

ITMA における NCD の実用性と UML を用いた 分析プロセスの可視化と共有について

藤本 悠

日本学術振興会特別研究員 PD(同志社大学)/
Affiliate Academic at Institute of Archaeology,
University College London

小野原彩香

同志社大学大学院文化情報学研究科
博士前期課程

理念型モデル化分析法(ITMA: Ideal Types Modeling and Analysis)とは、情報モデルによって文化研究者の認識構造をモデル化し、データベース化し、分析するための方法論で、研究者の認識構造を矛盾無く組み立てるための「メタモデル」と、判断と解釈を写像する「メソッド」と「インタフェース」を提供する。本稿ではITMAにおける新しいインタフェースとして「正規化圧縮距離(NCD: Normalized Compression Distance)」の可能性を検討すると同時に、一連の分析プロセスをUML アクティビティ図に記述し、分析の再利用性を向上させる方法についても検討する。

Availability of Normalized Compression Distance in ITMA and Documentation of Comprehension Process by UML Activity Diagram

Yu FUJIMOTO

JSPS Research Fellow(Doshisha University) /
Affiliate Academic at Institute of Archaeology,
University College London

Ayaka ONOHARA

Graduate School of Culture & Information Science,
Doshisha University

Ideal Types Modeling and Analysis (ITMA), which provides a 'meta-model', 'methods' and 'interfaces' reflecting researchers' perceptual comprehensions and semantic structures without logical discrepancy, is a methodology for cultural science. In this paper, we propose a new interface of ITMA based on Normalized Compression Distance (NCD). We additionally suggest a means for documentation of researchers' comprehension procedure by using activity diagram in Unified Modeling Language (UML).

1. はじめに

人文科学は主観的解釈の自己言及的な無限ループに陥るとみなされることがある。「じんもんこん 2010」が掲げる「人文工学」とは、それらに科学(理学)・工学的手法を応用することによって、共有可能な解釈へと導くことに挑戦するものである[12]。この考えは、藤本[16]によって位置づけられた、文化情報学における「理哲・工文」の考え方にも通じる。藤本は、この概念に従って、情報モデルによって文化研究者の認識構造をモデル化し、データベース化し、分析するための方法論、理念型モデル化分析法 (ITMA: Ideal Types Modeling and Analysis)[15]を考案した。この方法論は、個々の研究者の経験量の差を許容し客観的な方法によって論理に矛盾が無いと証明し得るものを文化現象研究における客観性とする。すなわち、個々の研究者の認識構造の違いは、主観的な解釈ではなく、同じメタモデルから派生したインスタンスに過ぎないという考えに基づく。ITMA は研究者の認識構造を矛盾無く組み立てるための「メタモデル」と、判断と解釈を写像する「メソッド」と「インタフェース」を提供する。

この方法論は、理系・工系の技術と人文科

学を結びつける上で重要な概念的枠組みを提供する一方で、メソッドやインタフェースが十分に整理されておらず、また、ITMA に基づく研究のプロセスが不明瞭であるといった課題がある。そこで、本研究では、一つ目の課題に対しては「正規化圧縮距離(NCD: Normalized Compression Distance)」を用いた新しいインタフェースの可能性を検討し、二つ目の課題に対しては、一連の分析プロセスをUML アクティビティ図に記述する方法について検討する。

2. ITMA の根本概念

2.1. ITMA とは

ITMA は、概念的には M. ウェーバー[8]の「理念型(Ideal Types)」の考えと、技術的にはオブジェクト指向 GIS(OOGIS)に示唆を受けて考案した。

ウェーバーは、現実における諸問題をより鮮明に浮き上がらすためには、究極的理想的状态(理念型)と現実の観察可能な状態を比較することが重要であると考えた。そうすることで、現実がその理想的な状態に対して過剰に現れていたたり、その逆に全く現れていなかったりする部分をより鮮明にすることができ

