

第 25 回ソフトサイエンスワークショップ 講演プログラム

3 月 10 日 (火)

12:30-	受付		
13:00 -14:10	特別講演 1 和多田 淳三 教授 (早稲田大学 情報生産システム研究科) タイトル:「クレディビリティに基づくポートフォリオ選択モデルの構築」 司会: 松本 義之 (下関市立大学) 日本経営システム学会中国四国支部・OR 学会中国四国支部 SCM&サービス工学研究部会 提供		
会場	A 会場	B 会場	C 会場
	TA1 司会: 中岡伊織(宇部工業高等専門学校)	TB1 司会: 小西幹彦(広島国際大学)	TC1 司会: 小田哲久(愛知工業大学)
14:20 -14:40	TA1-1 ファジィランダム重み付き Weber 問題に対する最適配置の導出 ○宇野 剛史(徳島大学), 片桐 英樹(広島大学), 加藤 浩介(広島工業大学)	TB1-1 ファジィ分割表を応用したメディア授業におけるニーズ分析 ○上江洲 弘明(早稲田大学)	TC1-1 画像内挿における微分ヒストグラムによる性能評価法 ○松田 充夫(大島商船高等専門学校), 但馬 文昭(横浜国立大学)
14:40 -15:00	TA1-2 複数台のカメラを用いたプリント基板検査経路の最適化 片桐 英樹, 郭 慶強, ●王 彬, 村中 智幸(広島大学), 羽森 寛(オー・エイチ・ティー株式会社), 加藤 浩介(広島工業大学)	TB1-2 狩野モデルによる大学数学講義の学生ニーズ分析 (2) ○高木 悟(工学院大学), 上江洲 弘明(早稲田大学)	TC1-2 色相に注目したイチョウの黄葉度算出法の提案 ●長家 瑛, 松田 修三, 小沢 和浩(法政大学), 但馬 文昭(横浜国立大学)
15:00 -15:20	TA1-3 タクシー勤務時間の最適配置 ○谷崎 隆士(近畿大学)	TB1-3 区間 UTA 法によるセミロバスト多基準意思決定支援 ○乾口 雅弘, 奥村 朗(大阪大学大学院)	TC1-3 トランプカードの画像認知・想起時脳波による推定 ートリックなしのカード 当てマジック ○山ノ井 高洋(北海学園大学), 豊島 恒(ジャパンテクニカルソフトウェア), 大西 真一(北海学園大学), 山崎敏正(九州工業大学), 菅野道夫(ソフトコンピューティグヨーロッパ研)
15:20 -15:40	TA1-4 マイクロアレイデータ解析に基づく腹膜偽粘液腫のバイオマーカー探索 片桐 英樹(広島大学), ○宇野 剛史(徳島大学), 佐々木 康平, 村上 和也(広島大学), 加藤 浩介(広島工業大学), 桑野 裕昭(金沢学院大学), 米村 豊(腹膜播種治療支援機構), 遠藤 良夫(金沢大学)	TB1-4 確率的労働市場における社会格差 ●陳 鶴, 井上 純一 (北海道大学大学院)	

注: セッション TA1 は日本経営システム学会中国四国支部・OR 学会中国四国支部 SCM&サービス工学研究部会のセッション, ●は学生による発表です。

15:40 -16:00	休憩		
	TA2 司会：河口万由香(北海道大学)	TB2 司会：藪内賢之(下関市立大学)	TC2 司会：乾口雅弘(大阪大学)
16:00 -16:20	TA2-1 Multiple human detection based on Hog feature and Statistical learning method ●Daisuke Aoki, Junzo Watada (Graduate School of Information, Production and Systems, Waseda University)	TB2-1 A Simplified Tracking Method based on Compressive Tracking ●Sun Fei, Watada Junzo (Waseda University)	TC2-1 GA-rough sets approach to fuzzy time-series based forecasting model ●Zhao Jing, Junzo Watada (Graduate School of Information, Production and Systems, Waseda University)
16:20 -16:40	TA2-2 Brain wave based on speech emotion recognition ●Kitagawa Hanayuki, Junzo Watada (Graduate School of Information, Production and Systems, Waseda University)	TB2-2 Detecting Pedestrians and vehicles in traffic scenes based on boosted-Hog feature and SVM ●Diqing Sun, Watada Junzo (Waseda University)	TC2-2 Sensitivity Analysis of Portfolio Selection Models Huiming ZHANG, ●Ye LI, Junzo Watada (Graduate School of Information, Production and Systems, Waseda University)
16:40 -17:00	TA2-3 A bilevel synthetic winery system for balancing profits and production quality ●Yu Haiyu, Junzo Watada, Li Jingru (Graduate School of Information, Production and Systems, Waseda University)	TB2-3 Human detection and tracking using panoramic image modeling ●Jiaxi Wang, Junzo Watada (Graduate School of Information, Production and Systems, Waseda University)	TC2-3 Fuzzy variable-based two-echelon supply chain model Junzo Watada, ●Xian Chen (Waseda University)
18:00 -20:00	懇親会		

3 月 11 日 (水)

09:30-	受付		
会場	A 会場	B 会場	C 会場
	WA1 司会：古殿幸雄(大阪国際大学)	WB1 司会：工藤康生(室蘭工業大学)	
09:30 -09:50	WA1-1 入力の線形結合をもつ SIRMs ファジィ推論モデルの能力について ●宮島 洋文, 重井 徳貴, 宮島 廣美(鹿児島大学大学院)	WB1-1 道路網上のボトルネックで分割したホップフィールドネットワークによる効率的な経路探索 ●仲田 紘平, 和久屋 寛, 伊藤 秀昭(佐賀大学)	
09:50 -10:10	WA1-2 RoboCup に対する SIRMs を用いた位置予測モデルの開発 ●山下 雄大, 三船 哲史, 中島 智晴(大阪府立大学), 秋山 英久(福岡大学)	WB1-2 階層型ニューラルネットワークによる音声信号分離～分離行列の多様性～ ●久木原健介, 和久屋 寛, 伊藤 秀昭(佐賀大学)	
10:10 -10:30	WA1-3 RoboCup におけるセットプレイのマーク割当に関する検討 田中 翔, 山下 雄大, 三船 哲史, ○中島 智晴(大阪府立大学), 秋山 英久(福岡大学)	WB1-3 図示化技術を用いた複数 DEA モデルからなる分析結果の解析支援手法 ●宮西 昌吾, 青木 真吾(広島工業大学)	
10:30 -10:40	休憩		
	WA2 司会：和久屋寛(佐賀大学)	WB2 司会：小西幹彦(広島国際大学)	
10:40 -11:00	WA2-1 相関ルール生成における 1 つの問題点について ●山口 直人, 劉 臣席(九州工業大学大学院), 中田 典規(城西国際大学), 酒井 浩(九州工業大学大学院)	WB2-1 ファジィ回帰モデルを用いた生活習慣病の分析 ○藪内 賢之(下関市立大学)	
11:00 -11:20	WA2-2 ラフ集合における Generalized Dynamic Reduct の並列抽出 ●高橋 智, 工藤 康生(室蘭工業大学), 村井 哲也(北海道大学)	WB2-2 ファジィ理論に基づく経営戦略分析 ○古殿 幸雄(大阪国際大学), 陳 法恩, 賈 志聖(大阪国際大学大学院)	
11:20 -11:40	WA2-3 差分進化による機会損失を考慮した多期間輸送問題の一考察 ●堀岡久喜, 安高真一郎(大阪国際大学)	WB2-3 SNS から収集した地域・観光情報の分析 ○松本義之, 藪内賢之(下関市立大学)	

11:40 -13:00	昼食		
13:00 -14:10	特別講演 2 中村 誠 教授（水産大学校） タイトル：「鮮魚の目利きの技の解析とモデル化」 司会：藪内 賢之（下関市立大学）		
14:10 -14:20	休憩		
	WA3 司会：上江洲弘明(早稲田大学)	WB3 司会：松本義之(下関市立大学)	
14:20 -14:40	WA3-1 ファジィ論理のほとんど全て(28－最終回)－今後解決すべき課題－ －含意，評価，評価のファジィ推論法－ ○中島 信之	WB3-1 社会ネットワーク分析による ICT 企業のコア硬直性に関する一考察 ○中岡 伊織(宇部工業高等専門学校), 朴 唯新(県立広島大学), 陳 韻如(滋賀大学)	
14:40 -15:00	WA3-2 ファジィ積分とファジィ推論 ○高萩 栄一郎(専修大学)	WB3-2 青年群、高齢者群における理学療法士・作業療法士・看護師などに望む対応 ○奥田 裕紀(金城大学)	
15:00 -15:20	WA3-3 BCK 代数の ordinal sum に関する一考察 ○河口 万由香, 青木 俊介(北海道大学大学院)	WB3-3 ファジィ事象とその応用 ○堀 芳樹, 椿 広計(統計数理研究所)	
15:30 -15:45	表彰式・閉会		

鮮魚の目利きの技の解析とモデル化

Analysis and Modeling of Expert's Quality Evaluation of Fresh Fish

○中村 誠¹

○ Makoto Nakamura¹

¹ 水産大学校水産学研究科

¹Graduate School of Fisheries Science, National Fisheries University

Abstract: This study aims to develop a quality evaluation system by analyzing the skills of experienced auctioneers and puffer cooks from well-established intermediary fish markets and wholesalers in order to create a models for quality evaluation. Six kinds of fresh fish and a total of 510 puffers dressed by the cooks (*Migaki*) were graded into five classes to determine what aspects of the surface color and meat freshness of the fish the auctioneers and the cooks use for evaluation. Analysis of the appearance evaluation by them, combined with information on fish coloration and meat freshness, indicated that the auctioneers and the cooks focus on four colors on the surface of the fish.

1. はじめに

魚介類は死とともに徐々に硬直していき、完全硬直の状態が一定時間続くと解硬、軟化、腐敗へと進む。一般に、解硬に至るまでの状態を「活き（生き）が良い」という。魚介類は死後直ちに筋肉の酵素による分解が始まり、解硬以降では筋肉中に増殖した細菌による分解が加わり、急速に鮮度が低下する。魚介類、特に魚類の鮮度の程度は、体表の色彩（退色の程度）や鰓の色彩（鮮紅色か？、暗濁色か？）、眼の状態（白濁の有無とその程度）、魚体の硬直や軟化の程度等の外観により推し量ることができる[1]。よって、産地市場（以下、魚市場とよぶ）の熟達した競り人[2-5]等は鮮魚の外観から、また老舗ふぐ仲卸のふぐ処理師[6-8]は身欠き後のフグ肉の外観から即座に品質の程度を見積もる。所謂「鮮魚の目利きの技」である。

水産物の流通に関わる目利きの知見は、古くから品質の評価や生産管理等に有効に活用されており、貴重な知的財産として見なすことができる。とりわけ水産物の流通の起点となる魚市場では、競り人の知識が食の安全の保証に対して重要な役割を果たしている。しかし、現在多くの魚市場では厳しい労働条件に起因して競り人の後継が育ち難い状況が続いている。同様に、長引く市況の低迷や国民の魚食離れに起因して習熟したふぐ処理師の数も減少している。よって、もし「目利きの技」の円滑な継承が滞れば、魚市場やふぐ仲卸の技術水準の維持のみならず食の安全に対しても看

過し得ない状況となる。

一方、化学分析により得られる鮮度判定指標としては、死後の筋肉中の ATP（アデノシン三リン酸）の経時的な分解過程に着目する K 値[9]があげられる。この指標は ATP の分解物総量中に占める、鮮度劣化に伴って増す HxR（イノシン）と Hx（ヒポキサンチン）の合計量の百分率から鮮度を判定するものである。K 値は 20 %以下で生食可、60 %以下が可食限界という目安を与える。K 値は初期鮮度を判定するために非常に有用な指標である。しかし、検査には魚肉採取を伴うこと、中和処理やクロマトグラフィーを経ることで結果を得る迄に長時間を要すること、専用の分析機器類や専門知識を要すること等を理由として、生鮮度を競う流通の現場では必ずしも普及が進んでいない状況にある。

以上を背景として、著者らは鮮魚の流通における技術水準の維持と食の安全の確保を目的として、競り人による鮮魚の外観評価とふぐ処理師による身欠きの品質評価を解析し、それらのモデル化を試みてきた。本稿では先ず目利きによる評価、即ち魚体体表の色彩および魚肉鮮度の関係を整理する。次に、目利きによる評価に相当するファジィ推論モデルと「目利きの技」の由縁について述べる。

2. 方法

2.1 実験

(a) 競り人、ふぐ処理師と試料魚

解析の対象とした競り人は、山口県漁業協同組

合萩地方卸売市場に所属する 50 代男性（勤続年数約 31 年）と 40 代男性（同約 23 年）の 2 名とした．試料魚には山口県が指定する重要魚種のうち市場において品質管理強化への要望が強い計 6 魚種を定めた．表 1(a)に体表の彩度が高いものから順に並べた試料魚の構成を示す．サワラは 1 尾毎の評価,その他の魚種は魚箱毎の評価とした．

また，解析の対象としたふぐ処理師は下関市内の老舗仲卸に勤務する 50 代の男性 2 名（経験年数 33 年と 28 年）とした．試料魚には下関市地方卸売市場南風泊市場で競りにかけられた「フグの王様」と呼ばれるトラフグ計 330 尾（天然 80 尾，養殖 250 尾）と「フグの女王様」のマフグ 230 尾を用いた．表 1(b)に身欠き処理後の試料の概要を示す．なお，身欠き処理は解析の対象としたふぐ処理師が中心となり行った．身欠き処理に要した時間は 1 尾あたり 3 分弱であった．身欠きは 1 尾毎の評価とした．

目利きによる試料の品質の評価尺度は，評価の高いものから順に第 1 群（優：Class 1），第 2 群（良：Class 2），第 3 群（並：Class 3），第 4 群（やや劣る：Class 4）および第 5 群（劣る：Class 5）とした．

(b) 体表の色彩と魚肉鮮度の測定

図 1(a)と図 1(b)に各々鮮魚と身欠きの体表の色彩の測定点を示す．鮮魚の測定点は背部 2 点，体幹部 3 点，腹部 2 点，頭部および眼各 1 点の合

表 1 試料

(a) 鮮魚 [mean ± SD]

Species	Total length [mm]	Weight [g]	n
Tilefish	346.9 ± 24.1	530.8 ± 115.9	320/40
White chestnut	309.1 ± 22.6	450.5 ± 94.2	350/45
Striped grunt	304.8 ± 29.1	436.2 ± 120.7	348/46
Horse mackerel	276.3 ± 42.7	229.5 ± 78.6	375/54
Filefish	263.7 ± 30.9	261.3 ± 93.8	294/50
J.S. mackerel	468.8 ± 33.0	514.5 ± 113.7	154/-

(b) フグ身欠き

Species	Length [mm]	Weight [g]	n
Tiger puffer (natural)	255.5 ± 14.6	257.8 ± 39.2	80
Tiger puffer (cultivated)	246.3 ± 16.4	196.3 ± 34.2	250
Purple puffer	254.4 ± 28.1	232.6 ± 75.0	230

鮮魚の試料数は，尾／箱を表す

計 9 点を，また身欠きの測定点には背部 2 点，体幹部 3 点，尾鰭 1 点，腹部 2 点，腹腔 2 点の合計 10 点を定めた．以下，測定点の名称は各測定線の交点（例えば，背部の頭部寄りの測定点では記号により（L1-C1）で，また鮮魚の頭部と眼球はそれぞれ（L2-C0），（E）で，身欠きの腹腔は頭部寄りより（acL3-C1），（acL3-C2）と表す．

体表の色彩の測定項目には CIE1976 表色系に規定される明度指数 L^* ，色座標 a^* ， b^* （以下，色度 a^* ， b^* とよぶ）を，また体表の保水性や粘性の程度を表す指標として輝度 L を定めた．各測定点の彩度 C^*_{ab} ，色相角 h 及び 2 点の測定点間の色差 ΔE^*_{ab} はそれぞれ式(1)から式(3)で計算した．なお，色彩の測定には測定径 8 mm の接触式色彩計（コニカミノルタ㈱：CR-400）を，また輝度の測定には非接触式輝度計（日本電色工業㈱：NL-1）を測定面積が色彩の測定と同程度になるように用いた．

$$C^*_{ab} = \sqrt{(a^*)^2 + (b^*)^2} \quad [-] \quad (1)$$

$$h_{ab} = \tan^{-1} \frac{b^*}{a^*} \quad [^\circ] \quad (2)$$

$$\Delta E^*_{ab} = \sqrt{(\Delta L^*)^2 + (\Delta a^*)^2 + (\Delta b^*)^2} \quad [-] \quad (3)$$

ここで， Δ は 2 点の値の差を表す．

魚肉の鮮度指標には初期鮮度との相関性が高い K 値を用いるものとした．式(4)に K 値の計算式

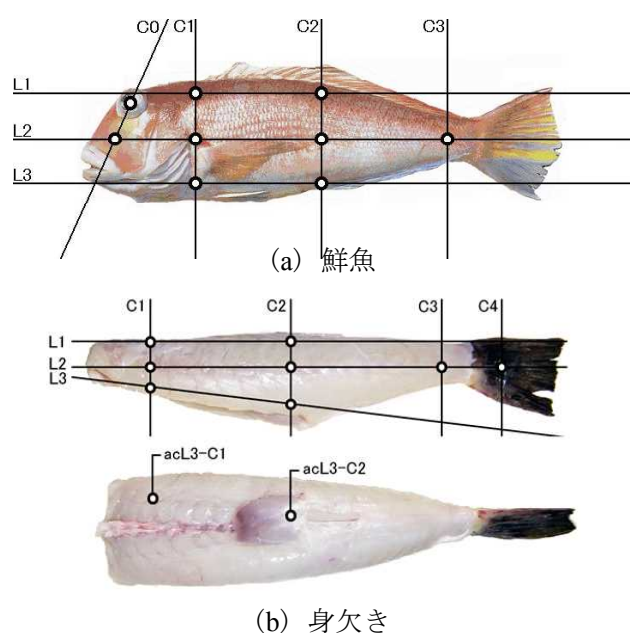


図 1 体表の色彩の測定点

を示す．K 値の測定法には HPCL (high performance liquid chromatography) 法を用いるものとし，試料の調整方法には過塩素酸抽出法を適用した[10]．なお分析用の試料肉は 1 尾あたり 1g として，各測定時毎に採取した．

K value =

$$\frac{100 \times (HxR + Hx)}{ATP + ADP + AMP + IMP + HxR + Hx} \quad [\%] \quad (4)$$

なお，色彩と K 値の測定時間は，鮮魚では水揚げから販売に至る迄の期間を考慮して競り終了後 2 時間，24 時間，48 時間および 72 時間の各経過時を，また身欠きでは身欠き処理後の水洗と水切りを兼ねる冷蔵期間を考慮して試料入手後 1 時間（身欠き処理後約 7 時間），24 時間，48 時間および 72 時間の各経過時と定めた．

魚体体表の色彩の測定と K 値分析における中和处理迄は，温度を 12.0 °C，湿度を 50 % RH 程度，照度を 300lx に保った恒温室内で，また試料魚の保管は発泡スチロール製の容器に氷納して温度を 2.0 °C に保った冷蔵室で行った．

2.2 解析

競り人の鮮魚に対する外観評価（以下，競り人の評価という）と K 値との関係，ふぐ処理師の身欠きに対する品質評価（以下，ふぐ処理師の評価という）と K 値との関係については，各測定時間における K 値の基本統計量を求めた後，評価群を 5 水準，測定時間を 4 水準とする二元配置の分散分析により確認するものとした．分析では等分散性の検定にルービン検定を，また多重比較の検定にはシェッフェの検定を用いた．なお，等分散性が棄却された場合には，改に評価群と測定時間に対してクラスカル・ウォリスの検定を行うものとした．

解析に用いる体表の色彩は，次の第 1) 項～第 4) 項に定めた鮮魚では計 23 点，身欠きでは計 29 点とした．

- 1) 部位（背部，体幹部，腹部，腹腔）内の各測定点の平均（鮮魚 3 点，身欠き計 4 点）
- 2) 各測定点（鮮魚 9 点，身欠き計 10 点）
- 3) 部位内の 2 点の測定点間の差（鮮魚 5 点，身欠き計 6 点）
- 4) 部位間の 2 点の測定点間の差（鮮魚 6 点，

身欠き計 9 点）

また，競り人の評価と体表の色彩，ふぐ処理師の評価と体表の色彩との関係は，先ず各評価群の色彩と K 値についての相関分析と無相関の検定により確認した．次に，各評価に対する体表の色彩の影響は前述の競り人の評価，ふぐ処理師の評価と K 値との関係の解析と同様に，二元配置の分散分析と多重比較を行い確認するものとした．

競り人の評価，ふぐ処理師の評価の構造は正準判別による判別分析法により解析した．式(5)に判別関数の基本式を示す．

$$Z = a_0 + \sum_{i=1}^n a_i x_i \quad (5)$$

ここで，Z は判別関数， x_i は説明変数， a_0 と a_i はそれぞれ定数と判別係数である．説明変数には，主に先の K 値との相関分析で関連が確認された体表の色彩のうち，分散分析と多重比較により 24 時間経過時迄継続して有意差 ($P < 0.01$ ，又は $P < 0.05$) を有するものを用いた．説明変数の数は判別分析式の複雑化を避けるために 4 個迄とした．判別分析の結果，判別率 70 % 以上を得た判別分析式のうち，魚肉の硬直，解硬および腐敗などの変性に伴う体表の色彩の挙動[11, 12]，体表の保水性と粘性の経時的変動[13, 14] など，および競り人とふぐ処理師の意見を考慮して最も良く反映するものを絞り込み，競り人の評価とふぐ処理師の評価の構造を検討するものとした．

2.3 モデル化と評価実験

ファジィ規則の前件部変数には，先の判別分析で絞り込んだ判別分析式の説明変数の組合せを用いた．また，後件部変数には品質評価のカテゴリーを定めた．各変数は各々 5 個のメンバーシップ関数で構成した[15]．モデルの推定精度を高めるためにシミュレーションを繰り返してメンバーシップ関数の形状をチューニングした．なお，確定演算方式には min-max 重心法[16]を用いた．

モデルの有用性を検討するために鮮魚の各魚種を 25 箱程度，天然トラフグを 20 尾，養殖トラフグとマフグを各々 150 尾ほど用いて評価実験を行うものとした．

3. 結果

競り人による鮮魚の外観評価の解析結果について

ては体表の色彩の退色により鮮度劣化が捉えにくい彩度の低い魚種を対象に述べる。同様に、ふぐ処理師による身欠きの品質評価の解析結果については鮮度劣化が捉えにくい養殖トラフグを対象とする。

3.1 評価と体表の色彩の関係

(a) 競り人の評価と体表の色彩

表 2 に競り終了後 2 時間経過時（以下、2 時間経過時と略す）における各試料魚の色彩を示す。マアジの色彩は背部、体幹部、腹部でそれぞれ暗い灰色、黄みの灰色、うすい灰色であった。ウマヅラハギの背部、体幹部、腹部はそれぞれ暗い灰色、緑みを帯びた黄みの暗い灰色、黄みの灰色を、またサワラではそれぞれ青みの暗い灰色、明るい灰色、白を呈した。各試料魚の眼 (E) を除いた体表の彩度 C^*ab の平均はマアジ、ウマヅラハギ、サワラでそれぞれ 6.23, 4.36, 2.06 となった。競り人の外観評価の特徴と傾向が確認された指標（魚体体表の色彩）、すなわち魚肉鮮度 K 値との関連が確認され、かつ群間に 24 時間経過時迄継続して有意差が確認された指標は有彩色魚種 [3,17] で 11 個から 9 個の範囲であったことに対してマアジとウマヅラハギが各 9 個、サワラが 8 個と少なかった。加えて、体表の彩度 C^*ab の低下に伴って減少する傾向を示した。

表 3 に判別分析と競り人との協議を経て選定した判別分析モデルの変数の構成を示す。各モデルとも計 4 個の変数で構成しており、モデルの判別的中率はマアジで 84.6 % (26 箱中 22 箱が正答)、ウマヅラハギで 75.0 % (24 箱中 18 箱が正答)、サワラで 71.3 % (80 尾中 57 尾が正答) となり、体表の彩度 C^*ab の低下に伴って低下した。競り人の評価と判別分析モデルでの判別結果が一致しなかったものは、マアジとウマヅラハギでは総て隣群への判別であったが、サワラでは第 4 群の 2 尾を第 2 群へと群を跨いで判別したものも混ざった。表 3 より、競り人の外観評価には体幹部を通る測定線 L2 上の色彩又は色彩の差と、魚体の体中央部に位置して背部から腹部迄を通る測定線 C2 上の腹部 (L3-C2) との色彩の差が共通して反映することが示された。また、競り人の評価のポイントが体表の彩度 C^*ab の低下に伴って頭部寄り（測定線 C1）から体中央部（測定線 C2）へと移る

傾向が確認された。

(b) ふぐ処理師の評価と体表の色彩

表 4 に 1 時間経過時における各測定点の色彩の一覧を示す。尾鰭を除く全試料の色彩の平均は明度 $L^* = 52.27$, 色度 $a^* = -0.64$, 色度 $b^* = -2.89$, 彩度 $C^*ab = 3.43$, 色相角 $h = 235.3^\circ$ (-124.7°) となり、体表全体の平均では青みの灰色を呈した。色彩の各指標間では色度 a^* と色度 b^* の間に正の相関 ($r = 0.499$, $P < 0.01$) が、また彩度 C^*ab

表 2 彩度の低い魚種の色彩

(a) マアジ

Measuring points	L^*	a^*	b^*	Colors
L1-C1	28.54	-1.00	1.81	dark gray
L1-C2	29.32	-1.19	2.19	dark gray
L2-C0	80.43	-0.12	8.27	yellowish gray
L2-C1	67.76	-0.86	10.93	grayish yellow
L2-C2	69.54	-1.83	9.48	yellowish gray
L2-C3	50.03	-3.40	14.35	grayish olive
L3-C1	86.01	-0.52	-0.05	white
L3-C2	85.16	-1.11	0.15	white
E	32.46	0.61	-0.31	dark gray

(b) ウマヅラハギ

Measuring points	L^*	a^*	b^*	Colors
L1-C1	28.12	-0.90	1.44	dark gray
L1-C2	28.87	-1.01	2.31	dark gray
L2-C0	36.48	-0.63	4.12	dark olive gray
L2-C1	44.36	-0.71	4.65	olive gray
L2-C2	44.48	-1.82	4.68	olive gray
L2-C3	37.31	-1.28	3.05	dark gray
L3-C1	57.51	-0.76	6.90	olive gray
L3-C2	55.92	0.23	6.52	brownish gray
E	35.51	-1.33	6.90	dark olive gray

(c) サワラ

Measuring points	L^*	a^*	b^*	Colors
L1-C1	35.58	-2.04	-2.07	dark greenish gray
L1-C2	40.32	-1.29	-4.37	bluish gray
L2-C0	77.99	-0.05	-0.58	light gray
L2-C1	66.99	-0.38	2.77	medium gray
L2-C2	73.15	-0.45	-1.03	light gray
L2-C3	66.43	-3.40	-2.63	bluish gray
L3-C1	86.54	-0.78	-0.48	white
L3-C2	93.46	-0.58	-1.01	white
E	43.66	-0.34	4.23	olive gray

と色相角 h との間にも正の相関 ($r = 0.589$, $P < 0.01$) が確認された. 明度 L^* と彩度 C^*ab との間および明度 L^* と色相角 h との間には関連がみられなかった.

表 5 にふぐ処理師の評価について判別分析を行った後, 結果についてのふぐ処理師との協議を経て最も適切と判断した判別分析モデルの構成を示す. モデルの変数には魚市場の競り人による鮮魚の外観評価に共通して表れた体幹部内(測定線 $L2$ 上)の色彩[3]は入らない結果となった. モデルの判別率的中率は 74 % であった. ふぐ処理師の評価とモデルの判別結果とが異なったものは, ふぐ処理師が第 3 群と評価したものをモデルで第 1 群へと判別した 2 尾を除いて, 総て隣接する群への判別であった. 表 5 よりモデルは体表の計 4 点の測定点の色彩で構成されており, 判別には特に腹部の頭部寄り ($L3-C1$) の色彩が貢献する様子がわかる. 腹部 ($L3-C1$) は身欠き処理において腹部の粘膜と内臓除去のための切り口付近にあたる部位である.

3.2 評価と魚肉鮮度の関係

(a) 競り人の評価と魚肉鮮度

図 2 に 2 時間経過時における各魚種の群別の魚肉鮮度 K 値を示す. 各魚種とも K 値は総ての群で生食の目安となる 20 % 以下となっており, 極

めて鮮度の良い状態で評価されていることがわかる. また, 図 2 に示すとおり, 各魚種とも K 値はほぼ良好に群順位を反映した.

表 3 に示す判別分析モデルを構成する変量と K 値との関係は, マアジでは腹部 ($L3-C2$) の色度 b^* が K 値と 72 時間経過時迄継続して強い正の相関 ($0.681 \leq r \leq 0.705$, $P < 0.01$) を示した. 特に評価の高い群 (第 1 群, 第 2 群) においてはかなり強い関連 ($0.800 \leq r$, $P < 0.01$) が確認された. 同様に, 背部一体幹部間 $\{(L1-C1) - (L2-C1)\}$ の明度差 ΔL^* と K 値は全群で正の相関 ($0.363 \leq r \leq 0.505$, $P < 0.01$) が, また体幹部-腹部間 $\{(L2-C2) - (L3-C2)\}$ の彩度差 ΔC^*ab は評価の高い群において負の相関 ($-0.610 \leq r \leq -0.433$, $P < 0.01$) を示した. ウマヅラハギでは評価の高い群の体幹部 ($L2-C1$) の色度 a^* が K 値と 48 時間経過時迄負の相関 ($-0.357 \leq r \leq -0.327$, $P < 0.01$) を, また, 評価の低い群 (第 4 群, 第 5 群) の体幹部 ($L2-C3$) の明度 L^* が 72 時間経過時迄正の相関 ($0.415 \leq r \leq 0.685$, $P < 0.01$) を示した. 加えて, 第 1 群から第 3 群迄の眼 (E) の色度 b^* が 72 時間経過時迄負の相関 ($-0.567 \leq r \leq -0.343$, $P < 0.01$) を示した. サワラでは全群の体幹部-腹部間 $\{(L2-C2) - (L3-C2)\}$

表 3 彩度の低い魚種の判別分析モデル

(a) マアジ

Measuring points	Indices
Abdominal region ($L3-C2$)	b^*
Trunk part ($L2-C1$)-($L2-C2$)	ΔL^*
Dorsal-trunk part ($L1-C1$)-($L2-C1$)	ΔL^*
Dorsal-abdominal region ($L1-C2$)-($L3-C2$)	ΔC^*ab

(b) ウマヅラハギ

Measuring points	Indices
Trunk part ($L2-C1$)	a^*
Trunk part ($L2-C3$)	L^*
Eye (E)	b^*
Dorsal-abdominal region ($L1-C2$)-($L3-C2$)	ΔC^*ab

(c) サワラ

Measuring points	Indices
Trunk part ($L2-C1$)	L^*
Abdominal region ($L3-C1$)	b^*
Trunk part ($L2-C2$)-($L2-C3$)	Δa^*
Trunk-abdominal region ($L2-C2$)-($L3-C2$)	Δa^*

表 4 養殖トラフグ身欠きの色彩

Measuring points	L^*	a^*	b^*	Colors
$L1-C1$	52.664	-0.59	-2.28	bluish gray
$L1-C2$	51.29	-1.06	-4.01	bluish gray
$L2-C1$	54.68	-0.33	-0.28	medium gray
$L2-C2$	49.80	-0.79	-2.11	medium gray
$L2-C3$	52.49	-0.44	-4.65	bluish gray
$L2-C4$	27.61	0.15	-1.32	dark gray
$L3-C1$	51.851	-1.13	-3.04	bluish gray
$L3-C2$	53.42	-0.63	-2.62	bluish gray
ac $L3-C1$	47.20	-0.97	-3.60	bluish gray
ac $L3-C2$	51.06	4.14	-1.47	reddish gray

表 5 養殖トラフグの判別分析モデル

Measuring points	Indices
Dorsal region ($L1-C2$)	L^*
Caudal fin ($L2-C4$)	a^*
Abdominal region ($L3-C1$)	a^*
Abdominal region ($L3-C1$)-($L3-C2$)	Δb^*

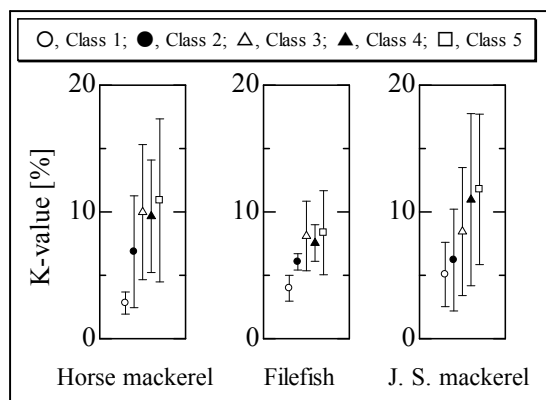


図2 彩度の低い魚種の K 値 (2 時間経過時)

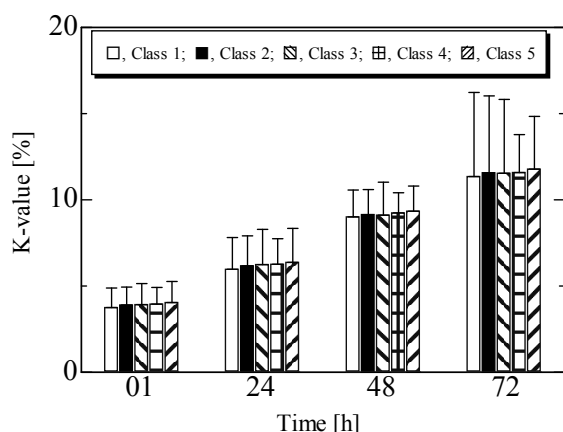


図3 養殖トラフグの K 値の経時変化

の色度差 Δa^* が K 値と 48 時間経過時迄負の相関 ($-0.359 \leq r \leq -0.270$, $P < 0.01$) を示した。

(b) ふぐ処理師の評価と魚肉鮮度

図3に魚肉鮮度 K 値の経時変化の様子を示す。1 時間経過時における K 値は全試料で 1.32 % から 6.74 % 迄の非常に狭い範囲に分布した。各群の K 値の平均は第 1 群から第 5 群迄で 3.75 % から 4.04 % の範囲にあり、評価の低下に伴ってやや上昇する様子を呈した。二元配置の分散分析では等分散性の検定で測定時間が有意 ($P < 0.01$) となり、等分散性が棄却された。クラスカル・ウォリス検定では測定時間の全水準間の比較で有意 ($P < 0.01$) となったが、各群間の比較では有意差は確認されなかった。なお、二元配置の分散分析でも多重比較の検定では全測定時間の比較において有意差 ($P < 0.01$) が確認された。また、評価群の各水準における測定時間の単純主効果の検定では各群とも有意 ($P < 0.01$) となり、多重比較でも各評価群における全測定時間の比較が有意 ($P < 0.01$) となった。

4. 考察

4.1 目利きの評価

4.1.1 競り人の評価

以上の結果より、低彩度魚種では体表の彩度 C^*ab の低下に伴って色彩と K 値との関連は低くなった。しかし、図2に示すとおり、最も体表の彩度 C^*ab の低いサワラでも競り人の外観評価には K 値が良好に反映した。このことから、競り人は彩度 C^*ab の低い魚種でも体表の色彩に表れる僅かな鮮度劣化の特徴を参考にして高い確度で品質を見定めているものと言える。

表3より、競り人の低彩度魚種の外観評価には体幹部を通る測定線 L2 上の色彩又は色彩の差と、魚体の体中央に位置して背部から腹部迄を通る測定線 C2 上の腹部 (L3-C2) との色彩の差が共通して反映することが示された。競り人の評価に体幹部内の色彩と色彩の差が反映する理由については前報まで[3,17]で述べた。よって本報では、先ず競り人の低彩度魚種の評価に背部から腹部迄を通る測定線 C2 上で腹部 (L3-C2) との色彩の差が共通して反映することについて考察する。次に、競り人の評価のポイントが体表の彩度 C^*ab の低下に伴って頭部寄り (測定線 C1) から体中央部 (測定線 C2) へと移る理由を明らかにする。

マアジの 2 時間経過時における背部 (L1-C2) の色彩は全尾の平均で明度 $L^* = 29.32$, 色度 $a^* = -1.19$, 色度 $b^* = 2.19$, 彩度 $C^*ab = 2.75$ となり全群が暗い灰色を呈する。背部 (L1-C2) では第 1 群の彩度 C^*ab で若干の経時的な低下が見られたものの、明度 L^* , 色度 a^* , b^* , 彩度 C^*ab の何れにも群間差は確認されなかった。しかし、第 1 群と他の群との色差 ΔE^*ab は 1.42 から 2.44 の範囲で群が離れる程拡大し、第 1 群と他の群とを並べて比較すると最低でも僅かに異なるものとなる。腹部 (L3-C2) の色彩は全尾の平均で明度 $L^* = 85.16$, 色度 $a^* = -1.11$, 色度 $b^* = 0.15$, 彩度 $C^*ab = 3.04$ となり、第 1 群が青みのうすい灰色を、また他の群がうすい灰色を呈している。腹部 (L3-C2) の彩度 C^*ab は 2.01 から 4.07 の範囲にあり、第 1 群と第 2 群及び第 3 群との間で、また第 2 群、第 3 群と評価の低い群 (第 4 群、第 5 群) との間で有意差 ($P < 0.01$) が確認された。腹部

(L3-C2) の群間の色差 ΔE^*ab は第 1 群と他の群との間で 2.14 から 3.06 の範囲にあり、並べて比較すると最低でも感知し得るほどに異なるものとなる。以上のマアジの背部 (L1-C2) 並びに腹部 (L3-C2) の色差 ΔE^*ab には色度の差、言い換えると彩度 C^*ab の差が貢献しており、これらの呈色の違いは主に漁獲後の経過時間や取扱いの良否に起因するものである。よって競り人は、比較的退色の遅い背部 (L1-C2) を基準として腹部 (L3-C2) の色彩を見ることで品質を見積もるものと推察する。

図 4 にウマヅラハギの背部 (L1-C2) と腹部 (L3-C2) の彩度 C^*ab の経時変化の様子を示す。図 3 より、両点間の彩度差 ΔC^*ab は主に腹部 (L3-C2) の彩度 C^*ab に影響を受けることがわかる。ウマヅラハギの 2 時間経過時における背部 (L1-C2) の色彩は、全尾の平均で明度 $L^*=28.87$ 、色度 $a^*=-1.01$ 、色度 $b^*=2.31$ となり暗い灰色を呈する。背部 (L1-C2) の色彩の群間の比較では、第 5 群と他の群との間の明度 L^* 、色度 b^* 及び彩度 C^*ab のみ有意差 ($P < 0.01$) が確認された。第 5 群と他の群との間の色差 ΔE^*ab は 2.38 から 4.43 の範囲にあり、並べて比較すると最低でも感知し得るほどに異なるものとなる。しかし、24 時間経過時以降には群間の有意差は確認されなかった。

よってウマヅラハギの背部 (L1-C2) は漁獲後の初期の段階における品質劣化の特徴を捉え易い箇所と言える。一方、腹部 (L3-C2) の色彩は全尾の平均で明度 $L^*=55.92$ 、色度 $a^*=0.23$ 、色度 $b^*=6.52$ となり、黄みの灰色を呈する。腹部 (L3-C2) では多くの群間で色度 a^* 、色度 b^* 及び彩度 C^*ab の有意差 ($P < 0.01$) が確認され、かつこれらの色彩は第 2 群を除いて良好に群順位を反映した。また、群間の色差 ΔE^*ab は 1.09 から 5.63 の範囲にあり、第 1 群と第 5 群との間で最大となった。よって、各群を並べて比較すると最低でも僅かに異なるものとなる。背部 (L1-C2) と腹部 (L3-C2) との明度差 ΔL^* は平均で -27.03 であり、群間に有意差は確認されなかった。以上により、両点の色差 ΔE^*ab には主に色度差 Δa^* と Δb^* の大きさ、言い換えると彩度差 ΔC^*ab が強く反映するものと見なすことができる。よって、競り人は主に漁獲後の初期の段階で品質を表し易い背部 (L1-C2) と、

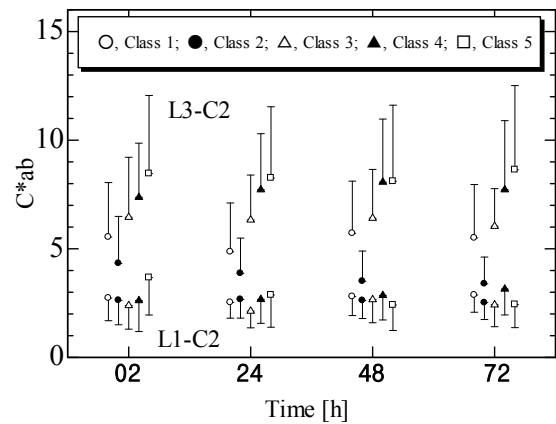


図 4 ウマヅラハギの背部内測定点 (L1-C2) と腹部内測定点 (L3-C2) の彩度 C^*ab の経時変化

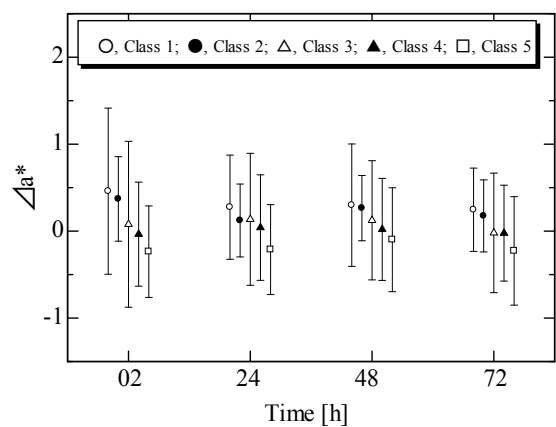


図 5 サワラの体幹部内測定点 (L2-C2) と腹部内測定点 (L3-C2) の色度差 Δa^* の経時変化

良好に品質を反映する腹部 (L3-C2) の彩度 C^*ab を捉えて品質を見積もるものと推察する。

図 5 にサワラの体幹部 (L2-C2) と腹部 (L3-C2) の色度差 Δa^* の経時変化の様子を示す。図 4 に示すとおり体幹部 (L2-C2) と腹部 (L3-C2) の色度差 Δa^* は極めて良好に群順位を反映して小さくなり、経時的にも低下する傾向を示した。両点間の色彩では色度差 Δa^* のみで評価の高い群 (第 1 群、第 2 群) と評価の低い群 (第 4 群、第 5 群) との間に有意差 ($P < 0.01$) が確認された。2 時間経過時における体幹部 (L2-C2) の色彩は全尾の平均で明度 $L^*=73.15$ 、色度 $a^*=-0.45$ 、色度 $b^*=-1.03$ 、彩度 $C^*ab=2.06$ となり、第 1 群から第 4 群迄が明るい灰色を、また第 5 群が青みの明るい灰色を呈した。体幹部 (L2-C2) の色彩は、色度 a^* のみが 24 時間経過時迄評価の高い群 (第 1 群、第 2 群) と第 5 群との間で有意差 ($P < 0.01$)

が確認された。群間の色差 ΔE^*ab には色度差 Δa^* が貢献して 0.64 から 2.52 の範囲にあり、並べて比較すると最低でも僅かに異なるものとなる。2 時間経過時における腹部 (L3-C2) の色彩は全尾の平均で明度 $L^* = 93.46$ 、色度 $a^* = -0.58$ 、色度 $b^* = -1.01$ 、彩度 $C^*ab = 1.53$ となり、第 1 群と第 2 群が白を、また第 3 群から第 5 群迄が青みの白を呈した。群間における色差 ΔE^*ab は 0.33 から 1.14 の範囲にあり、第 1 群と第 5 群との比較以外には色彩に大きな違いは表れなかった。以上のことより、サワラは品質劣化が進行すると幹部 (L2-C2) と腹部 (L3-C2) の色彩が青みを帯びる傾向を有している。また、その呈色は体幹部 (L2-C2) の方が腹部 (L3-C2) のものより迅速であり、相対的な変化の様子が両点間の色度差 Δa^* に表れる。競り人は経験的にこれらの部位間の経時的な色彩の変動と品質劣化の関係を捉えており、それを基に品質を見積もるものと推察する。

4.1.2 ふぐ処理師の評価

1 時間経過時においては、身欠きの魚肉鮮度 K 値は全試料の最大値を 6.74 % とする非常に狭い範囲に分布し、かつ平均値も 3.91% と極めて低い値を得た。身欠きの K 値は一般の即殺魚の目安の 5 % [18] や、品質の良さで定評のある山口県萩産のアカアマダイ等の鮮魚 [17] のものよりも低い値となった。これは試料の入手先である南風泊市場では養殖物のトラフグを活きた状態で競りに出すこと、競り落とした仲買（ふぐ処理師）が活きた状態で作業場に持ち帰って即座に身欠き処理をすること、1 尾あたりの処理時間が 3 分程度と非常に短いこと、身欠き処理直後に水洗して出荷まで低温で冷蔵保管すること等、活きたトラフグを扱い、迅速に内臓除去を伴う身欠き処理をして冷蔵することで筋肉酵素や細菌酵素による分解を進み難くしていることに起因する。即ち、下関のふぐ処理師の技は、地の利を活かし、長年の作業で培われた市場と仲買の双方のフグ流通における一貫した品質保持の態勢に由縁するものである。1 時間経過時では各群の K 値は評価の低下に伴ってやや上昇する様子を示したが、何れの群間にも母平均の有意差は確認されなかった。また、5 水準の評価群と 4 水準の経過時間を因子とする二元配置の分散分析の結果、 K 値に対する主効果が経

過時間であることが示された。よって、1 時間経過時における K 値の違いは、主として身欠き処理から K 値分析のための試料採取迄に要する時間の差に起因するものである。また、分析対象が活きた養殖トラフグで、かつ極めて厳しく品質管理された環境下では、 K 値が低く群間差も僅かとなるためふぐ処理師の品質評価には魚肉鮮度は反映しないものとする。

魚肉のかたさは 1 時間経過時に 37.5MPa から 118.6MPa 迄の広い範囲に分布し、何れの群間にも有意差は確認されなかった。魚介類の死後変化は、死後に徐々に魚肉の硬直が始まって完全硬直の状態に至り、完全硬直が一定時間続いた後に解硬が始まり軟化していく。これらの変化の速度は魚種により異なるが、本実験で確認したところ、触手により硬直を確認できたのはおよそ 4 時間から 5 時間経過時以降であった。また、養殖トラフグの生（切り身）の 100g 中における標準的な成分 [19] は水分とタンパク質で 98.2g を占めており、イサキ [3] やマアジ [20] 等の一般的な魚種と比較すると水分含量が高い。よって、ふぐ処理師が試料を品質評価した際には未だ触感で判る程度には硬直が進んでおらず、また魚肉の成分構成によりかなり柔らかい印象を受けることを理由として品質評価には魚肉のかたさが反映しないものと推察する。

以上のことから、ふぐ処理師の品質評価には体表の色彩が貢献するものと推察する。トラフグの身欠きは「身が透明で、絹のような光沢があるものが上質」 [21] とされる。これは身欠き全体から受ける印象について述べているものであり、対象は天然で鮮度の良い未凍結品の品質を指すものと推察する。しかし、天然のトラフグは出荷量としてはごく僅かのものである。本研究では流通の多い活きた養殖トラフグを身欠きにした未凍結品を対象としている。体表の色彩は灰色または青みの灰色を呈しており、光沢については体表全体の輝度 L と明度 L^* が共に 72 時間経過時迄継続して上昇する傾向を示す [20]。1 時間経過時では背部の頭部寄り (L1-C1) と体幹部の 2 点 {(L2-C1), (L2-C2)} が灰色を、背部の中央 (L1-C2)、体幹部の尾部寄り (L2-C3) および腹部 2 点と腹腔の頭部寄り (acL3-C1) が共に青みの灰色を呈し

た．また、腹腔の中央（acL3-C2）は赤紫みの灰色、尾鰭は暗い灰色であった．先に述べたように、養殖トラフグの生（切り身）の場合は魚肉成分のほとんどを水分とタンパク質で占めており、脂質は0.3gで他の魚種よりも極めて低い特徴がある．よって、身欠きでは脂質の酸化に伴う油焼けなどの現象は起こり難いことから、尾鰭を除いた体表の色彩は水分とタンパク質で成すものと言える．また、ふぐ処理師の品質評価には体幹部を除く計4箇所を表す色彩の組合せが最も対応するものとした．表6に1時間経過時におけるこれらの色彩と群間における色差 ΔE^*ab の最大値を示す．以下、ふぐ処理師の品質評価にこれらの色彩が反映する理由について考察する．

(a) 背部(L1-C2)の色彩

1時間経過時の背部中央（L1-C2）の色彩は全群とも青みの灰色を呈する．各群間の色差 ΔE^*ab は0.26から2.95の範囲にあり、評価の低い群（第4群、第5群）と他の群との間で拡大する傾向を示した．第1群と第5群との色差 ΔE^*ab が群間における最大値となり、両群の試料魚を並べて比較すると感知し得るほどに異なるものとなる．第1群の色彩は全試料の平均で明度 $L^*=53.76$ 、色度 $a^*=-1.11$ 、色度 $b^*=-3.50$ であり、第5群では明度 $L^*=50.90$ 、色度 $a^*=-0.91$ 、色度 $b^*=-4.17$ であった．よって両点の色差 ΔE^*ab には明度の差 ΔL^* が最も貢献する．背部中央（L1-C2）は背鰭に近い位置にある．よって身欠き処理において皮を剥ぐ際には表皮の内側にある真皮（薄皮）が僅かに残り易い箇所であり、薄皮が残る場合には明度 L^* は低下する．背部中央（L1-C2）の明度 L^* と魚肉鮮度 K 値との間には何れの群にも関連が確認されなかった．

以上から、ふぐ処理師は背部中央の真皮の残りの有無とその程度、即ち身欠き処理の丁寧さを評価しているものと推察する．

(b) 尾鰭(L2-C4)の色彩

1時間経過時の尾鰭（L2-C4）の色彩は全群とも暗い灰色を呈する．各群間の色差 ΔE^*ab は0.55から1.99の範囲にあり、評価の高い群（第1群、第2群）と評価の低い群との間での比較で大きくなる．各群の尾鰭（L2-C4）の色彩は平均で明度 L^* が26.26から27.96、色度 a^* が-0.22から

1.10、色度 b^* が-1.76から-0.81の範囲にあり、色差 ΔE^*ab には明度の差 ΔL^* と色度の差 Δa^* が共に良く貢献する．しかしこれらの明度 L^* と色度 b^* には何れの群間にも母平均に有意差が確認されず、また群順位への規則的な対応も見られなかった．尾鰭を含む体表の表皮には粘液細胞があり、粘性を保つことで体表の防御等を行う．粘液の成分はムチンと総称される糖タンパク質、無機塩類、糖類等から成るが、空気に触れると分子量の大きなタンパク質の変成で白化が進行する．即ち、尾鰭では粘性の強さ（糖タンパク質の含有量）の程度が明度 L^* に表れるものであり、個体差が大きいことから特に品質に対しての規則性を有するものではないと考える．

図6に尾鰭（L2-C4）の色度 a^* の経時変化の様子を示す．図6より色度 a^* は評価の高い群が高く、評価の高い群と評価の低い群との間では72時間経過時迄継続して母平均に有意差（ $P < 0.01$ ）が確認された．また、評価の高い群の色度 a^* は徐々に低下することから、経時的に評価の低い群との差が縮小される傾向を示している．トラフグの尾鰭は先端を除いて血管を配していることから、尾鰭の色彩は身欠き処理後の血液の経時的変動を表し易い．即ち、尾鰭の色彩は身欠き処理後の時間が経たない段階では赤みが強く、経時的に醒めていく．よって、尾鰭の色彩は身欠き処理後の経過時間、即ち鮮度の程度を反映し得る一つの指標となる．

以上から、ふぐ処理師は尾鰭の色彩の程度から鮮度の良さを見積もるものと推察する．

(c) 腹部(L3-C1)の色彩と腹部内{(L3-C1)-(L3-C2)}の色彩の差

1時間経過時の全試料の腹部の彩度 C^*ab の平均は3.23となり、他の部位と比較すると有意（ $P < 0.01$ ）に高い．腹部の頭部寄り（L3-C1）の色彩は明度 L^* が48.97から50.85、色度 a^* が-1.31から-0.92、色度 b^* が-3.59から-2.98の範囲にある．また、腹部の中央（L3-C2）の色彩は明度 L^* が53.72から58.16、色度 a^* が-1.31から-1.11、色度 b^* が-3.02から-2.11の範囲に分布する．色彩は両点とも青みの灰色を呈している．1時間経過時における腹部の頭部寄り（L3-C1）の群間の色差 ΔE^*ab の最大は1.95であり、第1群と第5

表6 ふぐ処理師の評価の判別分析モデルを構成する色彩の群間における最大色差 ΔE^*ab

Measuring points	Indices
Dorsal region (L1-C2)	2.95
Caudal fin (L2-C4)	1.99
Abdominal region (L3-C1)	1.95
Abdominal region (L3-C1) - (L3-C2)	4.46

群との間の比較であった。また、各測定時における群間の色彩の比較では、色度 a^* のみが 72 時間経過時迄継続して評価の高い群と低い群との間で平均の有意差を示し、かつ第 1 群の色度 a^* が最小となった。身欠き処理後の腹部は他の部位と比べて肉厚が薄く、特に腹部の頭部寄り (L3-C1) が最も薄くなる。

図 7 に腹部内 {(L3-C1) - (L3-C2)} の色度の差 Δb^* の経時変化の様子を示す。色度の差 Δb^* は 1 時間経過時にはほぼ群順位を反映するが、評価の高い群が経時的に徐々に低下することから群間差が縮小する傾向が見受けられる。1 時間経過時における腹部の頭部寄り (L3-C1) の色度 b^* は第 1 群が最大で、群順位の低下に伴って低くなる傾向が見られた。一方、腹部の中央 (L3-C2) の色度 b^* では群間における平均の差異は確認されなかった。よって、群間における腹部内 {(L3-C1) - (L3-C2)} の色度の差 Δb^* には腹部の頭部寄り (L3-C1) の色度 b^* が強く反映するものと言える。

以上の腹部の頭部寄り (L3-C1) の色彩の差は前述したとおり身欠き処理後の肉厚に起因するもので、この部位は身欠き処理前のトラフグの放置状態、身欠き処理後の水洗等の丁寧さおよび乾燥の進行に伴う魚肉の変性の程度等が最も身色に反映し易くなる箇所と言える。よって、ふぐ処理師は腹部のこれらの色彩を確認することで品質を見定めるものとする。

5. 目利きの技の所以

(a) 競り人の評価

図 8 に競り人による有彩色魚種（彩度の高い魚種）と低彩度魚種（彩度の低い魚種）の魚体上の評価のポイントをまとめる。両魚種とも体表上の 4 個の指標のうち体幹部（測定線 L2 上）の色彩、又は色彩の差と腹部（測定線 L3 上）の色彩が競り人の評価に反映したが、

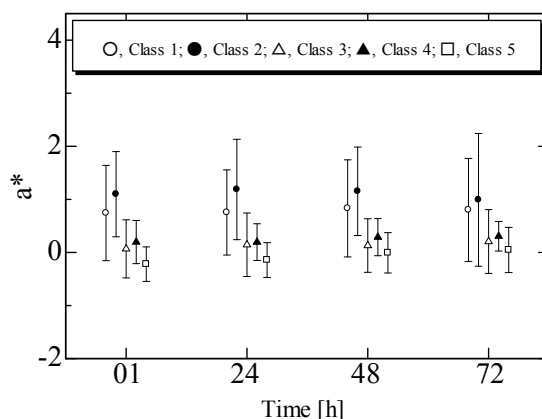


図 6 尾鰭 (L2-C4) の色度 a^* の経時変化

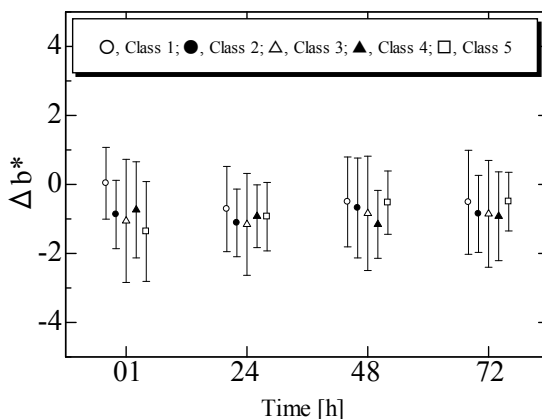


図 7 腹部内 {(L3-C1) - (L3-C2)} の色度差 Δb^* の経時変化

有彩色魚種で確認した頭部寄り（測定線 C1 上）の指標が低彩度魚種では彩度の低下に伴って体中央部（測定線 C2 上）へと移動した。

有彩色魚種と低彩度魚種の部位毎の色彩を比較した場合、明度 L^* の分布は両魚種共に背部（測定線 L1 上）、体幹部（測定線 L2 上）、腹部（測定線 L3 上）の順に高くなり、かつ背部と腹部との明度差 ΔL^* はほぼ同程度の大きさとなる。よって、体表における部位上の明度 L^* の分布の差異は小さいものと言える。次に部位毎の彩度 C^*ab の分布は、有彩色魚種では腹部、背部、体幹部の順に高くなり、腹部と体幹部との彩度差 ΔC^*ab は 10 程度となる。一方、低彩度魚種では背部、腹部、体幹部の順に高くなり、背部と体幹部との彩度差 ΔC^*ab は 3 程度と小さくなり、有彩色魚種と比較すると退色による品質劣化の確認は難しくなる。しかし、両魚種とも相対的に体幹部の色彩が最も高くなることから、体幹部（測定線 L2 上）の色彩、又は色彩の差で品質の程度を

見積もり易くなると考える。

有彩色魚種と低彩度魚種の頭部寄り（測定線 C1 上）と体央（測定線 C2 上）の色彩を比較した場合、有彩色魚種の明度 L^* は頭部寄りと体央共に若干低彩度魚種より高くなる。また、有彩色魚種では頭部寄りと体央の色彩に大きな差異はみられず、共に品質の程度を表し易い。一方、低彩度魚種では体表全体の彩度 C^*ab の低下に伴い、体央の明度 L^* 、彩度 C^*ab 共に頭部寄りより高くなる傾向が強まり、外観評価には体央の色彩が大きく貢献するものとなる。よって、低彩度魚種のなかで最も体表の彩度 C^*ab の高いマアジでは頭部寄りと体央の両方の指標が評価に反映し、マアジより彩度 C^*ab の低いウマヅラハギとサワラでは体央のみの色彩の差が反映するものと考えられる。

以上のことより、競り人の視点は魚体体表の色彩の分布に良く対応しており、魚種毎に品質に関する知識を組み立てている。よって、競り人は対象とする鮮魚の体表の色彩とその経時的変化の特徴を経験的に上手く把握しており、精度良く対象の品質を見積もっているものと言える。

(b) ふぐ処理師の評価

ふぐ処理師の品質評価には魚肉鮮度と身欠き処理の丁寧さが反映すること、またこれらが表れる計4個の色彩を前件部変数に定めてモデルを構成することで、良好な推定結果が得られることを示した。これは、身欠きはフグ肉の流通で最も消費者に近いものになることから、鮮度は当然のこととして、身欠き処理の丁寧さがふぐ処理師の技の高さや仲卸の信用と質の良さを示すことに由来するものと考えられる。

6. 目利きの技のファジィ推論モデル

(a) 競り人の評価モデル

表7にファジィ推論を用いた競り人の鮮魚に対する外観評価モデルの構成とモデルの評価結果を示す。表7中の的中率(hitting ratio)とは競り人の外観評価とモデルによる推論結果が一致した割合を示す。表7から明らかなように、有彩色魚種の的中率が低彩度魚種のものよりも高くなる傾向が確認された。これは第5章で述べたとおり、有彩色魚種では品質劣化の特徴が体表の色彩に表れ易

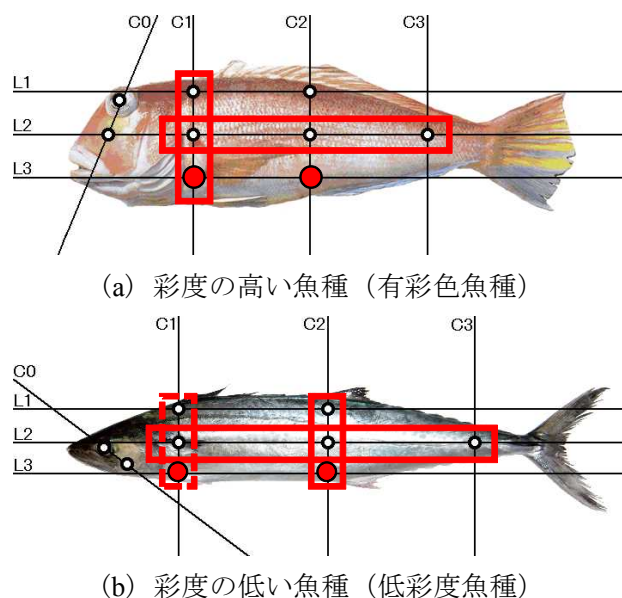


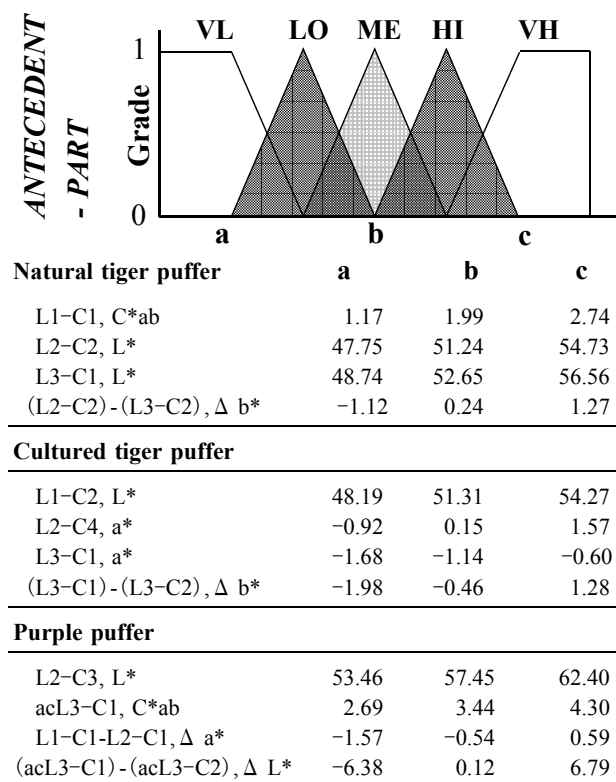
図8 競り人の評価のポイント

表7 競り人の外観評価モデルの構成と評価結果

Species	Antecedent variables	Hitting rates
Tilefish	L2-C3, L L3-C2, C^*ab (L2-C1)-(L2-C3), ΔC^*ab (L1-C1)-(L2-C1), ΔC^*ab	94.4% /100%
White chestnut	L1-C0, b^* L3-C2, L^* (L2-C1)-(L2-C3), Δb^* (L1-C1)-(L3-C1), Δb^*	95.1% /100%
Striped grunt	L3-C1, L^* E, L^* (L2-C1)-(L2-C3), ΔC^*ab (L1-C1)-(L2-C1), ΔE^*ab	96.0% /90.0%
Horse mackerel	L3-C2, b^* (L2-C1)-(L2-C2), ΔL^* (L1-C1)-(L2-C1), ΔL^* (L1-C2)-(L3-C2), ΔC^*ab	92.3% /89.3%
Filefish	L2-C1, a^* L2-C3, L^* E, b^* (L1-C2)-(L3-C2), ΔC^*ab	87.5% /84.6%
J.S. mackerel	L2-C1, L^* L3-C1, b^* (L2-C2)-(L2-C3), Δa^* (L2-C2)-(L3-C2), Δa^*	88.6% /90.7%

Hitting rate: Simulation / Experimental evaluation

いことに起因する。競り人の外観評価とモデルによる推論結果が不一致であったものの多くは、競り人による評価群の隣群に推定したものであつ



(a) 前件部

(b) 後件部

図9 身欠きの品質評価モデルの構成

た。これらについてはメンバーシップ関数のチューニングを丁寧にする必要がある。しかし、全般的に一致率が良好であることからモデルの有用性は高いものとする。

(b) ふぐ処理師の評価モデル[7]

図9にチューニング後のファジィ推論モデルのメンバーシップ関数の構成を示す。モデルの前件部変数は各身欠き共に計4個の色彩から成り、後件部変数は共通である。ファジィ規則の数は、天然トラフグが150個、養殖トラフグとマフグが各250個とした。

表8にモデルに対して全試料を用いて行ったシミュレーションと評価実験の結果を示す。天然トラフグではふぐ処理師の評価とモデルによる推定結果が一致したものが80尾中77尾で、一致率は

表8 ふぐ処理師の品質評価モデルの評価結果

(a) 天然トラフグ

		Estimation by fuzzy model				
		Cl. 1	Cl. 2	Cl. 3	Cl. 4	Cl. 5
Cook's evaluation	Cl. 1	16	0	0	0	0
	Cl. 2	0	14	2	0	0
	Cl. 3	0	0	16	0	0
	Cl. 4	0	0	1	15	0
	Cl. 5	0	0	0	0	16

Hitting rate = 96.3%

(b) 養殖トラフグ

		Estimation by fuzzy model				
		Cl. 1	Cl. 2	Cl. 3	Cl. 4	Cl. 5
Cook's evaluation	Cl. 1	45	1	4	0	0
	Cl. 2	5	38	7	0	0
	Cl. 3	2	4	43	1	1
	Cl. 4	0	3	7	39	0
	Cl. 5	0	0	9	4	37

Hitting rate = 80.8%

(c) マフグ

		Estimation by fuzzy model				
		Cl. 1	Cl. 2	Cl. 3	Cl. 4	Cl. 5
Cook's evaluation	Cl. 1	42	2	1	0	0
	Cl. 2	4	42	1	0	0
	Cl. 3	3	2	43	0	0
	Cl. 4	0	2	4	36	4
	Cl. 5	0	0	1	1	42

Hitting rate = 89.1%

96.3%と極めて高かった。一致しなかった3尾は全て隣群への推定となった。また、養殖トラフグでは250尾中202尾、マフグでは230尾中205尾が一致し、一致率はそれぞれ80.8%、89.1%となった。

養殖トラフグの一致率は他のものと比較して低い。これは養殖トラフグの身欠きの1時間経過時の各群のK値の平均が3.75%から4.04%の非常に狭い範囲にあり、他と比較すると鮮度の違いを見積もることが容易でなかったこと、よってふぐ処理師は養殖トラフグの品質を身欠き処理の丁寧さと鮮度の両面から評価をする傾向が強まることに起因するものとする[6,7,22]。また、養殖トラフグで不一致となった内容は、隣群に推定したものが29尾、2群離れたものが19尾であった。これらの計48尾の内、モデルがふぐ処理師の評価よりも高い品質であると推定したものが34尾、

逆に低い品質と推定したものが 14 尾であった。このようにふぐ処理師の評価がモデルの推定結果よりも厳しくなる理由としては、体表に定めた測定点以外の箇所でも評価を下げる要因、例えば他のフグから噛まれて鬱血した痕が見られたことなどが考えられる。本研究で試料とした養殖トラフグは、前述のとおり競りの直前まで市場生け簀に活かされていたものであり、市場における競り終了後にふぐ処理師により即座に身欠き処理をされたことから、1 時間経過時における K 値の各群の平均が 3.75 % (第 1 群) から 4.04 % (第 5 群) と非常に狭い範囲にあった。この養殖トラフグの K 値は、即殺魚の目安である 5 % 以内に入っており、例えば第 5 群の養殖トラフグであっても大変新鮮なものであったと言える。よって、この段階で品質を高く評価することのリスクは、思ったほど大きくないと考える。マフグでは隣群と 2 群離れたものがそれぞれ 18 尾と 7 尾であり、養殖トラフグの場合と同様に、ふぐ処理師の評価がモデルの推定結果よりも厳しくなった。

以上、特に養殖トラフグのモデルでは推定精度を改善するために更に丁寧なチューニングをすることを要するものの、天然トラフグとマフグではほぼ満足できる一致率を得た。よって、本モデルの有用性は高いものと考えられる。

7. 結言

「新鮮で安全な魚を美味しく食べたい」というのは魚食文化が普及している日本人に限らず、世界中の人々の願いだろう。本稿ではそれを叶える一方法として、先ず、所謂プロの目利きである魚市場の競り人の技に着目し、モデル化を試みた結果を紹介した。鮮魚に対する競り人の品質評価は、計 4 個程度の体表の色彩を変数としてモデル化が可能であることを示したが、イカ等の頭足類についても同様の結果[4]を得ている。

次に、本稿では下関市の老舗ふぐ仲卸のふぐ処理師による天然トラフグ、養殖のトラフグおよびマフグの身欠きについての品質評価を解析すると共に、ふぐ処理師の品質評価に相当するファジィ推論モデルの作成を試みた。その結果、ふぐ処理師の品質評価には魚肉鮮度と身欠き処理の丁寧さが反映すること、またこれらが表れる計 4 個の色

彩を前件部変数に定めてモデルを構成することで、良好な推定結果が得られることを示した。

これらの品質に差がある多様な魚類やフグ類を長年扱うことで、確かなプロの目も養われ、それを解析してモデル化することで、技術の維持と継承の一助となると考える。

現在、著者らはフグ類の身欠きを対象として、天然と養殖の鑑定や魚種の鑑別について取り組んでいるところである。

本研究にご協力を頂いた山口県漁業協同組合山口県漁協萩地方卸売市場(萩市)と株式会社酒井商店(下関市)の皆様には感謝いたします。本研究の一部は JSPS 科研費(21500232, 25330302)の助成を受けたものです。

参考文献

- [1] 渡邊悦生編著, 魚介類の鮮度と加工・貯蔵, pp.1-27, 成山堂書店, 1998.
- [2] M. Nakamura, H. Ohta, E. Morimoto, S. Ezoe, K. Harada, Quality Estimation Model of White Chestnut (*Lagocephalus wheeleri*) by Auctioneers, Fisheries Engineering, Vol.45, pp.119-128, 2008.
- [3] 中村 誠, 太田博光, 森元映治, 江副 寛, 前田俊道, 三好佳廣, 鮮魚の熟練的品質評価の解析—有彩色魚種の外観評価について—, 人間工学, Vol.45, pp.118-125, 2009.
- [4] M. Nakamura, K. Matsumoto, E. Morimoto, S. Ezoe, T. Maeda, T. Hirano, Model of Auctioneer Estimation of Swordtip Squid (*Loligo edulis*) Quality, Kansei Engineering International Journal, Vol.10, pp.99-107, 2011.
- [5] 中村 誠, 太田博光, 平 雄一郎, 森元映治, 江副 寛, 前田俊道, 三好佳廣, 鮮魚の熟練的品質評価の解析—彩度の低い魚種の外観評価について—, 人間工学, Vol.48, pp.142-149, 2012.
- [6] 中村 誠, 太田博光, 平 雄一郎, 森元映治, 江副 寛, 前田俊道, 中村堯史, 下関トラフグ身欠きの熟練的品質評価の解析, 人間工学, Vol.48, pp.304-312, 2012.
- [7] 中村 誠, 太田博光, 鴻上健一郎, 明田川雅子, 徳永憲洋, 前田俊道, ファジィ推論を用いたフグ類身欠きの熟練的品質評価モデル, 知能と情報, Vol.26, pp.781-792, 2014.
- [8] 中村 誠, 鴻上健一郎, 明田川雅子, 太田博光, 徳永憲洋, 色彩によるフグ肉の肉眼鑑別に関する基礎的検討, 人間工学, Vol.50, pp.294-299, 2014.
- [9] T. Saito, K. Arai, M. Matsuyoshi, A New Method for Estimating the Freshness of Fish, Journal of Japanese Society of Scientific Fisheries, Vol. 24, pp.749-750, 1959.
- [10] E. H. Lee, T. Oshima, C. Koizumi, High performance liquid chromatographic determination of K

value as an index of freshness of fish, Bulletin of the Japanese Society of Scientific Fisheries, Vol. 48, p.255, 1982.

[11] R. Fujii, Cytophysiology of fish chromatophores, Int. Rev. Cytol., Vol.143, pp.191-255, 1993.

[12] M. A. McNiven and K. R. Porter, Microtubule polarity confers direction to pigment transport in chromatophores, J. Cell Biology, Vol.103, pp.1547-1555, 1986.

[13] S. K. Whyte, The innate immune response of finfish – A review of current knowlege, Fish Shellfish Immunol., Vol.23, pp.1127-1151, 2007.

[14] 山中英明, 魚介類の死後変化と品質, 日本水産学会誌, Vol.68, pp.5-14, 2002.

[15] 中村 誠, 酒井義郎, 藤本正克, 高浪五男, 魚類乾燥予測モデルの構成, 計測自動制御学会論文集, Vol.31, pp.1778-1786, 1995.

[16] E. H. Mamdani and S. Assilian, An experiment in linguistic synthesis with a fuzzy logic controller, Int. J. Man Machine Studies, Vol.7, pp.1-13, 1974.

[17] 中村 誠, 平 雄一郎, 森元映治, 橋立知典, 池田一朗, 三好佳廣, 鮮魚の熟練的品質評価に関する基礎的検討ーアカアマダイの外観評価についてー, 人間工学, Vol.43, pp.261-267, 2007.

[18] 渡邊悦生編著, 魚介類の鮮度と加工・貯蔵, pp.61-80, 成山堂書店, 1998.

[19] 文部科学省, 五訂増補日本標準成分表, 食品コード 10236, 2005.

[20] 中村 誠, 他 5 名, 彩度の低い魚種の熟練的品質評価のモデル化, 人間工学, Vol.47(特別号), pp.266-267, 2011.

[21] 海沼 勝, 新ふぐ調理師必携 (下巻) 料理編, pp.20-23, 柴田書店, 東京, 2008.

[22] 中村堯史, 中村 誠, 太田博光, 平雄一郎, 明田川雅子, 森元映治, ファジィ推論を用いたフグ身欠きの品質評価について, 日本知能情報ファジィ学会第 17 回中国四国支部大会予稿集, 12 月 2 日, 山口市, pp.42-43, 2012.

連絡先

〒 759-6595 下関市永田本町 2 丁目 7 番 1 号

独立行政法人水産大学校水産学研究科

中村 誠

nakamura@fish-u.ac.jp

ファジィランダム重み付き Weber 問題に対する最適配置の導出

An Approach for Finding an Optimal Location of Fuzzy Random Weighted Weber Problems

○¹ 宇野 剛史, ² 片桐 英樹, ³ 加藤 浩介

○¹ Takeshi Uno, ² Hideki Katagiri, ³ Kosuke Kato

¹ 徳島大学

¹ Tokushima University

² 広島大学

² Hiroshima University

³ 広島工業大学

³ Hiroshima Institute of Technology

Abstract: This article considers a new location problem, which is extended Weber problems by adding fuzzy random weights. It can be formulated as a fuzzy random programming problem. By translating it to a deterministic programming problem based upon α -level sets and chance constraints, its solution method is shown.

