目する。図4-8に代表的な離散補選値

の分布は確率の定義に従って1に規格化される．ただし、 B^k は矛盾ルール k の後件部であり、 B_l^k は矛盾ルール k における l 番目の補選値であり、 $l \in \{1, 2, \dots, L^k\}$ 、 L^k は矛盾ルール k の後件部における補選値の数であり、 P_l^k は補選値の生起確率である．

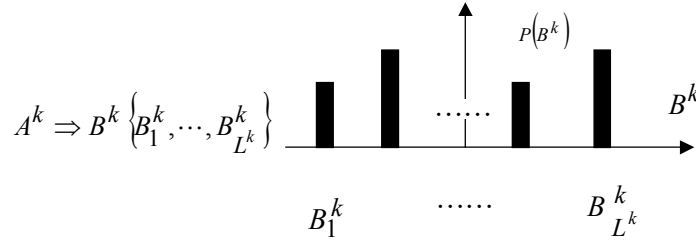


図4-8 矛盾ルールにおける補選値の生起分布

学習対象の相違により、離散補選値の生起確率の分布モデルが異なる．生起確率の分布モデルが不明の場合に、特に、人間の行為特性に対して、誤差分布のモデルが常に平均値を中心とした対称形ではないことを考慮して、矛盾ルールの後件部が最大生起確率 $\max\{P(B_1^k), P(B_2^k), \dots, P(B_{L^k}^k)\}$ を持つ補選値に属すると判断すれば、判断の誤りの誤差が一番少なくなる．従って、元のアロリズムにおけるルール更新処理を改善するため、ルールの前件部との距離誤差の範囲内に可能な後件部の補選値に生起確率を付けると、最大生起確率を持つ後件部の補選値を結果ルールの後件部とする．まず、学習と推論への距離型ファジィ推論法を以下に紹介する．

第1項 距離型ファジィ推論法

通常な適合度型ファジィ推論方法に比べて、距離型ファジィ推論法 (Distance-Type Fuzzy Reasoning method, DTFR) はファジィ集合間の距離を利用する上に推論を行い、より詳細な説明については文献[5]を参照されたい．

通常、次の推論を対象とする．

$$R^i : \text{if } x_1 = A^{i1}, x_2 = A^{i2}, \dots, x_m = A^{im} \text{ then } y = B^i \quad (4-7)$$

$$\frac{i = 1, 2, \dots, n}{\text{---}}$$

$$\text{事実: } x_1 = A^1, x_2 = A^2, \dots, x_m = A^m$$

結果: B

ただし, $x_1, \dots, x_m$ と y とはそれぞれ入力変量と出力変量を表す. $A^{i1}, \dots, A^{im}, B^i$ は相応なファジィ集合であり, n はルールの数であり, m は前件部の数である.
 $i = 1, 2, \dots, n, \quad j = 1, 2, \dots, m.$

距離型ファジィ推論法は次の四つのステップから構成される.

STEP1: 事実中のファジィ変数 A^j とルール中のファジィ変数 A^{ij} との距離 $d_{ij}(A^j, A^{ij})$ を計算する(距離の詳細計算については「付録」を参照される).

STEP2: 事実と i 番目ルールとの距離を計算する.

$$d_i = \sum_{j=1}^m d_{ij}(A^j, A^{ij}) \quad (4-8)$$

STEP3: $\forall \alpha \in [0, 1]$ に対して, 次のように推論結果 B の α -レベル集合 B_α を求める.

$$B_\alpha = [\inf(B_\alpha), \sup(B_\alpha)] \quad (4-9)$$

$$\inf(B_\alpha) = \frac{\sum_{i=1}^n \left[\inf(B_\alpha^i) \prod_{j=1, j \neq i}^n d_j \right]}{\sum_{i=1}^n \prod_{j=1, j \neq i}^n d_j}$$

$$\sup(B_\alpha) = \frac{\sum_{i=1}^n \left[\sup(B_\alpha^i) \prod_{j=1, j \neq i}^n d_j \right]}{\sum_{i=1}^n \prod_{j=1, j \neq i}^n d_j}$$

STEP4: 合成原理(分解定理とも呼ばれる)により推論結果 B を求める.

$$B = \bigcup_{\alpha} B_\alpha \quad (4-10)$$

距離型ファジィ推論法は次のような特徴を持っている.

- 1) 分離規則は満たされている. ルールの物理意味が明確であり, 学習により最適なルールを獲得することが容易である.

- 2) 結果Bは凸なファジィ集合となる．推論結果はファジィ数になるので，実数の拡張の形で数理論理の立場から厳密に議論できる．
- 3) 漸近特性を持っている．推論の方向性が大雑把に予測できる．
- 4) 疎なルールにも適用できる．実際のシステムにおいて少ないルールで同じ推論特性が得られる．
- 5) 推論結果に整合性がある．推論結果と後件部変数が同値性を持っていれば，推論結果としての重心値を求める積分計算は必要ではなく，簡単な代数演算で求められる．

第2項 学習アルゴリズムの提案

改善した学習アルゴリズムは次の5つのステップから構成され、 $\varepsilon \geq 0$ は任意に指定された誤差である。

STEP1: 与えられた教師データ $\{A_i^j(k), B_i(k)\}$ に対して、 $A_i^j(k)$ を事実として推論システムに入力して推論結果 $B(k)$ を求める。 $j = 1, 2, \dots, m$, $k = 1, 2, \dots, L$, m は前件部の数であり、 L は教師データの総数である。

STEP2: 教師データ $B_i(k)$ と推論結果 $B(k)$ との距離 $d(B_i(k), B(k))$ を計算する。

STEP3: もし、 $d(B_i(k), B(k)) \leq \varepsilon$ であれば、なおかつ、 $\exists q \in \{1, 2, \dots, n\}$ 等式

$$\sum_{j=1}^m d(A_i^j(k), A^{qj}) = 0 \text{ が成立すれば、} q \text{ 番目のルールの後件部が応じる補選値 } B_i(k)$$

の生起確率 P_l^q を更新する、 $l \in \{1, 2, \dots, L^q\}$, n は現有ルールの数であり、 L^q は q 番目のルールの後件部における補選値の数である。その他の教師データに対して STEP1 に戻る。

STEP4: もし、 $d(B_i(k), B(k)) > \varepsilon$ であれば、なおかつ、 $\exists q \in \{1, 2, \dots, n\}$ 等式

$$\sum_{j=1}^m d(A_i^j(k), A^{qj}) = 0 \text{ が成立すれば、} q \text{ 番目のルールの後件部を } B_i(k) \text{ に更新し、後$$

件部補選値 $B_i(k)$ の生起確率 P_l^q を更新してから STEP1 に戻る。

STEP5: $A_i^1(k), \dots, A_i^m(k) \Rightarrow B_i(k)$ を新しいルールとしてルール集合に追加し、後件部補選値 $B_i(k)$ の生起確率 $P_l^n = 1$ を初期化してから STEP1 に戻る。

STEP6: 矛盾ルールにおける後件部補選値の生起確率により、最大生起確率を持つ補選値を後件部とする。

図4-9は学習アルゴリズムの流れ図を示す。結果的に、IF-THEN型ファジィルールが生成される。

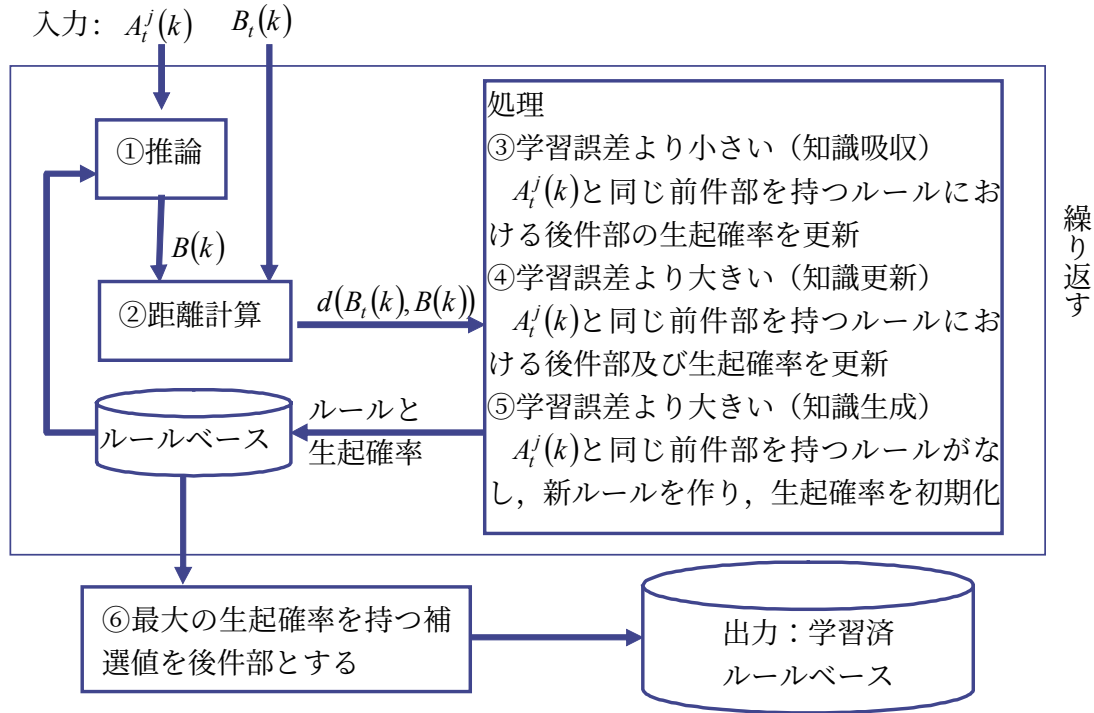


図4-9 学習アルゴリズムの流れ図

本アルゴリズムは元のアロリズムの特徴を継承して定理1と定理2に纏め、改善部分により定理3に説明する。

定理1: L 個の与えられた教師信号 $\{A_i^j(k), B_i(k)\}$ に対して、本学習アルゴリズムに基づいて有限時間内に新しいルールの追加、古いルールの構成、間違ったルールの訂正により、学習効果として式 $d(B_i(k), B(k)) \leq \varepsilon$ を満たす。ただし、 $\varepsilon \geq 0$ は任意に指定された誤差であり、 $B(k)$ は $A_i^j(k)$ を事実として推論システムに入力された時の推論結果を表す。 $d(B_i(k), B(k))$ はファジィ集合 $B_i(k)$ と $B(k)$ の距離を表す。 $j = 1, 2, \dots, m$, $k = 1, 2, \dots, L$ 。

定理2: n 個のルールにより推論された結果を B_n として、もし学習アルゴリズムが実行される過程において、新しいルール $A^{(n+1)1}, \dots, A^{(n+1)m} \Rightarrow B^{n+1}$ が追加される場合、同一の事実に対しての推論結果 B_{n+1} は新たに推論のSTEP1から求める必要がなく、(4-11)式 of ファジィ数の演算により簡単に求めることができる。

$$B_{n+1} = \frac{1}{1 + W_{n+1}} (W_n B_n + B^{n+1}) \quad (4-11)$$

ただし, $q = 1, 2, \dots, n+1$.

$$W_{n+1} = \frac{1}{d_{n+1}} \sum_{i=1}^n \frac{1}{d_i}$$

$$d_q = \sum_{j=1}^m d(A^j, A^{qj})$$

定理 3: 同様な前件部且つ幾つの異なる後件部補選値を持つ矛盾ルールに対し、本学習アルゴリズムは厳密に分離規則を満たし、学習過程において得た矛盾ルールの後件部補選値の分布状況により最大生起確率を持つ補選値を後件部として選択する。すると、結果的なファジィシステムは最大信用な発生事象を代表するルールから組み立て、現有データに基づいて分離規則を満たす最大信用なシステムである。

証明:

ここでは確率理論に基づいて、学習アルゴリズムにより獲得したルールの信用程度のみを議論している。

学習アルゴリズムにより得た n 個のルール $\{R^1, R^2, \dots, R^n\}$, ただし, ルール $R^i: A^i \Rightarrow B^i$, A^i, B^i はそれぞれ i 番目のルールにおける前件部と後件部を表す。それらの前件部や後件部の数は ≥ 1 である。

ルールの前件部が確率空間 Ω の分割であるとき、以下の等式が成り立つ。

$$\sum_{i=1}^n P(A^i) + \sum_{i=n+1}^N P(A^i) = 1 \quad (4-12)$$

ただし, $\{A^1, A^2, \dots, A^n\}$ は n 個の種類の前件部であり, $n+1$ から N までは前件部の未知種類である。且つ,

$$\lim_{n \rightarrow N} \sum_{i=1}^n P(A^i) = 1 \quad (4-13)$$

n 個のルールに基づく任意の出力事象 B の信用程度に対して、以下の式が成り立つ。

$$P(B) = P(A^1)P(B|A^1) + P(A^2)P(B|A^2) + \dots + P(A^n)P(B|A^n) \quad (4-14)$$

ただし,

$P(B)$: 任意の出力事象 B の信用程度

$P(A^i)$: 前件部 A^i としてのルールの信用程度, 即ち, 学習統計に基づく前件部 A^i としてのルールの生起確率

$P(B | A^i)$: 前件部 A^i による任意の出力事象の信用程度

i 番目のルールにおける後件部が応じる補選値 B_k^i の生起確率は $P(B_k^i | A^i)$ であり, 且つ補選値の独立生起が存在するので, 以下の式が成り立つ. ただし, $k \in \{1, 2, \dots, L^i\}$, L^i は i 番目のルールにおける後件部の補選値の総数.

$$\begin{aligned} P(B | A^i) &= P(B_1^i \cup B_2^i \cup \dots \cup B_{L^i}^i | A^i) \\ &= P(B_1^i | A^i) + P(B_2^i | A^i) + \dots + P(B_{L^i}^i | A^i) \\ &\leq 1 \end{aligned} \quad (4-15)$$

矛盾ルールが存在しない場合に, $P(B | A^i) = 1, \forall i \in [1, \dots, n]$ が成り立つ上に,

$$P(B) = \sum_{i=1}^n P(A^i).$$

矛盾ルールが存在する場合に, 元のアルゴリズムは, 矛盾ルールを更新するために最新の補選値をルールの後件部として選択する. これは手当たり次第に補選値から 1 つの補選値を選択することに等価する. 以下のように, $r \in [1, \dots, L^i]$, 補選値から 1 つの補選値 B_r^i を i 番目のルールの後件部として任意に選択する.

$$\begin{aligned} P(B | A^i) &= P(B_1^i \cup B_2^i \cup \dots \cup B_{L^i}^i | A^i) \\ &= P(B_1^i | A^i) + P(B_2^i | A^i) + \dots + P(B_{L^i}^i | A^i) \\ &\geq P(B_r^i | A^i) \end{aligned} \quad (4-16)$$

しかし, 改善した学習アルゴリズムは学習過程において得た矛盾ルールの後件部補選値の分布状況により最大生起確率を持つ補選値を後件部として選択する. 以下の式で表示する.

$$\begin{aligned} P(B | A^i) &= P(B_1^i \cup B_2^i \cup \dots \cup B_{L^i}^i | A^i) \\ &= P(B_1^i | A^i) + P(B_2^i | A^i) + \dots + P(B_{L^i}^i | A^i) \\ &\geq \max\{P(B_1^i | A^i), P(B_2^i | A^i), \dots, P(B_{L^i}^i | A^i)\} \\ &\geq P(B_r^i | A^i) \end{aligned} \quad (4-17)$$

改善した学習アルゴリズムは $P(B_r^i | A^i)$ を用いるのではなく, $P(B | A^i)$ の代わ

りに $\max\{P(B_1^i | A^i), P(B_2^i | A^i), \dots, P(B_{L_i}^i | A^i)\}$ を用いる。従って、 $P(B | A^i)$ の近似値が増える。同時に、元のアゴリズムは $P(A^i)$ を考慮しなくても、現有データに基づいて分離規則を満たしており、改善アゴリズムは元のアゴリズムと $P(A^i)$ に一致する、 $\forall i \in [1, \dots, n]$ 、更に式 (4-12) に導入して、 $P(B) \leq \sum_{i=1}^n P(A^i)$ がやっぱり成り立つとしても、 $P(B)$ の値が増える。従って、改善アゴリズムを用いて生成するファジィシステムの信用程度が増える。結果的なファジィシステムは現有データに基づいて分離規則を満たす最大信用なファジィシステムである。

第3項 シミュレーションによる検討

図 4-10 に示すように、静的障害物が 3 つ存在する環境において、操作者が障害物を回避して目標点を成功に到着した。そして、ファジィルールの獲得に対して、改善した距離型ファジィ推論法の学習アルゴリズムが採用された。本学習アルゴリズムにより、行動データから以下のようなルールが獲得された。ただし、符号は 2.2 節と同じ意味を持つ。

$$\text{if } [do, vo, dc, vc, vrx, vry]^l \cdots [do, vo, dc, vc, vrx, vry]^n [dg, vg] \text{ then } \text{force} / \text{direction} \quad (4-18)$$

学習前のパラメータ設定と学習後の結果は表 4-10 に示される。

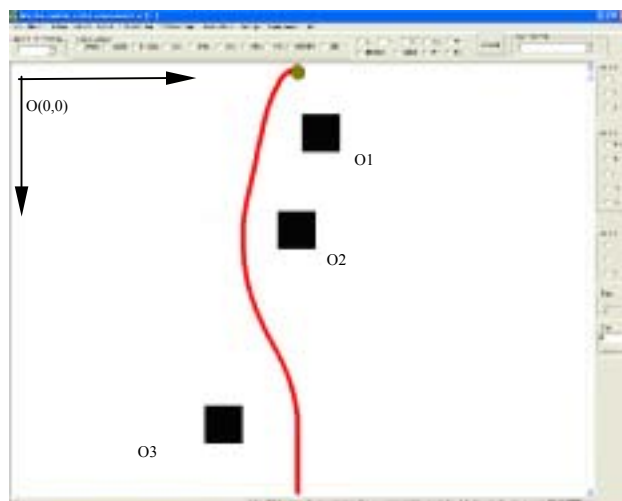


図 4-10 操作者の経路

表 4-10 学習過程のパラメータ

環境	3 個の静的障害物 o1(600, 100) o2(550, 200) o3 (400, 700)
データの数	660
学習誤差	回転角度 15 度 推進力 3N
学習時間	6.9 秒
ルールの数	回転角度 18 推進力 10

学習過程において、660 個のデータに対して、学習過程は 6.9 秒のみをかけた。1 番目の結果ルールに対しては後件部に 3 つの補選値があり、かつ補選値の

生起確率は図 4-11 に示される。学習アルゴリズムの改善部分を通して、最新の補選値ではなく、最大生起確率を持つ補選値-10 を1番目の結果ルールの後件部として選択する。繰り返しの操作がなかったため、初期の状態のみにおいて矛盾ルールが存在したことが分かる。繰り返しの操作が多くなる、または学習誤差が大きく設定されると、補選値の選択が多く発生する。

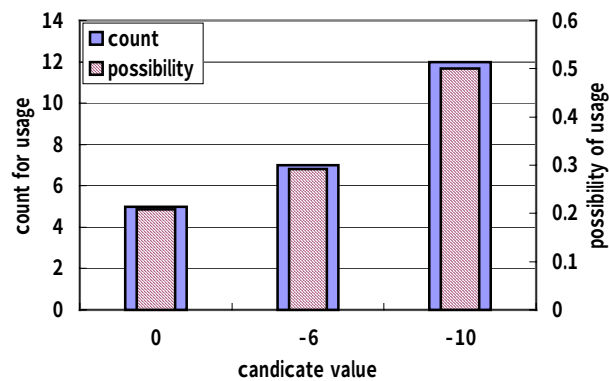


図 4-11 1 番目のルールにおける後件部の補選値の生起分布

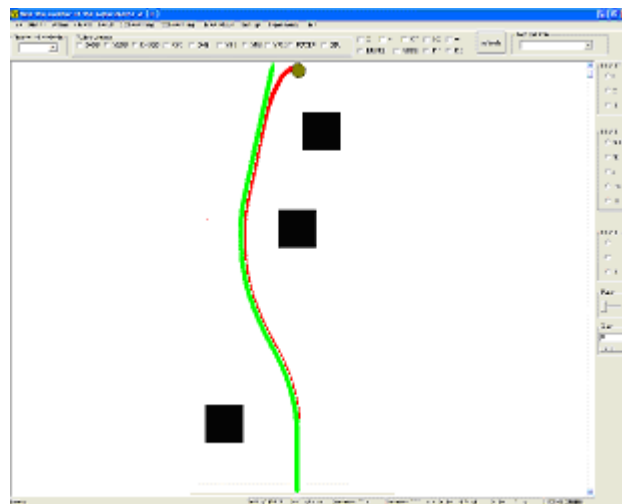


図 4-12 データ学習のルールによる推論効果

学習アルゴリズムにより獲得したルールを第2節の先験的ルールの代わりにすると、従来のプロダクションシステムのとおり、知能エージェントの推論エンジンとして代表的な Mamdani の推論法に基づいてエージェントの行動を決

める、つまりエージェントの推進力と回転角度とを決める。そして、エージェントの運動方程式(2-1)により、エージェントの運動軌道を計算する。

推論に基づいて計算されたエージェントの経路は、図4-12に示すように、結果的には、エージェントはゴールの周りで止まって成功にゴールまで到着しなかったものの、相似な経路は学習アルゴリズムにより知識獲得を行う有効性を示す。更に、正しい知識を獲得する可能性がある上に、適切な推論アルゴリズムと知識の用法に着目してより行動模倣効果の実現を行うことが可能となる。よりよい推論結果を得るために、適切な推論アルゴリズムと正確な知識と知識の用法とも必要からである。例えば、ある時に有効な模倣を達成することができなかった。不適合な知識の用法や推論アルゴリズムは理由の一つであると考えられる。もし本章で提案した知識獲得法を用いて有効なルールを抽出する上に、更に知識の用法や推論アルゴリズムを強調すれば、行動後の結果として残されたデータからもっとよい人間の行動知能の模倣効果を達成することが可能である。

第5節 まとめ

本章では、知識の表現法・獲得法の見方から模倣システムを構築した。まず、知識の表現方法を通して人間の障害物回避戦略を表示する有効性を確認した。そして、知識は問題解決能力をチェックするために、遺伝アルゴリズム GA を用いた知識の評価法を提案した。獲得されたルールに定量的評価基準を提供し、更なる知識獲得に指導する。つぎに、知識の不足と知識獲得の困難という問題点を確認したうえで、多変量で記述しにくい場合でも高速で知識獲得を実現するために、改善した距離型ファジィ推論法の学習アルゴリズムを提案した。距離型ファジィ推論法のオリジナルな学習アルゴリズムに知識（ルール）の生起確率を導入することにより、知識の最適化を行った。この学習アルゴリズムは、多変量、つまり言語で記述しにくい場合でも極めて速い知識獲得を実現することができる。主に、行動知能の模倣を実現するための知識獲得を解決することができる。更に、第2章において人間の障害物回避戦略の限界を分析した先行研究の結果をもとに、1つの静的障害物と処理可能な限界数以内の動的障害物の場合対して、提案した学習アルゴリズムを充分に使用して、豊富な障害物回避知識を抽出することができる。

参考文献

1. Micheal R. Genesereth (著), Nils J. Nilsson (著), 古川 康一(翻訳):人工知能基礎論, オーム社, 1993.
2. E. H. Mamdani : Applications of Fuzzy Algorithms for Control of Simple Dynamic Plant, Proc. IEE, Vol.121, No.12, pp.1585-1588, 1974.
3. T. Takagi and M. Sugeno : Fuzzy Identification of Systems and Its Applications to Modeling and Control, IEEE Transaction on SMC, Vol.15, No.1, pp.116-132, 1985.
4. 市橋秀友, 渡辺俊彦 : 簡略型ファジィ推論を用いたファジィモデルによる学習型制御, 日本ファジィ学会誌, Vol.2, No.3, pp.429-437, 1990.
5. 王碩玉, 土谷武士, 水本雅晴 : 距離型ファジィ推論法, ファジィ学会誌, Vol.1, No.1, pp.61-78, 1999.
6. J. Rasmussen: Information Processing and Human-Machine Interaction-An Approach to Cognitive Engineering, Elsevier, 1986.
7. J. Rasmussen: Skill, Rule and Knowledge; Signals, Signs and Symbols and Other Distinctions in Human Performance Models, IEEE Transaction on Systems, Man, and Cybernetics, Vol.13, No.3, pp.257-266, 1983.
8. L.X. Wang : Fuzzy System Are Universal Approximators, Proceedings of the First IEEE International Conference on Fuzzy Systems, San Diego, U.S.A, pp.1163-1170, 1992.
9. T. Takagi and M. Sugeno : Fuzzy Identification of Systems and Its Application to Modeling and Control, IEEE Transaction System Man and Cybernetic, pp.116-132, 1985.
10. L. Davis : Handbook of Genetic Algorithm, Van Nostrand Reynold, 1989.
11. S.Y. Wang, T. Tsuchiya, and M. Mizumoto : A Learning Algorithm for Distance-type Reasoning Method, Biomedical Soft Computing and Human Science, Vol.6, No.1, pp.61-68, 2000.