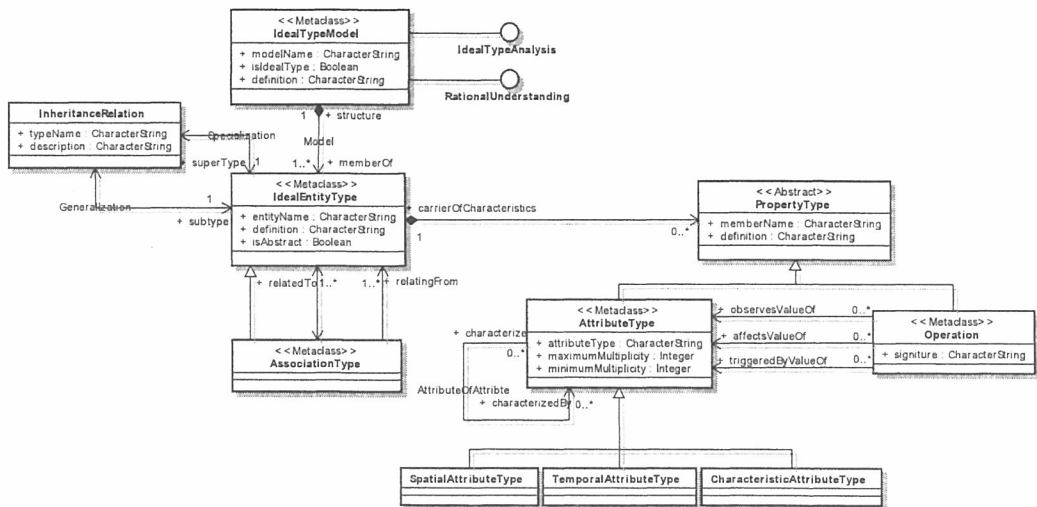


図 1 ITMA のメタモデル

る。つまり、理念型として構成される概念は「平均化された原則」ではなく、「一つあるいは幾つかの観点をそれだけ強く一面的に高める」ために機能する。理念型そのものは、本質ではないが、現実では目に見えない状態で潜んでいる諸問題を取り出すための方法となる。ITMA は、このウェーバーの理念型の考えをコンピュータ上で実現する方法として整理したもので、現実世界に対する認識構造を情報モデルによって記述し、その情報モデルに従って仮説的に構成される理念型モデルと、実際の調査で取得される現実型をコンピュータ上で比較するための「概念的枠組み(図 1)」をメタモデルとして提供する。

ITMA は、認識構造を記述するための具体的な情報モデル化言語と符号化規則の何れも定義していないが、現状においては、情報モデル化言語にはオブジェクト指向に基づく UML を、符号化には XML が最適であるとしている。この二つの技術を推奨している背景には、現実世界は個々の抽象的なオブジェクトから構成され、それらは相互に関連性を持っているという前提がある。この前提は、OOGIS の前提そのものであり、ITMA のメタモデルにおいても明示されている。

ITMA の前提は、OOGIS そのものであるため、OOGIS の実装の一つである『地理情報標準(JIS X7100 シリーズ)』の「一般地物モデル(GFM: General Feature Model)」と ITMA のメタモデルは、類似した構造を持つ。そもそも現実世界をオブジェクト指向で記述するための時空間オブジェクトの定義や、情報モデル化する際の規則を規程するには多大なコストをかけることになるが、GFM と類似構造を持っているため、時空間オブジェクトは同標準

の「空間スキーマ(ISO-19107)」および「時間スキーマ(ISO-19108)」で代用でき、設計方法は「応用スキーマの規則(ISO-19109)」に準じることができる。さらに、本来は UML クラス図で定義されたクラスのインスタンスを XML で符号化するためのルールは規定されていないが、地理情報標準に準じて UML クラス図を設計することで、「符号化規則(ISO-19118)」あるいは「GML(Geography Markup Language)(ISO-19136)」を用いて XML ドキュメントを作成できる。

GFM と ITMA のメタモデルは構造上はほぼ同じであるが、大きく異なる点もある。まず、GFM は、オブジェクト指向に基づいているが、主たる目的はデータベース構築であるためメソッド(図 1 における Operation クラス)に関する具体的な議論はほとんどされていない。一方、ITMA のメタモデルは、最終的に理念型を用いた現実世界の合理的理解を促進させることが主たる目的であるため、メソッドは重要な概念である。また、設計されたモデルそのものに適用可能なインタフェースを備えている点も GFM と大きく異なる。