1 はじめに

Weber 問題は最適配置問題の一つであり、施設から需要の存在する点 (以下では需要点と呼ぶ) までの距離の総和を最小化する配置を求める問題である。Weber 問題を含む最適配置問題の詳細については、栗田 [8] の書籍が詳しい。各需要点に付随する需要の大きさを考慮する場合、重み付き距離の総和の最小化を目的とする重み付き Weber 問題に拡張される。Weber 問題の現実への応用例として、配送センター、消防署などの施設配置が挙げられる。Weber 問題は多くの研究者に考察され、近年でも多くの関連論文が出版されている [1, 5]。

本研究では、需要が不確実性およびあいまい性を同時に伴う場合について考察する。例えば配送センターの配置において、店舗から要求される配送品に関する需要の頻度およびタイミングは不確実性であり、かつ発生した需要の見積もりにはあいまい性が含まれる。このような需要をファジィランダム変数 [7] として表すことで、Weber 問題をファジィランダム重み付き Weber 問題に拡張し、ファジィランダム計画問題 [4] として定式化されることを示す。

定式化された問題は目的関数値がファジィランダム変数で与えられることから、最適配置を直接求めることはできない。本研究では、 α -レベル集合および機会制約条件 [6] を用いることで確定計画問題に変換する

ことを提案する。さらに、変換された問題が凸計画問題となることから一般的な非線形計画法で求解可能であることを示す。

2 ファジィランダム重み付き Weber 問題の定式化

A を数直線 $\mathbf{R}$ 上のファジィ集合とし、そのメンバーシップ関数を $\mu_A : \mathbf{R} \rightarrow [0, 1]$ とおく。このとき、 μ_A が $\mathbf{R}$ 上で準凸かつ上半連続であるとき、 A はファジィ数 [2] であると定義される。また、定数 $\alpha \in (0, 1]$ に対して、 A の α -レベル集合を次式で定義する：

$$A_\alpha \equiv \{z \mid \mu_A(z) \geq \alpha\}. \quad (1)$$

このとき、 Ξ がファジィランダム変数 [7] であるとは、その α -レベル集合 $\Xi_\alpha(\omega)$ が任意の $\alpha \in (0, 1]$ に対してランダム区間となることと同値であると定義される。

以下では、平面上に m 個の需要点が存在する場合について、その平面内の凸集合 S 上への施設配置について考察する。 $i, i = 1, 2, \dots, n$ 番目の点について、その位置および需要を各々 v_i および w_i^{FR} で表す。本研究ではすべての $w_i^{FR}, i = 1, 2, \dots, n$ を、発生確率が $p_1, p_2, \dots, p_m$ と与えられる m 個のシナリオで、実現値がすべてファジィ数となるファジィランダム変数により与える。

S へ配置する施設の位置を x で表し、定式化される

問題の決定変数とみなす。施設から i 番目の需要点までの間の距離を $d_i(\mathbf{x})$ と表す。このとき、施設から全ての需要点までの距離の重み付き総和は次式で表される：

$$f^{FR}(\mathbf{x}) := \sum_{i=1}^n w_i^{FR} d_i(\mathbf{x}) \quad (2)$$

よって、ファジィランダム重み付き Weber 問題は次のファジィランダム計画問題として定式化される：

$$\left. \begin{array}{ll} \text{minimize} & f^{FR}(\mathbf{x}) \\ \text{subject to} & \mathbf{x} \in S \end{array} \right\} \quad (3)$$

3 解法

$f^{FR}(\mathbf{x})$ はファジィランダム変数で与えられることから、最適配置を求めるためには確定値に変換する必要がある。まず、意思決定者は自身の状況に応じて必要とする目的関数値の実現度を $\alpha \in (0, 1]$ で与える。このとき、 $f^{FR}(\mathbf{x})$ の α -レベル集合の最小値を $f_\alpha^R(\mathbf{x})$ で与えることにより、(3) は次の確率計画問題に変換される：

$$\left. \begin{array}{ll} \text{minimize} & f_\alpha^R(\mathbf{x}) \\ \text{subject to} & \mathbf{x} \in S \end{array} \right\} \quad (4)$$

つぎに、意思決定者は自身の状況に応じて保障したい目的関数値の達成確率を $\sigma \in (0, 1]$ で与える。このとき、機会制約条件 [6] を追加することで、(4) は次の確定計画問題に変換される：

$$\left. \begin{array}{ll} \text{minimize} & \beta \\ \text{subject to} & \Pr [f_\alpha^R(\mathbf{x}) \leq \beta] \geq \sigma \\ & \mathbf{x} \in S \end{array} \right\} \quad (5)$$

ここで、 $\Pr(\cdot)$ は達成確率を与える関数であり、 β は与えられた達成確率を保障する目的関数値である。

$j, j = 1, 2, \dots, m$ 番目のシナリオに対する $f_\alpha^R(\mathbf{x})$ の実現値を $f_\alpha^j(\mathbf{x})$ で表す。このとき、問題 (5) に対して、次の定理が成り立つ。

定理 1 任意に与えられた $\alpha \in (0, 1]$ および $\sigma \in (0, 1]$ に対して、(5) は凸計画問題である。

(証明) ファジィランダム変数の定義より、与えられたシナリオに対してすべての需要点に対してその需要に α -レベル集合をとってその最小値を取ることができ、(2) で重み付き距離の総和をとると $f_\alpha^j(\mathbf{x})$ が得られる。 $d_i(\mathbf{x})$ は凸関数であることからそれらの定数重み付き総和もまた凸関数となるので、 $f_\alpha^j(\mathbf{x})$ は凸関数である。よって、 $f_\alpha^j(\mathbf{x}) \leq \beta$ をみたす領域は凸集合で表され、(5) における機会制約条件もまた凸集合となる。以上より、(5) が凸計画問題であることが示された。

定理 1 より、(5) は逐次二次計画法などの一般の非線形計画法 [3] により求解可能である。

4 おわりに

本論文では、需要がファジィランダム変数で表される場合の重み付き Weber 問題について考察した。本問題の数値例への適用については、ワークショップ発表にて紹介する予定である。

参考文献

- [1] J. Brimberg, Z. Drezner, N. Mladenović, S. Salhi: A new local search for continuous location problems, *European Journal of Operational Research*, Vol. 232, No. 2, pp. 256-265, 2014
- [2] 古川 長太: ファジィ最適化の数理, 第 2 章 ファジィ集合とファジィ数, 森北出版, pp. 14-23, 1999
- [3] 茨木 俊秀: 最適性の数学 (共立講座 21 世紀の数学 13), 3.3 制約つき問題のアルゴリズム, 共立出版, pp. 112-132, 2011
- [4] 石井 博昭, 斎藤 誠慈, 森田 浩: 不確実・不確定性の数理 (大阪大学新世紀レクチャー), 大阪大学出版会, 2004
- [5] A. Jamalian, M. Salahi: Robust solutions to multi-facility Weber location problem under interval and ellipsoidal uncertainty, *Applied Mathematics and Computation*, Vol. 242, pp. 179-186, 2014
- [6] H. Katagiri, H. Ishii: Chance constrained bottleneck spanning tree problem with fuzzy random edge costs, *Journal of the Operations Research Society of Japan*, Vol. 43, No.1, pp. 128-137, 2000
- [7] R. Kruse, K.D. Meyer: *Statistics with Vague Data*, D Riedel Publishing Company, 1987
- [8] 栗田 治: 都市モデル読本 (造形ライブラリー 05), 共立出版, 2004

連絡先

〒770-8502 徳島市南常三島町 1-1

徳島大学 大学院 SAS 研究部

宇野 剛史

E-mail: uno.takeshi@tokushima-u.ac.jp

複数台のカメラを用いたプリント基板検査経路の最適化

Path optimization using multiple cameras for inspecting printed circuit boards

¹ 片桐 英樹 ¹ 郭 慶強 ○ ¹ 王 彬 ¹ 村中 智幸
¹Hideki Katagiri ¹Guo Qingqiang ○ ¹Wang Bin ¹Tomoyuki Murakami
² 羽森 寛 ³ 加藤 浩介
²Hiroshi Hamori ³Kosuke Kato

¹ 広島大学

¹Hiroshima University

² オー・エイチ・ティー株式会社

²OHT Inc.

³ 広島工業大学

³Hiroshima Institute of Technology

Abstract: In inspections of printed circuit boards (PCBs), a test probe has to be moved to a number of PCB wiring patterns in some order. Due to the miniaturization of PCBs, the procedure of “alignment” has to be done by cameras in order to know the exact position of each wiring pattern before each of PCB wiring patterns is electrically tested. This paper considers path optimization using multiple cameras for minimizing a total path length of testing PCBs arrayed on a plane, which is modeled as a kind of traveling salesman problems with a precedence constraint. In order to solve the problem with a practically reasonable time, an efficient heuristic algorithm is proposed.

1 はじめに

プリント基板は、携帯電話やパソコンなどの電子機器のほとんどに使用されており、プリント基板の生産の効率化は重要な課題である。

プリント基板の検査では、同一の配線パターンが多数並べられた基板シートの上を検査治具が移動しながら、1つ1つの配線パターンに欠陥がないかどうかをチェックする。近年、電子機器の小型化に伴い、プリント基板配線のピッチは数十 μm と非常に狭くなり、検査プローブを配線パターンに正確に合わせることが難しくなっている。そのため、配線パターンの正確な位置を知るために、カメラによるアライメント撮像を行う新しい検査システムが開発されている。このようなアライメント撮像を考慮したプリント基板検査経路最適化問題は、0-1 整数計画問題として定式化することが可能であり [2]、実用時間内に近似最適解を導出する解法も提案されている [3]。

従来研究 [2, 3] では、アライメント撮像を行うカメラが1台の場合のみを扱っていたが、更なる効率化を図るためには複数台のカメラを用いた検査システムを

考えることができる。本研究では、複数のカメラによるアライメント撮像を考慮したプリント基板の検査経路最適化問題を「先行順序制約付き一般化巡回セールスマン問題」としてモデル化し、高速に近似最適解を求める近似解法を提案する。

2 基板検査の方法と経路最適化

本節では、カメラによるアライメント撮像を含んだプリント基板検査及び従来解法について紹介する。

2.1 アライメント撮像を考慮した基板検査

図1はカメラが取り付けられたプローブ治具とプリント基板の模式図を表している。この例では1つの配線パターンに対して、2点のアライメント点が存在する。この配線パターンを検査する際には、プローブ治具の中心を検査点に合わせる前に、2つのアライメント点をカメラで撮像しなければならないという「先行順序制約」がある。したがって、本研究は先行順序制約付き巡回セールスマン問題と密接な関係がある。

通常、検査対象は同一の配線パターンが平面上に複数個並べられて1枚のシート（基板シート）になって

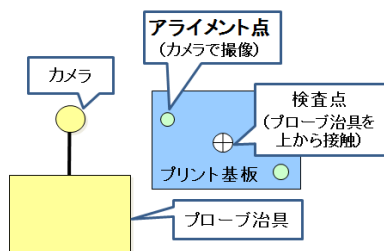


図 1: アライメント点と検査点

いる。1枚の基板シートにはおよそ4個～200個の配線パターンが並べられ、一台の検査機で1日あたり1,000枚以上の基板シートを検査することになる。したがって、基板シート1枚あたりの検査時間の改善（短縮）がたとえ数パーセントでも大きな効率改善につながる。

図2に配線パターン数4（縦2×横2）の基板シートの例を示す。この例では、検査前のプローブ治具の初期位置が基板シートの左側にある。検査ユニットの基準点はプローブ治具の中心点とする。

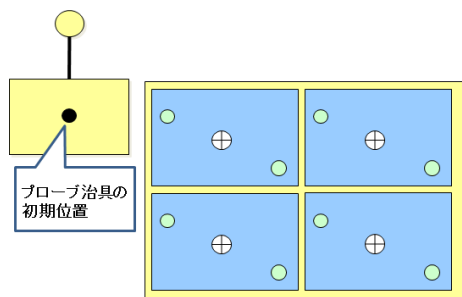


図 2: 基板シート（配線パターン数が4の例）

2.2 1台のカメラ用いた基板検査と経路最適化 [2, 3]

図2に対して、当初、現場で採用されていた訪問順序を図3に示す。

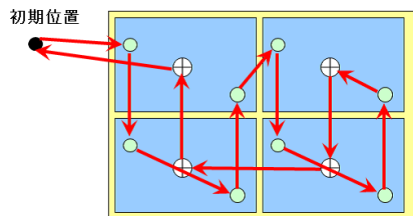


図 3: 現場で採用されていた訪問順序

図3においては、左上のアライメント点から開始し

て、順にアライメント点のみに移動し、先にすべてのアライメント点が訪問されている。その次に、最も近い検査点に向かい、アライメント点の訪問とは逆向きの順序で検査点のみを訪問して初期位置に戻るという手順で検査が行われていた。すなわち、先にまとめてアライメント点を訪問（カメラで撮像）し、後でまとめて検査点を訪問（プローブ治具を接触させて検査）するというように、アライメント点への訪問と検査点への訪問は別々に考えられていた。このようにすれば、先行順序制約は必ず満たされる一方で、一般には無駄が多い経路となるため、最適化技術により大幅な改善が見込める。そこで、片桐ら [2, 3] は、この最適化問題を「先行順序制約付き巡回セールスマン問題」としてモデル化し、実用的な近似解法を提案した。

3 複数台のカメラを用いた基板検査

プリント基板検査をより高速に行うためには、プローブ治具に取り付けられるカメラを複数台に増やすことが考えられる。本節では、カメラを複数台用いた場合のプリント基板検査経路最適化問題が「先行順序制約付き一般化巡回セールスマン問題」としてモデル化できることを示すとともに、高速に近似最適解を導出する近似解法を提案する。

3.1 先行順序制約付き一般化巡回セールスマン問題としてのモデル化

最も簡単な例としてカメラを2台用いた場合を考える。図4において、A点のアライメント点をカメラで撮像する場合、カメラ1とカメラ2のどちらかを用いればよい。すなわち、アライメント点Aの撮像はプローブ治具の中心をA'とA''のどちらかに合わせることで行える。

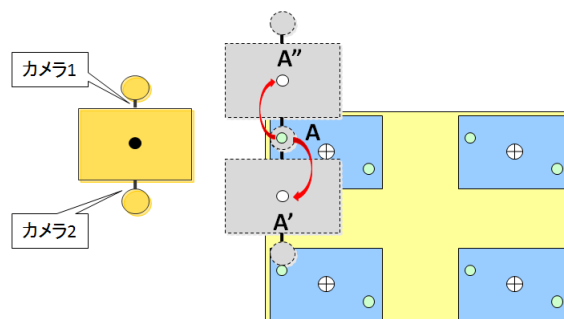


図 4: 2台のカメラを用いた場合のアライメント撮像時のプローブ治具の中心位置

ここで、アライメント点に対応する2点 A' , A'' をクラスタ（アライメントクラスタ）としてグループ化させると、この2点のどちらかを訪問すればよいことになる。他のアライメント点についても同様にプローブ治具の中心点でアライメント点を置き換え、要素2のクラスタを形成する（図5参照）。さらに、検査点を要素数1のクラスタ（検査クラスタ）とみなす。このとき、2台のカメラを用いた基板検査経路最適化問題は、「検査クラスタを訪問する前に（その検査クラスタに対応する）アライメントクラスタを訪問する必要があるという先行順序制約の下で、すべてのクラスタの中の少なくとも1点を訪問する最短巡回路を求める問題」となる。この問題は、先行順序制約付き巡回セールスマン問題と一般化巡回セールスマン問題を組合わせた「先行順序制約付き一般化巡回セールスマン問題」とみなすことができる。図5において、矢印は先行順序制約を表しており、ここでは、それぞれの検査クラスタ（要素数1）に対して、2つのアライメントクラスタ（要素数2）が先に訪問される必要があるという制約を示している。一般にカメラを k 台使用するときは、アライメントクラスタの要素数は k となる。

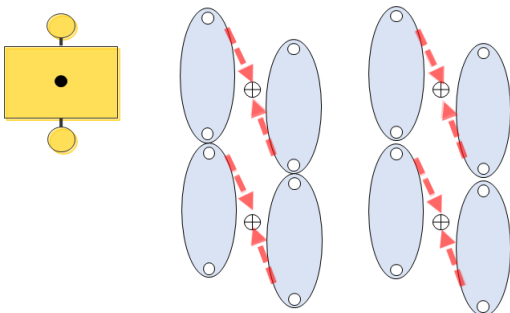


図 5: プローブ治具の中心が訪問するクラスタ

3.2 近似解法の提案

ここでは、実用的な計算時間で近似最適解を求める解法を提案する。提案解法の概要は次のようになる。

【提案解法の概要】

手順 1: 最近近傍法により、初期解を生成する。

手順 2: 現在の訪問点を固定した状態で、2-opt, Or-opt, 順序入替操作による局所探索を行う。

手順 3: 訪問するクラスタの順序を固定した状態で、各クラスタにおいて選択する訪問点の最適化を行う。終了条件を満たせば終了し、満たさなければ手順 2 へ戻る。

手順 1 における初期解の生成では、最近近傍法の他に、貪欲法、最近挿入法、最速挿入法の方法もあり、すべて試みた結果、最近近傍法で初期解を生成した場合が最も良い結果となった。

手順 2 では、手順 1 で得た初期解に対して 2-opt 法や Or-opt 法による局所探索を行う。ただし、2-opt や Or-opt を用いて検査経路の枝をつなぎ換えることにより、アライメント点と検査点の間の先行順序制約を満たさない解（検査経路）が生成される可能性がある。過去の研究では、先行順序制約を満たすようなつなぎ換えを高速に見つけて最良解を更新していく逐次的な局所探索法 [5] が提案されている。一方、Renaud ら [6] は通常の巡回セールスマン問題と同様の局所探索を行って、一旦は先行順序制約を考慮しないで解を生成した後、「各時点での最良解よりも良い（経路長が短い）解が見つかったときのみ」先行順序制約のチェックを行うことにより高速に良質な近似解を導出する方法を提案している。本研究でも Renaud らの研究と同様に、一旦は先行順序制約を考慮しないで解を生成した後、各時点での最良解よりも良い解が見つかったときのみ先行順序制約のチェックを行う。ただし、Renaud らの方法では先行順序制約を満たさない実行不可能解は探索候補から除外していたのに対し、本研究では、実行可能解に対して先行順序制約を満たすように修正する操作を組み込む。具体的には、実行不可能解において、先行順序制約を満たしていない検査点を対応するアライメント点の位置の後に移動する。ただし、一般に移動する位置は複数個存在するため、経路が最も短くなる位置に検査点を移動する。

手順 3 では、手順 2 で求まった解におけるクラスタの訪問順序は固定する。各アライメントクラスタに含まれる k 点（ k はカメラの数）の中で、いずれの点を選択すると経路が短くなるかを探索する。これは、複数あるカメラのいずれでアライメントを行うと経路が短くなるかを探索していることと同値である。すなわち、クラスタの訪問順序は固定した状態で、クラスタ内で選択する点のみを変えることにより解の改善を行う。ただし、すべての組合せを考えると多大な計算時間を要するため、ここでは、手順 2 で局所探索で順序が入れ替わったクラスタの部分集合のみについて最適

化を行う。

手順3を実施した後は、再び手順2に戻る。そして、何らかの終了条件（定められた計算時間、改善解が得られない時点など）が満たされるまで手順2と3の操作を繰り返す。

4 数値実験

提案手法の有用性を検証するために、現場での8種類の実データに対して前章にて提案した近似解法を適用し、数値実験を行った（表1参照）。 n は配列パターン数を表す。ここでは1パターンあたりのアライメントは全て2点とした。計算機環境は、Intel(R) Core(TM) i7-2640M(2.80GHz), RAM:4.00G, OS: Windows 7 (64bit), プログラミング言語: Microsoft Visual C++ 2010 Express を用いた。

表 1: 数値実験結果

n	現場	従来解法 [3] (時間: 秒)	提案解法 (時間: 秒)
12	2427.5	2100.8 (0.082)	1938.1 (0.016)
18	3606.1	2602.6 (0.12)	2236.0 (0.016)
30	7427.8	4583.8 (0.35)	3350.6 (0.31)
48	5342.2	3825.8 (0.72)	3526.3 (0.13)
50	8351.9	4722.5 (0.73)	4527.0 (0.51)
100	16150.5	10437.7 (6.5)	10256.1 (13.07)
150	23139.1	12130.8 (11.8)	11804.6 (26.1)
200	31508.0	12115.8 (7.5)	9789.3 (3.04)

表1に示される通り、現場の検査経路に比べて、平均で約45%と大きく改善されており、大規模問題でも30秒以内で近似最適解が導出されている。また、従来の1台のカメラを用いた場合と比較しても平均で13%以上改善されており、費用対効果を考慮したときに、カメラを1台追加するのに十分に高い改善率が得られている。提案手法は既に現場の検査システムに組み込まれており、その実用性が認められている。

5 まとめ

本研究では、複数のカメラによるアライメント撮像を考慮したプリント基板検査経路の最適化問題を「先行順序制約付の巡回セールスマン問題」としてモデル化した。また、高速な近似解法を提案し、現場のデータに基づく数値例を用いた数値実験により、その有用性を示した。

今後の課題としては、LK法やその拡張版[1]、進化計算[4]などに基づく高性能な近似解法や分枝カット法による厳密解法の構築が挙げられる。

参考文献

- [1] K. Helsgaun, “An effective implementation of the lin-kernighan traveling salesman heuristic,” *European Journal of Operational Research*, Vol. 126, No. 1, pp. 106–130, 2000.
- [2] 片桐, 呉, 羽森, 加藤, プリント基板の電気検査におけるアライメントを考慮した検査経路最適化, 日本オペレーションズ・リサーチ学会春季研究発表会アブストラクト集 2014, pp. 38–39, 2014.
- [3] H. Katagiri, Q. Guo, H. Wu, H. Hamori, and K. Kato, “Path optimization for electrically inspecting printed circuit boards with alignment marks,” *Proceedings of the International Multi-Conference of Engineers and Computer Scientists 2015* (to appear).
- [4] 永田裕一, 局所的な交叉EAXを用いたGAの高速化とTSPへの適用, 人工知能学会論文誌, Vol. 22, No. 5, pp. 542–552, 2007.
- [5] H.N. Psaraftis, “ k -interchange procedures for local search in a precedence-constrained routing problem,” *European Journal of Operational Research*, Vol. 13, pp. 391–402, 1983.
- [6] J. Renaud, F.F. Boctor, G. Laporte, “Perturbation heuristics for the pickup and delivery traveling salesman problem,” *Computers & Operations Research*, Vol. 29, pp. 1129–1141, 2002.

連絡先

広島大学大学院工学研究院

片桐 英樹

E-mail: katagiri-h@hiroshima-u.ac.jp

タクシー勤務時間の最適配置

Optimum Arrangement of Taxi Operation Time

○谷崎 隆士
○Takashi Tanizaki
近畿大学

Abstract: We propose optimum arrangement of taxi operation time. In Japan, income of taxi vehicle is decreasing about 11 thousand yen in the past 15 years. Then some taxi companies are investing to gain more customers. But there are many small taxi companies that are difficult to invest with much money. Therefore we have been researching the other method to gain more customers by little investment for small companies. In this paper, we analyze present situation of the Taxi Company which we research, research optimum arrangement of taxi drivers' working hours so as to increase sales amount using mathematical model, and verify validity of our method with numerical calculation.

1. はじめに

バブル崩壊後、わが国のタクシー利用者数は減少し続けている。1995年から2011年にかけて、年間利用者数は約23億4000万人から約15億3000万人へ約35%減少し、年間売上高も約2兆4300億円から約1兆5600億円へ約35%減少した。一方、タクシー台数は21万両から19万7000両と約6%の減少に留まった。この結果、日車営収(=総営業収入/延べ実働車両数)が、約37700円から約27200円へ約10500円(約28%)減少した[1][2]。従って、タクシー会社同士の競争が激化し、タクシー利用者の囲い込みや運行の効率化のために、様々な方法を導入している。例えば、GPSやスマートフォンを利用した運行管理システム等の導入などがある。しかし、日本のタクシー会社は殆どが資本金1千万円以下の中小企業にて、投資が必要なシステムの導入に踏み切れない企業が多い。以上より、少額の投資でタクシー利用者を増加する手法について研究してきた。

(1)タクシー配車依頼件数の多い時間帯のタクシーの稼働の増加を狙いに、単位時間・タクシー1台当たりの電話での配車依頼件数を最小化するタクシー乗務員の勤務の再配置策[3]。

(2)タクシー待機場所からタクシー必要者への総走行距離の最小化を狙いとした、タクシー待機場所におけるタクシー配置台数の決定方法[4]。

上記の研究におけるデータ分析を通じて、タクシーの勤務時間帯がタクシー売上高の高い時間帯と一致していない事が明らかになった。タクシー売上高

を増加するためには、売上高の高い時間帯に勤務するタクシー台数を多くした方が良い。本報告では、上記を狙いとして数理手法を用いたタクシー勤務時間の最適配置についての研究結果について報告する。

2. A社の現状分析

共同研究先の東広島市のA社は、前述の日本の大多数を占めるタクシー会社の範疇に属する企業である。A社の現状分析として、タクシー勤務時間および2012年の売上実績について以下に述べる。

2.1 A社のタクシー勤務体系

表1にA社の勤務形態別のタクシー勤務時間を示す。勤務形態はシフト勤務・日勤・定時制・夜勤である。シフト勤務は、1日おきに1勤～4勤と公出を変更する。公出は、本来休日であるが、希望があれば7時～24時の間で自由に勤務できる。日勤は月曜日～土曜日に勤務する。定時制は「月・水・金」勤務と「火・木・土」勤務で構成される。夜勤は「月・水・金」と「火・木・土」の勤務を一週間ごとに交代する勤務と、毎週「金・土」の勤務で構成される。

表1 勤務形態別のタクシー勤務時間

勤務形態		開始時刻	終了時刻	稼働時間
シフト 勤務	1 勤	7:00	1:00	18 時間
	2 勤	8:00	2:00	18 時間
	3 勤	8:00	3:00	19 時間
	4 勤	12:00	8:00	20 時間
	公出	7:00	24:00	17 時間
日 勤		7:00	19:00	12 時間
定時制		7:00	21:00	14 時間
夜 勤		20:00	4:00	8 時間

2.2 A社の売上高実績

図1, 図2に2011年12月から2012年11月のA社の売上高実績を示す。図1は、曜日別時間帯別年間平均売上高であり、図2は、週別時間帯別年間平均売上高である。平均売上高は、最低金額が1.0となるように正規化している。図1, 図2より年間平均売上高の特徴は以下の通りである。

- (1) 曜日によって異なる。金曜りと土曜日が高く、日曜日が低い。
- (2) 週による差はほとんどない。
- (3) 8時～9時と17時～24時は高く、10時～11時と13時～16時は低い。2時～6時は非常に低い。金曜りと土曜日の22時～2時は、特に高い。

表1と図1, 図2よりタクシー勤務時間帯と売上高の高い時間帯が一致していない部分がある。

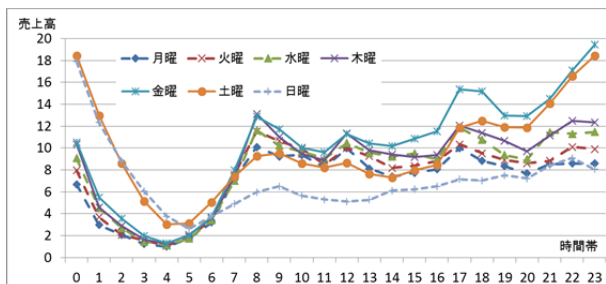


図1 曜日別時間帯別年間平均売上高

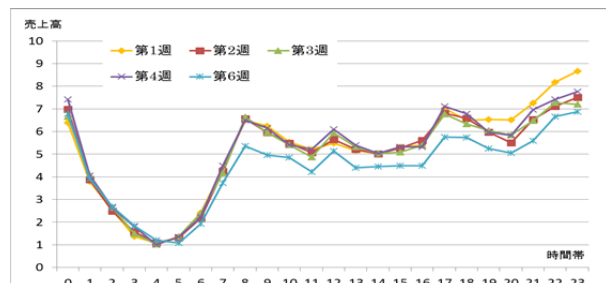


図2 週別時間帯別年間平均売上高

3. 定式化

2に記載の通り、タクシー勤務時間帯と売上高の高い時間帯が一致していない部分がある。売上高の増加には、売上高の高い時間帯の勤務台数を多くした方がよい。そこで、売上高の増大を狙いとしたタクシー勤務時間の最適配置について数理手法を用い研究を行った。定式化の基本的な考え方は、以下の制約条件を遵守し、1日のいずれかの時間帯にタクシー連続勤務時間の帯を割り付ける事である(図3)。なお、表1の公出は、勤務日が確定的ではないため、

考慮しない。また、1週目と2週目で夜勤の勤務曜日が異なるため、2週間分の定式化を行う。

- (1) 勤務形態ごとの連続勤務時間を遵守する。
- (2) 勤務形態ごとの勤務開始時刻終了時刻は毎日同じである。
- (3) どの時間帯も少なくとも1つシフト勤務を行う。

	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	21	22	23
1勤																								
2勤																								
3勤																								
4勤																								
日勤																								
定時制																								
夜勤																								

図3 勤務形態ごとのタクシー勤務時間

3.1 記号の定義

i : 勤務時刻($i=0, 1, \dots, 335$)

j : 勤務形態($j=0, 1, \dots, 7$)

但し、 $\left[\begin{array}{l} 0:1 \text{ 勤}, 1:2 \text{ 勤}, 2:3 \text{ 勤}, 3:4 \text{ 勤}, \\ 4: \text{日勤}, 5: \text{前半定時制}, 6: \text{後半定時制}, \\ 7: \text{夜勤} \end{array} \right]$

現状分析より午前と午後の間に売上高の低い時間帯がある事より、定時制を2つに分割した定式化を行うために、定時制を前半・後半に分割した

C_i : 勤務時刻 i における売上高

但し、 $C_i = \text{勤務時刻 } i \text{ の年間売上高} /$

勤務時刻 i の年間日数

N_{ij} : 勤務時刻 i 時の勤務形態 j のタクシー勤務可能台数

T_j : 勤務形態 j の連続タクシー勤務時間

x_{ij} : 勤務時刻 i における勤務形態 j のタクシー勤務の有無を表す 0-1 変数

$$x_{ij} = \begin{cases} 1 & (\text{タクシー勤務有}) \\ 2 & (\text{タクシー勤務無}) \end{cases}$$

3.2 定式化

$$\text{Maximize } \sum_{i=0}^{335} \sum_{j=0}^7 C_i \cdot N_{ij} \cdot x_{ij} \quad (1)$$

$$\text{subject to } \sum_{i=0}^3 x_{ij} \geq 1, \quad (2)$$

$$\sum_{i=0}^{23} x_{ij} = T_j, \quad (3)$$

$$\sum_{i=1}^{23} x_{ij} \cdot x_{i-1,j} = T_j - 1, \quad (4)$$

$$x_{ij} = x_{i+24 \cdot n, j} \quad (n=1, 2, \dots, 23). \quad (5)$$

目的関数(1)式は、売上高合計の最大化である。制約条件は、(2)式～(5)式である。(2)式は、どの時間帯も少なくとも1つシフト勤務を行う事を示す。(3)式と(4)式は、勤務形態jがT_j時間連続勤務である事を示す。(5)式は、勤務形態ごとの勤務開始、終了時刻が毎日同じであることを示す。

3.3 解法アルゴリズム

本問題は、非線形計画問題であるため、混合整数計画問題に対する求解法である分枝限定法を用いる事が出来ない。しかし、変数が0・1整数であることから重み付き制約充足問題(WCSP)に変換し、タブーサーチを用いて求解する。WCSPは、重要な制約条件をなるべく満足するためには値をどのように割り当てるとよいかを決定する問題である[5]。WCSPでは、制約条件の重要度は重み付けパラメーターで与える。従って、制約条件の遵守度に対するモデル設計者の意向を反映できる。目的関数 $f(x_1, x_2, \dots, x_n)$ を有する最適化問題の場合は Target 値 μ を導入し、下記の制約条件に変換する。

(1)最小化問題： $\mu \cdot f(x_1, x_2, \dots, x_n) \geq 0$

(2)最大化問題： $f(x_1, x_2, \dots, x_n) \cdot \mu \geq 0$

目的関数はソフト制約(必ずしも満たす必要はないが、できるだけ満たして欲しい制約)として扱う。本問題は(1)式を(6)式に変換し WCSP として扱う。

$$\sum_{i=0}^{335} \sum_{j=0}^7 C_i \cdot N_{ij} \cdot x_{ij} - \mu \geq 0. \quad (6)$$

WCSP の終了条件は、下記いずれかを満足することである。

- (1)すべての制約条件を満足する解が得られた時。
- (2)予め設定した反復回数を超えた時。
- (3)予め設定した計算時間を超えた時。

本研究では、(3)の終了条件を採用し、計算時間は3600秒とした。求解手順は以下の通りである。

- S.1 計算時間を設定。
- S.2 Target 値を設定。
- S.3 タブーサーチを用いて WCSP を求解。
- S.4 設定した計算時間以内に解が出た場合は、Target 値を増やし S.3 に戻る。
- S.5 設定した計算時間を超えた時点で計算を終了し、計算終了時点の解を解とする。

4. 数値計算結果

4.1 数値計算ケース

3.3に記述した手法を用いて準最適解を求解する。現状の勤務形態にこだわらず、売上高がより高くなる勤務形態を研究してほしいというA社からの要望に応えるため、新しい勤務形態となる制約条件を本問題に追加した問題の求解結果と現状との比較を行う(表2)。

表2 数値計算ケース

パターン	ケース	追加制約
0	0	変更なし
1 定時制の勤務時間を変更	1	8時～22時に勤務
	2	9時～23時に勤務
	3	10時～0時に勤務
	4	11時～1時に勤務
	5	12時～2時に勤務
2 定時制の勤務時間を分割	6	7時～14時, 17時～0時に勤務
	7	8時～15時, 17時～0時に勤務
	8	8時～14時, 18時～1時に勤務
	9	7時～13時, 16時～0時に勤務
	10	8時～14時, 16時～0時に勤務
3 定時制の勤務曜日を変更	11	8時～14時, 17時～1時に勤務
	12	A班：月金土, B班：火金土に勤務
	13	A班：月金土, B班：水金土に勤務
	14	A班：月金土, B班：木金土に勤務
	15	A班：火金土, B班：水金土に勤務
	16	A班：火金土, B班：木金土に勤務
	17	A班：水金土, B班：木金土に勤務

(1)パターン0

[内容]3の定式化通りに求解。

[狙い]定式化通りで準最適解を求解し、現状の勤務時間と比較。

(2)パターン1

[内容]勤務人数の最も多い定時制の勤務時間帯を変更して固定する制約条件を追加。

[狙い]顧客の多い20時以降の勤務台数を増加。

(3)パターン2

[内容]定時制の勤務時刻を分割する制約条件を追加。

[狙い]14～16時の売上高が低い時間帯のきぬ台数を減少し、20時以降の売上高が高い時間帯の勤務台数を増加。

(4)パターン 3

[内容]定時制の勤務する曜日を変更する制約条件を追加.

[狙い]売上高の高い金曜日と土曜日の勤務台数を増加.

4.2 数値計算結果

A社より入手した 2011 年 12 月から 2012 年 11 月のデータを用いた数値計算結果を表 3 に示す. 年間売上高が最も高い結果はパターン 3 のケース 17 であり, 全般的にパターン 3 の順位が高かった. さらに売上高が低い 9 時台~16 時台の時間帯の勤務台数を減らし, 売上高の高い 18 時台~0 時台の時間帯の勤務台数を増やす傾向の結果であった(図 4). 以上より, 今後の新しい勤務形態の指針として以下を得た.

(1)月曜~木曜日の勤務台数を減らし, 金曜と土曜の勤務台数を増やす.

(2)9 時~16 時台の勤務台数を減らし, 18 時~0 時台の勤務台数を増やす.

5. おわりに

少ない投資で売上高の増加を狙いとした最適なタクシー勤務時間帯の決定方法について研究した. 現状の売上高分析結果に基づき, 売上高の増加を狙いとした最適化問題として定式化を行った. 本問題は非線形計画問題であるため, WCSP に変換し, タブーサーチを用いて求解する手法を開発した. 数値計算にて本手法の有効性を確認すると共に,

(1)金・土曜の勤務台数を増加,

(2)18 時~0 時台の勤務台数を増加,

の新しい勤務形態についての指針を得た.

参考文献

- [1]国土交通省: 全国のタクシー事業の現状について, <http://www.mlit.go.jp/common/000226517.pdf>.
- [2]杉山雅洋, 山内弘隆, 山本雄二郎: 規制緩和時代のバス&タクシー, 地域科学研究会, 2002.
- [3]谷崎隆士, 片岡隆之: 配車依頼情報に基づくタクシー会社のサービス生産性向上, 日本経営システム学会誌, Vol.30, No.2, p.p.141-147, 2013.

表 3 数値計算結果

	ケース	シフト勤務				日勤	定時制	夜勤	売上高	順位	
		1	2	3	4						
	現状	7:00 1:00	8:00 2:00	8:00 3:00	12:00 8:00	7:00 19:00	7:00 21:00	20:00 4:00	ベース	15	
	0	8:00 2:00	12:00 7:00	6:00 1:00	6:00 2:00	18:00 6:00	13:00 3:00	20:00 4:00	106.0%	5	
パターン1	1								解なし		
	2	8:00 2:00	8:00 2:00	7:00 2:00	17:00 13:00	14:00 2:00	9:00 23:00	17:00 1:00	102.3%	13	
	3	3:00 21:00	7:00 1:00	7:00 2:00	8:00 4:00	14:00 2:00	10:00 0:00	13:00 21:00	99.2%	16	
	4	8:00 2:00	16:00 10:00	8:00 3:00	8:00 4:00	15:00 3:00	11:00 1:00	20:00 4:00	105.9%	6	
	5	7:00 1:00	14:00 8:00	1:00 20:00	5:00 1:00	17:00 5:00	12:00 2:00	23:00 7:00	105.5%	9	
パターン2	6	19:00 13:00	7:00 1:00	8:00 3:00	7:00 3:00	15:00 3:00	7:00 14:00	17:00 0:00	19:00 3:00	105.7%	8
	7	8:00 2:00	8:00 2:00	8:00 3:00	16:00 12:00	8:00 20:00	8:00 15:00	17:00 0:00	18:00 2:00	101.9%	14
	8								解なし		
	9	8:00 2:00	8:00 2:00	7:00 2:00	13:00 9:00	13:00 1:00	7:00 13:00	16:00 0:00	19:00 3:00	103.8%	11
	10	15:00 9:00	8:00 2:00	10:00 5:00	7:00 3:00	14:00 2:00	8:00 14:00	16:00 0:00	20:00 4:00	105.9%	7
	11	5:00 23:00	9:00 3:00	9:00 4:00	1:00 21:00	14:00 2:00	8:00 14:00	17:00 1:00	15:00 23:00	102.4%	12
パターン3	12	2:00 20:00	20:00 14:00	8:00 3:00	16:00 12:00	20:00 8:00	20:00 10:00	15:00 23:00	107.7%	2	
	13	14:00 8:00	22:00 16:00	6:00 1:00	4:00 0:00	14:00 2:00	11:00 1:00	7:00 15:00	106.5%	4	
	14								解なし		
	15	0:00 18:00	7:00 1:00	14:00 9:00	23:00 19:00	15:00 3:00	3:00 17:00	20:00 4:00	104.9%	10	
	16	8:00 2:00	8:00 2:00	11:00 6:00	5:00 1:00	18:00 6:00	13:00 3:00	20:00 4:00	106.7%	3	
	17	7:00 1:00	10:00 4:00	21:00 16:00	6:00 2:00	16:00 4:00	12:00 2:00	20:00 4:00	108.7%	1	

現状の勤務時間

	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	21	22	23
1勤																								
2勤																								
3勤																								
4勤																								
日勤																								
定時制																								
夜勤																								

ケース 17 の勤務時間

	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	21	22	23
1勤																								
2勤																								
3勤																								
4勤																								
日勤																								
定時制																								
夜勤																								

図 4 勤務時間の比較

[4]谷崎隆士, 片岡隆之: タクシー待機場所におけるタクシー配置台数の最適化, 日本経営システム学会誌, Vol.31, No.2, p.p.201-206, 2014.

[5]野々部宏司: メタヒューリスティクスによる汎用ソルバーの構築, オペレーションズ・リサーチ, p.p.257-262, 2011.

連絡先

谷崎 隆士

E-mail: tanizaki@hiro.kindai.ac.jp

マイクロアレイデータ解析に基づく腹膜偽粘液腫の バイオマーカー探索

A microarray gene expression data analysis for discovering biomarkers of pseudomyxoma peritonei

¹ 片桐 英樹 ○ ² 宇野 剛史 ¹ 佐々木 康平 ¹ 村上 和也
¹Hideki Katagiri ○ ²Takeshi Uno ¹Kohei Sasaki ¹Kazuya Murakami
³ 加藤 浩介 ⁴ 桑野 裕昭 ⁵ 米村 豊 ⁶ 遠藤 良夫
³Kosuke Kato ⁴Hiroaki Kuwano ⁵Yutaka Yonemura ⁶Yoshio Endo

¹ 広島大学

¹Hiroshima University

² 徳島大学

²The University of Tokushima

³ 広島工業大学

³Hiroshima Institute of Technology

⁴ 金沢学院大学

⁴Kanazawa Gakuin University

⁵ 腹膜播種治療支援機構

⁵Organization to support Peritoneal Dissemination Treatment

⁶ 金沢大学

⁶Kanazawa University

Abstract: Pseudomyxoma peritonei (PMP) is a rare intractable disease and its rate of incidence is estimated at 1 per 1,000,000 populations annually. Due to the rarity of PMP, a very few doctors encounter a patient who suffers from PMP. Therefore, it is very important to discover tumor markers. This research presents a microarray gene expression data analysis method for discovering biomarkers of pseudomyxoma peritonei (PMP). A set of probes that are good candidates for tumor markers of PMP is extracted by using ESTProfile and GO term. The proposed method finds 3 tumor markers of PMP.

1 はじめに

「腹膜偽粘液腫」とはその発生割合が100万人に1人といわれる難病性のがん種 [7] であり、発病すると粘液性の腫瘍が虫垂や卵巣などから発生、増殖し多量の液体を産出し腹腔内を圧迫する (図1参照)。



図1: 腹膜偽粘液腫の症状

現在の主な治療法としては外科的減量手術・外科手術であるCRS(cytoreductive surgery)と術中温熱化学療法であるHIPEC(hyperthermic chemotherapy)を組み合わせたものが多いが、標準術式はない [3]。

研究が十分に進んでいないため、腫瘍の特性、詳細な発生機構は未解明であり、発生由来の特定も困難となっている。また、希少性が極めて高い疾病であるため、腹膜偽粘液腫の診断ができる医師が少なく、バイオマーカーの特定は非常に重要な課題である。腹膜偽粘液腫に対する有力なバイオマーカーとして、現在のところ、MUC2, MUC5AC [8, 9] やCEA [1] などが知られているが、これらは腹膜偽粘液腫以外の組織でも

表 1: マイクロアレイデータ

プローブ名	PMP 1	...	PMP 8	Cancer 1	Cancer 2	...
プローブ 1	42	...	35	38	34	...
プローブ 2	55	...	44	10	12	...
...	...	...	...	...	...	...
...	...	...	...	...	...	...

発現を示す場合があるため、更なる研究が求められている状況である。

2 マイクロアレイデータ

生物は遺伝子情報を基にしてタンパク質の合成などを行っており、遺伝子発現量は遺伝子の発現に必要な mRNA の量を調べることで計測が可能となる。本研究では扱うデータは網羅的遺伝子発現解析法の一つであるマイクロアレイ法を用いたデータ、すなわち、マイクロアレイデータ (SurePrint G3 Human Gene Expression 8×60K, Agilent Technologies) である。具体的には、自前で用意した 8 つの腹膜偽粘液腫の組織、虫垂、大腸、盲腸の正常組織を各 1 つの他に、公共データベースの一つである GEO[2] において GSE12470, GSE21815, GSE32062, GSE32063 をダウンロードして大腸がん、卵巣がん、胃がんなど、腹膜偽粘液腫に関係の深いと考えられる組織の遺伝子発現データも解析する。表 1 に示される通り、各行はプローブと呼ばれる遺伝子断片、各列は腹膜偽粘液腫、各種がんなどの組織に対応し、数値が発現量を表す。

表 1 において、プローブ 1 は全ての組織で高い発現を示しており、特定の組織で特異的な発現は見られない。一方で、プローブ 2 は腹膜偽粘液腫の組織でのみ特異的に発現しており、このような遺伝子を見つけることが本研究の目的となる。

本研究では、マイクロアレイデータに加えて、EST Profile, GO term を用いてバイオマーカー候補遺伝子を絞り込むことを考える。また、組織特異的に発現していると判断される遺伝子を特定すると同時に、特異度を計算して上位から順にランキングする。

2.1 GO term による候補遺伝子の絞り込み

バイオマーカーとしての使いやすさを考えると、体内の奥深くよりは尿や血液の採取によって得られる遺伝子が望ましい。そこで、GO Term と呼ばれる遺伝子 (プローブ) の特徴 (部位や機能等) を表す用語を用いて、遺伝子 (プローブ) の絞り込みを行うことを考える。

GO とは、遺伝子オントロジー (Gene Ontology) の頭文字であり、生物学的概念を記述するための共通の

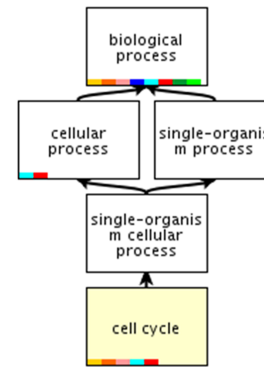


図 2: GO term の一例 (一部)

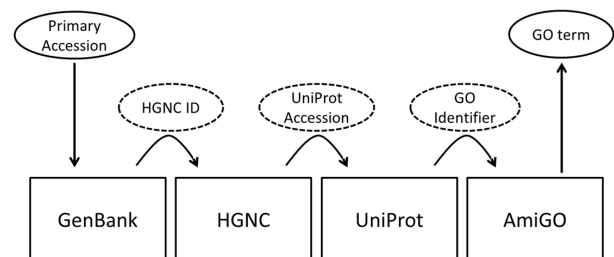


図 3: GO term に関する情報抽出

語彙を策定しようとするプロジェクトである。現代では生物学における実験手法の発達により、さまざまな種の遺伝子関連情報がデータベース化されており、これらを有効に活用するために統一された語彙が GO term である。一般に、各遺伝子 (プローブ) には複数の GO term が付けられている。図 2 は GO term の例である。

図 2 で示される通り、GO term は階層構造をもっており、子 term (下位 term) は親 term (上位 term) の情報よりも詳細な情報をもつ。したがって、本研究で特定の GO term が付加された (あるいは付加されていない) 遺伝子を選別する際には、子 term は自分自身の term だけでなく、その親 term や先祖 term も全て含んでいるとみなす。各 GO term に対して固有の ID とその親 term を記録したデータベースを作成し、GO term の階層構造を考慮した遺伝子を絞り込みを行う。

ただし、各プローブに付加されている GO term の情報については、プログラミング言語でそのまま扱えるデータをダウンロードすることはできないため、ここでは公共データベースである GenBank, HGNC, Uniprot, AmiGO の 4 つのデータベースを使ってプローブの Primary Accession から GO term を抽出している (図 3 参照)。

2.2 ESTProfile による候補遺伝子の絞り込み

ESTProfile は、あらかじめカタログ化された組織群の塩基配列内に、ある遺伝子の塩基配列が現れる数を読み取ったデータで、NCBI 内のデータベース Unigene[10] にて無償で提供されている。図 4 はヒトの遺伝子 Hs.647723 に対する ESTProfile のデータ画面を表している。

EST Profile

Hs.647723 - PRG4: Proteoglycan 4

Breakdown by Body Sites

		Hs.647723	
adipose tissue	155	2/12866	
adrenal gland	0	0/32940	
ascites	0	0/39834	
bladder	0	0/29860	
blood	0	0/122252	
bone	0	0/71618	
bone marrow	0	0/48737	
brain	3	4/1092688	
cervix	103	5/48486	
connective tissue	1106	165/149072	
ear	0	0/16100	
embryonic tissue	4	1/212896	
esophagus	0	0/20154	
eye	9	2/208840	
heart	22	2/89524	
intestine	4	1/231981	

図 4: ESTProfile のデータ画面（一部）

図 4 において、左端の列には組織の名前であり、155, 0,..., という中央付近に書かれた数字がこの遺伝子（ここでは Hs.647723）の各組織での発現レベルである。例えば、connective tissue（結合組織）での値は 1106 と高くなっている。

ESTProfile についても GO term と同様、プログラミング言語でそのまま扱えるデータをダウンロードすることはできない。そのため、あらかじめ全遺伝子に対して Unigene[10] のサイトへアクセスし、ソースから各組織に対する発現値を抽出することでデータベースを作成した。

本研究は、腹膜偽粘液腫で特異的発現をしている遺伝子の発見を目的としているため、腹膜偽粘液腫以外の組織では発現が無い、低いことが望ましい。そのため、組織群を腹膜偽粘液腫に対する関連の度合いにより次の 3 つのグループ (G1, G2, G3) に分類し、それぞれに対して閾値を設けることで絞り込みを行った。

G1 腹膜偽粘液腫の原発巣になる可能性のある組織および腹膜播種を起こしやすい原発組織（大腸（盲

腸および虫垂、結腸、直腸）、胃、卵巣の正常組織及びそれらの腫瘍）

G2 腹膜偽粘液腫の原発巣になる可能性はないが腹膜播種を起こしやすい原発組織（子宮、膵臓、肝臓、膀胱の正常組織及びそれらの腫瘍）

G3 腹膜播種と関連性が少ない原発組織（脳、乳房、肺、骨、血液およびそれらの腫瘍）

ESTProfile と GO term に基づいて遺伝子を絞り込んだ後、腹膜偽粘液腫の組織に特異的に発現している遺伝子を上位からランキングすることを考える。本研究では、エントロピーを用いて組織特異的な遺伝子を上位からランキングする ROKU[6] という手法を用いる。

3 実験結果

本研究で取り扱ったマイクロアレイデータは自前のデータと公共データの両方を用いているため、それぞれの解析で含まれているプローブの種類が異なっている。ここでは、全てのデータに共通する 15,340 個のプローブを解析の対象とした。

GO term と ESTProfile で遺伝子の絞り込みを行った後、ROKU を用いて腹膜偽粘液腫の組織で特異的に発現していると考えられる遺伝子（プローブ）のランキングを導出した。プログラミング言語として R[4] を使用した。

導出したランキングに基づいて、RT-PCR 実験を行ったところ、COL6A6 (NM_001102608), GALNT9 (NM_021808), IBSP (NM_004967) の 3 つの遺伝子で再現性が見られた。図 5 と図 6 は COL6A6n に対する RT-PCR 実験の結果である。2 つの図において下部の β -Actin はコントロールであり、上部が COL6A6 に対する結果である。

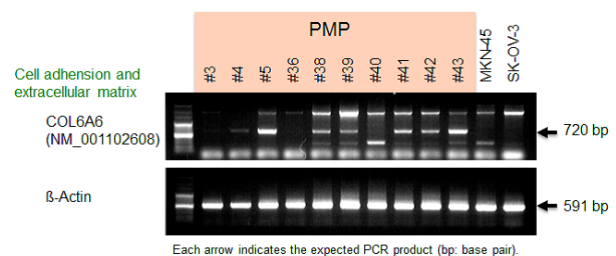


Fig. 6 Expression of COL6A6 (NM_001102608) in PMP tissues

図 5: COL6A6 (NM_001102608):腹膜偽粘液腫の組織での発現

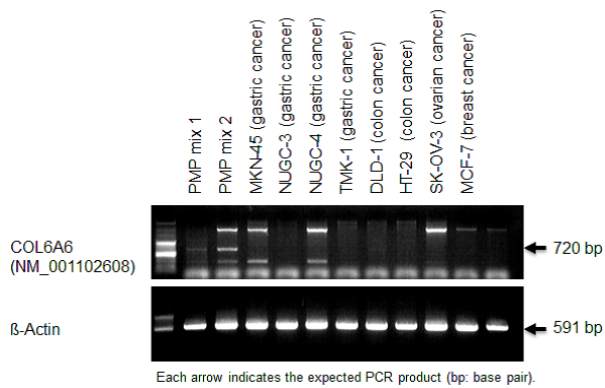


図 6: COL6A6 (NM_001102608) : がん組織での発現

図 5 において, #3, #4, #36, #39, #40, #41, #42, #43 が腹膜偽粘液腫であり, #4, #5, #38, #41, #42, #43 で高い発現が見られる. すなわち, COL6A6 は腹膜偽粘液腫の組織で明らかに高い発現が示すことがわかる. また, 図 6 から, COL6A6 は胃がん, 大腸がん, 卵巣がん, 乳がんでは発現していないことがわかる. したがって, COL6A6 は腹膜偽粘液腫でのみ発現が高い可能性があり, 腹膜偽粘液腫のバイオマーカー候補として有力な遺伝子であることが示された. GALNT9, IBSP についても腹膜偽粘液腫で発現が高く, 他のがん種で発現がほとんど見られないことがわかった.

4 おわりに

本研究では, 腹膜偽粘液腫や他の組織のマイクロアレイデータを解析して, 腹膜偽粘液腫のバイオマーカー探索を行った. GO term と ESTProfile による候補遺伝子の絞り込みを行った後, 組織特異的な遺伝子を抽出してランキングを行った. RT-PCR 実験の結果から, 腹膜偽粘液腫に対して有力な 3 つのバイオマーカー候補遺伝子を見つけることができた.

参考文献

- [1] E. Canbay, H. Ishibashi, S. Sako, A. Mizumoto, M. Hirano, M. Ichinose, N. Takao, Y. Yonemura, "Preoperative carcinoembryonic antigen level predicts prognosis in patients with pseudomyxoma peritonei treated with cytoreductive surgery and hyperthermic intraperitoneal chemotherapy," *World Journal of Surgery*, Vol. 37, No. 6, pp. 1271–1276, 2013.
- [2] Home - GEO - NCBI, <http://www.ncbi.nlm.nih.gov/geo/>.
- [3] M. Hirano, Y. Yonemura, E. Canbay, M. Ichinose, T. Togawa, T. Matsuda, N. Takao, A. Mizumoto, "Laparoscopic diagnosis and laparoscopic hyperthermic intraoperative intraperitoneal chemotherapy for pseudomyxoma peritonei detected by CT examination," *Gastroenterology Research and Practice*, Vol. 2012, 4 pages, 2012 (doi.org/10.1155/2012/741202).
- [4] R. Ihaka, R. Gentleman, "R: a language for data analysis and graphics," *Journal of Computational and Graphical Statistics*, Vol. 5, No. 3, pp. 299–314, 1996 (<http://www.R-project.org>).
- [5] 門田 幸二, トランスクリプトーム解析, 共立出版, 2014.
- [6] K. Kadota, J. Ye, Y. Nakai, T. Terada, and K. Shimizu, "ROKU: a novel method for identification of tissue-specific genes," *BMC Bioinformatics*, Vol. 7, pp. 294, 2006.
- [7] T. Kitai, T. Hirai, T. Fujita, Y. Yonemura, E. Canbay, "Survey on the incidence and management of pseudomyxoma peritonei in Japan," *癌と化学療法*, Vol. 40, No. 8, pp. 1043–1048, 2013.
- [8] J.T. O'Connell, J.S. Tomlinson, A.A. Roberts, K.F. McGonigle, S.H. Barsky, "Pseudomyxoma peritonei is a disease of MUC2-Expressing goblet cells," *American Journal of Pathology*, Vol. 161, No. 2, pp. 551–561, 2002.
- [9] 築山 吾朗, 米村 豊, 川村 泰一, 坂東 悦郎, 伊藤 寛晃, 草深 公秀, P.H. Sugarbaker, "腹膜偽粘液腫の Peritoneal Carcinomatosis Index と悪性度の評価," *日本消化器外科学会*, Vol. 40, No. 10, pp. 1739–1744, 2007.
- [10] Unigene, <http://www.ncbi.nlm.nih.gov/unigene>.

連絡先

広島大学大学院工学研究院

片桐 英樹

E-mail: katagiri-h@hiroshima-u.ac.jp

Multiple human detection based on Hog feature and Statistical learning method

○Daisuke Aoki¹, Junzo Watada²

Graduate School of Information, Production and Systems
Waseda University

E-mail: d.aoki.0330@akane.waseda.jp¹, watada@waseda.jp²

Abstract: We propose an object detection method that uses Histograms of Oriented Gradients (HOG) feature using two-stage boosting. There has been done many research works in recent years on statistical training methods and object detection methods that combine low-level features obtained. However in our proposed approach, low-level HOG features are combined by using Real AdaBoost to automatically generate features. In his paper, to evaluate the effectiveness of the proposed method, Three different experiments with different patterns are conducted for detecting humans. Moreover, a boosting classifier is used to represent the co-occurrence of the HOG features appearance for detecting a pedestrian.

I. INTRODUCTION

In this section, considering the above mentioned points, let us explain the past research works after clarifying factors that influence on human detection. The Object detection is one of the most important problems in computer vision, and many methods have already been proposed for achieving it. Image processing real technologies are required to provide convenience for people's life and to realize safety life and society. The image processing technologies include methodology of tracking human positions and recognizing human behaviors. Especially, recognizing a human position is essential in a pre-processing stage to track humans and understand human behaviors. In recent years, most object detection methods use object classification methods with statistical training methods for the selection of local features from considerable training samples. For the local features used by such methods, there are low-level features such as Haar-like features [1], Edge Orientation Histograms (EOH) [2] and Edgelet features [3]. For example, the Haar-like features have been used frequently in face detection, and it considers adjacent rectangular regions at a specific location in a detection window, sums up the pixel intensities in each region and calculates the difference between these sums. However, Haar-like features are difficult to use in human detection because the colors of clothes and the background are not always the same. So as to absorb such variations, Histogram of Oriented Gradients (HOG) [4] has been proposed as possible common information as a feature amount extracted. Among these features, the HOG has demonstrated to be robust in changes in illumination and local changes in geometry to detect a human. However, the recognition capability is limited in complex scenes with only low-level features. Therefore, recently proposed methods[5], [6] that generates features by using statistical learning algorithms, such as AdaBoost[7] [8] [9], combine low-level features based on the relations between features (relatedness) to achieve highly accurate object detection. As in late 1990s the progress of general computers technologies enables to process the huge amount of data in high speed, the main methods in image processing became to abstract high dimensional feature vectors from images and discriminate them by using statistical learning methods [10][11][12][13]. Even though the statistical learning

methods require huge learning samples with class labels, general purpose classifiers can learn without preparing previously some rules like a rule base system. In detecting faces, Neural Network method [11], SVM method [14]. Naive Bayes method [12], AdaBoost method [7] [13] are representative methods using statistical learning methods to show highly detecting performance. Especially in 2001 Viola and Jones proposed face recognition methods [7] [13] with highly accurate and high speed detection. After then, the method is widely used in object detecting research. So, the detection target was transferred to human that is more difficult to detect from face. Human detection took over the knowledge and technologies derived from several research works of the face detection, so it was approached image local features and based on statistical learning methods[15] [16]. The flow of learning classifier by statistical learning method is illustrated in Fig. 1.

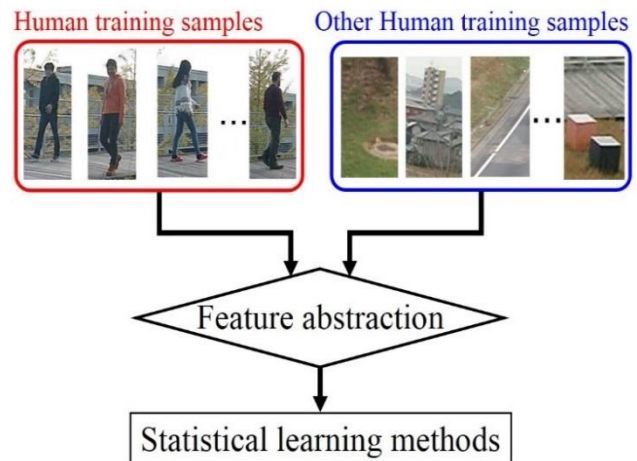


Fig. 1. Learning classifier by statistical learning methods.

However, the appearance of people image as a detection target largely changes as compared with face image to be due to clothing and body type, direction, change of perspective and change of lighting. Therefore, it is difficult to extract common feature from many people images, so accurate human detection becomes difficult. The various attempts have been done till now to absorb fluctuation of human images in processing of the classifier and the local feature image to solve this problem. However the detection target in a whole-body of person image is larger compared with a face image, since the clothing, body type, direction and change of perspective and effects of the lighting are

taken into account. Therefore, it is difficult to extract a common feature from numerous people images and achieve a high accurate human detection. Until now, there has been several attempts to solve this problem by absorbing the fluctuation of the number of humans that are visible in the processing of the classifier and the local feature.

In this paper, we proposed an approach to resolve the factors that make it difficult to detect a human. From the next section, we demonstrate the effectiveness of the proposed method experimentally, using HOG feature and Real Adaboost of three patterns.

II. METHOD

A. Histogram of Oriented Gradients

In this paper, the Histogram of Oriented Gradients (HOG) that was proposed by Dalal et al. is used for low-level features [17]. HOG features are calculated from the gradient orientations in local areas that are called cells and converted into histograms. They can capture the shape of an object. They are also robust to changes in illumination and local changes in geometry. The procedure for calculating the HOG features is described below. From the brightness I of each pixel, the gradient magnitude m and orientation θ are computed with the following equations:

$$m(x, y) = \sqrt{f_x(x, y)^2 + f_y(x, y)^2} \quad (1)$$

$$\theta(x, y) = \tan^{-1} \frac{f_y(x, y)}{f_x(x, y)} \quad (2)$$

$$f_x(x, y) = I(x + 1, y) - I(x - 1, y) \quad (3)$$

$$f_y(x, y) = I(x, y + 1) - I(x, y - 1) \quad (4)$$

The brightness gradient orientation histogram of each cell is generated from the calculation of the gradient magnitude m and orientation θ . Then, the obtained gradient orientations are divided into 20-degree groups to create a gradient orientation histogram. Finally, the features are normalized in each block area (Fig. 2) with the following equation:

$$v = \frac{v}{\sqrt{(\sum_{i=1}^k v_i^2) + \epsilon}} \quad (\epsilon = 1) \quad (5)$$

where, v is the HOG feature, k is the number of HOG features in the block, and ϵ is a coefficient for preventing the division by zero.

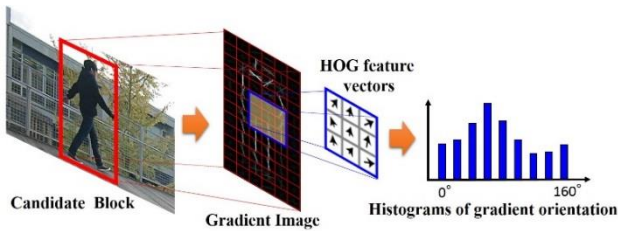


Fig. 2. HOG feature extraction.

B. Real Adaboost

Adaboost [17] realized a Boosting method with a highly accurate classifier by combining plural simple classifiers. This paper proposes an improved Real Adaboost to discriminate humans and other objects. Let us explain the detail of the algorithm in the following where M denotes the number of samples, and n the number of

learning cycles. These values are set previously. Learning data are denoted as $x_i (i = 1, \dots, n)$, i is number of learning sample and the weight of each datum is $D_t (i = 1, \dots, n)$. x_i is a sample obtained from each frame of stream. Each x_i is given with $y \in \{-1, +1\}$. y is the class label associated with a training sample x_i . That is, when x_i is a datum with a human, $y = +1$. Otherwise, $y = -1$. The probability density distributions, W_+ and W_- are represented by one-dimensional histogram that is generated from the training sample weights $D_t(i)$ with the following equations:

$$W_+^j = \sum_{i: j \in J \wedge y_i = +1} D_t(i) \quad (6)$$

$$W_-^j = \sum_{i: j \in J \wedge y_i = -1} D_t(i) \quad (7)$$

where, t is the number of training cycles, i is the number of training samples, j is the BIN number of the one-dimensional histogram, and y_i is the class label where $y_i \in \{+1, -1\}$. D_t is calculated in the same manner as the first-stage Real AdaBoost. The probability density distributions W_+^j and W_-^j can be generated by calculating the training samples x_i features and applying the training sample weights $D_t(i)$ to the BIN numbers j of the one-dimensional histograms that correspond to these feature values (Fig. 3).

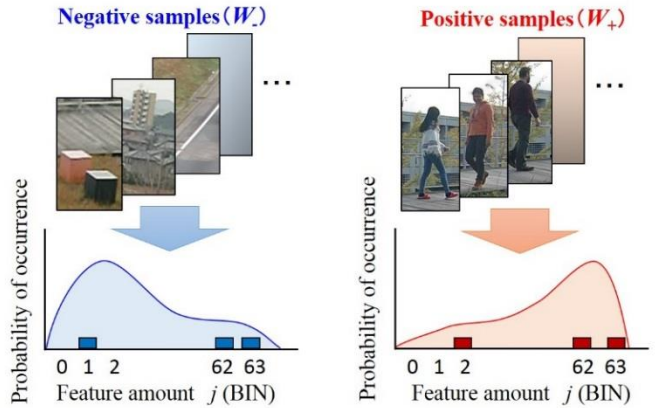


Fig. 3. Histogram construction of positive and negative features.

These probability density distributions are used to obtain the value of z , which represents the degree of separation between the distributions. Evaluation value:

$$z = 1 - \sum_j \sqrt{W_+^j W_-^j} \quad (8)$$

A small value of z indicates the separation amount between the positive and negative class distribution. The smaller the value is the larger the separation is. Therefore, the smallest value of z is used to select the weakest classifier at each cycle (Fig. 4).

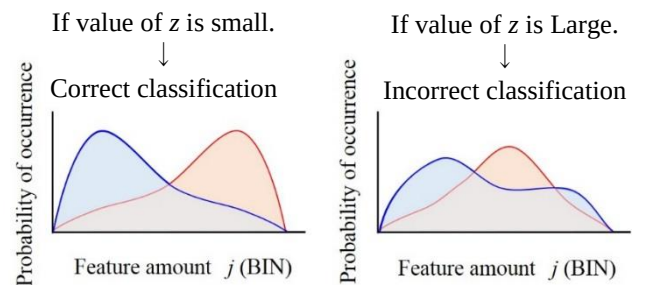


Fig. 4. Feature of classification.

Next, the resulting probability density distribution from W_+^j and W_-^j is used to calculate the output $h_t(x)$ in the second-stage Real AdaBoost previously selected classifier. From the HOG features values, the co-occurrence j can be obtained from the training samples. Then, the classifier output $h_t(x)$ is calculated from one-dimensional histogram BIN numbers j and their corresponding probability density distributions W_+^j and W_-^j by using the following equation:

$$h_t(x) = \frac{1}{2} \ln \frac{W_+^j + \epsilon}{W_-^j + \epsilon} \quad (9)$$

where, x is a serial number that represents the number of combinations of cells, ϵ is a coefficient for preventing the division by zero and takes the same value that was used in the Real AdaBoost, $\epsilon = 0.0000001$. The last classifier H is considered as the stronger classifier H is calculated:

$$H(x) = \text{sign} \left[\sum_{t=1}^T h_t(x) - \theta \right] \quad (10)$$

where θ is a threshold whose default value is zero. After determining the value of $H(x)$, if $H(x) \geq \theta$, then the input data x indicates the presence of a person in the image. Furthermore, if $H(x) \leq \theta$, then the input data x indicates the absence of a person in the image. An important feature in AdaBoost is its learning technique that is realized by increasing the value of the sampling weights in the erroneous data so the subsequent learning can be focused on the previous erroneous data.

The detected human can be visualized by using a rectangular pattern in the image called raster scanning window. After processing the image by AdaBoost, several number of raster scanning window are generated per individual. The overlapped windows indicate the zone or area where a human can be located. Therefore the Mean Shift is proposed to integrate these windows and have a positive detection of a human.

C. Mean Shift

In this section, a comprehensive overview of the basic theory and mean shift applications in object tracking is provided.

Mean shift is a popular robust framework for statistical data analysis using kernel density estimation that was originally proposed by Fukunaga and Hostetler in 1975 [18]. In this paper, the Mean shift is used for the integration of all the windows generated by the raster scanning. In the Mean shift, the density function maxima is located by sampling discrete data from the given function [19]. Moreover, it helps on the detection of the density modes. Let be x the initial estimate at the start of the iteration. The kernel function is given by $K \|x - x_i\|$. The weight of the adjacent points are calculated with this function by re-estimating the mean. Then a Gaussian kernel is used to determine the distance of the current estimate with the following equation:

$$K \|x - x_i\| = e^{-c \|x - x_i\|^2} \quad (11)$$

where K determines the weighted mean of the density.

Mean Shift Vector:

$$m(x) = \frac{\sum_{i=1}^n x_i K \left(\left\| \frac{x - x_i}{h} \right\|^2 \right)}{\sum_{i=1}^n x_i \left(\left\| \frac{x - x_i}{h} \right\|^2 \right)} - x \quad (12)$$

Kernel function:

$$K(x) = \begin{cases} c(1 - \|x\|) & \|x\| < 1 \\ 0 & \text{otherwise} \end{cases} \quad (13)$$

There is important to decide radius h using Mean shift. The optional parameter is only the radius h of the circle area, even though that it can be expanded the same radius of the sphere in the case of high-dimensional as shown in Fig. 5.

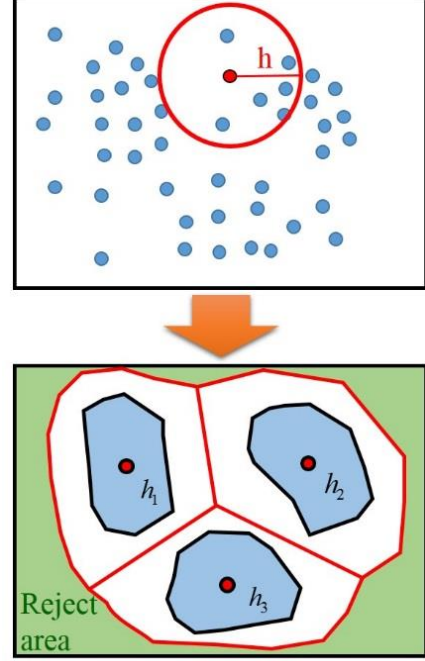


Fig. 5. Integration of the detection window

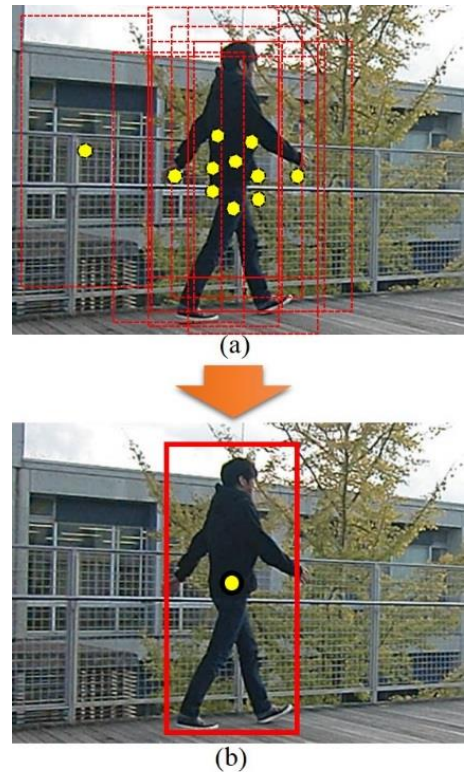


Fig. 6. Integration of the detection window.

D. Real-Adaboost with Mean shift

Fig.7 shows the detection results with a window raster scanning over the image. Fig. 6(a) shows the multiple windows in which the target was detected at multiple times at different scales. Then, in Fig. 6(b) all the windows are integrated by using Mean Shift clustering [20]. We can observed from Fig. 6 that our method can detect individual objects even if they are overlapped.

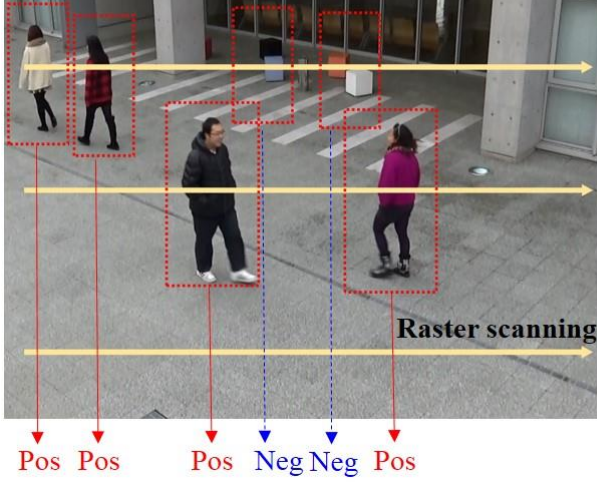


Fig. 7. Raster scanning of the detection window.

III. RESULTS

To test the performance of our proposed method, we conducted the following experiments where the detection of humans is involved.

A. Data Source

The dataset used for our experiment was compiled with the images taken with different individuals and different backgrounds. The individuals who participated in this experiment, differ from one image to another, so not only the body size differs but also the clothing, illumination and background. The images that are used in our database is shown in Fig. 8. One of the most popular database use for training is the Evaluation Dataset INRIA [21]. These dataset contains a total of 1774 images where 1671 images do not include only humans but also objects in the background. The INRIA dataset includes various kinds of samples from different points of view, colorful clothing, diverse poses, partial occlusion, sundry illumination and locations. A negative samples is considered when the random objects that appears in the background are selected. However the if the humans in the image are selected without any error, then it is considered as a positive sample. The database used for training contains 2000 positive and 5000 negative samples.

For our experiment, an OpenCV 2.4.9 and Windows 8 Intel® Core™ i7 4790 M CPU (3.60 GHz) computer was used.



Fig. 8. Human database.

B. Experiments

The image dataset for the evaluation of our human detection experiment was taken with a Video camera with a 12.1 Mega pixels resolution. The parameters for these experiments are listed in Table 1.

Table 1. Experiment parameter.

	Human
Orientations	9(0°~180°)
Patch size (pixels)	64×128
Cell size (pixels)	8×8
Block size (cells)	2×2
Cells	64

We take three different movies that involve the different subjects, background and camera angles as shown in Fig. 9. The Patterns that are considered are explained as follows:

- Pattern1: In this pattern a more simple background (not objects) is considered as shown in Fig. 9(a) & (b).
- Pattern2: In this pattern, a more complex background is considered. In this background objects as trees and fences are interact with the individual as shown in Fig. 9(c) & (d).
- Pattern3: In Pattern 3 not only objects interact with the individual but also the human size changes according with the perspective as shown in Fig. 9(e) & (f).

C. Experimental Results

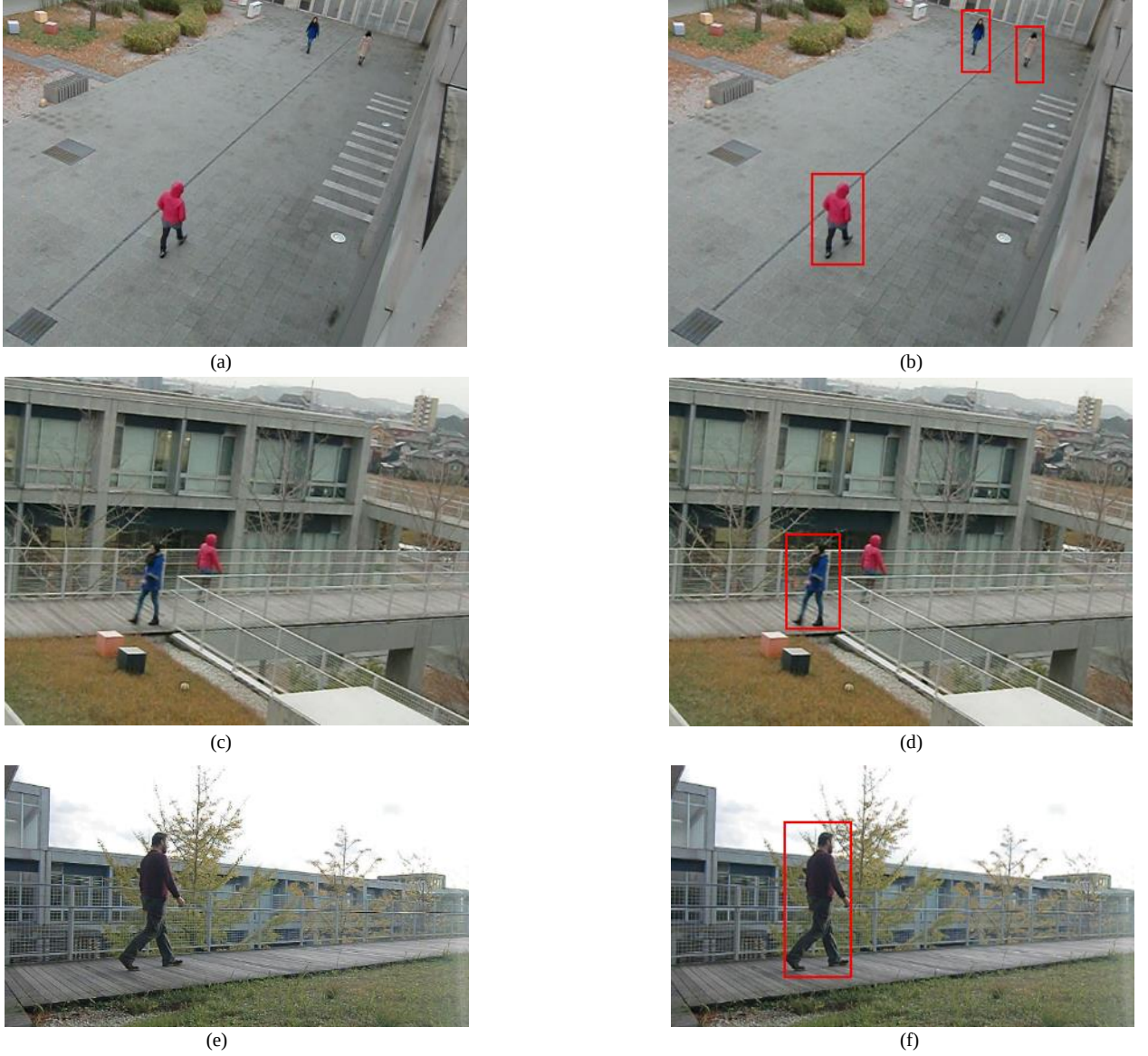


Fig. 9. Different results of human Detection.
(a)Simple Background image, (b) Detection Result of (a),
(c)Complex Background image, (d) Detection Result of (c),
(e)Variable human size image, (f) Detection Result of (e).

The detection error for Fig. 9(b), (d), (f) are shown in Table 2.

Table 2. Detection Result

	Pattern1	Pattern2	Pattern3
Detection Error[%]	19.0	49.5	38.7

The proposed method can detect a single individual with simple background our multiple individuals where the interaction with objects in the background is not involved. However, if the objects in the background had some interaction or overlap the human for its analysis, the classifier would consider as an object. Therefore this will be considered as a detection error. In the Pattern 3, the error rate is 38.7 %. That reason is the scanner window size is predetermined, so when the human body size change respective with the perspective angle the window does not change.

IV. CONCLUSIONS

In this paper we proposed an improved object detection method that combines Histograms of Oriented Gradients (HOG) feature with two-stage boosting and Real Adaboost to automatically generate features. Therefore, HOG low-level features obtained can be used traditional statistical training methods. To evaluate the effectiveness of our proposed method we conducted three different experiments. As shown in the results, we conclude that a lot of correct samples is required in human detection. Moreover, our method can overcome the disadvantages of the statistical methods as the situation demands. In future work, we plan to accomplish high accuracy human detection in images.

REFERENCES

- [1] Papageorgiou, C. P., Oren, M. and Poggio, T. "A general framework for object detection", IEEE ICCV, pp. 555–562 (1998).
- [2] K. Levi and Y. Weiss, "Learning Object Detection from a Small Number of Examples: the Importance of Good Features.", IEEE CVPR, 2, pp. 53-60 (2004).
- [3] B. Wu and R. Nevatia, "Detection of Multiple, Partially Occluded Humans in a Single Image by Bayesian Combination of Edgelet Part Detectors", IEEE ICCV, 1, pp. 90-97 (2005).
- [4] Dalal, N. and Triggs, B. "Histograms of Oriented Gradients for Human Detection", IEEE CVPR, pp. 886–893 (2005).
- [5] P. Sabzmeydani and G. Mori, "Detecting Pedestrians by Learning Shapelet Feature", IEEE CVPR, pp. 511-518 (2007).
- [6] Mita, T., Kaneko, T., Stenger, B., and Hori, O. "Discriminative Feature Co-Occurrence Selection for Object Detection", IEEE PAMI, 30, (7), pp. 1257-1269, (2008).
- [7] Freund, Y. and Schapier, R. E. "Experiments with a new boosting algorithm", Machine Learning, (1996).
- [8] P. Viola, et al., "Rapid Object Detection Using a Boosted Cascade of Simple Features," CVPR, pp. 511–518, (2001).
- [9] P. Viola, et al., "Robust Real-Time Object Detection," IJCV, 57, (2) pp. 137–154, (2004).
- [10] K.-K. Sung, et al., "Example-Based Learning for View-Based Human Face Detection," Technical Report MIT AILab, (1994).
- [11] H.A. Rowley, et al., "Neural Network-Based Face Detection," CVPR, pp. 203–208, (1996).
- [12] H. Schneiderman, et al., "A Statistical Method for 3D Object Detection Applied to Faces and Cars," CVPR, (2000).
- [13] P. Viola, et al., "Rapid Object Detection Using a Boosted Cascade of Simple Features," CVPR, pp. 511–518, (2001).
- [14] E. Osuna, et al., "Training Support Vector Machines: an Application to Face Detection," CVPR, pp. 130–136, (1997).
- [15] C. Papageorgiou, et al., "A Trainable System for Object Detection" IJCV, 38, (1) pp. 15-33, (2000).
- [16] P. Viola, et al., "Detecting Pedestrians Using Patterns of Motion and Appearance," ICCV, pp. 734 - 741, (2003).
- [17] Freund Y., and Schapire R. E., "A decision theoretic generalization of on-line learning and an application to boosting", Journal of Comp. and System Sci, 55, pp. 119-139, (1997).
- [18] K. Fukunaga and L. Hostetler. The estimation of the gradient of a density function, with applications in pattern recognition. IEEE Trans. Information Theory, 21(1):32–40, (1975).
- [19] Cheng, Yizong "Mean Shift, Mode Seeking, and Clustering". IEEE Transactions on Pattern Analysis and Machine Intelligence (IEEE) 17 (8): 790–799. (1995).
- [20] D. Comaniciu, P. Meer. "Mean Shift: A Robust Approach toward Feature Space Analysis", IEEE PAMI, vol. 24, (5), pp. 603-619, (2002).
- [21] Evaluation Dataset INRIA, <http://lear.inrialpes.fr/data>.

