

意味構造記述のための有意義に制約された図法を求めて —概念化の ID 追跡モデルの提案—^{1,2}

黒田 航
独立行政法人通信総合研究所
けいはんな情報通信融合研究センター
kuroda@crl.go.jp

Keyword: characteristic semantics, diagramming method, geometrical theory of meaning construction, ID tracking model of conceptualization, visualization of semantic structure

1 言語記述における図の役割

ラネカー流の認知文法 (Langacker 1987, 1991a, b) で意味構造の記述に用いられる図, あるいはダイアグラム (diagram) は (恣意的とは言わないまでも) 解釈が一定でないことはよく指摘される. この傾向は年々増加し, 止まるところを知らない. 当初の提案, 分析の不備を補うため, あるいは説明力を増加させるため, 様々な補助的デバイスが考案されている (中村 2001). このような状況の中, ラネカー流の図の解釈が収斂するようには見えない.

このような記法の拡散の原因の一つは認知言語学における図の位置づけ自体にある. ラネカーは *across* の意味を記述する図の一つについて, 次のように述べる:

Note that I regard these diagrams as **heuristic** in character, not as formal objects. They are analogous to the sketch a biologist might draw to illustrate the major components of a cell and their relative positions within it. (Langacker 1991: 22, Note 9, 筆者によるボールド体の強調)

¹この論文で取り扱われている様々な問題に関して、黒宮公彦氏 (大阪学院大学) との討論がたいへん参考になった。この場を借りて感謝の意を捧げたい。

²この論文は『言語研究』に投稿されたが、採択されなかった。その際、次のような評を査読者の一人から頂いた。

ラネカーの図法の意味の理解不足。言語に中立的といっているが、そうはいかない。たとえば、英語では自動詞と他動詞の break ですむが、日本語では『割る』『壊す』『破る』『崩す』などが対応し、どう対応するのかが明らかではない。また英語には自他対応がないが、日本では自他対応がある場合がある: 英語: decide 日本語: 『決める』『決める』 実例の部分の経験的な観察を網羅的に行い、その上で自らのモデルがそれらの事実を有効に捕らえることができることを示すべき。

この査読者は Langacker の枠組みの批判をしているのか、私の修正案の批判をしているのか、自分で判っていないようだ。彼の指摘する問題は、私の提案に固有の問題ではなく「認知文法」という枠組み自体の問題である。更に、私は言語中立性が概念化の比較的抽象的なレベルに成立するとは言ったが、それが語彙レベルにあるとは言っていない。

図の「発見的」な役割はどの分野研究でも極めて重要であり、筆者はその価値を否定するものではない。図の役割をその用途に限るなら、それはそれで構わない。だが問題は、そのような「素描」程度の中途半端な記述精度しかもたない図を基に、しばしば極めて大胆な主張（例えば 1997 で「統語構造は記号構造に還元しうる」という大胆な主張）がなされることである。これは少し野心が過ぎるというものである。

更に言うと、この傾向は「文法がイメージ的な基盤をもつ」という（少なからずイデオロギー的な）想定によって是認されており³、それに歯止めがかかるようにも思われない⁴。実際、文法の記述が「文法がイメージ的な基盤をもつ」という想定に乗りかかっているのは状況は、記述的な妥当性の観点からは決して好ましいものではない。そのような想定がもし成立しなかったら、それに基づく記述は多かれ少なかれ、ご破算である。そのよう信憑性の乏しい想定に全面的に依存する企図は、控えめに言っても賭博的である。もっと好ましいのは、そのようなメタ理論的な想定に依存しない、真に抽象的で有用な意味記述ためのモデルをもつことである。そして、それは以下に示すように、可能である。

また、その想定が仮に成立し、認知文法の枠組みが記述的に多大な貢献をもたらすとしても

³実際、ラネカーは次のように述べる。

Lexicon and grammar form a continuum of symbolic elements. Like lexicon, grammar provides for the structuring and symbolization of conceptual content, and is thus imagic in character. When we use a particular construction or grammatical morpheme, we thereby select a particular image to structure the conceived situation for communicative purposes. Because languages differ in their grammatical structure, they differ in imagery that speakers employ when conforming to linguistic convention. [...] The symbolic resources of a language generally provide an array of alternative images for describing a given scene, and we shift from one to another with great facility, often within the confines of a single sentence. The conventional imagery invoked for linguistic expression is a fleeting thing that neither defines nor constrains the contents of our thoughts.” (1991a: 12, 筆者のボードでの強調)

著者には、ここでラネカーが image, imagery と呼んでいるものが正確に何であるかを言うことはできない。概念化の記述の多くがそれに基づいてなされるにも拘わらず (conventional) imagery と彼が呼ぶ対象について彼が 1991a: 5 で与えている定義 “I refer instead [of sensory images *à la* Shepard (1978) and Kosslyn (1980)] to our manifest capacity to structure or construe the content of a domain in alternate ways.” は、それが既知のものであるという想定の下でしか意味をなさず、実質がない。従って、文法をそのような未定義項目で「説明」するのは、控えめに言っても循環的である。

⁴同様の方向づけは、些か別のニュアンスで形式の空間化仮説 Spatialization of Form Hypothesis にも見受けられる。ここでは、例示のため、Lakoff 1987 と Deane 1992 から引用する。

Strictly speaking, the Spatialization of Form Hypothesis requires a metaphorical mapping from physical space into a “conceptual space.” Under this mapping, spatial structure is mapped into conceptual structure. More specifically, image schemas (which structure space) are mapped into the corresponding abstract configurations (which structure concepts). The Spatialization of Form Hypothesis thus maintains that conceptual structure is understood in terms of image schemas plus a metaphorical mapping. (Lakoff 1987: 283)

The Spatialization of Form Hypothesis claims that grammatical knowledge is grounded on conceptual metaphor—and conceptual metaphor operates by projecting embodied (image) schemas on other cognitive domains. (Deane 1992: 251)

, 解釈の定まっていない図式の使用は, それ自体として十分に大きな問題である. 曖昧な記述は, 事実の解明に貢献するより, 隠蔽に貢献する.

このような問題意識から出発し, この論文は意味記述のために有効な図法を ID 追跡モデル (ID Tracking Model: IDTM) という枠組みの下に提案する. 2節で本論に入る前に, IDTM の特徴を短く素描しておこう.

1.1 IDTM の特徴

1.1.1 特徴的意味記述

IDTM の目的の一つは, 表層の言語表現の背後, あるいは基底にあって, その解釈可能性を保証している抽象的構造のモデル化である. その構造は便宜的に, **特徴的意味** characteristic semantics, あるいは**特徴的意味構造** characteristic semantic structure と呼ばれる. IDTM の意図するのは, 特徴的意味構造のモデル化である. これゆえ, IDTM は厳密には言語構造のモデルではなく, その背後にある**概念化** conceptualization と呼ばれる抽象的構造 (ないしはプロセス) のモデルである.

IDTM における図は, 意味そのものではない. それは文法はイメージそのものではないという理由からも, 明らかである. この点で, IDTM における図の位置づけは, 例えばラネカー派の図の位置づけとは微妙に異なる.

特徴的意味記述は, 言語の表層の多様性に対し, 中立的である. 文法役割の標識の有無, 有意意味な単位の出現順序という因子には依存しない. この点で, 特徴的意味構造は, 最小主義プログラム (Chomsky 1995) 以前の生成文法で想定されていた深層構造 (Chomsky 1965), D構造 (Chomsky 1982) に類似した記述対象であるが, その技術的詳細は大きく異なる. まず, 特徴的意味構造からの**派生** derivation のような操作は想定されていない. 特徴的意味構造を表層の音韻・音声構造を関係づけるために想定しなければならない操作は, 連想的な写像で十分であろう. もちろん, これはその写像が簡単に定式化しえることは意味しない. 更に, 意味構造と音声構造とのあいだに中間的な表示レベルが存在する可能性に対しても, IDTM は中立的である. 実際, IDTM は, 比較的「浅いレベル」での意味構造の表示であり, それがどれぐらい「意味寄り」の表示レベルなのかは, 明らかではない. 「身体化された意味」のレベルは, 明らかにもっと深いレベルにある.