第5章

模倣に向ける知識の使用と推論法

第1節 はじめに

模倣システムの構築に対しては、第4章は知識の表現法と獲得法の角度から展開されたものである。豊富な知識量、つまり豊富な障害物回避戦略は模倣システムを構築するための重要な要素であるが、模倣システムを構築するためのもう一つの重要な要素は障害物回避戦略の運用方法としての推論方法である。本章では、距離型ファジィ推論法[1]を使用する。理由としては、ファジィ集合間の距離を利用して推論を行うので、推論用ルールの物理意味が明確であり、推論の方向性が大雑把に予測できる。または、前章に説明したよう、距離型ファジィ推論法のアルゴリズムを利用すれば、学習アルゴリズムとの結合も簡単にでき、かつ獲得した障害物回避戦略に関する知識の最適化も容易に行うことができる。

これまでの研究では、知識の用法については殆ど言及していない。しかし、人間が推論により何かの結論を下す時に、脳にある全ての知識を同時に利用するのではなく、状況に応じて選択的に知識を利用している。知識の選択的利用は、最も関連性のある知識を利用してすばやく結論を下すことが背景にあって、自然に最適化されている自然的な行動だと考えられる。本章では、第4章の学習法により得られた if-then 型宣言的知識が正確なものとして、脳内に行われる知識使用の選択策略を表現する一手法を与え、知識を選択的に利用する推論法を提案する。具体的に、まず事実にも最も関連性のある知識の範囲を意味する知識半径の概念を導入することにより、知識の選択的使用行為を表現する。次に知識半径を考慮した距離ファジィ推論アルゴリズムを提案した。更に、知識半径という概念を静的な知識半径と動的な知識半径に拡張し、もっと柔軟性をもつように知識を利用するに検討した。

第2節 知識使用を融合した推論法

推論を行う際に、使用される知識が多ければ多いほど、時間が掛かるだけではなく、主要なつまり最も関連性のある知識の役割が返って薄れる。事実が与えられた人間の脳では、少ない労力で素早く結論を下すために、脳中にあるすべての知識を同等に利用するのではなく、事実と最も関係性のある知識だけを使用している。この機能は、恐らく脳の長い進化の歴史においては、省エネルギーと高速性を追求し、推論の効率化に寄与する目的で最適化された必然的な結果であろう。本節では脳における知識の選択的使用策略を表現する一手法を与える。

第1項 知識半径の概念

知識の選択的使用策略を表すには、事実と知識との関係を定量化する必要がある。つまり、事実と知識との関係性の大きさを数値により表現することができれば、諸知識が事実に対しての位相関係が分かり、「事実と最も関係性のある知識を使用する」脳の知識使用行為をモデル化することが可能となる。ここでは、推論に用いる知識はif-then 型宣言的ルールであるとして、与えられる事実と知識の前件部との距離値により、事実と知識との関係性の大きさを表す。理由としては、距離情報を利用すれば、数学的に厳密に議論しやすいだけではなく、物理的な概念も明確である。たとえば、前件部と事実との距離値が大きければ大きいほど、それらのルールは事実との関連性が少ないことを意味する。ファジィ集合間の距離計算としては、文献[2]では連続値のファジィ集合間の距離計算法を、文献[3]では離散値のファジィ集合間の距離計算法を与えているが、距離の公理（非負性、対称性、三角性）を満たしていれば、位相関係が数学的に保証されるので、どれを使っても構わない。

つづいて、知識半径の概念に基づいて、知識を選択的に利用する脳の行為を表現する。知識半径を事実と最も近いルールの数と定義する。事実とルールの近さは前述した事実とルールの前件部との距離値によって計算され、知識半径

は整数で表される。例えば、知識半径は q であれば、今現在の事実に最も近いルール 수는 q 個であることを意味する。したがって、知識半径 q は範囲($2 \leq q \leq \text{ルールの数}$)内の整数である。知識半径を用いれば、「事実と最も関係性のある知識を使用する」という知識の選択的利用行為を表現することが可能になる。具体的には、推論あるリズムに、知識半径 q を導入すれば、事実と最も関係性のある q 個のルールを使用して、推論を行うことになる。

原理上では、直接に距離情報を用いて、知識の選択的利用行為を表現することができるが、知識半径を導入したメリットとしては、範囲が明確であり、何らかの最適な基準で最適な知識半径を搜索する場合、少ない計算量で最適な知識半径が探せる。逆に距離情報は直接に利用するとすれば、距離値の範囲の設定が非常に難しく、かつ連続範囲内では距離の最適値を搜索するには膨大な計算量が要る。したがって、整数の知識半径が与えられた離散的ルール形式に適すると考える。また、知識の使用法については幾つかの手法が提案されている。例えば、文献[4]と文献[5]では、それぞれルールの活性化とルールの重みの設定の立場から知識使用法の重要性を示している。本章では、事実とルールとの距離情報に着目し、知識半径の概念を導入することにより、脳中で行われている知識の選択的使用行為をモデル化する。

第2項 知識半径を導入した距離型ファジィ推論法

フィジィ論理は人間の曖昧な言語表現に基づく論理的思考過程をモデル化する特徴を考慮して、人間の選択的な知識利用行為を真似るために、知識使用の運用策略に対して、いわゆる知識半径という概念を採用して、推論による問題解決における知識使用の選択策略を着目する。知識半径が最初に文献[6]に提案されたが、知識半径の性質がまだ明かされないし、知識半径の応用も難しくなる。具体的に言えば、知識半径とは、ファジィ集合間の距離空間における事実と各ルールの前件部との距離値を計算する上に、事実を中心、各ルールの前件部を離散な点、距離値を半径の度量長さとして円が囲む離散な点の数である。事実の変化とともに、事実と各ルールの前件部との距離値が変わり、選択されたルールも変わるかもしれない。従って、知識半径を用いて、ファジィ推論過程において推論ルールが制限的に使用される。結果的ルール表現として、知識半径は、事実と各ルールの前件部との距離値から小さい順に指定的に選ばれるルールの数である。ここでは、基本の距離型ファジィ推論法(DTFR)を元にして、知識半径を導入した距離型ファジィ推論法を提案する。

STEP1：事実中のファジィ変数 A^j とルール中のファジィ変数 A^{ij} との距離 $d_{ij}(A^j, A^{ij})$ を計算する。

STEP2：事実と i 番目ルールとの距離を計算する。

$$d_i = \sum_{j=1}^m d_{ij}(A^j, A^{ij}) \quad (5-1)$$

STEP3：事実と各ルールの前件部との距離値により d_i を単調増加の順に並べ替えて知識半径の変化範囲を決める。そして、指定的知識半径の値 q により d_i を相応な距離閾値 δ で変換する。

$$d'_i := d_i \cdot \varepsilon^{\text{sgn}(\delta - d_i) - 1}, \quad \varepsilon \rightarrow 0 \quad (5-2)$$

ただし、 $\text{sgn}(\bullet)$ は符号関数であり、 $\text{sgn}(x) = \begin{cases} 1, & x \geq 0 \\ -1, & x < 0 \end{cases}$

STEP4: $\forall \alpha \in [0,1]$ に対して, 次のように推論結果 B の α -レベル集合 B_α を求める.

$$\begin{aligned}
 B_\alpha &= [\inf(B_\alpha), \sup(B_\alpha)] \tag{5-3} \\
 \inf(B_\alpha) &= \frac{\sum_{i=1}^n \left[\inf(B_\alpha^i) \prod_{j=1, j \neq i}^n d'_j \right]}{\sum_{i=1}^n \prod_{j=1, j \neq i}^n d'_j} \\
 &= \frac{\sum_{i=1}^n \left[\inf(B_\alpha^i) \prod_{j=1, j \neq i}^n \left\{ d_j \cdot \varepsilon^{\operatorname{sgn}(\delta - d_j) - 1} \right\} \right]}{\sum_{i=1}^n \prod_{j=1, j \neq i}^n \left\{ d_j \cdot \varepsilon^{\operatorname{sgn}(\delta - d_j) - 1} \right\}} \\
 \sup(B_\alpha) &= \frac{\sum_{i=1}^n \left[\sup(B_\alpha^i) \prod_{j=1, j \neq i}^n d'_j \right]}{\sum_{i=1}^n \prod_{j=1, j \neq i}^n d'_j} \\
 &= \frac{\sum_{i=1}^n \left[\sup(B_\alpha^i) \prod_{j=1, j \neq i}^n \left\{ d_j \cdot \varepsilon^{\operatorname{sgn}(\delta - d_j) - 1} \right\} \right]}{\sum_{i=1}^n \prod_{j=1, j \neq i}^n \left\{ d_j \cdot \varepsilon^{\operatorname{sgn}(\delta - d_j) - 1} \right\}}
 \end{aligned}$$

STEP5: 合成原理により推論結果 B を求める.

$$B = \bigcup_{\alpha} B_\alpha \tag{5-4}$$

知識半径の変化範囲について, 以下に詳しく説明される. STEP3 で計算された事実と各ルールの前件部との距離値, つまり集合 $D = \{d_1, \dots, d_n\}$ の上に, 同じ距離値を除いて単調増加の順に並べ替えた結果は $D' = \{d'_1, \dots, d'_m \mid d'_j < d'_{j+1}, j = 1, 2, \dots, m-1\}$ であり, 同じ距離値を持った要素の数は $L = \{l_j \mid l_j \geq 1, j = 1, 2, \dots, m\}$, $\sum_{j=1}^m l_j = n$. 従って, 知識半径の変化範囲

$k = \{k_1, \dots, k_i, \dots, k_R\}$ と相応な距離閾値 $\delta = \{\delta_1, \dots, \delta_i, \dots, \delta_R\}$ については, $l_1 = 1$ の場合に, $R = m - 1$, $k_i = \sum_{p=1}^{i+1} l_p$, $\delta_i = d'_{i+1}$, $i = 1, \dots, R$. またはそうなく, $l_1 > 1$ の場合に, $R = m$, $k_i = \sum_{p=1}^i l_p$, $\delta_i = d'_i$. 結果的に, 同じ距離値が存在するので, 知識半径の値が連続的に存在しないかもしれない. たとえば, 図 5-1 の示すように, ルールベースに五個のルールがある. 推論時に, 事実と各ルールの前件部との距離値は d_1, d_2, d_3, d_4, d_5 である. 距離値の単調増加の順で並べ替えた後で, $d_1 = d_2 < d_3 < d_4 = d_5$ の場合に, 知識半径の変化範囲は $\{2 \ 3 \ 5\}$ である. $d_1 < d_2 = d_3 < d_4 = d_5$ の場合に, 知識半径の変化範囲は $\{3 \ 5\}$ である.

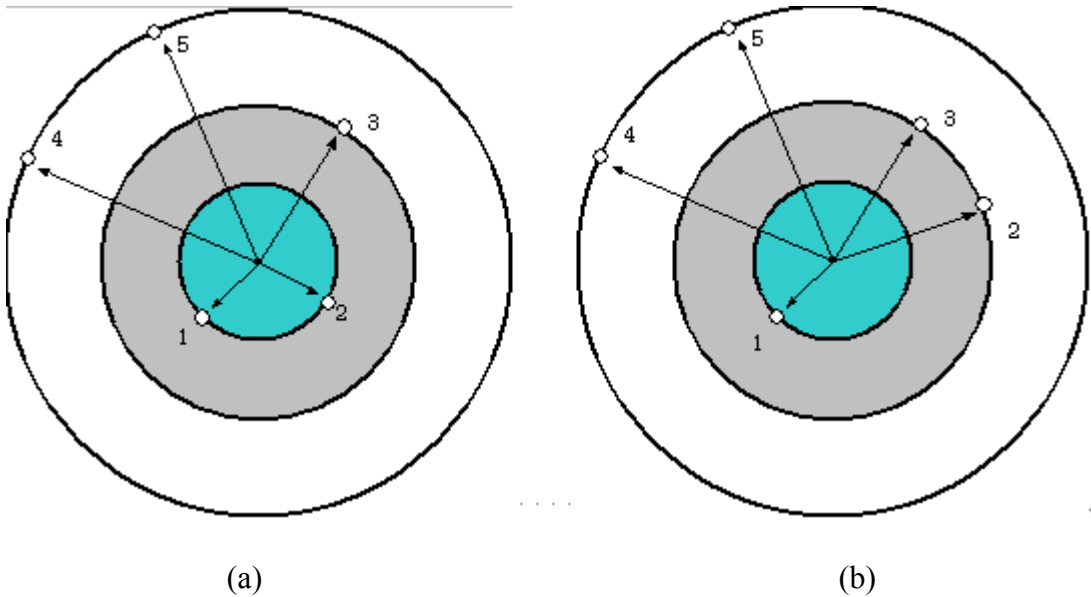


図 5-1 (a) $d_1 = d_2 < d_3 < d_4 = d_5$ (b) $d_1 < d_2 = d_3 < d_4 = d_5$

上述の推論アルゴリズムの明白な説明として, 図 5-2 には知識半径を導入した DTFR 推論法のブロック線図を示す. 知識半径 q に対して, 事実と前件部との距離値が最も近いルール集合 S のみを利用して高速的に推論を行う. ただし, $|S|$ は知識半径により確定される推論用ルールの数である.

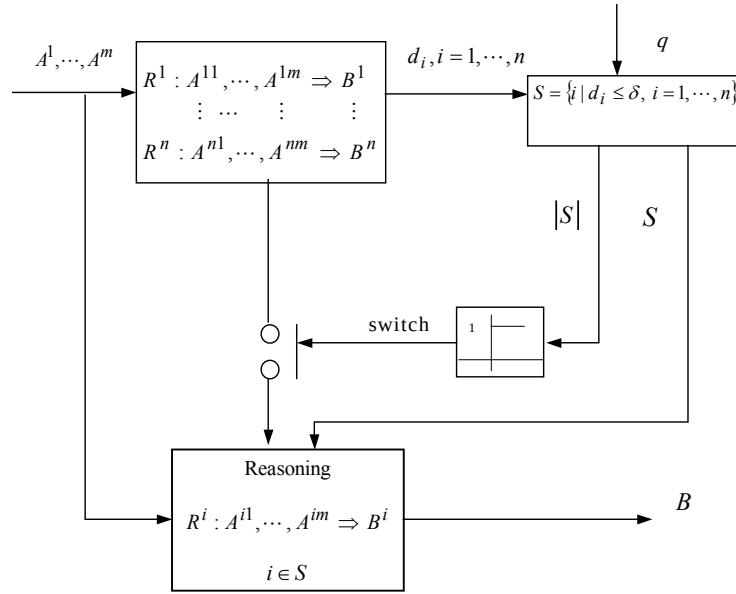


図 5-2 知識半径を導入した DTFR 推論法のブロック線図

数値型データ学習に適合するために、シングルトン型ファジィ集合が知識表示として採用される．シングルトン型の知識表示による推論効果がメンバーシップ関数の影響に受けないので，知識利用の有効性を考察するに適合する．従って，知識半径を導入した距離型ファジィ推論法の特例として，前件部と後件部を実数に簡略化した簡略型推論法について述べる．

具体的に，知識半径 q を指定すれば，知識半径を導入した簡略型距離型ファジィ推論法は以下の 4 つのステップから構成される．

STEP1：式(5-5)により事実 A^j と i 番目のルールにおける j 番目の前件部 A^{ij} との距離値 d_{ij} を求める．ただし， $i = 1, 2, \dots, n$ ， $j = 1, 2, \dots, m$ ．

$$d_{ij} = d_{ij}(A^{ij}, A^j) = |a^{ij} - a^j| \quad (5-5)$$

式(5-5)は 2 つのファジィ集合間の距離計算公式をシングルトン A^{ij} とシングルトン A^j との距離計算に適用する場合に，簡略した距離の計算式である．図 5-3 の示すように， A^{ij} と A^j と B^i はシングルトン型ファジィ集合であり，パラメータ a^{ij} と a^j と b^i で表す．

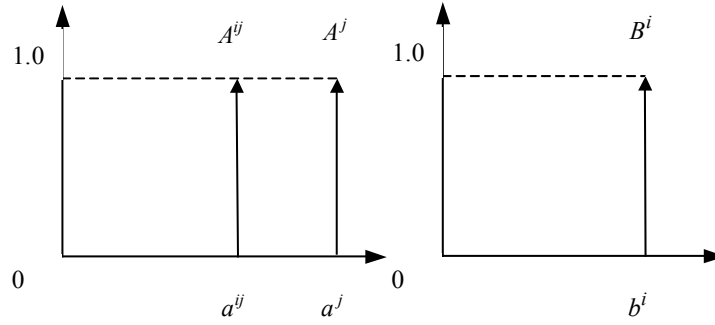


図 5-3 シングルトンファジィ集合の例

STEP2: 式(5-6)より i 番目のルールの前件部全体と事実との距離値 d_i を求める.

ただし, w_j と ρ_j は $w_j > 0$, $\rho_j \geq 1$ を満たすパラメータである

$$d_i = \sum_{j=1}^m w_j (d_{ij})^{\rho_j} \quad (5-6)$$

STEP3: 事実と各ルールの前件部との距離値により d_i を単調増加の順に並べ替えてから, 式(5-7)より d_i を指定的知識半径の値 q の応じる距離閾値 δ で変換する.

$$d'_i := d_i \cdot \varepsilon^{\text{sgn}(\delta - d_i) - 1}, \quad \varepsilon \rightarrow 0 \quad (5-7)$$

ただし, $\text{sgn}(\bullet)$ は符号関数であり, $\text{sgn}(x) = \begin{cases} 1, & x \geq 0 \\ -1, & x < 0 \end{cases}$

STEP4: 式(5-8)により出力 y_0 を求める. ただし, b^i はシングルトンの後件部 B^i のパラメータである.

$$y_0 = \frac{\sum_{i=1}^n \left[b^i \prod_{j=1, j \neq i}^n d'_j \right]}{\sum_{i=1}^n \prod_{j=1, j \neq i}^n d'_j} = \frac{\sum_{i=1}^n \left[b^i \prod_{j=1, j \neq i}^n \left\{ d_j \cdot \varepsilon^{\text{sgn}(\delta - d_j) - 1} \right\} \right]}{\sum_{i=1}^n \prod_{j=1, j \neq i}^n \left\{ d_j \cdot \varepsilon^{\text{sgn}(\delta - d_j) - 1} \right\}} \quad (5-8)$$

相対的に, 元のアルゴリズムによる出力は式(5-9)に示される.

$$y_0 = \frac{\sum_{i=1}^n \left[b^i \prod_{j=1, j \neq i}^n d_j \right]}{\sum_{i=1}^n \prod_{j=1, j \neq i}^n d_j} \quad (5-9)$$

第3項 知識半径の性質

知識半径を導入した距離型ファジィ推論法は、もとの距離型ファジィ推論法に比べて、もとの特徴を持つだけでなく、知識半径を導入して推論過程に知識の選択的利用を強めることで新しい特徴を引き出す。従って、知識半径に関する一般性質を纏める必要がある。

性質1：知識半径 q の範囲: $2 \leq q \leq n$ 。

q は知識半径を、 n はルールの個数を表す。

性質2：知識半径 q の参与で、部分ルールの推論用の重み w_i が増大されるが、他のルールの推論用の重み w_i がゼロになる。従って、全部のルールの重要性が知識半径の参与につれて変わる。

証明：

前件部と後件部を実数に簡略化した簡略型推論を例として、元のアルゴリズムによる出力は式(5-9)から式(5-10)に変化される。ただし、 w_i はルールの推論用の重みを、 m は前件部の数を、 n はルールの数を表す。

$$y_0 = \frac{\sum_{i=1}^n \left[b^i \prod_{j=1, j \neq i}^n d_j \right]}{\sum_{i=1}^n \prod_{j=1, j \neq i}^n d_j} = \sum_{i=1}^n \left[\frac{b^i \prod_{j=1, j \neq i}^n d_j}{\sum_{i=1}^n \prod_{j=1, j \neq i}^n d_j} \right] = \sum_{i=1}^n [b^i w_i] \quad (5-10)$$

$$w_i = \frac{\prod_{j=1, j \neq i}^n d_j}{\sum_{i=1}^n \prod_{j=1, j \neq i}^n d_j} \quad (5-11)$$

同様に、提案アルゴリズムによる出力は式(5-8)から式(5-12)に変化される。

$$y'_0 = \frac{\sum_{i=1}^n \left[b^i \prod_{j=1, j \neq i}^n d'_j \right]}{\sum_{i=1}^n \prod_{j=1, j \neq i}^n d'_j} = \sum_{i=1}^n \left[\frac{b^i \prod_{j=1, j \neq i}^n d'_j}{\sum_{i=1}^n \prod_{j=1, j \neq i}^n d'_j} \right] = \sum_{i=1}^n [b^i w'_i] \quad (5-12)$$

$$w'_i = \frac{\prod_{j=1, j \neq i}^n d'_j}{\sum_{i=1}^n \prod_{j=1, j \neq i}^n d'_j} \quad (5-13)$$

1) 事実 $A^1, \dots, A^m$ を推論方法に入力する際に, $\exists k \in \{1, 2, \dots, n\}$ 等式

$$d_k = \sum_{j=1}^m d(A^j, A^{kj}) = 0 \text{ という条件を満たす場合に, 知識半径により変換された}$$

距離値は必ず $d'_k := d_k \cdot \varepsilon^{\text{sgn}(\delta - d_i) - 1} = 0 \cdot \varepsilon^0 = 0$ になる. このルールの推論用の重みは $w'_k = w_k = 1$ なりながら, ほかのルールの推論用の重みは $w'_i = w_i = 0, i = 1, \dots, n, i \neq k$ になる. 結果的に, k 番目のルールは必ず知識半径に単に選択されると, $y'_0 = y_0$, 分離規則を満たす. 従って, 元のアルゴリズムに比べて, 全部のルールの推論用の重みが変わらない.