Brainwave and Speech Emotion Recognition : Survey

Kitagawa Hanayuki¹, Junzo Watada²

^{1, 2} Graduate School of Information Production and System, Waseda University

¹kitagawa.hanayuki@yahoo.co.id, ² junzow@osb.att.ne.jp

Abstract : Brainwaves are predicted with links between spatial distribution of synaptic activity and temporal frequencies in our body. The electroencephalogram (EEG) signal can be used for the purpose of recognizing speech emotion. The scalp recorded EEG provides the estimation of synaptic action at the large scales closely related to behavior and cognition. The emotion classification system can be used to recognize the support states. In this survey, the most research works related will be summarized. The discussion in this paper, draws the conclusion given with some clues for the future works.

1. Background & Motivation

Research in speech emotion recognition is improving in the field of human computer interaction. Emotion recognition is an important research area and plays a pivotal role in handling information of emotion people can communicate with each other more smoothly. The emotion recognition systems in speech or facial expressions include several emotional states such as joy, fear, sadness, disgust, anger, surprise and neutral [1]. In recent years, electroencephalogram (EEG) research has proven to be useful for a multitude of new methods of communication such as emotion recognition besides the well known clinical applications[2]. EEG uses an electrical activity of the neurons inside the brain. Some areas of the brain are related with the human emotional behavior: brainstem, hypothalamus, thalamus, prefrontal area and limbic system [3]. EEG based technology has become more popular in games design, EEG-based technologies were used in neurofeedback games and brain computer interfaces [4].

2. Theory of Brainwave Signals and Related Technique

The emotion recognition system based on brain signal or EEG can be applied to many areas such as education, medicine, gaming, entertainment and intelligent tutoring system, an example of EEG based technology in the field of entertainment is music player. In this application, the current emotional state of the user is identified, and a music related to this state is played. The songs from the music player are clasified into six emotion: sad, frustated, fear, happy, pleasant, and satisfied[3].



Figure 1. The mindset device [5]

The device gives different brainwaves strengths alpha, gamma and delta along with the meditation and attention calculated states. Brainwaves with fewer frequencies than alpha brainwaves are named slow brainwaves and the higher frequencies than alpha brainwaves are fast brainwaves [6]. From all of those data with data analysis we can try to guess the emotion of the user. The figure 2 below mentioned waveforms when the person falls asleep, the dominant EEG frequency decreases.

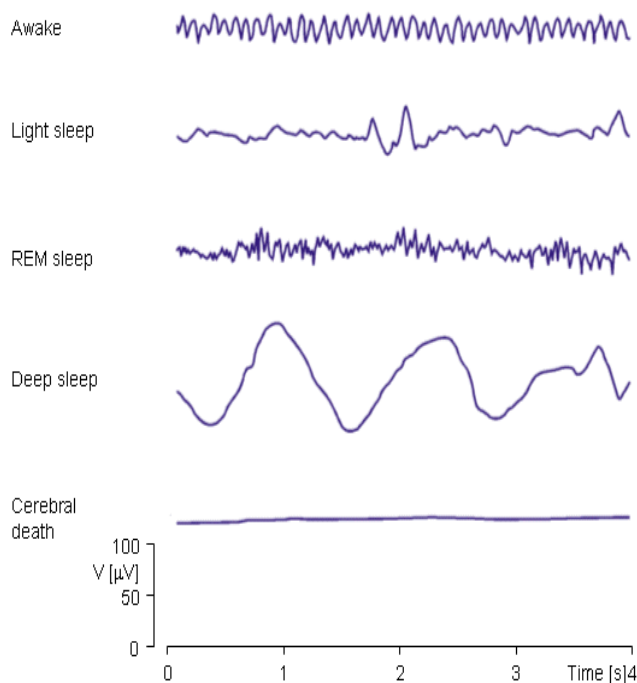


Figure 2. EEG activity is dependent on the level of consciousness[7] .

In a certain phase of sleep, rapid eye movement called (REM) sleep, the person dreams and has active movements of the eyes, which can be seen as a characteristic EEG signal. In deep sleep, the EEG has large and slow deflections called delta waves. No cerebral activity can be detected from a patient with complete cerebral death. As the activity increases, the EEG shifts to higher dominating frequency and lower amplitude.

When the eyes are closed, the alpha waves begin to dominate the EEG[7].

Table 1. The Brain states associated with different frequency bands

Type	Freq. (Hz)	Amp (μV)	Brain States
δ	0.1 – 0.3	100-200	Deepest, dreamless sleep, unconscious state.
θ	4 – 7.5	< 30	REM sleep, dreaming, intuitivity.
α	8 – 12	30 – 50	Relaxed, but non-sleepy state.
β	13 – 30	< 20	Sensory and emotional influences, thinking, alertness.

The table above lists the relationship of frequency spectrum of EEG and brain states. The profile of the spectrum can indicate specific brain state. The brain wave is changing all the time. It depends on the age, sex, emotion, activity and so on[8].

3. Theory of Speech Recognition

Research in emotion recognition by speech or image remains one of the important research challenge. The various types of speech recognition can be indicated as [9]:

- **Connected Words** : It accepts single word or utterance at a time and allows a proper pause between the utterances.
- **Continuous Speech** : The users speaks smoothly and with the help of methods the utterances is calculated.

- Isolated Words : Accepts single utterance and word at a time

The performance of speech recognition systems is usually specified in terms of accuracy and speed.

4. Method and Experiment

EEG database can be used for the source of brain signal, and the database will be recorded by using music videos to induce emotions in the participants of the experiment[10]. The purpose of this paper is to investigate whether it is possible to reliably recognize speech based on EEG signals. The figure below shows EEG signals, it is possible to differentiate alpha (α), beta (β), delta (δ), and theta (Θ) waves as well as spikes associated with epilepsy.

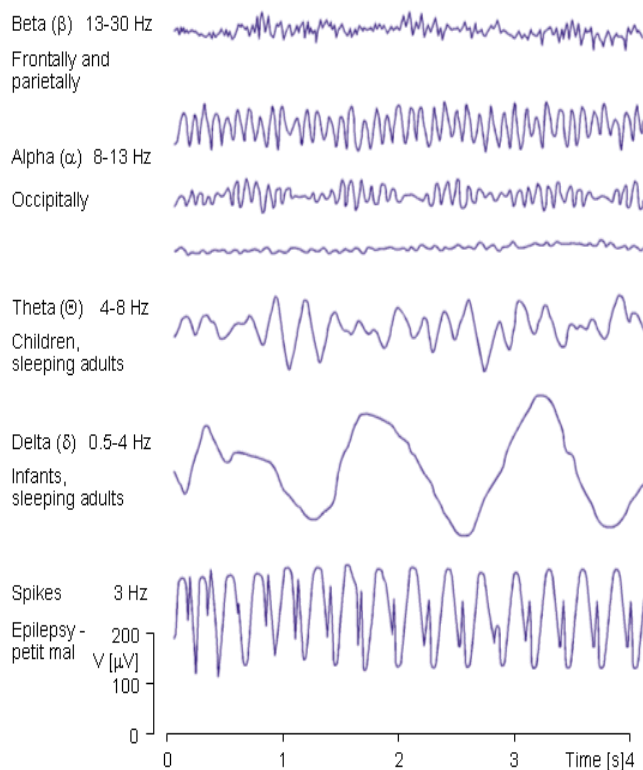


Figure 3. Some examples of EEG waves [7]

The alpha waves have the frequency spectrum of 8-13 Hz and can be measured from the occipital region in an awake person when the eyes are closed. The frequency band of the beta waves is 13-30 Hz; these are detectable over the parietal and frontal lobes. The delta waves have the frequency range of 0.5-4 Hz and are detectable in infants and sleeping adults. The theta waves have the frequency range of 4-8 Hz and are obtained from children and sleeping adults. The EEG signals are basically and closely related to the level of consciousness of a person. When the eyes of the subject are closed then alpha waves begin to have a very strong influence over EEG [11].

5. Conclusions

Brainwave will be used to recognize speech emotions. The survey evaluated and attempted to emerge with the EEG signal for detection. The future research on the speech emotion recognition and brainwave signal can benefit for many technologies.

References

- [1]. Kazuhiko Takahashi, "Remarks on Emotion Recognition from Bio-Potential Signals" 2nd International Conference on Autonomous Robots and Agents, Palmerston North, New Zealand, 13-15, 2004.
 - [2] Porbadnigk, Anne and Wester, Marek and Jan-p Calliess, Tanja Schultz, "EEG-based speech recognition impact of temporal effects", 2009.
 - [3] Taciana Saad Rached and Angelo Perkusich, Emotion Recognition Based on Brain-Computer Interface Systems, <http://dx.doi.org/10.5772/56227>, chapter 13, 2013.
- Accessed on January 2015

- [4] Qiang Wang, Olga Sourina, and Minh Khoa Nguyen, EEG-Based “ Serious” Games Design for Medical Applications, International Conference on Cyberworlds, DOI 10.1109/CW.2010.56, IEEE computer society, 2010.
- [5] Mindwave device,
<http://store.neurosky.com/products/mindwave-1>, Accessed on January 2015
- [6] Junzo Watada, Masato Takagi, Naoyoshi Yubazaki Hideyasu Hirano,; Realization of the Comfortable Space using Brainwave-Signals,; Journal of Systems and Control Engineering (JSCE), vol. 22, no. 8, pp.667-673 (2006.12).
- [7]Electroencephalography
<http://www.bem.fi/book/13/13.htm>
Accessed on January 2015
- [8] Yung-Chin Hsiao, Junzo Watada, Hanayuki Kitagawa,; Measurement of Attentions for Auto-lighting System by Brainwave and Eye Tracking: Survey, International Society of Management Engineers (ISME) , July 2014.
- [9]Priyanka Abhang, Shashibala Rao, Bharti W. Gawal, Pramod Rokade, Emotion Recognition using Speech and EEG Signal – A Review, International Journal of Computer Applications (0975 – 8887), vol. 15, no.3, February 2011.
- [10] DEAPdataset: a dataset for emotion analysis using eegphysiological and video signals.
<http://www.eecs.qmul.ac.uk/mmv/datasets/deap/>
Accessed on January 2015
- [11] Rafal Bogacz, Urszula Markowska-Kaczmar, Androzej Kozik, “ Blinking Artefact Recognition in EEG signal using Artificial Neural Network”, Computer Science and Management Department, Wroclaw, 2002.

ワイナリー・システムの最適化のための二層モデル

A Bilevel Model for Optimizing Winery System

○うかいぎょく¹, 和多田淳三², 李鏡儒³

○Yu Haiyu¹, Watada Junzo, Li Jingru

^{1 2 3}早稲田大学大学院情報生産システム研究科

^{1 2 3}Graduate School of Information, Production and System, Waseda University

Abstract: Recent research works show an increasing trend in analyzing agriculture problem because of its good applicability and usability. Also, the rapid development of mathematical model mitigates the difficulty to solve such abstract and uncertain problem. In this paper, the questions are discussed about a wine system, which contains some grape blocks and winery together. In the paper, a bilevel model is developed for obtaining the balance between profits and production quality.

1. Introduction

Recent research results show an increasing trend in analyzing agriculture problem because of its good applicability and usability. In this paper, wine related issues are well discussed and a profitable model is developed based on the provided background. In the wine industry, as grapes are most characterized by its perishability after harvested, it is a standard practice for wine companies to develop a system which contains the grape blocks and winery together to reduce the loss in transportation process. In the harvesting process, the grape blocks of the same wine company are interrelated and interdependent, which enables the company to get the best quality ingredients to produce wine by designing a favorable harvesting route among the blocks. That is to say, as the harvest dates differ in blocks, then a favorable harvesting route will contribute to obtain the optimum quality in the whole system. As a result, the promising wine rating can be effected by the harvesting route (the high rating profits from the appropriate harvesting date), and also benefited from human intervention if necessary. A wine unit is defined as “case” in actual production process. For standard bottle (750ml), one case contains 12 bottles. The relationship between profits and quality is not only association but also confliction. Since it is standard practice for wine companies to develop a system which contains the grape blocks and winery together, the set background for this paper is that a wine system contains two wineries as well as several blocks. Although several determining factors influence on wine

quality as elaborated in the review part, for the blocks in the same area, harvest route among blocks still plays the determinant role. Usually, a good harvest scheduling among the blocks is due to improve the wine quality and reduce labor costs in the meanwhile. However, when taking the capacity of winery into consideration, it is unwarranted to harvest all the grapes totally on the optimal harvest date. In addition, if severe weather comes without expectation, more serious problems will ensue. The core issue is how to strike a balance of contradiction by using mathematical methods. In the early 1960s, an optimum mathematical model called bilevel programming came studied. Bilevel optimization is a special kind of optimization where one problem is embedded (nested) within another. If the amount of profits is assumed as the upper level, while the wine quality accounts for the lower level, then the balanced result can be calculated by using mathematical method. The paper is structured as follows. In Section 2 we state the formulation of the bilevel model and discuss how the harvest route is realized. In Section 3 the application is worked out based on the bilevel model. At last, the conclusion will be drawn on after calculation.

2. Bilevel formulation

The general form of bilevel programming problem can be mathematically stated as

$$\min_{\mathbf{x}} F(\mathbf{x}, \mathbf{y})$$

$$s.t. \quad (\mathbf{x}, \mathbf{y}) \leq 0$$

$$\min_y f(\mathbf{x}, \mathbf{y})$$

$$s. t. \quad \mathbf{g}(\mathbf{x}, \mathbf{y}) \leq 0,$$

where $\mathbf{x} \in \mathbb{R}^{n1}$ and $\mathbf{y} \in \mathbb{R}^{n2}$ are respectively the upper level variables and lower level variables. Similarly, the function $F : \mathbb{R}^{n1} \times \mathbb{R}^{n2} \rightarrow \mathbb{R}$ is the upper level objective function and $f : \mathbb{R}^{n1} \times \mathbb{R}^{n2} \rightarrow \mathbb{R}$ is the lower level objective function, while the vector-valued functions $\mathbf{G} : \mathbb{R}^{n1} \times \mathbb{R}^{n2} \rightarrow \mathbb{R}^{m1}$ and $\mathbf{g} : \mathbb{R}^{n1} \times \mathbb{R}^{n2} \rightarrow \mathbb{R}^{m2}$ are respectively called the upper level and lower level constraints. As discussed in the introduction part, the goal functions of upper level and lower level are respectively to maximize the profits and maximize the quality of harvested grapes. For the lower level, the quality of wine is measured by the deviation from the optimum date and the penalty for harvest delay. There is a clear agreement that some deterioration is produced if we harvest ahead of time. Thus, delaying harvesting seems to be less dangerous, but it exposes the crop to an increasing danger of rain (Bordeu et al 2002) [1]

3 . Application model

Based on the description in Section 2, we make use of the following notation described in Table 1.

Table 1. Definition of Notation Involved in the Bilevel Model

1.1 sets of indexes	
B_1	set of blocks suitable for mechanical harvest
B_2	set of blocks suitable for manual harvest $B = B_1 \cup B_2$
k	set of harvest modes, $k \in K = \{1,2\}$
$k = 1$	is manual harvest
$k = 2$	is mechanical harvest

1.2 model variables

M_r set of actual amount of r rating grape harvested one day, $r \in R = \{1,2,3\}$

x_{jtkb} kilograms harvest in block j , at period t , in mode k , routed to winery b

y_{jtkb} binary variable with value 1 when block j initiate harvesting at period t , in harvest mode k , routed to winery b , and 0 otherwise.

D_{ib} the actual harvest n th date in harvest period when the quality is the best in block i and with destination winery b .

v_{jtk} binary variable indicating if there is harvest or not in block j , during period t , and in harvest mode k

z_{ijt} binary variable that is 1 when an operation moves from block i to block j in period t and with destination winery b , 0 otherwise

τ_{jtb} variable used to eliminate subtours in the MTZ formulation. It presents the position inside the cycle of block j , at period t , routed to winery b

w_t^H number of workers hired at period t

w_t^F number of workers hired at period t

u_{jtb} productive resources used in block j in period t .

It is a continuous variable such that for $k = 1$ it is machine hours and for $k = 2$ it is number of workers

1.3 model parameters

p_r price for different wine ratings, $r \in R = \{1,2,3\}$

α_r wine juice yield for rating r , $r \in R$

β_r maximum allowable production deviation rate from planning amount each day; for rating r

c_m cost of initial equipment investment for different factory scale m , $m \in M = \{1,2,3,4,5\}$

c_k unit cost of productive resource k . For $k = 1$ it is cost per hour of machine and for $k = 2$ it is cost per worker

H hiring cost of adding one unit of labor

F saving cost for reducing one unit of labor

T days available to produce wine among one year

Q_{jt}	quality cost for block $j \in B$ in period t
D_{ij}	distance between block i and $j, i, j \in B$
P_{kj}	productivity of a unit of productive resource k in kg/day, $j \in B$
G_{jkb}	estimated kg of grapes in block $j \in B$ harvest in mode $k \in K$, with destination to winery b
W_{jt}	harvest window parameter indicating whether block j can be harvested in period t . It is 1 if possible, and 0 otherwise
L_{kbt}	processing capacity of winery b for grapes harvest in mode k in period t
B_{jtk}	maximum volume that can be harvested in block j at time t using harvested mode k
d_{ib}	the n th date in harvest period when the quality is the best in block i and with destination winery b .
A_t	machine availability, in hours, in period t
M	a “big M ” value
M_{Tr}	total planning grape production amount each year of rating r
R	operations relocation cost in \$ per kilometer
V_k	minimum volume to be harvested from a block using k harvest mode
λ	penalty parameter used for the relative balance between operational and quality cost
N	minimum number of contractors needed
$card(B)$	number of elements in the set B , number of blocks

Based on the previously defined parameters, the whole model is

upper level

$$\begin{aligned}
 \max_{M_r, w_t^H, w_t^F, u_{jtk}, v_{jtk}, z_{ijtb}, \tau_{jtb}} & \sum_{r=1,2,3} p_r M_r \alpha_r - \sum_{k,j,t} c_k u_{jtk} - H \sum_t w_t^H \\
 & - F \sum_t w_t^F - R \sum_{i,j,t,b:i \neq j} D_{ij} z_{ijtb} - c_m \\
 s. t. & (1 - \beta_r) \frac{M_T}{T} \leq M_r \leq (1 + \beta_r) M_r
 \end{aligned}$$

$$c_m = \begin{cases} c_1, & M_T = 2000 case \\ c_2, & M_T = 5000 case \\ c_3, & M_T = 10000 case \\ c_4, & M_T = 15000 case \\ c_5, & M_T = 20000 case \end{cases}$$

R_r

$$= \begin{cases} R_1 & (d_{ib} - D_{ib})^2 = 0, p_1 = 3.229 \times 10^2 \$/750ml \\ R_2 & 0 < (d_{ib} - D_{ib})^2 \leq 784, p_2 = 2.021 \times 10^2 \$/750ml \\ R_3 & 784 < (d_{ib} - D_{ib})^2 \leq 1521, p_3 = 1.179 \times 10^2 \$/750ml \end{cases}$$

$$\sum_t \sum_k v_{jtk} \geq 1,$$

$$\sum_{j \in B_2} u_{jt2} = \sum_{j \in B_2} u_{j(t-1)2} + w_t^H - w_t^F,$$

$$\sum_{j \in B_1} u_{jt1} \leq A_t,$$

$$u_{jt2} \geq N v_{jt2},$$

lower level

$$\min_{x_{jtkb}, y_{jtkb}, D_{itb}} \sum_i (d_{ib} - D_{ib})^2 + \lambda \sum_{j,t,k,b} Q_{jt} x_{jtkb}$$

$$s. t. \quad \tau_{jtb} \geq \tau_{itb} + 1 - (card(B) - 1)(1 - z_{ijtb}),$$

$$\forall t, b; i, j \in B_2; i \neq j; W_{it} \neq 0; W_{jt} \neq 0,$$

$$\sum_{j \neq i} z_{ijtb} = y_{iy2b}, \forall t, b, j \in B: W_{it} \neq 0; W_{jt} \neq 0,$$

$$\sum_{i \neq j} z_{ijtb} = y_{it2b}, \forall t, b, i \in B: W_{it} \neq 0,$$

$$\tau_{1tb} = y_{1t2b}, \forall t: W_{1t} \neq 0, \forall b,$$

$$\tau_{jtb} \geq 2y_{jt2b}, \forall j, t \in B: W_{jt} \neq 0, \forall b,$$

$$\tau_{jtb} \leq card(B), \forall j, t \in B: W_{jt} \neq 0, \forall b$$

$$z_{ijtb}, v_{jtk}, y_{jtkb} \in \{0,1\}, \forall t, b, k, i, j$$

$$w_{tj}^H, w_{ti}^F, u_{jtk}, x_{jtkb}, \tau_{jtb} \geq 0, \forall t, b, k, i, j$$

$$\sum_{j \in B} x_{jtkb} \leq L_{kbt}, \forall k, b, t,$$

$$\sum_t x_{jtkb} = G_{jkb}, \forall k, b, j,$$

$$x_{jtkb} \leq G_{jkb} v_{jtk}, \forall t, b, k, j,$$

$$x_{jtkb} \geq V_k v_{jtk}, \forall t, b, k, j,$$

$$\sum_b x_{jtkb} \leq B_{jtk},$$

$$\sum_b y_{jtkb} \leq M v_{jtk},$$

$$v_{jtk} \leq \sum_{b, s \leq t: W_{js}=1} y_{jskb},$$

$$y_{jtkb} \leq \sum_{s \leq t: W_{js}=1} v_{jsb},$$

$$\sum_b \sum_{t: W_{jt}=1} y_{jtkb} \leq 1,$$

$$y_{jtkb} \leq W_{jt},$$

$$\sum_b x_{jtkb} \leq P_{kj} u_{jtk},$$

Fenologica: Experiencia en Chile. Topicos de Actualizacion en Viticultura y Enologia 2002, 190-192

連絡先

北九州市若松区ひびきの2-7-S261

E-mail: yuhaiyu@ruri.waseda.jp

In the model, the harvest route is in the constraint part of the lower level while the promised grape rating criteria is in the constraint part of upper level. The first constraint in upper level also gives the limit for daily production limit according to the annual production planning.

4. Conclusions

Bilevel model was applied in this paper to work out the optimum result when considering with profits and wine quality. In the model presented, several methods such as principal component analysis, routing strategy were efficiently used to calculate some portions. Upper level and lower level are well connected by routing strategy and the evaluation method of wine quality. The final result shows the real routing and the maximum profits this year in the system.

references

[1] Bordeu, E., Mira, P., Rivadeneria, R., 2002, Madurez

ファジィ分割表を応用したメディア授業におけるニーズ分析

Needs Analysis for Media Lectures Applying Fuzzy Contingency Table

上江洲 弘明
Hiroaki Uesu
早稲田大学
Waseda University

Abstract: With development of the Internet, media lessons has also developed greatly. Recently, student's needs is increasing for media lessons, we have to satisfy student's needs. In this paper, we propose a needs analysis applying Kano model.

1. はじめに

パソコンやタブレット、大容量インターネット通信の普及により、メディア授業に対する大学生の認識が大きく変わってきた。よりよいメディア授業を展開するためには、教材や教育技術の評価の他に学生が授業に対して感じる評価を定量的に把握することが重要となる。

今日、メディア授業に対するニーズは、教授法や教材が良ければ満足を感じ、悪ければ不満をもつという一元的なものだけではない。学生のニーズの中には、フィードバックやさまざまな形態の課題など求めているものは数多くある。本論文では、このようなニーズに対応するため、ファジィ分割表を応用したアンケート分析を提案する。

2. ファジィ分割表

定義1. ファジィ集合の濃度

全体集合を $U = \{x_i \mid i = 1, \dots, n\}$ とする。このとき、 U 上のファジィ集合 A の濃度 $|A|$ を以下のように定義する。

$$|A| = \sum_{i=1}^n \mu_A(x_i)$$

ただし、 $\mu_A(x)$ はファジィ集合 A のメンバーシップ関数とする。

定義2. ファジィ $m \times n$ 分割表

全体集合 $U = \{x_i \mid i = 1, \dots, n\}$ とする。 U 上のファジィ集合 $A_1, \dots, A_n, B_1, \dots, B_m$ の $m \times n$ ファジィ分割表を以下のように定義する。

	A_1	$\dots$	A_n	Sum
B_1	f_{11}	$\dots$	f_{1n}	$ B_1 $
$\vdots$	$\vdots$		$\vdots$	$\vdots$
B_m	f_{m1}	$\dots$	f_{mn}	$ B_m $
Sum	$ A_1 $	$\dots$	$ A_n $	$n(U)$

ただし、

$$\sum_{i=1}^n \mu_{A_i}(x_k) = 1, \quad \sum_{i=1}^m \mu_{B_i}(x_k) = 1$$

であり、

$$f_{ij} = |A_j \cap B_i|$$
$$\mu_{A_j \cap B_i}(x) = \mu_{A_j}(x) \cdot \mu_{B_i}(x)$$

である。

この $m \times n$ ファジィ分割表を拡張し、type-2 ファジィ分割表を定義する。この type-2 ファジィ分割表を定義するときに type-2 ファジィ集合演算が必要になるため、ファジィ数の平均・ファジィ数の積・type-2 ファジィ集合の共通集合を次のように定める。

定義3. ファジィ数の平均

サポートが有界なファジィ数 $x_1^*, x_2^*, x_3^*, \dots, x_n^*$ に対し α -カット ($\alpha \in \mathbb{R}, 0 \leq \alpha \leq 1$) が次のように与えられているとする。

$$C_\alpha(x_i^*) = [a_{\alpha,i}, b_{\alpha,i}]$$

このとき、ファジィ数の平均 $\bar{x}^*$ は以下の式で与えられる。

$$\bar{x}^* = \bigcup_{\alpha \in (0,1]} \alpha C_\alpha(\bar{x}^*)$$

$$C_\alpha(\bar{x}^*) = \left[\frac{1}{n} \sum_{i=1}^n a_{\alpha,i}, \frac{1}{n} \sum_{i=1}^n b_{\alpha,i} \right]$$

定義4. ファジィ数の積

サポートが有界な任意のファジィ数 u_1^*, u_2^* に対し、ファジィ数の積 $u_1^* \cdot u_2^*$ は以下の式で与えられる。

$$u_1^* \cdot u_2^* = \bigcup_{\alpha \in (0,1]} \alpha C_\alpha(u_1^* \cdot u_2^*)$$

$$C_\alpha(u_1^* \cdot u_2^*) = \left[\min_{(x_1, x_2) \in C_\alpha(u_1^*) \times C_\alpha(u_2^*)} x_1 \cdot x_2, \max_{(x_1, x_2) \in C_\alpha(u_1^*) \times C_\alpha(u_2^*)} x_1 \cdot x_2 \right]$$

定義5. type-2 ファジィ集合の共通集合

全体集合 $U = \{x_i \mid i = 1, \dots, n\}$ とする。 U 上の type-2 ファジィ集合 $\tilde{A}, \tilde{B}$ が次のように与えられているとする。

$$\tilde{A} = \{(x_i, u_i^*) \mid i = 1, \dots, n\}, \tilde{B} = \{(x_i, v_i^*) \mid i = 1, \dots, n\}$$

ただし、グレード u_i^*, v_i^* はサポートが有界なファジィ数とする。このとき、2つの type-2 ファジィ集合 $\tilde{A}, \tilde{B}$ の共通集合 $\tilde{A} \cap \tilde{B}$ は以下のように与えられる。

$$\tilde{A} \cap \tilde{B} = \{(x_i, u_i^* \cdot v_i^*) \mid i = 1, \dots, n\}$$

これらにより、次の type-2 ファジィ $m \times n$ 分割表を定義する。

定義6. type-2 ファジィ $m \times n$ 分割表

全体集合 $U = \{x_i \mid i = 1, \dots, n\}$ とする。 type-2 ファジィ集合 $\tilde{A}_1, \dots, \tilde{A}_n, \tilde{B}_1, \dots, \tilde{B}_m$ の type-2 ファジィ $m \times n$ 分割表を以下のように定義する。

	$\tilde{A}_1$	...	$\tilde{A}_n$
$\tilde{B}_1$	$\bar{f}_{11}$	...	$\bar{f}_{1n}$
$\vdots$	$\vdots$		$\vdots$
$\tilde{B}_m$	$\bar{f}_{m1}$	...	$\bar{f}_{mn}$

ただし、 $\bar{f}_{ij}$ は type-2 ファジィ集合 $\tilde{A}_i, \tilde{B}_j$ の共通集合 $\tilde{A}_i \cap \tilde{B}_j$ のグレードの平均 $u^* \cdot v^*$ とする。

3. 適用事例

2014 年度秋学期開講のオープン科目「数学基礎プラスシリーズ」を受講した学生 2,566 名に対し、図 3 のようなメディア授業に対する機能についてアンケートを行った（設問数は 12）。有効回答数は設問によってばらつきはあるが、820 前後である。

本論文では設問 1, 2 についての分析を紹介する。

(充足質問)

設問1

下記の質問に対し、あなたの感覚にもっとも近い箇所にチェックしてください。

各回必須課題について、期間内 再受験可能 だったらどう思いますか？

- ☐ とても嬉しい
- ☐ 当然だろう
- ☐ 特に何とも思わない
- ☐ 仕方がない
- ☐ それは困る

(不充足質問)

設問2

下記の質問に対し、あなたの感覚にもっとも近い箇所にチェックしてください。

各回必須課題について、期間内 再受験不可能 だったらどう思いますか？

- ☐ とても嬉しい
- ☐ 当然だろう
- ☐ 特に何とも思わない
- ☐ 仕方がない
- ☐ それは困る

図 3 アンケート例

このアンケートより、表 4 のようなデータが得られる（表は 1 人目から 27 人目のみ表示している）。

表1 アンケート結果

	設問1	設問2	設問3	設問4	設問5	設問6	設問7	設問8	設問9	設問10	設問11	設問12
1	2	3	1	4	4	1	0	3	2	2	0	2
2	0	3	0	3	2	2	0	3	0.33	2	0	3.33
3	0	3	1.67	1	3	1	0	3.67	2	2	0.67	3.67
4	0	2	2	2	3.33	1.33	0	2	0.67	2	0	2
5	0.33	2	0.33	2	1	2	0.67	2	0	4	2	2
6	0	3	1	3	3	2	2	2	1	3	0	3
7	0	3	0.67	2	3.67	0.33	0.33	2.67	2	2	0.67	2
8	0	3	1	3	4	1	0	3	2	3	2	2
9	0	1	0.33	2	3.33	0.33	0	3	2	2	0.33	2
10	0	3.67	0	3.67	3.33	2	2	3	2.33	3.67	0	2
11	0	3	0.67	3.67	3.33	2	0.33	2	1.67	2.67	1.67	2
12	0	4	0	4	4	0	0	2	2	2	0	4
13	1.33	3	0	4	2	1.33	1	3.33	2	2	0.67	2
14	0.33	3	2	3	3	2	0	4	0	4	2	2
15	0	2	0	2	3	0	0	3	0	2	0	3
16	0	3	0.33	#N/A	4	0	0	2	2	2	0.33	2
17	0	4	0	4	4	0	0	4	2	2	2	2
18	0	4	0	4	2	2	0	2	0	2	0	2
19	0	3	0	3	4	0	1	2	1	2	0	2
20	0	4	0	4	3.33	0	0	3.33	1	3	0	4
21	1.67	2.67	1.67	2.67	3.33	1.33	2	2	0.67	3.67	2	2
22	0	3	0	4	1.33	2	1.33	3	0.67	2	1.33	2
23	0	3	0	3	2	2	2	2	2	2	2	2
24	0.67	3	0	3	3	0.67	2	2	0.33	2.33	0.33	3.33
25	0	3	0	3.33	3.67	0	1.33	2	1.33	2.33	0.67	2.67
26	0.67	2.67	0.67	3	3.33	0.67	2	2	1.67	2	0.67	2
27	0.33	1	2.67	1	4	0	3	1	0	4	1	3

なお，表中の数字は

- 0 . . . とても嬉しい
 1 . . . 当然だろう
 2 . . . 特に何とも思わない
 3 . . . 仕方がない
 4 . . . それは困る

のように対応する．

このアンケート結果から，設問 j に対し，ファジィ集合「とても嬉しい： $A_0^{(j)}$ 」「当然だろう： $A_1^{(j)}$ 」「特に何とも思わない： $A_2^{(j)}$ 」「仕方がない： $A_3^{(j)}$ 」「それは困る： $A_4^{(j)}$ 」に対するグレードを算出する．ファジィ集合 $A_k^{(j)}$ ($k = 0, \dots, 4$) に対するグレードは以下の式で算出される．

$$\mu_{A_k^{(j)}}(x_{ij}) = \begin{cases} 1 - x_{ij} + [x_{ij}] & , k - 1 < x_{ij} < k + 1 \\ 0 & , otherwise \end{cases}$$

このアンケートはインターネットを利用して行うため，中間値を取り扱うことができない．そこで，グレードにもさらに曖昧性を持たせ type-2 ファジィ集合として扱う（つまりグレードをファジィ数とする）．グレード s をファジィ数にしたものを s^* とする．このとき，ファジィ数 s^* のメンバーシップ関数を以下のように定める．

$$\mu_{s^*}(u) = \max(0, 1 - 3|u - s|)$$

$$= \begin{cases} 3u - 3s + 1 & , s - \frac{1}{3} < u < s \\ -3u + 3s + 1 & , s < u < s + \frac{1}{3} \\ 0 & , otherwise \end{cases}$$

この演算により type-2 ファジィ集合 $\widetilde{A_l^{(2k-1)}}, \widetilde{A_j^{(2k)}} (i, j = 0, \dots, 4, k = 1, \dots, 6)$ の共通集合 $\widetilde{A_l^{(2k-1)}} \cap \widetilde{A_j^{(2k)}}$ のグレードが得られて，グレードの平均により type-2 ファジィ 5×5 分割表が求められる．

表2 type-2 ファジィ 5×5 分割表

	$\widetilde{A_0^{(2k)}}$	$\widetilde{A_1^{(2k)}}$	$\widetilde{A_2^{(2k)}}$	$\widetilde{A_3^{(2k)}}$	$\widetilde{A_4^{(2k)}}$
$\widetilde{A_0^{(2k-1)}}$	$\widetilde{f_{00}}$	$\widetilde{f_{01}}$	$\widetilde{f_{02}}$	$\widetilde{f_{03}}$	$\widetilde{f_{04}}$
$\widetilde{A_1^{(2k-1)}}$	$\widetilde{f_{10}}$	$\widetilde{f_{11}}$	$\widetilde{f_{12}}$	$\widetilde{f_{13}}$	$\widetilde{f_{14}}$
$\widetilde{A_2^{(2k-1)}}$	$\widetilde{f_{20}}$	$\widetilde{f_{21}}$	$\widetilde{f_{22}}$	$\widetilde{f_{23}}$	$\widetilde{f_{24}}$
$\widetilde{A_3^{(2k-1)}}$	$\widetilde{f_{30}}$	$\widetilde{f_{31}}$	$\widetilde{f_{32}}$	$\widetilde{f_{33}}$	$\widetilde{f_{34}}$
$\widetilde{A_4^{(2k-1)}}$	$\widetilde{f_{40}}$	$\widetilde{f_{41}}$	$\widetilde{f_{42}}$	$\widetilde{f_{43}}$	$\widetilde{f_{44}}$

ファジィ分割表から各カテゴリーに対する帰属度表を作成する（表 3）．

表3 各カテゴリーに対する帰属度表

一(無)	当(無)	魅(無)	無	魅(有)	当(有)	一(有)
c	$\bar{b}^*$	$\bar{c}^*$	$\bar{d}^*$	$\bar{e}^*$	$\bar{f}^*$	$\bar{g}^*$

ただし， $\bar{f}_{ij}^*$ の α -カットを

$$C_\alpha(\bar{f}_{ij}^*) = [a_{\alpha,ij}, b_{\alpha,ij}]$$

とし，

$$\bar{a}^* = \bar{f}_{51}^*$$

$$\bar{b}^* = \bigcup_{\alpha \in (0,1]} \alpha C_\alpha(\bar{x}_b^*), C_\alpha(\bar{x}_b^*)$$

$$= \left[\sum_{j=2}^4 a_{\alpha,5j}, \sum_{j=2}^4 b_{\alpha,5j} \right]$$

$$\bar{c}^* = \bigcup_{\alpha \in (0,1]} \alpha C_\alpha(\bar{x}_c^*), C_\alpha(\bar{x}_c^*)$$

$$= \left[\sum_{i=2}^4 a_{\alpha,i1}, \sum_{i=2}^4 b_{\alpha,i1} \right]$$

$$\bar{d}^* = \bigcup_{\alpha \in (0,1]} \alpha C_\alpha(\bar{x}_d^*), C_\alpha(\bar{x}_d^*)$$

$$= \left[\sum_{i=2}^4 \sum_{j=2}^4 a_{\alpha,ij}, \sum_{i=2}^4 \sum_{j=2}^4 b_{\alpha,ij} \right]$$

$$\bar{e}^* = \bigcup_{\alpha \in (0,1]} \alpha C_\alpha(\bar{x}_e^*), C_\alpha(\bar{x}_e^*)$$

$$= \left[\sum_{j=2}^4 a_{\alpha,1j}, \sum_{j=2}^4 b_{\alpha,1j} \right]$$

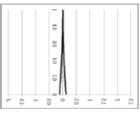
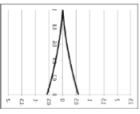
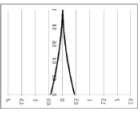
$$\bar{f}^* = \bigcup_{\alpha \in (0,1]} \alpha C_\alpha(\bar{x}_f^*), C_\alpha(\bar{x}_f^*)$$

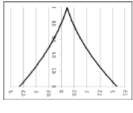
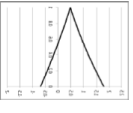
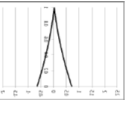
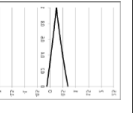
$$= \left[\sum_{i=2}^4 a_{\alpha,i5}, \sum_{i=2}^4 b_{\alpha,i5} \right]$$

$$\bar{g}^* = \bar{f}_{15}^*$$

これにより，設問 1,2 の各カテゴリーに対する帰属度表は表 4 のようになる．

表4 各カテゴリに対する帰属度表(設問 1,2)

カテゴリ	一(無)	当(無)	魅(無)
ファジィ数の グラフ			
値	0.0012092	0.0036276	0.0045681

無	魅(有)	当(有)	一(有)
			
0.2196695	0.4832729	0.0264678	0.2464060

この各カテゴリの帰属度から、ファジィ加重平均を求めこれを総合評価とする。各ウェイトはファジィ数とし、各メンバーシップ関数は図4のように定める。

表5 ウェイト

ウェイト	入力	結論
-1	一(無)	この機能は積極的に実装しない(G)
-2/3	当(無)	この機能はできるだけ実装しない(F)
-1/3	魅(無)	この機能は余裕があれば実装しない(E)
0	無	どちらでもよい(D)
1/3	魅(有)	この機能は余裕があれば実装する(C)
2/3	当(有)	この機能はできるだけ実装する(B)
1	一(有)	この機能は積極的に実装する(A)

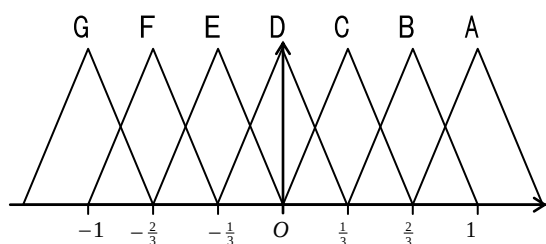


図4 ウェイトとなるファジィ数のメンバーシップ関数

これにより、推論結果 y は以下の式で求められる。

$$y = \bar{a}^* \cdot A + \bar{b}^* \cdot B + \bar{c}^* \cdot C + \bar{d}^* \cdot D + \bar{e}^* \cdot E + \bar{f}^* \cdot F + \bar{g}^* \cdot G$$

よって、設問 1,2 に対する推論結果は図5のようなファジィ数となる。

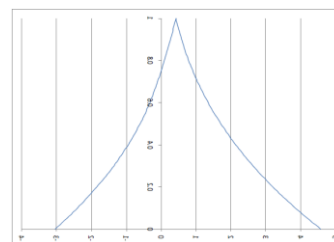


図5 推論結果

このファジィ数に対し重心法で非ファジィ化を行うため、重心を計算すると

$$z^* = 0.649669559$$

となる。

3. おわりに

今回、type-2 ファジィ 5×5 分割表を応用し、学生のニーズ分析を実行してみた。この結果、手法としての有効性のある程度確認できたので、今後はさらにアンケートの取り方や分析方法の改善などを改良していきたい。

参考文献

- [1] 狩野，瀬楽，高橋，辻：魅力的品質と当たり前品質，日本品質管理学会会報『品質』14 (2), pp. 39-48, 1984.
- [2] M.Rashid, J.Tamaki, A.M.M. S.Ullah, and A.Kubo: A Kano Model Based Linguistic Application for Customer Needs Analysis, International Journal of Engineering Business Management, Vol.3, No.2, pp. 30-36, 2011.
- [3] 狩野モデルを応用したメディア授業における学生のニーズ分析，日本知能情報ファジィ学会ソフトサイエンスワークショップ，pp.89-90, 2014.
- [4] 狩野モデルを応用したメディア授業におけるニーズ分析，日本知能情報ファジィ学会ファジィシステムシンポジウム，ME1-3, 2014.

連絡先

早稲田大学

上江洲 弘明

(Tel.:03-5286-2966, E-mail:uesu@suou.waseda.jp)

狩野モデルによる大学数学講義の学生ニーズ分析 (2)

An Analysis of Students' Needs for Undergraduate Mathematics Lectures through the Kano Model, II

○高木 悟¹, 上江洲 弘明²
○Satoru Takagi¹, Hiroaki Uesu²

¹ 工学院大学

¹Kogakuin University

² 早稲田大学

²Waseda University

Abstract: At universities, mathematics lecturers have to change their teaching materials and methods according to students' mathematical skills, and also have to make them understand mathematics. In this paper, we modify the previous questionnaire analysis and apply it in some undergraduate mathematics lectures for the method of students' needs analysis using the Kano model, and report the results.

1. はじめに

大学での数学講義における受講生のニーズは、ただ単に教材や教授法が良ければ満足し、悪ければ不満を持つといった「一元的な要素」だけではなく、例えばあったら満足するが無くても仕方ないと思う「魅力的な要素」や、無ければ不満を持つがあっても当然と思う「当たり前要素」もある。このようなニーズを分析するとき、1つの要素に対して「充足質問」と「不充足質問」を問うことにより、ニーズ分類の手法を提案した狩野モデル[1]が有用である。第二著者は、フルオンデマンド授業における受講生ニーズ分析にこの狩野モデルを利用し、さらにファジィ推論によって分析した[2]。また、第二著者のこの手法を用いて、筆頭著者が担当する大学での数学講義（フルオンデマンドではなく対面形式の授業）において、講義のさまざまな要素を狩野モデルで分類し、ファジィ推論で分析した[3]。本稿では、前回実施の反省点を踏まえ、より詳細な分析なできるよう改良したアンケートで実施した結果を報告する。また、今後の改善案についても述べる。

2. 分析方法

狩野モデルでは、顧客に対し「充足質問（「この要素があるとどう思いますか？」という質問）」と「不充足質問（「この要素が無いとどう思いますか？」という質問）」を行ない、その回答によって分類する。質問の回答は「とてもうれしい」「当然だろう」「特に何とも思わない」「別にそれでも構わない」「それは困る」の5つの中から1つを選択するものである。前回のアンケートでは、この狩野モデルにおける五者択一の回答方法を採用したが（図1）、「別にそれでも構わない」の解釈が難しく、回答者が混乱する様子が前回のアンケート結果で見られたため、この文言は「仕方ない」に変更した。また、この五者択一ではあらかじめこちらで用意した5つの中から1つを選ばないといけないが、例えば「それは困るけど、仕方ないところもある」といった具合に、五者択一では回答できない気持ちも回答に反映できるよう、図2のように「とてもうれしい」から「それは困る」までを一直線で表して、その等分間隔に「当然だろう」「特に何とも思わない」「仕方ない」を配置し、回答者の思う気持ちのところに縦線を引いてもらうような回答方法とした。このように、回答者の気持ちを一直線上で表せるのかという問題もあるが、今回はその点は考慮せず実施した。今後への課題としたい。

「それは困る」の5つの中から1つを選択するものである。前回のアンケートでは、この狩野モデルにおける五者択一の回答方法を採用したが（図1）、「別にそれでも構わない」の解釈が難しく、回答者が混乱する様子が前回のアンケート結果で見られたため、この文言は「仕方ない」に変更した。また、この五者択一ではあらかじめこちらで用意した5つの中から1つを選ばないといけないが、例えば「それは困るけど、仕方ないところもある」といった具合に、五者択一では回答できない気持ちも回答に反映できるよう、図2のように「とてもうれしい」から「それは困る」までを一直線で表して、その等分間隔に「当然だろう」「特に何とも思わない」「仕方ない」を配置し、回答者の思う気持ちのところに縦線を引いてもらうような回答方法とした。このように、回答者の気持ちを一直線上で表せるのかという問題もあるが、今回はその点は考慮せず実施した。今後への課題としたい。

授業改善アンケート

今後の授業改善のために、ご協力ください。もちろん、成績とは無関係です。
また、結果については教育に関する研究にも使わせていただきます。

★工学院大学へは、一般入試(S/A/B/M/センター試験利用)、あるいは推薦入試(AO/指定校/付属/公募/海外など)のどちらで合格し、入学しましたか？ 当てはまるもの1つに○を付けてください。

	一般入試	推薦入試
設問1	毎回、提出必須の宿題があったら、どう思いますか？	とてもうれしい 当然だろう 特に何とも思わない 別にそれでも構わない それは困る
設問2	毎回、提出必須の宿題がなかったら、どう思いますか？	とてもうれしい 当然だろう 特に何とも思わない 別にそれでも構わない それは困る
設問3	あなたからのコメントに対して、担当教員からのフィードバックがあったら、どう思いますか？	とてもうれしい 当然だろう 特に何とも思わない 別にそれでも構わない それは困る
設問4	あなたからのコメントに対して、担当教員からのフィードバックがなかったら、どう思いますか？	とてもうれしい 当然だろう 特に何とも思わない 別にそれでも構わない それは困る

図1 改良前のアンケート用紙（一部）

授業改善アンケート【無記名式】2014年度 前期

今後の授業改善のため、ご協力よろしくお願い申し上げます。
結果については、教育改善のほか、関連する研究にも使わせていただきます。

★工学院大学へは、一般入試(S/A/B/M/センター試験利用)、あるいは推薦入試(AO/指定校/附属/公募/海外など)のどちらかで合格して入学しましたか？ 当てはまるもの1つに○を付けてください。

一般入試	推薦入試
------	------

★以下の設問について、あなたの感覚にもっと近い箇所「縦線」を入れてください。

【回答例】	とても 少ない	普通 だろう	特に 多い と思う	仕方 ない	それは 困る
設問1 毎日、提出必須の宿題があったら どう思いますか？					
設問2 毎日、提出必須の宿題がなかったら どう思いますか？					
設問3 あなたからのコメントに対して、 担当教員からのフィードバックがあったら どう思いますか？					
設問4 あなたからのコメントに対して、 担当教員からのフィードバックがなかったら どう思いますか？					

図 2 改良後のアンケート用紙（一部）

このアンケート回答結果を、まずは数値化する。具体的には、「とてもうれしい」を1,「当然だろう」を2,「特に何とも思わない」を3,「仕方ない」を4,「それは困る」を5としたときに、回答された縦線の位置を小数第1位あるいは小数第2位までの数値で表すこととする。

続いて、数値化した結果に対して、整数ではない数値については、次のような評価を行なう。例えば、その回答における読取数値が 4.7 であったとき、4は「仕方ない」、5は「それは困る」であるから、回答者は「それは困る」が70%,「仕方ない」が30%の気持ちであると評価する。

この評価値をもとに、表1のような分類マトリックスで集計を行なう。なお、表1における記号の意味は、以下のとおりである。

魅(有)：あると魅力的な要素

当(有)：あって当然の要素

一(有)：あればあるほど満足度があがるような要素

無：無関心で、あってもなくても気にならない程度の要素

魅(無)：無いと魅力的な要素

当(無)：無いことが当然の要素

一(無)：無ければ無いほど満足度があがるような要素

懐：懐疑的回答、回答が不整合

表 1 狩野モデルの分類マトリックス

充足\不充足	とてもうれしい	当然だろう	特に何とも思わない	仕方ない	それは困る
とてもうれしい	懐	魅(有)	魅(有)	魅(有)	一(有)
当然だろう	魅(無)	無	無	無	当(有)
特に何とも思わない	魅(無)	無	無	無	当(有)
仕方ない	魅(無)	無	無	無	当(有)
それは困る	一(無)	当(無)	当(無)	当(無)	懐

このように著者らは、狩野モデルを大学の数学講義でのアンケート分析に応用し、ファジィ推論で分析する手法の改良版を提案する。

- (1) 受講生に対し「充足質問」と「不充足質問」で構成されたアンケートを行なう。
- (2) 表1のクロス集計表を作成する。
- (3) 表1における「一(有)」「当(有)」「魅(有)」「無」「魅(無)」「当(無)」「一(有)」の箇所に関して度数表を作成する(表2)。

表 2 度数表

一(無)	当(無)	魅(無)	無	魅(有)	当(有)	一(有)
<i>a</i>	<i>b</i>	<i>c</i>	<i>d</i>	<i>e</i>	<i>f</i>	<i>g</i>

- (4) 表2の度数表から各カテゴリに対する帰属度表を作成する(表3)。

表 3 各カテゴリに対する帰属度表

一(無)	当(無)	魅(無)	無	魅(有)	当(有)	一(有)
<i>a/m</i>	<i>b/m</i>	<i>c/m</i>	<i>d/m</i>	<i>e/m</i>	<i>f/m</i>	<i>g/m</i>

(ただし、 $m = \max\{a, b, c, d, e, f, g\}$)

(5) 帰属度から、ファジィ推論を実行する。

ファジィルールを表 4 のように定め、そのメンバーシップ関数は図 3 のように定める。

表 4 ファジィルール

入力	結論
一(無)	この要素は積極的に取り入れない (G)
当(無)	この要素はできるだけ取り入れない (F)
魅(無)	この要素は余裕があれば取り入れない (E)
無	どちらでもよい (D)
魅(有)	この要素は余裕があれば取り入れる (C)
当(有)	この要素はできるだけ取り入れる (B)
一(有)	この要素は積極的に取り入れる (A)

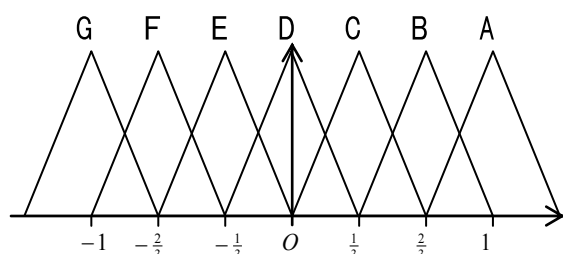


図 3 ファジィ推論における結論のメンバーシップ関数

(6) 加算法により結論の合成を求め、重心法により非ファジィ化を行ない、これを結論とする。

3. 適用事例

2014 年前期に、筆頭著者が担当した工学院大学の「情報数学 I」「代数学」「数学 I」「微分積分 I」「線形代数学 I」「基礎数学」を受講した学生 235 名に対し、以下の 6 要素においてそれぞれ充足質問と不充足質問をした。

- (1) 提出必須の宿題
- (2) コメントへのフィードバック
- (3) 補足プリント
- (4) 雑談
- (5) 質問タイム
- (6) 学習支援センター

このアンケート結果に対し、各設問に対する帰属度表を作成し、それにファジィ推論を適用して分析した結果と、それによる授業改善優先順位を以下に示す(表 5)。また、参考までに前回調査(2013 年度後期)の結果を表 6 に掲載する。結論の数値は、1 に近いほどあった方が良く、-1 に近いほど無い方が良いことを表す。

表 5 分析結果(今回: 2014 年度前期)

	今回(2014 前)	優先順位
(1) 提出必須の宿題	0.011369	6
(2) コメントへのフィードバック	0.267740	4
(3) 補足プリント	0.455698	2
(4) 雑談	0.206695	5
(5) 質問タイム	0.516409	1
(6) 学習支援センター	0.413105	3

表 6 分析結果(前回: 2013 年度後期)

	前回(2013 後)	優先順位
(1) 提出必須の宿題	-0.015152	6
(2) コメントへのフィードバック	0.377919	4
(3) 補足プリント	0.628450	2
(4) 雑談	0.292994	5
(5) 質問タイム	0.670940	1
(6) 学習支援センター	0.527778	3

この結果から、学生の授業に対するニーズの中で「質問タイム(担当教員が授業中に質問に答える時間・環境を作る)」が一番高く、次に「補足プリント(授業内容を補足するプリントを担当教員が作成し、配布する)」、「学習支援センター(主に基礎学力が不足した学生をフォローする専属の部署である学習支援センターで質問相談できる)」であることが分かった。このことから、学生が大学数学講義に対し求める付加的なものは、授業時間内外を問わず、担当教員や学習支援スタッフによる質問対応やフォローが一番であることが分かった。

この結論は, 前回のアンケートと同じであったが, 前回に比べて数値が低い, つまり要求度合いが前回よりも低いことがわかる. これは, 前回のアンケートでは回答が五者択一だったことで, どちらかの気持ちを回答する必要があったが, 今回この回答方法を改良したことにより, 詳細な気持ちを反映できるようにになった影響が大きいと考えられる. 前回と今回とで受講生も異なり, クラスの雰囲気も異なるので一概には言えないが, より学生のニーズを表した結果になったと判断でき, 少なくとも前回の回答方法よりは「改良」されていると考える.

4. おわりに

前回, そして今回と, 狩野モデルを学生のニーズ分析に応用し, 手法としての有効性が十分に確認できた. 今後の検討事項としては, アンケート回答結果の読取について, スキャナ等を利用するなどして時間短縮を図りたいこと, またクラスごとの集計も必要であることが挙げられる. 前者については, スキャンする際にどの用紙も必ず同じ位置に来るように読取ができるのかどうか, などの懸念事項があり, 今後も並行して検討はしていくものの実現化は難しいと考えている. ただ, どうかして読取作業を早める工夫は考えたい. 後者については, 学部や学年, 科目などクラスによって授業中の雰囲気がかなり異なることから, クラスごとの分析をすることで, そのクラスに合ったよりきめ細やかな対応が出来るのでは, という期待がある. そのためには, 早い段階でアンケートを実施し, 分析した結果をフィードバックしないといけなく, 読取作業の高速化が一番の問題であると考えている.

参考文献

- [1] 狩野紀昭, 瀬楽信彦, 高橋文夫, 辻新一: 魅力的品質と当たり前品質, 日本品質管理学会会報『品質』14 (2), pp. 39-48, 1984.
- [2] 上江洲弘明: 狩野モデルを応用したメディア授業における学生のニーズ分析, 第24回ソフトサイエンスワークショップ講演論文集, pp.89-90, 2014.

[3] 上江洲弘明, 高木悟: 狩野モデルによる大学数学講義の学生ニーズ分析, 第24回ソフトサイエンスワークショップ講演論文集, pp.87-88, 2014.

[4] S. Takagi, H. Uesu: A questionnaire analysis by fuzzy reasoning for undergraduate students in mathematics lectures applying the Kano model, Proceedings of the Joint 7th International Conference on Soft Computing and Intelligent Systems and 15th International Symposium on Advanced Intelligent Systems (SCIS-ISIS 2014), pp.851-854, 2014.

連絡先

工学院大学

高木 悟

E-mail: satoru@cc.kogakuin.ac.jp

早稲田大学

上江洲 弘明

E-mail: uesu@suou.waseda.jp

区間 UTA 法によるセミロバスト多基準意思決定支援

Semi-robust multiple criteria decision aiding by interval UTA method

○ 乾口 雅弘, 奥村 朗
○ Masahiro Inuiguchi, Akira Okumura
大阪大学 大学院基礎工学研究科

Graduate School of Engineering Science, Osaka University

Abstract: In this paper, we propose an interval UTA method as an intermediate approach between the conventional and robust ordinal regression methods. The proposed method improves the dependability of the arbitrary estimation by the conventional ordinal regression and the tractability and estimation ability of the robust ordinal regression. In the proposed approach, an interval utility function compatible with given preference information is obtained by solving a linear programming problem. The estimation of the overall preference relation is done by simple calculations except special cases. The application to a numerical example demonstrates that the preference relation estimated by the proposed approach has an intermediate property between those by the conventional and robust ordinal regressions.

1 はじめに

本研究では、多基準決定問題における一意思決定支援手法を提案する。従来法として、UTA 法のような順序回帰法 [1] や、 UTA^{GMS} 法のようなロバスト順序回帰法 [2] などが挙げられるが、これらでは意思決定者の選好は加法的効用関数で表されると仮定され、参照代替案の対比較情報に基づき加法的効用関数が推定される [3]。前者では加法的効用関数が 1 つ同定され、それを用いて評価される。効用値を計算する簡潔な評価となる反面、その効用関数に依存した評価になるという欠点がある。一方、後者では意思決定者の選好情報に整合するすべての加法的効用関数の集合により、代替案の評価を行うため、安全な評価を得られる反面、比較不能となりやすく、一対比較の度に線形計画問題を解く必要がある。

本研究では、これら 2 手法の中間的な性質を持つ区間順序回帰を提案する。前報 [4] では、各基準の改善分に対する効用の増分を区間として表現するモデルを提案している。このモデルでは、効用関数の正規化に対処するため、幅を小さくする区間演算を導入していた。このため、増分効用の区間の扱いが一樣ではなかった。また、選好関係は強い区間の不等号関係で表現した反面、無差別関係は緩い区間の類似関係として表現していたというアンバランスがあった。

本研究では、正規性や無差別関係を増分効用間の相互関係としてとらえ、相互関係のある区間モデルとして、区間効用関数を求める方法を提案する。一つの 0-1 混合線形計画問題を解く必要があった従来モデルに対し、本提案モデルでは、線形計画問題を小数回解くだけで代替案間の選好関係が求められるという利点をもつ。

2 提案モデル

本研究で用いる区間効用関数について述べる。提案モデルでは、値が大きいほど望ましい基準 g_i の取る範囲 $[c_i^0, c_i^{m_i}]$ に区分点 $c_i^1, c_i^2, \dots, c_i^{m_i-1}$ を考え、基準 g_i が c_i^{k-1} から c_i^k へ改善されたことに対する効用の増加量 $\delta_{ik} \geq 0$ を考える。 δ_{ik} は、与えられた選好情報を満足し、かつその変動範囲が区間 $\Delta_{ik} = [d_{ik} - w_{ik}, d_{ik} + w_{ik}]$ で制限される。このモデルを区間増分効用モデルと呼ぶことにする。ここで、区間 Δ_{ik} で各 δ_{ik} を制限しなければ、 δ_{ik} の集合は仮定するロバスト順序回帰 [2] で考えられてきたモデルと等価になることを補足する。

与えられた選好情報に整合する δ_{ik} の集合を Con と表し、区間増分効用モデルによる代替案 a_1 と a_2 との比較を考える。 $g_i(a_j)$ を代替案 a_j の基準 g_i の値とし、 k_j を $g_i(a_j) = \min\{k \mid g_i(a_j) \leq c_i^k\}$ となる k の値とする。また、 $V_1 = \{i \mid g_i(a_1) > g_i(a_2)\}$, $V_2 = \{i \mid g_i(a_2) > g_i(a_1)\}$ と定め、 $f(\delta_{ik} | a_1, a_2)$ を

$$\begin{aligned} f(\delta_{ik} | a_1, a_2) &= \sum_{i \in V_1} \sum_{k=k(a_2)}^{k(a_1)-1} \delta_{ik} + \sum_{i \in V_1} \frac{g_i(a_1) - c_i^{k(a_1)-1}}{c_i^{k(a_1)} - c_i^{k(a_1)-1}} \delta_{ik(a_1)} \\ &\quad + \sum_{i \in V_2} \frac{c_i^{k(a_2)} - g_i(a_2)}{c_i^{k(a_2)} - c_i^{k(a_2)-1}} \delta_{ik(a_2)-1} \end{aligned} \quad (1)$$

と定める。次の計画問題の最適値が非負となるとき、 a_1 を a_2 以上にきつと好む ($a_1 \succ^S a_2$) と推定できる。

$$\begin{aligned} \min & f(\delta_{ik} | a_1, a_2) - f(\delta_{ik} | a_2, a_1) \\ \text{s.t.} & \{\delta_{ik} \mid k = 1, \dots, m_i, i = 1, \dots, p\} \subseteq Con \\ & \delta_{ik} \in \Delta_{ik}, k = 1, \dots, m_i, i = 1, \dots, p \end{aligned} \quad (2)$$

なお、 Con は δ_{ik} に関して線形領域になるので、式 (2) は線形計画問題となる。

式 (2) の最初の制約条件を除くと、この問題の最適値は $f(d_{ik} - w_{ik}|a_1, a_2) - f(d_{ik} + w_{ik}|a_2, a_1)$ と容易に求められる。この値は式 (2) の最適値以下となるので、非負であれば $a_1 \succsim a_2$ と結論付けることができる。したがって、 $f(d_{ik} - w_{ik}|a_1, a_2) - f(d_{ik} + w_{ik}|a_2, a_1) < 0$ かつ $f(d_{ik} - w_{ik}|a_2, a_1) - f(d_{ik} + w_{ik}|a_1, a_2) < 0$ である場合を除き、式 (2) を解く必要はない。

$a_1 \succsim^S a_2$ かつ $a_2 \succsim^S a_1$ が成立する場合もあり、この場合は $a_1 \sim^S a_2$ (a_1 と a_2 はきつと無差別) と推定される。一方、 $a_1 \not\succsim^S a_2$ かつ $a_2 \not\succsim^S a_1$ となることも多く、 a_1 と a_2 とを確信して比較できないことを表す。

このように、区間増分効用モデルが与えられると、代替案間の選好関係が得られる。この関係は反射性、推移性を満たすので前順序となる。

3 選好情報からのモデル同定法

ロバスト順序回帰では、いくつかの代替案対、代替案ペア対、基準値対、あるいは基準値ペア対に関する選好情報と整合する加法的効用関数の集合を考える。本研究で提案するモデルでも同様にこれらの選好情報を取り扱い、それらに整合する加法的効用関数の集合を考えることができるが、本稿では基本的な代替案ペアの選好情報のみが与えられた場合を取り上げ、区間増分効用モデルの同定法について述べる。すなわち、 Δ_{ik} の中心値 d_{ik} 、幅の半分 w_{ik} と、 $\delta_{ik} \in \Delta_{ik}$ が満たすべき条件である Con を定める方法を議論する。

本研究では、表明される選好情報は、意思決定者がある程度確信して判断したものとしてとらえる。ある代替案ペア a_1, a_2 に関して、 $a_1 \succsim^S a_2$ という選好情報が与えられた場合、すべての選好情報から定められる d_{ik}, w_{ik} および Con に対して式 (2) の最適値が非負とならなければならない。しかし、モデル同定にこの条件を課せば、同定問題が制約条件に最適化問題を含む 2 レベル計画問題になり、容易解けるとは限らない。そこで、式 (2) の最適値以下の値となる $f(d_{ik} - w_{ik}|a_1, a_2) - f(d_{ik} + w_{ik}|a_2, a_1)$ を用いて、選好情報 $a_1 \succsim^S a_2$ を次式で扱うことにする。

$$a_1 \succsim^S a_2 \Leftrightarrow f(d_{ik} - w_{ik}|a_1, a_2) - f(d_{ik} + w_{ik}|a_2, a_1) \geq 0 \quad (3)$$

また、 $a_1 \succsim^S a_2$ より、 UTA^{GMS} 法と同様に、 Con の条件に $f(\delta_{ik}|a_1, a_2) - f(\delta_{ik}|a_2, a_1) \geq 0$ を加える必要があるが、式 (3) の条件は式 (2) の最適値の非負条件

より強く、 $f(\delta_{ik}|a_1, a_2) - f(\delta_{ik}|a_2, a_1) \geq 0$ を Con の条件に追加する必要がないことになる。

次に、 $a_1 \sim^S a_2$ であるという選好情報が与えられた場合、すべての選好情報から定められる d_{ik}, w_{ik} および Con に対して式 (2) の最適値、および a_1 と a_2 の順序を入れ替えた式 (2) の最適値がともに非負になるという条件を考えることになる。しかし、 $f(\delta_{ik}|a_1, a_2) - f(\delta_{ik}|a_2, a_1) \neq 0$ となる $\delta_{ik} \in \Delta_{ik}$ が存在する限りこの条件は成立しない。したがって、 $f(\delta_{ik}|a_1, a_2) - f(\delta_{ik}|a_2, a_1) = 0$ を Con の条件に含める必要がある。

一方、 $f(\delta_{ik}|a_1, a_2) - f(\delta_{ik}|a_2, a_1) = 0$ が Con の条件に含まれると、この条件に絡む δ_{ik} の範囲 $\Delta_{ik} = [d_{ik} - w_{ik}, d_{ik} + w_{ik}]$ の上下限値に次の制約が生じる。

$$\left. \begin{aligned} & f(d_{ik} - w_{ik}|a_1, a_2) - f(d_{ik} + w_{ik}|a_2, a_1) \\ & + 2w_{ik_1} \leq 0, \quad k_1 = k(a_2), \dots, k(a_1) - 1 \\ & f(d_{ik} - w_{ik}|a_1, a_2) - f(d_{ik} + w_{ik}|a_2, a_1) \\ & + \frac{2(g_i(a_1) - c_i^{k(a_1)-1})}{c_i^{k(a_1)} - c_i^{k(a_1)-1}} w_{ik(a_1)} \leq 0 \\ & f(d_{ik} - w_{ik}|a_1, a_2) - f(d_{ik} + w_{ik}|a_2, a_1) \\ & + \frac{2(c_i^{k(a_2)} - g_i(a_2))}{c_i^{k(a_2)} - c_i^{k(a_2)-1}} w_{i(k(a_2)-1)} \leq 0 \end{aligned} \right\} i \in V_1 \quad (4)$$

$$\left. \begin{aligned} & f(d_{ik} + w_{ik}|a_1, a_2) - f(d_{ik} - w_{ik}|a_2, a_1) \\ & - 2w_{ik_1} \geq 0, \quad k_1 = k(a_2), \dots, k(a_1) - 1 \\ & f(d_{ik} + w_{ik}|a_1, a_2) - f(d_{ik} - w_{ik}|a_2, a_1) \\ & - \frac{2(g_i(a_1) - c_i^{k(a_1)-1})}{c_i^{k(a_1)} - c_i^{k(a_1)-1}} w_{ik(a_1)} \geq 0 \\ & f(d_{ik} + w_{ik}|a_1, a_2) - f(d_{ik} - w_{ik}|a_2, a_1) \\ & - \frac{2(c_i^{k(a_2)} - g_i(a_2))}{c_i^{k(a_2)} - c_i^{k(a_2)-1}} w_{i(k(a_2)-1)} \geq 0 \end{aligned} \right\} i \in V_1 \quad (5)$$

$$\left. \begin{aligned} & f(d_{ik} - w_{ik}|a_1, a_2) - f(d_{ik} + w_{ik}|a_2, a_1) \\ & - 2w_{ik_1} \leq 0, \quad k_1 = k(a_1), \dots, k(a_2) - 1 \\ & f(d_{ik} - w_{ik}|a_1, a_2) - f(d_{ik} + w_{ik}|a_2, a_1) \\ & - \frac{2(g_i(a_2) - c_i^{k(a_2)-1})}{c_i^{k(a_2)} - c_i^{k(a_2)-1}} w_{ik(a_2)} \leq 0 \\ & f(d_{ik} - w_{ik}|a_1, a_2) - f(d_{ik} + w_{ik}|a_2, a_1) \\ & - \frac{2(c_i^{k(a_1)} - g_i(a_1))}{c_i^{k(a_1)} - c_i^{k(a_1)-1}} w_{i(k(a_1)-1)} \leq 0 \end{aligned} \right\} i \in V_2 \quad (6)$$

$$\left. \begin{aligned} & f(d_{ik} + w_{ik}|a_1, a_2) - f(d_{ik} - w_{ik}|a_2, a_1) \\ & + 2w_{ik_1} \geq 0, \quad k_1 = k(a_1), \dots, k(a_2) - 1 \\ & f(d_{ik} + w_{ik}|a_1, a_2) - f(d_{ik} - w_{ik}|a_2, a_1) \\ & + \frac{2(g_i(a_2) - c_i^{k(a_2)-1})}{c_i^{k(a_2)} - c_i^{k(a_2)-1}} w_{ik(a_2)} \geq 0 \\ & f(d_{ik} + w_{ik}|a_1, a_2) - f(d_{ik} - w_{ik}|a_2, a_1) \\ & + \frac{2(c_i^{k(a_1)} - g_i(a_1))}{c_i^{k(a_1)} - c_i^{k(a_1)-1}} w_{i(k(a_1)-1)} \geq 0 \end{aligned} \right\} i \in V_2 \quad (7)$$

例えば、式 (4) の最初の不等号が $k' \in [k(a_2), k(a_1) - 1]$, $i' \in V_1$ について成立しなければ、 $\delta_{i'k'} = d_{i'k'} + w_{i'k'} \in \Delta_{i'k'}$ に対して、 $f(\delta_{ik}|a_1, a_2) - f(\delta_{ik}|a_2, a_1) = 0$ となる $\delta_{ik} \in \Delta_{ik}$ は存在しない。すなわち、 $f(\delta_{ik}|a_1, a_2) - f(\delta_{ik}|a_2, a_1) = 0$ を Con の条件に含む限り、 $\delta_{i'k'} = d_{i'k'} + w_{i'k'}$ となることはなく、 $\Delta_{i'k'}$ に取りえない値が含まれることになる。式 (4)~(7) は $f(\delta_{ik}|a_1, a_2) - f(\delta_{ik}|a_2, a_1) = 0$ という条件下で、この条件に絡む Δ_{ik} に無駄な領域を含まないことを意味する。したがって、選好情報 $a_1 \sim^S a_2$ が与えられれば、 Con の条件に $f(\delta_{ik}|a_1, a_2) - f(\delta_{ik}|a_2, a_1) = 0$ を含めるとともに、式 (4)~(7) の制約条件を d_{ik}, w_{ik} に課す必要がある。

UTA 法や UTA^{GMS} 法では、最大の効用値が 1 となるように正規性条件を課しているが、正規化は本質的ではない。必要があれば、 $\delta_{ik} \in \Delta_{ik}$ の実現値が与えられれば総和を 1 に正規化すればよいので、本研究では、正規性条件を Con の条件に含めないものとする。

目的関数としては、できるだけ多くのクリスプな効用関数が含まれるように、各 Δ_{ik} の広がり w_{ik} を均等に最大化する。最適解が複数存在する場合には、それらの中で Δ_{ik} の広がり w_{ik} の総和を最大化するものを選ぶ。しかし、 Con の条件に正規性条件を含めていないので、広がり w_{ik} は無限に大きくなる。そこで、 Δ_{ik} の中心 d_{ik} の総和に対するこれらの広がりを最大化する。すなわち、 q を十分小さい正数として、次の目的関数を用いる。

$$\max \frac{q \min_{i=1, \dots, p} \min_{k=1, \dots, m_i} w_{ik} + (1-q) \sum_{i=1}^p \sum_{k=1}^{m_i} w_{ik}}{\sum_{i=1}^p \sum_{k=1}^{m_i} d_{ik}} \quad (8)$$

以上より、増分効用の非負性、 $d_{ik} - w_{ik} \geq 0$ 、与えられたすべての選好情報に対する式 (3), (4)~(7) の制約条件下で、式 (8) を目的関数とする線形分数計画問題を解くことにより、 d_{ik}, w_{ik} すなわち、 Δ_{ik} が定められる。この問題は線形計画問題に帰着される。

表 1: 人材評価データ

番号	名前	g_1	g_2	g_3	選好情報
1	Alexievich	4	16	63	$9 \succ 10$
2	Bassama	28	18	28	$6 \succ 10$
3	Calvet	26	40	44	$10 \succ 8$
4	Dubois	2	2	68	$8 \succ 7$
5	El Marbat	18	17	14	$10 \sim 15$
6	Feeret	35	62	25	
7	Fleishman	7	55	12	
8	Fourny	25	30	12	
9	Frechet	9	62	88	
10	Martin	0	24	73	
11	Petron	6	15	100	$10 \sim 15$
12	Psorgos	16	9	0	
13	Smith	26	17	17	
14	Varlot	62	43	0	
15	Yu	1	32	64	

4 数値例

参考文献 [3] で用いられている表 1 の人材評価データを用いて、従来法と提案法とを比較する。このデータは、15 人の人材について生産管理能力 (g_1)、国際経験 (g_2)、人柄 (g_3) の得点を示したもので、表 1 の最右列に示すように、人材間の五つの選好情報が与えられている。各基準の最小値と最大値により範囲を定め、これを 10 等分して 10 区分からなる区分的線形な周辺効用関数を考える。

まず、UTA 法を適用する。整合する多くの効用関数から任意の一つの効用関数が選ばれるため、選好情報に関係しない 10 等分された各基準の最後の区分で効用値が大幅に変化する周辺効用関数が得られた。各基準の重みに相当する最良値での効用値は、生産管理能力 (g_1) が 0.333333、国際経験 (g_2) が 0.001634、人柄 (g_3) が 0.665033 となった。国際経験 (g_2) の影響が弱い結果が得られた。実際、国際経験 (g_2) の周辺効用関数は、第 5 区間で僅かに増大するのみで、この区間を除けば概ねフラットになり、求められた加法的効用関数では、国際経験 (g_2) は評価に大きく影響しない。得られた加法的効用関数から求められる人材間の選好関係のハッセ図を図 1(a) に示す。図中の細線は与えられた選好情報および属性値から得られる選好とそれらから推移律を用いて得られる選好である。太線は加法的効用関数により得られた選好である。弱順序関係となるが、任意に選んだ効用関数による評価であるため、太線の選好の信憑性は疑わしい。

次に、UTA^{GMS} 法を適用する。整合するすべての効

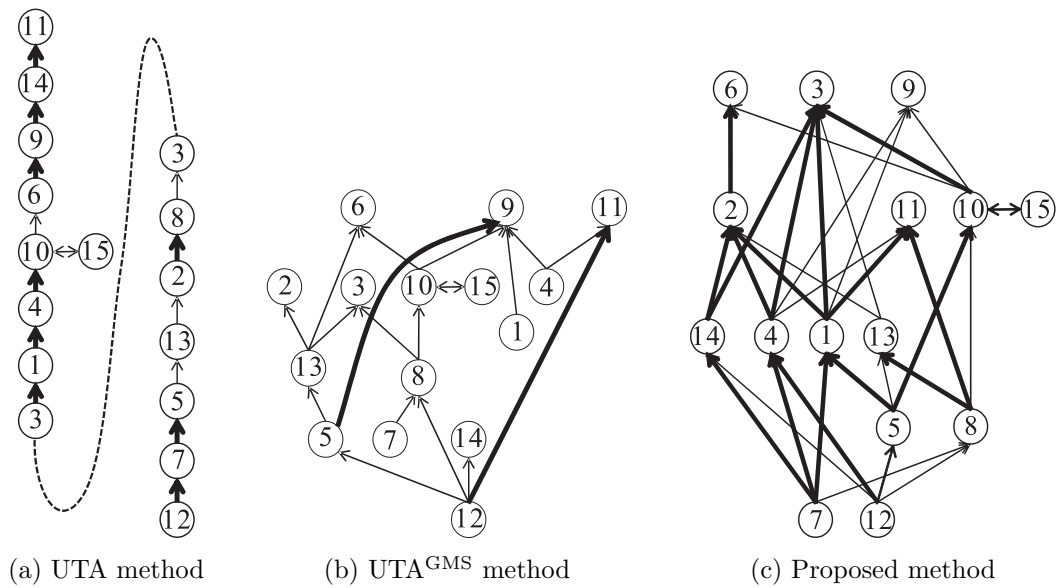


図 1: 各方法より得られる選好関係のハッセ図

用関数に関して選好が成立する場合しか受け入れないので、得られる選好は少なく、選好関係は図 1(b) に示すハッセ図のようになった。細線は選好情報やデータから求められる選好であり、太線は UTA^{GMS} 法により推定されるものである。この場合、 $9 \succ 5$ と $11 \succ 12$ のみが求められた。この選好は、UTA 法による結果および後に述べる提案法によるものと一致する。 UTA^{GMS} 法は非常に安全で控えめな評価を与えるため、得られる選好は少なく代替案を十分に絞り込めないことが多い。

最後に、 $q = 0.01$ として提案法を適用する。10 等分された各基準値の区分で考えられる増分効用の範囲をできる限り一様に大きくするように、整合するすべての効用関数の集まりを切り出すことになる。増分効用の範囲をできる限り一様に大きくするので、各区分で効用が増加する加法的効用関数の集まりが求められる。中央値の総和で正規化すると、各基準の最大効用値は生産管理能力 (g_1) が 0.305908、国際経験 (g_2) が 0.191332、人柄 (g_3) が 0.502764 となった。やはり、国際経験 (g_2) の影響力が小さいが UTA 法の結果より十分に大きい。生産管理能力 (g_1) と人柄 (g_3) の影響度は UTA 法の結果と似ているが、人柄 (g_3) の重みがやや小さい。得られた区間効用関数から求められる人材間の選好関係のハッセ図を図 1(c) に示す。細線と太線の区別は先と同様である。UTA 法のように弱順序関係にならないが、 UTA^{GMS} 法よりかなり多くの人材間で選好が成立している。これらの選好が確実に成立するとは限らないが、(i) 情報として得られた選好対に対し

て区間効用の強い大小関係の成立を要請し、(ii) すべての区分で一様に効用が大きくなるように区間効用関数を求め、(iii) 区間効用の強い大小関係が成立する場合のみ、選好が成立すると推定しているという 3 点で、意味づけられた推定になる。

謝辞 JSPS 科研費 26350423 の助成に謝意を表します。

参考文献

- [1] E. Jacquet-Lagrange, J. Siskos, “Assessing a set of additive utility functions for multicriteria decision-making, the UTA method,” *European Journal of Operational Research*, 10, pp.151–164 (1982)
- [2] S. Greco, et al., Robust ordinal regression, in: M. Ehrgott et al. (eds.), *Trends in Multiple Criteria Decision Analysis*, Springer, NY, pp.241–283 (2010)
- [3] M. Kadzinski, S. Greco, R. Słowiński, “Extreme ranking analysis in robust ordinal regression,” *Omega*, 40, pp.488–501 (2012)
- [4] M. Inuiguchi, et al., Interval ordinal regression for multi-attribute decision aiding, in: *Proceedings of 8th NACA*, Yokohama Publishers, pp.177–206 (2015)

連絡先

大阪大学大学院基礎工学研究科 乾口 雅弘

E-mail: inuiguti@sys.es.osaka-u.ac.jp

確率的労働市場における社会格差

Social inequalities in probabilistic labor markets

○¹ 陳 鶴, ² 井上 純一

○¹ He Chen, ² Jun-ichi Inoue

^{1,2} 北海道大学大学院

^{1,2} Hokkaido University

Abstract: As an inequality manifestation, job-matching problems widely exist in today's labor markets. In order to simulate this kind of inequality, we improve the probabilistic model for Japanese university graduates, which we have introduced in our previous studies. We introduce the concept of feature vectors to measure the complex selection abilities of students. And by model simulation, we evaluate some inequalities in labor markets from both micro and macro views, such as acceptance ratio, unemployment, and Gini index.