なお、これら ITMA の諸概念に関しては、[16]で詳しく説明されている。

2.2. ITMA のメソッドとインタフェース

ITMA におけるメソッドは、GFM におけるメソッドと同様に属性値に作用する操作である。ITMA のメソッドは、人が五感を通して認識した情報を、認識の通りにコンピュータが理解できるように属性値を操作するために機能する。具体的には、音、匂い、色、形状、感触といった属性値に対して、現在の属性値を返す「オブザーバ操作(observeValueOf)、属

性値の変更を行う「変更操作(affectsValueOf)」,そして定義されたクラスの初期化を行う「コンストラクタ操作(triggeredByValueOf)」の何れかの操作を提供する。人文科学では、既に認識されている人工物と、未分類の人工物との知覚上の類似性を評価を試みることがあり、メソッドはそうした認識上の類似性をコンピュータ上で扱うために重要である。

ITMAのメソッドとして最初に実装したのは、形状を符号化するための「空間形状分析(SSA: Spatial Shapes Analysis)」である。メソッドは無限に存在するが、現状ではSSAは唯一の実装されたメソッドである。このメソッドは、人が認識する通りにコンピュータが形状の類似性を扱えるように属性値を符号化する。例えば、考古学では従来から形状を数値化して扱う方法に対する感心が強く、円形度(Circularity)、方形度(Quadrature)、不規則度(Irregularity)および伸長度(Elongation)といった方法や、最近ではニューラルネットワークを用いた手法なども議論されている[1]。こうした試みの全てはSSAの実装対象となり得るが、実装済みのSSAは、ある属性の値が2次元の面データとして取得されると同時に、その形状を「楕円フーリエ記述子(EFD: Elliptic Fourier Descriptor)」として取得する方法である。フーリエ記述子を用いて形状の類似性を評価する方法は、形態測定学(Morphometrics)において用いられていて、同手法を実装するプログラムは多い[4]。ITMAのメソッドではこの手法をベクトル型データにも適用できるように改良してある。

ITMAにおけるメソッドが主として属性に機能するのに対して、インタフェースは理念型を用いた分析のための具体的な分析手法を提供する。この概念は、ITMA独自のものでGFMには類似した概念は無い。インタフェースには、「合理的理解(RationalUnderstanding)」と「理念型分析(IdealTypeAnalysis)」があり、前者は、ウェーバが提唱した理念型との論理的矛盾を解消するために存在し、後者は実際の分析を定義するために存在している。現状においては、後者のみが定義されていて、形質構造分析(CSA: Character Structure Analysis)」と「時間位相分析(TTA: Temporal Topology Analysis)」の二つが実装されている[17]。まず、CSAは、理念型と現実型の形質的特徴の差異から個々の資料を「系統」として分析する。この分析手法は、分子進化遺伝学における系統樹作成法に示唆を受けたもので、「置換」と「欠失」の二種類の変化から形質的变化を検出し系統樹や系統ネットワークを作成する。そもそも、この手法を用いるためには遺伝子にお

ける「座位(Site)」に相当するものが必要となるが、CSAでは理念型モデルによって設計されたクラスをリスト化して座位の代わりとしそのクラスのインスタンス化の可否から欠失と置換の検出を行う。一方、TTAは、時間関係が定義されている場合に、インスタンス間の時間的な位置関係を位相としてみなし、歴史的に重要なイベント間関連に関する仮説導出を促すために機能する。時間位相は、地理情報標準における「時間情報スキーマ(ISO 19108)」において定義されている概念の一つで、同標準の「時間位相複体要素(TM_TopologicalComplex)」はネットワーク構造を持つ。TTAでは、この構造を利用してグラフ理論に基づく中心性評価(Centrality Measures)の算出を行う。