2) 分離規則を満たさない, 即ち, $\forall i \in \{1, 2, \dots, n\}$ に対して d_i は事実とルール R^i の前件部との距離として零にならない場合に,

元のアルゴリズムによる推論用の重みは式(5-14)に説明される.

$$\begin{aligned} & \because \forall i \in \{1, 2, \dots, n\}, d_i > 0 \\ \therefore w_i &= \frac{\prod_{j=1, j \neq i}^n d_j}{\sum_{i=1}^n \prod_{j=1, j \neq i}^n d_j} = \frac{1/d_i}{\sum_{j=1}^n (1/d_j)} \end{aligned} \quad (5-14)$$

相対的に, 知識半径に影響される推論用の重みは式(5-15)と(5-16)に説明される.

$$w'_i = \frac{\prod_{j=1, j \neq i}^n d'_j}{\sum_{i=1}^n \prod_{j=1, j \neq i}^n d'_j} = \frac{1/d'_i}{\sum_{j=1}^n (1/d'_j)} \quad (5-15)$$

$S = \{i \mid d_i \leq \delta, i = 1, \dots, n\}$ を満たす $d'_i = d_i$ と $\{i \mid d_i > \delta, i = 1, \dots, n\}$ を満たさない $d'_i \rightarrow \infty$ が成立する. ただし, δ は知識半径の値 q の応じる距離閾値である. 従って, 知識半径に限定されるルールの重みが以下に示される.

$$w'_i = \frac{1/d'_i}{\sum_{j \in S} (1/d'_j)} = \frac{1/d_i}{\sum_{j \in S} (1/d_j)} \geq w_i \quad (5-16)$$

またはそうなく、 $d'_i \rightarrow \infty$ のために、知識半径に限定されないルールの重みは $w'_i \rightarrow 0$ になる。

性質 3：知識半径 q の調整は現有知識を最大程度で運用することであり、結果的効果は少なくとも元の効果と一致すると保証する。

第4項 知識半径の確定法

第2, 3項で知識半径を導入した距離型ファジィ推論法及び知識半径の性質を述べたが, 知識半径を導入した距離型ファジィ推論法を問題解決に適用する際に, 知識半径 q の確定が新しい問題として現れる. 知識半径の選び方は問題解決のための効果評価基準に関係があり, かつ2つの意味を含む. まず, 異なる問題に対して, 効果評価基準の相違のために, 適合な効果を満たす知識半径が異なるかもしれない. 次に, 同じ問題に対しても, 効果評価基準の相違のために, 適合な効果を満たす知識半径も異なるかもしれない. 具体的に, ある問題に対して, 全局で評価基準を達成する知識半径の値は推論過程において統一する, つまり変わらない. 逆に, 局所で評価基準を達成する知識半径の値は推論過程において動的に変わるかもしれない. 手法となる知識半径の選び方を以下に与える.

全局で評価基準を達成する場合に, 推論過程において知識半径が一致するので, タスクによる評価関数 $J(\bullet)$ で知識半径の有効性を考察する. あるタスクに対して, $x_1 = A^1(t)x_2 = A^2(t)\cdots x_m = A^m(t)$ は時刻 t の事実を, $B_k(t)$ は時刻 t において知識半径の値 k を用いた推論結果を, T_k は知識半径の値 k を用いた推論ステップの数を表す. ただし, $t \geq 1$, $k = 2, \dots, n$, n はルールの数であり, m は前件部の数である. 具体的に, 知識半径の選び方は以下の4つのステップから構成される.

STEP1: 知識半径の範囲 $[2, \dots, n]$ に, $k = 2$ を初値として $B_2(1), \dots, B_2(T_2)$ を計算する.

STEP2: 評価値として $J_k = J(B_k(1), \dots, B_k(T_k))$ を計算する.

STEP3: 次の知識半径 $k+1$ を取ってSTEP2に戻る. またはそうなく, STEP4に進む.

STEP4: 小さい評価値はよい効果に応じると仮定すると, $\min(J_2, \dots, J_q, \dots, J_n)$ を持つ q を適合な知識半径として選択する.

局所で評価基準を達成する場合に, 推論過程において知識半径が時刻あたり

の評価により選択される．時刻あたりの評価関数 $J(\bullet)$ で知識半径の有効性を考察する． $x_1 = A^1(t)x_2 = A^2(t)\cdots x_m = A^m(t)$ は時刻 t の事実を， $B_k(t)$ は時刻 t において知識半径の値 k を用いた推論結果を表す．ただし， $t \geq 1$ ， $k = 2, \dots, n$ ， n はルールの数であり， m は前件部の数である．具体的に，知識半径の選び方は以下の5つのステップから構成される．

STEP1：知識半径の範囲を $[2, \dots, n]$ に， $t = 1$ の時に， $B_2(1), \dots, B_n(1)$ を計算する．

STEP2：評価値としての $J(B_2(t)), \dots, J(B_n(t))$ をそれぞれ計算する．

STEP3：小さい評価値はよい効果に応じると仮定すると，時刻 t において， $\min(J(B_2(t)), \dots, J(B_q(t)), \dots, J(B_n(t)))$ を持つ q を適合な知識半径として推論する．

STEP4：必要なら，計算時間を短くするため， $\{k \mid J(B_k(t)) < \varepsilon\}$ を満たす集合を新しい知識半径の範囲とする．そうならなければ，STEP5 に進む．ただし， ε は指定の閾値である．

STEP5：時刻 $t+1$ において，STEP2 に戻る．

通常，問題解決の効果 E を評価するために，閾値 δ を指定する必要がある．また，便利な分析のために，知識半径に関するいくつかの概念が定義される．有効な知識半径とは，知識半径を導入した距離型ファジィ推論法に基づく問題解決の効果 ($E \leq \delta$) を達成した知識半径である．もし有効な知識半径が存在すれば，有効な知識半径の個数 ≥ 1 ．更に，最適な知識半径とは，問題解決の最適効果を達成した ($\text{minimize } E \leq \delta$) かつ最も小さい値を持つ有効な知識半径である．もし最適な知識半径が存在すれば，最適な知識半径の個数は1である．一方，もし $\text{minimize } E > \delta$ であれば，最適な知識半径が存在しないが， $\text{minimize } E > \delta$ かつ最も小さい値を持つ知識半径はトレードオフ知識半径であると思われる．トレードオフ知識半径の個数は1である．

第 3 節 知識半径を用いた障害物回避行動の模倣

人間の選択的な知識利用行為を真似るために、いわゆる知識半径という概念を採用して、知識半径を導入した距離型ファジィ推論法を提案した。しかし、本推論方法で再現した時に得られる効果が如何なるものであるかは未知のものである。従って、推論方法の有効性を検討するため、エージェントに知識を付加する行動再現システムにおいて、行動後の結果として残されたデータから障害物回避行動を再現することに着手する。

同時に、1 つの実現目標が仮定される：任意の操作者が操作道具で操作したら、行動後の結果としてデータを元にして知識表示が生成する。そして、それらの知識と知識半径を導入した距離型ファジィ推論法に基づいて、操作者の経路に似る経路が生成される、つまり操作者の障害物回避戦略が有効に模倣される。

仮定目標を実現するために、従来の方法のように、メンバーシップ関数の調整やルールの増加削除により模倣効果を改善することができる。その最も代表的な手法は様々な知識獲得方法[7-10]と相関の最適化調整方法など[11,12]である。ここでは、獲得された知識の正確性を信じる上に、知識の利用に重視する知識半径を導入した距離型ファジィ推論法を採用して障害物回避行動模倣を行う。具体的模倣手法として、以下の 4 つのステップから構成される。

STEP1：操作者は運転シミュレータの特定のタスク環境において、障害物を回避しながら、できるだけゴールに到着する。

STEP2：タスクが終わったら、ファジィ推論学習則に対し、第 4 章の学習アルゴリズムを用いて行動後の結果として残されたデータから推論用のファジィルールを操作者の障害物回避戦略として獲得する。

STEP3：エージェントに知識を付加する行動再現システムにおいて、STEP2 で獲得されたルールをエージェントの知識、知識半径を導入した距離型ファジィ推論法をエージェントの推論モデルとして、STEP1 と同じタスク環境において、エージェントを行動させる。

STEP4：行動後の結果として残されたエージェントのデータをもとにして，効果評価基準により，最適な知識半径や有効な知識半径を確定してその知識半径に基づく推論効果を障害物回避行動の模倣効果とする．

上述の模倣手法を適用する過程に，問題解決のための効果評価基準と知識半径の確定は以下に詳しく紹介する．

第1項 経路面積誤差の定式化

知識半径を導入した距離型ファジィ推論法を用いて，エージェントの走行経路が生成される．推論による模倣効果を評価するために，もし推論によるエージェントの走行経路と人によるエージェントの走行経路との囲む面積誤差を利用すれば，経路間の相似程度を定式化することができる．式(5-17)は経路間に経路面積誤差を示す．

$$S = \int s(k)dk \quad (5-17)$$

$$\text{制約条件} \begin{cases} Y(0) = R(0) \\ |Y(L) - R(OL)| \geq 0 \end{cases}$$

ただし，

$R(k)$ ：人によるエージェントの走行経路曲線

$Y(k)$ ：推論によるエージェントの走行経路曲線

k ：軌道位置

S ： $R(k)$ と $Y(k)$ という2つの曲線の囲む面積誤差

$s(k)$ ：位置 k において， $R(k)$ と $Y(k)$ という2つの曲線の囲む標本面積

OL ： $R(k)$ の軌道長さ

L ： $Y(k)$ の軌道長さ

もし人によるエージェントの走行経路 $R(k)$ と推論によるエージェントの走行経路 $Y(k)$ という2つの曲線の囲む面積誤差 S がゼロであれば，模倣効果が最適である． S が小さければ小さいほど，模倣効果が良いと思われる．

面積 S の積分計算については，有限要素法の汎用性を結び付けて， S を多数

の三角形要素に分割して面積 S を計算する．このように，余計な面積が計算されないし，経路間の軌道の長さ差異が制限されない．具体的に，三角形要素の分割については，2 つの経路から，それぞれ 1 つの点を取り，最小距離方法により三番目の点を 2 つの経路における隣接点から選択して標本三角形を作る．結果的に， S が三角形分割に形成され，式(5-18)のように表される．

$$S = \int s(k) dk = \sum_{i=1}^N s_{\Delta}(k) \quad (5-18)$$

ただし，

$s_{\Delta}(k)$: 分割された三角形の面積．

N : 分割三角形の個数．

その中に，図 5-4 の斜線部で表される 1 個の標本 $s_{\Delta}(k)$ は次式のように表される．

$$s_{\Delta}(k) = \begin{cases} \frac{1}{2} \|e(k) \times \Delta R(k+1)\|, & e(k) = Y(k) - R(k), \\ & \|Y(k+1) - R(k)\| > \|R(k+1) - Y(k)\| \\ \frac{1}{2} \|e(k) \times \Delta Y(k+1)\|, & e(k) = R(k) - Y(k), \\ otherwise \end{cases} \quad (5-19)$$

Δ : 一階の後退差分演算子

$\|\cdot\|$: ベクトルのユークリッドノルム

$\times$: 外積

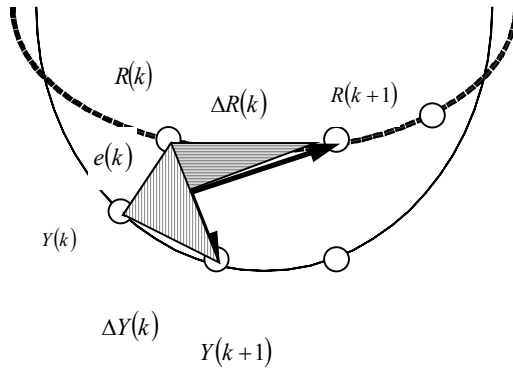


図 5-4 最小距離による三角形の選択

同時に、経路境界の処理、つまり軌道の長さ差異については、人によるエージェントの走行経路と推論によるエージェントの走行経路の相対関係を考慮すべきである。人によるエージェントの走行経路を元にして、推論によるエージェントの走行経路は3つの種類の相対関係を持つ：人によるエージェントの走行経路に近い、人によるエージェントの走行経路より短すぎる、人によるエージェントの走行経路より長すぎる。経路情報を損失しないように、以下の2つの原則を守れば、経路面積誤差を利用して模倣程度をうまく評価することが可能となる。

- 1) 経路末端のつながりで閉領域の面積を計算する。
- 2) 人によるエージェントの走行経路に比べて、より長い推論によるエージェントの走行経路は良い模倣程度に対して優位に占める。

結果的に、推論によるエージェントの走行経路の長さを S の計算に付けると、 S が式(5-20)に代わる。

$$S = \sum_{i=1}^N s_{\Delta}(k) / \eta L \quad (5-20)$$

ただし、

η : 指定された係数, 1 とする。

L : 推論によるエージェントの走行経路の長さ

N : L により分割された三角形の個数

式(5-20)を用いた経路面積誤差の計算については、一方、図5-5(a)のように同じ末端を持つ経路面積誤差の比較を解決できる。他方、図5-5(b,c,d)のように元の開領域であれば、経路末端のつながりで閉領域を形成する。経路が長くなるとともに、積分面積が増えることがわかる。式(5-20)により、図5-5(b)のように開領域且つ接近の長さを持つ経路にかかわらず、図5-5(c,d)のような開領域且つ不釣合いの長さを持つ経路に対しては経路面積誤差も合理的に計算される。

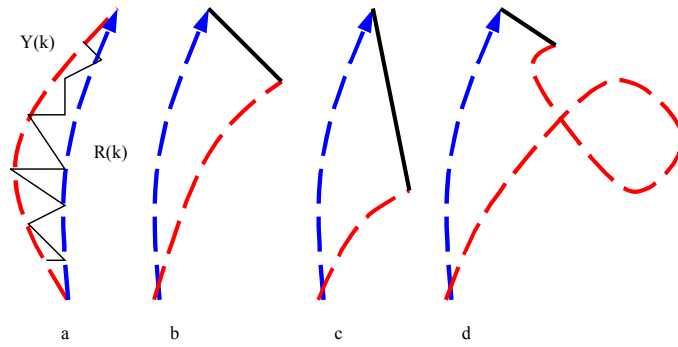


図 5-5 閉鎖的な経路(a),人によるエージェントの走行経路と比べて,近似な長さを持つ非閉鎖的な経路(b)や短い(c)や長い(d),矢印付き点線は人によるエージェントの走行経路を,矢印なし点線は推論によるエージェントの走行経路を,直線は閉鎖のためのつながりを表す.)

第2項 知識半径の確定

現有ルールにおいて知識半径の値を調整することにより，人によるエージェントの走行経路と推論によるエージェントの走行経路の最適模倣効果，すなわち，2つの経路の囲む最小面積誤差を求めるという知識半径の最適化問題に対しては Lagrange の未定乗数法で非線形固有値問題に定式化できるが，非唯一の局所的最適値が存在するという問題を解決し，及び Lagrange 乗数を避けるために，我々は知識半径の値 q という制約条件を式第2章のエージェントの運動方程式に導入したら，推論によるエージェントの走行経路 $Y(k)$ は次式のように知識半径を考慮して軌道位置を計算する．

$$m\ddot{Y} + D\dot{Y} = F(q) \quad (5-21)$$

ただし，

$F(q)$: 知識半径の値 q を考慮した距離型ファジィ推論によるエージェントが受ける合力

他の符号: 第2章と同じ意味

第2節の知識半径の確定法を参考して，行動模倣のための効果評価基準を結びつけると，有効な知識半径を探索する連続的な繰り返し計算方案を設計する．具体的に，この方案は次の4つのステップから構成される．

STEP1: 知識半径の値を1つ取り，人によるエージェントの走行経路と同じ初期条件を設定する．

STEP2: 知識半径を用いた推論過程を完成したら，結果的経路と人によるエージェントの走行経路との経路面積誤差を計算する．

STEP3: 次の知識半径の値を取って STEP1 に戻る．またはそうなく，STEP4 に進む．

STEP4: 最小経路面積誤差を持つ知識半径の値 q を選択する．

第3項 シミュレーションによる検討

1) 学習過程

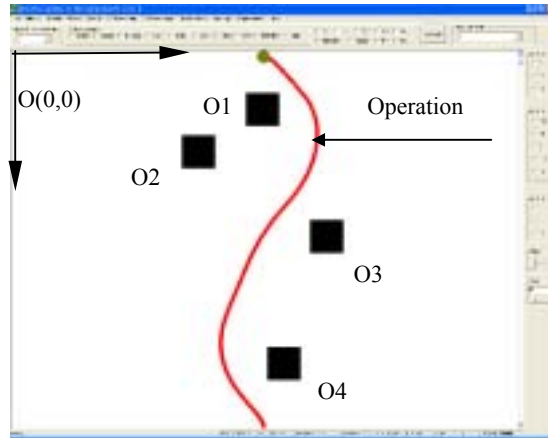


図 5-6 操作者によるエージェントの走行経路

図 5-6 のような実験場面において，左上角の座標は(0,0)であり，右向きと下向きは正方向であると仮定される．まず，操作者がエージェントを制御し，障害物を回避しながら，ゴールへの到着に成功した．そして，ファジィ推論ルールに対し，第 4 章の学習アルゴリズムを採用して行動データからルールを獲得した．具体的に，学習前のパラメータ設定と学習後の結果は表 5-1 に示される．

表 5-1 学習過程のパラメータ

環境	4 個の静的障害物 O1(550,100) O2(400,200) O3(700,400) O4(600,700)
データの数	371
学習誤差	回転角度 15 度 推進力 3N
学習時間	6.9 秒
ルールの数	回転角度 102 推進力 15

結果的ルール形式は以下に示される．ただし，符号は 2.2 項と同じ意味を持つ．

$$\text{if } [do, vo, dc, vc, vrx, vry]^1 \dots [do, vo, dc, vc, vrx, vry]^{n_o} [dg, vg] \text{ then } force / direction \quad (5-22)$$

2) 有効な知識半径の確定

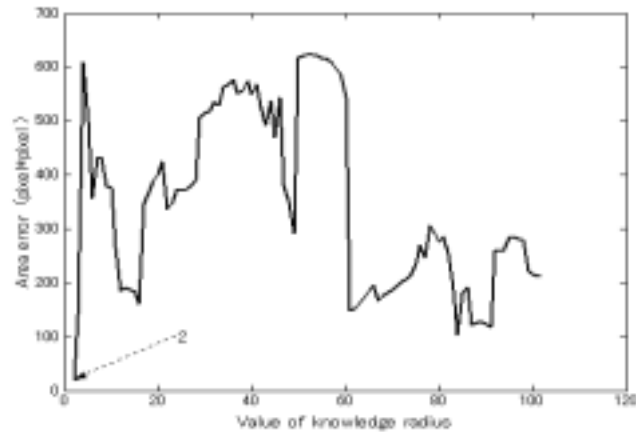


図 5-7 目的関数

表 5-2 有効な知識半径(探索時間:98 秒)

番号	有効な知識半径 q (判定閾値 $\varepsilon = 120$)	面積誤差	結果
1	2(最適)	16.4	成功
2	84	102.7	失敗
3	91	117.8	失敗
他の成功場合($\geq \varepsilon$)			
4	22	337.2	成功
5	23	349.5	成功

安定の推進力変化を守る上に、回転角度ルールに対して知識半径の調整を行った。第2項の繰り返し計算方案により計算した経路面積誤差は問題の目的関数として図5-7に示される。その目的関数は1つだけではない波の谷を持つ非凸曲線であると分かる。更に、判定閾値を考慮したら、表5-2のような3つの有効な知識半径と、最小面積誤差を満たすかつ成功経路を生成する最適な知識半径2を得た。最適な知識半径を導入した距離型ファジィ推論法に基づく効果は図5-8の示すようなエージェントの走行経路が生成された。その走行経路は操作者によるエージェントの走行経路を有効に模倣し、かつ知識半径を用いな

かった場合より良い経路を生成したことがわかる．最適な知識半径を通して，現有知識を最大程度で運用することを示す．

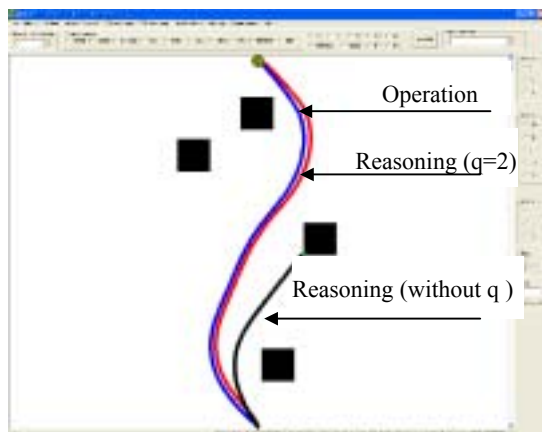


図 5-8 人によるエージェントの走行経路と推論によるエージェントの走行経路

繰り返し計算方案により経路面積誤差を計算したと同時に知識半径の値の応じる推論過程における推論用の距離値を全般に記録した．ただし，回転角度ルールに対して知識半径の調整を行ったので，知識半径の値の変化範囲は 2～102 である．距離型ファジィ推論法については，距離値 d_i は事実とルールの前件部との相違を定量に表す，つまり $1/d_i$ は事実と前件部との相関性を描く．知識半径を導入した距離型ファジィ推論法でも，知識半径を推論法に付けると，知識半径における事実と前件部との距離値 d'_i は，やっぱり事実と選択したルールの前件部との相違を定量に表す．そこで，知識半径における事実と前件部との距離値を利用して，知識半径と知識相関性の関係を解析することが可能となる．

知識半径における事実とルールの前件部との上限距離 $\max_{i \in n}(d'_i)$ が使用知識との不相関程度として定義される．更に，推論過程の時刻 t における上限距離を $\max_{i \in n}(d'_i(t))$ で，推論過程における使用知識との平均の不相関程度を

$\bar{d} = \text{average}_t \left(\max_{i \in n}(d'_i(t)) \right)$ で，それぞれ推論過程における使用知識との最大不相関

程度と最小不相關程度を $d_{\max} = \max_t \left(\max_{i \in n} (d'_i(t)) \right)$ と $d_{\min} = \min_t \left(\max_{i \in n} (d'_i(t)) \right)$ で表す.