1 はじめに

本研究では、我々が既に提案した労働市場の確率モデル [1, 2, 3] における企業と学生のジョブマッチング・プロセスを改良する。そして、計算機実験を通じて、大卒労働市場における格差現象を考察する。

2 労働市場のモデリング

2.1 提案モデル

まずは、我々が既に提案したモデルを簡単に説明する。この「人工労働市場」は以下のジョブマッチング・プロセス条件を課す。

- ① 各企業の募集定員を定め、さらに年次ごとに変動させない。
- ② 各年度で定員を大きく越えるまた下回る企業は次年度の集学生能力が低下傾向にある。
- ③ 各企業はランキングを付けられ、学生が企業を選択の際に共有する。
- ④ 学生の企業選択は条件付き情報論的エントロピー最大化の意味で多様性を持つ。

上記仮説に基づく企業ごとに定義される労働市場の集学生確率分布は次式で与えられる。

$$P_k(t) = \frac{\exp[\gamma \log(1 + \frac{k}{K}) - \beta |v_k^* - v_k(t-1)|]}{\sum_{k=1}^K \exp[\gamma \log(1 + \frac{k}{K}) - \beta |v_k^* - v_k(t-1)|]} \quad (1)$$

ここで、システムを構成する K 社企業と N 人学生は $k = 1, 2, \dots, K$ と $i = 1, 2, \dots, N$ で区別し、企業 k の募集定員を v_k^* と年度 t において企業 k に希望する学生数を $v_k(t)$ で与えられる。従って、社会総求人数を $V = \sum_{k=1}^K v_k^*$ と求人倍率を $\alpha = V/N$ である。さらに、 γ が企業のランキング重要度、 β が市場履歴重要度である。

2.2 企業の学生選考基準

従来の提案した数理モデルでは、企業は学生から回収したエントリの中ランダムまたは成績順位に募集定員数まで内定を出す単純な設定であった。しかし、現実社会では従業員に多様性をもたせる企業が多数である。この多様性選考を対応するため、ここでは、各企業 k 固有の評価基準ベクトルを導入する。具体的には評価は互いに直交する M 個のベクトル： $e_1 = (1, 0, \dots, 0), \dots, (0, \dots, 0, 1) = e_M$ を用いて

$$\mathbf{Y}^{(k)} = y_1^{(k)} e_1 + \dots + y_M^{(k)} e_M = (y_1^{(k)}, \dots, y_M^{(k)}) \quad (2)$$

で与えられる。ここで、 e_1 が学生の評価種目であるとする、 $y_1^{(k)}$ は企業 k が各成分 e_1 の重視度である。一方、学生 i も個人の「特徴ベクトル」を

$$\mathbf{X}^{(i)} = x_1^{(i)} e_1 + \dots + x_M^{(i)} e_M = (x_1^{(i)}, \dots, x_M^{(i)}) \quad (3)$$

で与えられる。ここで、 $x_1^{(i)}$ は学生 i が各成分 e_1 の点数化評価である。従って、学生 i が企業 k から受ける「評価」はベクトル $\mathbf{X}^{(i)}$ の $\mathbf{Y}^{(k)}$ への射影を

$$\epsilon^{(i,k)} = \mathbf{X}^{(i)} \cdot \mathbf{Y}^{(k)} = x_1^{(i)} y_1^{(k)} + \dots + x_M^{(i)} y_M^{(k)} \quad (4)$$

で与えられる。学生 i の点数 $\epsilon^{(i,k)}$ が企業 k にエントリーした総学生数の中上位 vk^* 番目にランクされていれば、学生 i は企業 k から内定を受ける。そして、本研究では $y_m^{(k)}$ が独立に平均 1, 分散 σ^2 の正規分布：

$$f(y_m^{(k)}) = \frac{1}{\sqrt{2\pi}\sigma} \exp \left[-\frac{(y_m^{(k)} - 1)^2}{2\sigma^2} \right] \quad (5)$$

に従い、 $x_m^{(i)}$ も独立に平均 0, 分散 τ^2 の対数正規分布：

$$f(x_m^{(i)}) = \frac{1}{\sqrt{2\pi}\tau x_m^{(i)}} \exp \left[-\frac{\log^2 x_m^{(i)}}{2\tau^2} \right] \quad (6)$$

に従う。このとき、パラメータ (σ, τ) が我々の労働市場モデルにおける格差評価指標である。

3 計算機実験

3.1 内定獲得率

各学生の内定獲得率 r_i を採用し、

$$r_i = \frac{\text{獲得内定数}}{\text{提出エントリー数}} = \frac{s_i}{a_i} \quad (7)$$

$$s_i = \left[\lim_{M \rightarrow \infty} \frac{1}{M} \sum_{t=1}^M \sum_{k=1}^K s_{ik}(t) \right] \quad (8)$$

$$a_i = \left[\lim_{M \rightarrow \infty} \frac{1}{M} \sum_{t=1}^M \sum_{k=1}^K a_{ik}(t) \right] \quad (9)$$

で定義する。そして、内定獲得率と学生成績順位の関係 $r_i (i = 1, \dots, N)$ を図 1 に示す。この図より、各種企業の学生選考基準では、内定獲得率 ($r = 1$) の学生の割合は平均提出エントリーの増加とともに減少する。紙面の制約のため、すべての計算結果そして詳しい説明を載せることはできない。

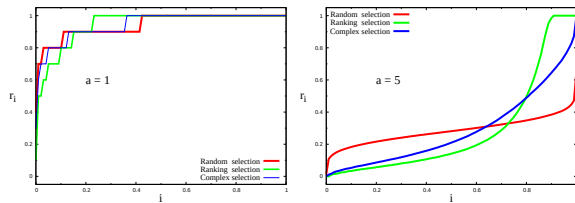


図 1: 内定獲得率の学生成績順位依存性。 $\alpha = 1.61$, $\beta = 1$, $\tau = 1$, $\sigma = 1$ と選び、システムサイズ $N = 42,000$, $K = 1,350$ とした。各点は 10 回独立試行による平均。 a は学生の平均提出エントリー数である。(以下同じ)

3.2 ジニ係数

我々過去の研究では、ジニ係数 [4] が $G = \frac{1}{N^2\mu} \sum_{i=1}^N (2i - N - 1)r_i$ から算出できる。 μ は特徴パラメータとして $\mu \equiv (1/N) \sum_{i=1}^N r_{i^*}$ で満たす。図 2 では、ジニ係数と企業ランキング重要度の関係を示す。

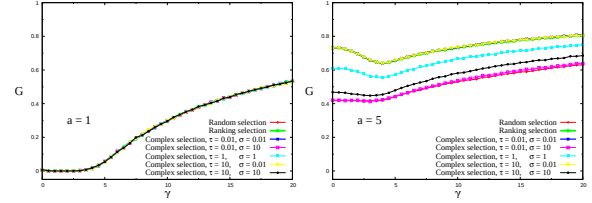


図 2: ジニ係数の企業ランキング重要度依存性。

4 おわりに

本稿では従来の労働市場の数値モデルを改良し、内定獲得率やジニ係数などに関する大きな格差が創発することを示した。紙面の制約のため、詳細は講演当日に報告したい。

5 謝辞

本研究は JSPS 科研費 26・1059 の助成を受けたものです。

参考文献

- [1] H.Chen, J.Inoue, Statistical Mechanics of Labor Markets, Econophysics of systemic risk and network dynamics, New Economic Windows Series, Springer-Verlag(Italia), pp.157-171, 2012.
- [2] H.Chen, J.Inoue, Dynamics of probabilistic labor markets: statistical physics perspective, Lecture Notes in Economics and Mathematical Systems 662, Managing Market Complexity, Springer, pp.53-64, 2012.
- [3] H.Chen, J.Inoue, Learning Curve for Collective Behavior of Zero-Intelligence Agents in Successive Job-Hunting Processes with a Diversity of Jaynes-Shannon's MaxEnt Principle, Evolutionary and Institutional Economics Review, Vol.10, No.1, pp.55-80, 2013.
- [4] C. Gini, Concentration and dependency ratios(1909)(in Italian), Rivista di Politica Economica 87, pp.769-789, 1997.

連絡先

北海道大学 大学院情報科学研究科
〒 060-0814 北海道札幌市北区北 1 4 条西 9 丁目
E-mail: ¹chen@complex.ist.hokudai.ac.jp
E-mail: ²j_inoue@complex.ist.hokudai.ac.jp

A Simplified Tracking Method based on Compressive Tracking

○Sun Fei¹, Watada Junzo²

^{1,2}Waseda University

¹sunfei@fuji.waseda.jp

²watada@waseda.jp

Abstract: The objective of the paper is to develop effective and efficient models for real-time tracking. However, it is difficult due to pose variation, illumination change, occlusion, and motion blur. Traditional tracking algorithms often update models by using samples extracted from nearly received frames. It is a weak point that appearance models are data-dependent. Intending to mitigate this weak point, we found tracking algorithms of an appearance model based on features extracted from the multi-scale image feature space with data-independent basis. One of these algorithms is named compressive tracking. Compressive tracking algorithm bases on the theory of compressive sensing. This algorithm I proposed gives a complete framework to solve tracking problem. The proposed algorithm realizes real-time performance favorably against state-of-the-art algorithms on challenging sequences in terms of efficiency, accuracy and robustness. This paper analyzes this method and tries to find out which factor makes this algorithm perform well, some improvement will be added based on the analyzed results.

1. Introduction

Tracking a moving object in video stream, as an important application in computer vision area, has variety of uses in medical imaging, intelligent human computer interaction, traffic guidance, and other areas. Traditional tracking methods usually can be divided into two groups. Based on their appearance models, one kind is generative tracking algorithm [2-6] and the other kind is discriminative tracking algorithm [7-11]. The first one learns a template to represent the target object, then uses the template to search for the video areas. KLT method, Meanshift method, and IVT method are included in this kind of algorithms. The second one regards tracking problems as a binary classification problem.

In Zhang's paper [1], they propose an effective and efficient tracking algorithm with an appearance model based on features extracted in the compressed domain. The main components of our compressive tracking algorithm are shown by Figure 1. Their appearance model is generative as the object can be well represented based on the features extracted in the compressive domain. It is also discriminative because they use these features to separate the target from the surrounding background via a naive Bayes

classifier. In their appearance model, features are selected by an information-preserving and non-adaptive dimensionality reduction from the multi-scale image feature space based on compressive sensing theories [12, 13]. It has been demonstrated that a small number of randomly generated linear measurements can preserve most of the salient information and allow almost perfect reconstruction of the signal if the signal is compressible such as natural images or audio [12-14]. They use a very sparse measurement matrix that satisfies the restricted isometry property (RIP) [15], thereby facilitating efficient projection from the image feature space to a low-dimensional compressed subspace. For tracking, the positive and negative samples are projected (i.e., compressed) with the same sparse measurement matrix and discriminated by a simple naive Bayes classifier learned online. The proposed compressive tracking algorithm runs at realtime and performs favorably against state-of-the-art trackers on challenging sequences in terms of efficiency, accuracy and robustness.

2. THEORY OF COMPRESSIVE TRACKING

A. Measurement matrix

Random Gaussian matrix $R \in \mathbb{R}^{n \times m}$ where $r_{ij} \sim N(0,1)$

is recently used as typical measurement matrix in numerous work. However, this kind of matrix is dense, if use this in our tracking case, it would cost huge quantity of memory and computational loads. As a solution Zhang adopts a very sparse random measurement matrix which is proved by Achlioptas [17].

$$r_{ij} = \sqrt{s} \times \begin{cases} 1 & \text{with probability } \frac{1}{2s} \\ 0 & \text{with probability } 1 - \frac{1}{s} \\ -1 & \text{with probability } \frac{1}{2s} \end{cases}$$

This type of matrix satisfies the Johnson-Lindenstrauss lemma with $s = 2$ or 3 . In addition, this matrix is friendly for computing cause it requires only one uniform random generator. Especially when $s = 3$ this type of matrix is sparse enough that two thirds of the computation can be avoided. What's more, Li et al. show that for $s = O(m)$ ($x \in \mathbb{R}^m$) or even $s = m / \log(m)$ the random projections are as accurate as the random Gaussian matrix. So in consideration of efficiency, CT algorithm sets $s = m / 4$. For each row there are at most four entries need to be computed. That is to say, only nonzero entries of R can make sure the accuracy rate, and the computational complexity is only $O(cn)$ which is very low, the memory requirement also light.

B. Random projection

A random matrix $R \in \mathbb{R}_{n \times m}$ whose rows have unit length projects data from the high-dimensional image space $x \in \mathbb{R}_m$ to a lower-dimensional space $v \in \mathbb{R}_n$

$$v = R_x \quad (1)$$

where $n \gg m$. Ideally, we expect R provides a stable embedding that approximately preserves the distance between all pairs of original signals. The Johnson-Lindenstrauss lemma [16] states that with high probability the distances between the points in a vector space are preserved if they are projected onto a randomly selected subspace with suitably high dimensions.

Baraniuk et al. [18] proved that the random matrix satisfying the Johnson-Lindenstrauss lemma also holds true for the restricted isometry property in compressive

sensing. Therefore, if the random matrix R in (1) satisfies the Johnson-Lindenstrauss lemma, we can reconstruct x with minimum error from v with high probability if x is compressive such as audio or image. We can ensure that v preserves almost all the information in x . This very strong theoretical support motivates us to analyze the high-dimensional signals via its low-dimensional random projections.

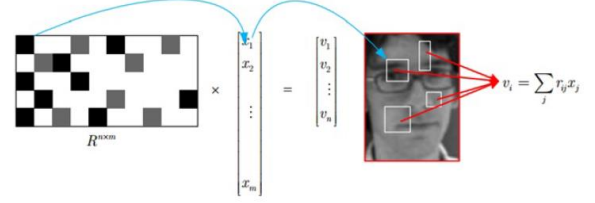


Fig 1 Graphical representation of compressing a high-dimensional vector x to a low dimensional vector v .

C. CT framework

We can learn from the compressive sensing theory, if we want to compress a high dimensional signal X into a low dimensional signal through the help of measurement matrix R , and the low dimensional signal contain almost all the information of X in the same time. In that case, we should make sure the measurement matrix R satisfies restricted isometry property (RIP). So before tracking, we should find an appropriate matrix R with RIP, and this matrix should be not dense in consideration of the memory and computational loads. In CT, Zhang adopted a very sparse matrix which has already been introduced above.

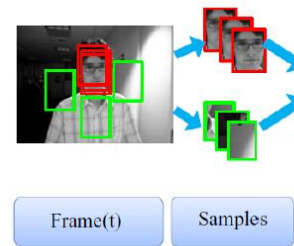


Fig 2 Get positive sample at tracking area and randomly get negative sample around the tracking area.

Firstly, choose an area or object you want to track. Then get positive sample of the tracking object which is in red rectangle and negative samples will be selected around positive sample which is in green rectangle.

Then we analyze the features of different samples. After get the image features, we do random projection to compress the feature data. After compressing, we build a classifier based on compressed data to check next frame.

Finally, get an area has the maximal classifier response. That area is the result in this frame. Then do the same work again, and update the classifier based on this frame. In the original method, a large set of Haar-like features are compressively sensed with a very sparse measurement matrix.

3. MY WORK

A. ORB feature

I recommend ORB [20] features instead of Haar-like features. The original tracking method do have some situation of losing the wanted target, especially in these cases: the camera zoom in or out, the target object has fiercely rotation and so on.

However, the ORB feature has good quality against the rotation and shape change of object. What's more, the calculation speed is quite well.

This feature is a combination of Oriented FAST and Rotated BRIEF. Both these techniques are attractive because of their good performance and low cost. That is the reason why I choose this feature.

B. Simplified Classifier

For each sample $Z \in \mathbb{R}^m$ $v = (v_1, \dots, v_n)^T \in \mathbb{R}^n$ with $m \gg n$ is the low-dimensional representation. Assuming all elements in v are independently distributed. Then we can build a naive Bayes classifier [19].

$$H(v) = \log \left(\frac{\prod_{i=1}^n p(v_i|y=1)p(y=1)}{\prod_{i=1}^n p(v_i|y=0)p(y=0)} \right) = \sum_{i=1}^n \log \left(\frac{p(v_i|y=1)}{p(v_i|y=0)} \right) \quad (2)$$

In this equation, we assume $p(y=1) = p(y=0)$; and $y \in 0,1$ is a binary variable represents the sample label.

As the random projections of high dimensional random vectors are nearly Gaussian mentioned before. So the distributions

$p(v_i|y=1)$ and $p(v_i|y=0)$ in (1) can be considered as Gaussian distributed with four parameters where

$$p(v_i|y=1) \sim N(\mu_i^1, \sigma_i^1), \quad p(v_i|y=0) \sim N(\mu_i^0, \sigma_i^0) \quad (3)$$

The scalar parameters in (2) are gradually updated

$$\begin{aligned} \mu_i^1 &\leftarrow \lambda_i^1 + (1-\lambda)\mu^1 \\ \sigma_i^1 &\leftarrow \sqrt{\lambda(\sigma_i^1)^2 + (1-\lambda)(\sigma^1)^2 + \lambda(1-\lambda)(\mu_i^1 - \mu^1)^2} \end{aligned} \quad (4)$$

where $\lambda > 0$ is a learning parameter. The above equations can be easily derived by maximal likelihood estimation.

C. method in pseudo-code

input : image sequences;

output : the output image sequences with a target frame according to the tracking algorithm

1. Init initiate
2. Gain the random measurement matrix R
3. Sample two sets of image patches. One set is positive set, the other is negative sample set.
4. Extract the features with these two sets of samples by the nonzero element of the R to the sample image pixel values of the weighted sum and update the classifier parameters according to equation (4)
5. for $i = 2$: num(num is the length of image sequence)
6. Sample a set of image patches. Extract the features with these samples.
7. Obtain a maximum classification score r with a Batesian classifier using equation(1)
8. perform step 3 and step 4;
9. End for

4. Conclusion

This paper introduced an improved compressive tracking method and analyzed some elements in this algorithm to see how do these elements effect the tracking result. We used random measurement directly project the high-dimensional. we proposed a simple yet

robust tracking algorithm with an appearance model based on non-adaptive random projections that preserve the structure of original image space. A very sparse measurement matrix was adopted to efficiently compress features from the foreground targets and background ones.

参考文献

- [1] Zhang K, Zhang L, Yang M H. 2012. Real-time compressive tracking. *Computer Vision-ECCV 2012*. Springer Berlin Heidelberg, 864-877.
- [2] Black, M., Jepson, A.: Eigenttracking: Robust matching and tracking of articulated objects using a view-based representation. *IJCV* 38, 63-84 (1998)
- [3] Jepson, A., Fleet, D., Maraghi, T.: Robust online appearance models for visual tracking. *PAMI* 25, 1296-1311 (2003)
- [4] Ross, D., Lim, J., Lin, R., Yang, M.-H.: Incremental learning for robust visual tracking. *IJCV* 77, 125-141 (2008)
- [5] Li, H., Shen, C., Shi, Q.: Real-time visual tracking using compressive sensing. In: *CVPR*, pp. 1305-1312 (2011)
- [6] Mei, X., Ling, H.: Robust visual tracking and vehicle classification via sparse representation. *PAMI* 33, 2259-2272 (2011)
- [7] Avidan, S.: Support vector tracking. *PAMI* 26, 1064-1072 (2004)
- [8] Collins, R., Liu, Y., Leordeanu, M.: Online selection of discriminative tracking features. *PAMI* 27, 1631-1643 (2005)
- [9] Grabner, H., Grabner, M., Bischof, H.: Real-time tracking via online boosting. In: *BMVC*, pp. 47-56 (2006)
- [10] Grabner, H., Leistner, C., Bischof, H.: Semi-supervised On-Line Boosting for Robust Tracking. In: Forsyth, D., Torr, P., Zisserman, A. (eds.) *ECCV 2008, Part I*. LNCS, vol. 5302, pp. 234-247. Springer, Heidelberg (2008)
- [11] Babenko, B., Yang, M.-H., Belongie, S.: Robust object tracking with online multiple instance learning. *PAMI* 33, 1619-1632 (2011)
- [12] Donoho, D.: Compressed sensing. *IEEE Trans. Inform. Theory* 52, 1289-1306 (2006)
- [13] Candes, E., Tao, T.: Near optimal signal recovery from random projections and universal encoding strategies. *IEEE Trans. Inform. Theory* 52, 5406-5425 (2006)
- [14] Wright, J., Yang, A., Ganesh, A., Sastry, S., Ma, Y.: Robust face recognition via sparse representation. *PAMI* 31, 210-227 (2009)
- [15] Candes, E., Tao, T.: Decoding by linear programming. *IEEE Trans. Inform. Theory* 51, 4203-4215 (2005)
- [16] Yan He, An improved real-time compressive tracking method, *ICIMCS '13 Proceedings of the Fifth International Conference on Internet Multimedia Computing and Service Pages* 380-385 (2013)
- [17] Achlioptas, D.: Database-friendly random projections: Johnson-Lindenstrauss with binary coins. *J. Comput. Syst. Sci* 66, 67-687 (2003)
- [18] Baraniuk, R., Davenport, M., DeVore, R., Wakin, M.: A simple proof of the restricted isometry property for random matrices. *Constr. Approx* 28, 253-263 (2008)
- [19] Ng, A., Jordan, M.: On discriminative vs. generative classifier: a comparison of logistic regression and naive bayes. In: *NIPS*, pp. 841-848 (2002)
- [20] Rublee, E. ; Willow Garage, Menlo Park, CA, USA ; Rabaud, V. ; Konolige, K. ; Bradski, G.: ORB: An efficient alternative to SIFT or SURF. In: *ICCV* , pp. 2564- 2571(2011)

連絡先

Sun Fei

Waseda University S261,Hibikino,Fukoka,Japan

E-mail:sunfei@fuji.waseda.jp

Detecting Pedestrians and Vehicles in Traffic Scenes Based on Boosted-HOG Features and SVM

Diqing Sun, Junzo Watada

Waseda University

Watada Lab

Email: sundiqing@moegi.waseda.jp junzow@osb.att.ne.jp

Abstract: This paper presents a popular method called boosted hog features to detect pedestrians and vehicles in static images. We compared the differences and similarities of detecting pedestrians and vehicles, then we selected boosted hog features to get an satisfying result. In the part of detecting pedestrians, Histograms of Oriented Gradients (HOG) feature is applied as the basic feature due to its good performance in various kinds of background. On that basis, we create a new feature with boosting algorithm to obtain more accurate result. In the part of detecting vehicle, we select the shadow underneath vehicle as the feature, so we can utilize it to detect vehicles in daytime. The shadow is an important feature for vehicles in traffic scenes. The region under vehicle is usually darker than other objects or backgrounds and can be segmented by setting a threshold. Finally, experimental results show that the detection of pedestrians and vehicles using boosted hog feature and linear SVM combines the advantages of hog feature and adaboost classifier, and can achieve better detection results than the detector by using conditional HOG features. At the end, the paper shows its efficiency and effectiveness using to application in real situations.

1. Introduction

Nowadays, traffic accidents are universally decreasing due to many advanced safety vehicle systems. To prevent the occurrence of a traffic accident, the first function that a safety vehicle system should accomplish is the detection of objects in a traffic situation. Therefore, the detection of pedestrians and vehicles in image plays a more and more important role in today's computer vision science.

The most popular method of detecting pedestrians is the HOG(Histograms of Oriented Gradients) feature[1]. We have investigated various methods such as HOG, haar-like, SIFT(The Scale Invariant Feature Transform) and DPM(deformable part model)[2][3]. In these methods, haar-like feature can be easily influenced by various factors such as illumination and movement of object. Owing to its low accuracy, haar-like feature can be only applied to detection of static objects. SIFT feature has higher accuracy and robustness, but it need longer run-time to realize computation. DPM is an advanced algorithm of HOG feature. It has the highest accuracy in pedestrians' detection area. But training samples in DPM is difficult than other method. Consequently, boosted-HOG method was chosen for this research due to its simplicity and high efficiency.

Although HOG feature achieves excellent detection rates, the high dimensionality of its feature vector poses as one of its disadvantages. The large size of the feature vector limits the number of training samples and increases the computation cost in classification. To Improve the accuracy of the research and shorten computation time, a modified hog feature called boosted hog feature is proposed[4]. Our method adopts HOG as the basic feature, takes the advantage of its highly discriminative power, and creates a much simpler feature with adaboost.

We then train a linear SVM to perform the classification with the obtained feature. As a result the computation cost and storage requirement in the SVM classification is much lower.

For the aspect of detecting vehicles, less related papers were published comparing with the detection of pedestrians. We investigated some local features like Haar-like features[5][6] and control points[7]. Different local features have different classification abilities. The HOG feature has been widely used in pedestrian detection because it wasn't sensitive to variance of light, direction, and size[8][9]. Unlike pedestrians, vehicles could do plenty of scale transformation and rotation transformation, but hog feature does not work well in scale invariant and rotation invariant, so we decided to use the shadow underneath vehicle as the feature in recognition[10]. The shadow was segmented accurately by using boosted hog method. The initial candidates were generated by combining horizontal and vertical edge feature of shadow, and these initial candidates were further verified by using adaboosted algorithm. Finally, linear support vector machine is used to estimate whether this object is a pedestrian or vehicle.

The remaining content is written as follows: section II explains the backgrounds of HOG , boosted HOG and SVM. Section III analyzes the performance and results of the experiment, and finally Section IV summarizes this paper.

2. Theory of detection

2.1 HOG Feature

The main method of the Histogram of Oriented Gradient descriptors is that local object appearance and shape within an image can be described by the distribution of

intensity gradients or edge directions. The implementation of these descriptors can be achieved by dividing the image into small connected regions, called cells, and for each cell computing a histogram of gradient directions or edge orientations for the pixels within the cell. The combination of these histograms can represent the descriptor at last. As shown in Fig.1, HOG feature can be extracted as the follow four steps:

- 1) the first step is to compute image gradients for each pixel, and the gradients are used to calculate their magnitudes and orientations;
- 2) the second step is to accumulate weighted votes for gradient magnitude into orientation bins over spatial cells. The weighted votes are interpolated bili-nearly between the neighboring bins in both orientation and position in order to avoid aliasing. Linear interpolation distributes the weight into two nearest neighboring bins and tri-linear interpolation distributes the weight into eight surrounding bins;
- 3) the third step is to normalize contrast for each block;
- 4) the last step is to collect the histograms for all overlapping blocks over detection window.

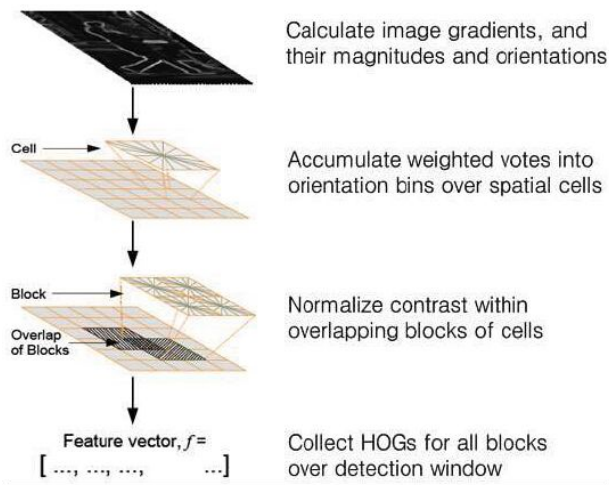


Figure 1. Overview of HOG Feature Extraction

2.2 Boosted HOG feature

The resulted HOG feature in each cell contains important information on how to separate pedestrians and vehicles from other objects, yet redundant information may also be included in the feature. Now the adaboost[11] is applied to learn a new feature from the HOG feature at hand. For each cell, the set of weak classifiers should be firstly created. As the HOG is a histogram with bins indicating local gradient distribution, we compare the value on one bin with a threshold to determine whether the image contains the objects. The adaboost algorithm starts with assigning weights to the training samples. In each iteration, the weak classifier with the least error rate is selected, and is given a weight to determine its

importance in the final classifier. Before the next iteration begins, the weights of those misclassified

samples are increased so that the algorithm can focus on those hard samples. The algorithm of Adaboost is given in Table 1.

2.3 Linear SVM training

The output of the adaboost training on each cell is a strong classifier and the strong classifiers from all the cells of an input image patch are combined to construct a feature vector. A linear SVM classifier[12] is then trained over the feature vectors for the final classification.

The procedure of SVM training for classification is summarized as follows.

Step 1) Train the adaboost classifier using the training samples and compute the boosted HOG features.

Step 2) With the obtained features, train the SVM classifier with all the samples. Test the trained SVM on the same training samples. Find all the false positive and false negative samples.

Step 3) Use the false positive samples and false negative samples to train the adaboost classifiers as in step 1).

Step 4) Combine the obtained adaboost classifiers to build the feature vector and train the final linear SVM classifier to classify objects.

The final classifier is constructed as a weighted combination of weak classifiers selected in each iteration just work as fig 2.

2.4 Pedestrian-Vehicle Context

A lot of detection of pedestrians are located around vehicles in traffic scenes, . It is possible to use the pedestrian-vehicle relationship to infer whether the detection is true or false positive. For example, if we know the location of vehicles, the detection above a vehicle, and detection at the wheel position of a vehicle can be safely removed. Fortunately, the vehicles are more easier to be localized than pedestrians[13][14]. After the detection of vehicles, we can check the result of pedestrian detection and remove some negative result to improve the detection rate.

3.Experimental results and analysis

3.1 Training Database

There are a lot of training sample libraries in the pedestrian detection field. We choose Caltech Pedestrian Detection Benchmark[15][16] as the database. The Caltech Pedestrian Dataset consists of approximately 10 hours of 640x480 30Hz video taken from a vehicle driving through regular traffic in an urban environment. About 250,000 frames (in 137 approximately minute long segments) with a total of 350,000 bounding boxes

Table 1 RealAdaboost algorithm

1. Suppose $S = \{(x_1, y_1), (x_2, y_2), \dots, (x_N, y_N)\}$ as a sample space, where $x \in X$ are feature vectors and $y \in \{-1, +1\}$ are labels.
2. D_t is the initial distribution
$$D_1(n) = \frac{1}{N}$$
3.
 - for $t=1$ to T do
 - for $m=1$ to M do
 - Compute probability distribution W_t of weak classifier $h_i(x)$

$$W_+^j = \sum_{i,j \in J \wedge y_i = +1} D_t(i)$$

$$W_-^j = \sum_{i,j \in J \wedge y_i = -1} D_t(i)$$
 - Compute estimation Z_m

$$Z_m = 2 \sum_j \sqrt{W_+^j W_-^j}$$
 - Select the smallest Z_m of the weak classifier
$$h_t = \arg \min Z_{t,m}$$
 - Renew sample weight $D_t(i)$

$$h_t(x_i) = \frac{1}{2} \ln \frac{W_+^j + \varepsilon}{W_-^j + \varepsilon}$$

$$D_{t+1}(i) = D_t(i) \exp[-y_i h_t(x_i)]$$
4. The final strong classifier is given by
$$H(x) = \text{sign} \left(\sum_{t=1}^T h_t(x) \right)$$

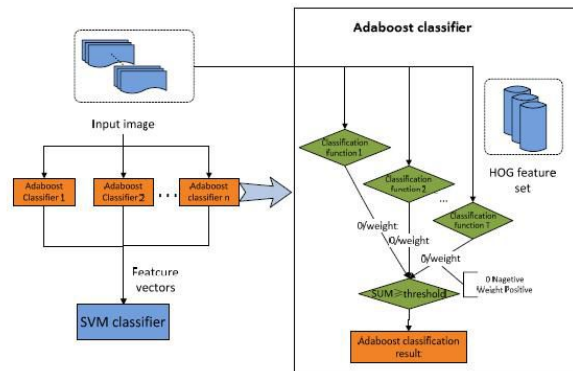


Figure 2. The process of the SVM classifier: input images are divided into cells, each of which corresponds to a strong classifier trained by adaboost using HOG features. And all adaboost outputs are combined to establish the final SVM feature vector.

and 2300 unique pedestrians were annotated. The

annotation includes temporal correspondence between bounding boxes and detailed occlusion labels.

At last, we choose 100 images in which pedestrians and vehicles both exist as testing materials.

3.2 Pedestrian Detection

We use Caltech Pedestrian Dataset as the training date. The training data (set00-set05) consists of six training sets (~1GB each), each with 6-13 one-minute long seq files, along with all annotation information. we choose 500 positive image and 300 negative image from the video and all their size is 75*75 pixels. After training, we select 100 scenes from testing data(set06-set10). Besides, we test these images using conditional hog feature, the result is as table 2.

Table 2 Performance of the pedestrian detection in traffic scenes

Feature Type	Testing Number	Positive Result	Negative Result	Detected Rate
Conditional HOG	100	86	14	86%
Proposed Method	100	75	25	75%

3.3 Vehicle Detection

In the detection of vehicles, The shadow is the main characteristic for the detection of vehicles. The region under vehicle is usually darker than other objects or backgrounds and could be segmented by setting a threshold. Our work is to detect these shadow so we can get the contour of vehicles.



Figure 3. Shadow Segmentation in the Traffic Scene

As the Figure 3 shows, we can detect the shadows easily due to the shadow is brighter than other background in grey-scale image. After obtaining the shadow regions, we can generate the initial contour of vehicle by combining horizontal and vertical edge feature of shadow. The images of vehicles that are used the experiment are compiled partly from the internet and partly from the images in Caltech Pedestrian Dataset. A total of 525 images of the vehicles are collected and 125 of them are

negative samples. These cropped sample images are resized to 45*45 in the training. We also apply 100 images from the testing set on Caltech Pedestrian Dataset as the testing images. Some of the samples are shown in Figure 4.

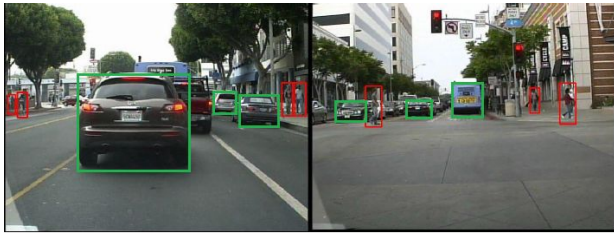


Figure 4. Detection of Pedestrian and Vehicle on Caltech Pedestrian Benchmark

4. Conclusions

In this paper we proposed a combined method for pedestrian detection and vehicle detection in static images. The original HOG feature is adopted first and we trained it with adaboost algorithm to obtain the boosted HOG feature. A linear SVM was trained to accomplish the detection. Shadow feature underneath vehicle is the important feature of vehicle detection in road image and can be used to generate hypotheses quickly and reliably. HOG feature and SVM has good generalization ability and can remove the non-vehicle regions such as buildings, clouds, flowers, fence presuming effectively. The detection has limitations that it can only work in daytime and we can't solve the overlapping problem because we detect the vehicles using the shadow underneath vehicles. The experiments on Caltech Pedestrian Dataset showed our proposed method has advantage in detection accuracy comparing with conditional HOG method. In the future, we will propose some new methods to help us detect pedestrians and vehicles more accuracy based on the relation between pedestrians and vehicles in traffic scenes.

References

- [1] N. Dalal, B. Triggs, Histograms of oriented gradients for human detection, Proc. IEEE Conference on Computer Vision and Pattern Recognition, 886-893, 2005
- [2] M. Enzweiler and D. M. Gavrila, Monocular pedestrian detection: Survey and experiments [J]. IEEE Transactions on Pattern Analysis and Machine Intelligence, vol. 31, pp. 2179-2195, 2009
- [3] Felzenszwalb, P.F., Cascade object detection with deformable part models, Computer Vision and Pattern Recognition, 2241 - 2248, 2010
- [4] Zhen-Rui Wang, Pedestrian Detection Using Boosted HOG Features, Intelligent Transportation Systems, 1155 - 1160, 2008
- [5] X.Wen, H.Yuan,C.Yang, et al. Improved Haar-Wavelet Feature Extraction Approachs for Vehicle Detection. IEEE Intelligent Transportation Systems Conference, 1050-1054,2007
- [6] R.Lienhart, J.Maydt. An extended set of haar-like features for rapid object detection. Proceedings of International Conference on Image Processing, 900-903, 2002
- [7] Yotam Abramson, Bruno Steux, Hicham Ghorayeb. YEF Real-time objetet detection. Proceedings of International Workshop on Automatic Learning and Real-time, 102-112, 2005
- [8] Q. Zhu, M. C. Yeh, and K.T. Cheng, Fast Human Detection Using a Cascade of Histograms of Oriented Gradients, Proc. IEEE International Conference on Computer Vision and Pattern Recognition, vol.2, pp.1491-1498, 2006.
- [9] F. Suard, A. R. Akotomamonjy, A. Bensrhair et al, Pedestrian Detection Using Infrared Images and Histograms of Oriented Gradients, Proc. Intelligent Vehicle Symp, pp. 206-212, 2006.
- [10] S. Han, Y. Han and H. Hahn.Vehicle Detection Method using Haar-like Feature on Real Time System[J]. Engineering and Technology, vol. 59, pp.455-459, 2009.
- [11] R. E. Schapire and Y. Singer, "Improved boosting algorithms using confidence-rated predictions", Machine Learning, 37, 3, 97-336, 1999
- [12] V.Vapnik, John Wiley & Sons , An overview of statistical learning theory,IEEE Transactions on Neural Networks, 988-999, 1998
- [13] P. Felzenszwalb, R. Girshick, D. McAllester, and D. Ramanan. Object detection with discriminatively trained part-based models. Pattern Analysis and Machine Intelligence, 1627-1645, 2010
- [14] A. Geiger, P. Lenz, and R. Urtasun. Are we ready for autonomous driving the kitti vision benchmark suite. Computer Vision and Pattern Recognition, 3354-3361, 2012
- [15] Piotr Dollar, Christian Wojek, Bernt Schiele, and Pietro Perona, Pedestrian Detection: An Evaluation of the State of the Art, Pattern Analysis and Machine Intelligence. 743-761, 2012
- [16] Dollar, P. Wojek, C. ; Schiele, B. ; Perona, P. Pedestrian detection: A benchmark, Computer Vision and Pattern Recognition, 304 - 311, 2009

Human detection and tracking using panoramic image modeling

Jiayi Wang
Graduate School of Information, Production
and System Waseda University
2-7 Hibiya, Wakamatsu Kitakyushu,
Fukuoka 808-0135 JAPAN
E-mail: wangjiayi@akane.waseda.jp

Junzo Watada
Graduate School of Information, Production
and System Waseda University
2-7 Hibiya, Wakamatsu Kitakyushu,
Fukuoka 808-0135 JAPAN
E-mail: junzow@osb.att.ne.jp

Abstract: Panoramic image mosaic is a technology to match a series of images which are overlapped with each other. Panoramic image mosaics can be used for different applications. Image mosaic has important values in various applications such as photogrammetry, computer vision, remote sensing image processing, medical image analysis and computer graphics. Image mosaics also can be used in moving object detection with a dynamic camera. After getting the panoramic background of the video for detection, we can compare every frame in the video with the panoramic background, and finally detect the moving object. To build the image mosaic, SURF algorithm is used in feature detection and OpenCV is used in the programming.

1. INTRODUCTION

Nowadays, in our daily life computers and cameras become cheaper and more widely used, and we are now using digital image almost everywhere. Panoramic image mosaic acts an pivotal role in this large use of digital equipment. A panorama is a single wide-angle image of the environment around the camera. However, as the size of a single photo and a single frame in a video is limited, image mosaic may be necessary in getting a panoramic background. With a panoramic background of a video or a series of images, we can get a complete image of the scene or get some help in object tracking with a dynamic camera. Image mosaic has important value in applications such as photogrammetry, computer vision, remote sensing image processing, medical image analysis, and computer graphics. Image mosaics also can be used in moving object detection with a dynamic camera. After getting the panoramic background of the video for detection, we can compare every frame in the video with the panoramic background, and finally detect the moving object.

We can use various methods to build the image mosaicking. Even though there is some difference in several steps, the general process is the same.

After image acquisition, we should do the preprocessing, such as the geometry calibration and denoising. Then we do the image registration. Image registration is to find the corresponding position of the template or feature point in the reference image with some matching algorithm.

Image fusion, or image stitching and image mosaicking, is the process of combining relevant information from two or more images into a single image. After detecting the pairs of corresponding points, their coordinates can be used to calculate the transformation matrix, get the transformation

relationship between the two frames and build the image mosaicking at last.

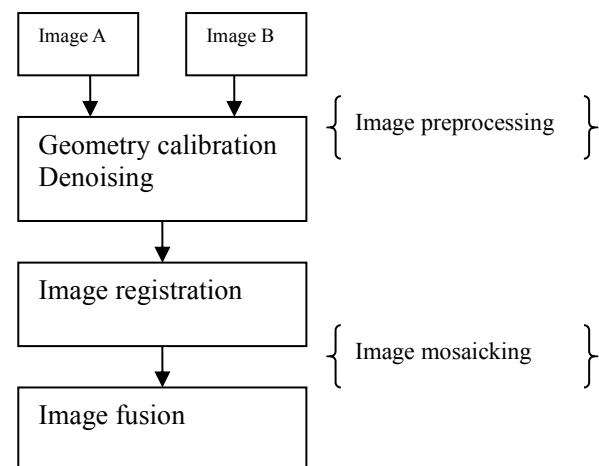


Figure 1 Steps of image mosaic

2. Image transformation

In image acquisition, the state of the image will be affected by the motion state of the camera. Generally speaking, the movement of the camera can be divided into the following basic movement:

Translation: While focal length is fixed, the camera only moves in the direction parallel to the imaging plane.

Zooming: Focal length is changed, that is to say, the distance between the camera and the imaging plane has been changed.

Panning: The camera is rotating around the y axis (vertical axis).

Tilting: The camera is rotating around the x axis (horizontal axis).

In general, while taking the video, the camera is not always on a tripod. For example, we may hold the camera in our hands. So the movement of the camera can not be described by only one of the above moves. Considering the movement of the camera in handheld shooting is very complex and the uncertainty of translation, zooming, rotation, the movement is always the composition of the motions above.

To deal with this and achieve image mosaicking of the image sequence in a same scene, we should determine the spatial relationship of the frames in the video. This is also the main work of image stitching. To determine the corresponding relationship between images, we should first know the corresponding relationship model and then the problem is changed to determining the parameter of the model. The most commonly used relationship models are rigid transformation model, affine transformation model and projective transformation model. Figure 1 shows the relationship transformation models.

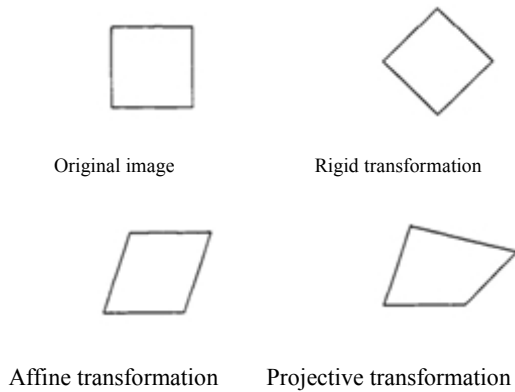


Figure 2 Relationship transformation models.

We know that a transformation model shows the relationship between two pictures and we can get the general information from the following two situations:

<1>. The two images are known, and we know that they are related. Then our job is to find the matrix M

<2>. The original image and relation matrix M are known, and we want to obtain the image after transformation. In this case, we can use Equation 1 to do the transformation easily and it can be clearly expressed by the matrix M or using the geometric relationship between the points in the two images.

We use Figure 3 as an example, the three point A, B and C is one-to-one mapping, and they still form a triangle even though the shape is not the same as

before. We can find three such pairs of points to calculate the transformation, and then apply it into the images.

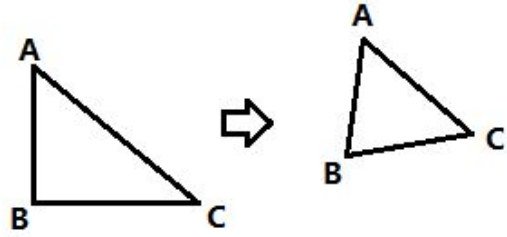


Figure 3 Image transformation

We can use homogeneous coordinates to describe the above relationship models with matrix model.

$$\begin{bmatrix} x' \\ y' \\ 1 \end{bmatrix} = \begin{bmatrix} m_0 & m_1 & m_2 \\ m_3 & m_4 & m_5 \\ m_6 & m_7 & 1 \end{bmatrix} \begin{bmatrix} x \\ y \\ 1 \end{bmatrix} \text{ or } X' = MX \quad (1)$$

The function of every parameter in matrix M is as follows:

Table 1 Function of parameter in matrix M

Parameter	Function
m_2	horizontal displacement
m_5	vertical displacement
m_0, m_1, m_3, m_4	scale and amount of rotation
m_6, m_7	deformation in horizontal and vertical directions

Using the coordinates of some pairs of corresponding points we can get in the following parts, parameters in the transformation matrix can be calculated.

In OpenCV, functions related with affine transformation generally refer to warpAffine and getRotationMatrix2D. The function warpAffine is used to realize some simple remapping and getRotationMatrix2D is used to get the rotation matrix

A. warpAffine

The function warpAffine is to do the affine transformation of the image according to Equation 2.

$$\text{dst}(x, y) = \text{src}(M11x + M12y + M13, M21x + M22y + M23) \quad (2)$$

The function prototype is :

```
void warpAffine(InputArray src, OutputArray dst, InputArray M, Size dsize, int flags=INTER_LINEAR, int borderMode=BORDER_CONSTANT, const Scalar & borderValue=Scalar())
```

The first argument *src* is the input image, the source image in other words.

The second parameter *dst* is where the calculation result after the function call is stored. Same size and type as the source image is required.

The third parameter *M* is a 2×3 transformation matrix.

The fourth parameter *dsize* represents the size of output image.

The fifth parameter *flags* is the identifier of interpolation method.

The sixth parameter *borderMode* is the boundary pixel mode and the default value `BORDER_CONSTANT`.

The seventh parameter *borderValue* is the value while boundary is constant and the default is value `Scalar()`, that is to say, 0.

`WarpAffine` requires that input and output images have the same data type, great resource cost and the output portion of the image can be unchanged

B. `getRotationMatrix2D`

The function `getRotationMatrix2D` calculates a two-dimensional rotation transformation matrix and the transformation will map the rotational center to itself.

```
Mat getRotationMatrix2D(Point2f center, double angle, double scale)
```

The first argument *center* means the rotation center of the source image.

The second parameter *angle* is the angle of rotation. While *angle* is a positive value, it indicates the image rotates counterclockwise (coordinate origin is the upper-left corner).

The third parameter *scale* is the scaling factor.

The function calculates the matrix below :

$$\begin{bmatrix} \alpha\beta(1-\alpha) \cdot center.x - \beta \cdot center.y \\ -\beta\alpha\beta \cdot center.x + (1-\alpha) \cdot center.y \end{bmatrix}$$

Where

$$\alpha = scale \cdot \cos angle,$$

$$\beta = scale \cdot \sin angle$$

3. Feature point detection algorithm

C. SUSAN corner detector

SUSAN is an acronym standing for the smallest univalue segment assimilating nucleus. SUSAN places a moving circular mask on the image. We compare the gray level of each pixel inside the circular mask with the centroid of the mask. If the difference of gray level between one pixel and the centroid pixel is less than a certain value, we consider that the pixel and the centroid has the same (similar) gray level and area consists of pixels satisfying this condition is called Univalue Segment Assimilating Nucleus (USAN).

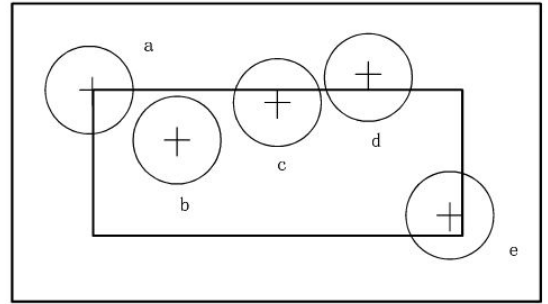


Figure 4 Relationship transformation models

As Figure 4 shows, when the circular mask is completely in the background or goals, USAN area gets its maximum; when it moves toward the edge of the target object, USAN area becomes smaller; when the centroid is at the edge of the target, USAN area becomes very small; when the centroid of the mask is at the corner point, USAN area gets its minimum.

Therefore, by calculating the USAN value of every pixel and comparing it with the set threshold, if the USAN value is less than the threshold value of the pixel, then the point will be considered as an edge point.

D. Harris corner detector

Harris and Stephens improved upon Moravec's corner detector by considering the differential of the corner score with respect to direction directly, instead of using shifted patches.

The corner is the most obvious and important feature of an image. For the first derivative, the angle change in each direction is the largest and the edge region of a significant change in one direction only.

The basic mathematical formula is as follows:

$$E(u, v) = \sum_{x,y} w(x, y) [I(x+u, y+v) - I(x, y)]^2$$

Where $W(x, y)$ represents a moving window, $I(x, y)$ represents the intensity of the pixel gray value range of 0 to 255. According to Taylor series, calculate an order to N order partial derivatives, eventually getting a Harris matrix formula:

$$A = \sum_u \sum_v w(u, v) \begin{bmatrix} I_x^2 & I_x I_y \\ I_x I_y & I_y^2 \end{bmatrix}$$

Calculate the eigenvalues according to the Harris matrix, and then calculate the Harris angle response value:

$$R = \det M - k(\text{trace } M)^2$$

$$\det M = \lambda_1 \lambda_2$$

$$\text{trace } M = \lambda_1 + \lambda_2$$

Where K is the coefficient value, typically in the range that between 0.04 to 0.06.

E. SURF feature detector

SURF, Speeded Up Robust Features in other words, is first proposed by the Bay in 2006. SURF is the speeded up version of scale invariant feature transform algorithm (SIFT algorithm). In general, the standard SURF operator is much faster than SIFT operator and has better stability while using multiple pictures. The most important feature of SURF is the use of harr features and integral image, which greatly speeds up the running time of the program. SURF can be applied to object recognition computer vision and 3D remodeling.

4. Analysis of OpenCV source code of SURF

In OpenCV, SURF algorithm is often related to SURF, SurfFeatureDetector, SurfDescriptorExtractor classes.

We can find such two definitions:

```
typedef SURF SurfFeatureDetector;
typedef SURF SurfDescriptorExtractor;
```

That is to say, SurfFeatureDetector class and SurfDescriptorExtractor class which we often use, in fact, are both SURF class, the three are equivalence. The further relationship is shown in Figure 5.

What's more, we use the function drawKeypoints for drawing key points:

```
void drawKeypoints(const Mat&image, const vector<
```

```
KeyPoint>& keypoints, Mat& outImage, const Scalar
& color=Scalar::all(-1), int flags=DrawMatchesFlag
s::DEFAULT)
```

The first argument *src* is the input image.

The second parameter *keypoints* is the feature point we get from the source image; it is an output parameter.

The third parameter *outImage* is the output image.

The fourth parameter *color* is the color of key point and it as the default value `Scalar :: all (-1)`.

The fifth parameter *flags* is feature identifier in drawing key points and it has the default value `DrawMatchesFlags :: DEFAULT`.

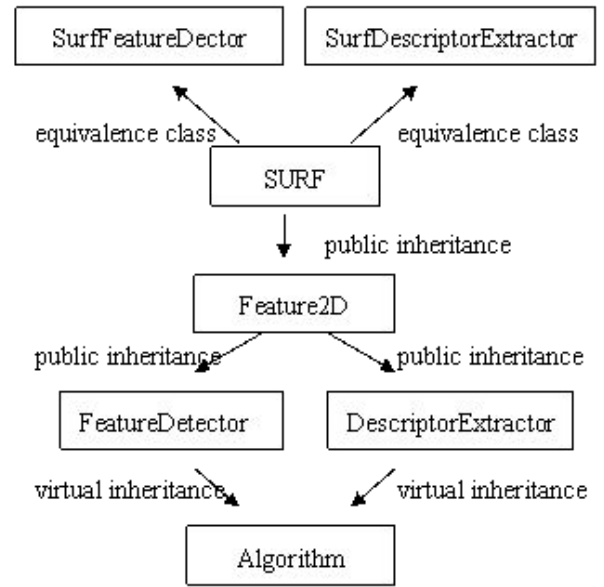


Figure 5 Relationship between the classes about SURF

5. Image mosaicking

While we need to stitch two frames, we first choose one corner point X in the reference frame. In the target frame, around the same coordinates as point X, select a rectangular area of certain size as the search area. Then we compare X with all the feature points of the next frame in the search area.

Here X' is one of the feature points of the target frame in the search area. We use the following equation to compare the two points:

$$\rho(x, x') = \frac{\sum_{i=-n}^n \sum_{j=-m}^m [I(u+i, v+j) - E(x)][I'(u'+i, v'+j) - E(x')]}{\sqrt{\sum_{i=-n}^n \sum_{j=-m}^m [I(u+i, v+j) - E(x)]^2 [I'(u'+i, v'+j) - E(x')]^2}}$$

$$E(x) = \frac{\sum_{i=-n}^n \sum_{j=-m}^m [I(u+i, v+j)]}{(2n+1)(2m+1)}$$

Where $I(u, v)$ is a gray level of (u, v) in the frame $E(x)$ is the average value of gray histogram at point $X(u, v)$ and $\rho(x, x')$ is the coefficient correlation of x and x' .

The value range of coefficient correlation is usually $[-1, 1]$. If the coefficient correlation is larger than the given threshold, the two feature points would be detected as corresponding points.

After detecting the corresponding points, we can use their coordinates to get the transformation matrix. There are four methods for image mosaicking: frame to frame mosaicking, frame to mosaic mosaicking, mosaic to frame mosaicking and mosaic to mosaic mosaicking.

5.1. Frame to frame mosaicking

In this method, we first work out the transformation matrix of adjacent frames and then use these matrixes to work out the transformation matrixes of any two frames. If we choose one frame as the reference image, other frames in the image sequence can be transformed into the reference coordinate system using the previously calculated transformation matrixes and then we can do the image mosaicking.

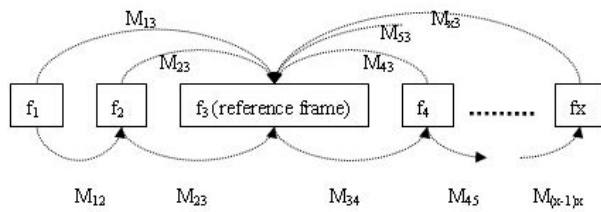


Figure 6 Mosaicking transformation

5.2. Frame to mosaic mosaicking

In this method, every image frame for stitching is transformed into the same coordinate as the mosaicking image and the current frame image is used to update the mosaicking image. If there is larger displacement between the current frame and mosaicking image, we can use the transformation parameter of the previous frame and mosaicking image as the initial value to do the registration.

5.3. Mosaic to frame mosaicking

Mosaic to frame mosaicking uses the current frame as the reference coordinate. Use the mosaicking image to register with the current frame and the same transformation matrix as the one from the previous frame to the current frame. Because in this method we always use the current frame as the reference frame, the coordinate is always changing. As the current frame that is being transferred does not need any

transformation, this image mosaicking method is particularly suitable for Real-time video transmission and stitching of dynamic image.

5.4. Mosaic to mosaic mosaicking

Mosaic to mosaic mosaicking is a tree-structured image mosaicking method. In this method, image sequence is segmented according to certain rules, stitched and we get sub images and at last get the final image.

6. Example



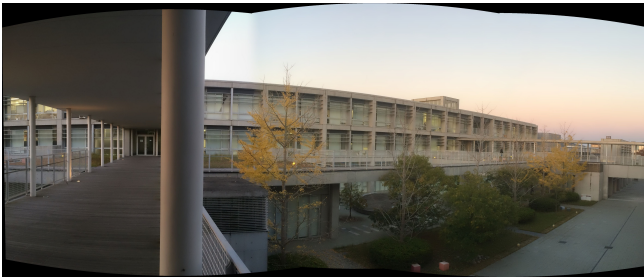
(a) Image 1



(b) Image 2



(c) Image 3



(d) Panoramic image created from image 1 to 3.
Figure 7 Example

Figure 7 (d) is a panoramic image stitching by Figure 7 (a), (b) and (c)

CONCLUSIONS

This paper introduces some methods of feature point detection and panoramic image mosaic using OpenCV. To do the panoramic image mosaic, after image acquisition, preprocessing is necessary so that the following registration and fusion are easier and more accurate. As the camera is moving, transformation becomes necessary in stitching and to get the transformation matrix of two images, corner points are detected first, and then do the comparing so that some pairs of corresponding points can be found and finally we use their coordinates to get the transformation matrix. After establishing a good background model, by comparing the current image with the background we can get the foreground and track the moving object.

References

[1] Donggyu Sim, Yongmin Kim, "Detection and

compression of moving objects based on new panoramic image modeling," *Image and Vision Computing*, Vol.27, pp. 1527–1539, 2009

- [2] Dae-Hyun Kim, Yong-In Yoon, Jong-Soo Choi, "An efficient method to build panoramic image mosaics," *Pattern Recognition Letters*, Vol.24, pp. 2421–2429, 2003
- [3] Yasuyuki SUGAYA, Kenichi KANATANI, "Extracting Moving Objects from a Moving Camera Video Sequence," *Memoirs of the Faculty of Engineering*, Vol. 39, pp. 56–62, 2005
- [4] Rao, N.I., Huijun Di, Guangyou Xu, "Panoramic Background Model under Free Moving Camera," *Fuzzy Systems and Knowledge Discovery*, Vol.1, pp. 639 - 643, 2007
- [5] Duke Gledhilla, Gui Yun Tiana, Dave Taylora, David Clarke, "Panoramic imaging—a review," *Computers & Graphics*, Vol.27, pp.435-445, 2003
- [6] Seong Jong Ha, Sang Hwa Lee, Nam Ik Cho, Soo Kyun Kim, Byungjun Son, "Embedded Panoramic Mosaic System Using Auto-Shot Interface," *IEEE Transactions on Consumer Electronics*, Vol. 54, No. 1, 2008
- [7] Drew Steedly, Chris Pal, Richard Szeliski, "Fast construction of panoramic images for cystoscopic exploration," *Computerized Medical Imaging and Graphics*, Vol.34, No.7, pp. 579–592, 2010
- [8] Steedly, D. Pal, Chris Szeliski, R, "Efficiently registering video into panoramic mosaics," *name of journal*, Vol.2, pp.1300-1307, 2005
- [9] Yongzhen Huang, Zifeng Wu, Liang Wang, Tieniu Tan, "Feature Coding in Image Classification: A Comprehensive Study," *Pattern Analysis and Machine Intelligence*, VOL. 36, NO. 3, MARCH 2014, Vol.36, No.3, pp. 493 - 506, 2014
- [10] Mittal A., Paragios, N, "Motion-Based Background Subtraction using Adaptive Kernel Density Estimation," *Computer Vision and Pattern Recognition*, Vol.2, pp.302-309, 2004
- [11] Elgammal, A., Duraiswami, R., Harwood, D., Davis, L.S., "Background and Foreground Modeling Using Nonparametric Kernel Density Estimation for Visual Surveillance," *Proceedings of the IEEE* Vol.90, No.7, pp. 1151 - 1163, 2002
- [12] Cucchiara, R., Grana, C., Piccardi, M., Prati, A., "Detecting Moving Objects, Ghosts, and Shadows in Video Streams," *Pattern Analysis and Machine Intelligence*, Vol.25, No.10, pp. 1337 - 1342, 2003

Contact address

Graduate School of Information, Production and System
Waseda University
2-7 Hibikino, Wakamatsu Kitakyushu,
Fukuoka 808-0135 JAPAN
E-mail:wangjiaxi@akane.waseda.jp

画像内挿における微分ヒストグラムによる性能評価法

Performance evaluation of image interpolation using on differential histogram

○松田 充夫¹, 但馬 文昭²

○Nobuo Matsuda¹, Fumiaki Tajima²

¹大島商船高等専門学校

²横浜国立大学

Abstract: An interpolation of image by using fuzzy inference is one of techniques suitable for keeping the edge features of the image. The mean square error (MSE) has been widely used in evaluating the performance of the image interpolation. However the mean square error is inadequate as performance evaluating for level of preservation of edge features in the image. In this paper, a measure using a differential histogram for evaluating the preservation level is proposed.

1. はじめに

画像内挿では sinc 関数の補間原理による最近隣補間法, 共 1 次補間法, 3 次畳み込み補間法などの手法^[1]がよく使われる。これらの手法はエッジまで滑らかに拡大するという欠点があるため, 補間値前後の画素状態に着目したエッジ情報を保存する Carrato らの補間法^[2]やファジィ推論を用いた画像拡大法^[3,4]が提案されている。

一般的に画像の内挿法の性能は平均二乗誤差 (MSE) の指標で評価される。しかしエッジの保存性を MSE の評価だけでは不十分に評価できない。そのため一部の画像の断面の画素値でエッジ情報の保存性を評価することが多いが, この評価方法では画像全体のエッジ情報の評価にはならない。

本論文は微分ヒストグラムを用いて内挿画像のエッジ特徴の保存性を評価する方法を提案する。

2. 画像の内挿法の性能評価

画像内挿の性能評価の手順を図 1 に示す。フィルタ処理にはガウシアンフィルタを使う。内挿処理として 6 種類 (共 1 次, 共 3 次, 3 次畳み込み, Carrato 法, 2 種類のファジィ推論による方法) を採用する。性能指標の計算は, 原画像と出力画像から平均二乗誤差 (MSE) と提案する微分ヒストグラムの情報エントロピーによる性能評価指数 (DHE と表記) とする。

2. 1 平均二乗誤差 (MSE)

平均二乗誤差 (MSE) は式 (1) で定義される。

$$MSE = \frac{1}{M \times M'} \sum_{i=1}^M \sum_{j=1}^{M'} [y(i, j) - s(i, j)]^2 \quad (1)$$

ここで $y(i, j)$ は出力画素値, $s(i, j)$ は原画素値, M と M' は画像サイズである。

2. 2 微分ヒストグラムによる性能評価指標

対象画像の微分ヒストグラムによる性能指標の算

出は図 2 の手順で行う。まず対象画像に微分フィルタをかける。微分フィルタには Prewitt 又は Sobel フィルタを用いる。得られた微分値の勾配から微分ヒストグラムを作成する。この微分ヒストグラムの情報エントロピーを算出して性能指標 (DHE) とする。

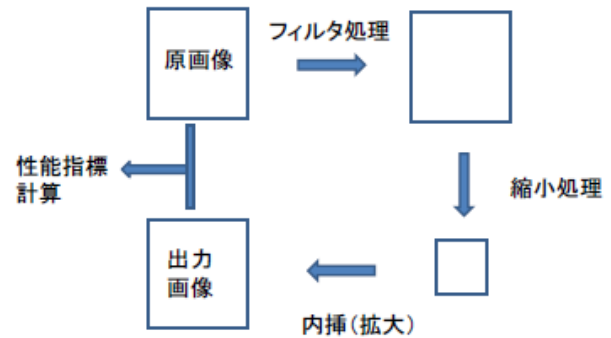


Figure 1 Process of performance evaluation on an interpolation of image.

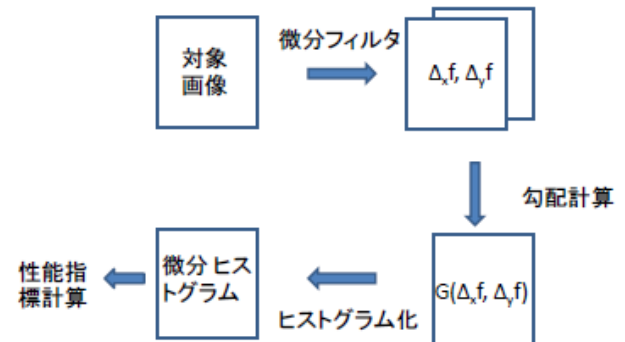


Figure 2 Evaluation of entropy from a differential histogram of target image.

3. 実験結果

図 3 に示すような 256×256 画素, 256 階調グレースケールの 13 種類の標準画 (Airplane, Barbara, Boat, Building, Cameraman, Girl, Text, LAX, Lena, Light house, Woman (紙面の関係で一部省略)) を用いて, 微分ヒストグラムによる画像内挿の性能評価の有効性の実験を行った。図 4 は Lena 画像の拡大画像 (上段) とその画像ヒストグラム (中段), 微分ヒストグラム (下

段) である。

表 1 は、主要な内挿画像についての、5 種類 (共 1 次, 共 3 次, 3 次畳み込み, Carrato 法, ファジィ推論による方法) の内挿による出力画像の MSE である。MSE の平均値は 13 種類の画像の平均値である。MSE の平均値からみると共 3 次内挿法が最も誤差が少ない。これに続きファジィ推論を用いた内挿法の誤差が少ない。



Figure 3 Tested Images; from top left to bottom right, airplane, Barbara, boat, girl, bridge, lena.

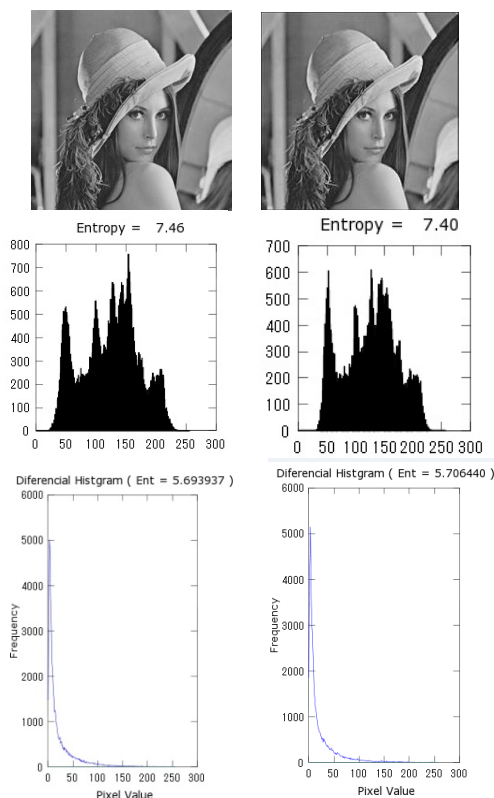


Figure 4 Top left: an original image; top right: the image after interpolation by using Fuzzy inference; middle left: histogram of the original image; middle right: histogram after interpolation; bottom left: a differential histogram of the original image; bottom right: the differential histogram after interpolation.

表 2 は主要な画像の微分ヒストグラムの情報エントロピーによる性能評価指数 (DHE) の結果である。この指標が原画像の値に最も近いのは 3 次畳み込みであるが、MSE の値は悪い。2 つの表を総合してみるとファジィ推論による内挿法は最も誤差が少なくかつエッジ特性を良好に保存する

方法であることがわかる。

以上の結果から、内挿画像の平均誤差は MSE で、一方画像のエッジ特徴の保存性は DHE から画像内挿法の性能を評価すれば、採用する画像内挿法が適切であるかどうか確認することができる。

4. まとめ

画像の内挿において、内挿画像の平均の誤差を MSE で、一方画像のエッジ特徴の保存は微分ヒストグラムの情報エントロピーを用いた評価指標 (DHE) で、画像内挿法の性能を評価することができることを 13 種類の標準画像で検証した。

Table 1 MSE of interpolation images

	Bi-linear	Bi-cube	Cubic	Fuzzy 1	Carrato
Airplane	215.2	101.2	242.9	123.4	137.6
Barbara	357.0	264.2	338.2	245.6	247.3
Boat	151.9	66.7	152.9	70.7	76.0
Girl	71.2	32.1	58.3	37.4	41.8
Bridge	413.0	310.3	317.1	378.5	327.0
Lena	172.4	63.7	73.3	130.9	80.9
Mandole	247.4	154.2	227.8	167.1	178.5
Average*	264.7	163.5	230.3	170.7	181.6

Table 2 DFE of interpolation images(Prewitt)

	Original	Bi-linear	Bi-cube	Cubic	Fuzzy 1	Carrato
Airplane	5.680	5.200	5.520	5.528	5.526	5.417
Barbara	6.700	5.794	6.027	6.139	6.111	5.859
Boat	5.865	5.124	5.480	5.592	5.555	5.499
Girl	5.475	5.020	5.167	5.244	5.310	5.334
Bridge	6.699	6.124	6.145	6.161	6.211	5.975
Lena	5.694	5.313	5.619	5.631	5.706	5.449
Mandole	6.671	5.666	6.172	6.219	6.222	6.027
Average*	6.049	5.486	5.763	5.884	5.798	5.596

*Average is values for all of 13 images.

参考文献

- [1]高木幹雄, 下田陽久 監修: 新編 画像解析ハンドブック, 東京大学出版, pp.1363-1373, 2004
- [2]S.Carrato,G.Ramponi,S.Marsi: "A Simple Edge-Sensitive Image Interpolation Filter", International Conf. on Image Processing, IEEE, pp.711-714, 1996
- [3]木村誠聡, 田口亮, 村田裕: "ファジィ推論を用いたデジタル画像の拡大に関する一手法", 第 13 回デジタル信号処理シンポジウム講演論文集, pp.255-260, 1998.
- [4]麻生隆史, 末竹規哲, 山川烈: "SOM に基づく学習型ファジィ推論とその画像拡大への応用", 第 18 回ファジィシステムシンポジウム講演論文集, pp.325-328, 2002.

連絡先

〒742-2193 山口県大島郡周防大島町小松 1091-1
大島商船高等専門学校 電子機械工学科
Tel/Fax 0820-74-5581
E-mail: :matsuda@oshima-k.ac.jp

色相に注目したイチョウの黄葉度の算出法

The calculation method of the yellow degree of the ginkgo using a hue expression

長家瑛
Akira Nagaie
法政大学
Hosei Univ.

松田修三
Shuzo Matsuda
法政大学
Hosei Univ.

小沢和浩
Kazuhiro Ozawa
法政大学
Hosei Univ.

但馬文昭
Fumiaki Tajima
横浜国立大学
Yokohama Nat'l. Univ.

Abstract: The authors proposed a simple observation method to grasp the state of microclimate of suburban area, using the plant indicator. Yellow leaves of ginkgo that is one of the season's indicators, serve as environmental barometer. We observed the yellow leaves of ginkgo which was planted in the suburbs of Tokyo. In this observation, we represented the degree of yellow leaves using the sense of human. In this study, from an image obtained by photographing the ginkgo leaves, we extract the RGB components, and then it was converted to the hue. In particular, we focus on the Hue component of the color. Then we quantitatively extract the Ginkgo's leaves from the image of Ginkgo tree. Finally, we can calculate the degree of yellow from the Hue component of leaves of Ginkgo. It is also possible to process in a short time a large amount of data by applying this technique to a moving image.

Keywords: plant indicator, ginkgo, yellow leaves, Hue, image processing

1. はじめに

イチョウの葉の黄葉の主な要因は秋になってからの気温の低下であることは一般に知られていることである。しかし気温の低下のしかたによって黄葉の進み方にも差が出ることが従来の研究でわかってきている[1]。また、比較的狭い地域でも人工物、すなわち熱源となる建造物付近のイチョウとそうでない場所にあるイチョウとでは黄葉の仕方に差が現れる[2]。つまり、イチョウの黄葉の様子を調べることはその地域の気温環境を知るための有効な方法であると考えられる。

我々は、毎年イチョウの黄葉が始まる9月下旬から12月上旬まで観測を行い、葉の緑から黄への変化の程度を黄葉度として5段階で評価し、黄葉の進み方を調べてきた。観測は動画像や静止画像として記録するとともに目測により黄葉度を判定する方法を用いてきた[3][4]。目測による方法の利点は、人の視覚の補正効果（知識や経験に基づく）により、イチョウの周辺や背後にある樹木などと区別し、逆光などによって暗く見える場合でも黄葉度の判定が可能であることである。しかし、観測期間や観測区域が長くなればなるほどデータ量も飛躍的に増加し、黄葉度の定量化処理に時間がかかるという問題点もある。

本研究は、イチョウの葉の緑から黄への色の変化がHSV表色系における色相（Hue）の連続した範囲に対応していることに注目し、撮影した静止画像または動画像から特定範囲の色情報を色相の値としてもとめ、さらに黄葉度として算出するための方法および一連の処理の自動化を試みたものである。

2. イチョウの黄葉の色相による表現

イチョウの樹全体の静止画を撮影し、その画像からイチョウのできるだけ多くの葉の部分抽出し、色すなわちRGB値の平均を求める。イチョウの葉が気温の低下とともに緑から黄色へ変化する様子は、HSV表色系で考えれば図1に示すように、色相の値が連続的に変化することになる。このときのイチョウの葉の色の色相の変化はおよそ30°から120°の間の値になると考えられる。色相の値は緑の葉の場合およそ120°、黄色の葉の場合はおよそ30°となる。すなわち黄葉が進むにつれて色相はおよそ120°から30°付近へ推移することになる。

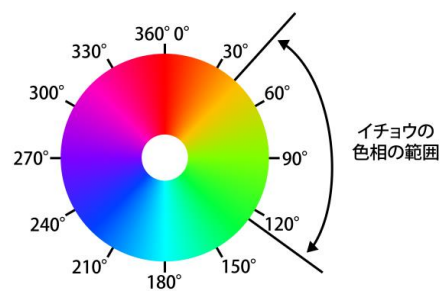


図1 色相環

2.1 色相表現によるイチョウの葉の抽出法

この特徴を利用すると、様々な背景とともに撮影されたイチョウの樹の画像からイチョウの葉以外の部分を取り除くことができる。撮影した画像内のそれぞれのピクセルのRGBの値から色相を求め、およそ30°から120°の範囲内に収まるピクセルのみを抽出することが可能となる。さらに、彩度にも注目した。イチョウの葉は彩度が50%以上であることがほとんどであり、こ

の彩度をしきい値として、イチョウ以外の不要な部分を取り除く。抜き出されたイチョウの葉の平均色から色相を求める。

2.2 RGB 表色系と HSV 表色系

RGB 表色系から HSV 表色系への変換には式(1), (2), (3), (4)を用いる。R, G, B のうち、最大のものを max 、最小のものを min 、色相の値を H とする。

- RGB の値のうち、 R が最大のとき ($max=R$)

$$H = 60 \times (G - B) / (max - min) + 360 \quad (1)$$

- RGB の値のうち、 G が最大のとき ($max=G$)

$$H = 60 \times (B - R) / (max - min) + 120 \quad (2)$$

- RGB の値のうち、 B が最大のとき ($max=B$)

$$H = 60 \times (R - G) / (max - min) + 240 \quad (3)$$

HSV 表色系の彩度(S)とは、色の鮮やかさの尺度である。彩度は 0% から 100% で表される。イチョウの葉の彩度はおよそ 50% を上回ることがわかった。幹などの色相の値がイチョウの葉に近い部分でも彩度の違いにより識別が可能となる。彩度は式(4)を用いて求める。 max と min は R, G, B のうち最大値と最小値を示す。

$$S = (max - min) / max \times 100 \quad (4)$$

2.3 色相と黄葉度の関係

従来、黄葉度は黄葉の進み具合を目測により観測し、0 から 4 の 5 段階で判断した[3][4]。一例を図 2 に示す。



図 2 黄葉の進み方の一例と黄葉度
(昭和の森公園(昭島市)2014/9 から 2014/12)

はじめに、イチョウの黄葉が進むにつれて、色相の値がどの範囲で変化するかを調べる。そのために、未黄葉と完全黄葉を含むいろいろな黄葉度のイチョウの葉を 32 枚、スキャナーより画像として取り込み、葉ごとの色相の平均値を

求めた。図 3 はその一部である。求めたイチョウの葉の色相(h)と黄葉度の関係を表 1 に示す。



図 3 採取した代表的なイチョウの葉
(2014 年 9 月から 12 月昭和の森公園で採取)

表 1 色相と黄葉度の関係

色相 $h/^\circ$	黄葉度
$h \leq 44^\circ$	4
$44^\circ < h \leq 53^\circ$	3
$53^\circ < h \leq 62^\circ$	2
$62^\circ < h \leq 71^\circ$	1
$71^\circ < h$	0

イチョウの葉の色相を求めると、その値は 35° から 80° の範囲内に分布した。これを 5 段階に等分し、それぞれを黄葉度 4, 3, 2, 1, 0 とした。さらに 44° から 71° の間の黄葉度の変化を線形と仮定し、色相(h)と黄葉度の関係を求めた。その結果を図 4 に示す。

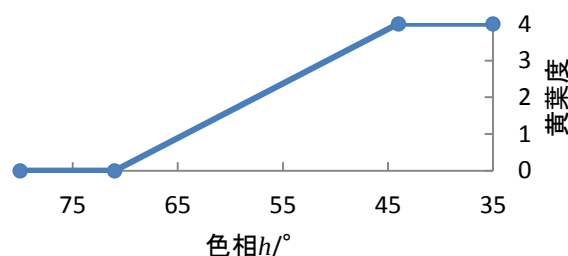


図 4 色相の値と黄葉度の関係

また、図 4 より色相(h)の 71° から 44° の範囲の黄葉度(y)は式(5)のように示すことができる。

$$y = -0.15h + 10.5 \quad (5)$$

この式(5)を用いて色相から黄葉度の値を求める。

3. 画像処理と色相の値から黄葉度への変換

画像からイチョウの葉の部分抽出する処理を行う前に、画像内にあるイチョウの木以外の部分(植え込みなど色相値が 35° から 80° の範囲にあると考えられる他の植物など)を図 5 のように大まかに取り除く。図 5 に示す黒い部分がこれにあたる。



植込みを取り除く前（左）と取り除いた後（右）

図 5 画像の前処理

処理に用いるデータの形式は JPEG 画像である。つぎに、色相の範囲と彩度の値を指定し、イチョウの葉以外の部分のピクセルを黒く塗りつぶす。この処理を全てのピクセルで実行し、残ったピクセルの平均の色の色相を求め、式(5)より黄葉度を算出する。最後に、葉を抜き出した画像と平均色の画像を出力する。ここでは、抜き出す色相の範囲を 30° から 100° 、彩度を 30%以上とした。図 6, 7, 8 はそれぞれ、画像処理を実行する前の画像、葉を抽出した画像、その抽出されたピクセルの平均色である。

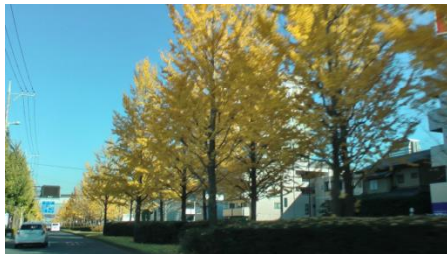


図 6 処理前の画像

図 7 を見ると、図 6 中のイチョウの葉が抽出されたことが確認できる。また、図 8 よりこの画像内のイチョウの葉の色相の平均は式(1),(2),(3)のいずれかより 52.2° となった。



図 7 イチョウの葉を抽出した画像



図 8 抽出したピクセルの平均色(色相 52.2°)

最後に、式(5)を用いて色相から黄葉度への変換を行うと、黄葉度は 2.7 となった。

4. 黄葉度の算出と従来方法との比較

画像処理により算出された黄葉度と従来の目測による黄葉度の比較を行う。用いた画像は昭和の森公園(東京都昭島市)で撮影したものである。イチョウを定点で未黄葉の状態から完全に黄葉するまで(2014 年 9 月から 2014 年 12 月)をおよそ 1 週間おきに撮影した。黄葉度は黄葉の進み具合によって図 2 のように 0 から 4 の 5 段階に定められる[3][4]。ここでは、2 つのイチョウの木(図 9)の黄葉度を求め、比較した。この 2 つのイチョウは歩道を挟んで向い合っており、西側をイチョウ 01、東側をイチョウ 02 とする。



イチョウ 01 イチョウ 02
図 9 昭和の森公園 (2014 年 9/27)

画像処理の後、式(5)を用いて求めた結果と目測で求めた黄葉度とを比較したものを図 10 (イチョウ 01)、図 11 (イチョウ 02) にそれぞれ示す。

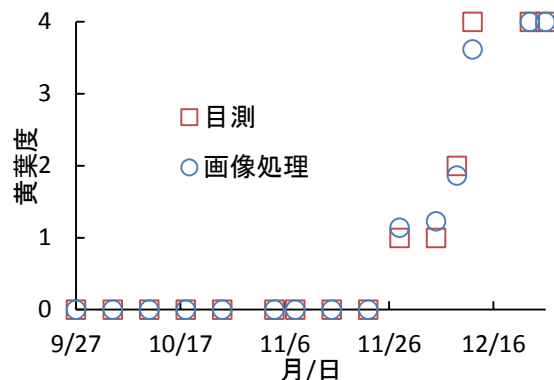


図 10 目測の値を加えた昭島のイチョウ 01 の黄葉度の進み具合(2014 年)

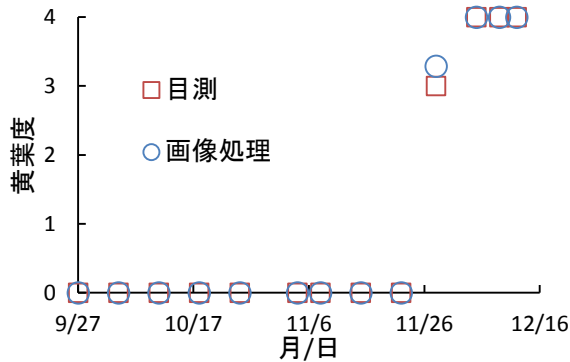


図 11 目測の値を加えた昭島のイチョウ 02 の黄葉度の進み具合(2014 年)

黄葉度を色相に注目して画像処理を用いて求めたものと、目測による黄葉度を比較した結果、図 10, 11 のように多少の差異は認められるが、一致している。わずかな差は画像処理の分解能が目測よりも高いためであると考えられる。

5. 動画画像への応用

我々はこれまで、小気候区の気温環境を把握するため、東京都八王子市および多摩市を東西に走る都道 158 号線（通称多摩ニュータウン通り）のイチョウに注目し観測を行ってきた[4]。これまでは、目測でイチョウの木の黄葉度を求めた。交差点と交差点の間を 1 区間(小気候区)として区間の黄葉度を求めるときにも、イチョウの木一本ずつの黄葉度を目測により求め、その平均値を区間の黄葉度としていた。本稿で提案した手法を用いれば、車にカメラを搭載し、道路に植栽されている街路樹としてのイチョウを撮影しながら走行することで、小気候区の動画データからの黄葉度の算出が可能である。ここでは予備実験として、多摩ニュータウン通りの交差点「多摩中央署入り口」から交差点「多摩山王橋」の間のおよそ 300m の区間で撮影した動画に 3. と同じ方法で処理した結果の動画の中の 1 コマを図 12 に示す。黒い背景の中に全体的に浮き上がっているのが抽出されたイチョウの葉である。左下隅にあるものが元の撮影された動画、その隣が抽出されたイチョウの葉の平均色、さらにその隣がその色相の値、一番右の値が黄葉度である。これらの色相の値と黄葉度の値は動画像として表示される。

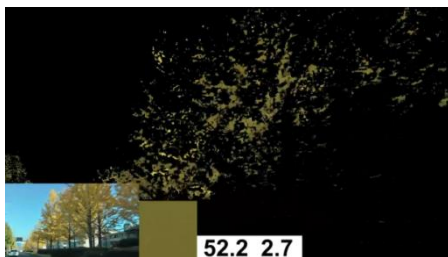


図 12 動画を用いた黄葉度算出

6. おわりに

本研究では、イチョウの黄葉の進み具合（黄葉度）が色相の連続した変化と関係があることに注目して画像処理を行い、その後黄葉度の算出を行った。その結果、処理に用いた画像データの範囲では従来の目測による方法と比べてよく一致することが確認できた。さらに多くの画像データについて検証する必要があるが、ここで提案する色相に注目した画像処理による黄葉度の算出方法は、とくに画像データが多い場合においては効率的で有効な方法であると考えられる。また目測による 5 段階評価に比べ、黄葉度の分解能が高くなり、黄葉度の滑らかな変化を調べることができるためイチョウ周辺の環境変化を敏感に捉えられるのではないかと考えられる。また今後の展望として、カメラからリアルタイムで得られる情報から黄葉度を瞬時に判別するシステムの開発なども可能である。

参考文献

- [1]丹下, 池ヶ谷, 松田, 小沢:HSV 表色系によるイチョウの黄葉の表現と分析, 24 回日本知能情報ファジィ学会・ソフトサイエンス・ワークショップ, 講演論文集 pp.73 から 76, 2014 年 3 月
- [2]松田, 小沢, 但馬, 宮武:植物指標としてのイチョウの葉の色相変化と気温の関係
- [3]松原, 松田, 小沢, 但馬, 宮武:植物指標から見た都市小気候区の様子
- [4]松田, 宮武, 角田, 小沢:イチョウの黄葉度を利用した都市微気候の可視化(Ⅱ), 法政大学多摩研究報告第 25 巻 pp.13-20, 2010 年 5 月

連絡先

長家瑛:法政大学理工学部電気電子工学科

E-mail:akira.nagaie.4c@stu.hosei.ac.jp

〒184-8584 小金井市梶野町 3-7-2

TEL.042-387-6230

松田修三:法政大学理工学部電気電子工学科

E-mail:matuda@hosei.ac.jp

〒184-8584 小金井市梶野町 3-7-2

TEL.042-387-6230

小沢和浩:法政大学経済学部

E-mail:kozawa@hosei.ac.jp

〒194-0298 町田市相原町 4342

TEL.042-783-2580

トランプカードの画像認知・想起時脳波による推定 —トリックなしのカード当てマジック—

Discrimination of EEGs on Recognizing and Recalling of Playing Cards -Playing Cards Magic without a Trick-

○山ノ井高洋¹, 豊島 恒², 大西真一¹, 山崎敏正³, 菅野道夫⁴
○Takahiro Yamanoi¹, Hisashi Toyoshima², Shin-ichi Ohnishi¹, Toshimasa Yamazaki³,
and Michio Sugeno⁴

¹北海学園大学

¹Hokkai-Gakuen University

²ジャパンテクニカルソフトウェア

²Japan Technical Software

³九州工業大学

³Kyushu Institute Technology

⁴ソフトコンピューティングヨーロッパ研

⁴European Centre for Soft Computing

Abstract: The authors measured electroencephalograms (EEGs) from subjects on recognizing and on recalling 13 types of playing card images (from Ace to King) presented on a CRT. Each presented image card is the club. The canonical discriminant analysis was applied to these single trial EEGs. Four channels of EEGs at the right frontal and temporal were used in the discrimination. They were Fp2, F4, C4 and F8 according to the international 10-20 system. Sampling data were taken from 400ms to 900ms at 25ms intervals from each single trial EEG. The number of variates is twenty one by four channels; so the data are eighty four dimensional vectors. The number of external criteria is thirteen, from Ace to King, and a number of explanation variates is thus eighty four. Results of the canonical discriminant analysis by use of so called jack knife method were more than 90 % for three subjects. We could perform a playing cards magic without a trick.