ITMAにおけるメソッドとインタフェースに採用されている手法は一般的であるが、現実世界に対する認識構造を反映した分析を可能とすることを目的に整理されている。

2.3. ITMAにおける課題

ITMAは、人文科学における客観的方法を追求すると同時に、OOGISのための分析手法を提供する。これまでの研究成果を経て、ITMAの基礎的な枠組みは固定されつつあるが、実用化のための課題は多い。

まず、実装済みの二つのインタフェースは、インスタンスそのものの分析に対しては不十分な点がある。まず、CSAは、インスタンス化時の実装の有無を検出するため、その属性値を考慮に入れることはできず、属性や多重度を認めるクラスや属性の問題を解決できない。そのため、実際に適用する場合には、個々のインスタンス構造のみを抽出して分析することになる。例えば、XMLドキュメントからXMLスキーマを自動的に生成するコードを用いて、個々のインスタンスから個別にスキーマを再生成し、本来のスキーマとの比較によって分析を行う必要がある。TTAに関しては、インスタンス間の時間的な関係を分析することはできるが、インスタンスの特性を加味した分析には十分に対応できない。したがって、現在のITMAでは、実質的に、複雑な構造を持つインスタンスを直接的に分析する手法が存在せず、メソッドによって符号化された情報を解釈に反映することもできないことになる。CSAとTTAそのものの有効性を否定はしないが、インスタンス間の類似性を直接的に分析できる手法が必要である。

第二の課題は、分析プロセスの可視化の問題である。ITMAは、方法論としての枠組みを提供するのが主たる目的であるため、メソッドとインタフェースの拡張に関する自由度は高く、様々な分析手法を組み合わせることも可能である。しかしながら、メソッドとイ

インタフェースを合わせても数えるほどしか存在しないにも関わらず、分析プロセスは複雑になりがちで、結果として透明性にかけている。この課題に対しては、ITMA を用いた研究を記述するためのプロトコルが必要である。

3. NCD を用いたインタフェース

3.1. NCD の算出法

まず、インスタンス間の類似性を直接的に分析できる手法について「正規化圧縮距離 (NCD: Normalized Compression Distance)」と呼ばれる類似距離について注目した。NCD とは、正規情報距離 (NID: Normalized Information Distance) をベースに提案された類似距離算であり、計算不可能なコルモゴロフ記述量に基づいて以下の式で表現される[3]。

$$NCD_{xy} = \frac{C_{(xy)} - \min(C_{(x)}, C_{(y)})}{\max(C_{(x)}, C_{(y)})} \quad \dots (1)$$

また、ラジスキー距離との対応関係から以下のように変形した式を好む場合もある[2]。

$$NCD_{xy} = \frac{2 \cdot C_{(xy)} - (C_{(x)} + C_{(y)})}{C_{(xy)}} \quad \dots (2)$$

ここで、 $C_{(x)}$ および $C_{(y)}$ は文字列 x および文字列 y を圧縮したファイルサイズで、 $C_{(xy)}$ は x と y を結合して圧縮したサイズである。この式が示す通り、仮に文字列 x と文字列 y が著しく異なる場合、 $C_{(xy)}$ は大きくなり、結果として NCD 値も大きくなる (1 に近づく)。逆に文字列 x と文字列 y が限りなく等しい場合、 $C_{(xy)}$ は小さくなり、NCD 値は 0 に近づく。NCD は二つのオブジェクトに対して用いるため、全オブジェクトの関係は距離行列 (NCD 行列) で表すことができる。なお、文字列 x および y は、実際にはファイルオブジェクトに置き換えることができるため、BMP などの画像ファイルや MIDI などの音声ファイルなどでも用いることができる。NCD 値は、圧縮アルゴリズムに大きく依存するが、 $C_{(xy)}$ のサイズを限りなく小さくできるような圧縮アルゴリズムが理想的である。NCD に用いることができる標準圧縮としては GZIP, BZIP, PPMZ などが定義されているが、本研究で開発した「百花老式 (Ver.1.0.0)」は GZIP 圧縮のみを、「百花老式 (Ver.1.0.0) for JAVA」は ZIP 圧縮をサポートし、(2) 式を用いて NCD を算出する。現バージョンでは両者とも、テキストファイルおよび XML ファイルにのみ対応している。