1.1.2 図を制約する必要性

IDTM に基づく意味構造の記述は, 通常, 一群の制約下で生成される図の形式で与えられる. これは, IDTM の狙いの一つが意味構造の効果的な**視覚化** visualization の手法を与えることだ

からである。

IDTM は言語記述, とりわけ言語の意味記述における図=ダイアグラムの役割を真剣に考える。IDTM で理解される図は, ことばによる廻りくどい説明を簡略化する単なる小道具ではない。図は, それ自体が記述言語であり, 一定の制約の集合, あるいは「文法」をもつと理解される。その制約の集合は図法 diagramming method と呼ばれる。図法が明示的に与えられない限り, 図の記述的有効性には限界がある。

IDTM が図に期待する効果は, 単なるヒューリスティクス以上のものである。実際, IDTM が強く要求することの一つは, 図の利用は効果的でなければならない, ということである。解釈の一定でない図しか与えられていない状況は, 場合によっては, 図がない場合よりも始末が悪い状況である。単に「分った気にさせる」程度の図が, 学術論文で主要な役割を演じるとしたら, その分野の研究に学術的な価値があるのか, 非常に怪しいと思われる。

理想的には, IDTM の図は分子式の図が抽象的分子構造の視覚化であるのと同じ意味で, 抽象的な構造の視覚化だと理解される。これから直ちに, 図は効果的なものでなければならないということが帰結する。実際, IDTM の記述は, この意味で効果的なものとして提案される。

IDTM が提供するのは個々の図ではなく, それらが満足すべき制約や条件の体系としての図法である。これにより, 一定の解釈をもつような図のみが許され, 恣意的な図の使用は大きく制限される。これは図の利用価値を向上させ, 効果的にするためにである。

1.1.3 特徴的意味構造の記述と最大主義

特徴的意味記述は, 言語表現の意味を利用者が(再)構成する際に利用する情報のすべてではない。それは, 発話内容が解釈可能となるために必要最小限の情報のみを扱う。それは語彙の意味の部分集合でしかない。

それゆえ, 状況的にによってのみ定まる意味は, 特徴的意味構造の記述には含まれない。これは一面, ラネカーの最大主義 maximalism (Langacker 1988) とは相い容れない想定かも知れないが, 基本的な記述のレベルで高い信頼性を提供するためには不可欠な条件だと考える。

状況的に与えられる意味は, 実際, 言語表現の解釈の非常に多くの部分を占める。この状況的意味の総体は文脈 context と呼びならわされている対象に相当する⁵。

IDTM では理論的整合性のため, 文脈を処理している機構は特徴的意味記述が処理されるよりも低次の(あるいは「深い」)レベルにあると想定する。つまり, IDTM は解釈に層化され

⁵文脈の効果は非常に古くから存在が知られているが, しかし, それがこれまで計算論的に満足のゆく形で定義されたことはないように思われる。従って, 文脈の効果の存在を指摘するだけでは, 実は何の解決でもない。筆者の理解する限りで, 生成的辞書 generative lexicon (Pustejovsky 1995) が限定された場合に関して, 文脈を明示的に処理しようとするアプローチであるように思う。

た処理レベルを想定し、意味の肉付け *elaboration* の程度はそのレベルの深度に応じて決まると想定する⁶。

意味記述に説明的な役割を与えるためには、解釈を得るために必要な情報の利用可能性に格差があると認識されるべきである。仮に最大主義が、意味解釈のために必要な情報のタイプやレベルを区別せず、それらが皆同じ「経費」で利用可能だと考えるならば、それは明らかに説明力に欠ける。情報の利用可能性に格差がなければ、解釈が収束するという事実すら説明し難い。これを説明する一つの可能性は、基本的なレベルの情報が「割安」で、状況的な情報が「割高」だと考えることである。

IDTM では、この利用可能性の格差の想定の下で、特徴的意味記述は、比較的割安な情報を提供し、状況的な意味は割高だと考える。解釈のための体系は、割高な高次の処理を必要に応じて、消極的に駆動されるような体系の方が効率的であり、望ましいという洞察に基づく。

1.2 議論のあらまし

以下、2節では、ラネカーの認知文法の枠組み (Langacker 1987, *et seq*) で提案されている図法の批判的検討を行なう。3節では、その修正的拡張として IDTM を提案し、その詳細を説明する。4節では、IDTM が提供する具体的分析を示す。

\$4.1では、英語の *break* を例に採り、*X break Y ~ Y break* の自他形の交替を IDTM の観点から分析する。ここでの記述は、IDTM の技術的詳細、特に IDTM において図がどのように制約されるかに関して有益な情報を含む。

\$4.2では、*break* に関して *X break Y into Z* にある *into Z* 結果述語の付加、*X break Y with W ~ W break Y* の道具主語現象を、IDTM で分析する。この節では、玉突きモデルに基づく同様の構文のラネカー自身の分析 (1987, 2000) を、IDTM の分析と直接比較し、その結果に基づいて、IDTM が玉突きモデルに基づく記述に対し上位互換性をもつことを主張する。

\$4.3では、*X make Y ~ X make YZ ~ X make YVP* の対比を分析する。この際、基本モデルの拡張のため ID 階層と属性地図 *attribute geometry* の概念が導入される。これにより、ID の内部構造の記述が可能となり、IDTM の記述力は大幅に向上する⁷。

\$4.4では、\$4.2, \$4.3での *break* の挙動に対応する日本語での現象の分析を示す。“壊れる” ~

⁶ここで想定している層別処理システムが文法のモジュール性に関してどのような意味をもつのか、現時点では判定し得ない。文法が（意味部門、語用論部門、統語部門、語形成部門のような）幾つかの「部門」をもつというのは、筆者が想定していることではない。そうではなく、想定しているのは、非常に数多くの「部分」が相互作用しながら上位レベルを構成する「複雑系」としての文法である。黒田 (2003) は形態論の扱いに関して、この相互作用的、創発的文法観を明示的に示している。

⁷本論文では詳しく扱わなかったが、これは IDTM をメンタルスペース理論 (Fauconnier 1985, 1997) に対し上位互換性であるような強力な理論とする。

“壊す”の対比が分析の対象となる。これは IDTM の記述の言語中立性を例示するためである。

§4.5は $X \text{ give } Y \text{ to } Z \sim X \text{ give } Z \text{ Y}$ の交替現象の記述に当てられる。認知文法の枠組みで同様の現象を分析した中村 2001 と IDTM の提供する分析の比較を行う。結論的には、両者は、ほとんど同様の洞察を幾つか提供しているが、中村の依拠しているラネカー流の図法は、それらを首尾一貫した形で表現するほど制約されていない、と主張する。このもとに、IDTM の提供する図法は、認知文法の図法に対し上位互換性をもつことを主張する。§4.6では、現時点での IDTM の問題点の一つとして、それが認可する図の恣意性の一つを指摘する。

2 認知文法の図法の問題点

この節では、ラネカーの認知文法の枠組みで用いられる図に内在する自己矛盾的な問題を幾つか指摘し、その修正を提案する。具体的には、図法には一貫性が欠如し、冗長性が高すぎることを指摘する。図法の恣意性を制約するために、適性確定の要請を提案する。

2.1 ラネカー流の図法に内在する記述力の欠如の問題

Langacker (2000: 149, Fig. 5.2) は、例えば(1)の意味構造をFig. 1にある図で記述できるとする。

- (1) *Alice saw Bill.*

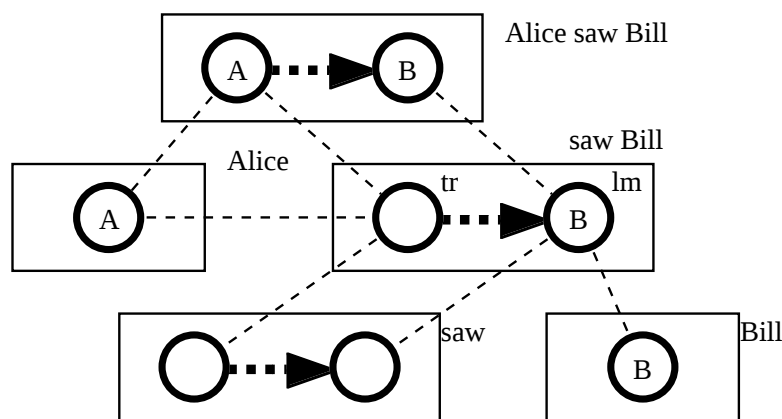


Fig. 1

この主張には、少なくとも次のような問題がある。

- (2) a. *Alice, Bill* のベースに何も無いのに理由が与えられていない。
- b. *saw* と *Bill* の組みあげ (assembly) が *Alice* と *saw* の組み立てよりも優先され

- る理由がまるで明らかではない。つまり, VP 現象の存在理由が示されていない
- c. 破線で示されている関係の意味が不明瞭. 実際, 破線は, この図では統合 (integration) の効果と同一指示 (co-indexing) の両方の効果を担っていて, 一貫性に欠ける
 - d. 主語名詞句, 目的語名詞句, 動詞のおおのの個別のプロファイルが, 節全体のプロファイルに貢献する仕方は, まったく冗長で, 無駄である (動詞としての SAW と節としての x SAW y を区別するのは A, B のラベルだけで, それは結局 $x=A$, $y=B$ という情報を表しているだけである)
 - e. V, VP, S レベルの肉づけ elaboration を $x=A$, $y=B$ を示すアノテーションなしに区別が不可能である
 - e'. 部分の意味の全体の意味への貢献が (それが必ずしも構成的 (compositional) 関係とは言えなくても) 正しく表現されているとはいえない
 - f. 同一指示がどういう認知的・計算的資源に基づいて決定されるのか, そのことがまったく示されていない. これが提供されないかぎり, このような記述が自明な形で統語構造を与えているとは考えられない

特に, (2)d は, ラネカー自身 (1987: 184) がベース/プロファイルの区別を用いて円弧 (arc) の概念に与えている記述と矛盾している. それでは, 円がベースにあり, その周囲の一部が太線で示され, プロファイルが当たっているとされている. **プロファイルが当たっていない部分が存在することが重要なのに**, Fig. 1では動詞(例えば *see*)の意味記述では, *tr*, *lm* を含めた全体にプロファイルが当たっているかのように示されている. これでは, 節レベルの意味と語レベルの意味とを区別するのが, 補助的手段である A, B のラベルづけ有無のみになる. 同様のことがベースに何もない名詞句に関しても言える. 結論として言えるのは, 認知文法で実践されている図法は, 明らかに正しく制約された, 効果的な図法ではない.