ただし, $t = 1, \dots, T$, T は推論過程の時間である.

図 5-9 には前述例題における知識半径と推論用の距離値との関係を示す. 知識半径の値 q の変化につれて, $d_{\max}$ と $d_{\min}$ と $\bar{d}$ はほとんど同じ変化の傾向があるが, $d_{\max}$ と $d_{\min}$ より $\bar{d}$ のほうが知識半径の値の増大につれて明白に増えることがわかる. 知識半径の値 q と $\bar{d}$ の変化により, $\bar{d}$ は知識半径 q にほとんど線形関係がある.

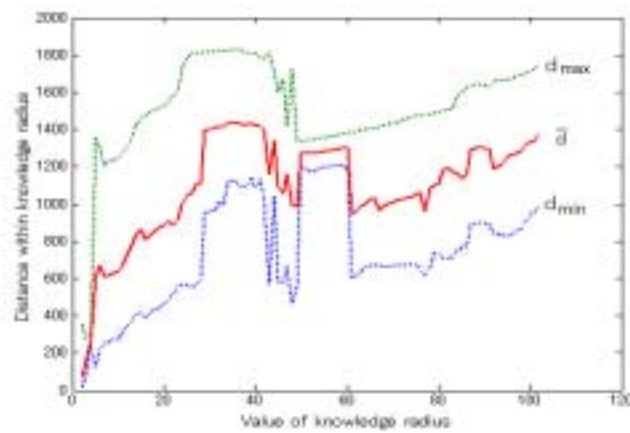


図 5-9 知識半径と推論用の距離値との関係

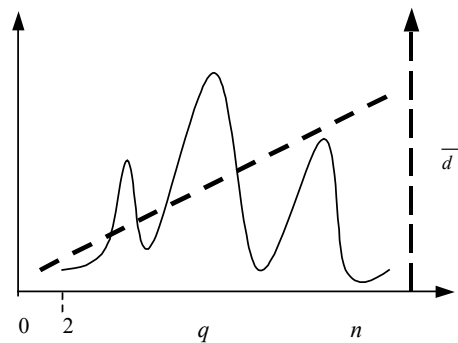


図 5-10 目的関数(点線は近似な知識半径範囲における平均最大距離の変化を表す)

一方, 図 5-9 に示すように知識半径 q と $\bar{d}$ の線形関係が成り立つが, 有効な知識半径は $\bar{d}$ に直接に関係がない. また, 表 5-2 の有効な知識半径の探索結果から, 幾つかの有効な解が存在したことがわかる. 従って, 性能評価として,

問題解決における目的関数は q に非凸関数の関係がある可能である。具体的に、図 5-10 のように、知識半径の最適化に関する非線形計画問題に帰着されることが出来る。

$$\begin{aligned}
 &\text{最小化(あるいは最大化)} \quad f(q) \\
 &\quad g_i(q) \leq 0, \quad i=1,2,\dots,k \\
 &\text{制約条件 } h_j(q)=0, \quad j=1,2,\dots,m \\
 &\quad 2 \leq q \leq n
 \end{aligned} \tag{5-23}$$

ただし、 $f(q)$ は目的関数であり、 $g_i(q)$ と $h_j(q)$ はそれぞれ不等式制約と等式制約である。

3) 模倣有効性の説明

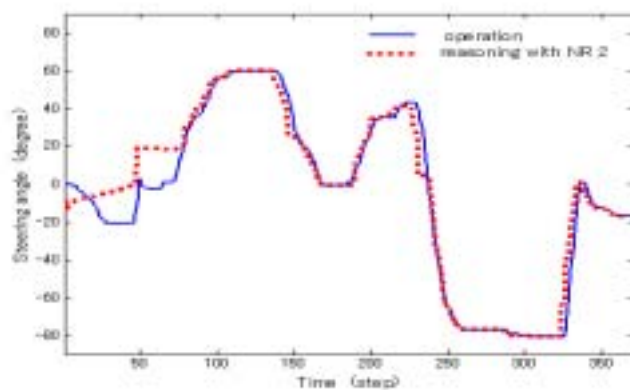


図 5-11 操作者と推論法による回転角度

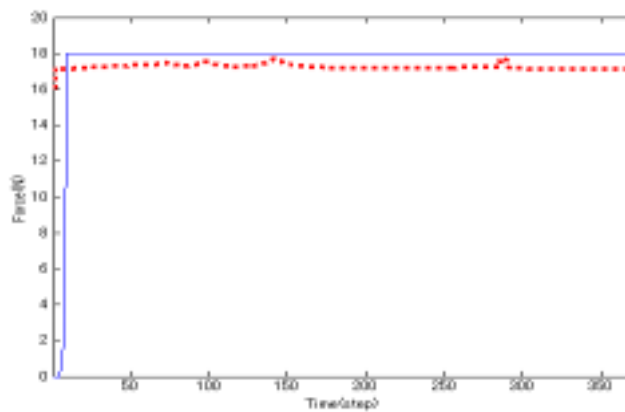


図 5-12 操作者と推論法による推進力

最適な知識半径を導入した距離型ファジィ推論法に基づいて、良い模倣効果を保証するが、経路の相似程度にかかわらず、局所的行動戦略としては、図 5-11 と 5-12 には模倣過程において異なる場合変化につれて制御用の推進力と回転角度の変化を示す。それから、最適な知識半径を導入した距離型ファジィ推論法に基づく効果は、安定推進力を保ちながら、相似の回転方向を追従したことがわかる。従って、最適な知識半径を利用して障害物回避行動戦略を有効に模倣すると確認した。

また、知識半径の適用は知識半径以外のルールが別の時刻に使用されたことを制限しないので、図 5-13 には全てのルールがほとんど一定の使用頻度で使用されたことを示す。それから、全部のルールは行動模倣に有用であるが、局所で最も有用なルールのみを注目して有効な模倣効果を達成することができる。有効な知識半径が推論過程において正確な知識利用を保証することを意味する。

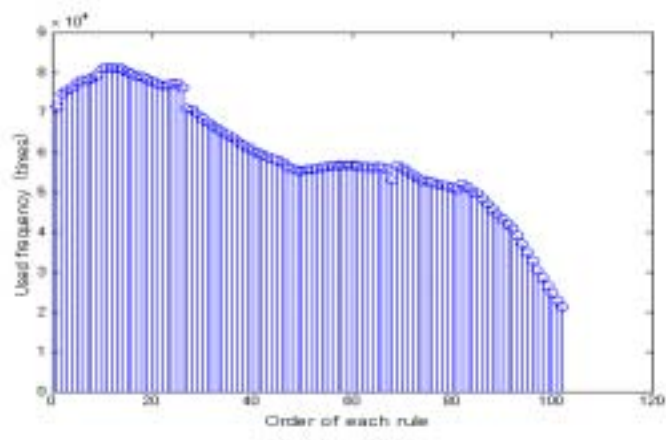


図 5-13 全てのルールの使用頻度

4) 知識のロバスト特徴

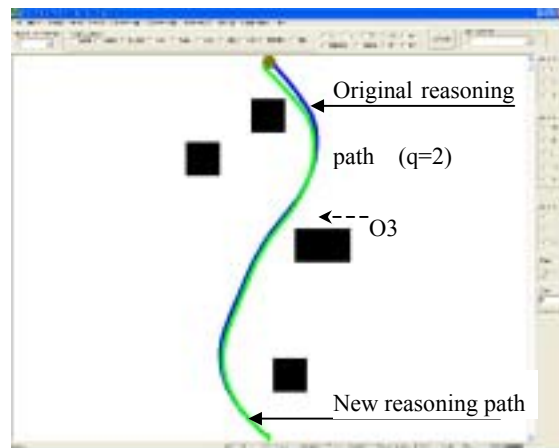


図 5-14 推論によるエージェントの最初の走行経路と新しい走行経路

模倣とはコピーと等しくないが、ファジィルールを用いて操作者の行動知識を表示する際に、知識はあいまいさとロバストという特徴を持つ。環境パラメータは少し変化しても、現有知識の許可範囲において、現有知識と知識半径に基づいて、やっぱり有効な模倣も実現できる。図 5-14 と 5-15 の示すように、障害物 03 は(700, 400)から(650, 400)に変わり、障害物 04 は(600, 700)から(600, 650)に変わったと、推論によるエージェントの走行経路も同じ傾向を持って少し変化したことがわかる。従って、ファジィルールの抽象的知識表示は拡張性を持つ特徴が分かる。

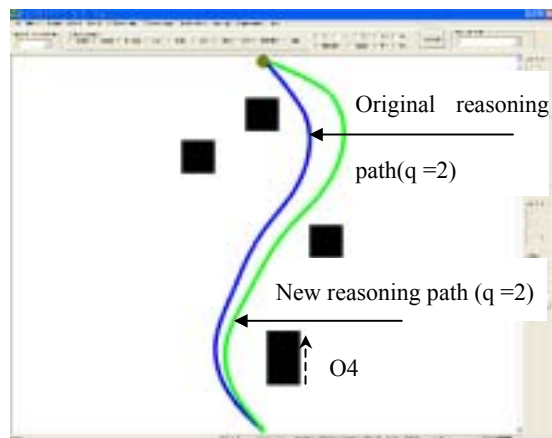


図 5-15 推論によるエージェントの最初の走行経路と新しい走行経路

第4節 動的な知識半径の導入による障害物回避行動の模倣

模倣システムの構築に対しては、本章は推論方法の角度から展開されたものである。知識半径の概念を導入することにより、知識の選択的利用を推論方法に融合するファジィ推論法を提案・検討した。しかし、知識半径を導入すると同時に知識半径を如何確定するかという新しい問題を引き出す。障害物回避行動の模倣を実現するために、第3節での連続的な繰り返し計算方案を通して有効な知識半径を確定するという方法は、知識半径の性質を定性に了解しやすいが、全局的評価基準を採用するのではなく、つまりより少ない評価情報のみを利用する場合に、苦手な面が明白に現れる。そこで、より汎用的・効率的知識半径の確定法が期待されている。従って、知識半径を考慮した距離型ファジィ推論法を採用し、障害物回避行動を再現するに着目し、より汎用的・適応的知識半径の確定法を開発することを目指す。つまり、知識の利用を強調し続けながら、もっと柔軟性をもつように知識を利用することを目的とする。具体的に、障害物回避行動を模倣するに着目し、知識半径の概念を導入する距離型ファジィ推論法を行動モデルとして、まず、予見制御理論を参考することで、局所で評価を達成する模倣効果の評価基準として予見ステップを有する経路面積誤差を定式化する。そして、予見ステップを有する経路面積誤差を利用して、より汎用的・適応的知識半径の確定方法を提案する。結果的に、単純な評価基準を採用する知識半径の確定方法を通じて、知識半径が広く適用される可能性が大きくなる。

第1項 動的な知識半径の導入

推論過程において同一な知識半径を利用したため、局所ごとに異なる知識半径の利用を必要とされる問題に対処しきれないことが有り得る。もし知識半径をダイナミックに変化させることができれば、人間のように状況に応じて知識を臨機応変に利用でき、行動模倣をより適切に実現することが可能である。前節の連続的な繰り返し計算方案により確定した知識半径比べて、使用される知

知識半径の値は、推論過程において一致しないかもしれない。更に、知識半径の定義を静的な知識半径と動的な知識半径を分かれるわけである。ただし、静的な知識半径とは、推論過程に一致する知識半径であると定義される。動的な知識半径とは、推論過程に適応に変わる知識半径であると定義される。従って、本節では、動的な知識半径という概念を行動知能モデルに導入し、知識半径を活用することで、障害物回避行動の模倣を実現するに着眼している。

具体的模倣手法として、以下の4つのステップから構成される。

STEP1：操作者は運転シミュレータの特定のタスク環境において、障害物を回避しながら、できるだけゴールに到着する。

STEP2：タスクが終わったら、ファジィ推論学習則に対し、第4章の学習アルゴリズムを用いて行動後の結果として残されたデータから推論用のファジィルールを操作者の障害物回避戦略として獲得する。

STEP3：エージェントに知識を付加する行動再現システムにおいて、STEP2で獲得されたルールをエージェントの知識、知識半径を導入した距離型ファジィ推論法をエージェントの行動モデルとして、局所的評価基準により確定する最適な知識半径や有効な知識半径を用いて、STEP1と同じタスク環境において、エージェントを行動させる。

上述の模倣手法を適用する過程に、問題解決のための局所的評価基準と知識半径の確定は以下に詳しく紹介する。

第2項 予見機能を導入した評価基準

局所で効果評価を達成するために、局所で可能な情報しか利用しない。しかし、目標経路として操作者によるエージェントのデータと、エージェント行動モデルと、知識半径の変化範囲という既知条件の上に、もしエージェント行動の過程において、過去情報のみにかかわらず、現在情報と未来情報も結び付けて知識半径の確定に向ける効果評価基準を引き出せば、知識半径をダイナミックに確定することが可能となると考えられる。つまり、未来にエージェントの行動にどのような要求がなされるのかを知っていることを利用して、エージェ

ントに対する現在の知識半径の確定を決定する。

制御系設計の中に未来情報を利用することは非常に有益なので、予見&予測機能を取り入れ、理論的な考察を行ったものとしていろいろな研究がある[13,14]。予見とはこれまでの話でわかるように、目標値や外乱信号の未来状況がはっきり分かる場合を意味する。予測とは制御対象の出力などの未来値がはっきりとはわからない場合で、何らかの方法によって未来の状況を予想することを意味する。ここでは、目標経路と明確な行動知能モデルのわかる上に、もともと予見制御理論が局所的評価基準の確定および知識半径の確定に役に立つだろうと思われる。予見制御理論とは過去及び現在における目標値だけでなく、予見ステップを通じて未来における目標値をもみながら、全制御期間に渡ってある評価関数を最小にするという立場による最適制御理論である。予見機能を表現するものとして予見ステップは、制御系の基本的なロバスト性を確保するのではなく、目標値の未来値を利用することにより追従性能を大幅に改善する。また、予見ステップに関する極限の性質については、未来目標値は1ステップ未来までで十分である、かつある程度以上遠くの未来の情報はほとんど効果のないと指摘された。

前節での全局的評価基準による有効な知識半径の確定に比べて、予見ステップを有する経路面積誤差を局所的評価基準として、現在時刻における有効な知識半径の値を決め、そしてエージェント行動モデルに基づいてエージェントの制御量及び軌道位置を計算する。

具体的に、図 5-16 の示すように、推論によるエージェントの走行経路と人によるエージェントの走行経路との予見ステップにおいて囲む面積誤差を利用して、予見ステップを有する経路面積誤差は式(5-24)に定式化される。

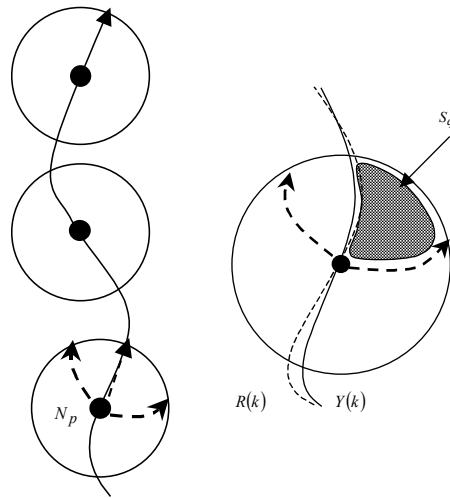


図 5-16 予見ステップを有する経路面積誤差

$$S_q = \int_k^{k+N_p} s_q(t) dt \quad (5-24)$$

ただし,

q : 知識半径の値

N_p : 予見ステップ

k : 軌道位置

$s_q(t)$: 軌道位置 t における $Y(k)$ と $R(k)$ という 2 つの曲線が囲む標本面積

$Y(k)$: 推論によるエージェントの走行経路

$R(k)$: 人によるエージェントの走行経路

もし人によるエージェントの走行経路 $R(k)$ と推論によるエージェントの走行経路 $Y(k)$ という 2 つの曲線が予見ステップにおいて囲む面積誤差がゼロであれば, 模倣効果が最適である. その面積が小さければ小さいほど, 模倣効果程度が良いと思われる.

予見ステップが計算時間に影響を与えるので, より小さい予見ステップを用いて最適な知識半径や有効な知識半径の確定を加速することができる. そこで, 知識半径をダイナミックに確定すると同時に, 予見ステップの値を確定してみる.

第3項 知識半径の確定

第2節に提供された知識半径の確定への通用な枠組を基にして、予見ステップを有する経路面積誤差を局所的評価基準とすると、知識半径の値が適合な知識運用のために適応に変わることが可能となる。

具体的知識半径の確定方法は、次の5つのステップから構成される。

STEP1: 人によるエージェントの走行経路から比較起点を選択する。

STEP2: 知識半径の範囲における全ての値をエージェント行動モデルに導入すると、STEP1で選択した比較起点から、人によるエージェントの走行経路と推論によるエージェントの走行経路の間に予見ステップを有する経路面積誤差を計算する。

STEP3: 最小の予見ステップを有する経路面積誤差を持つ知識半径の値 q を導入した距離型ファジィ推論法に基づいてエージェントの制御量及び軌道位置を計算する。

STEP4: 閾値を指定して、知識半径の範囲を減らす。

STEP5: 次時刻においてSTEP1に戻る。

知識半径の同定法について、以下に詳しく説明される。

- 1) STEP1で、人によるエージェントの走行経路と推論によるエージェントの走行経路との長さ相違にかかわらず、エージェントが軌道位置 k に位置すれば、(5-25)式により人によるエージェントの走行経路からエージェントの現在位置と一番近い点を比較起点として選択する。

$$Start\ point = \{i \mid \min(distance(R_i, Y_k)), i = 1, \dots, OL\} \quad (5-25)$$

ただし、

$R(k)$: 人によるエージェントの走行経路

$Y(k)$: 推論によるエージェントの走行経路

OL : 人によるエージェントの走行経路における軌道長さ

- 2) STEP2 で、ある知識半径の値をエージェント行動モデルに導入してエージェントの軌道位置を予見する際に、現在時刻におけるエージェントの位置 (x, y) と速度 (V_x, V_y) を予見の初期条件とする。
- 3) STEP4 で、知識半径の初期範囲は $[2, \dots, n]$ である。ただし、 n はルールの数である。計算を減少するために、知識半径の値を確定したら、効果評価を達成する指定閾値で無効な知識半径の値を抜け、次回の知識半径の確定に知識半径の範囲を減らす。

第4項 シミュレーションによる検討

1) 学習過程

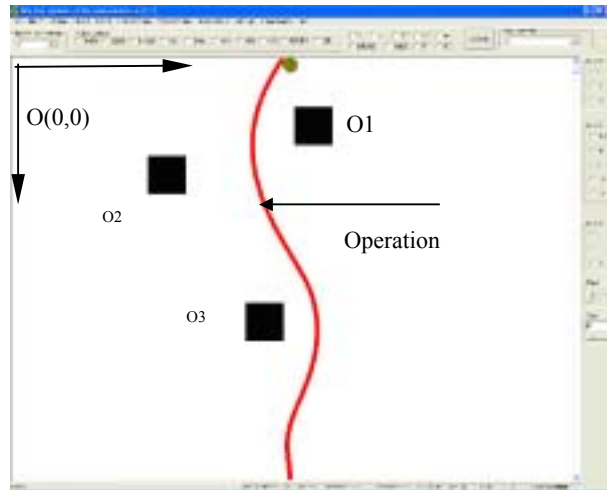


図 5-17 操作者によるエージェントの走行経路

図 5-17 のような実験場面において，左上角の座標は(0,0)であり，右向きと下向きは正方向であると仮定される．まず，操作者がエージェントを制御し，障害物を回避しながら，ゴールの到着に成功した．そして，ファジィ推論学習則に対し，第4章の学習アルゴリズムを採用して行動データからルールを獲得した．具体的に，学習前のパラメータ設定と学習後の結果は表 5-3 に示される．

表 5-3 学習過程のパラメータ

環境	4 個の静的障害物 01(600,100) 02(300,200) 03(500,500)
データの数	333
学習誤差	回転角度 15 度 推進力 3N
学習時間	5.4 秒
ルールの数	回転角度 101 推進力 17

2) 静的知識半径と動的知識半径による結果

安定の推進力変化を守る上に，回転角度ルールに対して知識半径の調整を行った．前節の繰り返し計算方案により計算した経路面積誤差は問題の目的関数

として図 5-18 に示される．最小経路面積誤差を満たす最適な知識半径 19 を決めた．最適な知識半径を導入した距離型ファジィ推論法に基づく効果は図 5-19 の示すようなエージェントの走行経路が生成された．その走行経路は操作者によるエージェントの走行経路を有効に模倣したことがわかる．また，局所的行動戦略としては，図 5-20 と 5-21 には模倣過程において異なる場合変化につれて制御用の推進力と回転角度の変化を示す．それから，安定推進力を保ちながら，相似の回転方向を追従したことがわかる．従って，静的な知識半径を利用して障害物回避戦略を有効に模倣すると確認した．