1. はじめに

ヒト脳機能に関する様々な研究によると、視覚刺激に関する処理は後頭葉で行われ、高次処理の機能は左脳と右脳とで分化されていると言われている[1][2]。著者らは先行研究[3][4][5]において言語および空間認知に関する脳活動部位の推定を行い、詳細な時空間的脳活動のモデルを得た。さらに逆の向きに直線移動する視覚刺激に対する脳活動の比較の結果、ERPのピークで極性が逆であること(図1)、また高次の脳機能処理が中前頭回で集中することを確認した[4]。

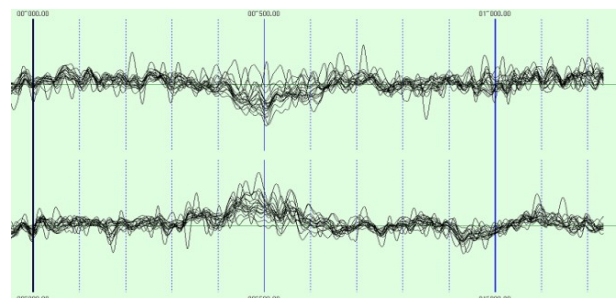


図1 Comparison between ERPs in imaging Rightward (Upper) and Leftward (Lower) in case of Kanji.

さらに、著者らは先行研究[6][7]において、被験者が向きを示す単語(漢字：上、下、左、右)および記号(矢印：↑, ↓, ←, →)を観察および黙読する際の脳波(electroencephalograms: EEG)を計測し、等価電流双極子推定(equivalent current dipole source localization: ECDL) 法 (SynaCenterPro: NEC Corporation)[8]を試み、脳活動の時空間的な比較を行った。この結果、黙読時に潜時 400ms 前後で漢字と矢印とで共通した脳活動が中前頭回で観察された。

著者らが着目した中前頭回は、ワーキングメモリの「中央実行系」と飛ばれる機能に関係していることが、病巣研究および機能画像(fMRI)などから明らかになっている。ワーキングメモリの中央実行系とは、外界の情報の取捨選択、一時的な保持、後続する他の機能への指令、その結果の評価、先に一時的に保存していた情報の消去の決定と指令等、いわば、高次の脳機能の司令部の役割をしているシステムと想定されている。

上述のように、これらの潜時で逆の向きを示す視覚刺激に対して得られた事象関連電位(event-related potential: ERP)を比較し、ERP のピーク極性が反転していることを確認した(図 1)。この結果をもとに、著者らは、ピーク極性の反転を用いることで、被験者がイメージした矢印の種類を判別できると考え、ブレイン・コンピュータ・インタフェース(brain computer interface: BCI)への応用を試みた[9]。その結果、全被験者で 80%程度の判別率を得た。しかしながら、この学習データには判別すべきデータ自身が含まれていた。

本研究では、EEG の判別率向上を目指し、潜時 400ms 以降の区間で、トランプのクラブ A から K までの判別に有効なチャンネルを探り、上前頭回に対応する 2 (Fp2), 4 (F4), 6 (C4), 12 (F8)の合計 4ch からの脳波信号を利用した。各カードはランダムに提示した。サンプリング潜時区間と間隔は先行研究[10]と同様とし、400ms から 900ms までの EEG を 25ms 間隔でサンプリングし、84 次元ベクトル空間の多変量データとした。このデータに対し、多変量統計的データ解析の一手法である正準判別分析を試みた。この結果、3 名の被験者に対して提示・想起時の EEG に対して 90%を越える判別率を得た。

2. 判別用 EEG データの前処理

著者らは従来 EEG 計測において、国際 10-20 法に基づいた 19ch で計測を行い、得られた EEG に等価電流双極子推定法を適用し、脳内処理の時空間的推移を明らかにしてきている。しかしながら、BCI の目的には、19ch のデータ全てを判別に使用することは効率的でない。著者らの一連の BCI 先行研究においては、ERP のピーク極性反転と中前頭回活動の関連性が示されたことに着目し、この付近の 4ch のデータを利用している。本研究で新たに提示画像として用いたトランプ画像の一例を図 2 に示す。

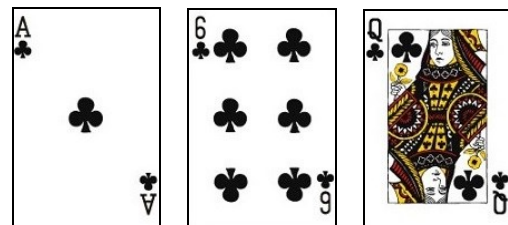


図 2 Presented example images of playing card

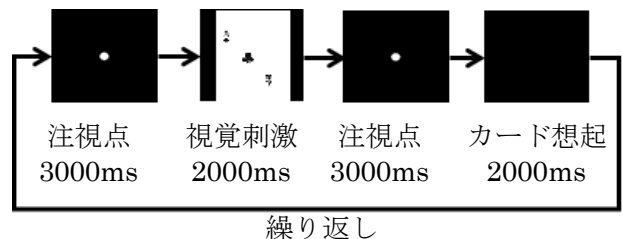


図 3 Flow of experiment

上述のチャンネル選択に従い、シングルトライアル EEGs の 400ms から 900ms の間を 25ms 間隔でサンプリングし、各チャンネルから 21 個の要素を抽出した。さらに、これらを 1 トライアルごとに横 1 列のデータとし、84 次元のベクトルデータとした。EEG 本来は時系列データであるが、今回の解析ではチャンネルごとに 21 点のデータが順不同で対応しているとみなした。したがって、1 つのシングルトライアル EEGs が 84 次元データに対応する。これをデータ 1 とする。

これまでに著者らは、いわゆるジャックナイフ(クロスバリデーション) 法を適用するために、不偏性

を保証する目的で、データを増やすことを考えてきた。データ 1 のサンプリング時点を-1ms 前方にずらして 25ms 間隔でサンプリングを行った。さらに-1ms 前方にずらし、同様に 25ms 間隔でサンプリングを行った。これらをデータ 1 に加え、データ数をデータ 1 単独の 3 倍にサンプリングされたデータをデータ 2 と呼ぶ。

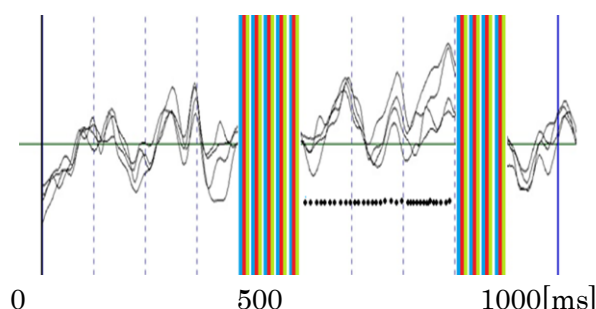


図 4 Data sampling from single-trial data

3. 前処理後 EEG データへの正準判別分析

使用したデータは、正常な視覚を有する 22 歳の被験者女性 YI, YT, YY に対して、トランプカードのクラブのエースからキングまでの 13 枚を被験者に提示し、図 3 の手順で提示時と想起時の EEG を計測したものである。なお、3 名の利き手は右であった。それぞれの EEG データを 84 個の説明変量を持つデータとして、13 群に対する正準判別分析を試みた。外的基準がトランプの画像 13 種類、説明変量が 4ch の脳波電極から計測された EEG データからサンプリングされた 84 個のデータである。

この際に、それぞれのデータに対してジャックナイフ統計を行った。判別すべき 1 トライアルのデータを除き残りの EEG を学習データとして判別係数を求めて、得られた判別係数を用いて除いた 1 トライアルのデータの判別を試みた。この結果、いずれの被験者に関しても、データ 1 の判別結果の方が 3 倍のデータ数のデータ 2 の結果よりも提示・想起時共に良好で、判別率は 90% を越えた（表 1, 表 2, 表 3, 表 4）。データ 2 に関しては、いずれの被験者に対しても提示時に関しての判別率が約 88%, 想起時に関しても約 88% となる結果を得た。データ

表 1 Result of the discriminant analysis for playing card recognition (Discriminant ratio 92.31 % Subject YI)

Obs./Pred.	A	2	3	4	5	6	7
A	4	0	0	0	2	0	0
2	0	5	0	0	0	0	0
3	0	0	6	1	0	0	0
4	0	0	0	8	0	0	0
5	0	0	0	0	8	0	0
6	0	0	0	0	0	8	0
7	0	0	0	0	0	0	8
8	0	0	0	0	0	0	0
9	0	0	0	0	0	0	0
10	0	0	0	0	0	0	0
J	0	0	0	0	0	0	0
Q	0	0	0	0	0	0	0
K	0	0	0	0	0	0	0
Obs./Pred.	8	9	10	J	Q	K	
A	0	0	0	1	0	0	
2	0	0	0	0	1	2	
3	0	0	0	1	0	0	
4	0	0	0	0	0	0	
5	0	0	0	0	0	0	
6	0	0	0	0	0	0	
7	0	0	0	0	0	0	
8	8	0	0	0	0	0	
9	0	8	0	0	0	0	
10	0	0	8	0	0	0	
J	0	0	0	8	0	0	
Q	0	0	0	0	8	0	
K	0	0	0	0	0	8	

Discriminant Ratio : 96/104 = 92.31

表 2 Example of result of the discriminant analysis for playing card recognition (Discriminant ratio 92.31% Subject YT in case of Data1)

Obs./Pred.	A	2	3	4	5	6	7
A	5	0	0	0	1	0	0
2	0	5	0	0	0	0	0
3	0	2	6	0	0	0	0
4	0	0	0	8	0	0	0
5	0	0	0	0	8	0	0
6	0	0	0	0	0	8	0
7	0	0	0	0	0	0	8
8	0	0	0	0	0	0	0
9	0	0	0	0	0	0	0
10	0	0	0	0	0	0	0
J	0	0	0	0	0	0	0
Q	0	0	0	0	0	0	0
K	0	0	0	0	0	0	0
Obs./Pred.	8	9	10	J	Q	K	
A	0	0	1	0	1	0	
2	0	0	2	1	0	0	
3	0	0	0	0	0	0	
4	0	0	0	0	0	0	
5	0	0	0	0	0	0	
6	0	0	0	0	0	0	
7	0	0	0	0	0	0	
8	8	0	0	0	0	0	
9	0	8	0	0	0	0	
10	0	0	8	0	0	0	
J	0	0	0	8	0	0	
Q	0	0	0	0	8	0	
K	0	0	0	0	0	8	

Discriminant Ratio: 96/104 =92.31

表 3 Example of result of the discriminant analysis for playing card recalling (Discriminant ratio 92.23 % Subject YI)

Obs./Pred.	A	2	3	4	5	6	7
A	4	0	0	0	2	0	0
2	0	5	0	0	0	0	0
3	0	0	6	1	0	0	0
4	0	0	0	8	0	0	0
5	0	0	0	0	8	0	0
6	0	0	0	0	0	8	0
7	0	0	0	0	0	0	8
8	0	0	0	0	0	0	0
9	0	0	0	0	0	0	0
10	0	0	0	0	0	0	0
J	0	0	0	0	0	0	0
Q	0	0	0	0	0	0	0
K	0	0	0	0	0	0	0
Obs./Pred.	8	9	10	J	Q	K	
A	0	0	0	1	0	0	
2	0	0	0	0	1	2	
3	0	0	0	1	0	0	
4	0	0	0	0	0	0	
5	0	0	0	0	0	0	
6	0	0	0	0	0	0	
7	0	0	0	0	0	0	
8	8	0	0	0	0	0	
9	0	8	0	0	0	0	
10	0	0	8	0	0	0	
J	0	0	0	8	0	0	
Q	0	0	0	0	8	0	
K	0	0	0	0	0	8	

Discriminant Ratio: 95/103 = 92.23

表 4 Example of result of the discriminant analysis for playing card recalling (Discriminant ratio 93.20 % Subject YT in case of Data1)

Obs./Pred.	A	2	3	4	5	6	7
A	5	3	0	0	0	0	0
2	0	5	0	1	0	0	0
3	0	1	6	0	0	0	0
4	0	0	0	8	0	0	0
5	0	0	0	0	8	0	0
6	0	0	0	0	0	8	0
7	0	0	0	0	0	0	8
8	0	0	0	0	0	0	0
9	0	0	0	0	0	0	0
10	0	0	0	0	0	0	0
J	0	0	0	0	0	0	0
Q	0	0	0	0	0	0	0
K	0	0	0	0	0	0	0
Obs./Pred.	8	9	10	J	Q	K	
A	0	0	0	0	0	0	
2	0	2	0	0	0	0	
3	0	0	0	0	0	0	
4	0	0	0	0	0	0	
5	0	0	0	0	0	0	
6	0	0	0	0	0	0	
7	0	0	0	0	0	0	
8	8	0	0	0	0	0	
9	0	8	0	0	0	0	
10	0	0	8	0	0	0	
J	0	0	0	8	0	0	
Q	0	0	0	0	8	0	
K	0	0	0	0	0	8	

Discriminant Ratio: 96/103 = 93.20

4. おわりに

本研究で提示に使用したデータは、図 2 に例示した様に、黒のトランプカードのクラブ 13 種類である。これらのトランプカードを被験者 3 名に提示し、その際の認知時および想起時の被験者の EEG を計測した。計測後の EEG の 2(Fp2), 4(F4), 6(C4), 12(F8)チャンネルに着目し、従来矢印(上, 下, 左, 右)黙読時の EEG に対して、BCI の応用で試みてきた正準判別分析法を適用した。

潜時 400ms から 900ms までの EEG を 25ms 間隔でサンプリングし、4 チャンネル分をベクトルの要素として並べ、84 要素からなるベクトル空間の多変量データとした。

正準判別分析法にジャックナイフ法を適用すると、一般に判別率が悪化することから、著者らは今までに、これに対処するため学習データの数の増加を試みた。計測したシングルトライアル EEG データをさらに 399ms から 899ms までを 25ms 間隔で同様にサンプリング、さらに 388ms から 898ms まで同様にサンプリングし、データ数を 3 倍として判別率を高めてきた。

本研究でも、トランプ認知・想起時の EEG 計測データに対し、従来と同様に多変量統計的データ解析の一手法である正準判別分析を試みたところ 13 種類の想起時 EEG の判別率はすべて約 90%であった。しかしながら、3 名の被験者ともに、元のデータであるサンプリングデータ 1 そのもので判別すると判別率が 90%を越え、3 倍にすると判別率がやや低下した。この結果は、我々の先行研究の結果とは逆の結果であった。

この結果、事前に各々 10 回程度の学習実験を行うことが必要であるが、トリックを使って行っていたカード当てマジックと同じことを、対象者に対してカード認知時に脳波計測し、これを判別することに

より 90%を越える確率で当てることができた。さらに、認知時だけではなく、「何のカードを見ましたか」と質問した時のカードを想起した脳波からでも同程度の確率で当てることが可能となった。

謝辞

本研究の一部は、平成 25 年度で終了した文部科学省私立大学戦略的研究基盤形成支援事業に伴う「電磁・光センシングを主体とする生体関連情報の先進的計測・処理技術の開発と応用」において購入済みの実験機器を用いて行われた。

脳波計測実験には、平成 26 年度山ノ井研究室電子情報工学科卒業研究生井上祐貴さん、泊侑華さん、安田由佳さんの 3 名の協力を得た。ここに謝意を表す。

参考文献

- [1] R. A. McCarthy, E. K. Warrington (相馬芳明, 本田仁視 監訳), 認知神経心理学, 医学書院, 1996
- [2] 中川敦子, 単漢字の処理における左右半球機能の検討 ―漢字の構造と音訓を通して―, 神経心理学, 第 18 巻, 第 2 号, pp.120-128, 2002
- [3] T. Yamanoi, T. Yamazaki, J. L. Vercher, E. Sanchez, M. Sugeno, Dominance of recognition of words presented on right or left eye -Comparison of Kanji and Hiragana-, Modern Information Processing, From Theory to Applications, B. Bouchon-Meunier, G. Coletti and R.R. Yager Eds., Elsevier Science B.V., pp.407-416, 2006
- [4] Takahiro YAMANOI, Hisashi TOYOSHIMA, Toshimasa YAMAZAKI, Shin-ichi OHNISHI: Localization of brain activity during perception of circle movement by use of equivalent current dipole analysis, 2004 IEEE international Conference on Fuzzy Systems Proceedings VOLUME 1, pp.321-324, 2004
- [5] 豊島恒, 山ノ井高洋, 山崎敏正, 大西真一: 時間等価電流双極子推定法による RDS 両眼立体視の脳内処理部位の解析, 知能と情報(日本知能情報ファジィ学会誌), Vol.18, No.1, pp.102-110, 2006
- [6] 豊島恒, 山ノ井高洋, 山崎敏正, 大西真一, 菅

野道夫: 向きを表す単語と記号に対する時空間的脳活動の比較, 知能と情報(日本知能情報ファジィ学会誌), Vol.18, No.3, pp.425-433, 2006

- [7] 山ノ井高洋, 森高篤司, 高柳浩, 山崎敏正, 菅野道夫, 野中秀俊: “向きを表す記号イメージング時 EEG によるマイクロロボット制御 II”, 第 25 回ファジィシステムシンポジウム, 1B2-01, pp.1-4, 2009
- [8] 山崎敏正: 32 チャネル電極キャップによる脳内等価電流双極子推定, CLINICAL NEUROSCIENCE, 18-2, pp.186-190, 2000
- [9] T. Yamanoi, H. Toyoshima, T. Yamazaki, S. Ohnishi, M. Sugeno, and E. Sanchez: Micro Robot Control by Use of Electroencephalograms from Right Frontal Area, Journal of Advanced Computational Intelligence and Intelligent Information, vol. 13, No. 2, pp. 68-75, 2009
- [10] T. Yamanoi, H. Toyoshima, T. Yamazaki, S. Ohnishi, M. Sugeno and E. Sanchez: “Brain Computer Interface by use Electroencephalograms from Right Frontal Area”, The 6th International Conference on Soft Computing and Intelligent Systems, and The 13th International Symposium on Advanced Intelligent Systems, pp.1150-1153, 2012.

連絡先

山ノ井 高洋

北海学園大学工学部生命工学科

〒064-0926 札幌市中央区南 26 条西 11 丁目 1-1

TEL: 011-841-1161 (内線 7874)

E-mail: yamanoi@hgu.jp

GA-rough sets approach to fuzzy time-series based forecasting model

Zhao Jing¹, Junzo Watada²

¹²Graduate School of Information, Production and Systems,
Waseda University.

Abstract: Fuzzy time-series (FTS) has been applied to handle non-linear problems, such as enrollment, weather and stock index forecasting. In the forecasting processes, fuzzy logical relation (FLR) plays a pivotal role in forecasting accurately. Usually FTS uses an equal interval to obtain forecasted values. But in this paper, we use Genetic algorithms (GA) to optimize the interval at first. Based on this, then rough sets (RS) method is used to recalculate the values. In the empirical analysis, Japan stock index is used as experimental data sets and one fuzzy time-series method, as a comparison model. The experimental results showed that the proposed method is more efficient than the FTS method.

Keywords: Fuzzy time series; Genetic algorithm; Rough set; Forecasting; Stock Index

I. Introduction

FTS analysis consists of the fuzzification, the set of fuzzy relationships, and defuzzification. Some studies only focus on the fuzzification or setting the fuzzy relationships, since they are not relevant in forecasting. Therefore, in this paper, a setting of fuzzy relationship will be proposed. On the thesis Darwin's evolution theory, GA proposed by Holland will be used in the model. Besides, RS method will be used to analyze FTS data [1][2]. The RS theory is proposed by Z. Pawlak to acquire rule-based knowledge from a set defined with plural attributes [3]. The main objective of this paper is to forecast the stock closing price when analysing the trend of the stock. We could identify trends of similar patterns from stock data to predict development of new data in the future.

II. Method review

This part briefly reviews fuzzy time-series, genetic algorithms and rough set method.

A) Fuzzy time series

The main difference between the fuzzy time series and conventional time series is that the values of the former are fuzzy numbers while the values of the latter are real numbers [4]. Let U be the universe of discourse, where $U = \{u_1, u_2, \dots, u_b\}$. A fuzzy set A_i of U is defined as $A_i = f_{A_i}(u_1/u_1) + f_{A_i}(u_2/u_2) + \dots + f_{A_i}(u_b/u_b)$, where f_{A_i} is the membership function of the fuzzy set A_i , $f_{A_i}(u_a) \in [0, 1]$. u_a is an element of fuzzy set A_i , $f_{A_i}(u_a)$ is the degree of belongingness of u_a to A_i , $f_{A_i}(u_a) \in [0, 1]$ and $1 \leq a \leq b$. $Y(t) = (t = \dots, 0, 1, 2, \dots)$ is a subset of a real number. Let $Y(t)$ be the universe of discourse defined by the fuzzy set $f_i(t)$. If $F(t)$

consists of $f_i(t) (i=1, 2, \dots)$, $F(t)$ is defined as a fuzzy time series on $Y(t) = (t = \dots, 0, 1, 2, \dots)$.

If there exists a fuzzy relationship $R(t-1, t)$, such that $F(t) = F(t-1) \times R(t-1, t)$, where $\times$ represents an operation, then $F(t)$ is said to be caused by $F(t-1)$. Let $F(t-1) = A_i$ and $F(t) = A_j$. The relationship between two consecutive observation $F(t)$ and $F(t-1)$, referred to as a fuzzy logical relationship (FLR), can be denoted by $A_i \rightarrow A_j$, where A_i is called the left-hand side (LHS) and A_j the right-hand side (RHS) of the FLR [5].

B) Genetic algorithms

In genetic algorithms, a population consists of chromosomes, and a chromosome consists of genes [6]. There are three basic operations of genetic algorithms, including crossover, mutation and reproduction [7-9].

- 1) Crossover: The system randomly selects two chromosomes from a population and randomly generates a real number between zero and one comparing it with a user defined parameter called the crossover rate [CR] between zero and one to decide whether the two selected chromosomes perform the crossover operations.
- 2) Mutation: The system randomly selects a chromosome from a population and uses a predefined parameter called the mutation rate [MR] between zero and one to decide the mutation probability of the selected chromosome. The system randomly generates a real number between zero and one and compares it with the value of MR to decide

whether the selected chromosome performs the mutation operation.

- 3) Reproduction: A genetic algorithm performs the reproduction operation according to the fitness value of each chromosome. The chromosomes with higher fitness values will be reproduced, and the chromosomes with lower fitness values will be abandoned. Through the reproduction process, the selected chromosomes will become members of the next generation. If the number of chromosomes selected by the system is smaller than the number of chromosomes of a population, then the system randomly generates some chromosomes and lets the selected chromosomes and the generated chromosomes form the new population of the next generation [10].

Procedure of genetic algorithms:

Step 1: Randomly choose two chromosomes from the population for performing crossover operations and randomly choose a crossover point. After performing the crossover operations, if the values of the genes of a chromosome are not in an ascending sequence, then the system sorts the values of the genes in the chromosomes in an ascending sequence.

Step 2: Randomly select a chromosome from the population and randomly select a gene from the selected chromosome to perform the mutation operation. If the random number generated by the system is smaller than or equal to the mutation rate, the mutation operation mutates the selected gene of the selected chromosome.

Step 3: Calculate the MSE value of each derived chromosome obtained from Step 2 based on the algorithm presented previously.

Step 4: Select 20 chromosomes from the population that have lower MSE values, and put 10 chromosomes randomly generated by the system into them to form the new population of the next generation. In the same way, after repeatedly performing Step 1 to Step 4 to evolve 300 generations, the chromosome that has the lowest value of MSE is the best solution to be used to forecast the stock index of Japan [10].

C) Rough sets

A rough set is especially useful for domains where the data collected are imprecise and/or incomplete about the domain objects. It provides a powerful tool for a data analysis and data mining of imprecise and ambiguous data. A reduction is the minimal set of attributes that preserves the

indispensability relation, that is, the classification power of the original dataset [11]. The rough set theory is a very useful method for decision support systems, especially when hybrid data, vague concepts and uncertain data were involved in the decision process [12]. In the biggest advantage of using rough set method dealing with uncertain problem is that it doesn't need to advance or additional information about the data. The attribute reduction and rule extraction are an important application in rough set theory in data mining.

Let $I = (U, A)$ be an information system (attribute-value system), where U is a non-empty set of finite objects (the universe). A is a non-empty, finite set of attributes. So that $\alpha: u \rightarrow V_\alpha$ for every $u \in U$. V_α is the set of values that attribute α may take.

The rough sets theory denotes such rough representation as approximation. This is a knowledge acquisition method. There are two kinds of approximations: one is an upper approximation to take an element of a rough set into consideration from possibility points of view and the other is a lower approximation to take an element of a rough set from viewpoints of necessity. The visual illustration of upper and lower approximations is shown in Fig.1 [3].

s1	s2	s3	s4	s4
s6	s7	s8	s9	s10
s11	s12	s13	s14	s15
s16	s17	s18	s19	s20
s21	s22	s23	s24	s25

Fig. 1: Upper and lower approximations

It is named “reduction” It is possible to decide a decision table, if features of a set can be divided into two subsets of condition features and decision features, respectively. The decision table can be understood as decision rules that correspond a value of conditional feature to a value of decision feature. For instance, a decision table shown in Table I illustrates a decision rule for sample x1 [3]. If a=1 and b=1 and c=1 then d=1.

TABLE I: Actual closing price(Nikkei Index Data)

Sample	Conditional Feature			Decision Feature
	a	b	c	d
x1	1	1	1	1
x2	2	1	2	2
x3	2	1	1	1
x4	1	1	2	1

This decision table has 5 samples features and 3 conditional. It is possible to derive 5 rules with 3 conditions. But the decision rules have redundancy in the conditional portion. By employing a reduction method in the rough set theory, it is possible to derive the minimal rules required for expressing the same decision rules.

In the case of a decision table shown in Table II, rules illustrate the decision feature $d=1$ as follows: If $a=1$ then $d=1$, If $b=1$ and $c=1$ then $d=1$. It is possible to derive decision rules for decision class=2 in the same way [3].

TABLE II: Decision matrix

$A^*(d=1)/d=2$	x2	x4
x1	a1,c1	a1,b1
x3	c1	b1
x5	a1	a1,b1,c2

III. Proposed method for forecasting stock index with FTS based on GA and rough set

In FTS forecasting method, FLR is one of the critical factors that influence the forecasting accuracy. For improving the forecasting processes, and many advanced algorithms such as Neural Networks, PSO are applied to mine FLRs from time-series. In this paper, we provide another approach, which uses GA and RS methods to deal with this issue. The flow chat is shown as Fig.2.

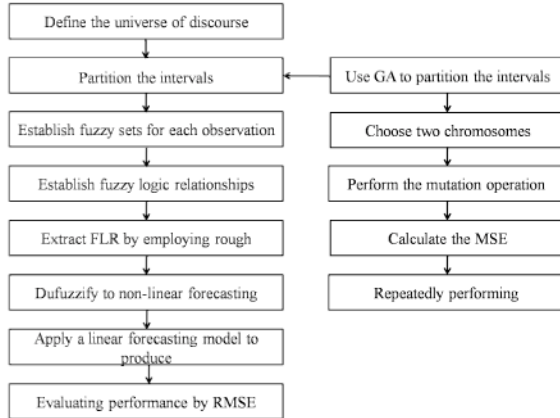


Fig. 2: Proposed method flow chart

IV. Proposed algorithms

In this section, we present a new method for forecasting stock index using FTS GA and RS. From Table III, we can see that the minimum D_{min} and the maximum D_{max} of the Japan stock index are 13885.11 and 17869.74. We let the universe of discourse $U=[13880,17880]$, which $D_{min} \in U$,

$D_{max} \in U$. The universe of discourse U can be divided into intervals $u_1, u_2, u_3, u_4, u_5, \dots, u_n$ where $u_1=[13880, x_1], u_2=[x_1, x_2], u_3=[x_2, x_3], \dots, u_n=[x_{(n-1)}, 17880]$, and $x_1 < x_2 < x_3 < \dots < x_{(n-1)}$. Each chromosome consists of $n-1$ genes are showed as Table IV. The algorithm (called Algorithm-GA) for forecasting and calculating the mean square error (MSE) based on the above chromosome is presented as follows.

TABLE III: Actual closing price (Nikkei Index Data)

Date	Closing price
2014/01/06	15864.44
2014/01/07	15784.25
2014/01/08	15906.57
...	...
2014/11/04	16720.99
2014/11/05	16778.28
2014/11/06	16725.45
2014/11/07	16824.08
2014/11/10	16713.37
2014/11/11	16855.30
2014/11/12	17170.00
2014/11/13	17099.95
...	...
2014/12/10	17308.82
2014/12/11	17043.46
2014/12/12	17298.19
2014/12/15	17037.21
2014/12/16	16714.83
2014/12/17	16672.94
2014/12/18	17136.41
2014/12/19	17471.60

TABLE IV: Chromosome consists of $n-1$ genes

$\bar{x}_1$	$\bar{x}_2$	$\bar{x}_3$	...	$\bar{x}_{n-1}$
-------------	-------------	-------------	-----	-----------------

TABLE V: Optimized chromosome

13880	13942	14016	...	17482	17613	17482
-------	-------	-------	-----	-------	-------	-------

Step1: Construct the membership functions corresponding to the genes $x_1, x_2, x_3, \dots, x_{n-1}$ of a chromosome. For example, assume that there is a chromosome shown as Table V.

Step2: Partition the universe of discourse into fifty intervals. In this paper, the universe of discourse $U=[13880,17880]$. And the equal intervals $u_1, u_2, u_3, u_4, u_5, \dots, u_n$,

Where $u_1 = [13800, 13960], u_2 = [13960, 14040], u_3 = [14040, 14120], \dots, u_{26} = [15880, 15960], u_{27} = [15960, 16040], \dots, u_{39} = [16920, 17000], u_{40} = [17080, 17160], \dots, u_{49} = [17720, 17800], u_{50} = [17800, 17880]$.

Step3: Define the universe of discourse and the intervals for the observations. Fuzzify the observations shown as equation (1).

$$\begin{aligned}
A_1 &= 1/u_1 + 0.5/u_2 + 0/u_3 + 0/u_4 + 0/u_5 + \dots + 0/u_{50} \\
A_2 &= 0.5/u_1 + 1/u_2 + 0.5/u_3 + 0/u_4 + 0/u_5 + \dots + 0/u_{50} \\
A_3 &= 0/u_1 + 0.5/u_2 + 1/u_3 + 0.5/u_4 + 0/u_5 + \dots + 0/u_{50} \\
&\dots \\
A_{49} &= 0/u_1 + 0/u_2 + 0/u_3 + 0/u_4 + 0/u_5 + \dots + 0/u_{47} \\
&\quad + 0/u_{48} + 0/u_{49} \\
A_{50} &= 0/u_1 + 0/u_2 + 0/u_3 + 0/u_4 + 0/u_5 + \dots + 0/u_{48} \\
&\quad + 0/u_{49} + 0/u_{50}
\end{aligned} \tag{1}$$

Step4: Obtain the fuzzy logical relationships. In the FTS method, different Japan stock index values can be fuzzified into corresponding fuzzy sets.

For example, the 2014/11/06 was 15864.44, which is fuzzified to A_{25} , 2014/11/12 was 17170.00, which is fuzzified to A_{42} . The fuzzification process is depicted in Table VI. Some fuzzified Japan stock index for the year 2014 are listed in Table 2. But when we use GA to optimize the intervals, the fuzzy sets have been changed obviously. For example, the 2014/11/06 was 15864.44, which is fuzzified to A_{36} , 2014/11/12 was 17170.00, which is fuzzified to A_{46} . The fuzzification process is depicted in Table VII.

TABLE VI: Fuzzified closing price calculated by FTS (Nikkei Index Data)

Date	Closing price	Fuzzified price c
2014/01/06	15864.44	A_{25}
2014/01/07	15784.25	A_{24}
2014/01/08	15906.57	A_{26}
...	...	...
2014/11/04	16720.99	A_{36}
2014/11/05	16778.28	A_{37}
2014/11/06	16725.45	A_{36}
2014/11/07	16824.08	A_{37}
2014/11/10	16713.37	A_{36}
2014/11/11	16855.30	A_{38}
2014/11/12	17170.00	A_{42}
2014/11/13	17099.95	A_{41}
...	...	...
2014/12/10	17308.82	A_{43}
2014/12/11	17043.46	A_{40}
2014/12/12	17298.19	A_{43}
2014/12/15	17037.21	A_{43}
2014/12/16	16714.83	A_{36}
2014/12/17	16672.94	A_{35}
2014/12/18	17136.41	A_{41}
2014/12/19	17471.60	A_{45}

Step5: Divide the derived fuzzy logical relationships (FLRs) based on the current states of the closing price of fuzzy logical relationships. In the FTS method, the fuzzy logical relationships of the closing price can be obtained from Table VI which is shown in Table VIII. But the fuzzy logical relationships of the closing price used by GA have been changed, too. We can be obtained from Table VII which is shown in Table IX.

TABLE IX: Fuzzy logical relationships calculated by GA

$A_{38} \rightarrow A_{35}$	$A_{35} \rightarrow A_{37}$	...	$A_{31} \rightarrow A_{30}$	$A_{30} \rightarrow A_{28}$	$A_{28} \rightarrow A_{24}$
$A_{24} \rightarrow A_{23}$	$A_{23} \rightarrow A_{17}$	$A_{17} \rightarrow A_{21}$	$A_{21} \rightarrow A_{24}$	...	$A_{47} \rightarrow A_{46}$
$A_{48} \rightarrow A_{47}$	$A_{47} \rightarrow A_{46}$	$A_{46} \rightarrow A_{44}$	$A_{44} \rightarrow A_{43}$	$A_{43} \rightarrow A_{46}$	$A_{48} \rightarrow A_{47}$

TABLE VIII: Fuzzy logical relationships calculated by FTS

$A_{25} \rightarrow A_{24}$	$A_{24} \rightarrow A_{26}$	...	$A_{21} \rightarrow A_{20}$	$A_{20} \rightarrow A_{18}$	$A_{18} \rightarrow A_{16}$
$A_{18} \rightarrow A_{15}$	$A_{15} \rightarrow A_{11}$	$A_{11} \rightarrow A_{14}$	$A_{14} \rightarrow A_{18}$	...	$A_{43} \rightarrow A_{40}$
$A_{40} \rightarrow A_{43}$	$A_{43} \rightarrow A_{43}$	$A_{43} \rightarrow A_{38}$	$A_{38} \rightarrow A_{35}$	$A_{35} \rightarrow A_{41}$	$A_{41} \rightarrow A_{45}$

TABLE VII: Fuzzified closing price calculated by GA (Nikkei Index Data)

Date	Closing price	Fuzzified price
2014/01/06	15864.44	A_{36}
2014/01/07	15784.25	A_{35}
2014/01/08	15906.57	A_{37}
...	...	...
2014/11/04	16720.99	A_{44}
2014/11/05	16778.28	A_{44}
2014/11/06	16725.45	A_{44}
2014/11/07	16824.08	A_{44}
2014/11/10	16713.37	A_{44}
2014/11/11	16855.30	A_{45}
2014/11/12	17170.00	A_{46}
2014/11/13	17099.95	A_{46}
...	...	...
2014/12/10	17308.82	A_{47}
2014/12/11	17043.46	A_{46}
2014/12/12	17298.19	A_{47}
2014/12/15	17037.21	A_{46}
2014/12/16	16714.83	A_{44}
2014/12/17	16672.94	A_{43}
2014/12/18	17136.41	A_{46}
2014/12/19	17471.60	A_{47}

Step6: Divide the derived FRLs into groups based on the current states of the closing price of fuzzy logical relationships. In FTS method, we can be obtained from Table X, while we used GA, we got Table XI.

TABLE X: Fuzzy logical relationships groups calculated by FTS

Group1	$A_1 \rightarrow A_1$	$A_1 \rightarrow A_2$					
Group2	$A_2 \rightarrow A_2$	$A_2 \rightarrow A_3$	$A_2 \rightarrow A_4$				
Group3	$A_3 \rightarrow A_3$	$A_3 \rightarrow A_4$					
Group4	$A_4 \rightarrow A_2$	$A_4 \rightarrow A_8$	$A_4 \rightarrow A_8$				
Group5	$A_5 \rightarrow A_1$	$A_5 \rightarrow A_5$	$A_5 \rightarrow A_8$	$A_5 \rightarrow A_7$	$A_5 \rightarrow A_8$	$A_5 \rightarrow A_9$	
...							
Group44	$A_{44} \rightarrow A_{42}$						
Group45	$A_{45} \rightarrow A_{45}$	$A_{45} \rightarrow A_{47}$	$A_{45} \rightarrow A_{50}$				
Group46	-						
Group47	$A_{47} \rightarrow A_{50}$						
Group48	$A_{48} \rightarrow A_{50}$						
Group49	$A_{49} \rightarrow A_{44}$	$A_{49} \rightarrow A_{50}$					
Group50	$A_{50} \rightarrow A_{43}$	$A_{50} \rightarrow A_{49}$	$A_{50} \rightarrow A_{50}$				

TABLE XI: Fuzzy logical relationships groups calculated by GA

Group1	$A_1 \rightarrow A_1$	$A_1 \rightarrow A_2$			
Group2	$A_2 \rightarrow A_2$	$A_2 \rightarrow A_3$	$A_2 \rightarrow A_4$	$A_2 \rightarrow A_5$	
Group3	$A_3 \rightarrow A_2$	$A_3 \rightarrow A_4$			
Group4	$A_4 \rightarrow A_3$	$A_4 \rightarrow A_8$			
Group5	$A_5 \rightarrow A_1$	$A_5 \rightarrow A_3$	$A_5 \rightarrow A_8$	$A_5 \rightarrow A_9$	$A_5 \rightarrow A_{11}$
...					
Group44	$A_{44} \rightarrow A_{43}$	$A_{44} \rightarrow A_{44}$	$A_{44} \rightarrow A_{45}$		
Group45	$A_{45} \rightarrow A_{48}$				
Group46	$A_{46} \rightarrow A_{44}$	$A_{46} \rightarrow A_{46}$	$A_{46} \rightarrow A_{47}$		
Group47	$A_{47} \rightarrow A_{45}$	$A_{47} \rightarrow A_{46}$	$A_{47} \rightarrow A_{47}$	$A_{47} \rightarrow A_{48}$	
Group48	$A_{48} \rightarrow A_{49}$				
Group49	$A_{49} \rightarrow A_{50}$				
Group50	$A_{50} \rightarrow A_{47}$	$A_{50} \rightarrow A_{50}$			

Step7: Establish fuzzy relationship groups based on RS theory. Assign trend weights. The FLRs are grouped into the trend to which they belong. For example, $A_1 \rightarrow A_2$ will be grouped into the “up trend,” $A_1 \rightarrow A_1$ into the “no-change trend,” and $A_2 \rightarrow A_1$ into the “down trend” as Table XII listed. These weights should be standardized to obtain the weight matrix $w(t)=[w_1, w_2, \dots, w_j]$. The standardized weight matrix equation is defined as equation (2). The rule extraction based on rough set is listed in Table XIII which employed by FTS, and Table XIV employed by GA.

$$w(t) = [w_1, w_2, \dots, w_j]$$

$$= \left[\frac{w_1}{\sum_{k=1}^j w_k}, \frac{w_2}{\sum_{k=1}^j w_k}, \dots, \frac{w_k}{\sum_{k=1}^j w_k} \right] \quad (2)$$

TABLE XII: Assign weights to different trends

FLR	Trend	Weight
(=1)A ₂₅ → A ₂₄	Down trend	Assign weight1
(=2)A ₂₄ → A ₂₆	Up trend	Assign weight1
(=3)A ₂₆ → A ₂₅	Down trend	Assign weight1
(=4)A ₂₅ → A ₂₄	Down trend	Assign weight1
(=5)A ₂₄ → A ₁₉	Down trend	Assign weight1
(=6)A ₁₉ → A ₂₂	Up trend	Assign weight1
...	...	...

TABLE XIII: Rough set rule induction Japan stock index of year 2014 based on FTS

Values	Recurrence	Decision rules
A ₁	2	A ₁ → {A ₁ [1], A ₂ [1]}
A ₂	5	A ₂ → {A ₂ [2], A ₃ [1], A ₄ [3]}
A ₃	2	A ₃ → {A ₃ [1], A ₄ [1]}
A ₄	6	A ₄ → {A ₂ [1], A ₆ [4], A ₈ [1]}
A ₅	7	A ₅ → {A ₅ [2], A ₆ [2], A ₇ [1], A ₈ [1], A ₉ [1]}
...	...	...
A ₄₅	2	A ₄₅ → {A ₄₇ [1], A ₄₈ [1]}
A ₄₆	-	-
A ₄₇	1	A ₄₇ → {A ₅₀ [1]}
A ₄₈	1	A ₄₈ → {A ₅₀ [1]}
A ₄₉	2	A ₄₉ → {A ₄₄ [1], A ₅₀ [1]}
A ₅₀	5	A ₅₀ → {A ₄₃ [1], A ₄₉ [1], A ₅₀ [3]}

TABLE XIV: Rough set rule induction Japan stock index of year 2014 based on GA

Values	Recurrence	Decision rules
A ₁	2	A ₁ → {A ₁ [1], A ₂ [1]}
A ₂	5	A ₂ → {A ₂ [1], A ₃ [1], A ₄ [2], A ₅ [1]}
A ₃	3	A ₃ → {A ₂ [1], A ₄ [1]}
A ₄	4	A ₄ → {A ₃ [1], A ₆ [3]}
A ₅	6	A ₅ → {A ₁ [1], A ₃ [1], A ₆ [3], A ₉ [1], A ₁₁ [1]}
...	...	...
A ₄₅	2	A ₄₅ → {A ₄₆ [2]}
A ₄₆	7	A ₄₆ → {A ₄₄ [1], A ₄₆ [1], A ₄₇ [1]}
A ₄₇	10	A ₄₇ → {A ₄₅ [1], A ₄₆ [5], A ₄₇ [3], A ₄₈ [1]}
A ₄₈	1	A ₄₈ → {A ₄₉ [1]}
A ₄₉	1	A ₄₉ → {A ₅₀ [1]}
A ₅₀	4	A ₅₀ → {A ₄₇ [1], A ₅₀ [3]}

Step8: Calculate the forecasted outputs using equation $Forecast(t) = L_{df}(t-1) \cdot W_n(t-1)$. $L_{df}(t-1)$ is defuzzified matrix, $W_n(t-1)$ is weight matrix. In the same way, the Japan stock index of year 2014 can be forecasted. The mean square error (MSE) of a chromosome can be calculated using equation (3) as below.

$$MSE = \frac{\sum_{i=1}^m (Forecasted_{ClosingPrice_{ofYear_i}} - Actual_{ClosingPrice_{ofYear_i}})^2}{m} \quad (3)$$

Let the value of MSE be the fitness value of the genetic algorithm, where m denotes the number of historical data. The system will choose chromosomes with smaller MSE values for evolution. For example, from the experimental result, we can see that the MSE value of the chromosome.

Step9: Evaluate the performance. The RMSE is used to evaluate the forecasting performance.

$$RMSE = \sqrt{\frac{\sum_{i=1}^n (actual_t - defuzzification(t))^2}{n}} \quad (4)$$

TABLE XV: Rough set rule induction Japan stock index of year 2014

Date	Actual closing price	Forecasted closing price by FTS	Forecasted closing price by the proposed method
2014/11/21	17108.20		
2014/11/25	17372.71	17600.00	17237.63
2014/11/26	17335.99	17360.00	17237.67
2014/11/27	17212.48	17253.33	17084.67
2014/11/28	17330.84	17360.00	17237.67
2014/12/01	17474.27	17621.67	17237.67
2014/12/02	17476.42	17621.66	17676.50
2014/12/03	17696.36	17840.00	17586.75
2014/12/04	17813.01	17626.67	17810.00
2014/12/05	17759.36	17626.67	17810.00
2014/12/08	17869.74	17626.67	17810.00
2014/12/09	17773.30	17626.67	17810.00
2014/12/10	17308.82	17093.33	17237.67
2014/12/11	17043.46	17280.00	17084.67
2014/12/12	17298.19	17093.33	17237.67
2014/12/15	17037.21	17093.33	17084.67
2014/12/16	16714.83	16840.00	16772.00
2014/12/17	16672.94	16640.00	16632.00
2014/12/18	17136.41	17520.00	17123.00
2014/12/19	17471.60	17626.67	17237.67
RMSE		217.64	112.16

V. Empirical analysis

This study used the Japan stock index as the forecasting target. The data from January to October were used to form the estimation while those for November were used for the forecasting. For comparison purposes, we calculated the RMSE which is shown in the Table XV. Comparing the performance of proposed model with that of conventional model we found that the RMSEs from our proposed model is better than the conventional model. We can see that our proposed method achieve a higher forecasting accuracy in Fig.3. In the future, we are considering improve the heuristic algorithm calculation affective.

VI. Conclusions

In this article, we have presented the new method to forecast the Japan stock index based on fuzzy time series genetic algorithms and rough set. The year of 2014 Japan stock index was used as the forecasting target. The empirical analyses showed that the RMSEs in the case of proposed model are all smaller than FTS model. That is, the proposed method can achieve a higher forecasting accuracy rate than the existing methods.

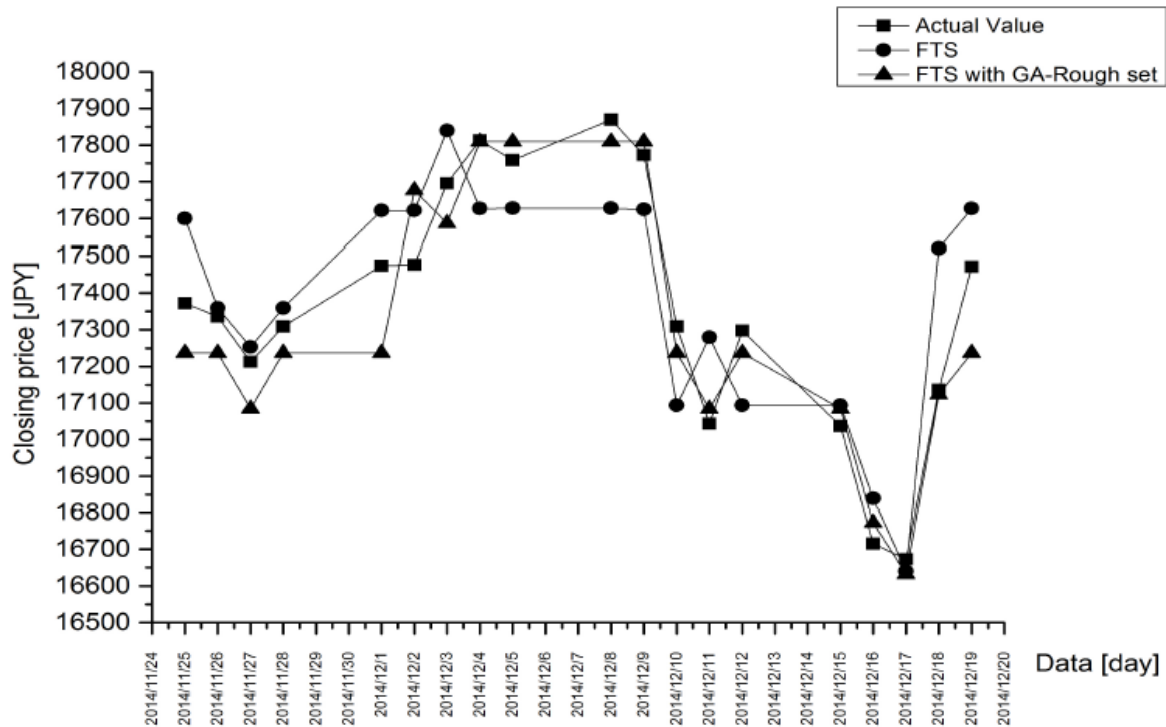


Fig. 3: Performance comparison for forecasting Nikkei Index Data

VII. References

- [1] Z. Pawlak, Rough Sets, International Journal of Computer and Information Science, (1982) 341-356.
- [2] Z. Pawlak, Rough Sets – Theoretical Aspects of Reasoning about Data, Kluwer Academic Publishers, (1991).
- [3] Y. Matsumoto, J. Watada, Rough set model based knowledge acquisition of market movements from economic data, Information Granularity, Big Data, and Computational Intelligence, Studies in Big Data Volume 8, (2015) 375-388.
- [4] L.A. Zadeh, Fuzzy sets, Inform and control 8, (1965) 338-353.
- [5] S.-M. Chen, Forecasting enrollments based on fuzzy time series, Fuzzy Sets and Systems 81, (1996) 311-319.
- [6] Holland JH. Adaption in natural and artificial systems. Cambridge, MA: MIT Press, (1975).
- [7] Davis L. Handbook of genetic algorithms. New York: Van Nostrand Reinhold, (1991).
- [8] Gen M, Cheng R, Genetic algorithms and engineering design. New York: John Wiley & Sons, (1997).
- [9] Goldberg DE, Korb B, Deb K. Messy, Genetic algorithms: Motivation, analysis, and first results, Complex Syst, (1989) 493-530.
- [10] Shyi-Ming Chen, Nien-Yi Chung, Forecasting Enrollments Using High-Order Fuzzy Time Series and Genetic Algorithms, Journal International Journal of Intelligent Systems archive, (2006) 485 - 501.
- [11] Tan, S., Cheng, X., Xu, H, An efficient global optimization approach for rough set based dimensionality reduction. Int. J. Innov. Comput., Info. Control 3(3), (2007) 725-736.
- [12] Ching-Hsue Cheng, Hia Jong Teoh and Tai-Liang Chen, Forecasting stock price index using fuzzy time series based on rough set, Fuzzy Systems and Knowledge Discover, (2007).

Sensitivity Analysis of portfolio selection models

Huiming ZHANG ¹, Ye LI ², Junzo WATADA ³

^{1,2,3} The Graduate School of Information, Production and Systems, Waseda University,
Kitakyushu, Fukuoka 808-0135, Japan.

E-mail: ¹ (huimingde@gmail.com), ² (liye@akane.waseda.jp), ³ (junzow@osb.att.ne.jp).

Abstract—Sensitivity Analysis is a method to evaluate the influence of each variable change. In portfolio selection model (PSM), it is essential to evaluate the sensitivity of each stock or security return rate in investment decision making. Investors look for selecting stable stocks or securities. For this purpose, sensitivity analysis should play a pivotal role. Under consideration of the sensitivity, risk (stability) and return of each selection, this paper proposes a new portfolio-selection model called the sensitivity-based portfolio selection model (SA-PSM).

Keywords: Sensitivity analysis, portfolio selection model.

I. INTRODUCTION

Portfolio selection model has been well developed based on a mean-variance approach proposed by Markowitz [2]. More than 50 years have been spent to study the portfolio under uncertain environment random selection development. The key principle of the mean-variance model is to use the expected returns of the portfolio as well as the investment risk in decision making. Each asset's monetary value may affect the risk and reward ratio of the portfolio and is referred to as the asset allocation of the portfolio. When determining a proper asset allocation, we target to maximize the expected return and minimize the risk. It is impossible to simulate the influence of all factors well. Therefore, in today's stock markets, it is not reasonable to evaluate future security returns only based on the historical data. In order to forecast more accuracy, the expert's opinions from related fields need to be considered. The model SA-PSM is mainly used in the stock market, analyzing and finding the balance between the value and risks. In finance, security returns in conventional models are always determined by historical data. However, such data are not always available. On the other hand, securities are always

affected by infinitive factors. It is impossible to simulate the influence of all factors well. Therefore, in today's stock markets, it is not reasonable to evaluate future security returns only based on the historical data. In order to forecast more accuracy, the expert's opinions from related fields need to be considered [4]. The model SA-PSM is mainly used in the stock market, analyzing and finding the balance between the value and risks.

The sensitivity analysis has been used in a range of purposes [3]. Including (1) Testing the robustness of the results or the uncertainty of the system (2) Increased understanding of the relationships between input and output variables in a system or model (3) Reducing the uncertainty: identifying model inputs (4) Searching for errors in the model (by encountering unexpected relationships between inputs and outputs) (5) Model simplification-fixing model inputs which have no effect on the output, or identifying and removing redundant parts of the model structure (6) Enhancing communication from modelers to decision makers (e.g. by making recommendations more credible, understandable, compelling or persuasive) (7) Finding regions in the space of input factors for which the model output is either maximum or minimum or meets some optimum criterion (see optimization and Monte Carlo filtering). For example, in a decision problem, the analyst may want to identify cost drivers as well as other quantities for which we need to acquire better knowledge in order to make an informed decision. On the other hand, some quantities have no influence on the predictions, so that we can save resources at no loss in accuracy by relaxing some of the conditions. However, the sensitivity analysis cannot analyze the uncertainty and the instability.

In this study, proposing a new portfolio-selection model (PSM) called the Sensitivity Analysis based on portfolio-selection-models (SA-PSM). This mod-

el use both types risk measure methods: variance is utilized to evaluate the portfolio stability. So this model can improve the stableness of the capital allocation. Therefore, the obtained selection result could be more helpful in decision making.

The remainder of this paper is organized as follows. Section 2 provides some basic mathematical preparation on portfolio selection model. In the section 3, we explain the method of sensitivity analysis and sensitivity-based portfolio selection model. Section 4 illustrates a numerical example and Finally, conclusions are drawn in Section 5.

II. PORTFOLIO SELECTION MODEL

Sensitivity analysis plays a pivotal role in uncertainty quantification and propagation of uncertainty.

This paper introduces two steps of the sensitivity analysis. Let us define the portfolio selection model before discussing the sensitivity based portofolio selection model (SPSM).

Model 2.1: (Portfolio Selection Model)

In order to construct the sensitivity, let us consider the following models.

$$\begin{cases} \max F(x) & \{= \sum_{i=1}^n x_i a_i\} \\ \min G(x) & \{= \sum_{i=1}^n \sum_{j=1}^n x_i \sigma_{ij} x_j\} \\ \text{Subject to} & \sum_{i=1}^n x_i = 1, x_i \geq 0, \quad i = 1, 2, \dots, n \end{cases} \quad (1)$$

III. SENSITIVITY ANALYSIS

In order to minimize the sensitivity, let us consider the following values.

Let us define the sensitivity of the linear function $O(x)$ of two objectives $F(x)$ and $G(x)$ at variable x_i in the following, where F_{x_i} and G_{x_i} denote the partial differentiation of the objectives $F(x)$ and $G(x)$ by variable x_i , respectively. Then, the sensitivity of $O(x)$ by t can be calculated as

Definition 3.1: Sensitivity

The sensitivity TS of the objective functions $O(x) = \omega_F F(x) + \omega_G G(x)$ at x_i can be defined as follows:

$$\begin{aligned} TS &= \frac{dO(x)}{dt} \\ &= \sum_{i=1}^n \frac{\partial O(x)}{\partial x_i} \frac{dx_i}{dt} \\ &= \omega_F \frac{dF(x)}{dt} + \omega_G \frac{dG(x)}{dt} \\ &= \sum_{i=1}^n \left(\left\{ \omega_F \frac{\partial F(x)}{\partial x_i} + \omega_G \frac{\partial G(x)}{\partial x_i} \right\} \frac{dx_i}{dt} \right) \\ &= \sum_{i=1}^n (w_F a_i + w_G 2 \sum_{j=1}^n \sigma_{ij} x_j) \frac{dx_i}{dt} \quad (2) \end{aligned}$$

where

$$\begin{aligned} \alpha_i = F_{x_i} &= \frac{\partial \sum_{i=1}^n x_i a_i}{\partial x_i} \\ &= a_i \quad (3) \end{aligned}$$

$$\begin{aligned} \beta_i = G_{x_i} &= \frac{\partial \sum_{i=1}^n \sum_{j=1}^n \sigma_{ij} x_i x_j}{\partial x_i} \\ &= 2 \sum_{j=1}^n \sigma_{ij} x_j \quad (4) \end{aligned}$$

and ω_F , ω_G denote the weight of the objectives, respectively and $\omega_F = \omega_G = 1$.

That means, the total sensitivity is calculated using α_i , β_i and the change of return rate per time.

After processing equations Equations (2) we can get the sensitivity of each objective function, which are defined as F_{x_i} , G_{x_i}

Note that from here ω_F , ω_G are set to 1.

Take this ten stocks as the return, we can calculate out the sensitivity of each variable, and also, we can calculate the total sensitivity which will be shown below:

Table II shows the result of sensitivity analysis.

Then we can understand which variables are more sensitive and which variables are more stable. But we would better select portfolio with more stable investment, that is, less sensitive variables (stable stocks). Then we formulate the problem including the sensitivity in the objective, The total sensitivity TS is defined as Equations (2).

TABLE I
DATA EMPLOYED FROM NYSE [4]

Stock	STV	CNEP	CNTF	CBEH	CHL	ACH	CAST	DANG	CNAM	CBAK
	x_1	x_2	x_3	x_4	x_5	x_6	x_7	x_8	x_9	x_{10}
a_i	-0.01503	0.00380	0.01031	0.02711	0.00371	-0.00903	-0.00118	-0.01145	-0.03755	-0.02906
$\sigma(i, j)$	1	-0.9505	-0.5349	-0.756	-0.5535	-0.3448	0.2036	-0.8166	-0.6358	0.6679
	-0.9505	1	0.3157	0.5492	0.7762	0.6183	-0.3065	0.8517	0.6734	-0.8331
	-0.5349	0.3157	1	0.9504	-0.3257	-0.4419	0.6381	0.5863	-0.2216	-0.1812
	-0.756	0.5492	0.9504	1	-0.0977	-0.2783	0.3737	0.6746	0.0917	-0.296
	-0.5535	0.7762	-0.3257	-0.0977	1	0.9609	-0.5993	0.5361	0.6883	-0.8017
	-0.3448	0.6183	-0.4419	-0.2783	0.9609	1	-0.4833	0.4668	0.4863	-0.7962
	0.2036	-0.3065	0.6381	0.3737	-0.5993	-0.4833	1	0.2212	-0.8821	0.0197
	-0.8166	0.8517	0.5863	0.6746	0.5361	0.4668	0.2212	1	0.1862	-0.9011
	-0.6358	0.6734	-0.2216	0.0917	0.6883	0.4863	-0.8821	0.1862	1	-0.2832
	0.6679	-0.8331	-0.1812	-0.296	-0.8017	-0.7962	0.0197	-0.9011	-0.2832	1

TABLE II
COMPUTATION RESULTS

Stock	STV	CNEP	CNTF	CBEH	CHL	ACH	CAST	DANG	CNAM	CBAK
	x_1	x_2	x_3	x_4	x_5	x_6	x_7	x_8	x_9	x_{10}
<i>PSM</i> Capital Allocation				50.4%	40.8%			9.0%		
<i>PSM</i> Total (TS_{x_i})	0.0004	0.0001	-0.0004	0.0111	0.0717	0.0001	0.0001	0.0065	0.0005	-0.0007
<i>SPSM</i> Capital Allocation	32.7%			31.2%	27.9%					8.2%
α_i	0.0004	0.0001	-0.0004	0.0010	0.0002	0.0001	0.0001	-0.0001	0.0005	-0.0007
β_i	-0.1475	0.1658	-0.0459	0.0201	0.1754	0.1382	-0.1769	0.0732	0.2072	-0.1018
$\frac{dx_i}{dt}$	-0.0250	0.0311	-0.0345	0.0352	0.0437	-0.0094	-0.0575	0.0053	-0.0142	0.0253
<i>SPSM</i> Total (TS_{x_i})	-0.0479	0.0001	-0.0004	0.0072	0.0491	0.0001	0.0001	-0.0001	0.0005	-0.0091

IV. BUILDING SENSITIVITY-BASED PORTFOLIO SELECTION MODEL

Using the definition 3.1, we can define the sensitivity-based portfolio selection model as follows:

Model 4.1: (the sensitivity-based portfolio selection model)

$$\text{Return} : \max \sum_{i=1}^n x_i a_i \quad (5)$$

$$\text{Risk} : \min \sum_i \sum_j \sigma_{ij} x_i x_j \quad (6)$$

$$\text{Sensitivity} : \min \sum_{i=1}^n \{a_i \frac{dx_i}{dt} + 2 \sum_{j=1}^n (\sigma_{ij} \frac{dx_i}{dt} x_j)\} \quad (7)$$

Subject to

$$\sum_{i=1}^n x_i = 1 \quad (8)$$

$$x_i \geq 0 \quad i = 1, 2, \dots, n \quad (9)$$

where ω_F , ω_G are given as a weight for each objective function. In this discussion, let us set them to 1

Let us approximate $\frac{dx_i}{dt}$ by the following:

$$\frac{dx_i(t)}{dt} \simeq \frac{x_i(t + \Delta t) - x_i(t)}{t + \Delta t - t} = \frac{x_i(t + \Delta t) - x_i(t)}{\Delta t} \quad (10)$$

This value $\frac{dx_i(t)}{dt}$ can be approxiamtely evalluted by regression analysis or Change rate of return $\frac{x_i(t)}{dt} \simeq \frac{x_i^{max} - x_i^{min}}{7}$ where x_i^{max} and x_i^{min} are the maximum and minimum values in a week.

The calculation can be performed based on Table II. We may using regression analysis prepare the changing rate of return rate x_i , that is, $\frac{dx_i}{dt}$.

V. ILLUSTRATIVE APPLICATION

Let us analyze the stocks of NYSE in December 2011. Table II shows the result obtained by the portfolio selection model 2.1 and the sensitivity-based portfolio selection model 4.1. In the portfolio selection model 2.1, CBEH and CHL were allocated 50.4% and 40.8% of the total investment. But in the sensitivity-based selection model 4.1 these are rather high sensitive variables with 0.0072 and 0.0491, respectively comparing with other stocks. Therefore, these stocks are not stable.

The sensitivity-based portfolio selection model 4.1 changed the allocation for stocks STV, CBEH and CH with 32.7%, 31.2% and 27.9%. So the total investment allocation is distributed to the less sensitive stock STV with -0.0479 as well.

VI. CONCLUSIONS

In this paper, we have introduced a sensitivity analysis for the portfolio selection model, which shows both the sensitivity of each variance and the total sensitivity of the model. To the existing PSMs, the approach is more acceptable for general investors, so the calculation of the sensitivity will be quite necessary. To calculate out sensitivity of each variance and the total sensitivity, we have improved the basic PSM models to get the definition of sensitivity which can be expressed in a more clear mathematic way. Also from the experiment result we can find out that, to the PSMs, different variance has different sensitivity which will make influence to the analysis of the whole model, to the total sensitivity, it will show a overall description of whether the model is stable or sensitive. To the proposed PSM, the proposed sensitivity analysis is proved useable, for calculating out both the sensitive variances and the stable variances. In the further research, we will introduce the proposed model into some other applications, and also, we will do some more analysis not only in the value of sensitivity but also the relationship between different variance's sensitivity.

REFERENCES

- [1] X. Huang, "Mean-semivariance models for fuzzy portfolio selection," J.Comput. Appl. Math., 217(1) 1-8 (2008).
- [2] H. Markowitz, "Portfolio selection," J. Finance, 7(1) 77-91 (1952)
- [3] Pannell, D.J. (1997). Sensitivity analysis of normative economic models: Theoretical framework and practical strategies, Agricultural Economics 16: 139-152.
- [4] Z. Zmeskal, "Value at risk methodology of international index portfolio under soft conditions," Int. Rev. Financial Anal., 14(6) 263-275 (2005)
- [5] R. A. Jabr, "Robust self-scheduling under price uncertainty using conditional value-at-risk," IEEE Trans. Power Syst., 20(4) 1852-1858 (Nov. 2005)
- [6] HM ZHANG, J.WATADA, Ye LI, You LiI and Bo WANG, "Building A Sensitivity-Based Portfolio Selection Models," Proceedings, KES-IDT2015, Italy, 15-18, June 2015.
- [7] D. Dubois and H. Prade, Possibility Theory. New York: Plenum, 1988
- [8] B. Liu and Y.K. Liu, "Expected value of fuzzy variable and fuzzy expected value model," IEEE Trans. Fuzzy Syst., 10(4) 445-450 (Aug. 2002)

Fuzzy Variable-Based Two-echelon Supply Chain Model

Junzo Watada¹ and OXian Chen²

^{1,2} Waseda University

The Graduate School of Information, Production and Systems, Kitakyushu, Fukuoka 808-0135, Japan. e-mail:

¹ (junzow@osb.att.ne.jp)

² (chenxian@fuji.waseda.jp)

Abstract—Game theory is applied widely to solving various problems to get optimal decisions. This paper analyzes the Stackelberg behaviors between a manufacturer and two retailers. As the market demand used to be uncertain, in order to describe the market demand fitting to real situations, this paper applies fuzzy variable in the two echelon supply chain problem. Then, numerical examples are used to illustrate the result in comparison between fuzzy variable and crisp number. Fuzzy variable and crisp number are compared with mean absolute percentage difference rate of fuzzy demand (MAPDR-FD).

Index Terms—Two-Echelon Supply Chain, Fuzzy Variable, Optimal Decision, Stackelberg

I. INTRODUCTION

GAME theory becomes more and more common in the research on two-echelon supply chain and plays a pivotal role in optimal decision making.

We used to face how much amount to assign each of branches or of sections as a target amount in their region. The most important is the whole company's selling volume but on the other hand we have to encourage each of branches or of sections to work hard and obtain the best result.

In this paper, let us make a problem simplified. We have the manufacturing section M and two retailing sections 1 and 2. These three sections build a supply chain. In this case, there is a competition between the two retailing sections 1 and 2. In game situation, these two retailers 1 and 2 are Stackelberg situation. Therefore, when we give a leading position to 1, Section 2 will follow section 1's initial decision. The objective here is to obtain their own best result in the Stackelberg situation. Then, according to the best results of the two retailers got before, The manufacturer can know the decision made by retailers and decide the optimal assignment of price, amount, etc.

In order to solve this problem, let us consider the uncertain situation such as uncertain demand. This situation resulted in building a fuzzy two-echelon supply chain model. In this paper, fuzzy variable is used in building model of two-echelon supply chain.

The remaining consist of the following sections. In Section 2, past research works will be shortly reviewed and Section 3 provides prerequisite mathematical preparation for building a fuzzy two-echelon supply

chain model. In Section 4 we define fuzzy two-echelon supply chain model and its solution based on Stackelberg game. Then, Section 5 discusses the numerical real-life problem of the strategic decision-making of how to obtain the optimal assignment for each of the two retailer sections. At the end in Section 6 we summarize the paper as conclusions and explain the remaining problems.

II. PREVIOUS RESEARCH RESULTS

In two-echelon supply chain, quantity, cost and price are key factors in order to get optimal decision. For example, Pui-Sze CHOW, et al. [1] analyze the effect when taking a minimum order quantity (MOQ) in a two-echelon supply chain. Govindan KANNAN [2] uses multi-objective models in a two echelon supply chain to optimize the total cost. Ilham SLIMANI and Said ACHCHAB [3] concerned with the optimization of the inventory and transportation in a supply chain. Y.A. HIDAYT, et al. [4] analyze an inventory problem with only a single-supplier and a single-buyer. Cai Jian-hu and Wang Li-ping [5] analyze the influence of providing adequate pricing incentives before demand information is revealed. Geon Cho, et al. [6] take a quantitative model when there is a partnership in supply chain. Daogang Qu and Ying Han [7] consider a single manufacturer and a single retailer model which has dual distribution. Longfei He and Jianyong Sun [8] consider a single vendor and a single retailer originated from China market background. CAI Jian-hu, et al. [9] proposed an inventory model in a two-echelon supply chain. Yihong Hu and Jianghua Zhang [10] consider that

when there is a price-only contract to a retail, how will be the manufacturer's reaction to a new vendor. Fangxu Ren and ZhongYuan [11] analyze pricing and advertising model under equilibrium. Tiezhu Zhang and Longying Li [12] consider some contracts in order to get the win-win situation. Huilin Chen and Kejing Zhang [13] take the full return coordination mechanism in consider. Fangxu Ren [14] considers the relation between the size of slotting allowance with the members' co-operation in a two-echelon supply chain. Moreover, there are also many researches studying the demand function. Junping Wang and Shengdong Wang [15] study the effect to the chain members when the different forms of demand function are used. Jiang Meixian, et al. [16] consider a model of coordinating pricing and express that the market demand is effected by retailer price and promotion expenses. Lau and Lau

[17] analyze the difference of the optimal decision when facing the different demand curve in Stackelberg game situation.

III. PREPARATIONS

As the demand of the market is always uncertain. This paper takes the demand of the market as fuzzy variable so as to make the value more close to real situation. To illustrate the difference between the fuzzy variable and crisp number, Mean Absolute Percentage Different Rate of the Fuzzy Demand (MAPDR-FD) will be used.

A. Fuzzy Variable

Fuzzy Variable is defined as a function which has a possibility space(Θ , $p(\Theta)$,Pos),where the Θ is a non-empty set, $p(\Theta)$ is a power set of Θ ,and Pos is a possibility measure. The possibility measure satisfies the following conditions:

- (1) $\text{Pos}\{\emptyset\}=0$
- (2) $\text{Pos}\{\Theta\}=1$
- (3) $\text{Pos}\{\bigcup_{i=1}^m A_i\}=\sup_{1 \leq i \leq m} \text{Pos}\{A_i\}$

By taking the triangular fuzzy variable $D=(a,b,c)$, the membership function will be given as follows:

$$\mu(x) = \begin{cases} \frac{x-a}{b-a} ; & a \leq x \leq b \\ \frac{c-x}{c-b} ; & b \leq x \leq c \\ 0 ; & \text{otherwise} \end{cases}$$

Supposing A^c is the complement set of A , then the necessity measure of A can be given as follows:

$\text{Nec}(A)=1-\text{Pos}(A^c)$,then the credibility measure is denoted as

$$\text{Cr}\{A\} = \frac{1}{2}[1 + \text{Pos}(A) - \text{Pos}(A^c)];$$

$$\text{Cr}\{A\} = \frac{1}{2}[\text{Pos}(A) + \text{Nec}(A)]$$

Let D be a fuzzy variable. The expected value of D can be give as follows:

$$E[D] = \int_0^{\infty} \text{Cr}\{D \geq r\} dr - \int_{-\infty}^0 \text{Cr}\{D \leq r\} dr ;$$

When D is a triangular fuzzy number(a,b,c)

$$E[D] = \frac{a + 2b + c}{4}$$

B. Model Assumptions and Notations

This paper is to analyze the behavior of a supply chain with a manufacture and two retailers. The assumptions and notations are given as follows:

Q_i : quantity ordered by retailer i , ($i = 1,2$);

c : unit cost of manufacturing ;

w : the wholesale price per unit ;

p_i : the sale price set by retailer i , ($i = 1,2$);

D_i : the demand of retailer i if prices are zero.

Π_i :retailer i 's profit ($i = 1,2$);

Π_M :manufacture's profit;

Π_i^* :retailer i 's maximum profit ($i = 1,2$);

Π_M^* :manufacture's maximum profit.

The downward-sloping demand function is defined as follows:(McGuire and Staelin(1983)and Ingene and parry(1995) have used this type of demand curve with $a_i = 1$ successively). In this paper will also let $a_i = 1$. So the demand function is:

$$Q_i = (D_i - p_i + \theta p_j) \quad (i, j=1,2) \quad (1)$$

θ is the degree of substitutability between retailers.