2.2 図の記述力向上のための提案

すでに導入部で指摘した恣意性, および(2)で指摘した問題を回避, 解決するために, 次のFig. 2にあるような図法を提案する.

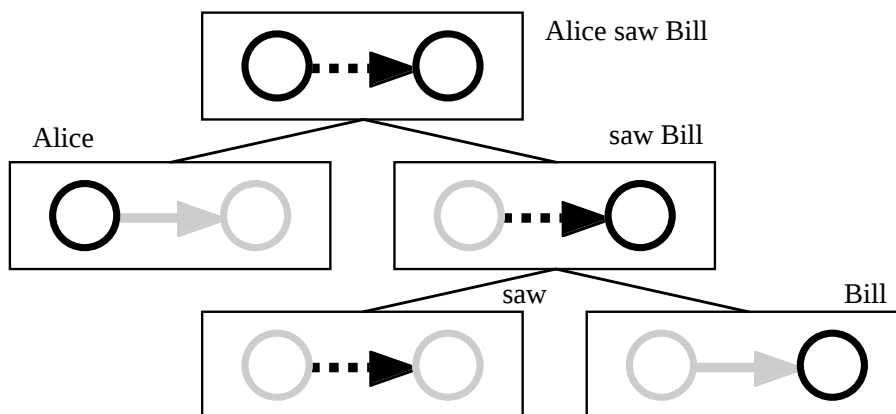


Fig. 2

このような図式の正当化には詳細な議論が必要であるが、この場はその目的のためには適切な機会ではあるまい。詳細な議論については, Kuroda (2000), 並びに, その基本概念の解説である Kuroda (2001) を参照されたい。Kuroda (2000) の基本的な主張は次のようなものであった:

- (3) a. 統合において, 全体の意味は部分の意味から厳密に合成されるのではなく, 部分に分散されている意味が重ねあわされる
- b. 主語のスキーマ, 目的語のスキーマ, 動詞のスキーマが存在し, 統合はそのようなスキーマの事例化 *instantiate* によって形成された部分に関して起こる
- c. 統合に利用される部分の表示は冗長性をもった形で与えられ, それらの重ねあわせの際には部分的な不整合も許容される
- b'. 従って, 単一化 (*unification*) は厳密には成立していない (この点で, 統合は Fauconnier 1997 のブレンドに近い性質をもつ)
- d. 矢印に非対称性を示す向きがあれば, *tr*, *lm* のアノテーションは単なる蛇足である

さて, 仮に Fig. 2 に提案された図式化が正しいとすると, 次のことが含意される:

- (4) a. 主語名詞句, 目的語名詞句には, 動詞的要素の主語, 動詞的要素の目的語という文法関係 *SRO* (ただし $R = \{V, P, \dots\}$) をベースにして成立するプロフィールが存在する⁸
- b. 動詞の項 (a.k.a., *elaboration sites*, e-sites) は, 項構造の記述でプロフィールされているが, それは表層で音形となって実現されているという情報とは区別されなければならない。そうでないと, プロフィールがほかの要素 (e.g., 主語NP, 目的

⁸この文法関係がラネカーの主張通り *trajector* と *landmark* の差異化に還元しうるかは明らかではない。

語NP) から受け継がれたのか, それとも始めから語彙的に指定されていたのか, 区別できない

- c. 統合の際にベースを共有する要素の間の一致=マッチングが前提となるので, 同一指示は自動的に満足され, 特に指定する必要がない⁹
- c'. ただし, 同一指示 (あるいは, 変項の拘束) と統合の効果は同じではなく, プロファイルの継承によって区別される
- d. 動詞句の効果は, 統合の順序とは独立に特徴づけられる必要があり, それは *tr/lm* の非対称性に基づくと考えられなければならない
- d'. ラネカー自身が (2000: 153, Fig. 5.3) で与えている (1) と *Bill saw Alice* との区別は, 認知的・計算的根拠の与えられていない同一指示 (あるいは変項拘束) が前提となっていて恣意的であり, 妥当だとは考えられない
- e. もっとも記述的に効果があるのは, ベース/プロファイル関係であり, それ以外の効果 (例えば同一指示を示すための破線) は, 結局, 蛇足に過ぎない
- f. 音形は意味 (正確にはプロファイル) に対して, ラネカーが考えている以上に類像的 (iconic) な関係にあるが, これは, 彼の推奨する記号的文法観 (symbolic view of grammar) には, より好ましい特質なはずである

ただし, 次の問題が新たに生じる:

- (5) 動詞の項構造において, 項の実現/未実現, 飽和/未飽和の区別は必要であるから, それをプロファイルの有無と別の仕方で表現しなければならない

この問題は, 実現されていない項のプロファイルをぼかすという形で解決されている. 実を言えば, e-sites の未実現は, 以前は (例えば 1991b: 36) 網かけで区別されていたのだが, この区別は, どのような訳か最近では失われている¹⁰. しかし, 網かけの効果によっては, 項のベースにある主語, 目的語のような潜在的な関係が表現できないので, ここで採用したぼかしの方が効果的であろう.

2.3 適確確定の要請

⁹別の観点からすると, これは挿入が環境指定成分をもつということである. ラネカーの対応 *correspondence* の定式化は, 生成文法の文脈で言う文脈自由 *context-free* で働く一般化変形 *generalized transformation* と実質的に同値であり, あまりに制約が緩い. なお, 認知文法で仮定されている操作は挿入ではないと言うのは, その正当化の理由にならない. 構造の生成に用いられる装置の「名称」は問題ではない.

¹⁰更に言うと, Langacker 1997 では同様の例を説明する際に *tr/lm* の区別が垂直方向に与えられていたが, ここでは水平に変えられている. これがどういう理由によるのかは, 判読不能である.

Fig. 2 に体现された理論的拡張を動機づけるもので最大級に重要なのは、次の要請である:

(6) **適確確定の要請 (Proper Specification Requirement: PSR)**

部分の意味 (プロフィール) の全体の意味 (プロフィール) への貢献の仕方が正確に確定されている必要がある

ラネカー風の図法で本質的に問題なのは、この PSR の条件を満足しないからだと考えられる。実際、この PSR の下では「典型的」なラネカー風の意味記述は、たいてい失格である。その理由は、部分の意味がどのように全体の意味に貢献するのか、図から自然に読み取れないからである。

実際、PSR の下では次のことが派生的に要請される:

(7) **理想的にはどんな文に対しても、その要素の一つ一つについて、それがどのような特徴的な意味をもっているのかを—例えばプロフィール化 (profiling) という手段を用いて—明示できなければならない。**

これは非常に厳しい要請であり、冠詞や否定辞の特徴づけなどを考慮に入れるだけで、その実現が困難なのはただちに明白になる¹¹。けれども、これは図を有意味に制約するために追及するには十分に値する指針だし、それを満足するような枠組みを構築できないか試してみる価値は、十分にある。以下ではその可能性を、IDTM という枠組みで素描する。

3 ID 追跡理論の導入

Fig. 2 の提案が正しい軌道に乗っているなら、それによって Langacker 学派の研究で提案されている図の性質が大きく変わる。それらを PSR を満足するように変更しなければならないからである。この節では、具体例を通じてそのことを例証する。

3.1 ラネカー流の概念構造のモデルの根本問題

ラネカー (Langacker 1987, 1991a, b) の玉突きモデル (billiard model) と呼ばれる図式法は**動作の連鎖** (action chain) の概念化が (メタファー的に) 取りこまれているという考えの上に成立している。それによれば、動作が基本的で、それによって生じる状態変化は派生的だと考えら