NCD は単に圧縮してファイルサイズを取得するだけなので、他の手法と比較して、データの事前クリーニングの手間が大幅に省略できる。これは、単に分析の事前処理に費やすコストを低減するだけでなく、OOGIS のデータセットのように複雑な構造を持ったデー

タセットに対して有用である。例えば、前述

の CSA などでは、原則的に座位の配列は固定されていなければならない。これに対して、NCD は原理上、どのようなファイルオブジェクトに対しても有効である。つまり、ITMA においては XML ドキュメントとしてインスタンス化される個々のオブジェクトは、構造上の違い、多重度の違い、属性値の違いを同時に考慮した分析が可能である。

NCD の可視化方法について、NCD の提唱者らは、NCD 行列をカルテット法によって分類する方法を推奨しているが、大規模なデータには不向きであるため、Partition Around Medoids (PAM) と Neighbor Joining (NJ) アルゴリズムに基づく Neighbor Net (NN) を用いた分類およびグラフを用いた可視化法を検討している。以下では、一連の手法が「第三のインタフェース」として有効であるかについて検討する。

3.2. NCD を用いた応用研究

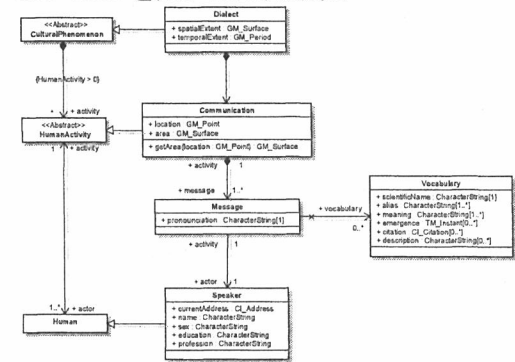


図2 方言のメタモデル

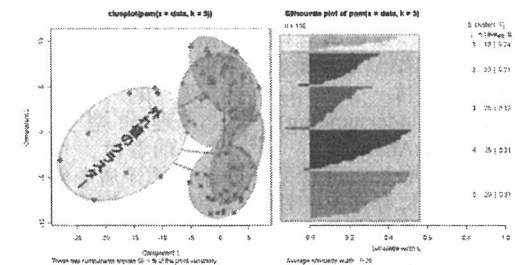


図3 PAMによる分類基数の確認結果

現在、NCD を中心とした一連の方法の有効性は「岐阜県における方言分布の研究」[10]において検証中である。この研究では文化現象としての方言を、図2に示した方言の「理想型モデル (Ideal Types Model)」として設計し、『岐阜県方言の研究』[9]に記載された岐阜県全域の方言語彙の地域差を示す10枚の分布図を元にインスタンス化した。参考にした10

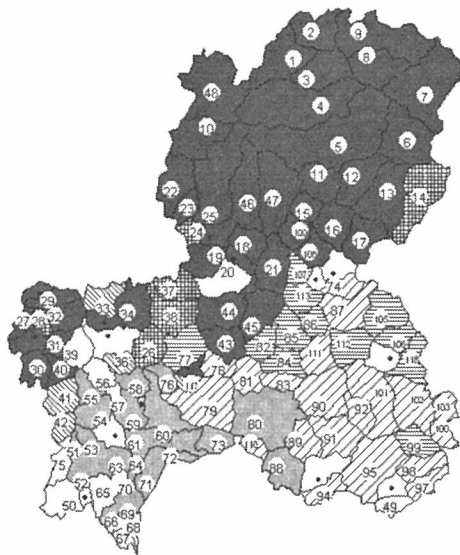


図4 NNによる分類を地理空間に投影

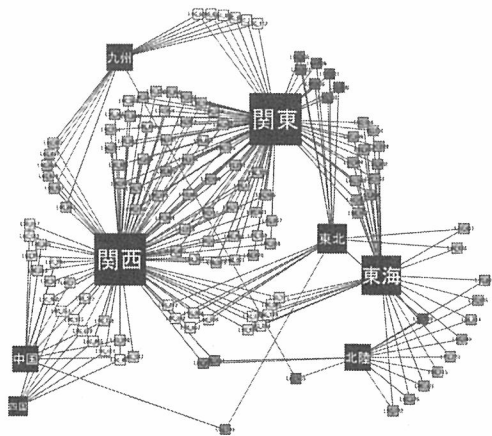


図5 「典型」との類似性を用いた分類
(ノードサイズは、度数中心性(Degree))