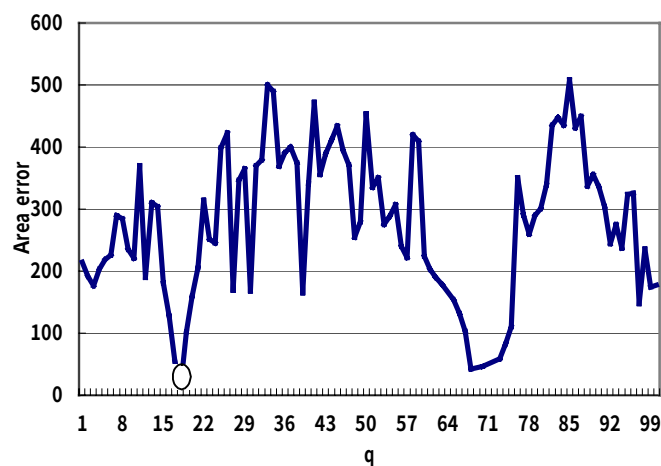


図 5-18 静的知識半径による評価関数の変化（最適な知識半径 $q=19$ ）

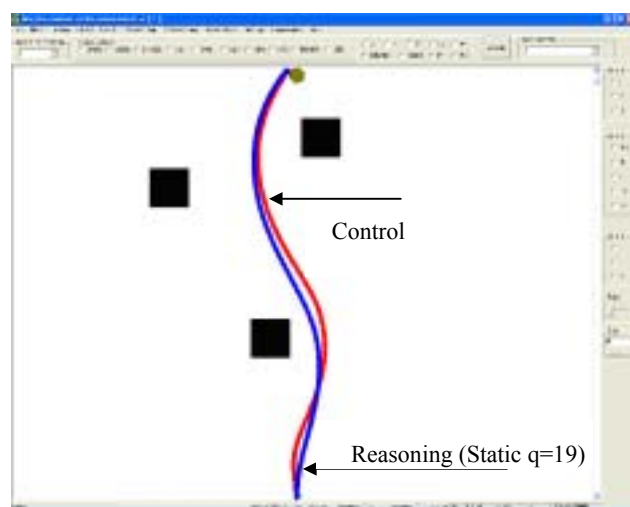


図 5-19 人によるエージェントの走行経路と推論によるエージェントの走行経路(静的な知識半径)

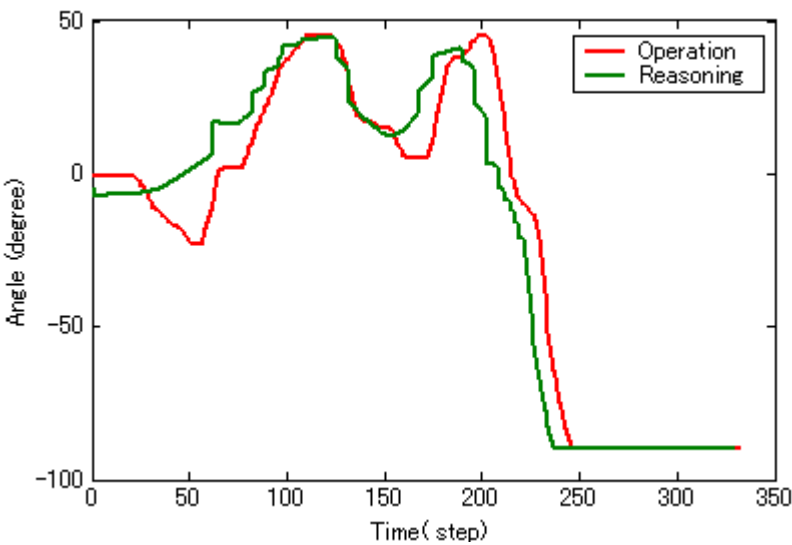


図 5-20 操作者と推論法による回転角度

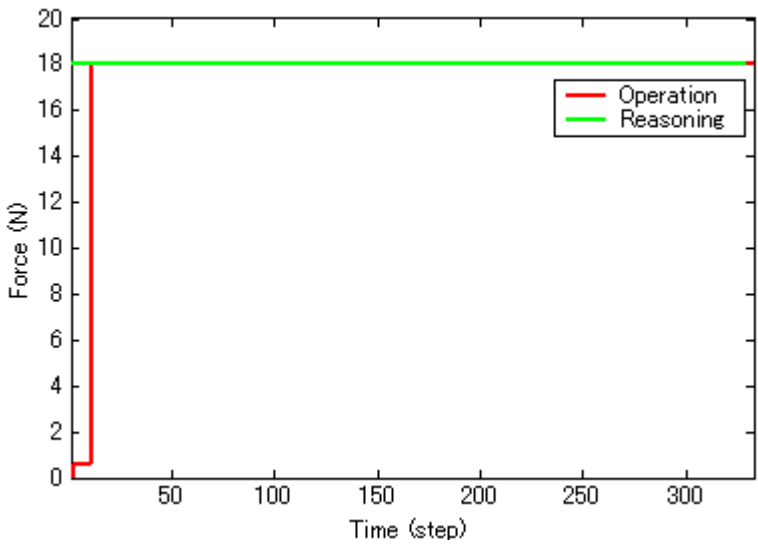


図 5-21 操作者と推論法による推進力

表 5-4 実験結果

知識半径 の種類	予見ステップ の種類	知識半径の値	評価値
動的	Step=2	最終 q=3 平均 q=10.79	2.23
	Step=5	最終 q=3 平均 q=11.57	1.79
	Step=10	最終 q=2 平均 q=9.88	99.59
静的	Step=0	q=19	13.49

また、予見ステップを有する経路面積誤差を局所的評価基準として、知識半径の確定方法に基づいて、異なる予見ステップに応じる動的な知識半径は表5-4と図5-22、5-23、5-24に示される。予見ステップが異なったが、動的な知識半径を利用して障害物回避行動を有効に模倣することができる。これから、局所で最適的な効果を達成するために、全部の知識を使う必要もないと確認した。

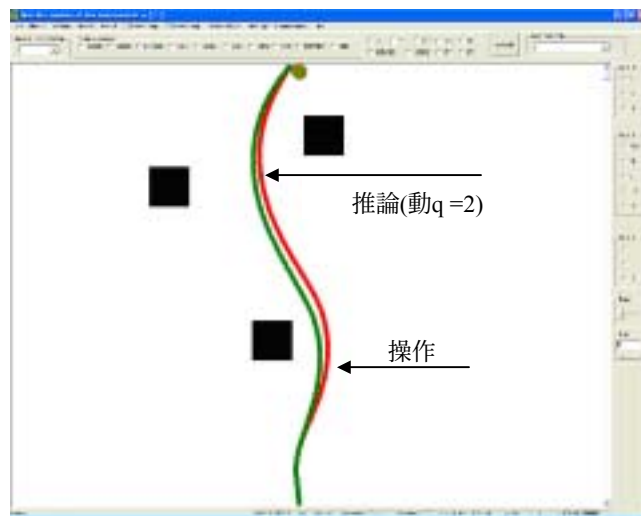


図 5-22 人によるエージェントの走行経路と推論によるエージェントの走行経路（予見ステップ2を持つ動的な知識半径）

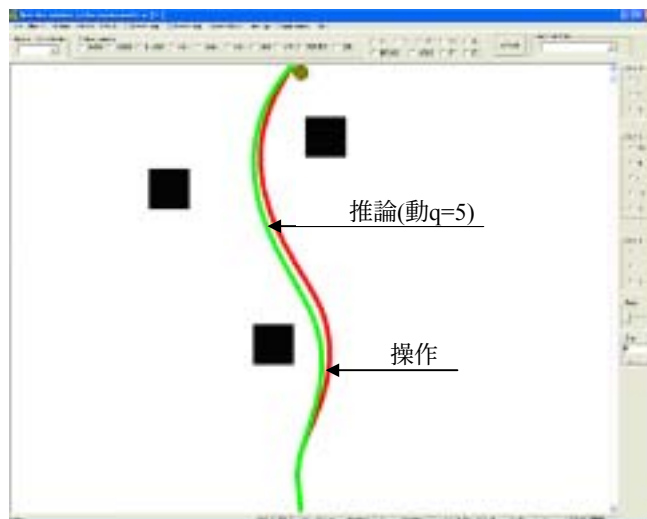


図 5-23 人によるエージェントの走行経路と推論によるエージェントの走行経路（予見ステップ5を持つ動的な知識半径）

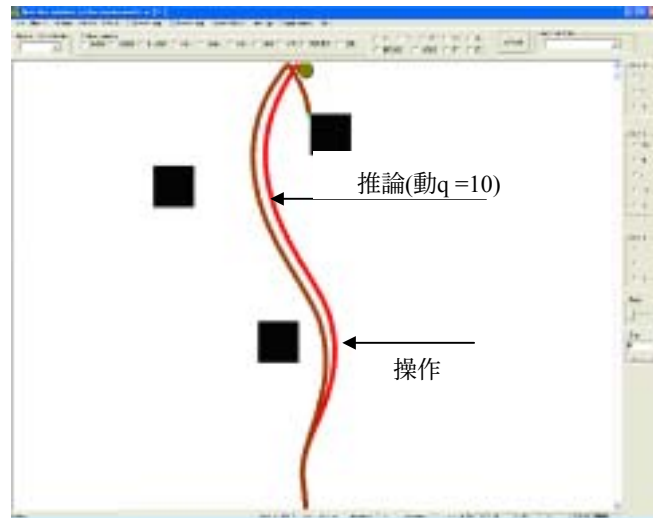


図 5-24 人によるエージェントの走行経路と推論によるエージェントの走行経路（予見ステップ 10 を持つ動的な知識半径）

さらに、表 5-4 により、動的な知識半径を用いて行動戦略を実現した過程において、予見機能を導入した評価基準による知識半径の確定方法で確定した知識半径の平均値は 10 ぐらいであり、その平均値は静的な知識半径に近づくことがわかる。もし適切な誤差閾値を選べて変化過程の終わりを判断できれば、動的な知識半径の平均値により静的な知識半径を計算し、そして静的な知識半径により障害物回避行動の模倣を速く実現することが可能となる。静的な知識半径の値を決めにくい場合に、動的な知識半径の平均値が指導的役割を果たすと少なくともいえる。

3) 予見ステップの検討

図 5-25, 5-26, 5-27 の示すように、異なる予見ステップを用いても、知識半径の同定過程において知識半径の振動状態を経過して最終的に急降安定になったことがわかる。しかし、知識半径の変化について、予見ステップ 2 の場合には急変な振動の個数は 6 であり、予見ステップ 5 の場合には急変な振動の個数は 5 であり、予見ステップ 10 の場合には急変な振動の個数は 2 である。それから、予見ステップの増大につれて、知識半径の同定過程において知識半径の急変な振動の個数が減少する傾向を示す。

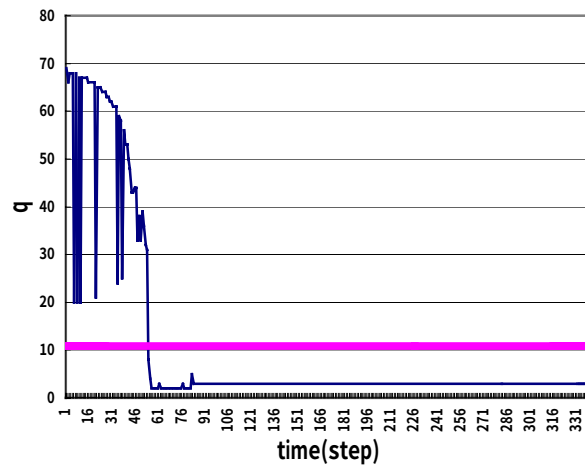


図 5-25 動的知識半径の変化(予見ステップ 2, 振動数 6)

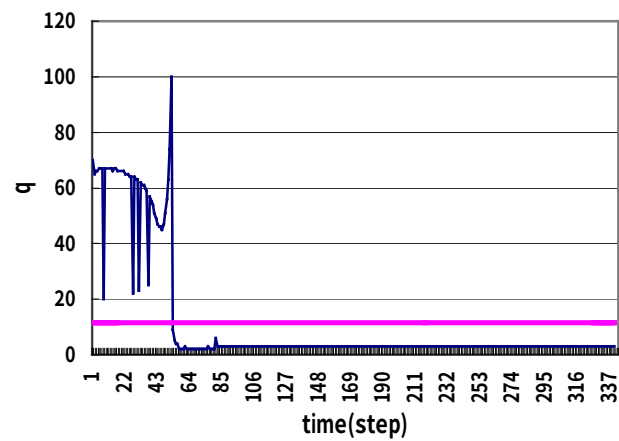


図 5-26 動的知識半径の変化(予見ステップ 5, 振動数 5)

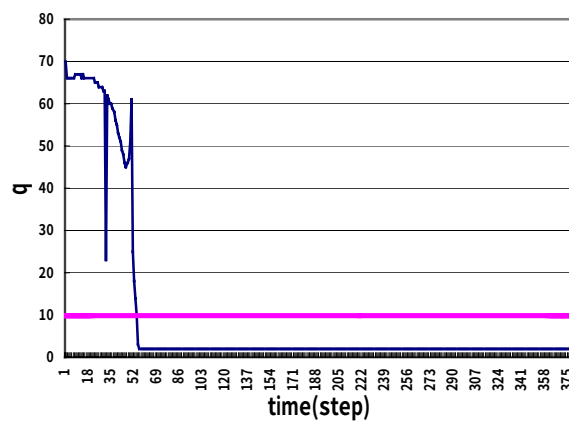


図 5-27 動的知識半径の変化(予見ステップ 10, 振動数 2)

一方、図 5-28 には異なる予見ステップを用いた知識半径の同定過程における評価関数の値の変化を示す。予見ステップの増大につれて、評価関数の値が増えた傾向を示す。

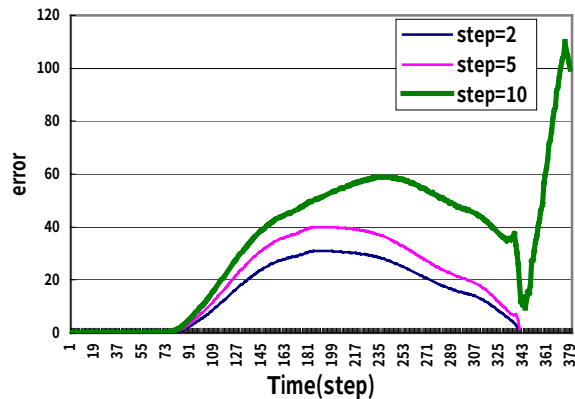


図 5-28 異なる予見ステップによる評価関数の変化

知識半径の同定過程において、少ない急変な振動が安定の制御状態を保証し、小さい評価関数値が経路模倣の有効性を保証するので、急変な振動数と評価関数値を総合的に考慮するうえに、適切な予見ステップを決めて障害物回避行動の模倣を実現することができる。

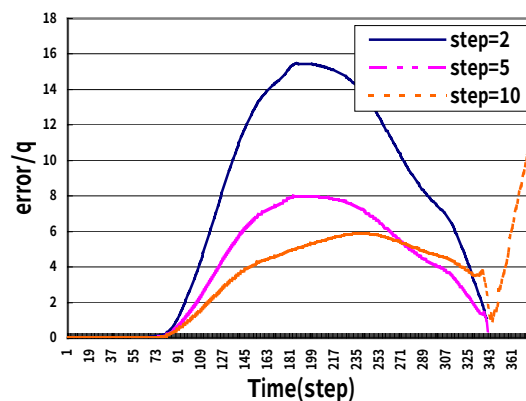


図 5-29 異なる予見ステップによる単位評価誤差の変化

更に、予見ステップの長さを考慮する単位面積誤差を評価関数値とすれば、異なる予見ステップによる単位面積誤差の変化は、全体的な評価基準として図 5-29 のように変化する。予見ステップの増大につれて、評価関数値が減ることがわかる。しかしながら、評価関数値の減少がますます明らかに変化しない。

予見ステップのある値までは評価関数値は減少するが、それ以上では変化のないことが示されている。予見ステップの大きさを選択すると、評価関数値が速く収束するということを意味する。

第5節 まとめ

第4章の知識の表現法・獲得法の角度から展開された模倣システムの構築に比べて、本章では、知識の使用と推論法の角度から模倣システムの構築を行った。何らかの学習法により得られたif-then型宣言的知識が正確なものとして、脳内に行われる知識使用の選択策略を表現する一手法を与え、知識を選択的に利用する推論法の構築について検討した。具体的に、基本の距離型ファジィ推論法を元にして、事実にも関連性のある知識の範囲を意味する知識半径の概念を導入することにより、知識の選択的使用行為を表現した。まず、知識の選択的利用を推論方法に融合する具体的な手法として知識半径を考慮した距離型ファジィ推論アルゴリズムを提案した。次に、知識半径に関する基本性質を明らかにしたばかりではなく、知識半径の確定への通用な枠組を提供した。さらに、提案推論法を検証するために、障害物回避行動の模倣を具体的な問題として取り上げ、具体的な問題解決を知識半径に関する最適化設計に定式化し、知識半径の有効性を検討した。知識半径という概念を静的な知識半径と動的な知識半径に拡張し、動的な知識半径を導入してもっと柔軟性をもつように知識を利用するに検討した。最後に、障害物回避問題のシミュレータを用いて実験を行うことにより、提案手法の有効性を示す。全部の知識を使う必要もないと結論した。

一括に、知識形式化を重視すると同時にファジィ推論における知識の利用を重視する上に、推論による問題をもっとうまく解決するかもしれないと思われる。今後、知識半径に関する研究の改善策として、いろいろな点の検討が以下に挙げられる。大域的な最適知識半径に収束するという保証を持つ最適化手法が望まれる。更に、経路面積誤差の代わりに新しい評価基準を使って知識半径の値を確定しやすいかもしれない。又、計測方法により、知識半径に関する性質を更に確定する予定である。

参考文献

1. 王碩玉,土谷武士,水本雅晴: 距離型ファジィ推論法, ファジィ学会誌, Vol.1, No.1, pp.61-78, 1999.
2. 坂和正敏: ファジィ理論の基礎と応用, 森北出版株式会社, 1992.
3. 田中英夫(監訳), 松岡浩(訳): ファジィ数理と応用, オーム社, 1992.
4. 山本康高, 川中普晴, 吉川大弘, 篠本剛, 鶴岡伸治: GA を用いた知識自動獲得問題におけるルール不活性化の効果について, 日本ファジィ学会誌, Vol.13, No.5, pp.496-505, 2001.
5. 山本隆, 石渕久生: ファジィ識別システムにおけるファジィルールの重みの設定方法, 日本知能情報ファジィ学会誌, Vol. 16, No.5, pp.441-451, 2004.
6. 王碩玉, 水本雅晴, 土谷武士: 距離型ファジィ推論法 Part17-ヒューマン推論エンジンの開発, 第19回ファジィシステムシンポジウム講演論文集, 9月8-10, 大阪, pp.433-436, 2003.
7. Hartmut Surmann and Michail Maniadakis: Learning Feed-Forward and Recurrent Fuzzy System: a Genetic Approach, Journal of System Architecture, Vol.47, pp.649-662, 2001.
8. 古橋武, 中岡謙, 前田宏, 内川嘉樹: 局所改善型遺伝的アルゴリズムの提案とファジィルールの発見, 日本ファジィ学会誌, Vol.7, No.5, pp.978-987, 1995.
9. 石渕久生, 新居学: 学習後のニューラルネットワークからのファジィ識別ルールの獲得, 日本ファジィ学会誌, Vol. 9, No.4, pp. 512-524, 1997.
10. M. C. Nechyba and Y. Xu: Human Control Strategy: Abstraction, Verification and Replication, IEEE Control Systems Magazine, Vol.17, No.5, pp.48-61, 1997.
11. 橋完太, 古橋武, 内川嘉樹: メンバーシップ関数の生成. 削除機能を有するファジィモデリング手法の一提案, 日本ファジィ学会誌, Vol.11, No.6,

- pp.1078-1088,1999.
12. 井上博行,宮阪憲治,塚本充: GA による超円錐形メンバーシップ関数を用いたファジィ識別ルールの獲得手法,第 19 回ファジィシステムシンポジウム講演論文集,9 月 8-10,大阪,pp.445-448,2003.
 13. 土谷武士,江上正: デジタル予見制御,産業図書株式会社,1992.
 14. 高橋安人: ニューラルネットワークによる非線形系及び非定常系の最適予測適応制御,科学技術社,1992.