¹¹この点は、黒宮公彦氏との討論の結果明らかになった。この場を借りて、感謝したい。

れる。これはクロフト (Croft 1991) とともに共有され、認知言語学で広く受け入れられている考えであるが、これは例えば、関係語一般の項構造の記述に十分なほど精緻なものであるとは考えられない。

以下に対案として素描されるモデルでは、このメタファー的概念を前提としない。そこでは、状態変化が中心的な役割を演じ、動作はその解釈のために導入される媒介的なものと理解される。

3.2 ID 経路の概念と固有 ID 仮説

このような概念的変更の基盤にあるのは、次の ID 経路 (ID Path) の概念である。

- (8) a. **固有 ID 仮説:** 概念化可能な任意の要素は、常に固有の ID をもつ
- b. **ID 経路仮説:** そのような ID は時間に沿った (抽象的な) 軌道、あるいは経路として概念化される

このモデルに従えば、概念化の本質は ID 経路に展開する変化を追跡することにあると考えることにある。このような理由から、筆者は自分の提案するモデルを **ID 追跡モデル (ID Tracking Model: IDTM)** と呼ぶ。

3.3 IDTM は何のモデルか？

IDTM は、言語構造のモデルではなく、その背景にある**概念化**のモデルである。この点は誤解のないように、はじめに強調しておく。

IDTM は上述の特徴的意味構造を定義する。言語の文法は、その特徴的意味構造と音声構造との対応づけの(記号的)体系だと理解される。

3.4 ID TRACK とは何か？

次に、IDTM の基本となる ID track の概念を説明する。

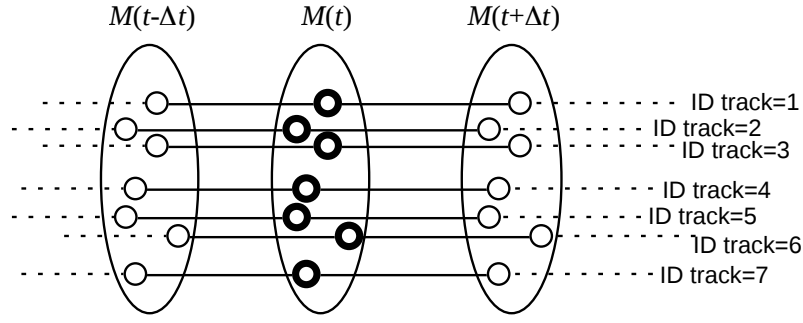


Fig. 3

この図は, ID=1-7 のトラック束 (track bundle) と, その三つの時点 $t-\Delta t$, t , $t+\Delta t$ での切断面を示している. $M(t-\Delta t)$, $M(t)$, $M(t+\Delta t)$ がおのおのの時間切断面である. その時間範囲にある ID 経路は実線で, 範囲外の経路は破線で示している. つまり, この図は七本の ID 経路を示している.

このように定義される M は Fauconnier (1985, 1997) のスペースに相当し, この点で IDTM はメンタル・スペース理論 (MST) の拡張という側面をもつ.

ID をもつ個物は相互作用する. その仕方を抽象的に表現することが, IDTM の基本となる. 次の図は, それを三体, X , Y , Z の場合に関して図式化したものである.

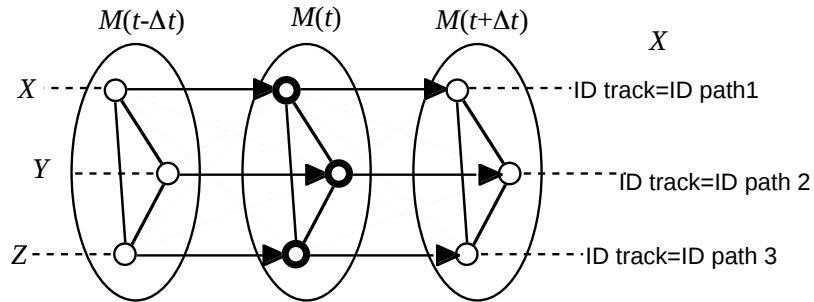


Fig. 4

この図で $M(t-\Delta t)$, $M(t)$, $M(t+\Delta t)$ 内部にあるのは静的な二項関係 (e.g., $R(X(t), Y(t))$, $R(Y(t), Z(t))$, $R(Z(t), X(t))$) のネットワークであり, 二つの時間切断の間の相互関係は, 多対多の動的な二項関係 (e.g., $R(X(t), X(t'))$, $R(X(t), Y(t'))$, $R(X(t), Z(t'))$) のネットワークである. 静的な関係は, 主に位置関係, 所有関係などであり, 動的な関係は, 使役などの相互作用である.

$R(x, y)$ は x, y 間の非対称的関係を表わすものとする ($R(x, y) \neq R(y, x)$). $R(x, y)$ は, 図では $x \rightarrow y$ ($\neq x \leftarrow y$) で表わされる. この条件の下で, 関係のネットワークは二項関係 $R(x, y)$ の有向グラフとして表現される.

一般に $M(t)$, $M(t')$ の対に成立する関係のネットワークの詳細は, 次のようになる.

- (9) I. 同一 ID のもとでの状態の(不)変化のクラス [再帰的]
- a. $R(X(t), X(t')), R(Y(t), Y(t')), R(Z(t), Z(t'))$
 - b. $R(X(t'), X(t)), R(Y(t'), Y(t)), R(Z(t'), Z(t))$
- II. 異なる ID をもつもの同士、同一時間内での相互作用のクラス
- c. $R(X(t), Y(t)), R(Y(t), X(t)), R(Y(t), Z(t)), R(Z(t), Y(t)), R(Z(t), X(t)), R(X(t), Z(t))$
- III. 異なる ID をもつもの同士、異なる時間内での相互作用のクラス
- d. $R(X(t), Y(t')), R(X(t), Z(t')), R(Y(t), Z(t')), R(Y(t), X(t')), R(Z(t), X(t')), R(Z(t), Y(t'))$

クラス I は再帰的な相互作用とも考えられる。

時間の非対称性から考えて、I.b は意味をもたないが、特殊な場合に、それが意味をもつ可能性はないとはいえない。事実、受身の基礎となる概念化は時間の向きを逆行していると考え、自然に取り扱える可能性も示唆されている。

3.5 特徴的な相互作用の選択＝プロファイル化

Fig. 4 にあるような ID 経路と動的・静的関係のネットワークからの有意な成分を（際立ちの相対的な差に基づいて）選択するプロセスが概念化だと考え、それがラネカー (Langacker 1987, 1991a, b) の用語でプロファイル化と呼ばれている説明項に相当すると考える。

これが正しいならば、動詞を始めとして、関係的な語彙の意味は ID 束内部のネットワークの关系的意味成分の組み合わせで表現できる。その具体例は 4 節以降にたくさん出てくるので、ここでは詳細は論じない。

3.6 IDTM の図法の基本要素

次の図に示したのが、IDTM が認可する図の基本要素である。これが IDTM を使って意味記述をするものにとっての「パレット」となる。

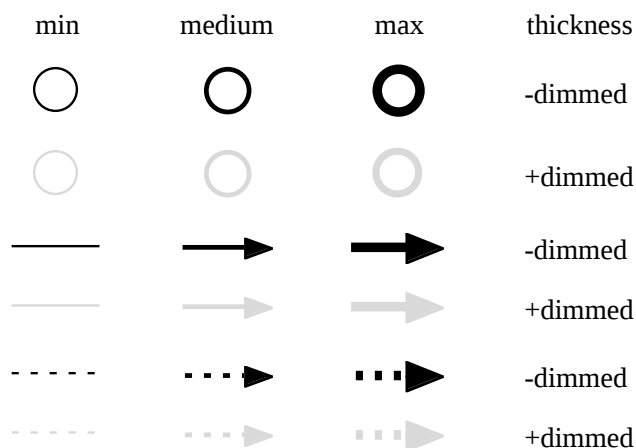


Fig. 5

この図では, 左から右へ際立ちのスケールがある.

現在導入されているパラメタは, (i) 三段階の線の太さ, (ii) 矢印の有無, (iii) ぼかしの有無, (iv) 破線と実線の対立がある. これらの図形と視覚的効果の意味は, 実際の使用の応じて随時説明する.

Fig. 5にある以外の要素を「非公式」に, 補足的な目的のために導入することは禁じられてはいない. だが, そのような拡張は公式の意味をもたず, 従って, そのような補足的な要素 (例えば, 丸の代わりに三角や四角を, 直線の代わりに波線を使う, 丸の内部に構造を入れこむ等々) が「説明的」な役割を担ってはならない.