の地図は、それぞれ「唐がらし」，「唐もろこし」，「薬指」，「旋毛」，「竹馬」，「お手玉」，「粳米」，「馬鈴薯」，「真綿」，「里芋」という10の調査語彙の分布を示していて、それぞれの分布図では各々の調査地点における語彙がどのように発音するかが地図上にプロットされている。この資料を元に、125の調査地点をインスタンスとした、125のXMLドキュメントを作成し、GZIP圧縮を用いた「百花壺式(Ver.1.0.0)」を用いてNCD行列を作成した。

この研究では、最初にPAMを用いた分類を通して、分類基数および地理空間的な傾向

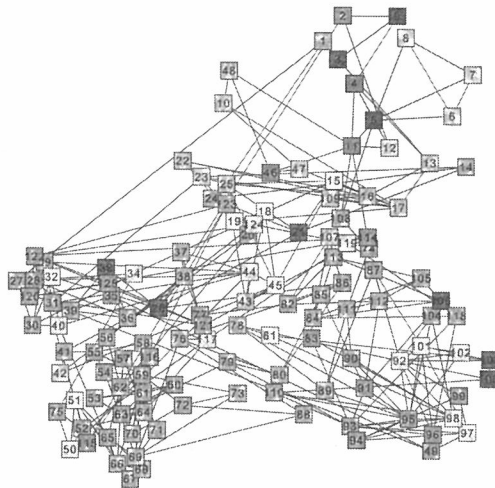


図6 地理空間に投影したグラフ
(図中左の部分に揖斐川)

を観察した(図3)。この分析の結果、PAMでは5クラスとし、はずれ値の1クラスを除く4分類とするのが適切であると解った。一方、NNによる分析では、関東系アクセントの戸入地域(調査地点28)が関西系アクセントの地域に囲まれているという状況や、大垣市における東西方言の対立といった先行研究の成果を最も忠実に再現していることが判明したことから([11], [13], [14]), 同手法を参考にした分類は適切であると言える(図4)。グラフを用いた分析では、理念型に見立てた「典型(Archetype)」として用意した日本全国の仮説的インスタンスを擬似的に実装し、近接する2点をエッジとするグラフを構築して、各インスタンスの地域的特性を観察した。その結果、関西系語彙と関東系語彙の中間的なタイプが揖斐川沿いに多く分布することが解った(図5および図6)。揖斐川上流域の旧徳山村では、関西系と関東系の両者の特徴を持っているとする先行研究[18]があることから、両研究結果についてはより慎重な検討を要する。

4. 分析プロセスの可視化

4.1. 分析プロセスの透明性

次に、分析プロセスの可視化について検討する。[10]の研究では、この課題についても取り組んでいている。この取り組みでは、UMLアクティビティ図を用いた、分析プロセスの記述法を検討している。

そもそも、分析対象の情報モデル化や処理工程の記述といった試みは、Individual Based Model(IBM)やAgent Based Model(ABM)に関連する分野では既に行われている。これら分野では、UMLクラス図とUMLアクティビティ図(あるいはUMLシーケンス図)を用いて分

Overview	Purpose
	State variables and scales
	Process overview and scheduling
Design concept	Design concepts
Details	Initialization
	Input
	Submodels

図7 ODDプロトコルの7つの要素([7]より転載)

ODD(Overview, Design concepts, Details)プロトコルと呼ばれるシミュレーションのドキュメンテーション標準を提案して(図7), このプロトコルでは Overview の記述に UML クラス図を用いることを推奨している。