IV. TWO-ECHELON SUPPLY CHAIN MODEL BASED ON STACKELBERG GAME.

In this Stackelberg game, we assume retailer 1 as a Stackelberg leader and retailer 2 as a Stackelberg follower. Then the profit function of retailer 1 and retailer 2 will be given as:

$$\Pi_{r1} = (p_1 - w)Q_1 = (p_1 - w)(D_1 - p_1 + \theta p_2) \quad (2)$$

$$\Pi_{r2} = (p_2 - w)Q_2 = (p_2 - w)(D_2 - p_2 + \theta p_1) \quad (3)$$

The reaction function of the retailer2 will be given as

follows(obtained by solving $\frac{d\Pi_{r2}}{dp_2} = 0$):

$$p_2 = \frac{(D_2 + w + \theta p_1)}{2} \quad (4)$$

Then retailer 1 knows the reaction function(4) of retailer 2 and substitute it into his profit function. Then the profit function of retailer 1 can be given as follows:

$$\Pi_{r1} = \frac{(p_1 - w)(2D_1 + \theta D_2 - 2p_1 + \theta^2 p_1 + \theta w)}{2} \quad (5)$$

To get the optimal price of retailer 1, setting $\frac{d\Pi_{r1}}{dp_1} = 0$, the function can be given as follows:

$$p_1^* = \frac{(2D_1 + \theta D_2 + 2w + \theta w - \theta^2 w)}{2(2 - \theta^2)} \quad (6)$$

Then the maximum profit of retailer 1 can be given as

$$\Pi_{r1}^* = \frac{(2D_1 + \theta D_2 + \theta w + \theta^2 w - 2w)^2}{8(2 - \theta^2)} \quad (7)$$

To get the optimal sale price and the maximum profit of retailer 2, the optimal sale price of retailer 1 can be substituted into reaction function of retailer 2. Then the function can be given as:

$$p_2^* = \frac{(2\theta D_1 + 4D_2 - \theta^2 D_2 + 4w + 2\theta w - \theta^2 w - \theta^3 w)}{4(2 - \theta^2)} \quad (8)$$

$$\Pi_{r2}^* = \frac{(\theta^2 D_2 - 2\theta D_1 - 4D_2 + 4w - 2\theta w - 3\theta^2 w + \theta^3 w)^2}{16(2 - \theta^2)^2} \quad (9)$$

The profit function of manufacture can be expressed as:

$$\Pi_M = (w - c)(Q_1^* + Q_2^*)$$

$$\Pi_M = (w - c)[D_1 + D_2 - (p_1^* + p_2^*) + \theta(p_1^* + p_2^*)] \quad (10)$$

Then the optimal price can be obtained by solving $\frac{d\Pi_M}{dw} = 0$.

$$w^* = \frac{c}{2} + \frac{B}{2J}, \quad (11) \quad \text{where}$$

$$B = 2[\theta + 2(D_1 + D_2)] - \theta[-2D_1 + \theta^2 D_2 + \theta(2D_1 + D_2)]$$

$$J = (8 - 4\theta - 7\theta^2 + 2\theta^3 + \theta^4)$$

So the maximum profit of the manufacture can be given as follows:

$$\Pi_M^* = \frac{(8 + 4\theta + 3\theta^2 - \theta^3)(c - c\theta - D_1 D_2)}{16(\theta - 1)(2 - \theta^2)} \quad (12)$$

V. NUMERICAL EXAMPLES

The numerical examples let us compare the difference of results between the fuzzy and non-fuzzy data sets. Setting $\theta = 0.5$ represents the moderate degrees of competition between two retailers.

TABLE I
ASSUMED PARAMETERS

c	θ
2	0.5

A. Mean Absolute Percentage Difference Rate of the Fuzzy Demand (MAPDR-FD)

To find the difference between fuzzy and nonfuzzy data sets, the mean absolute percentage difference rate of the fuzzy demand is introduced. The formula is expressed as follows:

$$MAPDR - FD = \frac{100\%}{n} \sum_{k=1}^n \left| \frac{\pi_k^l - \pi_{Fk}^l}{\pi_k^l} \right|$$

where π_k^l is the non-fuzzy value and π_{Fk}^l is the fuzzy value. Where $l=i, j, M$ and $k=1, 2, \dots, n$.

B. Fuzzy Variables

Let X_{D1} be a fuzzy variable of demand faced by retailer1. Assume $D_1 = (9, 19, 32)$

Let X_{D2} be a fuzzy variable of demand faced by retailer2. Assume $D_2 = (7, 18, 28)$

The possibility distribution of the triangular fuzzy variable are written as:

$$\mu_{D1}(x) = \begin{cases} \frac{x-9}{10} & ; 9 \leq x \leq 19 \\ \frac{32-x}{13} & ; 19 \leq x \leq 32 \\ 0 & ; \text{otherwise} \end{cases}$$

$$\mu_{D2}(x) = \begin{cases} \frac{x-7}{11} & ; 7 \leq x \leq 18 \\ \frac{28-x}{10} & ; 18 \leq x \leq 28 \\ 0 & ; \text{otherwise} \end{cases}$$

According to the possibility, necessity and credibility measure, the expected value can be calculated using following equation.

$$E[\xi] = \frac{a + 2b + c}{4}$$

TABLE II
FUZZY DEMAND OF MARKET FACED BY
SUPPLIER 1 AND 2

	Fuzzy Triangular	Expected Value
D_1	(9,19,32)	19.75
D_2	(7,18,28)	17.75

C. Comparing The Results Between Fuzzy and NonFuzzy Data Sets.

TABLE III
THE DIFFERENCE BETWEEN FUZZY NUMBER
AND NO-FUZZY DATA SETS.

	Fuzzy Data Set ($\times 10^2$)	Non-Fuzzy Data Set ($\times 10^2$)	MAPRD-FD
Π_{r1}^*	0.483	0.481	+4.3%
Π_{r2}^*	0.393	0.451	-13.2%
Π_M^*	2.653	2.588	+2.528%
P_1^*	0.200	0.248	-19.2%
p_2^*	0.187	0.240	-22.1%
Q_1^*	0.09	0.06	+45.6%
Q_2^*	0.09	0.07	+42.2%

From Tables 2 and 3, it is easy to find that by using fuzzy data set, the profit of manufacture is higher than the one without fuzzy data set. The quantity ordered by retailers using fuzzy number is more than the one using crisp number. On the other hand, the saleprices setted by retailers 1 and 2 using fuzzy number are lower than these using crisp number. When the fuzzy demand is less than crisp demand, the profit of retailer with fuzzy variable is lower than that without using fuzzy data set. On the other hand, when the fuzzy demand is greater than crisp demand, the profit of retailer using fuzzy data set is higher than that without using fuzzy data set. In addition, the MAPRD-FD shows the difference between the fuzzy data set and crisp number data set.

As a result, in this situation, when the retailer sections are competitive with each other, the retailer 1 and 2 can get the optimal price and amount based on the stackelberg game analysis. According to the result, the manufacture can also decide the optimal assignment of target price and amount. As the market demand is uncertain, we employed fuzzy variables in this calculation. Then we got the result that the quantity ordered by each of retailers 1 and 2 should be 0.09. The optimal profit of retailers 1,2 and manufacture will be 0.483, 0.393 and 2.653. Then the profit shown in Table 3 each section can get the price and amount.

D. Real Life Problem

The data of Coca-Cola firm is analyzed to explain the different profit by using fuzzy and crisp number data sets.

Let X_{D1} be a fuzzy variable of demand faced by LanZhou market. $D_1=(11736,20979,52689)$

Let X_{D2} be a fuzzy variable of demand faced by XiNing market. $D_2=(5844,15741,22140)$

Then the possibility distribution of the triangular fuzzy variable can be written as:

$$\mu_{D1}(x) = \begin{cases} \frac{x-11736}{9243} & ; 11736 \leq x \leq 20979 \\ \frac{52689-x}{31710} & ; 20979 \leq x \leq 52689 \\ 0 & ; otherwise \end{cases}$$

$$\mu_{D2}(x) = \begin{cases} \frac{x-5844}{9897} & ; 5844 \leq x \leq 15741 \\ \frac{22140-x}{6399} & ; 15741 \leq x \leq 22140 \\ 0 & ; otherwise \end{cases}$$

The same as above, according to the possibility, necessity and credibility, we can get the expect value. the results are shown in tables 4 and 5.

TABLE IV
FUZZY DEMAND FACED BY LANZHOU AND XI
NING MARKET

	Fuzzy Triangular number ($\times 10^4$)	Expected Value ($\times 10^4$)
D_1	(1.174,2.098,5.269)	2.660
D_2	(0.584,1.574,2.214)	1.487

TABLE V
THE DIFFERENCE BETWEEN THE FUZZY
NUMBER AND NO- FUZZY NUMBER

	Fuzzy Data Set ($\times 10^6$)	Non-Fuzzy Data Set ($\times 10^6$)	MAPRD-FD
Π_{r1}^*	97	60	+60%
Π_{r2}^*	28	34	-16.1%
Π_M^*	300	251	+19.7%
P_1^*	0.29	0.32	-8%
p_2^*	0.24	0.29	-18.1%
Q_1^*	0.092	0.038	+142%
Q_2^*	0.053	0.020	+162%

From Table5, we can get a conclusion that by using fuzzy number the profit of the manufacture is higher than using crisp numbers. When the fuzzy demand is less than crisp demand, the profit of retailer using fuzzy variable is lower than use crisp numbers. On the other hand, when the fuzzy demand is greater than crisp demand, the profit of retailer using fuzzy variable is higher than use crisp numbers.

The saleprices setted by retailer1 and 2 using fuzzy data set are lower than these using crisp number. On the contrary, under the fuzzy demand environment, the quantities ordered by retailers1 and 2 are more.

From the stackerberg game analysis, we can get the optimal assignment of price and amount. As the demand is not certain, we employ fuzzy variable to describe it. It is easy to find that though saleprices setted by retailers using fuzzy data set are lower than these using crispe number, the quantities ordered by retailers under fuzzy demand enviroment are more than using crisp number. As the profit by using fuzzy data set is higher, it is more benifit for manufacture to use fuzzy data set. We can also find that the the more the profit is, the bigger the difference between using the fuzzy data set and crisp number.

To simplify the real problem, this paper analysis the behaviors in a duopolostic situation. The result showed in the table is not affected by other manufactures. So this result is different from the real environment, but we can surely know that by using fuzzy data set , the profit of manufacture is more than using crisp number.

Moreover,according to the MAPRD-FD, the degree of difference by using the fuzzy number and no-fuzzy number can be calculated.

In conclusion, we regard retailer 1 as the leader and retailer 2 as the follower, after doing the analysis based on stackelberg game, then we can get the optimal amount and price shown in table 5. The quantity ordered by retailer 1 and 2 should be 9,214 and 2,027 respectively. Also, the optimal decision of manufacture can be got. As the market demand is uncertain mentioned before, fuzzy variables will be employed in this real life problem. Then we can finally get the result that the profits of retailer 1 2 and manufacture are 97, 28 and 300 million USD.

VI. CONCLUSIONS AND FUTURE WORK

In this paper, the competitive between a manufacture and two retailers are shown. The stackbelberg game always be used to solve this problem in order to make a best decision. By using the fuzzy variable, the profit of the manufacture will be higher than the one which use crisp number.

This paper analyze a manufacture and two retailers. As the future work, to analyze the market more accurately, there are more retailers need to be considered. Facing the market demand using fuzzy variable, the best decision of the manufacture and retailers should be discussed.

REFERENCES

- [1] Pui-Sze CHOW, Tsan-Ming CHOI, and T. C. E. CHENG, "Impacts of Minimum Order Quantity on a Quick Response Supply Chain," IEEE transactions on system, man, and cybernetics-part A: systems and humans. 42, (NO. 4), JULY 2012;868 - 879.
- [2] Govindan KANNAN, "A Multi Objective Model For TwoEchelon Production Distribution Supply Chain," Computers & Industrial Engineering, 2009. CIE 2009. International Conference on;6-9 July 2009;624 - 627.
- [3] Ilham SLIMANI, CSaid ACHCHAB, "Game theory to control logistic costs in a two-echelon supply chain " , Logistics and Operations Management (GOL), 2014 International Conference on;5-7 June 2014;168 - 170.
- [4] Y.A. HIDAYAT, SUPRAYOGI, D. T. LIPUTRA, S. N. ISLAM, "Two-echelon supply chain inventory model with shortage, optimal reorder point, and controllable lead time," Management of Innovation and Technology (ICMIT), 2012 IEEE International Conference on;11-13 June 2012;163 - 167.
- [5] CAI Jian-hu, WANG Li-ping, "Study on Dynamic Game Models in a Two-Echelon Supply Chain," 2009 International Conference on Computational Intelligence and Security;11-14 Dec. 2009;138 - 141.
- [6] Geon CHO, Soon-Hoo SO, JaeJon KIM, Yi Sook PARK, Hyong Ho PARK, Kyung Ho JUNG, Choon-Sup NOH; "An Optimal Decision Making Model for Supplier and Buyer ' s Win-Win Strategy in a Two Echelon Supply Chain, " Hawaii International Conference on System Sciences, Proceedings of the 41st Annual;7-10 Jan. 2008, 69
- [7] Daogang QU Ying HAN; "Optimal Pricing for Dual-channel Supply Chain with Fairness Concerned Manufacturer," Control and Decision Conference (CCDC), 2013 25th Chinese;25-27 May 2013;1723 - 1727.
- [8] Longfei HE, Daozhi ZHAO, Jianyong SUN; "Channel Coordination on Supply Chain with Unbalanced Bargaining Power Structure by Exogenous-force Impact: China Market Background," Service Operations and Logistics, and Informatics, 2008. IEEE/SOLI 2008. IEEE International Conference on 12-15 Oct. 2008;2018 - 2023.
- [9] CAI Jian-hu, WANG Li-ping , SHAO Zu-feng; "Study on an Advance Order Tactic in a Two-Echelon Supply Chain,"2009 International Conference on Multimedia Information Networking and Security;18-20 Nov. 2009;613 - 617.
- [10] Yihong HU, Jianghua ZHANG, Zui XU, "Asymmetric demand information ' s impact on supply chain performance and relationship under price-only contract," Automation and Logistics,

- 2007 IEEE International Conference on;18-21 Aug. 2007;2891 - 2896.
- [11] Fangxu REN; "Dynamic Pricing Decision in the Two-echelon Supply Chain with Manufacturer's Advertising and Dominant Retailer," Control and Decision Conference (CCDC), 2011 Chinese;23-25 May 2011;391 - 395.
- [12] Tiezhu ZHANG, Longying LI; "Study on a combined Supply Chain contracts of target rebate and quantity discount," 2009 Second International Conference on Intelligent Computation Technology and Automation;10-11 Oct. 2009;988 - 992.
- [13] Huilin CHEN, Kejing ZHANG; "Stackelberg Game in a Twoechelon Supply Chain under Buy-back Coordination Contract," Service Operations and Logistics, and Informatics, 2008. IEEE/SOLI 2008. IEEE International Conference on;12-15 Oct. 2008;2191 - 2196.
- [14] Fangxu REN, "Research on Slotting Allowance Decision in the Two-echelon Supply Chain with Retailer-led," Logistics Systems and Intelligent Management, 2010 International Conference on ;9-10 Jan. 2010;25 - 29.
- [15] Junping WANG, "Shengdong Wang;Optimal Manufacturing, Ordering, Pricing and Advertising Decisions in a Two-Echelon Supply Chain System," Proceedings of the Sixth International Conference on Management Science and Engineering Management Lecture Notes in Electrical Engineering;Vol. 185, 2013;589-597.
- [16] JIANG Meixian, YAN Lianlian, JIN Shousong, FENG Dingzhong, "Coordinating Price Model for Retailer Dominant in a Two-echelon Supply Chain," Industrial Engineering and Engineering Management, 2009. IE&EM ' 09. 16th International Conference on;21-23 Oct. 2009;1474 - 1477.
- [17] THISANA W, ROSARIN D., & Junzo WATADA, "A game approach to competitive decision making in logistics," ISME, 2012.
- [18] Shan-Lin YANG, & Yong-Wu ZHOU, "Two-echelon supply chain models: Considering duopolistic retailers' different competitive behaviors, Int. J. Production Economics, Pages 104-116, 2006.
- [19] Global Task Force ' Game Theory ' Comprehensive Law Publishing, 2003
- [20] LAU, A., & LAU, H. (2005) "Some two-echelon supply-chain games: improving from deterministic symmetric-information to stochastic asymmetric formation models, " European Journal of Operational Research, 161(1) 203-223.
- [21] LAU A, LAU HS, "Effects of a demand-curve's shape on the optimal solutions of a multi-echelon inventory/pricing model, " European Journal of Operational Research, Res2003; 147:530-548.
- [22] LAU, H., & ZHAO, A. (1993). "Optimal ordering policies with two suppliers when lead times and demands are all stochastic," European Journal of Operational Research, 68(05), 120-153.
- [23] H.H. DING, "Research of supplies chain management in enterprise cooperation based on game theory," Journal of Shijiazhuang Railway Institute (Social Sciences, 3(3), pp. 20-24, 2009.
- [24] Y. J. GAO, Ch. G. WU, Y. EN. WANG, "Enterprises cooperation game analysis on supply chain management based on strategy," Journal of Hefei University of Technology (Social Sciences), 19 (04), pp. 40-43. (2005)
- [25] Yang LI, Li-hai WANG, "A Study on Manufacturer and Suppliers for Product Innovation Coordinating to Avoid the Risk Based on Game Theory," Industrial Engineering and Engineering Management (IE&EM), 2010 IEEE17th International Conference on

CONTACT ADDRESS

Junzo Watada and Xian Chen are with the Graduate School of Information, Production and Systems, Waseda University, Kitakyushu, Fukuoka 808-0135, Japan. e-mail:(junzow@osb.att.ne.jp);(chen xian@fuji.waseda.jp).

入力の線形結合をもつ SIRMs ファジィ推論モデルの能力について

A Consideration on the SIRMs Fuzzy Inference Model
with Linear Transformation of Input Variables

○¹ 宮島 洋文, ² 重井 徳貴, ³ 宮島 廣美
○¹ Hirofumi Miyajima, ² Noritaka Shigei, ³ Hiromi Miyajima
鹿児島大学 大学院 理工学研究科

Graduate School of Science and Engineering, Kagoshima University

Abstract: In this paper, we propose a fuzzy inference system, apply to obstacle avoidance problem and show the effectiveness.

1 はじめに

ファジィ推論システムの自動構築に関する多くの研究が行われている。このシステムの特徴の一つは、言葉による情報処理を可能にすることである。この場合、入力変数が増えると扱いが困難になるという問題がある [1, 2]。この問題を解決するモデルの一つとして SIRMs (Single-Input Rule Modules) モデルを用いた研究も行われているが、必ずしも満足できるモデルとはなっていない [1]。本稿では、入力変数を線形変換した後に SIRMs モデルを構築する新しいモデルを提案する [2]。また、その有効性を示すために障害物回避問題に応用する。

2 SIRMs モデル

$\mathbf{R}$ をすべての実数の集合、自然数 k に対して $Z_k = \{1, 2, \dots, k\}$ 、入力 $\mathbf{x} = (x_1, \dots, x_N)$ とする。ここで、 $x_j \in \mathbf{R} (j \in Z_N)$ SIRMs は、次のような if...then~ 形式のルールにより構成される。

$$R_i^j : \{ \text{if } x_j \text{ is } A_{ji} \text{ then } y \text{ is } c_{ji} \}_{i=1}^r \quad (1)$$

ここで、 A_{ji} は入力要素 x_j に関するファジィ集合、 c_{ji} は実数値、 $j \in Z_N, i \in Z_r$ である。このときのガウス型メンバーシップ関数を用いた SIRMs の出力 y は、以下のようになる。

$$A_{ji}(x) = \exp \left(-\frac{(x - a_{ji})^2}{2b_{ji}^2} \right) \quad (2)$$

$$y = \sum_{j=1}^N h_j \frac{\sum_{i=1}^r A_{ji}(x_j) c_{ji}}{\sum_{i=1}^r A_{ji}(x_j)} \quad (3)$$

ここで、 a_{ji} と b_{ji} は標準偏差、 c_{ji} は後件部実数値、 $h_j (j \in Z_N)$ は各モジュールに対する重視度である。

3 提案手法

3.1 線形入力型 SIRMs ファジィ推論システム

SIRMs は複雑な入出力関係を表すのが難しいことが知られている [2]。そこで、本稿では入力変数を線形変換した後に SIRMs モデルを構築する新しいモデルを提案する (図 1 参照)。

はじめに、入力 $\mathbf{x} = (x_1, \dots, x_N)$ に対して式 (??) のような入力の線形変換を行う。

$$z_j = d_{j0} + d_{j1}x_1 + \dots + d_{jN}x_N \quad (4)$$

次に、変数 $\mathbf{z} = \{z_1, \dots, z_M\}$ を入力とした SIRMs システムの構築を行う。ここで、変換行列 $\mathbf{D} = (d_{ij}) (i \in Z_N, j \in Z_M)$ 、 M はシステムの近似精度により決定する定数である。

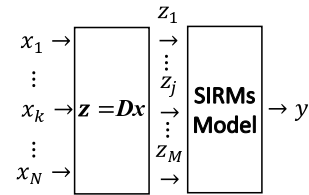


図 1: システムの概略

推論の出力 y は、 $\mathbf{z}$ を入力とした式 (3) により行われる。システムの出力を y 、理想の出力値を y^0 とする。システム中のパラメータ $\alpha \in \{a_{ji}, b_{ji}, c_{ji}, h_j\}$ は、以下で示す最急降下法を用いた学習により決定する [2]。ここに、 K_α は学習係数、 t は学習のステップを表す。

$$\alpha(t+1) = \alpha(t) + K_\alpha (y^0 - y) \frac{\partial y}{\partial \alpha} \Big|_{\alpha=\alpha(t)} \quad (5)$$

4 障害物回避問題への応用

ここでは、システムの有効性を示すために、図2に示すように障害物 (Obstacle) を回避しつつ目的地 (designated point) に到達するようなオブジェクト (Object) の行動を決定するファジィシステムを構築する。入力としてオブジェクトと障害物の距離 r_1 となす角 θ_1 および物体と目的地の距離 r_2 となす各 θ_2 の4入力を用い、出力としてオブジェクトの速度 $\mathbf{A}$ の成分 A_y を得る。なお、オブジェクトの速度 $\mathbf{A}$ の成分 A_x は一定である。図3に示す200個の点を、学習用データとして

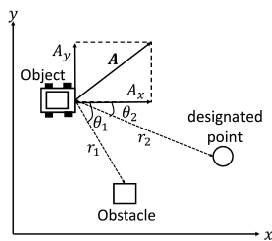


図2: 障害物回避問題

用い、提案法によりこの軌跡を実現する。

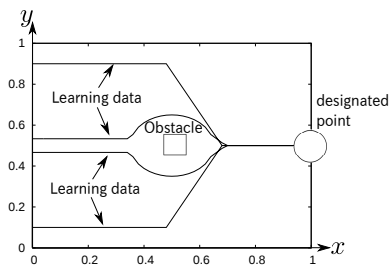


図3: 学習用データとその軌跡

学習の結果、図3に示すような軌跡を実現するファジィシステムが構築できた。このシステムを用い、未知の環境である次の実験を行った。

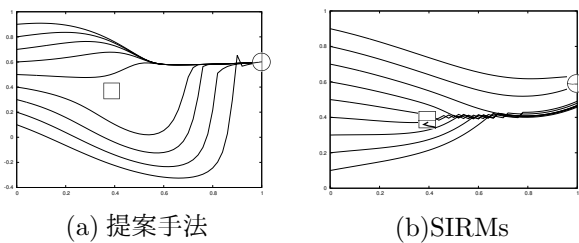


図4: 学習により得られたオブジェクトの軌跡

(1) スタート地点をさまざまに変えた場合のSIRMsと提案モデルの軌跡を図4に示す。通常のSIRMsでは目的地に到達しない場合がある。一方で、提案手法ではすべての場合において障害物を回避し目的地に到達している。

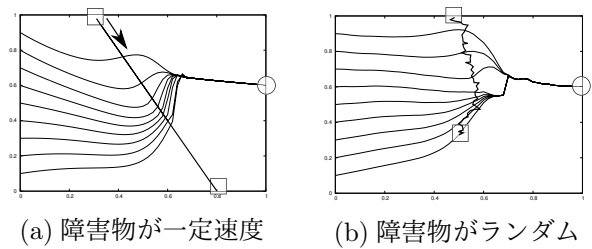


図5: 障害物の位置が時間とともに移動する実験

(2) 障害物の位置が変動する場合に提案手法を用いたとき、すべてのオブジェクトが障害物を回避し目的地に到達している。(図5参照)

最後に、なぜこのような汎化性が実現できたかをみるために、学習により得られたファジィルールの意味を考える。

入力に対する条件を“(r_1 が大きい) or (θ_1 が大きい) or (r_2 が小さい) or (θ_2 が大きい)”とすると、

- 1) 条件がほとんど成立しない → オブジェクトは y 軸正方向へ移動する。
- 2) 条件がある程度成立する → オブジェクトは y 軸負方向へ移動する。
- 3) 条件がかなり成立する → オブジェクトは y 軸負方向へ移動する。

これらは、人間の行動を言葉によるルールとして表現したものとなっている。

5 まとめ

本稿では、入力変数を線形変換した後にSIRMsモデルを構築する新しいファジィ推論モデルを提案し、制御問題の1つである障害物回避問題に適用し、その有効性を示した。

参考文献

- [1] N. Yubazaki, et al., "SIRMs(Single Input Rule Modules) Connected Fuzzy Inference Model", J. Advanced Computational Intelligence, 1, 1, pp.23-30, 1997
- [2] H. Miyajima, N. Shigei, H. Miyajima, A Proposal of SIRMs Model with Linear Transformation of Input Variables, pp.311-316, SCIS and ISIS 2014, 2014.

連絡先

鹿児島大学大学院理工学研究科
E-mail: k3768085@kadai.ac.jp

RoboCup に対する SIRMs を用いた位置予測モデルの開発

Developing a Position Prediction Model Based on SIRMs for RoboCup

○山下 雄大¹, 三船 哲史¹, 中島 智晴¹, 秋山 英久²
○Katsuhiko Yamashita¹, Satoshi Mifune¹, Tomoharu Nakashima¹, Hidehisa Akiyama²

¹大阪府立大学

¹Osaka Prefecture University

²福岡大学

²Fukuoka University

Abstract: In this paper, a prediction model is developed that estimates the position of RoboCup soccer players using Single-Input Rule-Modules (SIRMs) fuzzy systems. The purpose of the paper is to propose a prediction model that estimates the position of players mainly from the ball information. SIRMs fuzzy models are used to predict the position of opponent players. SIRMs fuzzy models reduce the number of fuzzy rules and parameters compared to common fuzzy systems even if the dimensionality of the problem is high. In the proposed method, training data are generated from the field information such as the position and the velocity of the ball, the position of the nearest opponent, and the expected cycles that the ball is intercepted by the opponent team. The SIRMs fuzzy models are trained after the first half using the data of the first half. Players use the trained SIRMs fuzzy models to predict the positions of opponent's players. A series of computational experiments are conducted to examine the performance of the proposed method.

1. はじめに

ロボット工学と人工知能の領域横断型研究プロジェクトとして RoboCup [1] が知られている。RoboCup には様々なリーグが存在しており、それぞれにおいて活発な研究、開発が行われている。RoboCup サッカーでは、競技で勝利することが重要視されているが、ただ単に勝利するだけではなく、ロボット工学や人工知能の発展につながるような技術を用いて勝利することが望まれている。特に人工知能的側面から考えると、チームでの試合運びを行うための戦術を知的に組み立てることが必要である。

戦術を組み立てる上で、プレイヤーが的確に意思決定を下すことは重要な要素である。より効果的な意思決定を下すためには、フィールド状態を正確に把握し、将来のフィールド状態を高い精度で予測する必要がある。しかし、プレイヤーはフィールド全体の情報を持つことができず、制限された情報から意思決定しなければならない。特に敵プレイヤーの位置を正確に予測することは困難である。

そこで本論文では、RoboCup サッカーシミュレーション 2D リーグを題材とし、プレイヤーの位置を予測する手法を提案する。位置予測をサッカーフィールドの状況からプレイヤー位置へのマッピングと定義する。位置予測の対象は敵チームのプレイヤーとする。提案手法では、位置予測に SIRMs ファジィシステ

ム[2] を用いる。数値実験では、SIRMs ファジィシステムを既存手法と比較し、性能を調査する。また、幾つかのチームに提案手法を組み込み、その予測精度を調査する。

2. RoboCup

RoboCup は、ロボット工学と人工知能の発展を目的とした、自律移動型ロボットによるサッカーなどを題材とした研究プロジェクトである。本論文では、RoboCup サッカーシミュレーションを研究の対象とする。サッカーシミュレーションには 2D リーグと 3D リーグがあり、本論文で使用する 2D リーグでは、プレイヤーとボールは二次元の仮想フィールド上における円で表現される。また、プレイヤーの他にもボールの位置と速度は全て二次元のベクトルとして表される。試合は 1 サイクル 0.1 秒とし、前後半 3000 サイクルずつ合計 6000 サイクルで試合が行われる。各プレイヤーはそれぞれ独立したエージェントとしてプログラムされており、制限された視覚情報や聴覚情報からドリブルやパス等の意思決定を行う。

3. 行動プランニング

プレイヤーが意思決定を的確にできるかどうかは、チーム戦略における重要な要素である。本論文で使用するプレイヤーは、行動連鎖と呼ばれる木構造を最

良優先探索することにより次の行動の意思決定を行う。行動連鎖を生成することで、プレイヤーの数手先の状態を考慮して、より戦略性の高い行動を選択することが可能になる。プランニング通りに試合を進めるためには予測状態の精度向上が必要である。その中でも、各プレイヤーの位置を正確に予測することが重要である。しかし、プレイヤーはフィールド全体の情報を持つことができず、制限された情報から意思決定しなければならないため、特に敵プレイヤーの位置を正確に予測する事は困難である。

オープンソースであるサンプルチームの `agent2d` [3]では、行動連鎖生成時の敵プレイヤーの予測位置は、最後に自分が認識した敵プレイヤーの位置としている。しかし、実際には敵プレイヤーはサイクル経過と共に動くため、予測状態での敵プレイヤー位置と、実際に行動を実行した後での敵プレイヤーの位置に大きな差異が生じてしまう。そのため、現状では生成された行動連鎖の質は低い。そこで本論文では、SIRMs ファジィシステムを用いて敵プレイヤーの位置を予測し、予測状態での敵プレイヤーの位置を考慮した行動をする手法を提案する。

4. 提案手法

4.1 SIRMs ファジィシステム

ファジィ推論システムの一つとして単一入力ルール群 (Single Input Rule Modules : SIRMs) ファジィシステム[2] がある。SIRMs ファジィシステムの概形を図1に示す。

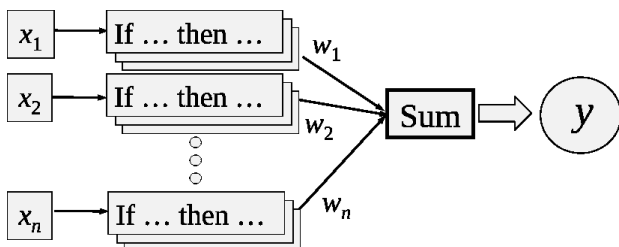


図1 SIRMs ファジィシステムの概形

各入力項目に対してひとつのルール群が用意される。ルール群は対象とする入力をファジィ分割することで得られるファジィ If-Then ルールから構成される。ルール群に含まれるファジィ If-Then ルールには対応する入力項目だけが前件部変数に含まれる。各ル

ール群には重視度と呼ばれる実数重みが割り当てられている。ルール群のファジィ推論結果の重視度付き総和を最終出力とする。SIRMs ファジィシステムは入力項目数が多い場合に、従来のファジィ推論モデルよりも大幅にファジィルール数とパラメータ数を削減することが可能であり、様々な制御問題へ応用されている。

4.2 位置予測手法

試合中における通常プレイが行われている状態を `playon` と呼ぶ。敵プレイヤーの位置予測を行う SIRMs ファジィシステムを、 x 座標と y 座標それぞれに1つ、すなわち1プレイヤーにつき2つ用意し、`playon` 時の敵プレイヤーの位置を学習させる。SIRMs ファジィシステムの入力として、

- 現在のボール位置 (x, y)
- 最も近い敵プレイヤーによるボールの予想インターセプト位置 (x, y)
- 敵プレイヤーがボールの位置に到達するまでに要すると予測されるサイクル数 (*cycle*)

を用いる。各入力データは $[0.0, 1.0]$ に正規化された後に SIRMs ファジィシステムに入力される。SIRMs ファジィシステムを学習させるには正確なフィールド情報が必要である。プレイヤーにはセンサ情報の制限があるため、試合中に正確な情報を持つ事は不可能である。一方、コーチエージェントは常に完全なフィールド情報を取得することができる。そこで、敵位置予測モデルの学習はコーチエージェントが行うものとする。コーチエージェントはハーフタイム中に各敵プレイヤーに対して SIRMs ファジィシステムを学習させる。学習済みの SIRMs ファジィシステムに予測状態時点での各敵プレイヤーの位置を出力させ、敵プレイヤーのフィールド上での位置を予測する。学習した SIRMs ファジィシステムの情報は、コーチから各プレイヤーに送られ、プレイヤーはその受け取った情報から SIRMs ファジィシステムを構築し、行動連鎖生成時の予測状態での敵プレイヤーの位置推定に使用する。これにより、より正確な予測状態を作成でき、戦術的価値のより高い行動連鎖を作成することが可能である。

5. 数値実験

数値実験では、提案手法とニューラルネットワークを用いた手法での位置予測の精度の比較、提案手法での入力パラメータによる精度の検証、提案手法を最新のチームに組み込んだ場合の性能の調査を行う。

5.1 従来手法との比較

まず、本節では提案手法である SIRM_s ファジィシステムによる位置予測と、ニューラルネットワークによる位置予測の精度を比較する。それぞれの手法を HELIOS2012 に組み込み、ベースチームとして公開されている agent2d [3]、RoboCup2012 に出場したチームに対してそれぞれ 10 試合行う。

提案手法を組み込んだ HELIOS2012 とニューラルネットワークを組み込んだ HELIOS2012 それぞれで試合を行い、予測精度を評価した。計測された誤差のうち、ゴールキーパーを除いたフィールドプレイヤー 10 人分の誤差の平均と標準偏差を表 1 に示す。なお、誤差の平均の単位はメートルである。以降の実験では同様の基準で評価を行う。

表 1 SIRM_s とニューラルネットワークの予測精度

対戦相手	SIRM _s	Neural Network
agent2d	7.80±6.56	10.35±7.89
WrightEagle	9.89±7.34	11.88±8.04
MarliK	9.92±7.26	11.07±9.01
Ri-one	8.80±6.51	11.26±6.84
Oxsy	9.63±6.21	13.62±9.21
YuShan	9.90±7.08	13.57±9.30
NADCO-2D	8.23±5.58	10.85±8.00

表 1 より、SIRM_s の誤差とニューラルネットワークの誤差を比較すると、どのチームに対しても SIRM_s ファジィシステムの誤差の値はニューラルネットワークよりも小さくなっていることが分かった。このことから、SIRM_s を用いた提案手法はニューラルネットワークによる手法と比べて予測の精度に改善がみられることがわかった。

5.2 入力属性の検証

次に、提案手法に用いる入力属性に関する検証を

行う。5.1 節で使用した入力属性以外に敵プレイヤーの位置予測をする上でより適切な入力属性のセットが存在するかどうかを調べるのが目的である。学習に用いる入力属性セットの形式として、以下の 3 種類を調べる。

- 2 入力
 - 現在のボールの位置 (x, y)
- 4 入力
 - 現在のボールの位置 (x, y)
 - 現在のボールの速度成分 (x, y)
- 5 入力
 - 現在のボール位置 (x, y)
 - 最も近い敵プレイヤーによるボールの予想インターセプト位置 (x, y)
 - 敵プレイヤーがボールの位置に到達するまでに要すると予測されるサイクル数 (*cycle*)

上記のうち、5 入力 が 5.1 節の数値実験に使用した入力属性セットである。その他の 2 つの入力属性セットが本節で追加されたものである。本研究室で開発している opuSCOM2014 に組み込み、agent2d、RoboCup2014 に出場したチームに対してそれぞれ 10 試合行う。

提案手法を組み込んだ opuSCOM2014 で試合を行い、入力属性として用いる情報による予測精度の違いについて検証を行った。全チームでの誤差の平均と標準偏差を表 2 に示す。

表 2 入力セットごとの SIRM_s での予測誤差

入力セット	平均	標準偏差
2 入力	12.87	10.41
4 入力	12.34	10.55
5 入力	10.14	8.29

表 2 から、基本的には 5 入力のセットが他 2 つのパターンよりも予測誤差が小さくなっていることが分かる。提案手法においては 5 入力のパターンが比較的安定して予測を行うことができていると言える。

5.3 性能の調査と結果の補正

次に、提案手法を最新のチームに組み込んだ場合

の位置予測の精度を調査する．学習後の出力で得られる座標の予測値が，敵プレイヤーの現在の座標から与えられたサイクル数以内に到達できない座標に位置している可能性がある．そこで，明らかに到達不可能と思われる位置を予測した場合には予測値を補正することを考える．敵プレイヤーが決められたサイクル数内に到達できるかどうかを判定し，到達不可能であれば，現在の座標から予測座標に向かうという仮定の下で，その方向へ入力の実予想サイクル数で到達することができる座標に出力結果を修正する．この補正する場合と補正しない場合の両者を検証する．HELIOS2014 に提案手法を組み込み，agent2d, RoboCup2014 に出場したチームに対してそれぞれ 10 試合行う．

提案手法を組み込んだ HELIOS2014 で試合を行い，補正の有無による予測精度の違いを調査するための数値実験を行った．数値実験の結果得られた誤差の平均と標準偏差を表 3 に示す．

表 3 補正の有無による予測精度の誤差

対戦相手	補正なし	補正あり
agent2d	7.19±5.36	5.43±5.81
WrightEagle	9.64±6.52	5.30±5.64
YuShan	10.82±8.46	5.71±6.02
AUTParsian	7.83±5.34	5.05±4.60
CYRUS	7.22±5.14	5.09±4.57
Enigma	7.89±5.77	5.67±5.32
FCP_GPR	7.44±5.86	4.97±4.53
Gliders	9.51±9.23	5.59±5.26
HERMES	14.14±11.23	6.11±7.55
HfutEngine	9.48±8.56	5.22±5.70
InfoGraphics	4.64±4.04	3.93±3.51
Oxxy	9.12±6.14	5.31±5.19
UFSJ2D	6.11±4.59	5.25±4.25
Ri-one	8.86±6.19	5.38±5.35
tokA1	8.27±6.30	4.96±5.25

表 3 から，最新のチームに提案手法を導入した場合は，特殊な場合を除き，補正を行わなくても 7～10m 程度の誤差の精度で位置予測を行うことができることが分かる．また，提案手法で得られた出力に

先述の補正をかけた場合は，5m 程度の誤差の精度で位置予測を行うことができていることがわかる．このことより，今回の決められたサイクル数内に予測位置に到達可能であるかどうかを考慮して出力結果を補正する効果があると言える．

6. おわりに

本論文では，SIRMs ファジィシステムに基づくプレイヤーの位置予測モデルを提案した．探索木に基づいた行動プランニングにおいて，プレイヤーができるだけ正確なフィールド情報に基づいて意思決定できることを目的とした．予測状態の精度向上のため，提案手法 SIRMs ファジィシステムを予測状態における敵プレイヤーの位置の予測に使用した．既存手法であるニューラルネットワークでの位置予測と比較し，提案手法である SIRMs ファジィシステムの優位性を示した．また，実データから学習することで敵チームの位置予測が可能であることを示した．

今後の課題として，予測状態における位置予測以外での SIRMs ファジィシステムによる位置予測モデルの応用や，予測誤差の傾向を敵チームが導入している戦術の分析に活用するための調査などが挙げられる．

参考文献

- [1] Hiroaki Kitano, Minoru Asada, Yasuo Kuniyoshi, Itsuki Noda, Eiichi Osawa and Hitoshi Matsubara, “RoboCup: A Challenge Problem for AI”, *AI Magazine*, Vol.18, No.1, pp.73-85, 1997.
- [2] 湯場崎直義, 易建強, 廣田薫, “複数入力ファジィ制御のための単一入力ルール群結合型ファジィ推論モデルの提案” *日本ファジィ学会誌*, Vol.9, No.5, pp.699-709, 1997.
- [3] 秋山英久, *RoboCup サッカーシミュレーション 2D 必勝ガイド*, 秀和システム, 2006.

連絡先

中島 智晴

大阪府立大学大学院 工学研究科

E-mail: tomoharu.nakashima@kis.osakafu-u.ac.jp

RoboCup における セットプレイのマーク割当に関する検討

A Study on Marking Assignment in Set-Plays for RoboCup

¹ 田中 翔, ¹ 山下 雄大, ¹ 三船 哲史
¹Sho Tanaka, ¹Katsuhiro Yamashita, ¹Satoshi Mifune
○ ¹ 中島 智晴, ² 秋山 英久
○ ¹Tomoharu Nakashima, ²Hidehisa Akiyama

¹ 大阪府立大学

¹Osaka Prefecture University

² 福岡大学

²Fukuoka University

Abstract: A marking assignment method in set-plays of RoboCup is proposed in this paper. The marking assignment is formularized as a transportation problem where the effort for marking is considered as a cost, which is to be minimized. The transportation problem is solved by first constructing a cost matrix, and then converting it using Hungarian method. Different cost functions and assignment methods are considered and their performance is examined through computational experiments.

1 はじめに

本論文では, RoboCup サッカーシミュレーション 2D リーグにおける, 相手側セットプレイ時のマーク割当手法を提案する. 現在開発を進めているチームでは, 開発者が手作業で調整を繰り返してセットプレイでの味方の位置を決定している. 相手チームとの相性が良くなければ味方の位置が適切でない可能性があるため, 試合中に相手チームに対応した味方選手の配置をする必要がある. そこで, 割当問題を解き, その結果を用いてマークを割り当てることを考える. 数値実験では, マーク割当手法の性能を調査する.

2 RoboCup

RoboCup は, ロボット工学と人工知能の発展を目的とした研究プロジェクトであり, 複数のリーグが存在する. 本論文では, RoboCup サッカーシミュレーション 2D リーグを研究の対象とする. サッカーシミュレーション 2D リーグでは, コンピュータ内に用意された二次元平面を仮想サッカーフィールドとし, 円形のエージェントをプレイヤーとしてサッカー競技を行う. 試合は前後半 3000 サイクルずつ合計 6000 サイクルからなる. 1 サイクルは 0.1 秒で離散化されている. プレイヤーは, それぞれ独立したエージェントとしてプログラムされており, 制限された視覚・聴覚情報から行動選択を行う.

3 セットプレイにおけるマーク割当問題

セットプレイとは, 一度ゲームを止めボールがセットされた状態で再開するプレイである. 本論文では, キックオフを除く敵チームがキックして始まるセットプレイのみ取り扱う. 現在開発しているチームにおいて, セットプレイでの味方選手の位置は過去の敵チームの動きをもとに手動で調整を行っている. そのため, 相手が戦術を変更した場合, 不利な状況に陥る可能性がある. そこで, 相手の戦術に合わせた配置を試合中に行う必要がある.

4 提案手法

提案手法では, マーク割当をオペレーションズリサーチにおける割当問題に定式化する. キーパを除く味方エージェントを $p_i (i = 1, 2, \dots, 10)$, 同じくキーパを除いた敵エージェントを $o_j (j = 1, 2, \dots, 10)$, 味方エージェント p_i が敵エージェント o_j をマークするのに必要なコストを c_{ij} とすると, コスト行列

$$C = \begin{pmatrix} c_{1,1} & \cdots & c_{1,10} \\ \cdots & c_{ij} & \cdots \\ c_{10,1} & \cdots & c_{10,10} \end{pmatrix} \quad (1)$$

を持つ割当問題を解くことでコスト最小のマーク割当を得ることが出来る. コストとして, 移動距離などが考えられる. 数値実験では, 4 つのコストを考える.

5 数値実験

5.1 実験設定

まず, HELIOS2014 と対戦相手 4 チーム (KU_BOST (2012), YuShan(2014), Oxsy(2014), Gliders(2014)) との試合中に発生したセットプレイ開始時の位置情報をチームごとに 100 種類用意する. この位置情報セットを使って数値実験を行い, 提案手法の性能を調べる.

マーク決定の際, 現状態のみに基づく決定と多数決による決定を考える. 現状態のみに基づく決定では, 4 章で示した割当法のコスト行列を現時点のみで計算し, その結果得られた割当を最終とするものである. 多数決では, 毎サイクル割当問題を解き過去の割当を考えた多数決によりマークの割当を考える. コストとマーク決定の組み合わせを表 1 に示す.

表 1: Combination of assignment and cost

Assignment	Marking assignment	Cost
1	Current state	1
2	Current state	2
3	Current state	3
4	Current state	4
5	Majority decision	1
6	Majority decision	2
7	Majority decision	3
8	Majority decision	4

Cost 1 : 現在の味方位置と現在の敵位置との距離

Cost 2 : ボールによる味方のフォーメーション位置と現在の敵位置との距離

Cost 3 : 現在の味方位置とボールから予測した敵位置との距離

Cost 4 : ボールによる味方のフォーメーション位置とボールから予測した敵位置との距離

ただし, 割当 4, 8 は全く同じであるため数値実験では割当 8 を実施せず, 割当 7 までを調査対象とする.

5.2 実験結果

用意した 100 種類の実験用データを, 表 1 の割当手法ごとに 20 回再現することで割当手法の性能を調査した. この際, 味方選手のマーク対象選手決定は毎サイクル行うものとする. 調査した性能指標は, 味方選手の移動距離, ボール奪取回数である.

セットプレイ開始時の味方選手の位置とボールが蹴られる際の味方選手の位置との距離の全味方選手の合計を味方選手の移動距離とする. 味方選手の移動距離の 1 セット当たりの平均を図 1 に示す. 図 1 より, 割当 3, 7 は, 移動距離が少ないことがわかる. これは, 敵選手の移動位置を予測し割当をしているためマークに無駄な移動が少ないことが要因であると考えられる.

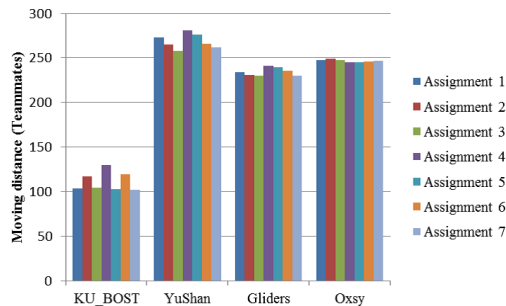


図 1: Moving distance (Teammates)

味方選手がボールを奪った 1 セット当たりの平均回数をチームごとに表したグラフを図 2 に示す.

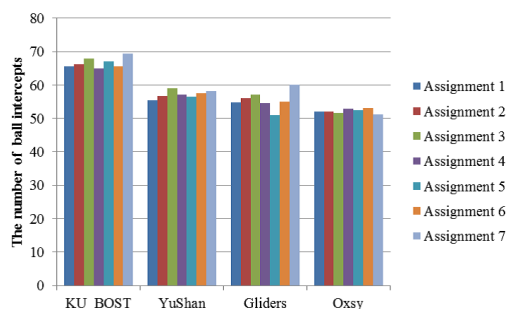


図 2: The number of ball intercepts

図 2 より, Oxsy 以外のチームでは, 割当 3, 7 の奪取回数が多いことがわかる. 割当 3, 7 は選手の移動が少ないことが要因であると考えられる.

6 おわりに

本論文では, セットプレイにおいてマーク割当手法を導入し, セットプレイの再現実験からコストによる性能変化を調査した. 今後の課題として, 他のコストの検討やプレイヤーの動きの洗練などが挙げられる.

参考文献

- [1] 秋山英久, RoboCup サッカーシミュレーション 2D 必勝ガイド, 秀和システム, 2006.
- [2] アーノルド・カウフマン, ロベール・フォール [著] 阿部統, 加藤茂夫 [訳] OR への招待, 好学社, 1972.

連絡先

田中 翔

大阪府立大学 工学部

E-mail: sho.tanaka@cs.osakafu-u.ac.jp

相関ルール生成における 1 つの問題点について

One problem for association rule generation in data mining

○ ¹ 山口 直人, ¹ 劉 臣席, ² 中田 典規, ¹ 酒井 浩
○ ¹Naoto Yamaguchi, ¹Chenxi Liu, ²Michinori Nakata, ¹Hiroshi Sakai

¹ 九州工業大学 大学院工学府

¹Graduate School of Engineering, Kyushu Institute of Technology

² 城西国際大学 経営情報学部

²Faculty of Management and Information Science, Josai International University

Abstract: This article considers a problem in association rule generation, which we faced with in the analysis of the actual table data sets. If we obtain an implication $\tau : \wedge_i [A_i, val_i] \Rightarrow [Dec, val]$ with higher criterion values, the criterion values for another implication $\tau' : \wedge_i [A_i, val_i] \Rightarrow [Dec, val']$ ($val \neq val'$) will be lower. Even though this seems trivial, if we are interested in the decision $[Dec, val']$ we do not obtain τ' in the standard data mining. In this paper, we consider this problem and describe our current solution.

Key words: Rough sets, Table data analysis, Major rule, Minor rule

1 はじめに

我々の研究室ではラフ非決定情報解析 (*RNIA*) と名付けた枠組みを提案し、非決定情報システム (*NIS*) におけるラフ集合理論の構築を通して、不完全なデータも対象とするデータマイニング手法の研究を進めている [3, 4]。第 24 回ソフトサイエンスワークショップにおいても現状を報告している [5]。

最近、構築中のデータマイニングツール *getRNIA* [6, 7] を用いて、実アンケートデータを解析した折、1 つの問題に直面した。本稿では、この問題に対する考察を述べる。

2 相関ルール生成と具体的問題

我々は、表データのタプルとして生じる含意式 $\tau : \wedge_i [A_i, val_i] \Rightarrow [Dec, val]$ の中で 2 つの基準、 $support(\tau) \geq \alpha$ かつ $accuracy(\tau) \geq \beta$ [1, 2] (α, β と決定属性 *Dec* はユーザが指定) を満たす含意式 τ をルールの候補として取り出す。*getRNIA* ではアプリオリアルゴリズム [1] の拡張版 NIS-アプリオリを用いている [6]。アンケートでは、2098 名がおおよそ 160 の項目 (属性とよぶ) に回答している。表 1 は決定属性 79 における回答者数の割合である。我々は属性 79 におけるよくできる人、全くできない人の相関ルールによる特徴付けを試み、次の問題を確認した。

(1) 属性値 1、4 とともに頻度が少ないために *support* の

表 1: 決定属性 79 における回答の比率

属性値	決定属性の意味	回答者の人数 (その割合)
1	よくできる	100 (4.7%)
2	まあできる	486 (23.1%)
3	あまりできない	1325 (62.9%)
4	全くできない	158 (7.5%)
5	無回答	39 (1.9%)

基準を弱める必要があり、 α の値を下げた。このことで、ルールの候補数が非常に増大した。

(2) 興味のない含意式、条件部 $\Rightarrow$ まあできる、条件部 $\Rightarrow$ あまりできない、のために、興味ある含意式、条件部 $\Rightarrow$ よくできる、条件部 $\Rightarrow$ 全くできない、の *accuracy* も低くなった。

(3) *support* を下げることは既に確認していたが [5]、*accuracy* も影響を受けることを確認した。

つまり、興味のない結論部が大多数を占める場合、通常の α と β の設定によるルール生成では、少数派の結論部を特徴付けることに難点が生じる。

3 上記問題への現状の取組

2 節で述べた内容を解決するための手法を現在検討している。興味ない結論部を持つ含意式 $\tau : \wedge_i [A_i, val_i] \Rightarrow [Dec, val]$ の *accuracy* が高い場合、その影響で τ' :

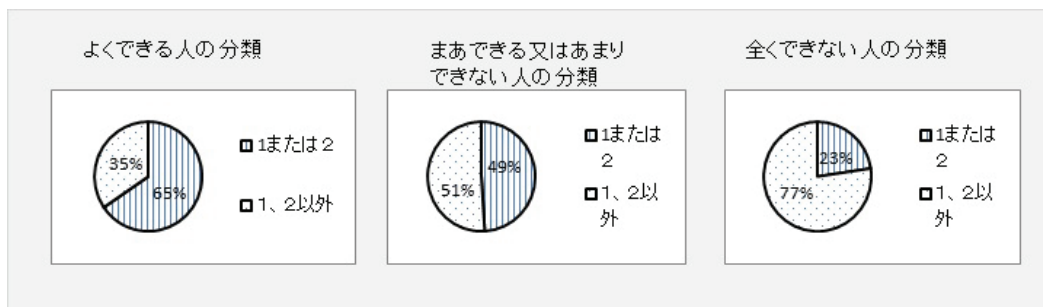


図 1: よくできる、まあできる又はあまりできない、全くできないの人を、属性 59（属性値、1 から 7）の 1 または 2 を選択した人、それ以外を選択した人で分類した円グラフ

$\wedge_i[A_i, val_i] \Rightarrow [Dec, val']$ ($val \neq val'$) の *accuracy* は必然的に下がる。もともと、 τ' がルールの候補にならない場合も考えられるが、 τ' が τ の陰に隠れてしまい取出しにくくなる場合も考えられる。現在、下記 2 つの手法により対応している。

(1) 手法 1 : 興味のない属性値のタプルを除去し、表を加工した後、通常のデータマイニング手法を適用する。表 1 において、興味ある属性値、1 と 4 を持つ 258 個のタプルを取出し、この 258 個の表データに対して、*getRNIA* を用いて解析した。この場合、いくつかの含意式を取り出すことができた。

(2) 手法 2 : 興味ある 1 属性値のタプルのみ扱う。表 1 において、属性値 1 を持つ 100 個のタプルから成る表データを作成、属性値 4 を持つ 158 個のタプルから成る表データを作成し、*support* 値のみで含意式の評価をおこなった。この場合、自動的に *accuracy* は 1 になる。*accuracy*=1 の設定でプログラムを動かすことはシステムに異常な負荷がおよびメモリエラー等が多々発生するために、*accuracy*=1 の設定は通常、実行が難しい。

手法 1 では、何故 256 個に限定したか、また 256 個のタプルに対する *support* と *accuracy* の解釈も明快でない。一方、手法 2 では、同一の結論を有するタプルの頻度で含意式の優位性を考えるために直観的でわかり易いと考ええる。

図 1 は、よくできる、まあできる又はあまりできない、全くできないの分類について属性 59 の属性値（1 から 7）の割合を手法 2 により実行した結果である。この例から、以下を得ることができると考えられる。

(1) よくできる人の 65% は属性 59 の回答で 1 または 2 を選択している。

(2) 全くできない人の 77% は属性 59 の回答で 1、2 以外を選択している。

(3) 属性 59 の回答により、よくできる人、全くできない人の判定を行える可能性が高い、しかしまあできる人とあまりできない人を特徴付けることは難しい。

4 まとめ

通常のデータマイニングでは処理しにくいルール（マイナールール）生成について、その問題点を示した。現状では、解釈の面から考えて手法 2 の方がより自然な手法に思える。しかし、両手法とも元データに加工を加えており、望ましくない。アプリオリや NIS-アプリオリに上記の性質を持つルールの抽出機能の追加ができれば望ましいと考える。

参考文献

- [1] Agrawal, R., Srikant, R.: Fast algorithms for mining association rules in large databases, VLDB'94, pp. 487–499, 1994.]
- [2] Pawlak, Z.: Rough Sets, Kluwer Academic Publishers, 1991.
- [3] 山口直人, 伍卯, 中田典規, 酒井浩: アンケートデータにおけるラフ集合情報解析の応用, 第 29 回ファジイシステムシンポジウム, MA1-2, 2013.
- [4] N. Yamaguchi, M. Wu, M. Nakata, H. Sakai: Application of rough set-based information analysis to questionnaire data, JACIII, Vol.18, No.6, pp.953-961, 2014.
- [5] 山口直人, 伍卯, 中田典規, 酒井浩: ソフトウェア *getRNIA* によるアンケートデータ分析, 第 24 回ソフトサイエンスワークショップ, pp.121–122, 2014.
- [6] M. Wu, M. Nakata, H. Sakai: An overview of the *getRNIA*-A system for non-deterministic data, *Procedia Computer Science*, Elsevier, Vol. 22, pp. 615-62, 2013.
- [7] M. Wu, H. Sakai: RNIA WEB software tool: <http://getrnica.org/>

連絡先

山口 直人

九州工業大学大学院 工学府

E-mail: kityn1124@gmail.com

ラフ集合における Generalized Dynamic Reduct の並列抽出

Parallel Extraction of Generalized Dynamic Reducts in Rough Set

○¹ 高橋 智, ¹ 工藤 康生, ² 村井 哲也
○¹ Satoshi Takahashi, ¹ Yasuo Kudo ² Tetsuya Murai
¹ 室蘭工業大学 ² 北海道大学
¹ Muroran Institute of Technology ² Hokkaido University

Abstract: In this paper, we discuss a parallelization of extracting generalized dynamic reducts (GDRs) in rough set. We have proposed a method for extracting GDRs from large-scale data, however, it takes high computational costs. We implement a parallel extraction method for GDRs and evaluate the method by experiments.

1 はじめに

ラフ集合 [6] を大規模なデータに適用する取り組みとして、確率的選択を用いた縮約手法 [2] や統計的手法に基づく縮約抽出 [1] が提案されている。また、著者らは両手法を組み合わせる縮約抽出を行う手法を提案している [3]。しかし、この手法では計算過程で膨大な計算時間を必要とする。本研究では、計算過程を並列化し高速に縮約を抽出することを目的とする。

2 背景

2.1 決定表と相対縮約

ラフ集合で扱うデータは一般的に、 $A = (U, C \cup D, V, \rho)$ で定義される決定表 A で表される [4]。ここで、 U は対象の空でない有限集合、 C は条件属性の空でない有限集合、 D は決定属性の空でない有限集合であり、 $C \cap D = \emptyset$ とする。すべての属性の集合を $AT \stackrel{\text{def}}{=} C \cup D$ と表す。 V は各属性 $a \in AT$ の値の集合、 $\rho: U \times AT \rightarrow V$ は対象 x の属性 a での値 $\rho(x, a) \in V$ を表す関数である。

相対縮約とは分類可能なすべてのデータを正確に分類するための必要最小限の条件属性の集合である。

2.2 確率的選択を用いた縮約手法

条件属性の確率的選択を用いて、大規模データにおける条件属性による対象の分類能力を保存しつつ、できるだけ条件属性の個数が少ない小規模決定表 [2] を生成する。この小規模決定表 $A' = (U, C' \cup D, V, \rho) (C' \subseteq C)$ の相対縮約は、元の決定表の相対縮約になることが保証されているため、多数の小規模決定表を生成し相対縮約を抽出することにより、元の決定表の相対縮約を多数抽出する。

2.3 Generalized dynamic reduct

決定表 A に対して、対象の集合 U をその任意の部分集合 $U' \subseteq U$ に制限した決定表 $B = (U', C \cup D, V, \rho)$ を A の部分表と呼び、 A のすべての部分表の集合を $P(A)$ 、部分表の任意の集合を $F \subseteq P(A)$ と表す。すべての部分表 $B \in F$ に対して相対縮約を求め、 F 中の $100 \times (1 - \epsilon)\%$ 以上 ($0 \leq \epsilon < 1$) の部分表で相対縮約となる属性集合 $R \subseteq C$ を、決定表 A の (F, ϵ) -generalized dynamic reduct (GDR) として抽出する [1]。しかし、大規模な決定表から GDR を求める場合、条件属性が多いと部分表 B のすべての相対縮約を求めるのは困難である。

2.4 大規模データからの GDR の抽出

対象の個数および条件属性の個数が多い決定表から GDR をできるだけ多く抽出するために、2.3 節で述べた手法に 2.2 節で述べた手法を組み合わせる抽出を行う [3]。決定表 A から、対象の集合 U を任意に制限した部分表 B を多数生成し、各部分表 B に対して小規模決定表 B' を多数生成し、各小規模決定表 B' から相対縮約を抽出する。その後得られたすべての相対縮約について、各部分表で相対縮約の条件を満たすかを調べる。生成した部分表の中で $100 \times (1 - \epsilon)\%$ 以上の割合で縮約となるものを A の (F, ϵ) -GDR として抽出する。

3 並列化による GDR 抽出の高速化

2.4 節で述べた手法では、GDR 抽出の際、小規模決定表から得られた相対縮約が GDR の条件を満たすかを確認する部分が計算時間のボトルネックとなっている。本研究では、GDR の抽出と確認の部分を並列化し高速に GDR を得ることを試みる。

3.1 実験

2.4 節で述べた手法の並列化を実装し実験を行った。実験環境は、CPU: Inter(R) Xeon(R)CPU E5-2650 v2 @ 2.60GHz(16cores), メモリ: 132GB, OS: CentOS release 6.5, Compiler: gcc 4.4.7 である。並列化には OpenMP [5] を使用した。

各種大規模データに対して、計算時間、GDR の個数、GDR による対象の分類精度を計測する。この計測を 5 回繰り返し、その平均をそれぞれのデータと比較する。なお、GDR による対象の分類精度の評価値として分割による近似の質 [4] を用いる。

使用したデータセットを表 1 に示す。データセットは UCI Machine Learning Repository [7] から取得した。

表 1: 使用したデータセット

データセット名	対象数	属性数
pure.spectra.matrix(PSM)	571	1230
GISETTE	3252	4999
letter-recognition(LR)	20000	16
TomsHardware(TH)	28178	97
Nomao	34465	119
adult	48842	14
diabetes	101766	49
tao-all	178080	11

3.2 実験結果

実験結果を表 2 に示す。なお、tao-all データセットと diabetes データセットに関しては、対象数が多く GDR の精度を計算することが出来なかった。

表 2: 実験結果

	計算時間	GDR の個数	GDR の精度
PSM	808 秒	4125.2 個	0.999
GISETTE	10891 秒	28658.6 個	0.998
LR	1028 秒	208 個	0.837
TH	362.6 秒	47.2 個	0.922
Nomao	3746 秒	498 個	0.938
adult	100 秒	20.6 個	0.898
diabetes	10226 秒	134 個	-
tao-all	759 秒	37.4 個	-

3.3 考察

表 2 の結果から、「対象数×属性数」と計算時間とのスピアマンの順位相関係数を求めると 0.67 となり、中

程度の正の相関が見られた。これは、データセットの規模が大きくなることに応じて計算コストも増加するためと考えられる。また、GDR の個数と属性数との順位相関係数は 0.88、GDR の精度と属性数との順位相関係数は 0.89 となり、どちらも強い正の相関が見られた。これは、属性数が多いほど GDR になりうる属性の組み合わせも増加し、正確に分類できる対象の個数も増加するためと考えられる。

4 まとめ

本研究では、GDR の抽出過程の全体を並列化することで、GDR の抽出を高速化することを試みた。並列化を実装し実験を行った結果、著者らの従来研究 [3] では扱えなかったデータ数 10 万件以上の規模のデータセットに対しても、約 3 時間程度で GDR を抽出することができた。また、平均的にデータ全体に対する高い分類能力を持つ GDR が得られた。

今後の課題として、実装した並列化を用いた、更に規模の大きなデータセットに対する GDR 抽出や実データの分析への応用などが挙げられる。

謝辞

本研究の一部は科研費基盤研究 (C) (25330315) の助成を受けたものである。

参考文献

- [1] J. G. Bazan: Dynamic Reducts and Statistical Inference, *Proc. of IPMU'96*, pp.1-5, 1996.
- [2] 工藤 康生, 丸山 太一, 村井 哲也: 確率的選択を用いた大規模データからの縮約抽出アルゴリズムについて, 第 26 回ファジィシステムシンポジウム講演論文集, pp.764-767, 2010.
- [3] Yasuo Kudo and Tetsuya Murai: An Attempt of Hybridization of Generalized Dynamic Reducts and A Heuristic Attribute Reduction Using Reduced Decision Tables, *Proc. of FUZZ-IEEE 2013*, pp.1-6, 2013.
- [4] 森 典彦, 田中 秀雄, 井上 勝雄編: ラフ集合と感性-データからの知識獲得と推論, 海文堂出版, 2004.
- [5] OpenMP: <http://openmp.org/wp/>
- [6] Z. Pawlak: *Rough Sets: Theoretical Aspects of Reasoning about Data*, Kluwer Academic Publisher, 1991.
- [7] UCI Machine Learning Repository: <http://archive.ics.uci.edu/ml/index.html>

連絡先

工藤 康生

〒050-8585 北海道室蘭市水元町 27-1

室蘭工業大学大学院しくみ情報系領域

Tel: 0143-46-5469, Fax: 0143-46-5499

E-mail: kudo@csse.muroran-it.ac.jp

ファジィ論理のほとんど全て (28 — 最終回)
— 今後解決すべき課題 —
— 含意, 評価, 評価のファジィ推論法 —
Almost All about Fuzzy Logics (Extra)

中島 信之
NAKAJIMA, Nobuyuki

1 はじめに

この原稿を書き始めた 1 月 31 日現在、『ファジィ論理のほとんど全て』はほぼ完成しています。ほんのあと数ページ — (これが問題!)。

だけど、この段階でまだよくわからないことがいくつも残っている。といって、それらを (わたしの力で) 解き明かすのを待っていたら、いつ「完成」するかわからない。とりあえずは、未解決問題を抱えたまま、印刷にまわすことにしよう。

未解決問題とは、①ファジィ含意、②評価、③評価のファジィ推論法の 3 つである。

2 ファジィ含意

まずはファジィ含意、あるいは (ファジィのつかない) 含意である。

わたしは論理学の専門家ではない。これまで論理学をきちんと学んだことはない。それなのに、大胆不敵にも『ファジィ論理のほとんど全て』を執筆してきた。最初は、含意についても漠然とわかったつもりでいた。古典 2 値論理における含意、多値論理における含意、ファジィ論理における含意 — と、別箇に展開したが、どうもうまいこといかん。

— 含意とは何か、がわからなくなったのである。

そこで、第 IV 部として含意をまとめて論じようと考えた。だが、結果はますます悲惨だった。含意とは何かはますます遠くなった。含意の満たすべき性質は何か、定義の根拠は何なのか、そういったことにはどこにも記述されていない。

古典 2 値論理の含意 $\rightarrow$ は、 $a \rightarrow b = \neg a \vee b$ または $\neg(a \wedge \neg b)$ で「定義」される。非古典論理でも、含意はあらかじめ定義されている。含意があたえられたとき、それか何が導かれるかが示されているだけ、それでいい。

じつのところ、含意が意識的に取りあげられるようになったのはファジィ論理においてが初めてだった。ファジィ論理自体が、誕生時 (というのがあれば、だが) には、まだかたちをなしておらず、したがって含意もなく、含意に何を求めるかについての合意もなかった。

含意って一体何だろうか。

いまわれわれに課せられているのは、含意とは何かを明らかにすること、含意を定義すること、あるいは定義の根拠を明らかにすること、あるいは他の演算と協力して推論を構築できるようにするか、だと思う。

3 評価

次は評価である。

評価をまず国語辞典を見た。言語的にはわかったが、われわれが考える数理的な評価についてはわからなかった。数理関係の事典にあたろうと思ったが、

現代百科事典編集委員会編、『現代百科事典 第 4 巻』、ぎょうせい、1981

OR 事典編集委員会編、『OR 事典』、日科技連、1975

ぐらいしかみつからなかったし、読んでもよくわからなかった¹。

これまで見た範囲でよかったのは『新明解国語辞典』(三省堂、2006) で、評価: ①物の価値や価格を (論じて) 決めること、②〔教育で〕児童・生徒の学習成果について測定すること。〔広義では、教育の全般についていう〕とあった。また『OR 事典』の「人の評価」の項には、「選別・選抜型の評価」と「育成型の評価」があげられており、前者は客観性や、社会的に認められた価値に立脚したものであることが要請される。それに対して後者は評価者と被評価者だけの閉じた関係

¹ファジィ学会には「評価問題研究会」なんてのがありますが、そもそも評価とは何か、という問題をあつかってくれますかねえ？

だから客観性や社会的な評価基準にもとづかなくてもよい、と述べられている。

評価の諸相

①項目別の評価と総合評価 1人の生徒の「成績」というとき、その成績とは総合成績のことである。それに対して箇々の数学や英語の成績は項目別の成績である。後者は、ペーパー試験や平常の授業での評価で定まる。それに対し前者は項目別の評価の組み合わせで定まる。組み合わせ方に評価者の思想が反映する。

②相対評価と絶対評価 特定の集団における相対的な位置を求めるのが相対評価である。絶対評価があたえられ、集団が定まると、相対評価を求めることができる。だが、逆は一般的には可能でない。

現在わかっているのは以上である。

4 評価のファジィ推論法

当初、第Ⅲ部として大々的に「ファジィ推論法」を取りあげようと思った。だが、モノが大きすぎ、時間的にも、量的にも、(そしてわたしの能力的にも) 手におえないことがわかった。いったんとりやめたものの、ファジィ論理の書としてまったくふれないのもどうかとも思った。そこで最近注目されはじめた評価の推論法を中心に取りあげることにした。

ファジィ推論法は制御のためのものというかたちで発達してきた。したがって、そのための蓄積は大きい。最近、ファジィ推論法を評価に応用しようという動きがある。具体的には、項目別の評価から総合評価を求めるのにもちいる。

評価のファジィ推論法は制御のファジィ推論法とは異なる。

①制御の場合は、(1) 目標値があり、(2) (たいてい) 1秒間に数回の推論をおこない、徐々に目標値に近づける(これを**接近原理**という)。1回1回の推論は大ざっぱでもかまわない。それに対し、評価では、推論のくりかえしはなく1回きりである。1回の推論できちんと答えを出さなければならない(これを**答え一発**という)。

②ある程度時間が経過すれば、制御値は目標値の近くにあるとみてよい。いいかえると、入力を入力空間のある場所で濃く他の場所で薄く……というように濃淡ができる。それに対し、評価の場合、一般に入力空間での入力に濃淡はない。

評価のファジィ推論法の設計

こうした違いは当然その設計法に違いを生む。①の答え一発については、推論の瑕瑾は許されない。制御

だと、推論設計にまずいところがあっても、2回目あるいはそれ以降の推論で修正が可能である。だが、評価では2回目はない。1回でちゃんとした答えを出さねばならない。……では、**ちゃんとした**答えとは何であるか。評価のファジィ推論法の満たすべき条件は何か。

われわれ(稲井田さん、高萩さんとわたし)はいまから8年前、メールを通じて、論じあったことがある。稲井田さんは、われわれの間に答えて、

(1) 単調性、(2) 連続性、(3) 境界条件

をあげた。はたしてこれで充分か。あるいはこれら全部は必要か。……こうした問いに対する答えはまだない。

②の入力の濃淡に関してはあまり問題はなかろうと思うが、これ以外に設計に要する評価のファジィ推論法の特徴はなにか。

5 おわりに

以上述べてきたように、『ファジィ論理のほとんど全て』を完成するためには、まだまだ解き明かさなければならない事柄がある。だが、ご勘弁願って、これまでまとめたもので世に問おうと思う。

[問い合わせ先]

中島 信之

930-0077 富山市磯部町 2-1-24

TEL & FAX:076-493-2104

ファジィ積分とファジィ推論

Fuzzy Integrals and Fuzzy Reasoning Methods

○ 高萩 栄一郎 (Eiichiro Takahagi)

専修大学商学部 (School of Commerce, Senshu University)

Abstract: There are close relation among fuzzy integral models and fuzzy reasoning methods. Fuzzy switching functions are essentially 3-valued logic and interpolate the 3-valued truth table. To interpolate a 2-valued truth table, we introduce Choquet integral model instead of Sugeno integral models that are max-min calculations. Fuzzy reasoning methods are also interpolate calculations among fuzzy rules using gravity calculations. The gravity calculations have properties similar to probability calculations. We show the Choquet integral type fuzzy reasoning method which use Choquet integral instead of gravity calculations.

Keywords: Choquet integral, Sugeno integral, Fuzzy Rules, Fuzzy Reasoning, Gravity Calculations

1 はじめに

ファジィ理論, 特にファジィ推論やファジィ制御が成功したのは, ファジィルール表にあるのではないかと考えている. ファジィルール表は, 入力値の状態を大きい, 中くらい, 小さいなどのいくつかの段階に分け, 各入力値についての各状態についての出力値 (制御量) を表の形でまとめて指定する. 入力値が各状態の中間の値であるときは, ファジィ推論によりうまく補間して出力値を決める. 表の形でルールが記載されているので, 分かりやすい. 何か変更しようとしたとき, ルール表を変更することにより対応できるなど扱いやすい. 表の形でルール表現は, 人間の思考形態 – 問題の整理のしかた – に近いのではないだろうか?

ファジィ推論での出力はファジィ集合であり, 重心計算などで, ルール表が単調 (入力値が大きくなれば, 出力値は必ず大きくなる関係) であっても入出力の関係で単調ではない場合があり [2], 単調性を満たす簡略化推論法が使われることが多い.

ファジィ積分も複数の入力値 ($[0, 1]$ 区間の値) から 1 つの出力値を計算するモデルである. ファジィ推論との対比で考えると, ルール表はファジィ測度に対応する. 計算方法は, max-min 演算 (菅野積分) や 減算, 積演算と和演算 (ショケ積分) など単純な演算で表現されているので扱いやすい. ファジィ測度は通常, 単調性が仮定されており, ファジィ測度が単調であれば, 入出力関係は単調である.

ファジィ積分は, 1 つの入力に対して, ファジィ推論のように入力値の状態を分けてルールを作ることはできない. ファジィ積分は, 入力値間の相互作用を扱

う. しかし, ファジィルール表でもこのようなルールを設定すれば, ある程度, 相互作用を表現できる.

本稿では, ファジィルール表による出力値の計算 (ファジィ推論) をファジィ積分, 特にショケ積分で行う方法を紹介し, 合わせて計算ツールを公開する.

2 ファジィ論理関数とファジィ積分

2.1 ファジィ論理関数

向殿 [3] は, ファジィ論理関数の本質は 3 値論理であると指摘した. 真, 偽と未知 (Unknown) の 3 値で, 1 が真, 0 が偽, 0.5 が未知に割り当てられる. $\wedge$ に min 演算, $\vee$ に max 演算, $\neg$ に $1 - x$ 演算が使われる.

ファジィ論理では排中律 $y = x \vee \neg x$ は成立しない. これは, 表 1 の 3 値の真理値表を補間するためである.

表 1: 排中律の真理値表

	x	y
(あ)	0	1
(い)	0.5	0.5
(う)	1	1

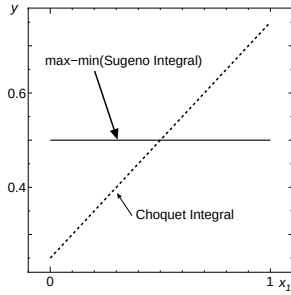
否定を用いない Fuzzy/C (定数係数を持ったファジィ論理関数) [4] でも, 2 値での入出力関係で狭義の単調性があっても, $[0, 1]$ のファジィ値の入出力間では狭義の単調性はない. 例えば,

$$y = (0.5x_1) \vee (0.5x_2) \vee (1x_1x_2) \quad (1)$$

では, 表 2 の真理値表となり, それを補間して, 図 1 の max-min になる.

表 2: 式 (1) の真理値表

	x_1	x_2	y		x_1	x_2	y
(か)	0	0	0	(さ)	0	0.5	0.5
(き)	0	1	0.5	(し)	0.5	0	0.5
(く)	1	0	0.5	(す)	0.5	0.5	0.5
(け)	1	1	1	(せ)	0.5	1	0.5
				(そ)	1	0.5	0.5

図 1: 式 (1), $x_2 = 0.5$

2.2 ファジィ積分

2.3 真理値表とファジィ測度

ファジィ測度 μ は集合関数と定義されるが, 真理値表として考えた方が理解しやすい. 例えば,

$$f(x_1, x_2, x_3) = 0.2x_1 \vee 0.3x_2 \vee 0.1x_3 \vee 0.6x_1x_2 \vee 0.7x_1x_3 \vee x_1x_2x_3 \quad (2)$$

のときのファジィ測度と真理値表は, 表 3 になる.

表 3: 式 (2) の真理値表とファジィ測度

A	x_1	x_2	x_3	$\mu(A) = f(x_1, x_2, x_3)$
$\emptyset$	0	0	0	0
$\{1\}$	1	0	0	0.2
$\{2\}$	0	1	0	0.3
$\{1, 2\}$	1	1	0	0.6
$\{3\}$	0	0	1	0.1
$\{1, 3\}$	1	0	1	0.7
$\{2, 3\}$	0	1	1	0.3
$\{1, 2, 3\}$	1	1	1	1

2.4 菅野積分

ファジィ積分の概念は, 菅野 [1] により提案された. 入力値 x_i を大きい順に並べ替えるような置換 σ を求

める. 菅野積分 f_μ^S は, 次式で定義され (式 (3)), ファジィ論理関数でも計算できる (式 (4)).

$$f_\mu^S(x_1, \dots, x_n) \equiv \bigvee_{i=1, \dots, n} [\mu(\{\sigma(1), \dots, \sigma(i)\}) \wedge x_{\sigma(i)}] \quad (3)$$

$$= \bigvee_{A \subseteq X} [\mu(A) \wedge (\bigwedge_{i \in A} x_i)] \quad (4)$$

菅野積分は Fuzzy/C で表現でき, 単調性などの条件を満たす Fuzzy/C は菅野積分で表現できる.

2.5 ショケ積分

もう一つのメジャーなファジィ積分として, ショケ積分 [5] がある. 菅野積分と同様の方法で, ファジィ論理関数をショケ積分で計算する.

表 3 のファジィ測度で, $x_1 = 0.5$, $x_2 = 0.7$, $x_3 = 0.2$ の場合, 表 4 のように計算される.

表 4: ショケ積分の計算例

i	$\sigma(i)$	$x_{\sigma(i)}$	差	A	$\mu(A)$	積
1	2	0.7	0.2	$\{2\}$	0.3	0.06
2	1	0.5	0.3	$\{1, 2\}$	0.6	0.18
3	3	0.2	0.2	$\{1, 2, 3\}$	1	0.20
計						0.44

2.6 論理関数をショケ積分で計算する解釈

菅野積分は, Fuzzy/C の max-min 計算と同じであった. 表 5 にショケ積分で計算するとどのような解釈になるのかを示す.

表 5: ショケ積分のファジィ論理関数による解釈

$x_{\sigma(i)}$	区間	幅	論理式	論理式の成立度	積
1	(0.7, 1]	0.3	—	0	0
0.7(x_2)	(0.5, 0.7]	0.2	x_2	0.3	0.06
0.5(x_1)	(0.2, 0.5]	0.3	x_1x_2	0.6	0.18
0.2(x_3)	[0, 0.2]	0.2	$x_1x_2x_3$	1	0.20
計					0.44

表 3 の 2 値の真理値表の値で考えると, 各論理式の成立度 (ファジィ測度の値) を区間幅 (0.3, 0.2, 0.3, 0.2) の加重平均で出力値を求めている. したがって, 2 値入力の真理値表を成立する論理式の成立度を区間幅で補間していることになる.