3.7 IDTM の目標

IDTM の目標は言語そのものの記述というより, 言語の背後にある抽象的な構造の特徴づけであるが, その構造が純粹に意味的なものなのか, あるいは統語的な情報を含むものなのか, という問題設定には中立的な立場をとる. 実際の面で, それをハッキリさせる必要は特にないし, その是非を延々と論争するのは, 時間のムダである. その問題の答えを曖昧にしたままで記述を進めることのできる枠組みがあれば, それに越したことはない.

3.8 状態の軌跡としての ID 経路

無用な誤解を招かないように, IDTM の抽象的な性質は, 始めに強調しておく必要がある.

IDTM の成立基盤は, エネルギーの受け渡しによって引き起こされる動作の連鎖ではない. それはむしろ, 様々な原因(そのほとんどはエネルギー伝達のメタファーではうまく概念化しえない)で生じる状態の軌跡の束であり, それを追跡する操作を視覚化することが, このモデルの根本的性質である. 玉突きモデルでは抽象的な「力」が根本的メタファーだったのに対し,

IDTM では抽象的な「移動」が中心的なメタファーになる。

IDTM の図の上では物理空間的な位置の変化は単なる状態の変化から区別されない。この点は特に注意が必要であろう。この正確な意味は、位置の移動は状態の変化の特殊な場合だと捉えられるということである。個物の状態の追跡は、それ自体が自律的な認知的構成物であって、空間的な位置移動の抽象化だとは必ずしも考えられていない。

3.9 IDTM と文法のメタファー的基盤

これは、IDTM をいわゆる「認知的」な諸理論から区別する基本的な性質の一つであるかも知れない。主流の「認知的」な諸理論は「説明」のため、多かれ少なかれ意味拡張を理論構築の基盤に置く (e.g., Deane 1992, Lakoff 1987, Langacker 1987, *et seq.*)。

IDTM はこれとは異なった方向性をもつ分析手法である。IDTM は、特に現象の「説明」を提供しようとはしない。それは IDTM が考案された理由ではない。IDTM が目指すものは、あくまでも精緻な意味記述のために有効な道具立てである。

これは退歩ではない。科学の進歩を見るかぎり、正しい記述にはたいてい正しい説明が自然についてくるように思われる。

文法がメタファー的基盤をもつかどうかは、現象の「認知科学的な説明」の観点からすれば興味深い問題であるが、その信憑性は疑わしく、それに関して中立的な記述の枠組みがあれば、それに越したことはない。

IDTM と同様に動詞の意味記述を力のメタファーに還元する方向づけに反対する主張には、例えば、定延 (1993), Sadanobu (1995) がある。実際、彼がラネカーの玉突きモデルへの代替案として提案した「カビ生え」モデルは、幾つかの点で IDTM の先駆となっている。

4 ID 追跡理論が提供する具体的分析

この節では IDTM が提供する分析を、幾つかの簡単な事例を通じて紹介する。\$4.1では、英語の *break* を例に採り、*X break Y ~ Y break* の自他形の交替を IDTM の観点から分析する。\$4.2では、*break* に関して *X break Y into Z* にある *into Z* 結果述語の付加、*X break Y with W ~ W break Y* の道具主語現象を分析する。\$4.3では、*X make Y ~ X make YZ ~ X make YVP* の対比を分析する。\$4.4では、\$4.2, \$4.3節での *break* の挙動に対応する日本語での現象の分析を示す。\$4.5は *X give Y to Z ~ X give Z Y* の交替現象の記述に当てられている。ここでは、中村 2001 と IDTM の提供する分析の比較を行う。\$4.6では、現時点での IDTM 恣意性の一つを指摘する。

4.1 自動詞形と他動詞形の交替: *BREAK* の分析

次の(10)にある *break* 自他形の交替を考察する.

- (10) a. $X \text{ break } Y$. e.g., *John broke the window*.
 b. $Y \text{ break}$. e.g., *The window broke*.

このような自他形の交替は、いわゆる**非対(象)格性** (unaccusativity) を示す動詞群 {*break*, *melt*, *move*, ...} に特徴的な振る舞いである.

IDTM は、この構文の交替を次のFig. 6とFig. 7にある図によって特徴づける.

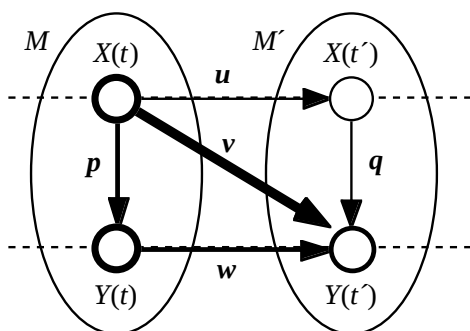


Fig. 6

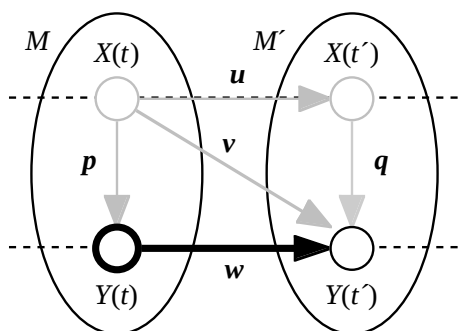


Fig. 7

Fig. 6に示した図は (10)a の, Fig. 7に示した図は (10)b の特徴的意味を記述するものである.

4.1.1 意味の幾何学的理論

u, v, w, p, q のような成分は、特に**意味成分** semantic components, あるいは**意味ベクトル** semantic vectors と呼ばれる。これらは組み合わせられて、複合的な意味構造を表現する。例えば, $v = x \text{ BREAK } y$ ($x \neq y$), $w = y \text{ BREAK } z$ ($y = z$) として、次のように書ける。

- (11) a. $v \Rightarrow p + w$
 b. $p + w \Rightarrow v$

(11)aは v が意味成分 p と意味成分 w とに分解可能であることを, (11)b は p と w とのベクトル和が意味成分 v であることを, おのおの表わす. $\alpha \Rightarrow \beta$ と書く時, α の長さが1以上なら $\Rightarrow$ は分解を, β の長さが1以上なら $\Rightarrow$ は合成を表わす. α, β の長さが共に1以上なら $\Rightarrow$ は変換を表わす.

意味成分の合成・分解のアイデアは、60年代後半から70年代前半にかけて**生成意味論** generative semantics 学派の分析手法として盛んだった**語彙分解** lexical decomposition の考

え (McCawley 1971) に非常に近いということは、補足的に指摘しておく¹²。黒宮 (2003) がこの点に関して、興味深い分析を行っている。

同一の表層形 *break* は、Fig. 6では *v* に、Fig. 7では *w* に対応する。Fig. 6で *p* は *Y(t)* の対(象)格性 (accusativity) をマークする意味成分である。いずれの場合でも、主語は意味成分の源である。

英語の場合、代名詞形 (e.g., *him* vs *he*) を除き、原則として *p* は主語と関係的要素 (e.g., *V*, *P*) との相対的位置をコードする情報に対応するが、いずれ§4.4の日本語の‘壊わす’～‘壊れる’の対比分析で明らかになるように、*p* は形態論的に実現することも可能である。従って、IDTMの記述は文法関係を抽象的に記述するレベルにあると主張することが適切であると考えられる。

Fig. 6とFig. 7にある図は基本構造を共有しているが、プロファイルの効果が異なっている。Fig. 7では、*u*, *v*, *p*, *q* の成分はベースにあるが、ぼかしが符号化しているように、表層形には顯示しない。これは(10)bの非対(象)格性の自然な記述となっている。

IDTMは多くの面で、動詞の意味を関与する項の相互関係の幾何学的関係として抽出しようとする。その点で、これをもし、認知文法が直感的に把握できる内容をもつ図の隠喩性に訴えるやり方に対比させるならば、IDTMは意味構成の幾何学的理論 (geometrical theory of meaning construction) あるいは、もっと野心的に幾何学文法 (geometrical grammar) の一翼と呼ばれても良いかも知れない。

4.1.2 IDTMは適性指定の要請を満足する

次のFig. 8に示す統合の操作によって、Fig. 6に示した図は(6)のPSRを満足する。統合の関係は破線で示した。

¹²語彙分解は近年、最小主義プログラム (Chomsky 1995) で、かなり矮小化された姿で再興された。詳細は、Pullum (1996) を参照。

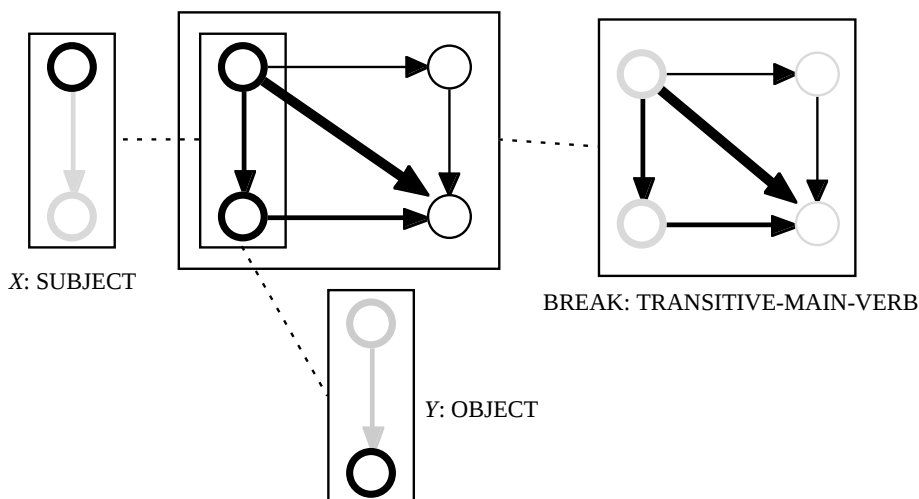


Fig. 8

因みに, X , Y , $BREAK$ はおののおの SUBJECT, OBJECT, Transitive-Main-VERB スキーマの実現/事例化である.