第6章 実験による模倣システムの検証

第1節 はじめに

開発した運転シミュレータを利用して、第2章では、異なる環境における人間の障害物回避戦略の計測により、運転環境における人間の行動知能は移動ロボットの自律行動を多様化に実現するに役に立つと確認した。第3章から第5章までは、知識の表現法・獲得法と知識の使用・推論法の構築という2つ角度から障害物回避戦略の模倣システムの構築を展開した。提案の諸手法及びシミュレータの有用性について実証するために、実験による模倣システムの検討が必要であると思われる。

本章では、実験による模倣システムの有効性を考察するために、知能ロボットの遠隔操作システムを構築した。全方向行動機能を持つ知能ロボットの遠隔操作システムを開発した。全方向移動ロボットを用いる理由としては、人間のように素早く障害物を回避することが出来るからである。実験の流れとしては、まず試験者が無線LANを介して、全方向移動ロボットを、障害物回避をしながらゴールまで運転する。人間の障害物回避行動をした結果として、どのような場面・環境に応じて、どのように操作されたのかがデータとして残されている。第4章の学習法を用いて、これらのデータから人間の回避戦略を知識として抽出する。得られた知識を利用して、第5章で述べた知識の使用法と推論法に基づいて、全方向移動ロボットの自律走行を実現する。実証実験を通じて、考案した概念、提案した諸手法、開発したシミュレータが、移動ロボットの自律走行に適用できることが分かった。

第2節 知能ロボットの遠隔操作システム

運転シミュレータにおけるエージェントの代わりに、実際の移動ロボットに知識を付加するような行動再現システムを構築することで、模倣システムによる障害物回避行動の模倣効果について考察した。第5章と同じ実現目標が仮定される：図6-1の示すような実験場面において、操作者が移動ロボットの作業環境の状況を把握し、動作の判断を行う。行動後の結果としてデータを元にして知識表示を生成する。そして、それらの知識とファジィ推論法に基づく移動ロボットが操作者の障害物回避戦略を模倣することで移動する。

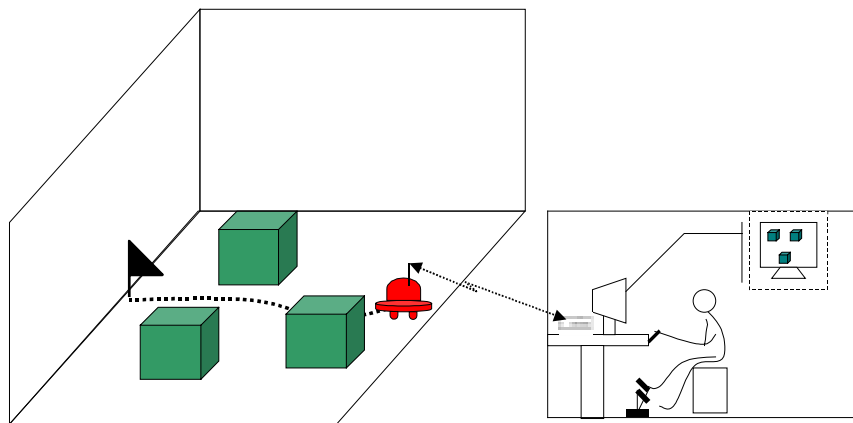


図6-1 実験場面

そのために、本節では、知能ロボット前期において、操作者への高度な臨場感の提示に重点を置いた知能ロボットの遠隔操作システムの開発を目的とする。遠隔操作作業では、操作者に環境状況を伝え、移動ロボットに操作者の意図を伝えることは重要である。そこで、遠隔操作システムの構成は、移動ロボット、操作者の意思を伝達するジョイスティック、移動ロボットに搭載されたカメラで撮られる環境情報を操作者に提示する三次元視覚提示と操作提示を含む操縦インタフェースから構成される。かつ、操作者は、ジョイスティックと操縦インタフェースを用いて、産業車両の代表的な運転姿勢の一つである「着座姿勢」で運転できる移動ロボットによる代行運転を実施した。

TCP/IP通信規約により、操縦インタフェースと移動ロボットを組み込む通信の仕組みとして図6-2にまとめる。各部分に11M無線LANカード(メルコ株式会社

WLI-PCM-L11)を付けると、10/100MスイッチングHub(メルコ株式会社製 WLR-L11-L)をゲートウエーとして無線通信を行う。

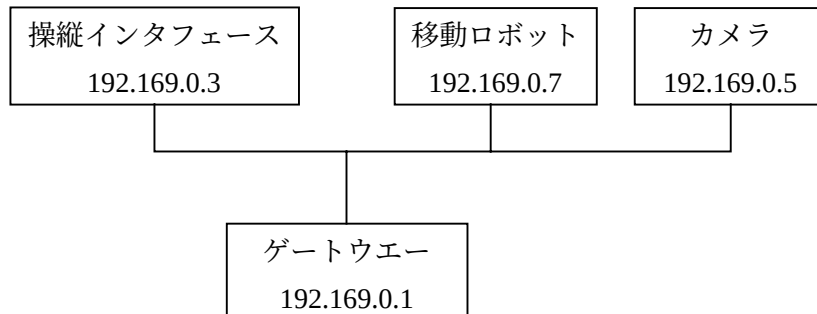


図6-2 操縦インタフェースと移動ロボットとの通信の仕組み

具体的に、開発した智能ロボットの遠隔操作システムを用いて障害物回避行動の模倣を行うための手法は、以下の4つのステップから構成される。

STEP1：操作者は特定のタスク環境において、操縦インタフェースから環境状況を了解し、移動ロボットを操縦する。同時に、センサによる環境情報と操作者による制御情報を行動後のデータとして保存する。

STEP2：タスクが終わったら、ファジィ推論学習則に対し、第4章の学習アルゴリズムを用いて行動後の結果として残されたデータから推論用のファジィルールを操作者の障害物回避戦略として獲得する。

STEP3：STEP2で獲得されたルールを移動ロボットの知識、知識半径を導入した距離型ファジィ推論法を移動ロボットの推論モデルとして、STEP1と同じタスク環境において、移動ロボットを行動させる。

STEP4：行動後の結果として残された移動ロボットのデータをもとにして、効果評価基準により、最適な知識半径や有効な知識半径を確定してその知識半径に基づく推論効果を障害物回避行動の模倣効果とする。

上述の模倣手法を適用する主要素として、移動ロボットと操縦インタフェースは以下に詳しく紹介する。

第1項 移動ロボットの構造

本移動ロボットの名前は「IMR」と称し、Intelligence imitation Mobile Robot

という言葉から省略して書いたものである。また、「IMR」も「I am Robot」という意味を含む。IMRは本研究室の開発した認知反応型ロボット[1]を元にして開発されたものである。

図6-3と6-4には移動ロボットの概観を示す。図6-5はロボットの車輪としてボール型アクチュエーターのより詳細な構造図である。仕様諸元を表6-1に示す。駆動部および制御用PCやそれらのバッテリーを含む全ての機器は本体に搭載されている。



図6-3 移動ロボット正面から見た写真

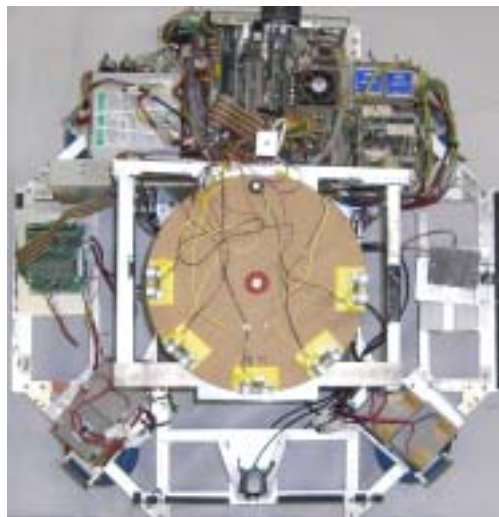


図6-4 移動ロボット真上から見た写真

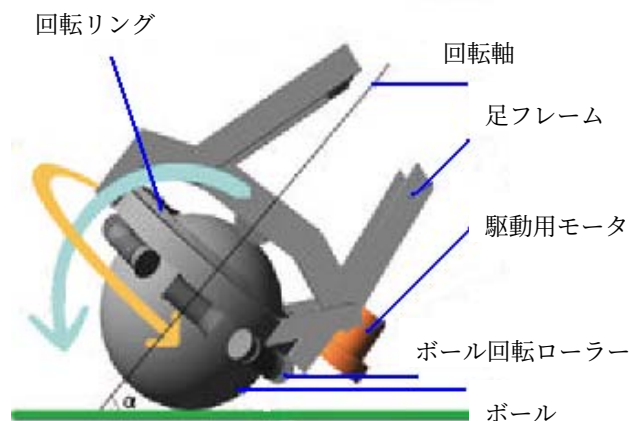


図6-5 ボール型アクチュエーター

表6-1 移動ロボットの仕様

項目	詳細
全高	600mm
全長・全幅	800mm×800mm
全重量	約30kg
移動速度	0.3m/sec

具体的構成を下記に示す。

1) ハードウェア

本体下部にある車輪駆動部は、DCモータとオムニホイールから成る駆動機構4組を90°ずつの間隔で配置することにより構成される。駆動用DCモータには定格出力90W、定格トルクは45kg・cmのものを使用している。本体中下部に制御用PCを搭載し、インタフェースボード(富士通製)によりセンサからの信号を入力、モータドライバ(岡崎産業製)によりモータを駆動している。図6-6にハードウェア構成の概要を示す。

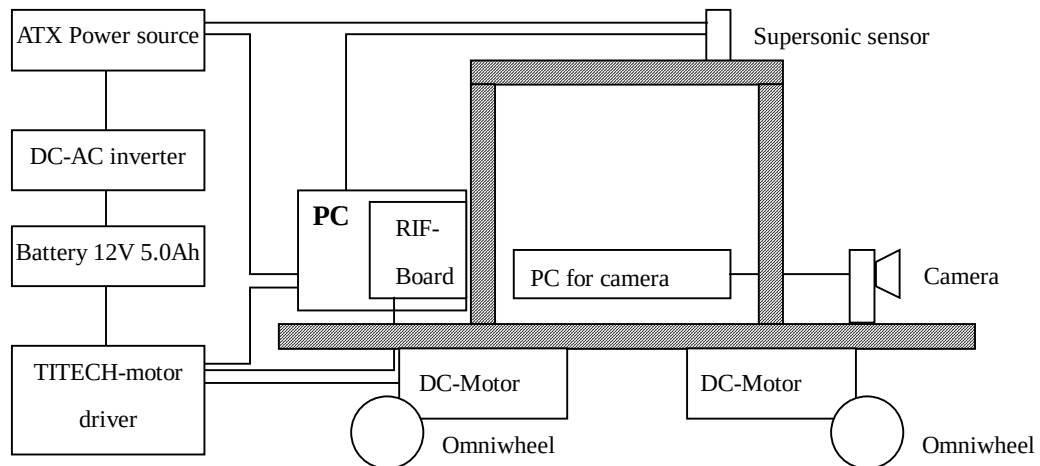


図6-6 ハードウェア構成

2) センサ

ロボット本体周囲に距離センサを設置する。距離センサには対象物の有無や距離を超音波の反射により距離の計測を行うUSセンサ2(BEST TECHNOLOGY製 BTE054)を使用している。各センサは、ロボットの周囲に放射状に合計5個が、床面から約600mmの高さに水平に配置されている。測定周期は34.7msecであり、最大測定レンジは3000mmである。センサ同士の間隔は、ロボット上での取り付け位置が最も空いている部分で60mm、人間の手程度の大きさのものや、壁状の障害物ならば確実に検出することができる。

3) ソフトウェア

全方向移動ロボットにおける手動と自動の統合を行う上で、複雑な実験の実現を保証することは非常に重要になる。人間が操縦するロボットにおいて、人間からの指令を正確な時間間隔で受け付けなければ、制御不能な状態が続いたり応答性にむらが生じるなどの事態が発生する。そこで、応答性を保証するうえで、制御を行うためのPC機(CPU Intel Celeron 633Mhz)のOSにWindows 2000を採用した。従来型104バス仕様の制御ボードより格段の処理能力を発揮できプログラミング面においてもより多くのメモリを使える環境下で開発が可能となる。

制御プログラムは、Visual C++6.0で記述されており、Windowsでリアルタイムに処理されている。特定のタスクの周期的な動作を保証することは非常に重

要になる。センサからの情報を処理するタスク(34.7msec毎)、入力を読みとってモータドライバへの出力を行うタスクが30msec毎に行う（Windows ping 命令で測定した無線LANの通信反応時間は<1msecであり、操作者の命令発送に影響を与えない。）。図6-7にソフトウェア構成の概要を示す。

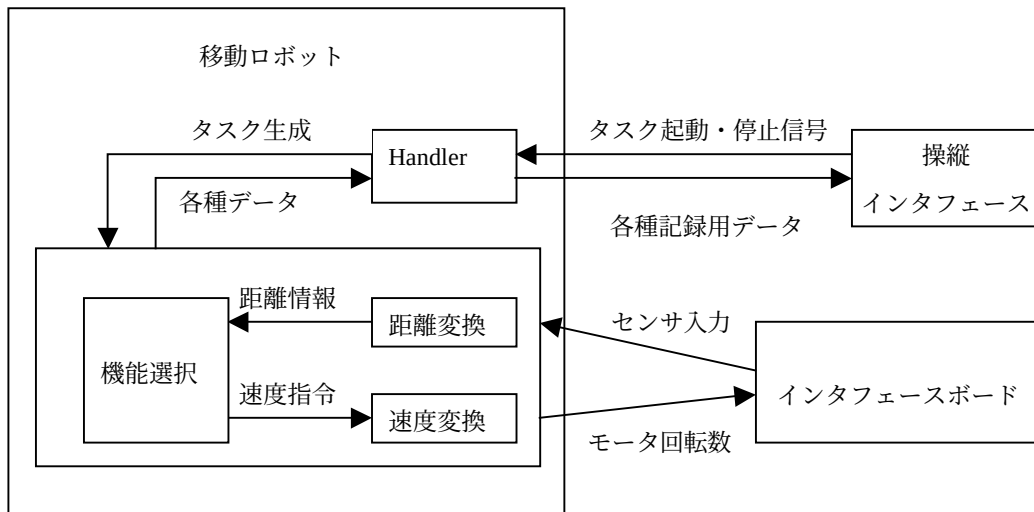


図6-7 ソフトウェア構成

第2項 移動ロボットの移動機能

全方向移動機能とは上部の姿勢を維持したまま、360度どの方向にも移動が可能となる機能である。全方向移動機能を実現することが試みられ、様々な研究開発が行われている[2-4]。ここでは、移動ロボットに人間のように全方向行動能力を持たせるために、図6-5の示すように、移動ロボットの4隅に対して、駆動輪に普通な車輪を使うのではなく、ボール型アクチュエーターを採用することにより、前進後退と方向転換だけでなく、横方向にも斜めにも、全方向に自由に動くことができる。回転軸は地面に対し45度傾いており、ボールを回転させ、ちょうど球の半径の半分の半径の車輪の如く推進力を出すことができる。それぞれの球に前向き、あるいは逆向きの回転、または停止を指示するとベクトル和の方向に動く。

実際に、ロボットにボデーの中心点の速度という制御命令を与えることでロボットを制御する。さらに、ボデーの中心点の速度を通じて各車輪の速度を変換する必要がある。座標系を図6-8に示す。

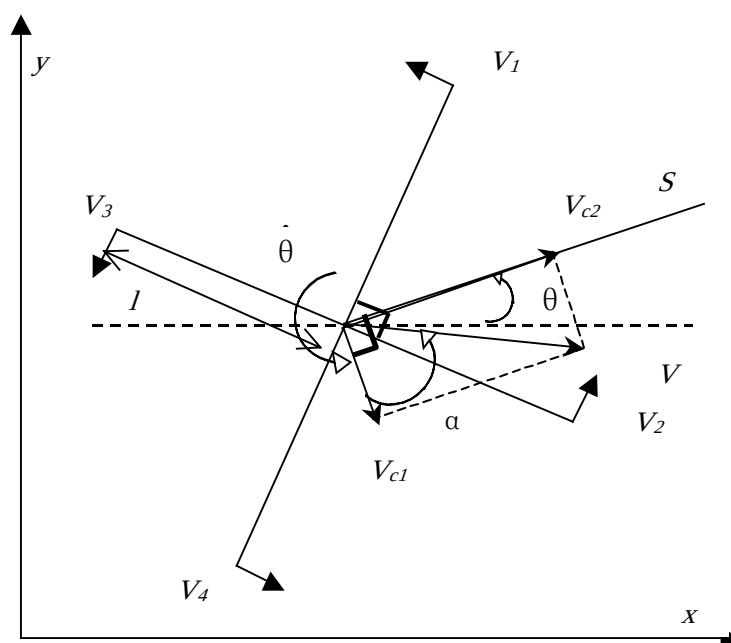


図 6-8 座標系の設定

ただし、

S : ロボットのボデーの前進方向

θ : 前進方向と X 方向の成す角

α : 操舵角度

V : ロボットのボデーの中心点の速度

V_i : ロボットの車輪の速度, $i=1,2,3,4$

l : 車輪とボデーの中心点との距離

V_{C1}, V_{C2} : V の分速度

ロボットの移動速度 (V_{C1}, V_{C2}) と各アクチュエータの速度 $V_1 \sim V_4$ との関係に対しては, C 点において, 式(6-1)と(6-2)が成立する.

$$\begin{bmatrix} V_{C1} & V_{C2} \end{bmatrix}^T = V \begin{bmatrix} \cos \alpha & \sin \alpha \end{bmatrix}^T \quad (6-1)$$

$$\begin{cases} (V_2 - V_3)/2 = -V_{C1} \cos 45^\circ + V_{C2} \cos 45^\circ & (a) \\ (V_4 - V_1)/2 = V_{C1} \sin 45^\circ + V_{C2} \sin 45^\circ & (b) \\ (V_2 + V_3)/2l = \dot{\theta} & (c) \\ (V_4 + V_1)/2l = \dot{\theta} & (d) \end{cases} \quad (6-2)$$

(6-2) 式により $V_1 \sim V_4$ を解くと, 運動学計算の(6-3)式を得る.

$$\begin{aligned} \begin{bmatrix} V_1 \\ V_2 \\ V_3 \\ V_4 \end{bmatrix} &= \begin{bmatrix} -\sin 45^\circ & -\sin 45^\circ & l \\ \cos 45^\circ & -\cos 45^\circ & l \\ -\cos 45^\circ & \cos 45^\circ & l \\ \sin 45^\circ & \sin 45^\circ & l \end{bmatrix} \begin{bmatrix} V_{C1} \\ V_{C2} \\ \dot{\theta} \end{bmatrix} \\ &= \frac{\sqrt{2}}{2} \begin{bmatrix} -1 & -1 & \sqrt{2}l \\ -1 & 1 & \sqrt{2}l \\ 1 & -1 & \sqrt{2}l \\ 1 & 1 & \sqrt{2}l \end{bmatrix} \begin{bmatrix} V_{C1} \\ V_{C2} \\ \dot{\theta} \end{bmatrix} \\ &= \frac{\sqrt{2}}{2} \begin{bmatrix} -1 & -1 & \sqrt{2}l \\ -1 & 1 & \sqrt{2}l \\ 1 & -1 & \sqrt{2}l \\ 1 & 1 & \sqrt{2}l \end{bmatrix} \begin{bmatrix} \cos \alpha \\ \sin \alpha \\ \dot{\theta} \end{bmatrix} V \end{aligned} \quad (6-3)$$

全方向移動ロボットを操作者が簡単に扱えるようにするために, 全方向の代わりに, 離散な 8 つの方向かつ 2 つの速度レベルを持つ移動手法が実現される.

具体的説明は図 6-9 と表 6-2 にまとめる.

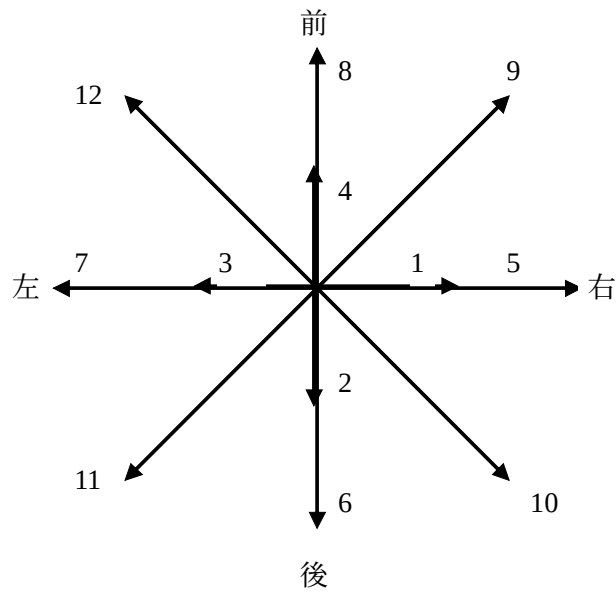


図6-9 離散な移動方向

表6-2 速度Vのモード

mode	説明
1	右方前進
2	後退
3	左方前進
4	前進
5	右方快進
6	快退
7	左方快進
8	快進
9	前右斜め
10	後右斜め
11	後左斜め
12	前左斜め

第3項 操縦インタフェースの構造

人の障害物回避戦略を便利に伝達するために、遠隔作業システムの見方から、運転シミュレータの代わりに、手動操縦と自動操縦を機能として組み込む操縦インタフェースはWindowsのOS環境においてVisual C++6.0で開発した。操縦インタフェースを通して移動ロボットを手動操縦により移動し、そして学習・推論を用いた自動操縦により移動することができる。図6-10には操縦インタフェースのソフトウェア構成の概要を示す。

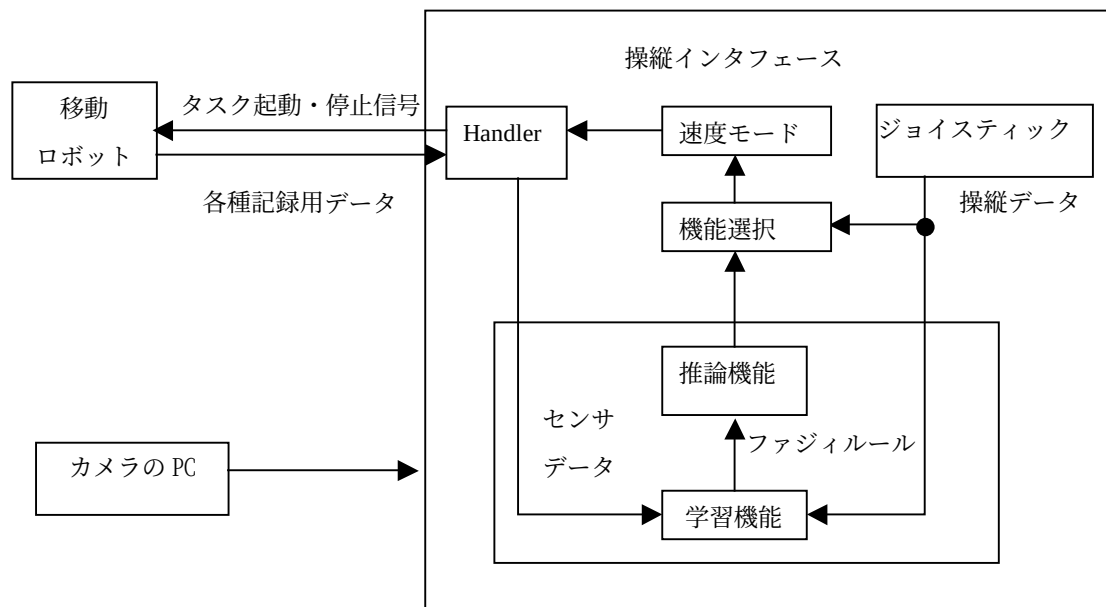


図6-10 ソフトウェア構成

手動操縦の際に、移動ロボットに搭載されたカメラ(Creative Technology 会社製 WEBCAM NX Pro)で撮られる三次元環境情報を操作者に提示する。同時に、全方向移動ロボットを操作者が簡単に扱えるようにするために、3自由度構造を持つジョイスティックをマニュアルインタフェースとして採用した。ジョイスティックの前後、左右、斜めの3自由度はそのまま移動ロボットの3自由度に1対1で対応している。具体的に、図6-11のような操縦インタフェースを利用して、操作者がロボットの作業環境の状況を把握し、コンピューターの入力装置として、ジョイスティックを上下左右に倒して、ロボットの位置を任意に移動させる。高い臨場感の提示とともに、操作者がロボットの中に入り込んだよう