Grimm [7]らによると, ODD プロトコルが登場した背景には, IBM が様々な分野で利用されている一方で, シミュレーションモデルを記述するための標準プロトコルが存在しないため, 結果として IBM の理解や再利用を阻害しているという問題がある。統計モデルのように一般的な数式によって表現可能な分析モデルの記述は一般的に, 完全かつ明瞭なので理解しやすいが, 対照的に IBM の記述は理解し辛く, 不完全で, 不明瞭で取っ付きにくいことが多い。さらに, ある IBM のモデルを批判的に利用したり再実装しようとするとき, 全体あるいは一部を実装する前に, 言葉で書かれた IBM のモデルを明示的な等式, 規則, 表に変換する必要がある。

ODD の Overview は 1.目的(Purpose), 2.状態の変数とスケール(State variables and scales), 3.処理の概要と工程(Process overview and scheduling), の三つのパートから成っていて, Overview を読み終わると IBM を実装するためのスケルトンコードを書く事ができるようになる。つまり, ドキュメント通じて再現可能なシミュレーションとなる。

ODD における Overview の考え方は, 人文科学全般にも通じるものがあり, ドキュメントから再構築可能できる方法は重要である。ITMA は, 主として UML クラス図を用いて認識構造を記述するため, ODD における Overview における 2 の部分はクリアできているが, 処理工程を記述する方法には言及していない。以下では, UML アクティビティ図を用いて分析プロセスを記述する方法について検討し, ITMA のためのプロトコルを定めるための基礎的研究とする。

4.2. UML アクティビティ図の利用

UML を構成する 13 の図式化言語において, UML アクティビティ図は, 機能的にはフローチャートと同じ役割を持っている。フローチ

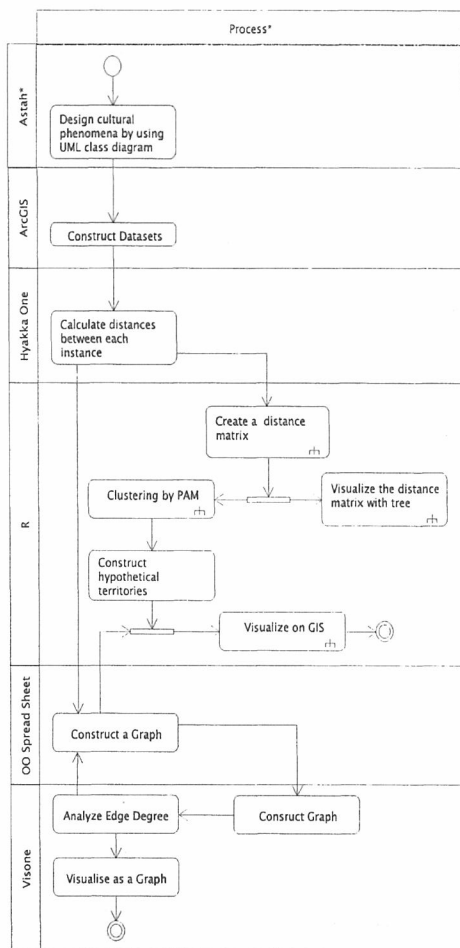


図8 UML アクティビティ図による分析プロセスの可視化

ャートと大きく異なる点は, 並列処理を記述できる点と, 処理(アクティビティ)を入れ子にできる点である。

図8は, 実際に, [10]の研究における一連の作業工程を記述したものである。この図において角が丸い矩形は, アクティビティ(作業)を表し, 矩形はオブジェクトを表す。また, この図においてはスイムレーンと呼ばれるシステム境界を使用してソフトウェアごとにフローを整理している。このように, 分析プロセスを記述することで, 一つの研究の中に含まれる処理プロセスは明確となり, どのような工程を経て解釈に至ったかを明示的に示すことができる。

ITMA は IBM/ABM と異なり様々な分析手法から構成されるため, 具体的な処理内容を記述しなければ再現可能なドキュメンテーションを作成することはできない。図9では, 通常は自然言語や制約言語によって記述するノ