このFig. 8は $X \text{ break } Y$ の特徴的意味構造に対し, X , Y , $BREAK$ の特徴的意味構造の貢献の具合を示すものである. この図では, 典型的なラネカー流の図と異なり, どの部分の意味がどのように全体の意味に貢献するかに関して, 曖昧性はない. これが PSR が満足されているという主張の根拠である¹³.

4.1.3 IDTM の基本: 図法の有意味な制約

ここで少し詳しく, IDTM がどんな図を認可するのか見てみよう.

例えば, IDTM は, 次の性質を満足する図のみを認可する.

(12) 関係の基本的非対称性

i. $(x \rightarrow y) \Rightarrow \neg(y \rightarrow x)$

ii. $(x' \rightarrow y') \Rightarrow \neg(y' \rightarrow x')$

つまり, 関係は常に一方から他方へのものである. これは矢印の向きの一意性を与える.

この基本非対称性の条件の下で, 更に次のような制約が可能なプロファイル化の範囲を厳しく限定する.

¹³ただし, IDTM が PSR を満足している仕方にも些か問題がないわけではない. 例えば, この図で $BREAK$ の $X(t)$, $Y(t)$ に相当する部分がぼかしになっているが, これはうまく現時点ではうまく定式化されていない規約を仮定せずには対処できない.

(13) プロファイルの大小の相対的一意性

- i. 太さ($x \rightarrow x'$) > 太さ($y \rightarrow y'$) OR 太さ($y \rightarrow y'$) < 太さ($x \rightarrow x'$) [プロファイルの度合いは, $x \rightarrow x'$ か $y \rightarrow y'$ かの一方が必ず他方より大きい]
- ii. 太さ($x \rightarrow y'$) > 太さ($y \rightarrow x'$) OR 太さ($x \rightarrow y'$) < 太さ($y \rightarrow x'$) [プロファイルの度合いは, $x \rightarrow y'$ か $y \rightarrow x'$ かの一方が必ず他方より大きい]
- iii. 太さ(x) > 太さ(x') OR 太さ(x) < 太さ(x') [プロファイルの度合いは, y か y' かの一方が必ず他方より大きい]
- iv. 太さ(y) > 太さ(y') OR 太さ(y) < 太さ(y') [プロファイルの度合いは, y か y' かの一方が必ず他方より大きい]

これらは, もちろん制約であり, プロファイル化に一意性を与えるものではない. プロファイル化の一意性はあくまで事例ごとに, 状況的に定まる.

条件の具体的数値化としては次のようなものがあり, 実際, これはFig. 6にある図の解析表現となっている.

- (14) i. 太さ(x) = 3, 太さ(y) = 3, 太さ(x') = 1, 太さ(y') = 2;
- ii. 太さ($y \rightarrow y'$) = 2 [$\Rightarrow$ 太さ($x \rightarrow x'$) < 3],
太さ($x \rightarrow x'$) = 1 [$\Rightarrow$ 太さ($y \rightarrow y'$) > 1],
太さ($x \rightarrow y$) = 2 [$\Rightarrow$ 太さ($x' \rightarrow y'$) < 3],
- iii. 太さ($x' \rightarrow y'$) = 1 [$\Rightarrow$ 太さ($x \rightarrow y$) > 1],
太さ($x \rightarrow y'$) = 3 [$\Rightarrow$ 太さ($y \rightarrow x'$) < 3],
太さ($y \rightarrow x'$) = 0 [$\Rightarrow$ 太さ($x \rightarrow y'$) > 1]

読者の便宜のために, Fig. 6を再掲する. これに対しFig. 9にあるような図は, 単独の要素の意味記述としては IDTM では排除される.

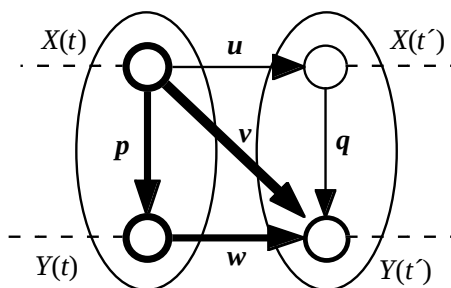


Fig. 6

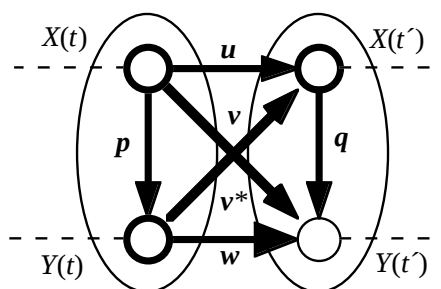


Fig. 9

Fig. 6では, p は q より, w は u よりプロファイルの度合いが大きい. また, v は $v^* = R(Y(t), X(t'))$ に対立しているが, v^* の太さがゼロであるため, それは明示的には現れていない. これに対し, Fig. 9では, $\langle p, q \rangle$, $\langle u, w \rangle$, $\langle v, v^* \rangle$ の対のおのおのについて, (13)にある相対的大小関係が満足されていない.

(13)にある制約のセットはまだ暫定的な性質のものだが, 仮にこのままでは成立しないとしても, このようにダイナミックな相互作用のネットワークに基づく記述は, それ自体が十分に抽象的で, 精密な意味記述に必要な解析力を十分に備えていると期待される.

IDTM の図の使用法は, ラネカー流のそれに比べると, 非常に制限されたものである. 何がどう書けるのか, はじめから厳しく制限されている. 実際, そのような制限なしに, ある記述的な枠組みに実践的価値があるとはとうてい主張できない.

すでに節3.6で明示したように, IDTM は図法に特殊な記号は極力用いない. 必要な区別は, 構造的, 関係的に表現されるように意図されている. 要素の性質の区別のために, 丸を三角や四角に代えたり, 直線を波線やジグザク線に代えたり, そういう場当たりの区別を導入することは, 意図的に避けている. そのような拡張には際限がないということがあるが, それ以上に, それは IDTM が本当の意味で精度の高い記述を狙っているからである. 許されている基本操作の集合が閉じていない, つまり「制約されていない」枠組みでの記述は, 万能の工具箱が与えられているのと同じで, 実質があるとは思われない.

4.2 具格要素, 結果述語の IDTM での扱い

(10)a, b ですでに見た **break** の自他形の交替に加えて, その結果述語, 具格性の問題を IDTM に基づいて記述する. 次の例を見よ.

- (15) a. $X\text{break } Y(\text{into } Z) [= (10)a + (\text{into } Z)]$
 b. $Y\text{break } (\text{into } Z)$
 c. $X\text{break } Y(\text{into } Z) \text{ with } W.$

c'. Xbreak Y with W into Z

d. Wbreak Y (into Z)

結果述語, 具格性の問題は, 次のFig. 10-Fig. 11で統一的に表される. $X(t), X(t')$ を X, X' と略記し, M, M' 外の ID 経路を示す破線は省略した. 以下, この規約に従う.