な感覚で操縦することができ、高い自己投射性を実現した。

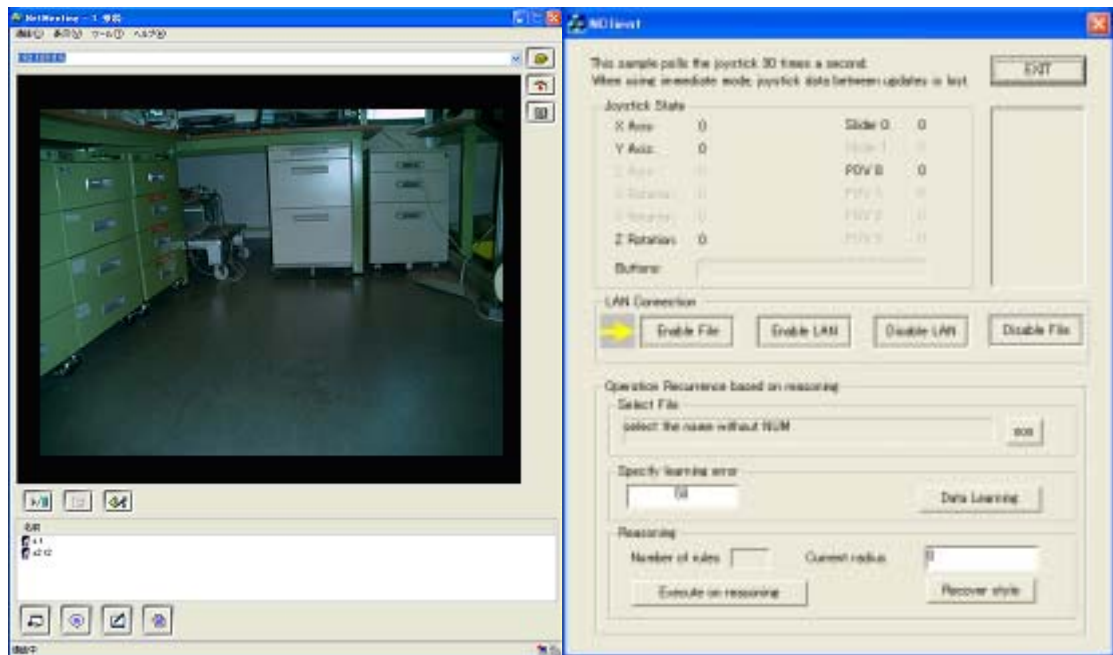


図6-11 操縦インタフェース

一方、模倣システムの有効性を検討するために、操縦インタフェースは、操作者の操縦データによる学習とルールによる推論という自動操縦機能も含む。

学習については、第4章で提案した改善された距離型ファジィ推論法の学習アルゴリズムが採用される。学習対象に対して、操作者が移動ロボットを操縦すると同時にセンサによる環境情報と操作者による制御情報とを獲得することができる。結果的ルールは以下の形式で示される。

$$\text{if } S1 \ S2 \ S3 \ S4 \ S5 \qquad \text{then } XYZ \qquad (6-4)$$

ただし、

$S1 \ S2 \ S3 \ S4 \ S5$: 五個のセンサの値による環境情報、センサの値の変化範囲 $[0,255]$

XYZ : ジョイスティックの位置による制御情報、 $X \in [-790,627]$

$Y \in [-930,215] \quad Z \in [-961,96]$

推論については、第5章で提案した知識半径を考慮した距離型ファジィ推論法が採用される。学習アルゴリズムを用いて得たファジィルールに基づいて推論を行うことで、移動ロボットがセンサの環境情報により自律に行動を行う。

また，ロボットの動作範囲内にロボットと干渉する領域が多いことから，ロボットの緊急状態を避けるために，推論部分を操縦インタフェースに置いてロボットを急停止させることができる．

第3節 走行実験による検討

開発した遠隔操作システムを用いて、移動ロボットによる手動操縦と自動操縦の実験を行う。衝突危険をなるべく回避するため、移動ロボットの走行速度を低減させる。

まず、図 6-12 のような実験場面において、操作者が操縦インタフェースを通じて移動ロボットを操縦した。最後にその環境において仮定ゴールに到着し終わった。

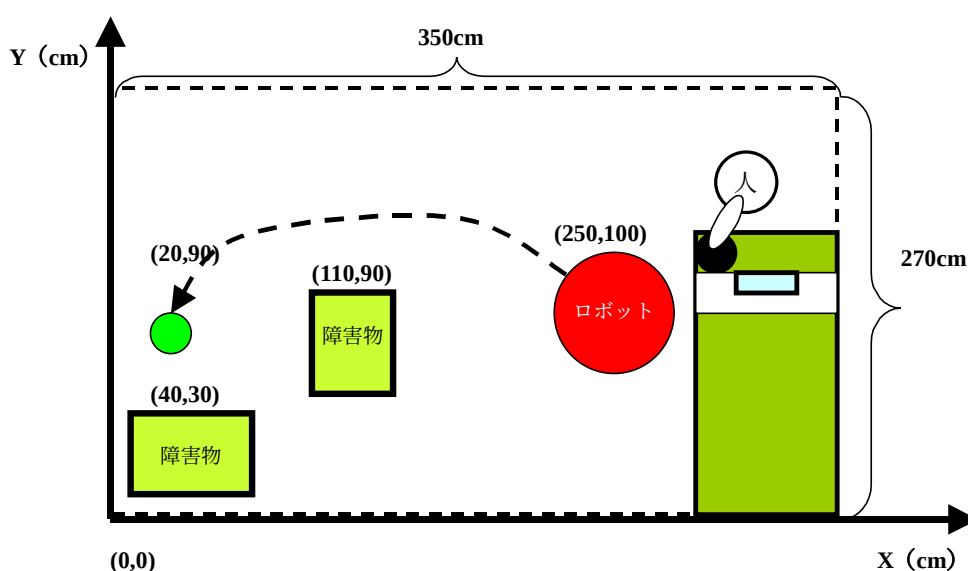


図 6-12 実験場面の説明図

つぎに、操縦インタフェースの学習機能を利用して、操作者の障害物回避戦略としてファジィルールを獲得した。学習前のパラメータ設定と学習後の結果は表 6-3 に示される。

表 6-3 学習過程のパラメータ

環境	静的障害物 2 個
データの数	435
学習誤差	50
学習時間	1.8 秒
ルールの数	96

そして、操縦インタフェースの推論機能を利用してロボットを自動操縦させる実験を行った。

図 6-13 には手動操縦による結果経路を示す。図 6-14～図 6-22 はそれぞれ知識半径なし，知識半径 30，20，16，15，14，12，10，5 の場合に自動操縦による結果経路を示す。同時に，模倣効果を考察するために，全ての結果経路は図 6-23 に表示する。

実験結果により，以下に纏めた。

- 1) 図 6-14～6-22 により，自動操縦によるロボットは，障害物を回避し，大体手動操縦による結果経路と似た結果経路を発生させたことがわかる。具体的な第 4 章の経路面積誤差を利用して評価を行える。そこで，提案の行動模倣法に基づくロボットは操作者の障害物回避行動を模倣する可能性を示す。
- 2) 知識半径 20 の場合は，最も小さい経路面積誤差があるので，最も模倣結果を得た。
- 3) 知識半径 20 の場合は，知識半径なしの場合より良い経路を生成したことがわかる。実験によって、推論過程において知識使用の重要性も確認した。

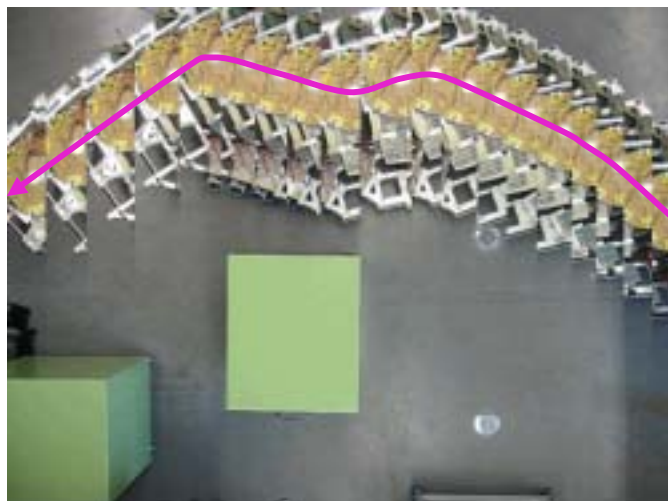


図 6-13 手動操縦による結果経路

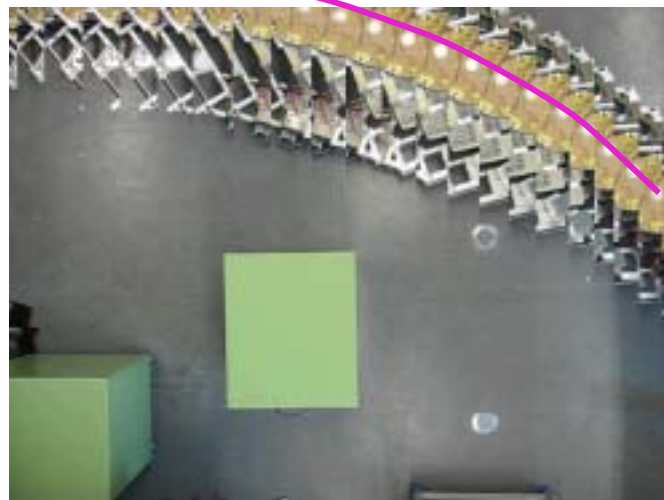


図 6-14 自動操縦による結果経路(知識半径なし)

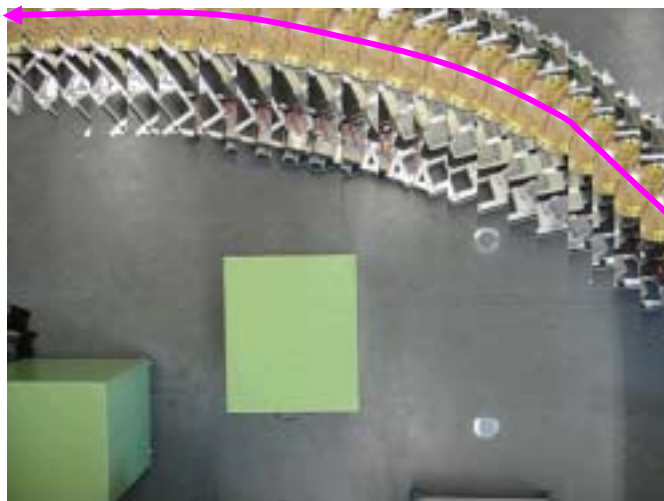


図 6-15 自動操縦による結果経路(知識半径 30)

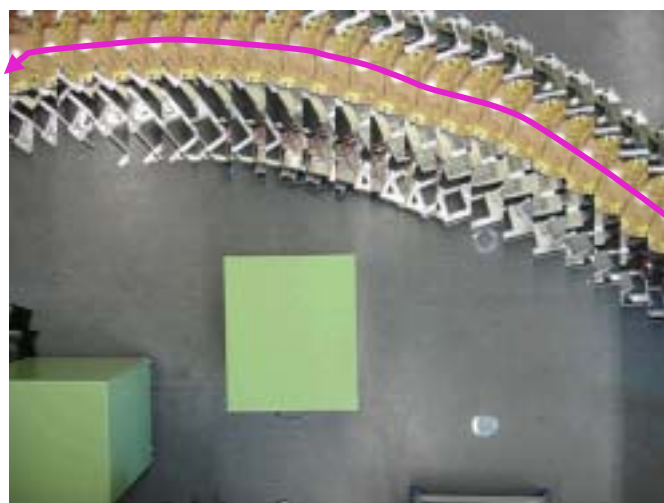


図 6-16 自動操縦による結果経路(知識半径 20)

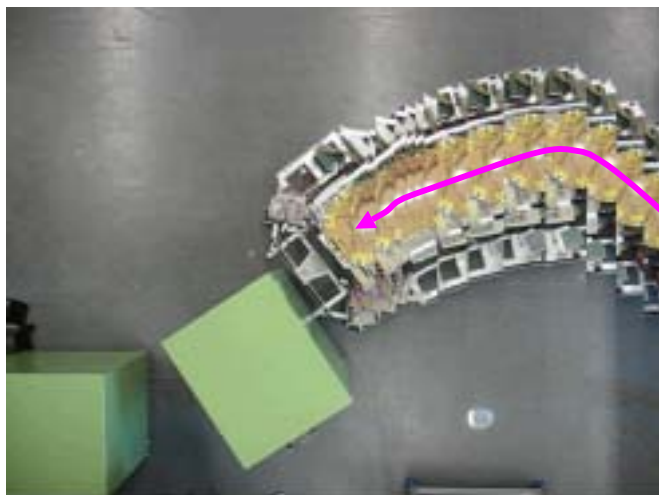


図 6-17 自動操縦による結果経路(知識半径 16)

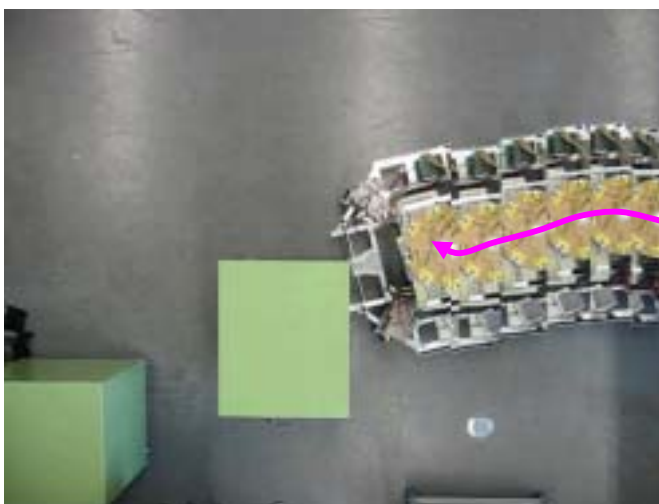


図 6-18 自動操縦による結果経路(知識半径 15)

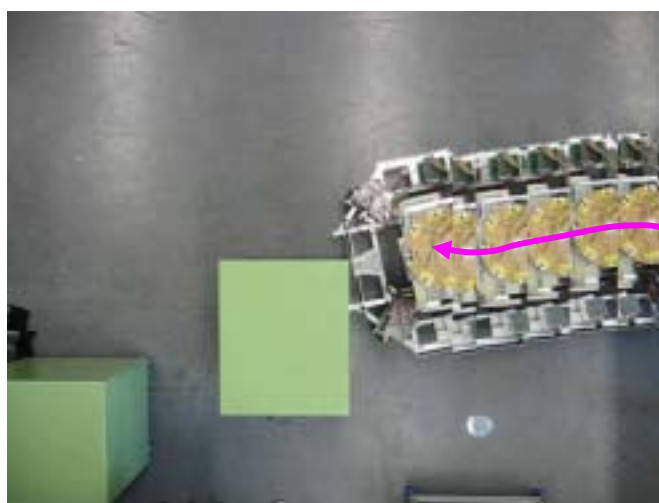


図 6-19 自動操縦による結果経路(知識半径 14)

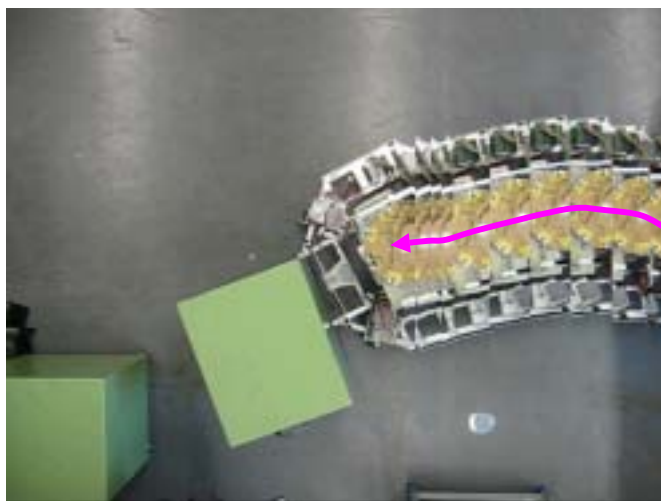


図 6-20 自動操縦による結果経路(知識半径 12)

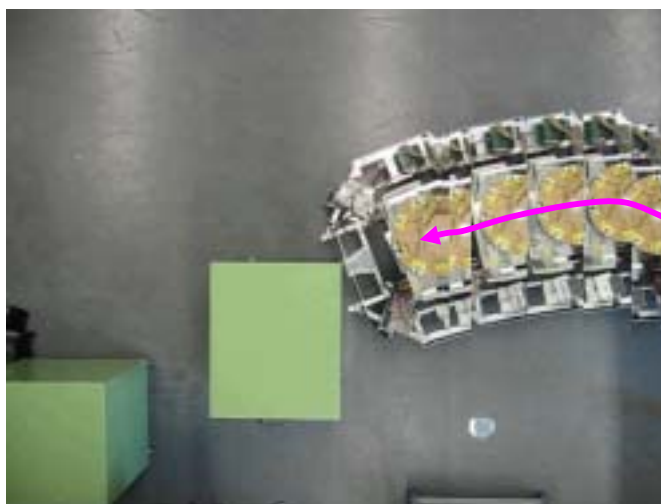


図 6-21 自動操縦による結果経路(知識半径 10)

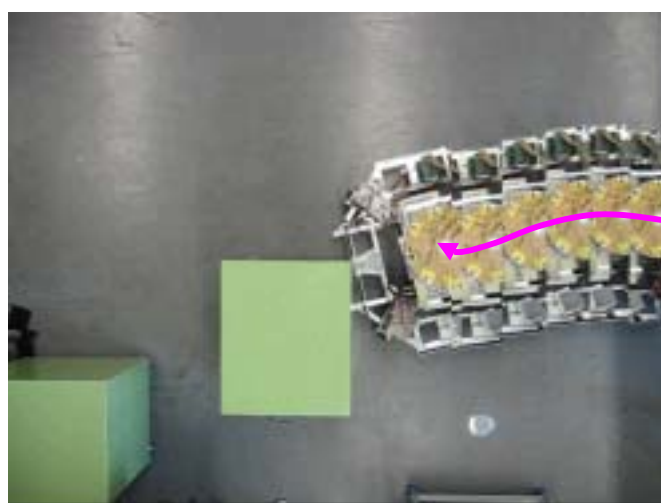


図 6-22 自動操縦による結果経路(知識半径 5)

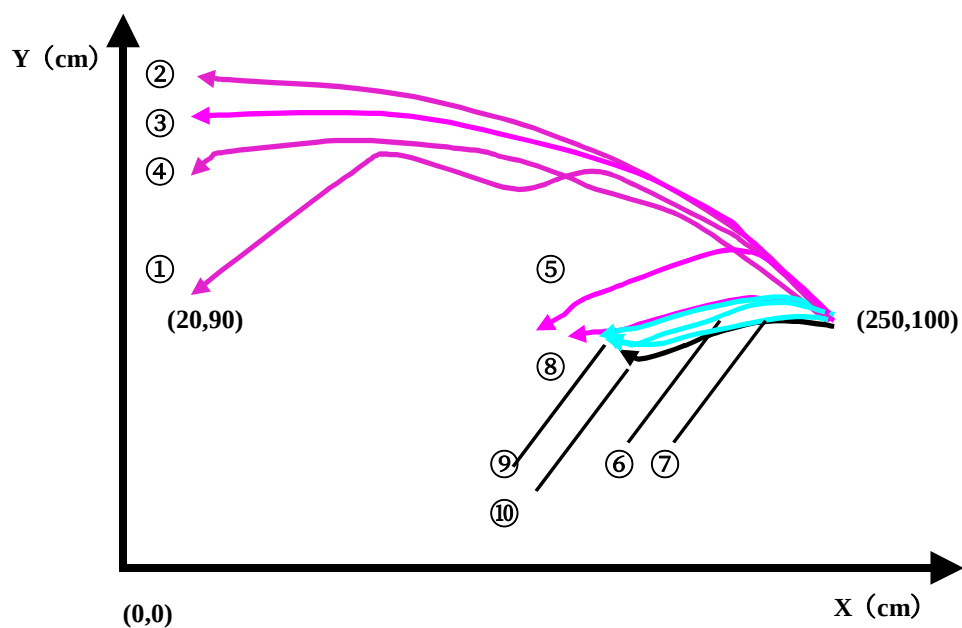


図 6-23 結果的経路の比較(①手動操縦による結果経路②知識半径なし③知識半径 30④知識半径 20⑤知識半径 16⑥知識半径 15⑦知識半径 14⑧知識半径 12⑨知識半径 10⑩知識半径 5)

第4節 まとめ

本章では、実験での模倣システムによる障害物回避行動の模倣効果を考察するために、知能ロボットの遠隔操作システムを構築した。まず、実験対象として、人間のように全方向行動能力を持ち、かつ認知機能を持つ移動ロボットを開発した。次に、移動ロボットに人の障害物回避戦略を便利に伝達するために、遠隔作業システムの見方から、手動操縦と自動操縦を機能として組み込む操縦インタフェースを開発した。移動ロボットは人の手動操縦により移動し、そして学習・推論を用いた自動操縦により移動することができる。最後に、開発した遠隔操作システムを用いて、移動ロボットによる手動操縦と自動操縦の初歩実験を行った。提案の行動知能模倣法を用いた実験により、ロボットの自律化を実現する可能性を確認した。しかしながら、知識半径の確定まだ難しい問題の一つであるが、今後の改善策として、いろいろな点が以下に挙げる。

- ロボットに定位機能をつけると、知識半径の確定に関する経路面積誤差を自動に計算できる。
- 画像処理による環境認識方法を開発すると、複雑な環境にロボットを移動させる。
- より汎用な知識半径の確定法を開発する。

本章の研究内容は、ロボットの汎用的な遠隔操作システムの開発にかかわらず、産業車両を運転するロボットの遠隔操作システムの開発に役に立つと思われる。遠隔操作システムは海中や宇宙空間、原子炉内など人間が直接作業できない場所での作業を行うロボットの操縦方法として、これまで数多く研究されてきた[5-7]。操作者が指示しなければならない部分のみ遠隔操作とし、残りはロボットの自律に任せる自律・遠隔融合が必要と考えており、今後、その方向に研究を進めていく。

参考文献

1. 浜口和洋：認知反応型ロボットの開発,高知工科大学知能機械システム工学
科学士論文,2000.
2. 大西隆之,高瀬國克：全方向移動椅子の操縦システムの研究—手動と自動の
統合化—,電学論 C, Vol. 123, No.6, pp.1109-1116,2003.
3. S. Kawasumi and K. Inagaki: Development of Electric Wheel Chair with Power
Assist System, Proc. of the 17th Annual Conference of the Robotics Society of
Japan, Vol.3, pp.1157-1158, 1999.
4. H. Kitagawa, T. Kobayashi, T. Beppu, and K. Terashima: Semi-Autonomous
Obstacle Avoidance of Omnidirectional Wheelchair by Joystick Impedance
Control, Proc. of IROS'2001, Vol.4, pp.2148-2153, 2001.
5. 高瀬國克：テレオペレーションの過去と未来, 日本ロボット学会誌,Vol. 21,
No.3, pp. 225-228,2003.
6. 蓮沼仁志,中嶋勝己,小林政巳：人間型ロボットのための遠隔操縦システム
の開発—人間型ロボットによる産業車両の代行運転への適用,日本ロボッ
ト学会誌, Vol. 22, No.1, pp.46-54,2004.
7. 堀口由貴男,榎木哲夫,桑各雅之：遠隔操作ロボットを用いた探索行為にお
ける技能解析と生態学の考察, 20th Fuzzy system symposium 論文集, pp.
394-397,20th Fuzzy system symposium,北九州,6月2～5, 2004.