謝辞

本研究を行う上で、様々な方々に有益なアドバイスを頂いた。まず、同志社大学の矢野環教授には、NCDについてのいくつかの技術的な問題について、University College London(UCL)の Enrico Crema 氏には ODD プロトコルの存在について教えていただいた。また、UCL 考古学研究所の Journal Club のメンバーとの議論や同クラブで取り上げられた文献のいくつかが本研究を進める上で大きなヒントとなった。こうした成果の背景には、日本学術研究会ならびに UCL の Andrew Bevan 氏の多大なる支援のおかげである。末尾ではあるが、以上の方々および組織に対して謝意を表する。

なお、本研究は、日本学術振興会特別研究員 PD/日本学術振興優秀若手海外派遣事業における研究成果の一部である。

参考文献

- [1] Barcelo, J. A. : *Computational Intelligence in Archaeology*, IGI Global, 2009.
- [2] Card, S.W.: *Information Distance Based Fitness and Diversity Metrics*. Genetic And Evolutionary Computation Conference , pp1851-1854. 2010.
- [3] Li, M.; Chen, X.; Ma, B. & Vitanyi, P. M. : *The Similarity Metric*. IEEE TRANSACTIONS ON INFORMATION THEORY, 1-13, 2004.
- [4] Huson, D.H. & Bryant, D.: Application of Phylogenetic Networks in Evolutionary Studies, *Molecular Biology and Evolution*, 23(2), pp.254-267, 2006.
- [5] Iwata, H & Ukai, Y. : *SHAPE: A Computer Program Package for Quantitative Evaluation of Biological Shapes Based on Elliptic Fourier Descriptors*, *Journal of Heredity*, 93(5), pp.384-385, 2002.
- [6] Open ABM: <http://www.openabm.org/>.
- [7] Grimm, V.; Berger, U.; Bastiansen, F.; Eliassen, S. ; Ginot, V.; Giske, J.; Goss-Custard, J.; Grand, T.; Heinz, S.K.; Huse, G.; Huth, A.; Jepsen, J.U.; Jørgensen, C.; Mooij, W.M.; Muller, B.; Prie, G.; Piu, C.; Railsback, S.F.; Robbins, A.M.; Robbins, M.M.; Rossmanith, E.; Ruger, N.; Strand, E.; Soussi, S.; Stillman, R.A.; Vabø, R.; Visser, U.; DeAngelis, D.L.: *A Standard protocol for describing individual-based and agent-based models*, *ECOLOGICAL MODELLING*, 198, pp.115-126, 2006.
- [8] ウェーバー M.: 社会科学の方法. (祇園寺信彦, 祇園寺則夫, 訳) 講談社学術文庫, 1994. (*Die 'Objektivität' sozialwissenschaftlicher und sozialpolitischer Erkenntnis*, 1904.)
- [9] 奥村三雄: 「岐阜県方言の研究」. 大衆書房, 1976.
- [10] 小野原彩香, 藤本悠: 方言分布から見た文化事象の流通と滞留-岐阜県における方言分布の事例-, *地理情報システム学会講演論文集 (CD-ROM)*, Vol.19, 2010.(採録決定済み)
- [11] 柴田武: 揖斐川上流のアクセント, 文字と言葉. 刀江書院, 1950.
- [12] じんもんこん 2010: <http://incomparable/sympo2010/>.
- [13] 徳山村の歴史を語る会: 徳山村のあけぼのを求めて -岐阜県揖斐郡徳山村遺跡分布調査中間報告-, 1984.
- [14] 服部四郎: 近畿アクセントと東方アクセントとの境界線, 音聲の研究. 第三輯. 文学社, 1930. pp.131-144.
- [15] 藤本悠: 地理情報標準応用スキーマ設計手法を用いた文化情報学的研究, *地理情報システム学会講演論文集 (CD-ROM)*, Vol.17, 2008.
- [16] 藤本 悠: 文化情報学および理念型モデル化分析法に関する基礎的研究 - 文化情報学における諸概念と方法について -, 博士論文 (同志社大学文化情報学部), 2010.
- [17] 藤本悠.: 学際分野研究としての文化情報学研究(2). *文化情報学*, 4 (1), pp.11-27, 2009.
- [18] 山田達也: 揖斐川上流域における基礎語彙について, 名古屋市立大学教養部紀要 人文社会研究, 22, pp.33-66