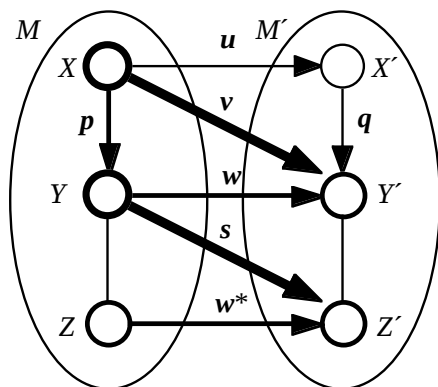


Fig. 10

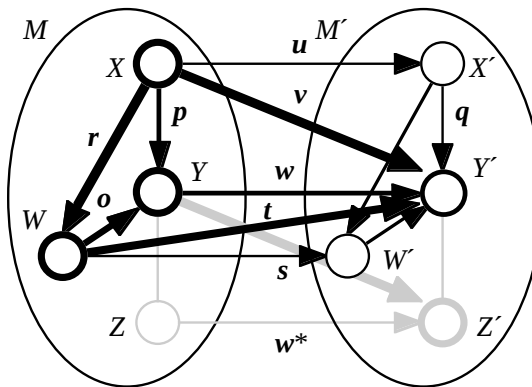


Fig. 11

(15)bにある自動詞形の *break* の特徴的意味成分は, Fig. 10-Fig. 11で w 成分 ($= R(Y, Y')$) がコードしている.

$R(Y, Z), R(Y', Z')$ に矢印がついていないのには, 意味がある. これは, 他の場合と関係のタイプが異なると思われるためである. この点は, §4.3.1で明確にする.

(15)a, cにある *break* の他動的成分は, v 成分 ($= R(X, Y')$) がコードしている. (15)a, b, c, d に随意的に現れる *into* の意味成分は, Fig. 10-Fig. 11で s 成分 ($= R(Y, Z')$) がコードしている. w^* は w の属性次元での対応物である. s はまた, 意味的には BECOME と等価な成分でもある. なお, Fig. 11で, s は随意性を示すため, *with W* との対比を明瞭にするために, ぼかしがかけられている.

(15)cに現れる *with* 特徴的意味成分は, Fig. 11で r 成分 ($= R(X, W)$) がコードしている.

(15)dにある道具主語をもつ *break* の特徴的意味成分は, Fig. 11の t 成分がコードしている. (15)dの特徴的意味を正確に記述する際には, u, v, p, q, r, X, X' には非実現を示すぼかしが入ることになる.

$\langle v, t \rangle, \langle p, o \rangle$ は正確な対応対であり, 相同的な関係にある. この現象の自然な記述は, $W = \tilde{X}$ とし, X, W の関係を属性関係の拡張にすることで達成できよう.

4.2.1 ラネカーの玉突きモデルでの分析との比較

ここで *break* の意味構造の記述について, IDTM の許すFig. 10-Fig. 11をラネカーの玉突きモデルの図法と比較してみよう.

玉突きモデルは(10)にある *break* の意味構造をFig. 12, Fig. 13にあるような図式で記述する

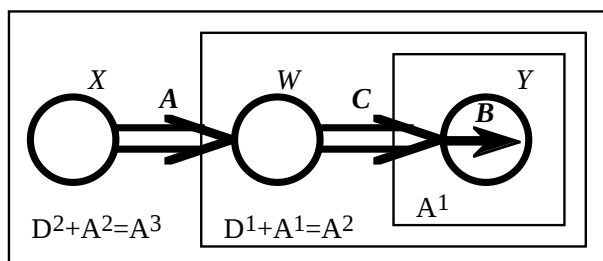


Fig. 12

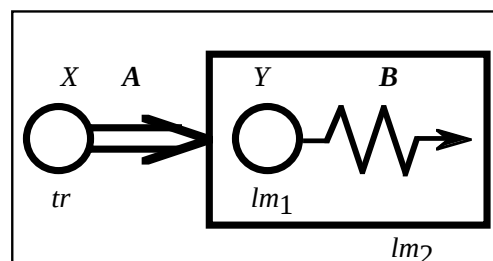


Fig. 13

Fig. 12は Langacker 2000: 85, Figure 3.5 に比較のため X, Y, W, A, B, C のラベルを付加したものの, Fig. 13は, Langacker 1987: 410, Figure 9.6 に X, Y, A, B のラベルを付加したものである.

Fig. 13は, Fig. 12の X, Y の関係を示すものらしいが, 二つの図法の関係の解釈は, 必ずしも明らかではない. ラネカー流の図には, 例えばFig. 13で, エネルギー伝達の矢印 A が外枠 lm_2 に接するか,あるいは, lm_1 に接するかなどに関して, 幾つかの奇妙な変異体があるが, そのような変異がどう正当化されているか,あるいはいないかは, ここでは論じない.

ラネカーの主張によれば, Fig. 12は次の容認性の分布を正しく記述するものである.

- (16) a. *The glass broke.* [= A^1]
- b. *A hammer broke the glass.* [= A^2]
- c. *Floyd broke the glass.* [= A^3]
- d. **Floyd {broke/caused}.* [*Floyd* = agent]
- e. **The hammer {broke/caused}.* [*hammer* = agent]
- f. **Floyd {broke/caused} the hammer.* [*hammer* = instrument]

ここにあるのは 2000: 85 の例文 (12) の引用である.

さて, Fig. 12をFig. 11に対比すると, 次のことが解る.

- (17) i. Fig. 12の X = Fig. 11の X ; Fig. 12の Y = Fig. 11の Y ; Fig. 12の W = Fig. 11の W ;
- ii. Fig. 12の A = Fig. 11の r , Fig. 12の B = Fig. 11の w , Fig. 12の C = Fig. 11の t
- iii. Fig. 11の X', Y', W', u, v, q, o に相当するものは, Fig. 12には存在しない
- iv. Fig. 12の tr, lm の補助記号は, Fig. 11ではプロファイルの相対的大小関係でコードされており, それ自体では大した意味がない
- v. Fig. 12の B は(内部の)状態変化であり, A, C のエネルギー,あるいは力と図形的

に区別されているが,そのような区別はFig. 11では図形的ではなく,関係的,構造的に符号化されている

- vi. Fig. 12の四角の枠入れ (lm_2 に相当) は, Fig. 11では相互作用の平面の区別として間接的に符号化されている

(17)iii は特に重要である. これは, $F = X \text{ break } Y \text{ with } W$ の構文で, X と Y のあいだの直接の相互作用がある場合, 玉突きモデルはそれを原理的に表現しえないことを明確にしている. これが自然な制約なのか, それとも図法の限界からくる恣意的な制約なのは議論の余地があるが, これまでの分析は後者の可能性を強く示唆する.

結論として言えることは, ラネカー流の図法では IDTM では区別されうる, そして, 正確な記述のためには区別されるべき意味成分が区別されていない. 単純に言うと, ラネカー流の図法には記述力が足りないということである. このような事実は, IDTM の提供する図法がラネカー流の図法に対し上位互換であることを強く示唆する.

4.2.2 玉突きモデルの分析上のバイアス

玉突きモデルは, 具格名詞句の特徴づけに関して, 一つ非常に信憑性に問題のある主張をなす. それによれば, X から W への, W から Y への働きかけには内在的順序がある, ということである. これはエネルギーの伝達の観点からは必然的なものであるが, 非常に信憑性に欠ける主張である. 実際, 玉突きモデルに基づく具格性の分析は, 動作主と道具とのあいだに, ありもしない内在順序を押しつけている可能性がある. Fig. 11の v 成分の存在がそれを示唆し, また, *with* に相当する語が様々な言語で道具性と随伴性の意味に関して曖昧であるという事実にもそれが強く示唆されている¹⁴.

玉突きモデルの本質的な限界は, 内在的に同時平行な事象をうまく記述できないという点にある. それは, しばしばありもしない順序を内在的なものとして, 意味構造に押しつける. 実際, これは, §4.5で $X \text{ give } Y \text{ to } Z$ と $X \text{ give } Z \text{ Y}$ の意味記述を提案する際に明確になることである.

4.3 ID は属性の束である: MAKE の場合

これまでの記述に加え, 次のような作成述語としての *make* の特徴的意味を区別するためには, ID の概念を拡張する必要がある.

¹⁴日本語では, 道具性と随伴性は格助詞「で」と「と」によって区別されるが, これはここで言及した傾向に対しては例外となる.

- (18) a. $X \text{ make } Y$. e.g., *Andy made money*.
 b. $X \text{ make } YZ$. e.g., *Bill made her sad*.
 c. $X \text{ make } YVP$. e.g., *He made his girlfriend cry*.