第7章 結論

本論文では、人間の障害物回避行動を模倣することにより、移動ロボットの自律走行の実現を目的として、データから人間の回避行動戦略の抽出法を提案した。そのため、まず障害物回避行動の運転シミュレータを開発した。次に運転シミュレータを用いて人間の障害物回避行動を計測することにより、障害物回避戦略と環境との関係を定性的および定量的に解明した。さらに得られた知見に基づいて、障害物回避の結果として残されるデータから戦略の抽出法を構築した。最後に、戦略の抽出法を移動ロボットに実装して、行動知能の模倣で実現された自律走行実験により、提案した戦略抽出法の有用性を検討した。これらの障害物回避戦略およびその限界特性を了解し、模倣システムを構築することによって、将来的には、移動ロボットの行動能力を向上させることが可能であると考えている。

以下、本論文における各章で得られた結論を整理する。

第2章では、人間の障害物回避行動に焦点を絞って、障害物回避戦略とその限界特性について報告した。まず日常の車に代わる運転シミュレータを構築した。次に運転シミュレータにおけるエージェント・ゴール・障害物のパラメータ組み合わせ方により配置される環境において、人間の障害物回避行動を計測することにより、障害物回避戦略と環境との関係を定性的および定量的に解明した。複雑な環境における人間の障害物回避戦略は以下に纏められる：静的障害物を有する環境では、人間は障害物回避戦略を予め決めておくという傾向があり、しかも障害物回避戦略の大局的再現性が明らかに存在する。人間の障害物回避戦略の局所的特徴については、障害物の個数にほとんど関係がない、人間は1つの最も危ない障害物に着目して障害物を回避すると少なくともと言える。動的障害物を有する環境において、成功率が低くなったために、障害物回避戦略の大局的再現性は見られなく、障害物回避能力に限界が明らかに存在することがわかる。人間の障害物回避戦略の局所的特徴については、障害物の個数に関係がある、人間は同時にいくつの危ない障害物、つまり、安定限界の値

を持つ障害物に着目して障害物を回避すると少なくとも言える．障害物回避能力の限界は障害物の個数と速度と大きさというパラメータで定量化でき，障害物の個数が最も重要な要因であると考えられる．

障害物回避戦略の抽出については，その能力限界，つまり処理可能な限界数を用いてなるべく多い有用な障害物回避戦略の獲得に役に立ち，ガイドラインとして次のように纏められる：静的障害物を有する環境において，少なくとも1つの障害物の場合に対して戦略の抽出を行う．動的障害物を有する環境において，少なくとも処理可能な限界数以内の障害物の場合に対して戦略の抽出を行う．

第3章では，人間の障害物回避行動を模倣することにより，移動ロボットの自律走行の実現を目的として，人間の障害物回避戦略の模倣問題を設定して問題解決手法として模倣システムの構築を着目した．具体的には，模倣システムの構築が推論方法や知識表現や知識ベースや学習・知識獲得に深くかかわる知識ベース型システム構成の問題となる．この考えに従い，知識ベース，学習方法，推論方法の構築という角度から人間の障害物回避戦略の模倣システムの構築を行えると考えられた．

第4章では，知識の表現法・獲得法の見方から模倣システムを構築した．if-then 型宣言的知識表現法を用いて，人間の障害物回避戦略を表すことにする．まず，知識の表現方法を通して人間の障害物回避戦略を表示する有効性を確認した．そして，知識は問題解決能力をチェックするために，遺伝アルゴリズム GA を用いた知識の評価法を提案した．獲得されたルールに定量的評価基準を提供し，更なる知識獲得に指導する．つぎに，知識の不足と知識獲得の困難という問題点を確認したうえで，多変量で記述しにくい場合でも高速で知識獲得を実現するために，改善した距離型ファジィ推論法の学習アルゴリズムを提案した．距離型ファジィ推論法のオリジナルな学習アルゴリズムに知識（ルール）の生起確率を導入することにより，知識の最適化を行った．この学習アルゴリズムは，多変量，つまり言語で記述しにくい場合でも極めて速い知識獲得を実現することができる．主に，行動知能の模倣を実現するための知識獲得を

解決することができる。更に、第2章において人間の障害物回避戦略の限界を分析した先行研究の結果をもとに、1つの静的障害物と処理可能な限界数以内の動的障害物の場合対して、提案した学習アルゴリズムを充分に使用して、豊富な障害物回避知識を抽出することができる。

第5章では、知識の使用と推論法の角度から模倣システムの構築を行った。何らかの学習法により得られたif-then型宣言的知識が正確なものとして、脳内に行われる知識使用の選択策略を表現する一手法を与え、知識を選択的に利用する推論法の構築について検討した。具体的に、基本の距離型ファジィ推論法を元にして、事実にも関連性のある知識の範囲を意味する知識半径の概念を導入することにより、知識の選択的使用行為を表現する。まず、知識の選択的利用を推論方法に融合する具体的な手法として知識半径を考慮した距離型ファジィ推論アルゴリズムを提案した。次に、知識半径に関する基本性質を明らかにしたばかりではなく、知識半径の確定への通用な枠組を提供した。また、知識半径という概念を静的な知識半径と動的な知識半径に拡張し、動的な知識半径を導入してもっと柔軟性を持って知識を利用できるように検討した。最後に、障害物回避問題のシミュレータを用いて実験を行うことにより、提案手法の有効性を示す。知識形式化を重視すると同時にファジィ推論における知識の利用を重視する上に、推論による問題をもっとうまく解決するかもしれないと思われると結論した。

第6章では、実験での模倣システムによる障害物回避の模倣効果を考察するために、知能ロボットの遠隔操作システムを構築した。まず、実験対象として、人間のように全方向行動能力を持ち、かつ認知機能を付く移動ロボットを開発した。次に、移動ロボットに人の行動知能を便利に伝達するために、遠隔作業システムの見方から、手動操縦と自動操縦を機能として組み込む操縦インタフェースを開発した。移動ロボットは人の手動操縦により移動し、そして学習・推論を用いた自動操縦により移動することができる。最後に、開発した操縦システムを用いて、移動ロボットによる手動操縦と自動操縦の初歩実験を行った。

提案の行動知能模倣法を用いた実験により、ロボットの自律化を実現する可能性を確認した。

最後に、今後の研究の方向性・展望についてまとめておく。

全体的な方向性は、本章の冒頭でも述べているが、局所的障害物回避戦略を更に了解し、より効果的な運転システムへと展開させていくことである。第2章で得られた知見と第5章で提案した知識半径を考慮した距離型ファジィ推論モデルを検証するための実用化を推進する一方では、新しい模倣システムの構築について追及していく。

1つには、第2章で言及した如く、3次元やVR技術を利用して運転シミュレータの視覚効果を改善することで障害物回避戦略、特に局所の特徴を持つ障害物回避戦略の分析を試みる。さらに、計測実験により、危険性が運転者に与える影響を定量化するなどの研究を進めていく。また、知能付きのエージェントとして利用できる計測環境を境界領域として捉えていく。

また、模倣システムの構築という観点からは、知識の利用を推論方法に融合する推論モデルの構築を進めていく。行動データやルールを蓄積するデータベースを構築しながら、模倣法を統計的な見方から行動知能の構築法を進めていく。例えば、最近のベイジアンネット方法を利用することも可能である。また、新しい評価基準などを工夫することで更なる効果を持つ知識進化方法が期待できる可能性も残されている。

さらに、実験に対する更なる実験を行っていく予定である。第5章における異なる評価基準を利用して実験における汎用な知識半径の確定方法の開発を進めていく。また、画像処理によりロボットに環境認識能力を向上させているとも考えられる。汎用な知識半径の確定方法を推進する一方では、ロボットに定位機能をつけると、知識半径の確定問題を緩和するまで展開できるものと期待している。

以上、本研究の障害物回避戦略の計測を通じ、模倣システムの方法を提供するという基本的な研究の枠組みを構築することができたと考えている。

謝辞

本学位論文をまとめるに当たり、多くの方々の御指導をいただきました。ここに記して謝意を表します。

本研究の貴重な機会をくださり、終始御指導御鞭撻をくださるとともに、研究者としての姿勢や心構え、信念についてご教示いただきました高知工科大学知能機械システム工学科 王碩玉教授に心より深く感謝いたします。振り返ってみると、知能模倣やロボット研究やファジィ推論などの言葉の意味も知らなかった筆者に対して、研究テーマの設定から具体的な手法の開発まで、熱心かつユニークな数多くの御指導ご援助をくださったおかげで、本論文における研究成果を挙げることができました。特に、研究興味の培養とともに、「面白い」研究を挙げて追求し続きたいです。ここに、改めて心より厚く御礼申し上げます。

また、本研究において実験装置の開発やプログラムの開発をお手伝い頂くとともに、研究に関して貴重な御応援御指導をいただいた知能ロボティクス研究室の皆様に深謝申し上げます。

本論文の副査を御担当下さった高知工科大学知能機械システム工学科 河田耕一教授、井上喜雄教授、岡宏一助教授および同大学情報システム工学科 篠森敬三教授から種々な御助言御援助をいただいて深く感謝の意を表します。

研究室の同輩の陳貴林氏、溝渕宣誠氏、干霞氏、陶衛軍氏、苗笛氏、姜銀来氏とはお互いの研究について意見を交換する多くの機会を持つことができ、筆者の興味の幅を広げることができました。さらに、多くの先輩・同輩・後輩に大変お世話になりました。特に、最終的博士論文修正に関して貴重な御応援をいただいた高下直也氏にこころより感謝いたします。

日頃、御討論御協力をいただいた大学院生や学部生の多くの方々、学会発表や論文査読の場などにおいて御討論御助言をいただきました方々に深謝致します。

本研究の遂行に当たり、日本高知工科大学から留学生特別コース SSP 奨学金の応援をいただき、全力を打ち込んで研究を続けることができました。心から

高知工科大学に深く感謝いたします。

また、本研究は、多くの方々の御指導、御支援により成し遂げることができました。

最後に、筆者が研究の道に進むことを快諾し、いつも応援してくれた両親の尚鳳春と鄒瑛に心より感謝いたします。

著者発表文献リスト

学術論文：

1. Tao Shang and Shuoyu Wang : Knowledge Acquisition and Evolution Methods for Human Driving Intelligence, International Journal of Innovative Computing, Information and Control, Vol.2, No.1, pp.221-236, 2006.
2. Tao Shang and Shuoyu Wang : Knowledge Acquisition Based on Human Adaptability, Journal of Advanced Computational Intelligence & Intelligent Informatics. (投稿中)
3. 尚濤, 王碩玉：仮想環境における人間の障害物回避戦略の解析, 機械学会論文誌C分冊. (投稿中)
4. 尚濤, 王碩玉, 水本雅晴：脳における推論機能の一表現法, 知能と情報. (投稿中)

国際会議：

1. Tao Shang and Shuoyu Wang : Knowledge Extraction and Evolutionary Methods for Drivers' Action Intelligence, International Workshop on Fuzzy Systems & Innovational Computing 2004, pp. 174-179, Kitakyushu, Japan, June 2~3, 2004.
2. Tao Shang and Shuoyu Wang : Acquisition and Analysis of Drivers' Action Intelligence, 2004 IEEE International Conference on Robotics and Biomimetics, pp.764-767, Shenyang, China, August 22~26, 2004.
3. Tao.Shang and Shuoyu.Wang : Analysis of Humans Tactical Limitation in Computer Driving Game, Joint 2nd Int'l Conference on Soft Computing & Intelligent Systems and 5th Int'l Symposium on Advanced Intelligent Systems, pp.21-24, Yokohama, Japan, September 21~24, 2004.
4. Tao Shang and Shuoyu Wang : Analysis of Environmental Factors Limiting Humans' Obstacle Avoidance Ability, The 2005 International Conference on Active Media Technology, pp.369-374, Kagawa, Japan, May 19~21, 2005.
5. Tao Shang and Shuoyu Wang : A Novel Imitation Approach on Human's Obstacle Avoidance Ability Considering Knowledge Radius, 2005 IEEE International

Conference on Robotics and Biomimetics, pp.736-741, Hong Kong SAR and Macau SAR, June 29~July 3,2005.

国内会議：

1. 尚濤, 王碩玉：運転手の行動知能の抽出法, 2003 年度精密工学岡山地方学術講演会論文集, pp.31-32, 岡山, 11 月 22 日, 2003.
2. 尚濤, 王碩玉：ファジイモデリングを用いた運転手の行動再現システムの開発, 13 回インテリジェント・システム・シンポジウム論文集, pp.44-47, 函館, 12 月 6~7, 2003.
3. 尚濤, 王碩玉：複雑な環境における人間の障害物回避能力の計測, 第 9 回知能メカトロニクスワークショップ論文集, pp.121-126, 和歌山, 8 月 5~6, 2004.
4. 尚濤, 王碩玉：運転ゲームにおける危険性の計測と評価関数の提案, 第 14 回インテリジェント・システム・シンポジウム論文集, pp.129-132, 高知, 10 月 9~10 日, 2004.
5. 尚濤, 王碩玉：複雑な環境における人間の障害物回避知識の獲得, 第 9 回日本知能情報ファジィ学会中国・四国支部大会論文集, pp.13-16, 広島, 12 月 4 日, 2004.
6. 尚濤, 王碩玉：知識半径を利用した距離型ファジィ推論法に基づく人間の操縦経路の模倣, 第 21 回ファジィシステムシンポジウム論文集, pp.239-244, 東京, 9 月 7~9 日, 2005.
7. 尚濤, 王碩玉：人間の行動策略模倣における知識半径の同定法, 第 15 回インテリジェント・システム・シンポジウム論文集, pp.453-458, 京都, 9 月 26~27 日, 2005.
8. 尚濤, 王碩玉：知能模倣型ロボットのための遠隔操縦システムの開発, 第 10 回日本知能情報ファジィ学会中国・四国支部大会論文集, pp.15-18, 広島, 12 月 3 日, 2005.
9. 尚濤, 王碩玉：知能模倣型ロボットのための遠隔操縦システムの開発, 中国四国支部第 44 期総会・講演会論文集, pp.469-470, 広島, 3 月 8 日, 2006.

付録

ファジィ集合間の距離計算法

ファジィ集合間の距離は2つのファジィ集合の類似性を表現する．ファジィ集合間の距離も距離の公理を満たさなければならない．すなわち，任意の2つのファジィ集合 $A, B \in F(R)$ に実数 $d(A, B)$ が対応し，(1)式～(3)式を満たすとき， $d(A, B)$ を $F(R)$ 上の距離関数という．

$$d(A, B) \geq 0; d(A, B) = 0 \Leftrightarrow A = B \quad (1)$$

$$d(A, B) = d(B, A) \quad (2)$$

$$d(A, B) \leq d(A, C) + d(C, B), \forall C \in F(R) \quad (3)$$

この定義に従えば， $d(A, B)$ の値が小さいほど，ファジィ集合 A と B との類似度が大きいといえる． $d(A, B)$ の具体的な計算法は以下に与えている．この計算法は，連続値メンバーシップ関数をもつファジィ集合を対象としているが，シングルトン値をメンバーシップ関数の幅を零に収束させた時の極限状態として考えられるので，シングルトンにも使える．

定義：連続メンバーシップ関数を持つ有界凸なファジィ集合 $A, B \in \bar{F}(R)$ に対して，(4)式で定義される実数関数 $d(A, B)$ を $\bar{F}(R)$ 上の距離関数という．

$$\begin{aligned} d(A, B) = & \left[\int_0^1 |\inf A_{M_\alpha} - \inf B_{M_\alpha}|^p d\alpha \right]^{\frac{1}{p}} \\ & + \left[\int_0^1 |\sup A_{M_\alpha} - \sup B_{M_\alpha}|^p d\alpha \right]^{\frac{1}{p}} \\ & + \left[\int_{-\infty}^{\infty} \left| \left(\frac{1}{M_A} - 1 \right) \mu_A(x) - \left(\frac{1}{M_B} - 1 \right) \mu_B(x) \right|^p dx \right]^{\frac{1}{p}} \end{aligned} \quad (4)$$

ただし， $1 \leq p \leq \infty$ ， $||$ は絶対値を表わすマークである． A_M は A のメンバーシップ関数の最大値 M_A で正規化されたファジィ集合を表わす． $\sup A_{M_\alpha}$ と $\inf A_{M_\alpha}$ はそれぞれファジィ集合 A_M の α -レベル集合 A_{M_α} の上限と下限を表わす． B についても同様である．

つづいて，距離関数(4)式において， p の値を 2 として，最もよく使われて

いる典型的なファジィ集合間の簡略化距離計算公式が以下に与えられる。

1) 三角型ファジィ集合の場合

$\bar{F}(R)$ における三角型ファジィ集合は、高さ、高さを取る座標、0-レベル集合の上限と下限により完全に表される。図 1 に示すように、ファジィ集合 $A, B \in \bar{F}(R)$ をパラメータの形で $T_A(a_1, a_2, a_3, a)$, $T_B(b_1, b_2, b_3, b)$ のように書くことにする。

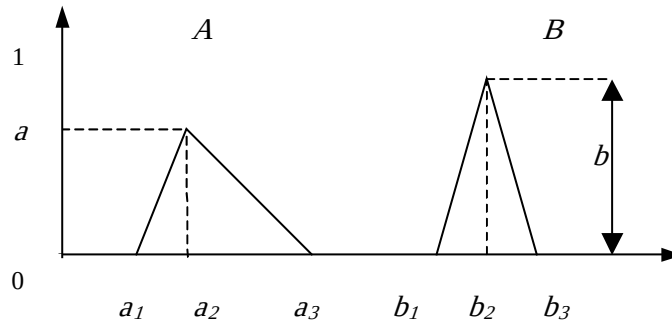


図 1. 三角型ファジィ集合

(4)式に基づいて、同図のような三角型ファジィ集合 A と B との距離計算式は(5)式になる。

$$d(A, B) = \frac{1}{\sqrt{3}} \sum_{i=1}^2 \left[\sum_{j=i}^{i+1} (a_j - b)^2 + \prod_{j=i}^{i+1} (a_j - b) \right]^{\frac{1}{2}} + \frac{1}{\sqrt{3}} \left[(1-a)[a_3 - a_1]^{\frac{1}{2}} + (1-b)[b_3 - b_1]^{\frac{1}{2}} \right] \quad (5)$$

2) 台形型ファジィ集合の場合

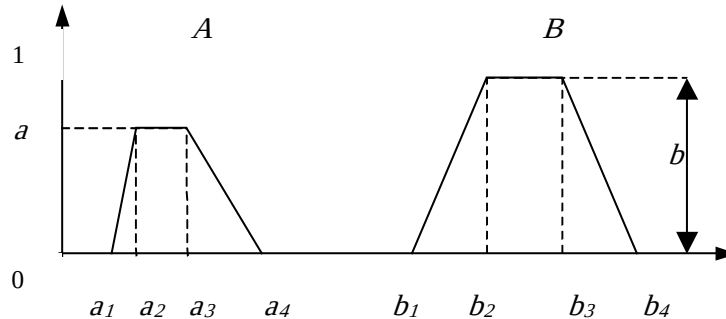


図 2. 台形型ファジィ集合

台形型ファジィ集合 $A, B \in \bar{F}(R)$ を表すには、図 2 に示すように、5つのパラメ

ータが必要である．ここでは，それぞれ $T_A(a_1, a_2, a_3, a_4, a)$ ， $T_B(b_1, b_2, b_3, b_4, b)$ のように書く．

(4)式に基づいて，与えられた台形型ファジィ集合 A と B との距離計算式は(6)式になる．

$$\begin{aligned}
 d(A, B) = & \frac{1}{\sqrt{3}} \sum_{i=1,3} \left[\sum_{j=i}^{i+1} (a_j - b_j)^2 + \prod_{j=i}^{i+1} (a_j - b_j) \right]^{\frac{1}{2}} \\
 & + \frac{1}{\sqrt{3}} (1-a) [a_4 + 2(a_3 - a_2) - a_1]^{\frac{1}{2}} \\
 & + \frac{1}{\sqrt{3}} (1-b) [b_4 + 2(b_3 - b_2) - b_1]^{\frac{1}{2}}
 \end{aligned} \tag{6}$$

3) シングルトンの場合

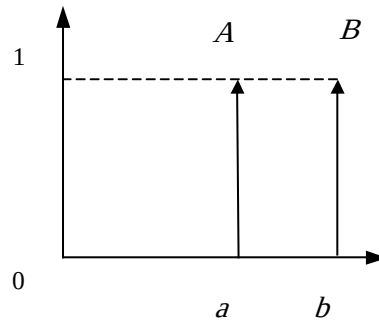


図 3. シングルトンとシングルトン

図 3 に示すようなシングルトンの場合に，メンバーシップ関数の最大値が 1 で， α -レベル集合の上限と下限が等しいので， A と B との距離計算式は次のようになる．

$$d(A, B) = |a - b| \tag{7}$$