確率との差異 max-min 計算の代わりに、確率の計算をするという考え方もある。しかし、程度の問題を議論するときは異なるのではないだろうか？ある学部に入學するとき、数学と化学の両方の能力が必要とされるとする。ある人が、数学の能力と化学の能力をそれぞれ 50% 有しているとき、その人は 25% でふさわしいのではなく、50% でふさわしいのではないだろうか？また、さらに物理の能力も必要になって、物理も 50% の能力を有しているとき、12.5% ではなく、50% でふさわしいのではないだろうか？

2.7 菅野積分とショケ積分の比較

図 1 と 図 2 は、菅野積分とショケ積分の比較である。

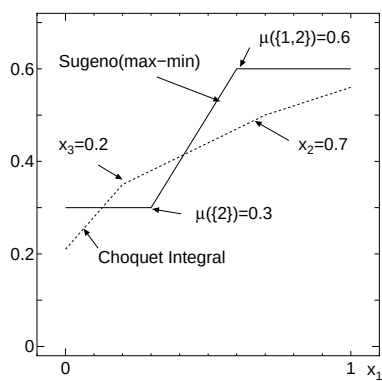


図 2: 表 3 での $x_2 = 0.7, x_3 = 0.2$

3 ファジィ推論とファジィ積分

ファジィ推論 [6] は、制御や評価のモデルで使われてきており、ファジィルールという理解しやすい表現が特徴になっている。ファジィ推論のルールを、入力値で補間（重心計算）するツールと考える。その補間の計算をファジィ積分で置き換えることが可能 [7] であり、ファジィ推論の重心計算と比較する。

3.1 ファジィ推論

ファジィ集合の重心演算を簡略にし、出力のファジィ集合をシングルトン（1 点からなるファジィ集合）とした簡略化推論を取り上げる。

総合成績を試験の点数 ($p_1 \in [0, 100]$) と平常点 ($p_2 \in [0, 100]$) の入力から求める例を考える。それぞれを低い (Low), 中くらい (Middle), 高い (High) の 3 つのファジィ集合 (試験, 平常点共通) で考える。メンバーシップ関数 (m_0, m_1, m_2) は共通で図 3 の三角型とする。

2 入力 3 分割の場合、ファジィルールは図 4 のように 9 つあり、すべての分割したファジィ集合の組が前件部になる。出力 $c(i, j)$ は、シングルトンファジィ集

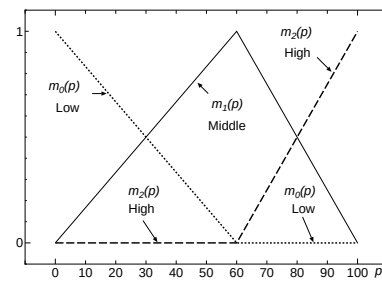


図 3: メンバーシップ関数の例

合であるのでスカラー値であり、代表点と呼んでいる。代表点に対して出力値が与えられたとき、代表点以外の入力値の場合を補間して出力するモデルとして、ファジィ推論やファジィ積分を捉える。図 4 は、代表点の値から補間すると考えたときの図解である。

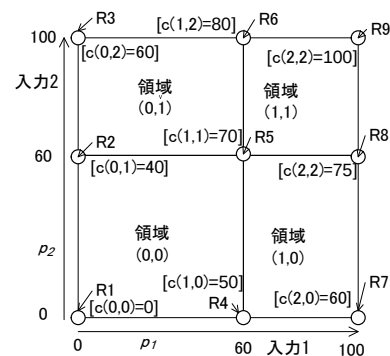


図 4: ファジィ推論

3.2 簡略型推論法（代数積加算重心法）

簡略型推論法（代数積加算重心法）は、すべてのファジィルール前件部の成立度を積で求め、代表点での出力値を前件部の成立度の加重平均で求めている。ある p_1 と p_2 を与えると、1 つの領域が選択され、4 つのルールのみが $m_i(p_1) \times m_j(p_2) > 0$ となりうる。その領域を選択領域と呼ぶ（境界線上の場合は任意の領域）。その選択領域の端点の代表点を前件部の成立度で重心演算を行う（表 6）。

3.3 ファジィ積分

前節のモデルのようにある 1 つの領域が選択されれば、その領域内の補間はファジィ積分で計算できる。 $c(i, j)$ の値からファジィ測度が決められ、各領域の左下の端点からのメンバーシップ関数の値が入力値になる。

3.3.1 菅野積分

菅野積分で計算する場合、ファジィ測度の条件を満たすために、ファジィルール表は入力値に対して単調

表 6: 簡略化推論法（代数積，論理積）

	c	$m_i(p_1) \times m_j(p_2)$	$m_i(p_1) \wedge m_j(p_2)$
R4	50	$\frac{3}{4} \times \frac{1}{6} = \frac{3}{24}$	$\frac{3}{4} \wedge \frac{1}{6} = \frac{1}{6}$
R5	70	$\frac{3}{4} \times \frac{5}{6} = \frac{15}{24}$	$\frac{3}{4} \wedge \frac{5}{6} = \frac{3}{4}$
R7	60	$\frac{1}{4} \times \frac{1}{6} = \frac{1}{24}$	$\frac{1}{4} \wedge \frac{1}{6} = \frac{1}{6}$
R8	75	$\frac{1}{4} \times \frac{5}{6} = \frac{5}{24}$	$\frac{1}{4} \wedge \frac{5}{6} = \frac{1}{4}$
重心		68.125	67.1875

増加とする．領域 (i, j) が選ばれたとき，

$$\mu(\emptyset) = 0, \mu(\{1\}) = \frac{c(i+1, j) - c(i, j)}{c(i+1, j+1) - c(i, j)}$$

$$\mu(\{2\}) = \frac{c(i, j+1) - c(i, j)}{c(i+1, j+1) - c(i, j)}, \mu(\{1, 2\}) = 1 \quad (5)$$

とし， $x_1 = m_{i+1}(p_1), x_2 = m_{j+1}(p_2)$ とする．出力値は，菅野積分の出力値を次式で戻す．

$$z = f_{\mu}^S(x_1, x_2)[c(i+1, j+1) - c(i, j)] + c(i, j) \quad (6)$$

3.3.2 ショケ積分

ショケ積分の場合も同様に，式 (5) のファジィ測度と $x_1 = m_{i+1}(p_1), x_2 = m_{j+1}(p_2)$ とする．出力値は，ショケ積分の出力値を菅野積分と同様な方法で戻す．

ショケ積分は，各代表点の出力値を直接補間することでも求めることができる．表 7 はその計算表である．

表 7: 集合関数に関するショケ積分の計算例

i	$\sigma(i)$	$x_{\sigma(i)}$	差	A	$\nu(A)$	積
0	0	1	0.1667	$\emptyset$	50	8.33
1	2	0.8333	0.5833	$\{2\}$	70	40.83
2	1	0.25	0.25	$\{1, 2\}$	75	18.75
計						67.92

4 計算ツール

本稿の作成で利用した計算ツールを公開する．

<http://www.isc.senshu-u.ac.jp/%7ethc0456/fuzzy/>

と日本知能情報ファジィ学会の「SOFT Computing Repository」から利用できるようにする予定である．

参考文献

- [1] M. Sugeno: *Theory of fuzzy integrals and its applications*, Doctoral Thesis, Tokyo Institute of Technology, 1974.

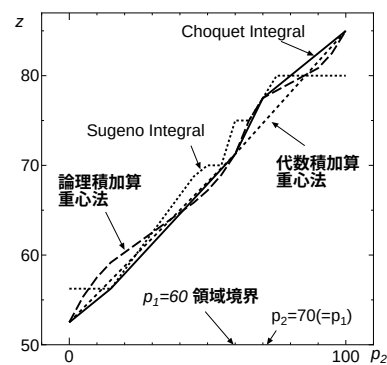


図 5: 比較 (図 4 のルール, $p_1 = 70$)

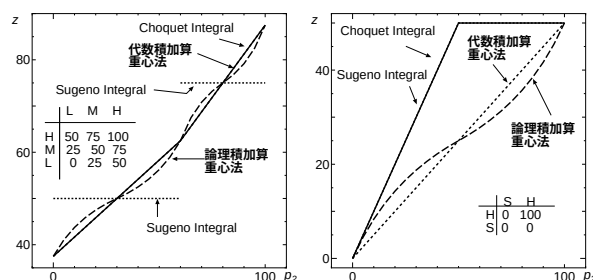


図 6: 左:連続性 ($p_1 = 80$), 右:程度 ($p_1 = 50$)

- [2] 奥田良治, 稲井田次郎: **ファジィ推論における単調性解析**, 早稲田大学数学教育学会誌, 21-1, pp.28-33, 2003.
- [3] 向殿政男: **ファジィ論理 (講座ファジィ 4)**, 日刊工業新聞社, 1993.
- [4] 荒木智行, 向殿 政男: **定数係数をもったファジー論理関数について**, 電子情報通信学会論文誌 D.I, J81-9, pp.1037-1047, 1998.
- [5] 菅野道夫, 室伏俊明: **ファジィ測度 (講座ファジィ 3)**, 日刊工業新聞社, 1993.
- [6] 水本雅晴他: **ファジィ集合 (講座ファジィ 2)**, 日刊工業新聞社, 1992.
- [7] 高萩栄一郎: **ファジィルール表によるショケ積分型総合評価法**, 知能と情報 (日本知能情報ファジィ学会誌), 21-4, pp.480-490, 2009.

連絡先

高萩 栄一郎

E-mail: takahagi@isc.senshu-u.ac.jp

BCK 代数の ordinal sum に関する一考察

A Consideration on the Ordinal Sum of BCK-Algebras

○河口 万由香 青木 俊介
○Mayuka F. Kawaguchi Shunsuke Aoki

北海道大学 大学院情報科学研究科 情報理工学専攻
Graduate School of Information Science and Technology, Hokkaido University

Abstract: In this research work, the ordinal sum method for totally ordered semigroups is extended to the case of partially ordered semigroups in order to apply it to a family of bounded BCK-algebras. The authors try to construct the ordinal sum by reforming BCK-algebras to semigroups via ajoininess property and then applying the method. In a conventional way, the ordinal sum has been constructed by applying the method and then reforming it to a semigroup via ajoininess property. As the main result of this work, it is verified that the above-mentioned two kinds of ordinal sum are coincided.

1. はじめに

本研究では、半群の族から新たな半群を生成する ordinal sum の方法を用いて、BCK 代数の族から新たな BCK 代数を生成する機序を調査する。ordinal sum の方法は t-ノルムの表現定理として全順序半群の俗から全順序半群を生成する際に用いられることが知られている[7], [9].

本稿では半順序半群の俗に対する ordinal sum の方法を定式化した上で、BCK 代数の ordinal sum について述べる。

2. Ordinal Sum

ordinal sum の概念はその起源を 1946 年の Climescu [3] まで遡る。

<Climescu [3]による ordinal sum の原型>

$(A; f)$ と $(B; g)$ を半群とし、 $A \cap B = \emptyset$ と仮定する。二項演算 h が $(A \cup B) \times (A \cup B)$ において

$$h(x, y) = \begin{cases} f(x, y) & (x \in A, y \in A) \\ y & (x \in A, y \in B) \\ x & (x \in B, y \in A) \\ g(x, y) & (x \in B, y \in B) \end{cases} \quad (1)$$

と定義されるとき、 $(A \cup B; h)$ は半群である。

一方 Clifford [1], [2] は半群の族の和集合がまた半群となる構造を考案した。ordinal sum という名称は Clifford によるものと思われる。

<Clifford [2]による ordinal sum の定義>

全順序 $<$ を有する半群 $(S; +, <)$ が単調性

$$a < b \Rightarrow a + c \leq b + c, c + a \leq c + b \quad (\forall c \in S) \quad (2)$$

を満足するとき S を ordered 半群 (o.半群) と呼ぶ。o.半群 S がさらに

$$a \leq a + b, b \leq a + b \quad (\forall a, b \in S) \quad (3)$$

を満足するとき S を positively ordered 半群 (pos.o.半群) と呼ぶ。また o.半群 S において二項演算 $+$ が可換であるとき S を ordered commutative 半群 (o.c.半群) と呼ぶ。

I を全順序集合とする。各 $i \in I$ に対して pos.o.c. 半群 $(S_i; +_i, <_i)$ を対応させ、 $S = \bigcup_{i \in I} S_i$ とおく。 $S_i \cap S_j = \emptyset$ ($i, j \in I, i \neq j$) と仮定する。このとき、 $\forall x, y \in S$ に対して

$$x < y \Leftrightarrow (x, y \in S_i, x <_i y) \text{ or } (x \in S_i, y \in S_j, i < j) \quad (4)$$

$$x + y = \begin{cases} x +_i y & (x, y \in S_i) \\ y & (x \in S_i, y \in S_j, i < j) \\ x & (x \in S_i, y \in S_j, j < i) \end{cases} \quad (5)$$

と定義されるとき $(S; +, <)$ は pos.o.c.半群である。このとき S を、pos.o.c.半群 S_i の全順序族 $\{S_i | i \in I\}$ の ordinal sum と呼ぶ。また pos.o.c.半群は、二つ以上の部分半群の ordinal sum で表現できないとき、ordinally irreducible であると呼ばれる。

Climescu の定義には半群 A, B の要素に順序は考慮されていない。唯一、 A と B については B に優先順位を与えていることがわかる。

一方、Clifford の定義において、各 S_i は全順序集合であり、その族 $\{S_i | i \in I\}$ も全順序構造を持ち、最終的に ordinal sum S は全順序集合である。実数体上の半群を意識した構造ではないかと考えられる。

本稿では、Clifford の ordinal sum における S_i が半順序半群 (partially ordered 半群, p.o.半群) である場合に拡張する。この場合も単調性(2)を仮定する。

【定義 1】 I を全順序集合とする. 各 $i \in I$ に対して pos.p.o.c. 半群 $(S_i; +_i, <_i)$ を対応させ, $S = \bigcup_{i \in I} S_i$ とおく. $S_i \cap S_j = \emptyset$ ($i, j \in I, i \neq j$) と仮定する. このとき, $\forall x, y \in S$ に対して

$$\begin{aligned} x < y &\Leftrightarrow (x, y \in S_i, x <_i y) \\ &\quad \text{or } (x \in S_i, y \in S_j, i < j) \\ x + y &= \begin{cases} x +_i y & (x, y \in S_i) \\ y & (x \in S_i, y \in S_j, i < j) \\ x & (x \in S_i, y \in S_j, j < i) \end{cases} \end{aligned}$$

と定義されるとき $(S; +, <)$ は pos.p.o.c. 半群である. このとき S を, pos.p.o.c. 半群 S_i の全順序族 $\{S_i | i \in I\}$ の ordinal sum と呼ぶ. また pos.p.o.c. 半群は, pos.o.c. 半群に場合と同様に, 二つ以上の部分半群の ordinal sum で表現できないとき, *ordinally irreducible* であると呼ばれる.

3. BCK 代数

本節においては, 以下の議論の準備として, BCK 代数に関する既知の理論[6]を概観する.

【定義 2】 代数系 $(X; *, 0)$ を考える. $*$ は X 上の二項演算, 0 は定数である. $\forall x, y, z \in X$ に対して以下の条件を満足するとき, $(X; *, 0)$ は BCK 代数であるという.

- (i) $((x * y) * (x * z)) * (z * y) = 0$
- (ii) $(x * (x * y)) * y = 0$
- (iii) $x * x = 0$
- (iv) $0 * x = 0$
- (v) $x * y = 0$ かつ $y * x = 0 \Rightarrow x = y$

【定義 3】 BCK 代数 $(X; *, 0)$ において, 二項関係 $\leq$ を以下のように定義する.

$$(vi) \quad x \leq y \stackrel{\text{def.}}{\Leftrightarrow} x * y = 0$$

BCK 代数 $(X; *, 0)$ において(vi)で定義される二項関係 $\leq$ は X 上の半順序関係であり, 定数 0 は半順序集合 $(X, \leq)$ の最小元である.

【定義 4】 BCK 代数 $(X; *, 0)$ において最大元 1 が存在するとき, この BCK 代数を有界であるという.

【定義 5】 BCK 代数 $(X; *, 0)$ において, $\forall x, y \in X$ に対して $s * x \leq y$ を満足する $s \in X$ の最大元が存在するとき, BCK 代数 X が条件(S)を満たすといい, その最大元を $x \circ y$ と表す.

BCK 代数 X が条件(S)を満たすとき, かつそのときに限り, 二項演算 $*$ と $\circ$ の間に以下の左随伴性が成り立つ.

$$\forall x, y, z \in X: x * y \leq z \Leftrightarrow x \leq y \circ z \quad (6)$$

以下, 条件(S)を満たす BCK 代数を BCK^S 代数と呼ぶこととする.

BCK^S 代数においては, $\forall x, y, z \in X$ に対して, 以下が成り立つ.

$$(x * y) * z = x * (y \circ z) \quad (7)$$

$$(x \circ y) * x \leq y \quad (8)$$

BCK^S 代数 $(X; *, 0)$ において, $(X; \circ, 0)$ は可換モノイドを成す. すなわち, X 上の二項演算 $\circ$ は可換律, 結合律を満たし, 0 を単位元として持つ. さらに次式で示される単調性を有する.

$$a \leq b \Rightarrow a \circ c \leq b \circ c \quad \text{for } \forall c \in X \quad (9)$$

BCK^S 代数 $(X; *, 0)$ が有界であるならば, 最大元 1 は次式を満足する. (証明は[5]を参照)

$$a \circ 1 = 1 \circ a = 1 \quad \text{for } \forall a \in X \quad (10)$$

すなわち, 有界 BCK^S 代数 $(X; *, 0, 1)$ において, 二項演算 $\circ$ は X 上の $\mathbf{t}$ -コノルムとなっている.

4. BCK 代数の ordinal sum

4.1 二個の BCK 代数からなる ordinal sum

二個の有界 BCK^S 代数 $\mathcal{X}_i = (X_i; \leq_i, *_i, \circ_i, 0_i, 1_i)$, $i \in \{1, 2\}$ を考える. 各二項演算 $\circ_i$ は結合的かつ可換であり, 単調性(9)を満足し, さらに 0_i が単位元であることから(3)を満たすことも容易に導かれる. 従って, $\mathcal{X}_i$ ($i \in \{1, 2\}$) は pos.p.o.c. 半群であり, [定義 1] を適用して ordinal sum を構成することができる.

【定義 6】 有界 BCK^S 代数 $\mathcal{X}_i = (X_i; \leq_i, *_i, \circ_i, 0_i, 1_i)$, $i \in \{1, 2\}$ において, $1_1 = 0_2$ かつ $(\mathcal{X}_1 \setminus \{1_1\}) \cap (\mathcal{X}_2 \setminus \{0_2\}) = \emptyset$ と仮定する. このとき, ordinal sum $\mathcal{X}_1 \oplus \mathcal{X}_2 = (X_1 \cup X_2; \leq, \circ, 0, 1)$ を次のように定義する.

$$x \leq y \Leftrightarrow (x, y \in X_i; x \leq_i y; i \in \{1, 2\}) \quad (11)$$

$$\text{or } (x \in X_1, y \in X_2)$$

$$x \circ y = \begin{cases} x \circ_i y & (x, y \in X_i; i \in \{1, 2\}) \\ y & (x \in X_1, y \in X_2) \\ x & (x \in X_2, y \in X_1) \end{cases} \quad (12)$$

$$0 = 0_1, 1 = 1_2 \quad (13)$$

Iorgulescu[4]は、有界 BCK 代数の ordinal sum を以下に示すように定義したが、BCK 代数の演算 $*$ は結合的ではなく、2 節で述べた ordinal sum を構成する要件を満たしていない。そこで有界 BCK^s 代数の場合に Iorgulescu の ordinal sum が定義 6 の場合に一致するか否かを吟味する。ただし、文献[4]では BCK 代数について異なる表記法を採用しており、さらに ordinal sum を “ordinal product” と呼んでいるが、本質的には全く等価な代数である。

<Iorgulescu[4]による有界 BCK 代数の ordinal sum>

有界 BCK 代数 $\mathcal{X}_i = (X_i; \leq_i, *, 0_i, 1_i)$, $i \in \{1, 2\}$ において、 $1_1 = 0_2$ かつ $(\mathcal{X}_1 \setminus \{1_1\}) \cap (\mathcal{X}_2 \setminus \{0_2\}) = \emptyset$ と仮定する。このとき、ordinal sum $\mathcal{X}_1 \oplus_1 \mathcal{X}_2 = (X_1 \cup X_2; \leq, *, 0, 1)$ を次のように定義する。

$$\begin{aligned} x \leq y &\Leftrightarrow (x, y \in X_i; x \leq_i y; i \in \{1, 2\}) \\ &\quad \text{or } (x \in X_1, y \in X_2) \\ x * y &= \begin{cases} 0 & (x \leq y) \\ x *_i y & (x \not\leq y; x, y \in X_i; i \in \{1, 2\}) \\ x & (x \not\leq y, x \in X_2 \setminus \{0_2\}, y \in X_1 \setminus \{1_1\}) \end{cases} \quad (14) \end{aligned}$$

$$0 = 0_1, 1 = 1_2$$

このとき $\mathcal{X}_1 \oplus_1 \mathcal{X}_2$ は有界 BCK 代数となる。また各
有界 BCK 代数 $\mathcal{X}_1, \mathcal{X}_2$ が条件(S)を満たすならば、
 $\mathcal{X}_1 \oplus_1 \mathcal{X}_2$ もまた有界 BCK^s 代数となる。

式(14)で定義された $*$ 演算から条件(S)を介して演算 $x \circ y = \max\{s \mid s * x \leq y\}$ が導出される。ordinal sum $\mathcal{X}_1 \oplus_1 \mathcal{X}_2$ における $*$ と $\circ$ の演算表をそれぞれ表 1, 表 2 に示す。表 2 の演算表が示す演算は式(12)の演算と一致していることが確かめられる。すなわち、Iorgulescu による ordinal sum は条件(S)を満足する場合には、pos.p.o.c.半群である演算 $\circ$ から生成される ordinal sum に一致することが判明した。

4.2 三個の BCK 代数からなる ordinal sum

本節においては、Iorgulescu[4]によって提案された BCK 代数の ordinal sum の拡張を試みる。

[定義 7] 二つの有界 BCK 代数

$\mathcal{X}_i = (X_i; \leq_i, *, 0_i, 1_i)$, $i = \{1, 2\}$ および BCK 代数 $\mathcal{X}_3 = (X_3; \leq_3, *, 0_3)$ を考える。ただし、 $1_1 = 0_2$, $1_2 = 0_3$, $X_1 \setminus \{1_1\} \cap X_2 \setminus \{0_2\} = \emptyset$, $X_2 \setminus \{1_2\} \cap X_3 \setminus \{0_3\} = \emptyset$ とする。このとき、BCK 代数 $\mathcal{X}_1, \mathcal{X}_2, \mathcal{X}_3$ の ordinal sum $\mathcal{X}_1 \oplus \mathcal{X}_2 \oplus \mathcal{X}_3$ を以下のように定義する。

$$\mathcal{X}_1 \oplus \mathcal{X}_2 \oplus \mathcal{X}_3 = (X_1 \cup X_2 \cup X_3; \leq, *, 0)$$

$$0 = 0_1$$

$$\begin{aligned} x \leq y &\Leftrightarrow (x, y \in X_i; x \leq_i y; i \in \{1, 2, 3\}) \\ &\quad \text{or } (x \in X_i; y \in X_j; i < j; i, j \in \{1, 2, 3\}) \end{aligned} \quad (15)$$

$$x * y = \begin{cases} 0 & (x \leq y) \\ x *_i y & (x \not\leq y; x, y \in X_i; i \in \{1, 2, 3\}) \\ x & (x \not\leq y; x \in X_i; y \in X_j; j < i; i, j \in \{1, 2, 3\}) \end{cases} \quad (16)$$

[定理 1] 定義 7 による、BCK 代数 $\mathcal{X}_1, \mathcal{X}_2, \mathcal{X}_3$ の ordinal sum $\mathcal{X}_1 \oplus \mathcal{X}_2 \oplus \mathcal{X}_3 = (X_1 \cup X_2 \cup X_3; \leq, *, 0)$ は BCK 代数となる。

[定理 2] 定義 7 において、BCK 代数 $\mathcal{X}_3$ が有界で最大元 1_3 をもつならば、 $\mathcal{X}_1, \mathcal{X}_2, \mathcal{X}_3$ の ordinal sum $\mathcal{X}_1 \oplus \mathcal{X}_2 \oplus \mathcal{X}_3 = (X_1 \cup X_2 \cup X_3; \leq, *, 0, 1)$ 、ただし $1 = 1_3$ 、は有界 BCK 代数となる。

[定理 3] 定義 7 において、BCK 代数 $\mathcal{X}_1, \mathcal{X}_2, \mathcal{X}_3$ が条件(S)を満足するならば、 $\mathcal{X}_1, \mathcal{X}_2, \mathcal{X}_3$ の ordinal sum $\mathcal{X}_1 \oplus \mathcal{X}_2 \oplus \mathcal{X}_3 = (X_1 \cup X_2 \cup X_3; \leq, *, 0)$ は BCK^s 代数となる。

このとき、条件(S)によって演算 $*$ から導出される演算 $x \circ y = \max\{s \mid s * x \leq y\}$ は次式で示される。

$$x \circ y = \begin{cases} x \circ_i y & (x, y \in X_i; i \in \{1, 2, 3\}) \\ y & (x \in X_i; y \in X_j; i < j; i, j \in \{1, 2, 3\}) \\ x & (x \in X_i; y \in X_j; j < i; i, j \in \{1, 2, 3\}) \end{cases} \quad (17)$$

この二項演算 $\circ$ は可換律、結合律を満たし、 0 を単位元として持ち、さらに単調性を有する。零元の役割を持つ最大元が存在しないことを除くと、ほぼ t-コノルムとみなすことができる。BCK 代数 $\mathcal{X}_3$ が有界であるならば t-コノルムである。

6. むすび

本研究においては、BCK 代数の表現定理の一種を表すための ordinal sum について、その期限と本質的な性質を調査した。全順序半群の族に限定されていた定義を半順序半群の場合に拡張し、さらに二個あるいは三個の BCK 代数からなる ordinal sum の構造を整理した。

本稿では Clifford [2] に倣って positively ordered 半群に基づいた記法で統一して記述したが, Iorgulescu [4] の記法は双対な形になっており, positively に対して negatively とでも呼ぶべき構造である. これは t -コノルムと t -ノルムの双対性に対応している.

参考文献

- [1] A.H. Clifford, “Naturally totally ordered commutative semigroups,” *American J. Mathematics*, Vol. 76, pp. 631-646, 1954.
- [2] A.H. Clifford, “Totally ordered commutative semigroups,” *Bull. Amer. Math. Soc.*, Vol. 64, pp. 305-316, 1958.
- [3] A.C. Climescu, “Sur l’équation fonctionnelle de l’associativité,” *Bull. École. Polytech. Jassy.*, Vol. 1, pp. 1-16, 1946.
- [4] A. Iorgulescu, “On BCK algebras: Part II: New algebras. The ordinal sum (product) of two bounded BCK algebras,” *Soft Computing*, Vol. 12, pp. 835-856, 2008.
- [5] K. Iséki, “A special class of BCK-algebras,” *Mathematics Seminar Notes*, Kobe Univ., Vol. 5, pp. 191-198, 1977.
- [6] K. Iséki, S. Tanaka, “An introduction to the theory of BCK-algebras,” *Mathematica Japonica*, Vol. 23, no. 1, pp. 1-26, 1978.
- [7] S. Jenei, “A note of the ordinal sum theorem and its consequence for construction of triangular norms,” *Fuzzy Sets and Systems*, Vol. 126, pp. 199-205, 2002.
- [8] 河口万由香, “離散 t -ノルムと BCK 代数,” 第 29 回ファジィシステムシンポジウム講演論文集, pp. 29-32, 2013.
- [9] E.P. Klement, R. Mesiar, E. Pap, *Triangular Norms*, Kluwer Academic Pub., 2000.
- [10] S. Tanaka, “Examples of BCK-algebras,” *Mathematics Seminar Notes*, Kobe Univ., Vol. 3, pp. 75-82, 1975.

連絡先

河口 万由香
 北海道大学 大学院情報科学研究科
 情報理工学専攻
 〒060-0814 札幌市北区北 14 条西 9 丁目
 Phone: 011-706-6805
 Fax: 011-706-7830
 E-mail: mayuka@ist.hokudai.ac.jp

表 1. ordinal sum $\mathcal{X}_1 \oplus_1 \mathcal{X}_2$ における $*$ の演算表

$x \backslash y$		X_1			X_2	
		0	y	$1_1 = 0_2$	y	1
X_1	0	0	0	0	0	0
	x	x	$x *_1 y$	0	0	0
X_2	$1_1 = 0_2$	$1_1 = 0_2$	$1_1 *_1 y$	0	0	0
	x	x	x	x	$x *_2 y$	0
	1	1	1	1	$1 *_2 y$	0

表 2. ordinal sum $\mathcal{X}_1 \oplus_1 \mathcal{X}_2$ における $\circ$ の演算表

$x \backslash y$		X_1			X_2	
		0	y	$1_1 = 0_2$	y	1
X_1	0	0	y	$1_1 = 0_2$	y	1
	x	x	$x \circ_1 y$	$1_1 = 0_2$	y	1
X_2	$1_1 = 0_2$	$1_1 = 0_2$	$1_1 = 0_2$	$1_1 = 0_2$	y	1
	x	x	x	x	$x \circ_2 y$	1
	1	1	1	1	1	1

道路網上のボトルネックで分割した ホップフィールドネットワークによる効率的な経路探索

An Efficient Route Search Method for Road System
Containing Bottlenecks by Multiple Hopfield Networks

○ 仲田 紘平, 和久屋 寛, 伊藤 秀昭
○ Kohei Nakata, Hiroshi Wakuya, Hideaki Itoh
佐賀大学

Saga University

Abstract: A Hopfield network is a good tool for solving combinatorial optimization problems. But it is well-known that as the task size becomes bigger, the performance becomes poorer. Then, in order to overcome this problem, a search method to find out an optimal solution through dividing a large-scale problem into several regions by bottlenecks, such as a railroad crossings, is tried. In this study, the proposed method is applied to the task with two bottlenecks. As a result of computer simulations, it is found that the proposed method is effective.

1 はじめに

組み合わせ最適化問題とは、多数ある組み合わせの中から与えられた制約条件を最も満たす解を求める問題である。この問題を効率的に解く手法として、ホップフィールドネットワーク [1] や遺伝的アルゴリズム [2] などが知られている。しかし遺伝的アルゴリズムでは、設定するパラメータが多いという問題点を持ち、またホップフィールドネットワークでは、問題の規模が大きくなるにつれ、解候補が増大し正答率が低下して、求解が困難となる傾向がある。そこで、

1. ネットワークの分割
2. 制約条件の分割
3. 問題の分割

という 3つの分割法を提案し、これまでに最初の 2つの分割法を検討した [3][4]。そこで今回は 3つ目の、問題を分割し一度に考慮する領域を小さくした問題を解き、それらの解を統合するという手法について検討する。ここでは、具体例として経路探索問題を取り上げる。

以下本稿では、2でホップフィールドネットワークを紹介し、3で経路選択問題の説明を行い、4で先行研究の概要を紹介する。そして 5で計算機シミュレーションの結果を示し、6で問題の分割について考察を行い、7で結論を述べる。

2 ホップフィールドネットワーク

これは、物理学との対応関係を参考にしたニューラルネットワークであり、与えられた制約条件に基づい

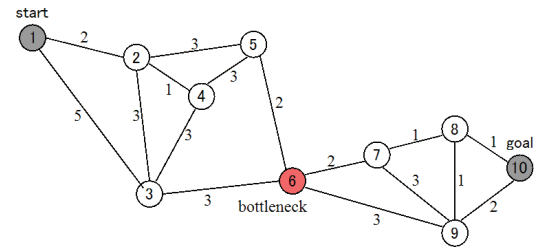


図 1: 経路選択問題の例

て設計したエネルギー関数から、結合荷重 W_{xy} と閾値 θ_x を求める。そして、式 (1)~(2) に従いエネルギーの減少する方向へ状態遷移を繰り返すことで、ニューロンの出力 V_{xy} から解が得られる手法である。このとき、ネットワークの初期状態によってはエネルギー極小値に陥り、局所解となることがある。

$$V_i(t+1) = f\left(\sum_{j=1}^n W_{ij} V_j(t) - \theta_i\right) \quad (1)$$

$$f(z) = \begin{cases} 1, & z \geq 0 \\ 0, & z < 0 \end{cases} \quad (2)$$

3 経路選択問題

ここでは、地図上の交差点や道に着目し、ある地点から目的地へのルート決めをする際に地図から与えられる道路網を、図 1 に示すようなグラフ構造に描き直して経路探索問題を解く。ここでノードとは、線路を結ぶ交差点のような場所のことを示し、図 1 の○を表す。また、図 1 の線路部分にある数字はノード間の距離を表している。

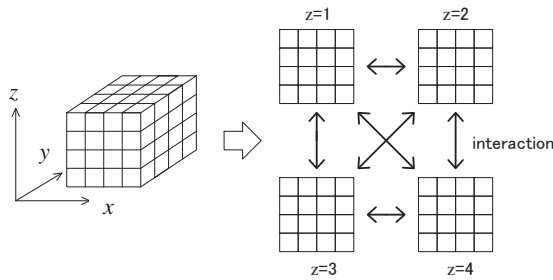


図 2: ネットワークの分割

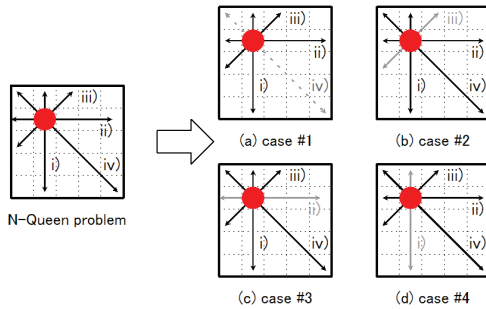


図 3: 制約条件の分割

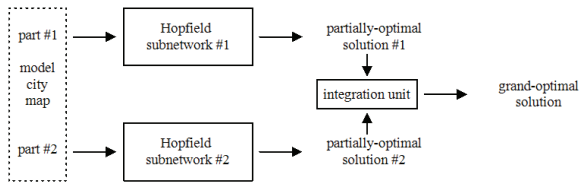


図 4: 問題の分割による解探索の手順

このとき、経路選択問題とは、

距離の与えられたノード配置において、ある2つのノード間を結ぶ経路の中で、総距離が最小となる経路を求める問題

である。この経路探索問題を解く際の制約条件は以下に示す3つである。

- i) 同時に複数のノードを訪れない。
- ii) 同じノードを複数回通過しない。
- iii) 経路長を最小化する。

上述の3種類の分割法のうち、1番目のネットワークの分割は、図2に示すように、全体のネットワークを複数のサブネットワークに分割し、それぞれのサブネットワーク毎に最適化を行う手法である。このとき、当初のネットワークが持つ他のネットワークからの影響を相互作用として与えることで、本来の問題に対する解が得られる。また、2番目の制約条件の分割では、図

表 1: 先行研究の概要
～ボトルネックが単一の場合～

	従来法 [%]	提案法 [%]
最短経路探索率	0	2.4
準最短経路探索率	0	23.1
経路探索率	0	86.4

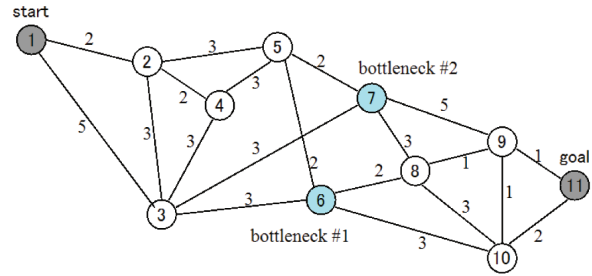


図 5: 経路探索問題の例
～ボトルネックが複数存在する場合～

3に示すように制約条件の一部を弱めた部分問題を複数作りだし、これらを繰り返し提示することで最適化を行う手法である。ここで、状態遷移を行う度に弱めた条件の影響を元へ戻すことで、本来の問題の解が求められる。

これらに対して今回は、例えば踏切のように経路選択の自由度が低い点を「ボトルネック」と命名し、これが存在する道路網について取り上げる。このような道路網について、ボトルネックに着目し、与えられた問題を分割して解く方法を検討する。

4 先行研究の概要

ボトルネックが単一の場合、図4に示すように、ここで分割した問題をそれぞれ個別のホップフィールドネットワークに割り当てて解き、それらの結果を統合することで本来の問題の解を求める。先行研究 [5] では、図1のようにノード6をボトルネックと定めた場合の検討を行った。ここでは1000通りの初期状態を準備し、経路探索を行ったところ表1のようになった。従来法では1例も経路を発見できていないが、提案法では86.4%が発見できた。最も距離の短い経路を発見した割合(正答率とする)が2.4%、次に距離の短い経路を発見した割合が23.1%となった。よって、問題を分割するというアイデアの有効性が認められた。なお、本手法では分割したそれぞれの問題に対するネットワークにおいて、同時に最適解が得られなければ、当初の問題に対する最適解を求めることができない。

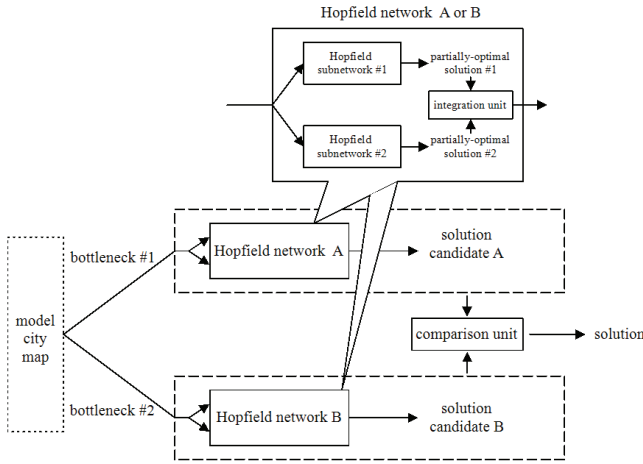


図 6: 問題の分割による解探索の手順
～ボトルネックが複数存在する場合～

表 2: 計算機シミュレーションの結果

(a) 実験データのまとめ方

		network B			
		correct B ₁	wrong B ₂	fail B ₃	total B ₀
network A	correct A ₁	n_{11}	n_{12}	n_{13}	n_{10}
	wrong A ₂	n_{21}	n_{22}	n_{23}	n_{20}
	fail A ₃	n_{31}	n_{32}	n_{33}	n_{30}
	total A ₀	n_{01}	n_{02}	n_{03}	n_{00}

(b) 実際に得られた数値データ

		network B			
		correct B ₁	wrong B ₂	fail B ₃	total B ₀
network A	correct A ₁	1	4	4	9
	wrong A ₂	3	256	131	390
	fail A ₃	3	350	248	601
	total A ₀	7	610	383	1000

5 計算機シミュレーション

5.1 実験方法

次に、ボトルネックが複数存在する場合について取り上げる。今回は図 5 のように、ボトルネックが 2 つの場合（ノード #6 と #7）について検討する。このときの解探索の手順は図 6 のようになる。

まず、ノード #6 を通る場合とノード #7 を通る場合に分けて考える。それぞれのネットワーク A と B では、ボトルネックが単一に定まるときと同様に、図 4 の手順にしたがって最適化を行い、最短経路を求める。その後、両者の経路長を比較することで問題の解を求める。ここでは、各ネットワークをペアとして 1000 通りの初期状態を用意し、i), ii), iii) の制約条件に対する重み付けを 10000:1000:1 と定めて計算機シミュレーションを行う。

表 3: 計算機シミュレーションの結果
～両ネットワークを独立試行とした場合～

(a) 実験データのまとめ方

	network A	network B
correct	n_{A1}	n_{B1}
wrong	n_{A2}	n_{B2}
fail	n_{A3}	n_{B3}
total	n_{A0}	n_{B0}

(b) 実際に得られた数値データ

	network A	network B
correct	9	7
wrong	390	610
fail	601	383
total	1000	1000

このときの結果は、例えば表 2(a) のようにまとめられる。ここで A_1, A_2, A_3 は、それぞれネットワーク A の最短経路、最短以外の経路、経路が求まらない場合を表し、 $A_0 = A_1 + A_2 + A_3$ とする。同様にして B_1, B_2, B_3 は、それぞれネットワーク B の最短経路、最短以外の経路、経路が求まらない場合を表し、 $B_0 = B_1 + B_2 + B_3$ とする。また、 $n_{11} \sim n_{33}$ は各項目の個数を表す。最短経路が求まる確率 P_1 は、

$$P_1 = \frac{n_{11}}{n_{00}} \quad (3)$$

である。

5.2 実験結果

このときの結果を、表 2(b) に示す。これは表 2(a) に各数値を代入したものである。このとき、最短経路が求まる確率（正答率） P_1 は、

$$P_1 = \frac{1}{1000} = 0.001 \quad (4)$$

となる。

なお、参考のため、ここで、A, B をそれぞれ独立に解探索を行ったとみなすと、表 3 より最短経路が求まる確率 P_2 は、

$$\begin{aligned} P_2 &= \frac{n_{A1}}{n_{A0}} \times \frac{n_{B1}}{n_{B0}} \\ &= \frac{9}{1000} \times \frac{7}{1000} = 0.000063 \end{aligned} \quad (5)$$

となる。このことから、本手法ではペアにして解探索することで最適解を求めやすくなることが分かる。

ところで、これまでは最短経路のみを考慮していたため、ネットワーク A, B の両方が最短経路の場合のみを考えてきた。ここで、図 6 の解候補 A, B と最終

的な解候補について距離に応じた経路数分布を確認した。表 2(b) よりネットワーク A により求めた経路数は 399 通りであり、またネットワーク B により求めた経路数は 617 通りとなる。また、本手法において、いずれかのネットワークで経路が求めたのは 752 通りとなる。このときの距離の分布を図 7 に示す。図 7(a), 図 7(b) から分かるように、ネットワーク A では距離の分布は短いところに集まっているが経路数は少なく、ネットワーク B では経路数は多いが距離の分布が広がっていることが分かる。しかし、図 7(c) に示すように、両者を比較すると距離の分布が短いところに集まって経路数も多くなる。ここで求まる経路数が多くなるのは、どちらか一方のネットワークで経路が求まっていれば、問題に対する解となるからであり、また、距離の分布が短いところへ移動するのは、比較を行うことでより経路長の短いものが選択されるからである。このように本手法を用いた解探索では、最短経路が求められていない場合であっても、比較的距離の短い経路を多く求められることが確認できた。

6 考察

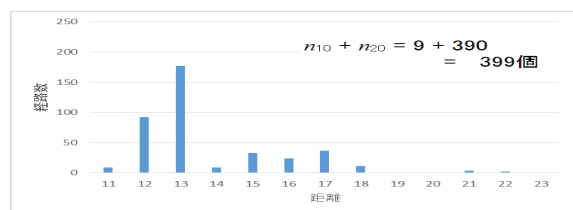
今回は、ボトルネックを 2 つ含んだ道路網の経路探索問題を取り上げ、問題の分割を用いた解探索法について検討した。その結果、最短経路を求めることができ、本手法は経路探索問題に有効であることが確認できた。しかし、分割した問題のそれぞれで求められる解候補の数が少なく、これが最終的な解候補の数にも影響するため、それぞれでの解探索率を上げる必要があると考えられる。

7 おわりに

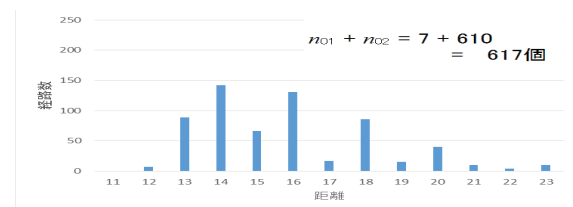
本研究では、経路探索問題を例に問題の分割を用いた解探索法について検討した。その結果、本手法を用いた経路探索において経路探索率が上昇することが確認できた。今後の展望として、本手法による解探索の性能をより向上させることが考えられる。

参考文献

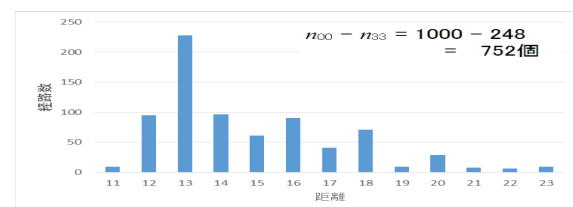
- [1] J.J.Hopfield and D.W.Tank : “ Neural computation of decisions in optimization problems ”, Biol.Cybern., Vol.52, pp.141-152, 1985
- [2] J.H.Holland : “ Adaptation in Natural and Artificial Systems ”, The University of Michigan Press, Ann Arbor, MI, 1975



(a) 解候補 A の分布



(b) 解候補 B の分布



(c) 最終的な解候補の分布

図 7: 距離に応じた経路数分布

- [3] 仲田紘平, 和久屋寛, 伊藤秀昭 : “ 分割されたホップフィールドネットワークを用いた組み合わせ最適化問題の解探索 ”, 第 3 回日本知能情報ファジィ学会九州支部学生会研究発表会予稿集, A104, pp.7-8, 2013
- [4] 仲田紘平, 和久屋寛, 伊藤秀昭 : “ 部分的な制約条件の繰り返し適用による組み合わせ最適化問題の求解能力改善 ”, 第 24 回ソフトサイエンス・ワークショップ講演論文集, P20-2, pp.31-34, 2014
- [5] H.Wakuya, K.Nakata, H.Itoh, S.H.Oh and Y.S.Oh : “ Evacuation route planning for disaster prevention design with multiple Hopfield networks ”, Proc. 1st Int. Symposium on ICT-based Disaster Prevention Design, S-2-4, pp.38-39, 2015

連絡先

和久屋 寛

佐賀大学大学院工学系研究科電気電子工学専攻

E-mail: wakuya@cc.saga-u.ac.jp

階層型ニューラルネットワークによる音声信号分離 ～分離行列の多様性～

Vocal Signal Separation by a Layered Neural Network
–Variety of Developed Separation Matrices–

○ 久木原 健介, 和久屋 寛, 伊藤 秀昭
○ Kensuke Kukihara, Hiroshi Wakuya, Hideaki Itoh
佐賀大学
Saga University

Abstract: Independent component analysis (ICA) is a signal separation technique inspired by famous psychological phenomena called “cocktail party effect”. Its basic idea has been applied to various kinds of tasks by a lot of researchers so far. In our previous studies, a vocal signal separation task by a layered neural network with structural pruning was tried, but all trials to extract only source signals was not completed successfully. Because relationship between the separation matrix and the mixing matrix have not been examined yet, some computer simulations focused on three points are also tried in this study. As a result, a signal separation is completed successfully, and correspondence between the separation matrix and the mixing matrix is confirmed analytically.

1. はじめに

心理学の分野で“カクテルパーティー効果”と呼ばれている信号分離能力を工学的に模倣する技術の一つとして、独立成分分析 (Independent Component Analysis : ICA) が知られている [1]。ICA における重要な点は、源信号の数と同等以上の観測信号を用意する必要があることと、復元信号は振幅の任意性と順序の入れ替えについて許容されるということである。

これまでに、カクテルパーティー効果が生体脳での情報処理に基づいているという考えのもと、入力層、隠れ層、出力層の3層から成る階層型ニューラルネットワークに、構造刈り込み学習を適用することによってICAを実現する手法が提案されている [2]。また、これを用いた適用事例の報告 [3][4] や構造刈り込み学習の詳細な性質に関する報告 [5] もある。

しかし、先行研究 [4] では、2種類の源信号以外にも、それらが混ざったままの信号が現れており、源信号のみを復元する完全な信号分離は実現できていなかった。そこで本研究では、構造刈り込み学習の効果の影響度を見直して、源信号のみを復元する完全な信号分離の実現を目指す。

また、これまでの研究を概観すると、従来は信号分離を実現する学習法に関するものが中心であり、学習によって獲得される分離行列と混合行列の関係性については議論が乏しかったようである。そこで本研究では、その点に着目した検討も行う。

以下本稿では、まず2で構造刈り込み学習に基づくICAの計算手法について紹介する。次に3で信号分離のための計算機シミュレーションを行う。また、4で学習で獲得した分離行列と混合行列の関係性を検討する。そして、5ではこれらについて考察を行い、最後に6で結論を述べる。

2. 構造刈り込み学習を用いたICA

2.1. 基本的な考え方

ICA は、 N 個の独立した未知の源信号 $s(t)$ が、任意の混合行列 A で線形混合されて得られた M 個の観測信

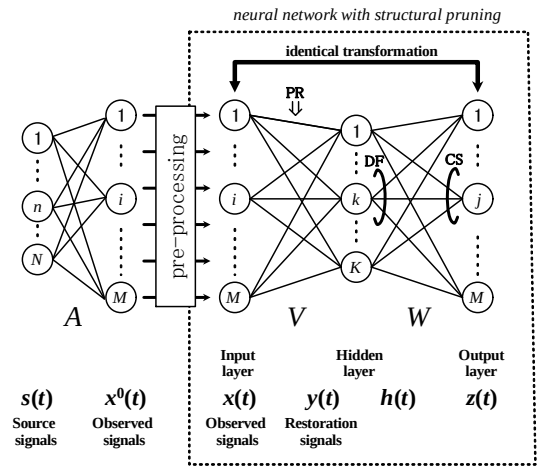


図1. 階層型ニューラルネットワークを用いた信号分離

号 $x(t)$ から、混合行列 A の逆変換[†]にあたる分離行列 V を推定して、復元信号 $y(t)$ を得るものである。

$$x(t) = As(t) \quad (1)$$

$$y(t) = Vx(t) = A^{-1}x(t) \quad (2)$$

本研究では、観測信号 $x(t)$ から復元信号 $y(t)$ に至る分離行列 V を、入力層、隠れ層、出力層の3層から成る階層型ニューラルネットワークを用いて獲得する。図1にニューラルネットワークによるICAでの信号の流れを示す。ここで、 N は源信号の数、 M は入力層と出力層のニューロンの数、 K は隠れ層のニューロンの数をそれぞれ表している。また V は入力層と隠れ層の間の結合荷重、 W は隠れ層と出力層の間の結合荷重をそれぞれ表している。図1における、入力層の出力 $x(t)$ 、隠れ層の活性化レベル $y(t)$ 、隠れ層の出力 $h(t)$ 、出力層の出力 $z(t)$ の入出力関係は以下のとおりである。

$$y_k(t) = \sum_{i=1}^M v_{ki}x_i(t) - \theta_{xk} \quad (3)$$

[†]一般的なICAでは逆行列である。混合行列 A が正方行列でない場合は一般化逆行列になる。

$$h_k(t) = \tanh \{y_k(t)\} \quad (4)$$

$$z_j(t) = \sum_{k=1}^K w_{jk} h_k(t) - \theta_{hj} \quad (5)$$

ここで、 θ_{xk} , θ_{hj} はニューロン $h_k(t)$, $z_j(t)$ の閾値を表している。なお、本手法での ICA の実現にあたっては、前処理として観測信号に中心化と標準化を施した後、入力層へ与える。

本手法では、入力層と出力層の間で恒等変換になるように学習を行う。このとき、階層型ニューラルネットワークで、入力層と出力層のニューロン数よりも隠れ層のニューロン数が少なければ ($M > K$)、いわゆる砂時計型のニューラルネットワークとなり、入力層と隠れ層の間で次元の圧縮、隠れ層と出力層の間で次元の復元が行われ、隠れ層に必要な信号のみが抽出されることになる。これはオートエンコーダと呼ばれており、本手法における復元信号は、隠れ層に現れることになる。

2.2. 構造刈り込み学習

本手法における学習は以下の式で行う。

$$e(t) = \frac{1}{2} \sum_{j=1}^M \{z_j(t) - x_j(t)\}^2 \quad (6)$$

$$w_{jk}^{[n+1]} = w_{jk}^{[n]} + \Delta w_{jk}^{[n]} \quad (7)$$

$$\Delta w_{jk}^{[n]} = -\eta \frac{\partial e(t)}{\partial w_{jk}^{[n]}} + \alpha \Delta w_{jk}^{[n-1]} + \varepsilon S_{jk} + \gamma F_{jk} \quad (8)$$

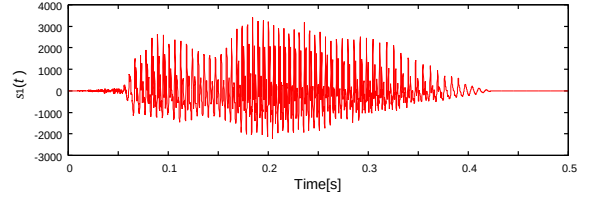
$$v_{ki}^{[n+1]} = \left(v_{ki}^{[n]} + \Delta v_{ki}^{[n]} \right) R_k \quad (9)$$

$$\Delta v_{ki}^{[n]} = -\eta \frac{\partial e(t)}{\partial v_{ki}^{[n]}} + \alpha \Delta v_{ki}^{[n-1]} \quad (10)$$

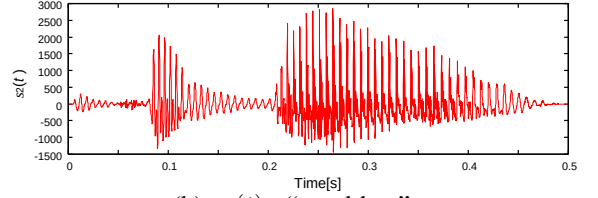
ここで、 n は学習回数を表す。また、 η は学習率、 α は慣性率、 ε は集束抑制率、 γ は発散促進率を定める定数である。このとき、式 (8), (10) の右辺第 1 項は、最急降下法に基づく結合荷重の修正量、第 2 項は慣性項を表す。これらは誤差逆伝搬法で用いるものである。

式 (8) の右辺第 3 項は集束抑制効果 S_{jk} である。これは、出力層のニューロン $z_j(t)$ が隠れ層から受け取る信号を、特定のニューロンへ集束する学習が行われる。また、式 (8) の右辺第 4 項は発散促進効果 F_{jk} である。これは、隠れ層ニューロン $h_k(t)$ から分かれる信号が出力層の多くのニューロンへ発散することによって、広く行き渡る学習が行われる。そして、式 (9) の R_k は比例縮小効果を表している。 R_k は、隠れ層ニューロン $h_k(t)$ から分かれる結合荷重 $w_{sk}^{[n]} (s = 1, 2, \dots, M)$ の平均的な強度を意味しており、その値が小さければ、不要なニューロンとして徐々に弱めていく学習が行われる。

以上の 3 つの構造刈り込み学習の効果の影響度を、学習の式の各パラメータによってバランスよく与えることで、2.1 で述べた砂時計型のニューラルネットワークを実現する。



(a) $s_1(t)$: “hello”



(b) $s_2(t)$: “good-bye”

図 2. 本研究における源信号

3. 音声信号分離の計算機シミュレーション

本手法の有効性を確認するため、音声信号を用いた信号分離実験を行った。今回は、図 2 に示すように、源信号として “hello” と “good-bye” という 2 種類の合成音声信号を採用した。これらは、サンプリング間隔が $125\mu\text{s}$ 、サンプル数が 4000 点から成る、長さ 0.5s のもので、フリーの音声読上ソフトウェア “SofTalk” [6] を用いて作成した。これらの源信号を事前に用意した混合行列 A によって混合し 4 種類の観測信号を作成した。

次に、入力層が 4 個、隠れ層が 6 個、出力層が 4 個のニューロンから成る階層型ニューラルネットワークに前処理を施した観測信号を与えた。なお、結合荷重の初期値は $[-0.5, 0.5]$ の閉区間からなる一様乱数で与えている。先行研究 [4] では、式 (8) における各定数を $\eta = 0.9$, $\alpha = 0.1$, $\varepsilon = 0.001$, $\gamma = 0.1$ と設定していた。しかし、この定数で学習した時には復元信号として源信号が混ざったままの信号も現れてしまった。そこで、先行研究 [5] の知見に基づいて、構造刈り込み学習の効果を強めるために、 ε と γ を大きくして、 $\eta = 0.9$, $\alpha = 0.1$, $\varepsilon = 0.005$, $\gamma = 0.5$ と設定した。

異なる 10 種類の初期状態を用意して、30,000 回ずつの学習を行った結果、源信号のみを復元する完全な信号分離に成功したのは全体の 20% であった[†]。そのときの各ニューロンの応答波形の一例を図 3 に示す。

図 3 における各波形の配置は、階層型ニューラルネットワークの各ニューロンに対応しており、入力層と出力層の間で恒等変換を学習できていることがわかる。このときの源信号 $s(t)$ と復元信号 $y(t)$ の類似度の目安である相関係数は表 1 のようになった。

これらの結果から、隠れ層の 1, 2, 4, 6 番目のニューロン $y_1(t)$, $y_2(t)$, $y_4(t)$, $y_6(t)$ が不要とみなされて削除され、3 番目のニューロン $y_3(t)$ では “hello” が、5 番目のニューロン $y_5(t)$ では “good-bye” に近い音声が見れたことがわかる。よって、本手法で源信号のみを復元する完全な信号分離を実現できた。なお、今回の計算機シミュレーションでは、源信号が 2 つであるということが既知であったため、先行研究 [4] の条件よりも強めればよい

[†] 源信号を 2 つとも復元できているが、混ざったままの信号も現れている例を含めると、源信号の復元率は 70% となった。これは先行研究 [4] よりも高い値である。

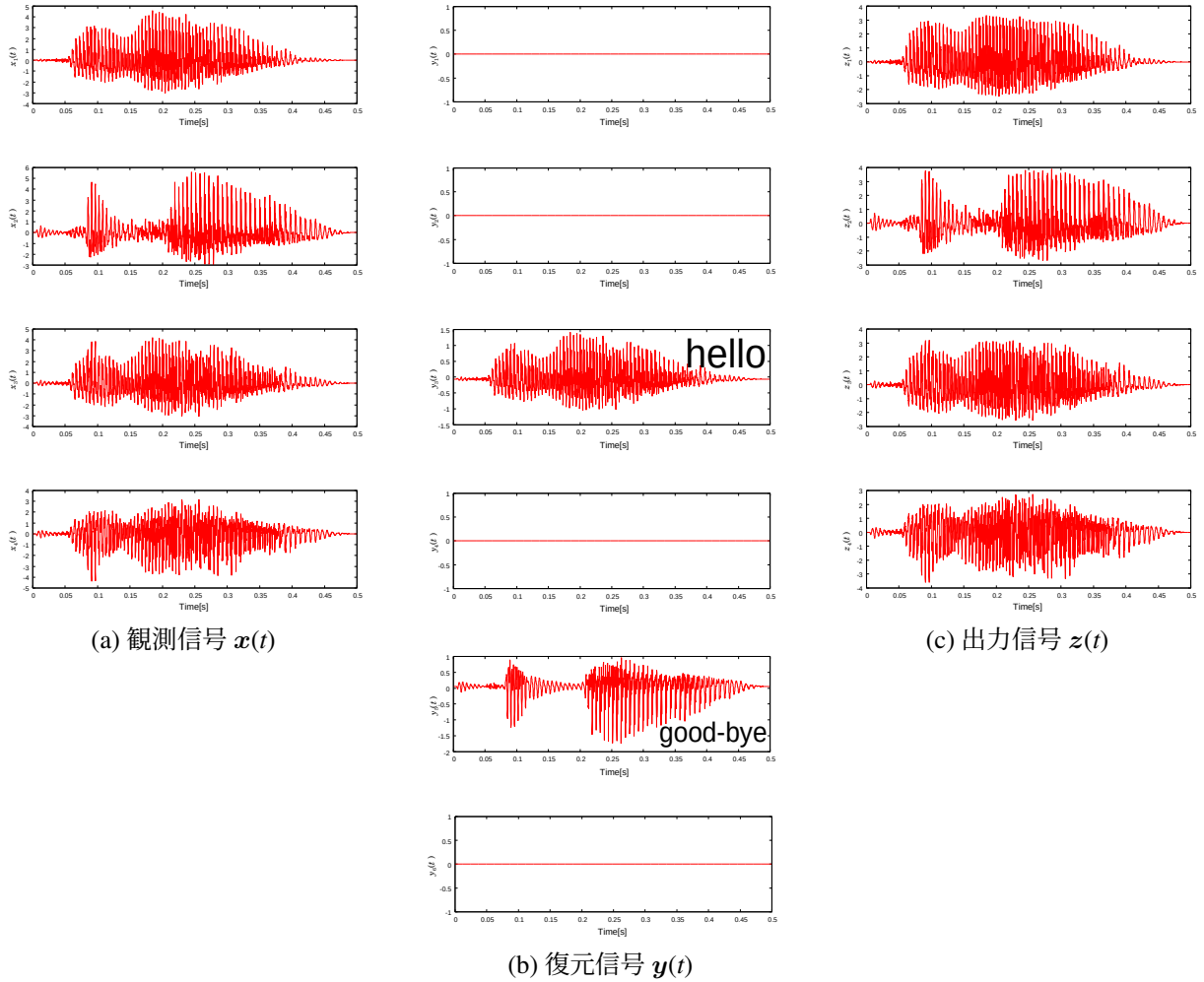


図 3. 学習終了時の各ニューロンの信号波形

表 1. 源信号 $s(t)$ と復元信号 $y(t)$ の類似度

		源信号	
		$s_1(t)$	$s_2(t)$
復元信号	$y_1(t)$	—	—
	$y_2(t)$	—	—
	$y_3(t)$	0.957	0.248
	$y_4(t)$	—	—
	$y_5(t)$	0.032	-0.999
	$y_6(t)$	—	—

という判断をすることができた。しかし、本来の ICA における源信号は未知であるため、パラメータの設定方法については、検討の余地がある。

4. 分離行列と混合行列の関係性

これまで、本手法における結合荷重 V と混合行列 A の関係性については、詳しく調査されていなかったようである。**3** の結果で、復元信号が源信号のみとなったことにより、本手法における分離行列 V と混合行列 A の関係性を解析的に調査することが可能になった。

一般的な ICA では、未知の混合行列 A の逆行列を求めることで、源信号 $s(t)$ を推定している。ここで、**3** で用いた混合行列 A は以下のとおりである。

$$A = \begin{pmatrix} 0.93 & 0.23 \\ 0.15 & 0.98 \\ 0.66 & 0.48 \\ -0.52 & -0.60 \end{pmatrix} \quad (11)$$

式 (11) は正方行列ではないため、通常の ICA における分離行列 V は混合行列 A の一般化逆行列に対応している。

一方、本手法における分離行列は入力層と隠れ層の間の結合荷重 V として獲得される。図 3 のときのネットワークにおいて、信号分離に関わっている部分は隠れ層の 3 番目のニューロンへ結合している結合荷重 $v_{31}, v_{32}, v_{33}, v_{34}$ 、隠れ層の 5 番目のニューロンへ結合している結合荷重 $v_{51}, v_{52}, v_{53}, v_{54}$ であるため、分離行列 V として該当する部分だけを書き出すと以下ようになる。

$$V = \begin{pmatrix} -0.054 & -0.146 & 0.232 & -0.240 \\ -0.044 & -0.390 & 0.045 & -0.109 \end{pmatrix} \quad (12)$$

ここで、例えば文献 [7] によると、一般化逆行列 G は以下のような特徴があると記述されている。

- ある行列 A に対して、一般化逆行列 G は一意には定まらない。
- 一般化逆行列 G と元の行列 A の積 GA を計算すると、単位行列に近い行列に成る。
- AGA を計算すると、元の行列 A に戻る。

これらの特徴から、分離行列 V と混合行列 A の関係性

表 2. 混合行列 A を基準とした時の $AV'A$ の変化分

		列数	
		1 列目	2 列目
行 数	1 行目	1%	57%
	2 行目	33%	26%
	3 行目	6%	21%
	4 行目	8%	3%

を調べる。

まず, VA を求めたところ, 次のようになった。

$$VA = \begin{pmatrix} 0.21 & 0.10 \\ -0.05 & -0.31 \end{pmatrix} \quad (13)$$

ICA では復元信号の振幅に任意性を認めるため, 式 (13) の (1,1) 成分が「1.00」となるように定数倍し, これを V' と置き, $AV'A$ を計算したところ以下のようになった。

$$AV'A = \begin{pmatrix} 0.92 & 0.10 \\ 0.10 & -1.33 \\ 0.62 & -0.38 \\ -0.48 & 0.62 \end{pmatrix} \quad (14)$$

式 (14) を見ると, 式 (11) の混合行列 A と一致しておらず, 一般化逆行列には完全には対応していないようである。ここで, 式 (11) を基準として, 式 (14) の変化分を計算した。その結果を表 2 に示す。

表 2 を見ると, 8 つの係数の内, 6 つの係数が 30% 以下の変動に収まっており, 獲得した各分離行列は混合行列 A に近い行列であると言える。なお, 式 (14) の二列目のほとんどの要素が, 式 (11) の各要素に対して符号が反転している。これは, 図 3 における隠れ層の 5 番目のニューロンが源信号 $s_2(t)$ の反転した波形を復元しており, $s_2(t)$ に対して高い負の相関を持っているためだと考えられる。

また, 3 におけるもうひとつの分離成功例についても同様の計算を行った結果, 同じように混合行列の一般化逆行列に近い行列であることがわかった。ただし, 両者は異なる行列である。よって, 本研究で獲得された 2 種類の分離行列は, それぞれ混合行列 A の一般化逆行列とは完全に対応していないものの, それに近い行列になっている。一般化逆行列の定義から, 通常の ICA における分離行列も一意には定まらないということが言えるため, その点に関する本研究との対応が確認できた。そして, 本手法における分離行列は唯一に定まらず, 多様性のある可能性が示唆された。

5. 考察

本手法では, 構造刈り込み学習の効果を強めるほど絞り込みが行われ, 隠れ層のニューロンの数が減っていく。その過程は, 今回の場合, 源信号及びそれらが混ざったもの (3 つ以上), そして源信号のみ (2 つ) であり, これに引き続いて急に学習不能に陥ると考えられる。これは, 効果を強めすぎると, 隠れ層に必要な情報が保持できなくなるためである。したがって, 適切なパラメータを設定する際には, 最初は影響を弱く設定し, ある程度の自由度を持たせることが重要である。そして, 学習の

傾向を確認しながら, 徐々に強めていく方法が有効であると考えられる。

また, 学習によって獲得される分離行列 V の関係性を検討したところ, 混合行列 A の純粋な一般化逆行列には対応していないことがわかった。これは, 今回獲得した復元信号は源信号にある程度高い相関を持っているが, 源信号と完全に一致する信号を復元しているわけではないためと考えられる。今回は, 相関係数の絶対値が「0.95」を超えていれば, 源信号を復元しているとみなし, 実際に音を聞いてみて源信号と似ている音声であることを確認していた。よって, 混合行列の一般化逆行列に完全に対応していなくても, ある程度近い行列であれば, 人間が聞いて源信号に近いと判断できる程の大きな相関を持った復元信号を獲得することが可能となると思われる。しかし, どの程度のずれまで許容できるかについては, まだ十分な検証が行えていないため, 追加の調査を進めていく必要がある。

6. おわりに

本稿では, 階層型ニューラルネットワークに構造刈り込み学習を適用し, ICA と同等の効果を得る手法を用いて, 源信号のみを復元する完全な信号分離を実現することができた。また, 学習によって獲得した分離行列と混合行列の関係性を, 一般化逆行列の定義から解析的に検討した。その結果, 本手法における分離行列は混合行列の一般化逆行列に完全には対応していないことが分かった。今後は, 学習に必要なパラメータの効率的な設定方法の検討や分離行列と混合行列のより詳細な関係性の調査が課題となっている。

参考文献

- [1] 村田昇: “入門 独立独立成分分析”, 東京電機大学出版局, 2004
- [2] Syozo Yasui: “Convergence suppression and divergence facilitation: Minimum and joint use of hidden units by multiple outputs”, *Neural Networks*, Vol.10, pp.353-367, 1997
- [3] 陸門将也, 和久屋 寛, 伊藤秀昭, 福本尚生, 古川達也: “構造刈り込み学習により不要ニューロンを削減するニューラルネットワークを用いた信号源分離”, 第 1 回日本知能情報ファジィ学会九州支部学生会研究発表会予稿集, A202, pp.9-10, 2011
- [4] H.Wakuya, K.Kukihira, H.itoh, H.Fukumoto and T.Furukawa: “An application to vocal signal decomposition by a neural network with structural pruning”, *ICIC Express Letters*, Vol.8, pp.999-1005, 2014
- [5] 和久屋 寛, 久木原健介, 伊藤秀昭, 福本尚生, 古川達也: “構造刈り込み学習を用いた音声信号分離の試み”, 第 23 回ソフトサイエンス・ワークショップ講演論文集, 2A-2, pp.26-29, 2013
- [6] <http://www35.atwiki.jp/softtalk/> (SofTalk のホームページ)
- [7] 新井康平: “独習 応用線形代数学 –基礎から一般化逆行列の理工学的应用まで–”, 近代科学社, 2006

連絡先

和久屋 寛
佐賀大学大学院工学系研究科電気電子工学専攻
佐賀大学人工頭脳医科学シナジー研究所
E-mail:wakuya@cc.saga-u.ac.jp

図示化技術を用いた複数 DEA モデルからなる 分析結果の解析支援手法

Decision support system for interpreting the DEA's results using visualization technology

○宮西 昌吾, 青木 真吾

○Syogo Miyanishi, Shingo Aoki

広島工業大学

Hiroshima Institute of Technology

Abstract: Data Envelopment Analysis (DEA) is a well-known technique for relatively evaluating the performance of Decision Making Units (DMUs). This study proposed the decision support system for interpreting the DEA's results using visualization technology. Through the case it's possible to illustrate the DEA's results with a few input/output data, this study picked out following important matters: (1) efficiency value, (2) diversity index, (3) index about the DMU which competes, and (4) the score of the feature the DMU has. Proposed method has consisted to express these indexes using a bubble chart. Numerical study has shown how the proposed method works.

1. はじめに

事業体(DMU: Decision Making Unit)の経営効率性を分析するための手法の一つとして包絡分析法(DEA: Data Envelopment Analysis)がある[1]. この手法のユニークな特徴は, 評価対象となる DMU の生産性測定に際し, 評価対象が属するグループから導き出されるパレート最適なライン(効率的フロンティア)との相対比較により, その評価値が算出されるところにある. 従って, 同じグループに属する DMU であっても, 相対比較される効率的フロンティア上のポイントは異なり, 各々の DMU の特徴を加味した評価を行うことができる. この特徴により, 近年では, 多面的な活動を行う企業や公共団体などの評価等に幅広く適用されている[2, 3].

DEA に関する研究論文に着目すると, 従来モデルを拡張した数理モデルの開発や, 適用する業界に応じて, 領域限定法といった DEA 研究者が持つ経験や勘が必要な方法論を組み込んだ実証研究例が多く見られる[4]. その一方で, DEA を利用する立場に立った, 結果を解釈するためのユーザビリティに関する研究例は数が少ない状況にある. ユーザビリティに着目した研究としては, 分析者の直観的な判断を支援するための方法として, 図示化技術を組み込んだものがある[5, 6]. 先行研究[5]は, 主成分分析を利用することで 2 次元マップ上に DMU の関係を図示化するものである. この成果により, 分析に用

いる入出力項目数が大きくなった場合の図示化をすることを可能としたが, 主成分分析の性質上, 項目数が増えすぎると, 情報量の損失が大きくなってしまい, 正確に DMU 間の関連性を図示化することが困難になってしまうという課題があった. 一方, 先行研究[6]は, 解析結果の全てを図示化するという方法論ではなく, 分析者が望むであろう 2 つの項目, 各事業体の評価値(効率値)と競合する DMU の数に絞って 2 次元マップ上に表現するというものであった. この方法は, また, 入出力の階層構造に着目したものであり, どの項目に特色があり, その項目に特化した場合に競合する DMU がどれくらいの数があるかを分析したものである. 従って図示化されるものは, 調査したい DMU の項目ごとの評価値と競合する DMU の数を図示化したもので, 全ての DMU 間の関係性を 2 次元マップ上に表現したものではない.

そこで, 本研究では, DEA を利用する分析者の立場に立ち, 分析結果の解釈を直観的に判断できる図示化技術による支援を行うフレームワークを構築することを目的とし, 先行研究において効果のあった, (1) すべての DMU 間の関連性を図示化すること, そして, 図示化に際しては, (2) 主成分分析等により, 情報の圧縮をするのではなく, DEA の持つ解析結果から得られる効果的な指標に絞り込むこと, を同時に達成することを目指す.

2. 包絡分析法 (DEA)

図示化を行う上でベースとなる DEA の紹介を行い、図示化するために必要な指標は何かを論じる。

2.1 DEA 概要と効率的フロンティア

DEA は、入出力項目数が極めて小さい場合は、その活動を 2 次元マップ上にプロットし、図から解を求めることが可能である。そこで、いま、7 つの DMU が存在し、それぞれが 1 つの入力項目(x)と 2 つの出力項目(y_1, y_2)によってその活動が特徴付けられたとする。このとき、 x 軸を y_1/x 、 y 軸を y_2/x とおくと、図 1 のように、それらの活動をプロットすることができる。少ない入力で、多くの出力を得るという効率性の概念に立つと、図 1 において、活動 A、B、C はパレート最適な活動と見なすことができる。このとき、DEA では、これらの活動を結んだラインを効率的フロンティアと呼ぶ。このフロンティア上にある活動は等しく効率的な活動であると判断され、100%という効率値(評価値)が与えられる。それ以外の活動は、非効率的な活動であると判断され、その評価値はフロンティアとの離れ具合によって測定され、100%未満の数字で与えられる。

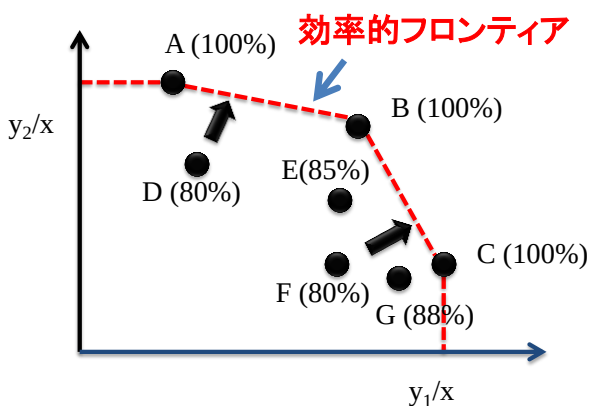


図 1 効率的フロンティア

2.2 DEA 結果を解釈するための要件

本研究では、図 1 で示されるような、極めて少ない入出力関係における図解から、DEA 結果を解釈するために必要な要件について洗い出すことにする。

(1) 効率値(評価値)：

図 1 において、各 DMU の横に括弧内で表現され

ている 0~1 までの数値である。これは、DEA 適用において、最も基本的な値となるもので、特に効率的な活動(100%の効率値)をしているか、あるいは、改善の余地があるとされる非効率的な活動(100 未満の効率値)かを知ることは重要な要件である。

(2) 多様性に関する指標：

DEA 特有のものとして、図 1 における活動 A、B、C を同等の活動と見なすことがある。活動 A は y_2 項目に特徴があり、活動 C は y_1 項目に特徴があり、そして、活動 B はそのバランスに優れているという特徴を持っている。この関係性は、図 1 を見ると明らかである。しかしながら、多入力-多出力といった項目が多くなった場合は、表ベースの解析結果から読み取る必要がある。

そこで、本研究では、これを多様性という指標に置きなおして、その関係性を表現する必要があると考える。具体的には、活動 A や活動 C に関しては、ある項目に特化したものであり、全ての項目に対して優れている訳ではない特徴的な DMU であると言える。一部の面だけが評価されるこれらの DMU は多様性が低いと考え、逆に、活動 B のように、どの項目においても、優れた面(y_1, y_2)をもつ場合に多様性が高いと考える。

(3) 競合他社に関する指標：

図 1 における活動 D と活動 F に着目する。この二つの活動は、共に効率値は 80%という値を持ち、改善の余地があることがわかる。しかしながら、それぞれが目指すべき方向は、効率的フロンティアに近い点が合理的であると考えられるため、異なるものとなる。活動 D は活動 A の方向に、活動 F は活動 C の方向へというものである。また、図解より明らかな点は、活動 Fの方が、効率的な活動へと改善活動をした場合に、その近辺に活動 E や活動 G といった競合他社が多いのに対して、活動 D に対しては存在しないということである。

本研究では、図解において把握できるこの情報は、改善方向を模索する分析者において重要であると考え、競合他社の指標を組み込むことが必要であると

考える.

(4) 特徴的な項目の抽出:

(3)の競合他社に関する指標で述べたように、図 1 より、各活動がどちらの軸によっているかを把握することが可能である.

この情報も、分析者にとって調査したい DMU が、どの面で優れているかを把握するための重要なものであると考えることができる.そこで、本研究では、この情報を組み込むために、次のような性質の利用を考える.いま、入出力項目の中から、 y_2 というデータ項目を取り除いたとする.すると、全ての活動を x 軸に射影し、右側にいるほど評価値が高い活動と判断されることになる.このとき、活動 C は変わらず効率的な活動と判断されるが、活動 A は効率的な活動から、非効率的な活動へと著しく変化することになる.すなわち、項目を取り除いていくという感度分析手法を利用し、その差を見ることで、調査したい DMU が、どの項目に特徴を持つのかを判断できると考える.

3. 各指標の算出方法

先に述べた要件に対しての計算の方法について述べる.

3.1 効率性に関する指標の測定

数学的な記述のために、 n 個の DMU が存在し、それぞれ m 個の入力と s 個の出力でその活動が特徴付けられるとする.このとき、 k 番目の DMU の入力データを $(x_{1k}, x_{2k}, \dots, x_{mk})$ とし、出力データを $(y_{1k}, y_{2k}, \dots, y_{sk})$ とおく. DEA では、数々の効率性モデルが開発されているが、ここでは、最も基本的なモデルである CCR モデルと BCC モデルの測定法について述べる. k 番目の DMU の効率性は、以下の線形計画問題を解くことによって得ることができる.

$(v^{(k)}, u^{(k)})$ は k 番目の DMU を評価する際に用いられるウェイトを示しており、(1)式は、各 DMU の評価値が 100%を越えないという条件のもとで、 k 番目の DMU の評価値を最大化するウェイトを探索するモデルということができる.なお、 L, U はパラメ

ータであり、 $L=U=0$ であれば、CCR モデルとなり、 $L=U=1$ であれば BCC モデルとなる.

どちらのモデルとも、(1)式における目的関数値が効率性を示す値となり、0~1 の数値として表現される.

$$\begin{aligned} \text{Max} \quad & \sum_{r=1}^s u_r^{(k)} y_{rk} + L\sigma^{(k)+} - U\sigma^{(k)-} \\ \text{s.t.} \quad & -\sum_{i=1}^m v_i^{(k)} x_{ij} + \sum_{r=1}^s u_r^{(k)} y_{rj} + L\sigma^{(k)+} - U\sigma^{(k)-} \leq 0 \\ & (j=1, 2, \dots, n) \\ & \sum_{i=1}^m v_i^{(k)} x_{ik} = 1 \\ & v_i^{(k)} \geq 0, u_r^{(k)} \geq 0, \sigma^{(k)+} \geq 0, \sigma^{(k)-} \geq 0. \end{aligned} \quad (1)$$

3.2 多様性に関する指標の測定

本研究では、多様性の測定に際し、多様性の高い DMU は、どのような評価基準でも高い評価値を得るという特色をもつ.逆に、多様性の低い DMU は、一部の面だけ、すなわち、他の評価基準で評価した場合は、その評価値が低下するという特色をもつと考え、クロス効率値[7]の方法論を適用する.(1)式を解き、得られるウェイトより、下記(2)式の関係式が最適性より必ず成立する.

$$\begin{aligned} & \left(\sum_{r=1}^s u_r^{(k)} y_{rk} + L\sigma^{(k)+} - U\sigma^{(k)-} \right) / \sum_{i=1}^m v_i^{(k)} x_{ik} \\ & \geq \left(\sum_{r=1}^s u_r^{(j)} y_{rk} + L\sigma^{(j)+} - U\sigma^{(j)-} \right) / \sum_{i=1}^m v_i^{(j)} x_{ik} \end{aligned} \quad (2)$$

(2)式における左辺は、 k 番目の DMU に対して、その評価値を最大にするウェイトによる生産性を示している.そして、右辺は、同じく k 番目の DMU に対して、 j 番目の DMU の評価値を最大にするウェイトでの生産性を示している.言い換えれば、 j 番目の DMU の評価基準に当てはめた場合の k 番目の DMU の評価値といえる.

すなわち、この左辺と右辺の差が大きければ大きいほど多様性が低いと判断でき、逆の場合、評価値が変動しなければ、どの評価基準でも高い評価値を得ることとなり、多様性が高いと判断できる.

そこで、本研究では、下記(3)式で示される標準偏差の値を多様性の指標として用いる.

Standard Deviation of

$$\left\{ \left(\sum_{r=1}^s u_r^{(j)} y_{rk} + L\sigma^{(j)+} - U\sigma^{(j)-} \right) / \sum_{i=1}^m v_i^{(j)} x_{ik} \ (j=1,2,\dots,n) \right\} \quad (3)$$

3.3 競合他社に関する指標の測定

改善方向に対する競合他社の多寡を把握するために、本研究では、調査したい DMU の評価基準で、他の DMU の生産性を求め、その値を比較したときに、調査したい DMU の上位に何個の DMU があるのか、その数を用いて競合他社に関する指標とする。

具体的には、k 番目の DMU の場合は、(4)式を満たす、すべての j の数である。この数が少ないほど、効率的フロンティアに改善されるまで、競合となる DMU が少ないことになる。

$$\left\{ \left(\sum_{r=1}^s u_r^{(k)} y_{rk} + L\sigma^{(k)+} - U\sigma^{(k)-} \right) / \sum_{i=1}^m v_i^{(k)} x_{ik} \right\} < \left\{ \left(\sum_{r=1}^s u_r^{(k)} y_{rj} + L\sigma^{(k)+} - U\sigma^{(k)-} \right) / \sum_{i=1}^m v_i^{(k)} x_{ij} \right\} \quad (4)$$

(j=1,2,...,n, j ≠ k)

3.4 特徴的な項目な抽出に関する数量化

各 DMU がどの面(項目)で、特徴を持っているかを数量化するために、感度分析を行い、その効率値の評価の差をもって測定するという方針をとる。

感度分析は、入出力項目の中から、一つとりのぞき、その際の効率値の差を見る方法である[8]。例えば、出力 1 という項目に特徴があるかを判断する場合は、(1)式において、出力 1 に対応するウェイト $u_1^{(k)}$ に着目し、 $u_1^{(k)} = 0$ として(1)式を計算する。

このとき、目的関数値が変化しなければ、その DMU を評価する際に、その出力 1 は計算に使用されなかったことになり、重要ではない項目ということができる。逆に、変化が起これば、その項目は評価に必要な項目、すなわち、その DMU を評価する際に外すことのできない重要な項目であると判定できる。本研究では、この変化量を見ることで、特徴的な項目であるかを判断する目安とする。

4. 数値実験

本研究で提案するアプローチの有効性を検証する

ために、サンプルデータを基に数値実験を行う。まず最初に、これまで述べた効率性、多様性、競合他社の 3 つの指標についての結果と、それらを可視化する方法を提示し、次に、感度分析による特徴的な項目を抜き出した時の変化の考察をする。

従って、本研究では、3.1 から 3.3 までの 3 つの指標を 2 次元マップ上にプロットし、3.4 の特徴的な項目の抽出については、図を見比べることで分析結果解釈の支援を行うものとする。

また、3.1 で述べた通り、生産性測定のための DEA モデルは数多く存在する。効率性の測定を CCR モデルと BCC モデルの 2 通り用意し、図解から 2 つのモデルの差が把握できるかについても検討する。

4.1 データセット

本研究で取り扱うサンプルデータの一部を表 1 にまとめる。50 の DMU が存在し、それらの活動が 3 つの入力(一人あたり人件費、労働分配率、自己資本比率)と 1 つの出力データ(今期経常利益)で特徴づけられるものと仮定する。

表 1 サンプルデータ

事業体	一人あたり 人件費	労働 分配率	自己資本 比率	今期 経常利益
A	1283	13	55	65.8
B	247	2	43	128
C	533	27	89	3
D	1186	21	94	12
E	437	24	89	3.4
F	553	56	61	2.08
G	836	20	80	1.43
H	1236	8	77	84.2
I	632	35	48	38.1
J	901	49	41	29.4
K	894	2	75	6.4
L	720	53	84	5.55
...	...	...	...	...
M	610	21	56	13.3
N	1221	33	77	2.21
O	766	44	81	10.5

4.2 3 つの指標の数量化

(1)式から(4)式までを解き、効率性に関する指標、多様性に関する指標、そして、競合他社に関する指標の計算結果を表 2 にまとめる。効率性に関する指標は 1 に近いほど効率的な活動をしていると判断できる。多様性に関する指標は、標準偏差の値を用い

ているため、0に近い方が高い多様性をもつという結果を示している。競合他社に関する指標は、自社より上位のもの数であるため、小さい方がライバルが少ないことを示している。なお、効率性の指標の値には、CCRで測定したものを採用している。

表2 3つの指標の値

事業体	効率性を示す指標	多様性を示す指標	競合他社を示す指標
A	0.294	0.058	4
B	1	0	0
C	0.011	0.002	40
D	0.036	0.005	17
E	0.015	0.002	24
F	0.010	0.002	41
G	0.005	0.001	48
H	0.310	0.052	3
I	0.190	0.040	7
J	0.128	0.031	9
K	0.050	0.005	7
L	0.019	0.004	28
...	...	...	...
M	0.063	0.011	14
N	0.007	0.001	47
O	0.036	0.007	18

4.3 3つの指標の図示化

本研究では、先にまとめた3つの指標をバブルチャートを用いることで図示化を行う。効率性については、バブルの大きさで、多様性の値をy軸に、そして、競合他社の値をx軸に表す。図示化したものを図2に示す。

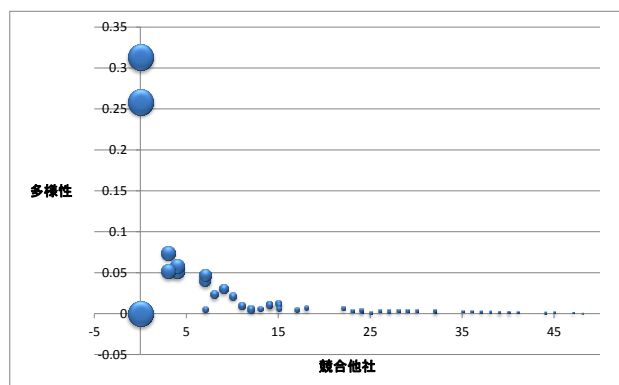


図2 指標図示化

図2において、横の軸では左にあるほど競合他社が少なく、自分より優れたDMUが少ないことが見

て取れる。また縦の軸では、下にあるほど他の評価軸で評価した時の差が少ないことが見て取れる。

図2より、左上のゾーンあるDMUはすべての事業体から相対的にみると、何かに特化した事業体が多いことが視覚的に見て取れる。左下のゾーンにいるDMUはバランスの良い堅牢な事業体であると見て取れる。右側のゾーンは生産効率値が低いため、競合他社が多くなり、多様性の値が低くなる傾向があることがわかる。

しかしながら、2次元上に図示化したにも関わらず右上のゾーンに属するDMU、多様性がなく、競合他社の数も多いというDMUがプロットされなかった。図示化における軸の範囲の設定や、多様性に関する指標の再考の必要性が課題として見つかった。

4.3 感度分析による特徴的な項目の抽出

図3は、元のサンプルデータから、3つ目の入力である自己資本比率の項目を抜き、再測定を行い、図示化を行ったものである。

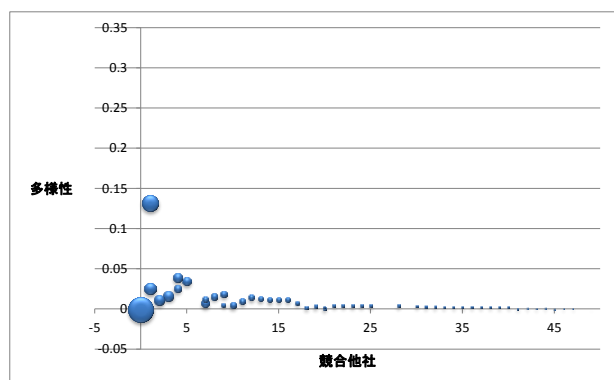


図3 自己資本比率を除いた図示化

図2と見比べると、左上に存在していた2つの事業体のバブルの大きさ（効率値）が小さくなっている、また全体的に多様性の値が大幅に下がったことが見て取れる。これは自己資本比率という項目がすべての事業体にとって、とても重要な項目であることが見てとれる。

4.4 CCRモデルとBCCモデルの変化

最後の検証ではCCRモデルからBCCモデルに移行することにより、今回のフレームワークでどのよ

うに変化するかを確認する。図 4 は BCC モデルを図示化したときのものである。

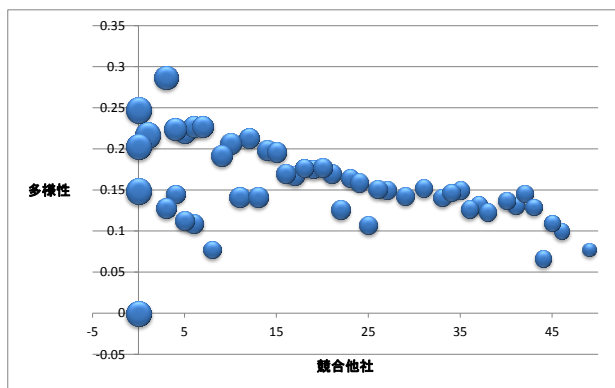


図 4 効率性を BCC モデルとした図示化

CCR モデルから BCC モデルに解析方法を移行したことにより、評価値がより高くなり全体的にバブルが大きくなることが確認された。また競合他社が 0 の事業体が新たに加わったことが見えてとれ、その競合他社が 0 の事業体は多様性の値が下がった。これは、条件が緩和され、苦手な分野での効率値が上がったことにより多様性の値が下がったためだと考えられる。そのほかの事業体では効率値が上がったため多様性の値が上昇した。BCC モデルは、CCR モデルよりも甘い評価をすることで知られているが、それだけでなく、多様性の指標にも大きな差が出ていることから、モデル間の差を把握するのにも役立つのではと考えられる。

5. おわりに

DMU 間の関連性を図示化することと DEA の持つ解析結果から得られる効果的な指標に絞り込むことを盛り込んで図示化し表示する方法論を提案した。図示化によって、同じように見える効率値 1 の事業体であっても多様性や競合他社という指標により、各 DMU がそれぞれどのような立場にあるのかが直感的に見て取れるようになった。また感度分析やモデルの変化などへの拡張性を持たすことにより、解析者がより効果的な指標に絞り込むことに成功した。

今後の展望としては、実データに当てはめた解析を行うことや、取り扱う DMU 数が増えた場合の可視化技術の検討などがあげられる。

参考文献

- [1] A. Charnes, W. W. Cooper and E. Rhodes: "Measuring the efficiency of decision making units", *European Journal of Operational Research*, Vol.2, pp429-444, 1978
- [2] N. Kamiyama: "Designing network topology using Data Envelopment Analysis", *Journal of the Operations Research Society of Japan*, Vol. 56, No. 3, pp199-220, 2013
- [3] S. Liu: "A fuzzy DEA/AR approach to the selection of flexible manufacturing systems", *Computers and Industrial Engineering*, Vol. 54, pp. 66-76, 2008
- [4] A. Parker and G. Tavares, "Evaluation of research in efficiency and productivity: A survey and analysis of the first 30 years of scholarly literature in DEA", *Journal of Socio-Economics Planning Science*, Vol. 42, No. 3, pp. 151-157, 2008
- [5] 青木真吾, 本多克宏: 変量選択を伴う主成分分析による包絡分析結果の視覚化, システム制御情報学会論文誌, Vol. 23, No. 9, pp. 235-225, 2010
- [6] 井上和重, 青木真吾, 福岡憲治, 辻洋: 包絡分析法における入出力項目の構造分析手法の提案, 電気学会論文誌 C, Vol. 134, No. 1, pp. 138-147, 2014
- [7] 枇々木規雄: DEA における修正クロス効率値を用いた評価法, *Journal of the Operations Research Society of Japan*, Vol. 41, No. 2, pp229-244, 1998
- [8] 内藤明夫, 青木真吾, 辻洋: 効率的フロンティアの設定に伴う DEA 感度分析, 電気学会論文誌 C, Vol. 130, No. 4, pp540-545, 2010

連絡先

広島工業大学 宮西昌吾

E-mail:m161404@cc.it-hiroshima.ac.jp

ファジィ回帰モデルを用いた生活習慣病の分析

Analysis of Lifestyle-related Disease using Fuzzy Regression Model

藪内 賢之

Yoshiyuki Yabuuchi

下関市立大学

Shimonoseki City University

Abstract: In 2012, 15.13% of the total fiscal medical care expenditure was for lifestyle-related health care costs, which was approximately 179 billion yen. Lifestyle-related diseases are not only the biggest factor in reducing healthy life expectancy but also have the most significant impact on the national medical care expenditure. In addition, lifestyle-related diseases can be prevented by moderate daily exercise, a well-balanced diet, and not smoking.

Our fuzzy robust regression model is used to analyze factors concerned with life-style disease.

1 はじめに

主要健康保険組合では、2012年度医療費の約15%が生活習慣病関係で支払われた[1]。生活習慣病は日常生活の適度な運動、バランスの取れた食生活、禁煙により防ぐことができる[5]。すなわち、生活習慣に関する食生活および日常生活の関係を明らかにすることが出来れば、生活習慣病の予防が可能である。

このため、本研究では厚生労働省の患者調査および国民栄養調査結果を用い、生活習慣病予防に挙げる適度な運動、食生活および患者数の関係を分析する。また、厚生労働省が指摘する適度な運動については体型を示す身長、体重、BMIなど、バランスが取れた食生活については摂取する食品、生活習慣の結果は外来患者数を用いる[12]。

また、ファジィ回帰モデルを用いると生活習慣と外来患者数の特徴を把握しやすいため、ファジィロバスト回帰モデルを用いて分析を行う[11]。

2 都道府県別医療費の分析

総務省統計局人口推計によると、2014年7月1日現在の総人口は1,2713万人である。年少人口(15歳未満)の人口は1,628万人で前年に比べ16万人減少、生産年齢人口(15歳以上65歳未満)は7,812万人で113万人減少した。一方、老年人口(65歳以上)の人口は3,272万人で前年と比べて109万人増加した。

図1に1950年から2011年までの日本人口の年齢区分別の推移を示す。65歳以上の老年人口は増加、15歳未満の年少人口が減少し続け、1997年に老年人口が年少人口を上回った。これにより社会保障のあり方も変

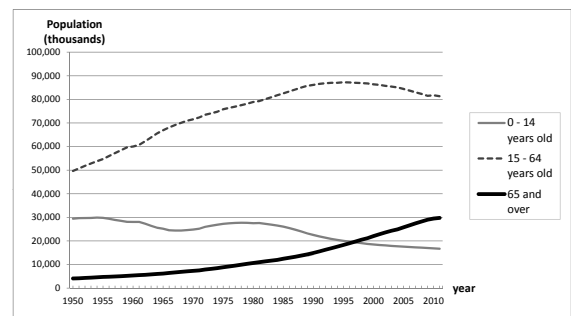


図1: 1950年から2011年までの日本人口の年齢区分別推移

化しつつあるが、十分に対応できているとは言えない。このような状況にあって、社会全体として健康に配慮するようになりつつある。厚生労働省によると、2011年度に病気やけがの治療で全国の医療機関に支払われた医療費の総額が、前年度比11,648億円増の385,850億円であった。国民1人当たりでは9,700円増の301,900円である。いずれも5年連続の上昇である。

ここで、平成24年度医療費総額約11,836億円のうち、生活習慣病関連は約1,791億円で全体の15.13%である。生活習慣病1人当たり医療費は14,407円である。

3 患者数と生活習慣の分析

2013年10月1日時点の65歳以上人口は約24%であり、国連の定義で日本は超高齢社会となる。生涯健康であり続けることは医療費とは別に、生活の質を低下させないためにも重要なことである。

本論文では、生活習慣病に関わる要因とその結果である外来患者数をファジィロバスト回帰モデルを用いて分析する。ここで、生活習慣は状態が継続し続けることであると考え、1995年から1999年の都道府県別データを説明変量、1999年の患者数を目的変量とし、回帰分析を行う[2, 3, 4]。標準化後の変数選択により、目的変量 Y として外来患者数、説明変量としてBMIを X_1 、米類摂取量を X_2 、肉類摂取量を X_3 、牛乳および乳製品を X_4 を用いることにした。ここでは最小自乗法による回帰モデル Y_S およびファジィロバスト回帰モデル Y_R を用いる。

得られた最小自乗モデル Y_S は次の通りである。

$$Y_S = -0.36X_1 - 0.39X_2 + 0.21X_3 + 0.18X_4$$

係数から、BMIと米類摂取量は患者数を減少させている。BMIと心臓疾患発症率はU字になり、25以上で肥満と循環器疾患のリスクが大きくなる[7, 10]。加えて肉類、牛乳および乳製品摂取量の増加が外来患者増加に繋がる。日本の主食は米であるが、肉や乳製品などのサイドディッシュから動物性タンパク質摂取量が増えた[7, 8]。このことを考慮すると、得られた統計モデル Y_S は納得できる。

次に得られたファジィロバスト回帰モデル Y_R を示す。

$$Y_R = (0.36, 0.30, 0.39) + (-0.46, 0.43, 0.32)X_1 \\ + (-0.30, 0.28, 0.34)X_2 + (0.27, 0.37, 0.30)X_3 \\ + (0.21, 0.31, 0.29)X_4$$

係数の中心は統計モデルと同じであるが、患者数への影響は米類摂取量よりもBMIの方が大きい。 X_3 は統計モデルと比べて大きな値である。肉に含まれるLカルニチンは動脈硬化や循環器疾患の原因となる[6]。また、日本人に対する調査では、乳製品から摂取されるカルシウムが脳梗塞や脳梗塞のリスクを低下させ、虚性精神疾患のリスクを増加させる[9]。

以上より、BMI、次いで米類摂取量が外来患者減少に影響することがわかった。そして肉類摂取量、牛乳および乳製品摂取量の多い都道府県は外来患者数が多い。

すなわち、適度な運動と日本的食事を食べる都道府県は外来患者数が少ない。また、肉類や牛乳などの洋食材の消費が多い都道府県は外来患者数が多い。

4 おわりに

本研究では生活習慣病を把握するため、外来患者を目的変量として回帰分析を行った。その際、説明変量にはBMI、米類摂取量、肉類摂取量、牛乳および乳製品摂取量を用いた。

分析結果をまとめると、一般的に言われるように和食を食べ、健康に心がけることが重要である。すなわち、和食が健康に繋がり、肉、牛乳、糖分などの摂りすぎは健康を害するのである。

参考文献

- [1] 健康保険組合連合会, 平成 24 年度生活習慣病医療費動向に関する調査分析, 2013.
- [2] 厚生労働省, 国民栄養の現状『平成 11 年国民栄養調査結果』, 1999.
- [3] 厚生労働省, 平成 11 年患者調査, 1999.
- [4] 厚生労働省, 国民栄養調査, 1995-1999.
- [5] 厚生労働省, 平成 23 年度厚生労働省白書, 2011.
- [6] R. A. Koeth et al., Intestinal microbiota metabolism of l-carnitine, a nutrient in red meat, promotes atherosclerosis, *Nature Medicine*, Vol.19, No.5, pp.576-585, 2013.
- [7] A. Morita, et al., Dietary Reference Intakes for Japanese 2010: Lifestage, *Journal of Nutritional Science and Vitaminology*, Vol.59, No.Supplement, pp.S103-S109, 2013.
- [8] S. Sasaki, Dietary Reference Intakes for Japanese 2010: Basic Theories for the Development, *Journal of Nutritional Science and Vitaminology*, Vol.59, No. Supplement, pp.S9-S17, 2013.
- [9] M. Umesawa, et al., Dietary Calcium Intake and Risks of Stroke, Its Subtypes, and Coronary Heart Disease in Japanese: The JPHC Study Cohort I, *Stroke*, Vol.39, Issue 9, pp.2449-2456, 2008.
- [10] G. Whitlock, et al., Body-mass index and cause-specific mortality in 900 000 adults: collaborative analyses of 57 prospective studies, *The Lancet*, Vol. 373, Issue 9669, pp.1083-1096, 2009.
- [11] Y. Yabuuchi and J. Watada, Japanese Economic Analysis by Possibilistic Regression Model which built through Possibility Maximization, *Journal of Advanced Computational Intelligence and Intelligent Informatics*, Vol.16, No.5, pp.576-580, 2012.
- [12] Y. Yabuuchi, T. Kawaura and J. Watada, Analysis of Medical Care Expenditure by Japanese Prefecture using Fuzzy Robust Regression Model, *Proceedings of KES International Conference on Innovation in Medicine and Healthcare*, 2014.

連絡先

藪内賢之

下関市立大学

〒751-8510 山口県下関市大学町 2-1-1

TEL&FAX: 083-254-8641

Email: yabuuchi@shimonoseki-cu.ac.jp

ファジィ理論に基づく経営戦略分析

Strategic Analysis based on Fuzzy Theory

○古殿幸雄¹, 陳法恩², 賈志聖²

○Yukio Kodono¹, Faen Chen², Chih Sheng Chia²

¹大阪国際大学

¹Osaka International University

²大阪国際大学大学院

²Graduate School of Osaka International University

Abstract: The strategy of business management, which is an external and internal environment analysis performed under the corporate philosophy, vision and then we make a comparison for evaluation of the present management strategy. Then, a new management strategy is developed, and to perform its management strategy. In order to manage the strategy making process, an external and internal environment analysis will be called the strategic analysis. In this paper, we apply the fuzzy theory to the strategic analysis, and discuss the possibility of more effective of management strategy.

1. まえがき

企業における経営戦略の位置づけは、図1のように表すことができる。

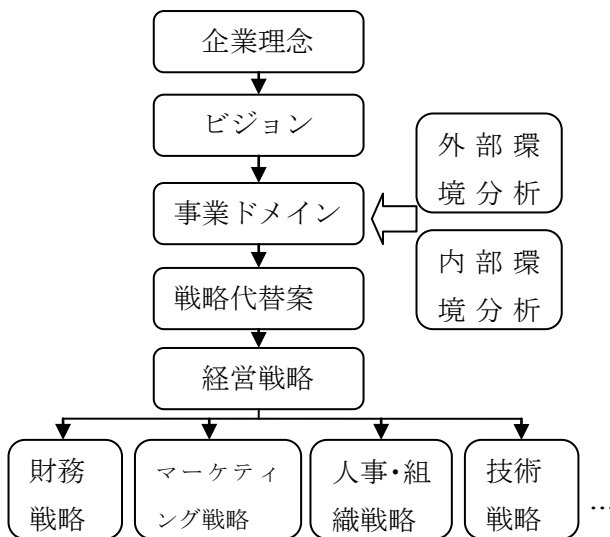


図1. 経営戦略の位置付け

図1のように、企業経営における戦略については、企業理念やビジョンの下、外部環境分析や内部環境分析を行い、現在の経営戦略を評価し、新しい経営戦略を策定して、その経営戦略を実行することになる。本研究では、経営戦略策定のための外部環境分析と内部環境分析および両者の混合分析を、経営戦略分析と呼ぶことにする。

さて、外部環境分析や内部環境分析の分析方法は、種々提案されているが、これらは主に二つの分析方法に分けることができる。一つは定量分析であり、

もう一つは定性分析である。

経営戦略における定量分析では、財務データや従業員数、販売数量などの具体的な数値データのもとで、企業の経営状況などを分析する。また、経営戦略における定性分析では、主観性やあいまいさなどの要因を含めて、企業の経営状況などを分析する。

経営戦略における代表的な定性分析として、外部環境分析では、PEST分析やポーターのファイブフォース分析[1]などがあり、内部環境分析では、バーニーのリソース・ベースト・ビュー[2]などがある。そして、通常は、外部環境分析と内部環境分析からSWOT分析が行われ、企業の強み、弱み、機会脅威を明らかにして、経営戦略の策定が行われる。

本研究では、定性的な経営戦略分析に、ファジィ理論を適用し、より有効な経営戦略分析の可能性について検討を行う。

先行研究としては、VRIO分析に対して、ファジィ概念が導入されたファジィVRIO分析が提案され[3]、中国のBYD自動車と奇瑞自動車に対して、ファジィVRIO分析とSWOT分析を行い、BYD自動車[3]と奇瑞自動車[4]の経営戦略の策定が示唆されている。また、SWOT分析に対して、ファジィ概念が導入されたファジィSWOT分析が提案されている[5]。SWOT分析は、企業の外部環境（機会と脅威）と内部環境（強みと弱み）を分析することで、企業の強みと弱みを把握し、企業の経営戦略を評価する。ファジィSWOT分析では、内部環境要因と外部環境要因の各々にファジィ概念を導入して、メンバシ

ップ関数で表現する。そして、三次元の SWOT 行列を用いて分析が行われる。

本研究では、VRIO 分析に対して、ファジィ概念が導入されたファジィ VRIO 分析を中心に検討することで、経営戦略策定におけるファジィ理論の有用性について明らかにする。なお、本研究のスタンスは、複雑な計算式や複雑なモデルを用いずに、できるだけシンプルな形での分析や問題解決への寄与である。

2. ファジィ VRIO 分析

ここでは、内部環境分析における定性分析の手法として、バーニーのリソース・ベースト・ビューを行う際の具体的なフレームワークである VRIO フレームワーク[2]に着目する。VRIO フレームワークにおいては、企業内部の資源に着目し、“価値があるか” “希少か” “模倣コストは大きい” “組織体制は適切か” の四項目に区分することで、競争優位をもたらすかについて判断する。このとき、各項目について、Yes, No という明確な区分を用いて、競争優位の判断が導かれる。ところが、企業経営戦略において、Yes, No という明確な区分は、相応の理由が無い限り難しいと考えられる。

そこで、VRIO 分析にファジィ概念を導入し、ファジィ VRIO 分析を提案することにより、明確に Yes, No と区分しなくても、その境界をあいまいなまま取り扱うことで、幅を持たせた経営状況の分析が可能になる。

表 1 は、VRIO フレームワークである。VRIO 分析は、企業の経営資源を分析するために用いられる。VRIO 分析のフレームワークは、Value（価値）、Rarity（希少性）、Imitability（模倣困難性）、Organization（組織）の四つに区分されており、その区分ごとに分析を行うことで、企業の経営資源が、競争優位をどれだけ持っているのかを把握できる。

表 1. VRIO 分析

価値があるか	希少か	模倣コストは大きい	組織体制は適切か	競争優位の意味合い
No	—	—	No	競争劣位
Yes	No	—	↑	競争均等
Yes	Yes	No	↓	一時的競争優位
Yes	Yes	Yes	Yes	持続的競争優位

そこで、ファジィ VRIO 分析を行う際の 4 つの問いに関して、整理しておく[2]。

①価値について

企業の製品やサービスなど様々な経営資源は、その企業に機会や脅威のような影響をもたらせるかどうかを検討する。

②希少性について

企業の製品やサービスなど様々な経営資源を有効に活用しているのは、ごくわずかな少数企業であるかどうかを検討する。

③模倣困難性について

製品やサービスなどの様々な経営資源を保有していない企業は、その経営資源を手に入れるまたは研究開発することで、易に保有できるかどうかを検討する。

④組織体制は適切かについて

企業の製品やサービスなど様々な経営資源（経営陣営の交代を含む）を活用するに当たって、企業の方針や経営手法などによって、証されているかどうかを検討する。

VRIO フレームワークにおいては、“価値があるか”、“希少か”、“模倣コストは大きい”、“組織体制は適切か” という 4 つの問いがある。これらの 4 つの問いに対して、Yes, No と答えることによって、企業の経営資源が競争優位をどれだけ持っているのかを判断する。しかし、4 つの問いに Yes, No とはっきりと区分するよりは、どちらかと言えば Yes であるとか、No に近いなどの答えが出てくるであろう。そこで、バーニーの VRIO フレームワークをもとに、ファジィ概念を導入し、新たなファジィ VRIO ルールが提案されている[3]。

ファジィ VRIO ルール

①IF “価値がない” and “組織体制は適切でない”
THEN “競争優位の意味合い” は “競争劣位”

②IF “価値がない” and “組織体制は適切である”
THEN “競争優位の意味合い” は “やや競争劣位”

③IF “価値がある” and “希少でない” and “組織体制は適切でない”

THEN “競争優位の意味合い” は “競争均等”

④IF “価値がある” and “希少でない” and “組織

体制は適切である”

THEN “競争優位の意味合い” は “一時的競争均等”

⑤IF “価値がある” and “稀少である” and “模倣コストがかからない” and “組織体制は適切でない”

THEN “競争優位の意味合い” は “やや一時的競争優位”

⑥IF “価値がある” and “稀少である” and “模倣コストがかからない” and “組織体制は適切である”

THEN “競争優位の意味合い” は “一時的競争優位”

⑦IF “価値がある” and “稀少である” and “模倣コストがかかる” and “組織体制は適切でない”

THEN “競争優位の意味合い” は “やや持続的競争優位”

⑧IF “価値がある” and “稀少である” and “模倣コストがかかる” and “組織体制は適切である”

THEN “競争優位の意味合い” は “持続的競争優位”

ファジィ VRIO ルールにおいて、前件部および後件部は、ファジィ集合で表現する。

3. ファジィ VRIO 分析の応用事例

ここでは、ファジィ VRIO 分析を用いて、中国の BYD 自動車と奇瑞自動車を分析した結果について述べる。ファジィ VRIO 分析を行うための BYD 自動車の4つの問いに対する定性分析の内容については、文献[3]にて報告され、奇瑞自動車については、文献[4]にて報告されているので、ここでは結果のみを示すことにする。

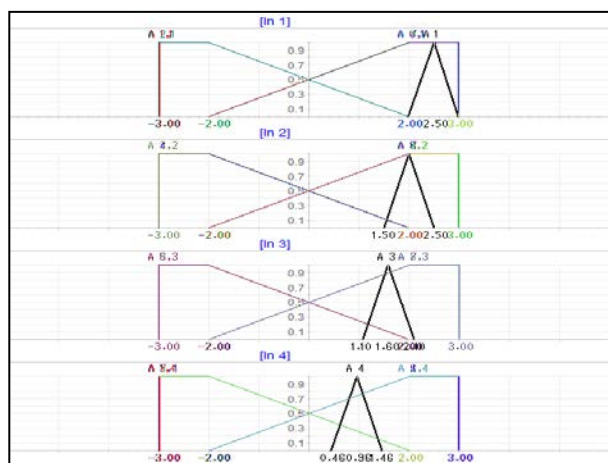


図 2. BYD 自動車に対するファジィ VRIO 分析
—前件部—

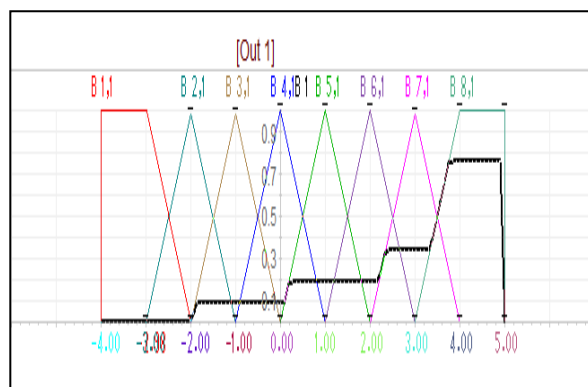


図 3. BYD 自動車に対するファジィ VRIO 分析
—後件部—

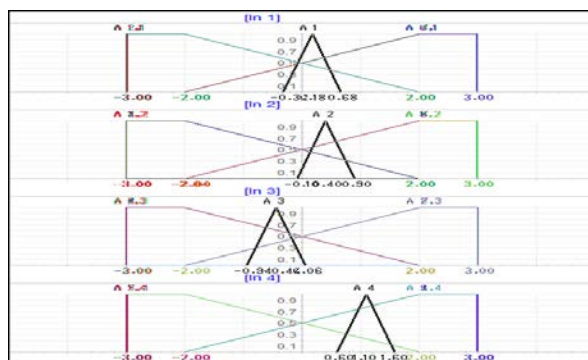


図 4. 奇瑞自動車に対するファジィ VRIO 分析
—前件部—

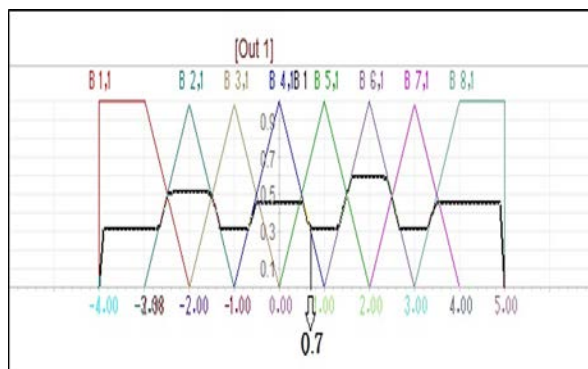


図 5. 奇瑞自動車に対するファジィ VRIO 分析
—後件部—

以上の結果より、“競争優位の意味合い” は、BYD 自動車は、“やや持続的競争優位” となり、奇瑞自動車は、“やや一時的な競争優位” となった。非ファジィ化の方法としては、min-max 重心法を用いている。

なお、通常の VRIO 分析であれば、表 1 を用いることになり、4つの項目に対して、仮にファジィ入力値の中心値が、Yes の領域にあれば Yes と判断し、

No の領域にあれば No と判断した場合、BYD 自動車は、4つの項目全てに対して Yes となるため、“競争優位の意味合い”は、“持続的競争優位”となり、奇瑞自動車は、4つの項目に対して順に Yes, Yes, No, Yes となるため、“一時的競争優位”となる。

したがって、ファジィ VRIO 分析を用いる方が、きめの細かい有用な経営戦略分析を行えることになる。

4. ファジィ VRIO 分析の評価

ファジィ VRIO 分析で得られた結果の妥当性について検討する。ここでは、紙面の都合上販売台数の推移のみを掲載するが、財務データなどを検討した結果については、発表時に行う。

図6は、BYD 自動車の販売台数の推移である。

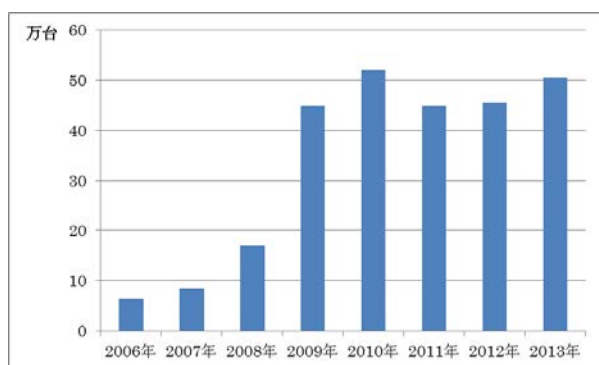


図 6. BYD 自動車の販売台数の推移

BYD 自動車は、2009 年から急激に販売台数を増やし、その後安定的に販売台数を維持してきている。また、BYD 自動車の営業収益は、2011 年～2013 年の間、毎年増え続けている。BYD 自動車は、2014 年に新エネルギー自動車に力を入れる方向であり、2014 年に入って、新エネルギー自動車の受注台数も大きく伸びている。

これら BYD 自動車の近年の販売台数と新エネルギー自動車の好調な販売、収益状況から考えて、現在の BYD 自動車の経営戦略は、やや持続的な競争優位を保っていると言える。

図7は、奇瑞自動車の販売台数の推移である。

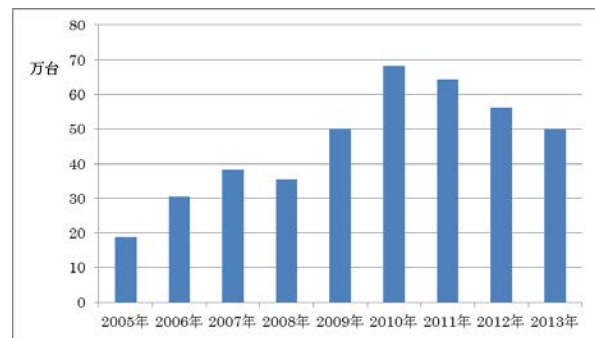


図 7. 奇瑞自動車の販売台数の推移

奇瑞自動車の販売台数は、2010 年をピークに年々減少傾向である。また、営業利益は、2012 年まで赤字が続き、2013 年に黒字に転化された。黒字に転じた一番大きな要因は、多ブランド経営を一本化することによって、車種の開発や車一台当たりの利益などを大きく変えたことである。

奇瑞自動車の近年の販売台数と収益状況から考えて、現在の奇瑞自動車の経営は、不安定な状況が続いており、“やや一時的な競争優位”と言える。

以上のことから、ファジィ VRIO 分析の結果は、より状況を的確に表現していると考えられる。

参考文献

- [1] M.E.ポーター（土岐坤・中辻萬治・服部照夫訳）：『競争の戦略』ダイヤモンド社，2005.
- [2] ジェイ・B・バーニー（岡田正大訳）：『企業戦略論【上】』ダイヤモンド社，2003.
- [3] 陳法恩・古殿幸雄：BYD 自動車のファジィ VRIO 分析，大阪国際大学紀要国際研究論叢，第 27 巻，第 3 号，pp.11-29,2014.
- [4] Faen Chen & Yukio Kodono：“Fuzzy VRIO and SWOT Analysis of Chery Automobile”，Journal of Advanced Computational Intelligence and Intelligent Informatics, Vol.18, No.3, pp.429-434, 2014.
- [5] S.Ghazinoory, A.E.Zadeh & A.Memariani: “Fuzzy SWOT analysis”, Jour. of Intelligent & Fuzzy Systems 18, pp.99-108, 2007.

連絡先

大阪国際大学 グローバルビジネス学部

古殿幸雄

E-mail:kodono@oiu.ac.jp

SNS から収集した地域・観光情報の分析

Analysis of Collecting Regional Information for Tourism retrieved from the SNS

○松本 義之¹, 藪内 賢之¹
○Yoshiyuki Matsumoto¹, Yoshiyuki Yabuuchi¹

¹ 下関市立大学
¹Shimonoseki City University

Abstract: Recently, people are collecting regional information from the Internet. Also, visitors are information dissemination using the SNS and blog the experience. It has become possible to disseminate information easily by the development of information technology. Describing personal experiences "word of mouth" information has become important. In this study, to collect local information from the big data that exists on the Internet. We focused on Twitter that can transmit information easily breaking property is high. Our objective is to acquire knowledge regional information by analyzing the Tweed.

1. はじめに

近年、地域情報あるいは観光・行楽情報を収集する手段として、個人が Web ページから収集する傾向が強くなってきている。

また、観光地を訪れた来訪者が、その体験・感想を blog や SNS などを使って情報発信している。情報技術の発展により、これまでに比べ簡単に情報を発信することが可能になっており、個人の経験を記述した「口コミ」情報が利用者・体験者の「生の声」として重視されるようになった。そのため、観光推進・地域振興を行うためには Web 上での情報発信・情報収集が不可欠である。

本研究では、Web 上に存在する膨大なデータから、観光推進や地域振興などで活用するために「下関」「関門地域」などの地域情報を収集し、当地域の特徴を抽出する事を目的としている。我々は、特に速報性が高く、手軽に情報発信が可能な Twitter に着目した。Twitter から発信される情報を収集・蓄積・分析することで、様々な知見を得ることを目的としている。

2. テキストマイニング

インターネット上に保存されているデータの大部分はテキストデータであり、これらのデータを分析しようとするれば、テキストマイニングと呼ばれる手法を用いることになる。テキストマイニングとは、小説・新聞・Web ページなどの文章を対象とするデータ分析技術である。テキストマイニングを用いる目的は、大量に存在するテキストデータの中から、

特定のテーマを抽出することや、文章の分類を行うこと[1]であり、様々な研究が行われている。テキストデータの総量が少なければ、コンピュータを使用せずに分析を行うことも可能であるが、インターネット上に保存されているような大量データを分析する場合は、必要不可欠な技術であるといえる。実際に様々なデータ分析が行われており、インターネット上における blog 情報を分析し、即時性・重要性が高い情報を発信する有用な blogger を発見するための分析[2]や、インターネット上でユーザが感心を持っている Web サイトや検索キーワードを抽出して、視覚化を行う研究[3]などが行われている。

3. 地域・観光情報の収集

これまで、地域情報や観光情報にアクセスすることは容易ではなかったが、近年はインターネットサービスが発達し、blog や SNS などに各種口コミ情報が書き込まれるようになった。このため、本研究では Web から地域・観光情報の収集を目的とし、Twitter を情報源にする。blog は日記のような形態となっているため情報量は豊富であるが、該当する Web ページを見つけ出す時間、Web マイニングのための処理時間などを考慮するとデータ発見の対象とするのは現実的ではない。Twitter は一件 140 文字以内のメッセージで、サイト内の検索に制限はあるものの blog と比べればデータ収集は効率的である。このため、本研究では Twitter の書き込みから地域・観光情報を収集することにする。分析には次の 2 つのデータを対象とした。

- 1) 長期データ：2013 年 7 月 9 日～2014 年 4 月 15 日 (5,646 件, リツイート含まず)
- 2) 短期データ：2014 年 3 月 5 日 14 時～3 月 6 日 14 時(1,478 件, リツイート含む)

長期データにおけるトライグラムの組み合わせは 129,870 であった。これにはハッシュタグ、絵文字などが多数含まれており、これらを取り除いた時の組合せ数は共起数 3 以上に限れば 3,320 となった。ネットワーク図はタームをノード、共起の前後関係を方向、共起数を線の太さで表示する。また、ネットワークはタームの隣接行列から作成するが、その隣接行列は{0,1}により繋がりを示す。図 1 は共起数 50 以上の組合せのみ用いたグラフである。

短期データにおけるトライグラムの組み合わせは 17,150 であった。ハッシュタグ、絵文字などを取り除いた時の組合せ数は共起数 3 以上に限れば 462 である。ここでは、共起数 20 以上の組み合わせのみを使用したグラフを図 2 に示す。トライグラムから見付かった情報はイベント、トラフグ、角島大橋である。長期データでも角島大橋はあったが、トラフグやイベントは見当たらなかった。

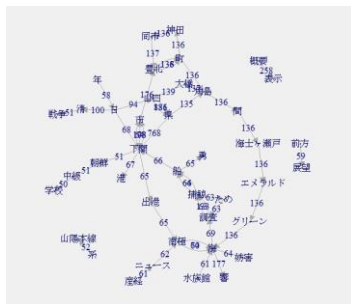


図 1 ネットワーク図（長期データ）

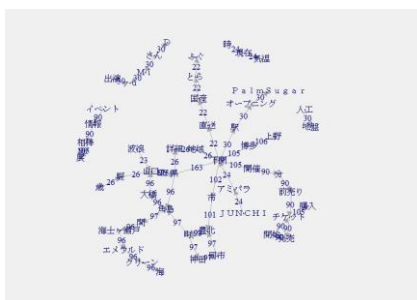


図 2 ネットワーク図（短期データ）

4. 結論

本研究では、インターネット上からの地域・観光情報収集し、収集したデータに対して分析を行い、その有用性を検討した。インターネット上に存在する膨大なデータの中から、地域や観光に関するデータを抽出し、観光推進・地域振興に役立つかどうかの検討を行った。

分析は、長期データ（2013 年 7 月～2014 年 7 月）と短期データ（2014 年 3 月 5 日）の 2 種類について行った。長期データについては、角島大橋・海響館・花火大会といった夏や暖かい季節の観光情報が多く含まれていた。短期データについては、その時期のイベント情報や観光情報・フグに関する発言が多く見られた。

参考文献

- [1]Manabu Nii, Shota Miyake, Kazunobu Takahama, Atsuko Uchinuno, Reiko Sakashita, "Consideration about Utilizing Text Architecture for Making Feature Vectors in Classifying Nursing-Care Texts", IEEE International Conference on Systems, Man, and Cybernetics, pp.1817-1821, 2013
- [2]中島 伸介、舘村 純一、原 良憲、田中 克己、植村俊亮、「重要な blogger 発見を目的とした blog スレッド解析手法」、知能と情報、Vol.19, No.2, pp.156-166、2007 年
- [3]村田 剛士、齋藤 皓太、「サイト・キーワードグラフを用いた Web ユーザの興味の抽出と視覚化」知能と情報、Vol.18, No.5, pp.701-710、2006 年

連絡先

下関市立大学 経済学部国際商学科

松本義之

E-mail:matsumoto@shimonoseki-cu.ac.jp

社会ネットワーク分析による ICT 企業のコア硬直性に関する一考察

Visualization of Core Rigidities of ICT Companies by Social Network Analysis

中岡 伊織^{*1} 朴 唯新^{*2} 陳 韻如^{*3}

Iori Nakaoka Yousin Park and Yun-Ju Chen

宇部工業高等専門学校^{*1} 県立広島大学^{*2} 滋賀大学^{*3}

Ube National College of Technology^{*1} Prefectural University of Hiroshima^{*2} Shiga University^{*3}

Abstract: Japanese ICT companies still have been struggling with poor performance in spite of various management reforms since collapse of the bubble economy in the early 1990s. Therefore, we recognize the necessity to understand issues of their technology innovation strategies. Specifically, this paper discloses and analyzes patent strategies and core rigidity of human resources of Japanese large ICT companies in smart phone market.

1. はじめに

1990年代初期まで強い競争力を保持していた日本の ICT 企業は、1991 年のバブル崩壊頃から多数の企業が業績不振に陥るようになった。日本の ICT 企業は競争力を回復させるために、直近 20 年間に人員削減や事業再編、協力企業との長期的な取引関係の見直しなどの様々な経営改革を実施してきたが、それらによって短期的には回復の兆しが現れたものの、持続的な競争優位を達成するには至らなかった。その理由としては、1.日本の ICT 企業が過度な品質を追求し、イノベーションのジレンマに陥った[1] 2. デジタル家電製品のアーキテクチャが従来のすり合わせ型からモジュラー型に変化した[2] 3. デジタル家電製品が急速にコモディティ化（低価格化）し、苦労して開発した新製品から価値獲得に失敗してしまった[3] ことなどが議論されている。その中で我々は、日本の ICT 企業が技術開発戦略の方向性を見誤っているのではないかと考える。米国市場調査機関である Thomson Reuters 社が特許を用いて調査した 2012 年の世界 100 大革新企業の中には、多数の日本の ICT 企業が上位に入っていることから、日本企業の技術開発能力が落ちたのではなく、日本の ICT 企業が実施していた技術開発が経営成果に結びつかなかったと考える。この仮説を検証するため、日本企業の技術開発戦略の方向性について他国の優良企業と比較検討する必要がある。本研究ではそのための探索的な取り組みの一つとして、日本の代表的 ICT 企業である SONY・SHARP・Panasonic および海外の優良 ICT 企業である Samsung に着目し、これら企業の携帯電話およびスマートフォンに関連する技術開発戦略を比較検討する。携帯電話及びスマートフォン市場は、2007 年の Apple の iPhone の発表および発売により、その市場構造が大きく変化した業界の一つである。よって 2000 年以後のこの業界における各社の特許戦略とその人員構成について明らかにし、その結果について検討する。

2. 特許情報

本研究では日本の ICT 企業の特許戦略を分析する。日本の特許情報は、全て特許庁における特許電子図書館にて管理されており、誰もが検索・閲覧す

ること可能となっている。各特許は IPC、IPC を細分化した FI、F タームにおける分類コードが複数付与されている。中でも F タームは、特許情報（特許公報類）中に記載されている技術的事項を把握した上で、種々の技術観点を付した F タームリストに照らして文献ごとに付与している。よって一つの文献には複数の F タームコードが付与されている。本研究では、特許情報分析（パテントマップ）からみた「スマートフォン」に関する技術開発実態分析調査報告書 CD-ROM 版[4]に記載されているスマートフォンに關係する特許に關する F タームを利用する。

3. SOM による技術開発分析

自己組織化マップを用いることにより、より多次元のデータから ICT 各社の戦略分析を行う。本研究では探索的な取り組みの一つとして、業界のイノベータとして存在する Apple が取得している特許において多く出現する F タームに対し、各社の特許におけるその F タームの出現回数を利用し分析を行う。出力されたマップは図 1 の通りである。

図中赤は 2005 年、青は 2009 年、緑は 2013 年の各社を示している。まず赤枠と青枠に着目するとその範囲が大きくなっていることから各社異なる戦略で進めていることがわかる。また、2009 年における日本企業は互いに距離が離れているが Samsung は Apple と近く出力されている。さらに、2013 年にはどの企業も似た分野の技術開発を行っている。この結果より、Apple のイノベーションに対して、海外優良 ICT 企業の一つである Samsung はいち早く対応しようとしたが、日本の ICT 企業群は全体的に見ると対応が遅れたということが明らかとなった。

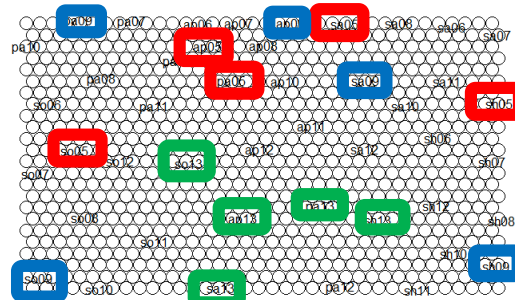


図 1 ICT 各企業の技術開発戦略推移

4. 社会ネットワーク分析によるコア硬直性分析

ICT 企業において、新製品の研究開発プロジェクトをいかに管理するかは非常に重要な課題の一つである。その際、既存のコア能力 (Core Capabilities) をベースに持続的に開発するか、あるいは、既存のコア能力を解体・破壊し、新たな能力を開発するかという難問に直面することになる[5]。これについて D. Leonard-Barton は、企業は研究開発プロジェクトにおいて既存のコア能力によるコア硬直性 (Core Rigidities) のパラドックス問題に直面すると指摘している[6]。つまり、企業は従来までの競争優位をもたらし続けていた既存のコア能力に膠着してしまうことで、既存のコア能力が環境変化に対応する新たなコア能力の構築を妨げる構造慣性と抵抗となり、企業にコア硬直性をもたらす危険性がある。

3 章において、日本 ICT 企業は Apple が行った技術革新への対応が遅れたことを明らかにした。本章ではその理由が上記に示したコア硬直性にあると考え、社会ネットワーク分析を用い、中心的な役割を担っている開発者の変遷を中心性分析にもとづき人員構成の変化について分析する。

技術開発に携わる人員の組織構造変化を表現するため、 i 年の媒介中心性上位 $j\%$ の人員が、 $i+1$ 年の上位 $j\%$ に存在したかという条件付き確率を図 2 から図 5 にヒートマップにて表す。なお、図中において色が薄いほど変化が激しいことを意味する。SONY は、iPhone ショックの 2007 年頃から硬直性が増していることがわかる。Panasonic および SHARP は、色の濃い部分が多いことより、全体的にコアとなる人員の入れ替わりが少ないことが読み取れる。対して、Samsung は、2007 年頃から特にコア人材が入れ替わっていることが読み取れる。これらより、スマートフォン市場は 2007 年に iPhone ショックがあり、技術開発の方向性が変わっていったにもかか

わらず、日本企業はコア硬直性が高く、必要となる技術領域に適応した人材にうまく転換できなかったことが見受けられた。

5. おわりに

本研究では日本の代表的 ICT 企業である SONY, SHARP, Panasonic および海外優良企業である Samsung におけるスマートフォン市場に対する技術開発戦略の推移を明らかにした。日本 ICT 企業は市場変化への対応への遅れが見受けられたがその理由には、組織としてコア硬直性が高い状態にあったことによると導かれた。今後は、より詳細な人員構成に関する分析、技術面のコアからの分析および ICT 企業の様々な事業を含めた分析を行いたい。

参考文献

- [1] Christensen C. M. “The Innovator's Dilemma: The Revolutionary Book That Will Change the Way You Do Business”, HarperCollins (2003)
 - [2] 藤本隆宏, 大鹿隆: 製品アーキテクチャ論と国際貿易論の実証分析, 経済産業研究所, 06-J-015 (2006)
 - [3] 延岡健太郎, 伊藤宗彦, 森田弘一: コモディティ化による価値獲得の失敗: デジタル家電の事例, 経済産業研究所, 06-J-017 (2005)
 - [4] パテントテック社. 特許情報から見たスマートフォンに関する技術開発実態分析調査(2012)
 - [5] M. Christensen and M.E. Raynor “The Innovator's Solution: Creating and Sustaining Successful Growth”, Harvard Business School Publishing Corp (2003).
 - [6] D. Leonard-Barton “Core Capabilities and Core rigidities: A Paradox in Managing New Product Development” Strategic Management Journal, Vol. 13, pp. 111-125 (1992).
- 連絡先
宇部工業高等専門学校 経営情報学科
中岡 伊織
(Tel.: 0836-35-5130, E-mail: nakaoka ube-k.ac.jp)

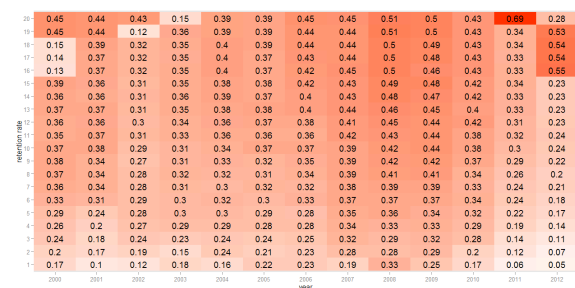


図 2 SONY のコア硬直性

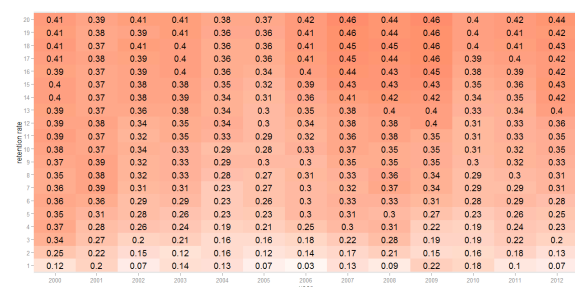


図 3 SHARP のコア硬直性

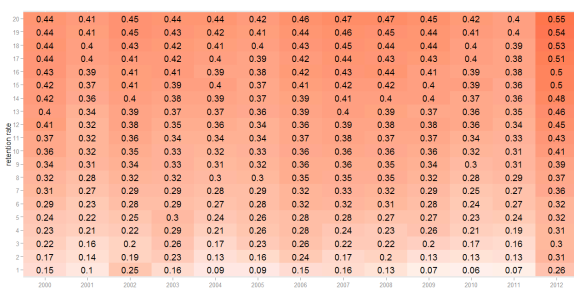
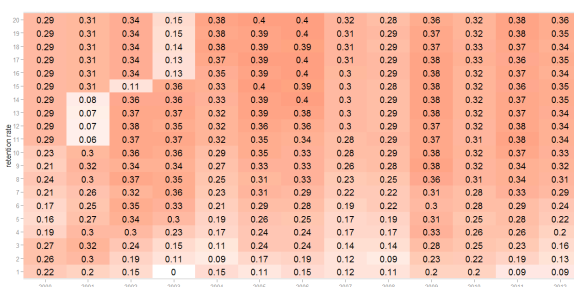


図 4 Samsung のコア硬直性



青年群、高齢者群における理学療法士・作業療法士・看護師などに望む対応

Research on the Desirable Degree of Self-Participation to the Determination of Medical Services for Ordinary People

奥田裕紀

Hiroki Okuda

金城大学

Kinjo University

Abstract: The purpose of this research is to examine the desirable degree of self-participation for ordinary people in determining the contents and the methods of medical services from medical service workers. The research participants are 100 ordinary Japanese people. 50 of the participants are university students and 50 are elderly people. 39 evaluation items are divided into three groups consisting of the same number of items according to the participation level of patients in the determination of the contents and the methods of medical services. The three groups are called low, medium, and high degree self-participation item groups. The participants were asked to assume three cases of mild, moderate and serious illnesses or injuries. In all of the three cases, the participants were asked to evaluate the degree of desirability of each item on the 7-point scale. The mean scores of the medium degree self-participation item group were higher than those of the other two item groups in all of the three cases. In the serious case, the medium degree self-participation item group showed the highest mean score. These results suggest that there are many people who think it desirable to carry out a medium degree of self-participation as patients in selecting and determining the contents and the methods of the medical services.

1. はじめに

理学療法士、作業療法士、看護師などの医療関係者にとって、理学療法、作業療法、看護などの対象となる方々（患者や施設利用者など、以下対象者とする）に対する対応や、対人コミュニケーションスキルは、極めて重要とされている。このため、これらの資格取得者を養成する教育課程や臨床の現場など多様な場面で、対象者への対応、対人コミュニケーションなどに関する研究が行われてきた。しかし、いまだに不明な点も多く、多くの課題が残されている。

例えば、Okuda (2014) ^[1]は、“リハビリテーション、医療の現場における医療関係者と対象者との一般的なコミュニケーションスキルに関する項目”を収集し、これらの項目を、“基本的スキル項目群”、“特定対象者項目群”、“対処困難項目群”、“一般項目群”に分類した。そして、

理学療法に関する教育や、臨床実習経験などが、これらの項目に関する認知に与える影響について検討を行った。

理学療法学を専攻する大学1年生群（理学療法に関する専門的教育は、未だほとんど受けていなかった）と、理学療法学を専攻する大学4年生群（基礎医学、理学療法学などに関する学内での学習と、3回（20週程度）の臨床実習を経験していた）において、これらの4項目群に関する認知を比較した。

具体的には、1年生群、4年生群の研究参加者に対して、各項目が理学療法士、作業療法士などの対象者に対する一般的な対人コミュニケーションスキルとして必要だと思われる程度（必要度）と、自分自身が現時点でそれらの項目に示されたスキルを上手くできると思う程度（可能性）について評定を求めた。

その結果、必要度に関する評定に関しては、1 年生群、4 年生群ともに、この研究で検討された 4 項目群ともに評定平均値は高かった。特定の対象者に対応する場合のみに必要となる“特定対象者項目群”や、医療関係者にとって特に対応が困難な“対処困難項目群”の必要度は低下する可能性があると思われたが、実際には、両群とも、“特定対象者項目群”の評定平均値が最高となり、“対処困難項目群”の評定平均値も他項目群と同じく比較的高かった。

一方、可能度に関しては 4 項目群とも、4 年生群の方が 1 年生群より評定平均値が有意に高く、大学における教育や臨床実習が、各項目の（特に可能度に関する）認知に影響を与えることを示唆している。

しかし、このような研究からは医療関係者や、医療関係資格の取得を目指す学生の立場から、対象者にどのように対応することが重要と認知しているのか、望ましいと認知しているのかなどに関する情報は得られても、対象者自身がどのような対応を医療関係者に望むのかに関しては、直接的な情報は得られない。

そこで、本研究では、この問題を検討するために、医療の現場において対象者となる可能性のある一般の人達が、理学療法士、作業療法士、看護師などの医療に関係する人達について、どのような対応が、どの程度望ましいと思うのかに関する評定を求めることとした。

ところで、一般の人達（対象者）が、医療関係者にどのような対応を望むのかは、当該の人のある時点での病気・ケガなどの重篤度（重症度）によって変化する可能性が考えられよう。

そこで、本研究では、研究参加者に軽度、中程度、重度の場合を想定してもらい、各々の条件における、理学療法士、作業療法士、看護師などの医療関係者の対象者に対する対応の望ましきなどに関して、評定を求め比較することとした。

2. 方法

① 研究参加者

研究参加者は、青年群（大学生）50 人、および高齢者群（60 歳以上の）50 人であった。これらの研究参加者には、医療従事者や医療関係資格の取得を目指す学生は含まれておらず、医療の対象者（患者や施設利用者）となる可能性のある一般の人達であった。研究参加者全体の平均年齢は、47 歳であった。

② 病気・ケガなどの重篤度条件

理学療法士、作業療法士、看護師などの医療関係者の対象者に対する対応、対人コミュニケーションなどの望ましきに関しては、各対象者の病気・ケガなどの重篤度（重症度）によって変化する可能性が考えられた。そこで、本研究では、軽度（短期間で完治する程度）、中程度（1 か月ほどの入院を必要とする程度）、重度（命にかかわるような場合）の 3 条件を想定してもらい、各条件について、理学療法士、作業療法士、看護師などの対応、対人コミュニケーションなどとして望ましきなどについて評定を求めることとした（従って、各研究参加者は同じ項目について、軽度条件、中程度条件、重度条件の 3 回評定を行った）。

③ 評定項目および項目群

評定項目は、全部で 68 項目であった。これらの項目は、Okuda (2014) による先行研究や、医療関係者の対象者に対する対応、コミュニケーションスキルなどに関する文献から収集した。

これらの項目の一部は、検査、治療方針・方法などに関する説明、選択・決定、確認・変更などに関する、対象者自身の決定関与度の大小によって、小自己決定項目群、中自己決定項目群、大自己決定項目群に分けられた（各 13 項目とした。なお、先行研究^[2]では、各群 14 項

目としていたが、項目について再検討した結果、13 項目に絞ることとした。また、各データの確認・修正を行った）。

小自己決定項目群は、例えば、“治療方針・方法などについては、詳しく説明されても自分には良くわからないので、ホドホドにしておいてもらいたい”、“療法士などが専門家の立場から必要だと思う検査や治療、生活習慣の改善などがある場合には、必ず行ってもらいたい”などの項目から構成されていた。

中自己決定項目群は、例えば、“治療方針・方法などについては、理解できるまででいいに説明してもらいたい”、“療法士などが専門家の立場から必要だと思う検査や治療、生活習慣の改善などがある場合には、何度も説得して納得させてもらいたい”などの項目から構成されていた

大自己決定項目群は、例えば、“治療方針・方法などについては、専門用語などを使っても良いので、自分に対してできる限り詳しく説明してもらいたい”、“たとえ療法士などが専門家の立場から必要だと思う検査や治療、生活習慣の改善などであっても、自分が拒否した場合はそれ以上勧めないでもらいたい”などの項目から構成されていた。

この他に、責任回避・押しつけ認知項目群（“検査や治療などに伴う危険性に関して事細かく説明されると、万が一の場合にそなえ療法士などが責任回避をしているように感じる”などの項目群）、一般項目群（一般的な対応、対人スキルなどに関する項目）などに分類された。

なお、本論文では、主として、小・中・大自己決定項目群の評定結果について示し、他の結果は別に報告することとする。

④ 心理的リアクタンスの定義および測定

“心理的リアクタンス”とは、人が自分の自由を外部から脅かされた時に生じる、自由を回

復しようとする動機づけ状態のことである。

人間は、たとえ自分が思っていたこと、そうしようと考えていたことと全く同じことであっても、他者から言われると反発することがある（例えば、そろそろ宿題をしようと思っていた子どもが、母親から“宿題をきなさい！”といわれると反発し、宿題をするのをやめてしまうことがある）。

心理的リアクタンスの強さは、人により異なり、自由を維持する重要性、自由を回復することによって得られる利得の大きさなどによっても変化すると考えられる。

本研究では、リハビリテーション・医療などの現場における医療関係者の対象者に対する対応、コミュニケーションスキルなどに関する認知に、研究参加者の心理的リアクタンスがどのように影響するか調べるために、研究参加者に、心理的リアクタンスに関する評定を求めた。

心理的リアクタンスに関する尺度は複数あるが、本研究では、妥当性・信頼性が比較的高いとされる今城（1993）^[3]の心理的リアクタンス関に関する尺度（“人に頼らず、自分が自由に意思決定することは、私にとって非常に重要だ”などの 18 項目からなっていた）に基づき、評定を求めた。

⑤ 各項目の評定方法

研究協力者に、小・中・大自己決定項目群を含む各評定項目が、リハビリテーション・医療の現場において、理学療法士、作業療法士、看護師などの医療関係者の対象者への対応、対人コミュニケーションなどとして望ましいと思う程度などについて、7 段階の尺度で評定することを求めた。

また、心理的リアクタンスについても、同じく 7 段階で評定を求め、各研究協力者における各項目の評定の合計を、心理的リアクタンスの強さの指標（心理的リアクタンス得点）とした。

3. 結果と考察

① 軽度、中程度、重度条件における小、中、大自己決定項目群の評定平均値

各評定項目の評定値は、0～1に変換した。

図1に示したように、軽度条件、中程度条件、重度条件とも、小自己決定項目群、大自己決定項目群の評定平均値よりも、中自己決定項目群の評定平均値が有意に高かった ($p < 0.01$)。また、重度条件の場合に、中自己決定項目群の評定平均値が最も高くなった。これは、先行研究において各群が14項目によって構成されていた場合の結果と、基本的には類似した結果であった。

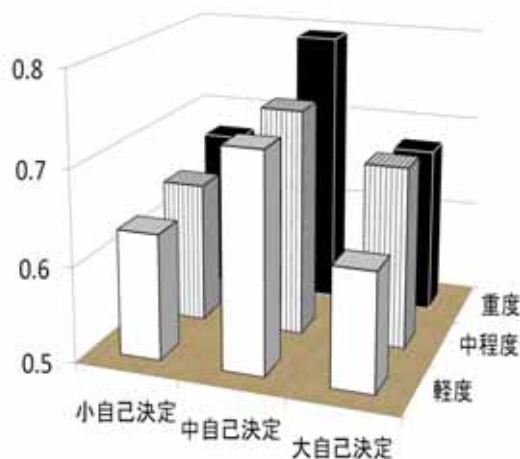


図1 各条件の小・中・大自己決定項目群の評定平均値

これらの結果から、青年群および高齢者群からなる本研究の研究参加者群全体においては、重度条件の場合(命にかかわるような場合)に、検査・治療の方針・方法などの重要な決定を行う場合などには、理学療法士、作業療法士などの専門家から十分な説明を受け、適切な選択が出来るように対応して欲しいと思う人が多いことが示唆される。ただし、本研究では、あくまで研究協力者に各条件を想定して評定することを求めており、実際の対象者、治療場面ではないことには、留意が必要であろう。

② 心理的リアクタンスとの関連と今後の課題

一方、軽度条件、中程度条件、重度条件における、各研究参加者の大自己決定項目群の評定平均値と、心理的リアクタンス得点との間には、有意な正の相関があることが示された ($P < 0.05$)。すなわち、本研究の研究参加者全体では、心理的リアクタンス得点が高い(自己決定を重んじる特性が高い)人ほど、大自己決定項目群を高く評定する傾向があることが示され、心理特性と自己決定に関する望ましさに関する認知の関連が示唆された。

本論文では、ページ数の制約などから、研究参加者の年齢などの要因による分析結果は示さなかったが、年齢要因も一般の人、対象者の医療の現場における医療関係者の対応の望ましさに関する認知に、影響を与える可能性が考えられるので、今後の検討課題となろう。

4. 参考文献

- [1] H. Okuda: Effects of education and clinical-training on the recognition of communication skills for physical and occupational therapist, International journal of modern education forum, Vol,3(3), pp.98-101, 2014
- [2] 奥田裕紀： 理学療法士、作業療法士、看護師等の対応に関する希望の検討 第19回曖昧な気持ちに挑むワークショップ, 3.3, 1-4, 2014
- [3] 今城周造: 心理的リアクタンスの検討(2), 日本教育心理学会第35回総会発表論文集, 502, 1992

5. 連絡先

〒924-8511 石川県白山市笠間町 1200
金城大学
TEL. 076-276-4400
FAX 075-236-4306
E-mail: okuda@kinjo.ac.jp

ファジィ事象とその応用

A Fuzzy Event and its Applications

堀 芳樹¹, 椿 広計²

Houju Jr. Hori Hiroe Tsubaki

1) 元三重大大学, 統計数理研究所

2) 統計数理研究所

Mie University (Former), the Institute of Statistical Mathematics

The Institute of Statistical Mathematics

0. 要約

ワルドの決定関数論の中の主観修飾は、後のファジィ理論の中のファジィ事象である。自然の状態を

ファジィ事象に積分変換することが、ワルドの主観修飾の概念である。ファジィ事象の確率やファジィ効用関数値の演算公式は、ワルドの積分変換から容易に導かれる。我々は、ワルドの積分変換にザデー拡張原理を適用し、多数の主観可能性分布が OR 結合と AND 結合の場合、ファジィ数学の適用が有効であり、ザデーの拡張原理が主観修飾の積分変換であることを示した。本稿では、ノーデータ問題に着目して、多数の OR 結合または AND 結合された主観可能性分布とメンバシップ関数が与えられている場合におけるファジィベイズの定理を紹介し、その意思決定問題の定式化を図る。

(キーワード)

主観修飾、ファジィ事象、メンバシップ関数、主観可能性分布、OR 結合、AND 結合、ファジィベイズの定理、ファジィベイズ意思決定法、モンテカルロシミュレーション

1. はじめに

ワルドの決定関数論[1]の中に、今まであまり注目されていなかった主観修飾がある。これは、後に提案されたファジィ集合[2]の中で、リスク関数を積分変換するファジィ事象と言える。ファジィ事象の確率[3]やファジィ効用関数値[4]の演算公式は、ワルドの積分変換公式から容易に導くことができる。その後、植村[5,6,7,8]は、ワルドの積分変換公式にザデーの拡張原理の適用し、ファジィ事象における意思決定法を構築した。さらに、堀と椿[9,10,11,12]は、主観可能性分布が OR 結合と AND 結合の場合、ザデーの拡張原理がワルドの積分変換公式となり、ファジィ数学の有効性と妥当性を示した。最後に、ファジィ事象を扱う場合、自然の状態上で状態が生起する可能性を表す主観可能性分布と自然の状態をファジィ事象に変換するメンバシップ関数が必要となる。ここで、ノーデータ問題においても、二つの分布が必要となることに注意しよう。まず、ノーデータ問題に注目し、多数の主観可能性分布とメンバシップ関数が、主観的に与えられていると仮定して、ファジィベイズの定理を定式化する。次に、OR 結合の場合は、効用関数、AND 結合の場合は、損失関数を導入し、ファジィベイズ意思決定法則を導く。最後に、事例研究として、相反する自然の状態の多次元化による新ゲーム理論[8]の紹介と、その新ゲーム理論に時間の概念を導入したファジィ推移行列を有するモンテカルロシミュレーション[13, 14, 15]について言及する。

2. ファジィベイズの定理 (OR 結合の場合)

自然の状態 s 上の主観可能性分布 $\Pi(s)$ は、メンバシップ関数 $\mu_{F_i}(s)$ により、自然の状態からファジィ事象 F に変換写像されていると考えられる。ザデーの拡張原理を用いれば、次式のように定式化できる。

$$\int \mu_{F_i}(s) / \Pi(s) \triangleq \max_{\{x=\Pi(s)\}} \mu_{F_i}(s) = \mu_{F_i}(\Pi^{-1}(x))$$

(1)

主観可能性分布自体、サバージ[16]のくじ引きで同定される。また、釣鐘型の分布の時は、ベイズ統計学と同値であることも示されている[17]。今、多数の主観可能性分布が、OR 結合である場合を考えているので、主観可能性分布論から、統合された分布は、(2) 式のように定義できる。

$$\text{or}_{\{x=\Pi(s)\}} \triangleq \max_{\{x=\Pi(s)\}} \mu_{F_i}(s) = \mu_{F_i}(\Pi^{-1}(x))$$

(2)

(1)式と(2)式は、同値となり、これが、OR 結合の場合のファジィベイズの定理である。

3. ファジィベイズの定理 (AND 結合の場合)

OR 結合の場合の双対系を考えれば、(3)式と(4)式を得る。

$$\int \mu_{F_i}(s) / \Pi(s) \triangleq \min_{\{x=\Pi(s)\}} \mu_{F_i}(s) = 1 - \mu_{F_i}(\Pi^{-1}(x))$$

(3)

$$\text{AND}_{\{x=\Pi(s)\}} \triangleq \min_{\{x=\Pi(s)\}} \mu_{F_i}(s) = 1 - \mu_{F_i}(\Pi^{-1}(x))$$

(4)

(3)式と(4)式は、同値となり、これが、AND 結合の場合のファジィベイズの定理である。

4. 意思決定問題への応用

OR 結合の場合は、効用関数 $U(s | d)$ を、AND 結合の場合は、損失関数 $L(s | d)$ を導入する。ここで、 d は行動を表記する。

4-1 ファジィ効用関数

ザデーの拡張原理から導かれたものを(5)式に、効用関数論から導かれたものを(6)式に示す。ここで、(5)式と(6)式は、同値であることに注意しよう。従って、効用関数論では、ファジィ数学が有効であるといえる。

$$\int \mu_{F_i}(s) / U(s | d) \triangleq \max_{\{x=U(s | d)\}} \mu_{F_i}(s) = \mu_{F_i}(U^{-1}(x | d))$$

(5)

$$\text{or}_{\{x=U(s|d)\}} \triangleq \max_{\{x=U(s|d)\}} \mu_{F_i}(s) = \mu_{F_i}(U^{-1}(x|d))$$

(6)

4-2 ファジィ損失関数

OR 結合の場合の双対系を考えれば、(7)式と(8)式を得る。

$$\int \mu_{F_i}(s)/L(s|d) \triangleq \min_{\{x=L(s|d)\}} \mu_{F_i}(s) = 1 - \mu_{F_i}(L^{-1}(x|d)) \quad (7)$$

$$\text{AND}_{\{x=L(s|d)\}} \triangleq \min_{\{x=L(s|d)\}} \mu_{F_i}(s) = 1 - \mu_{F_i}(L^{-1}(x|d))$$

(8)

この場合も、(7)式と(8)式が同値になり、損失関数論ではファジィ数学が有効であるといえる。

4-3 マックスプロダクト演算と最適行動

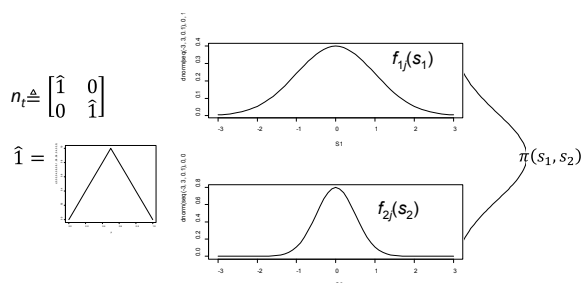
OR 結合の場合は、ファジィ効用関数に(1)式または(2)式を掛けて最大化を図った測度を採用し、これらの測度の中で最大のものを有する行動を最適行動とする。また、AND 結合の場合は、ファジィ損失関数に(3)式または(4)式を掛けて最大化を図った測度を採択し、これらの測度の中で最小のものを有する行動を最適行動とする。

5. ファジィ事象におけるモンテカルロシミュレーション

ノーデータ問題において、相反する多次元の自然の状態を考え、危険中立なファジィ事象の可能性

測度を評価関数として、多段階ファジィベイズ意思決定法則が定式化された[12]。ここでのファジィ推移行列は、危険中立的な状況を考えているので、単位行列における1を三角型のファジィ確率1と設定している。この単位行列でのファジィ確率の一方を0から1までのリニアーにした時、段階的に復活のモデルであり、1から0までのリニアーにした時は、段階的に消滅のモデルである[13]。また、反転の行列をファジィ化することで、段階的に、例えば、「Cool Head and Warm Heart」から「Dry Mind and Cool Head」に反転することが報告されている[14]。さらに、我々は、タイプ2型ファジィ推移行列を導入し、 α レベルが1の代表曲線を用いて、危険中立的な意思決定を行う方法を定式化した[15]。ここでは、基本モデルある多段階ファジィベイズ意思決定法則を紹介する。

(多段階ファジィベイズ意思決定法則)



$$M_t(s_1, s_2) := f(s_1, s_2) \times U(s_1, s_2) \times \pi(s_1, s_2)$$

$$= \{f_{1j}(s_1) \wedge f_{2j}(s_2)\} \times \{U(s_1) \wedge U(s_2)\} \times \pi_1(s_1) \pi_2(s_2) \quad (9)$$

$$M_{t+1}(F_j) = \int_0^1 \int_0^1 M_t(s_1, s_2) n_t ds_1 ds_2 \quad (10)$$

ここで、 f は 2 次元のメンバシップ関数、 U は 2 次元の効用関数、 π は、2 次元の主観可能性分布である。

6. おわりに

ファジィ事象を規定するメンバシップ関数は、自然の状態をファジィ事象に変換する関数で、ワルドの主観修飾の概念を取り入れたものである。主観可能性分布が、OR 結合と AND 結合の場合は、この主観修飾の積分変換にザデーの拡張原理が有効な手段となる。今後は、OR と AND が、混在している場合のファジィベイズの定理の構築とファジィ事象におけるモンテカルロシミュレーションの収束性・発散性を(ファジィ)固有値解析で行いたい。

(参考文献)

1. A. Wald, Statistical Decision Functions, Chelsea Publishing Company, 1950
2. L.A. Zadel, Fuzzy Sets, Information and Control, Vol. 8, pp. 338-353, 1965
3. L. A. Zadeh, Probability Measures of Fuzzy Events, J. of Mathematical Analysis and Applications, Vol. 22, pp.421-427, 1968.
4. H. Tanaka, T. Okuda, and K. Asai, Fuzzy Information and Decision in Statistical Model, Advances in Fuzzy Sets Theory and Applications, pp. 303-320, 1979.
5. 植村芳樹、ファジィ事象における意思決定法、日本ファジィ学会誌、Vol. 3, pp. 123-130, 1991.
6. 植村芳樹、坂和正敏、ファジィ事象の可能性分布に基づく簡易意思決定法、Vol. 5, pp. 528-536, 1993.
7. Y. Uemura, A Normal Possibility Decision Rule, Control and Cybernetics, Vol. 24, pp. 103-111, 1996.
8. Y. Uemura, Application of Normal Possibility Theory to Silence, Control and Cybernetics, Vol. 30, pp. 465-472, 2001.
9. 堀芳樹、椿広計、ファジィベイズ意思決定法則、応用統計学会年会、2014
10. 堀芳樹、椿広計、効用関数論のファジィ論的アプローチ、統計関連連合講演会、2014
11. 堀芳樹、椿広計、マックスプロダクト法とその応用、バイオメディカルファジィ学会、2014
12. 堀芳樹、椿広計、多段階ファジィベイズ意思決定法、バイオメディカルファジィ学会、2014
13. 植村芳樹、Case Study of Fuzzy Events, 応用統計学会年会、2010
14. Y. Uemura, Application of Fuzzy Decision Making to Smoking Problem, J. of Fuzzy Mathematics, Vol.18, pp.887-892, 2011
15. H. Hori Jr., and H. Tsubaki, Fuzzy Decision Making Rule under no-Data Problem, Applied Mathematics, (to appear), 2014.
16. L. J. Savage, The Foundations of Statistics, Dover Publications, Inc., 1972
17. 田中英夫、石渕久生、正規可能性分布による証拠理論、システム制御情報学会誌、Vol. 5, pp. 235-243, 1992.

(連絡先)

堀芳樹 (ほりほうじゅ)

統計数理研究所 椿広計研究室

〒190-8562 東京都立川市緑町 10 番地の 3

Phone: 050-5532-8562

E-mail: uemura0742@yahoo.co.jp