必要な拡張とは、次の仮説の導入である。

(19) ID 階層仮説

- a. ID は属性の束 (attribute bundle) であり, 内部構造をもつ。
 b. その属性はブール関数 $\{\text{TRUE}, \text{FALSE}\} = \{1, 0\}$ として表現しうる¹⁵

実際, 次の Fig. 14 が示すように, (18)の区別が示している *make* の多義は ID の内部構造に言及せずには正しく記述できない。

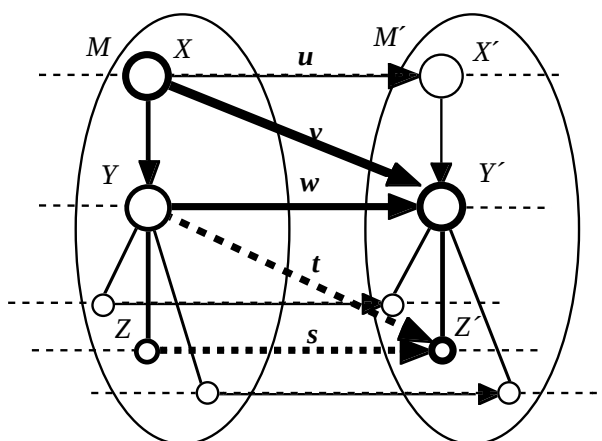


Fig. 14

この図で, 共通に $v = \text{make}$ である. Y の ID は属性の束だと考えられている. t は BECOME や $\text{break } Y \text{ into } Z$ の $Y \text{ into } Z$ に相当する意味成分である. w は GO, MOVE に相当する成分である.

(18)a の型については Y' (e.g., *money*) がプロファイルされる. (18)b の型については $Y(t')$ (e.g., *sad*) がプロファイルされるが, 起こっている変化は ID 内部の属性 $Z = \text{SAD}(Y)$ の変化 s である. Z はブール関数で, 時点 t の値は FALSE で, 時点 t' の値は TRUE である. (18)c の型については, 二つの成分にプロファイルがある. まず, $Y(t)$ (e.g., *his girlfriend*) に, 次いで w (e.g., *cry*) にである. つまり, 始点と軌道がプロファイルされている. (18)c に関しては, s (e.g., $Z = \text{SAD}(Y)$ ただし, 時点 t' で値が FALSE) が読み込まれる可能性もある.

¹⁵ここでは属性を二値関数だと考えているが, $\{-1, 0, +1\}$ の三値を取る関数を考えた方が適切に問題を扱える場合もある. ここでは簡素化のため, 二値とした.

4.3.1 ID 階層

$Z/X \text{ make } Y$ の性質を詳しく調べると, Z は次の性質を満足することが解る.

- (20) a. 一般に Y の属性を $\tilde{Y}$ で表すと, $Z = \tilde{Y}$
- b. 関係 $(Y, \tilde{Y})$ は, 部分と全体の関係であり, ID 同士の基本的非対称性と同じタイプの非対称性をもつかどうか, 現時点で定かではない. それゆえ, 関係に矢印は与えていない
- c. Y と $\tilde{Y}$ との関係は, 次のような形で階層をなしている

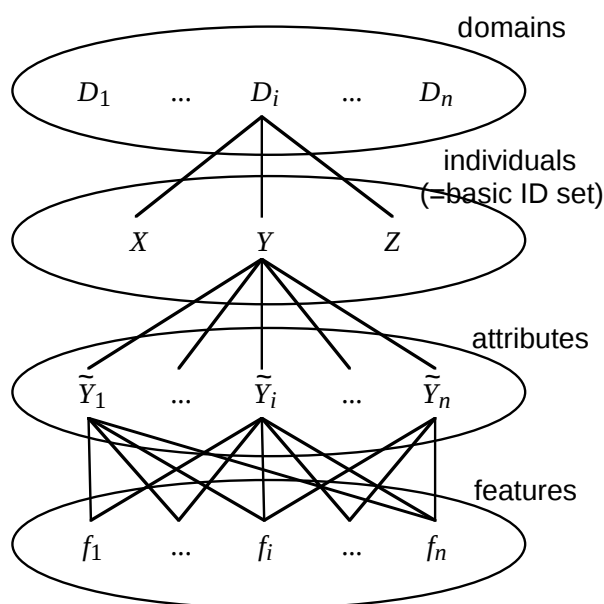


Fig. 15

つまり, Y は <属性, 値> の対の集合 $\{<\tilde{Y}_1, v_1>, <\tilde{Y}_2, v_2>, \dots, <\tilde{Y}_n, v_n>\}$ である. これが ID が属性の束であるという表現の正確な意味である.

ID 階層は時間軸上に展開され, 属性地図=属性幾何 (attribute geometry) と呼ぶ構造を形成する. 次の図は属性地図の概念を明確にするためのものである.

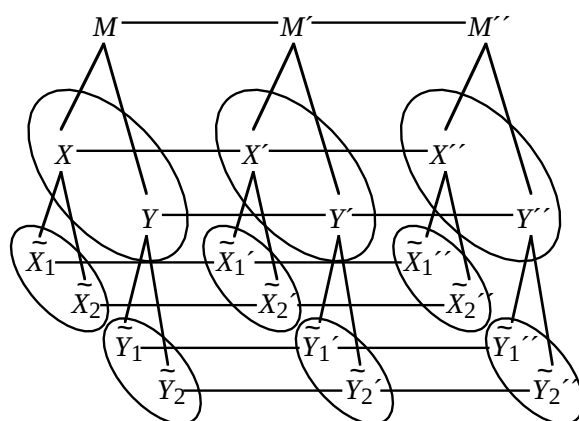


Fig. 16

この ID 階層の概念が生成音韻論で用いられる素性地図=素性幾何 (feature geometry) の概念と酷似しているという事実は、補足的に指摘しておきたい^{16, 17}.

4.3.2 属性地図と参照点構造との関連

素性地図の概念との類似がどのような意味をもつものであれ、ID 階層仮説の導入により、IDTM は非常に記述力を増す。それは実際、結果としてラネカーの参照点構造 (reference-point structure) の理論 (Langacker 1993) に新しい解釈を与えるほどの記述力をもつ。

実際、ID=Xを参照点、その属性集合 $\{\tilde{X}_1, \tilde{X}_2, \dots\}$ への投射を X の支配域 **dominion** だと解釈すれば、属性地図内への到達は参照点構造と同一のものと考えられる。ただ、この際、ラネカーの定式化に色濃く出ている視覚メタファーの性質は消失し¹⁸、到達可能性 **accessibility** のみが抽象的に表現されることになる。

このような類似は理論的に大変に興味深い問題であるが、紙面の関係から、この点に関しては、本論文はこれ以上論じない。

4.3.3 IDTM に容器メタファーは成立しない

Fig. 14の別の表記法として、次のFig. 17が考えられる。

¹⁶現時点では、属性地図と素性地図との関係は単なる類比以上のものではない。素性地図に関しては、Clements and Hume 1995 や、そこにあげられている文献を参照されたい。

¹⁷補足的に言うと、このような解釈の下では、ドメインやスペースというのは、プログラミング言語の記述で用いられる名称空間 **name space** の概念と等価なものであるとも想像される。

¹⁸実際、ラネカーの見地では、支配域は探索領域 **search domain** の概念の拡張だと見なされている。ID 階層の理論で支配域は内部構造と等価であり、この種の視覚メタファーは成立しない。実際、この理由で Viewer

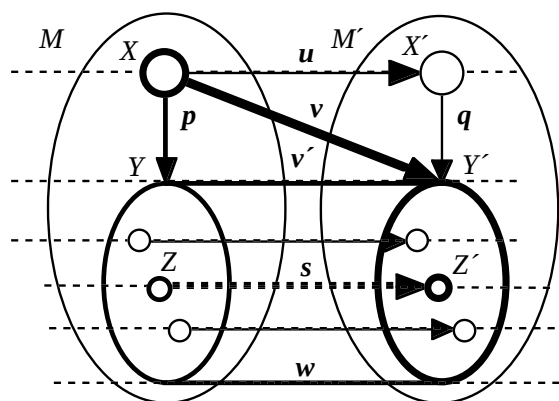


Fig. 17

これが意味しているのは, Y, Y' はそれ自体がスペースやドメインと等価なものとして表現されるということである. これは実際, 属性地図の概念が許す正当な解釈である. ただし, この表現法の意義は図の簡略化のためのものであり, 正確な解釈はあくまでも Fig. 14 によって与えられる.

特に注意して欲しいのは, Fig. 17 のような図で, Y の境界性に「容器メタファー」(Lakoff 1987) が読みこまれてはならないという点である. このような容器性は, 明らかに図を主観的に読みこむ際に発生するイメージの干渉の結果であり, 結果として, 図の解釈に, その本来の抽象的な性質を損なうような概念的混乱をもたらす. これは, 図に, 実際にはありもしない性質を付与する効果に繋がり, 緻密な記述の目的には明らかに危険である.

4.4 IDTM の記述の言語非依存性: 日本語動詞の特徴的意味記述

様々な研究者の開発の努力にも関わらず, ラネカーの玉突きモデルは, 英語以外の言語の特徴的意味記述にはうまく妥当しない. 実際, それは日本語の動詞の意味記述にはうまく妥当しない. その最大の理由は, 動詞以外の関係的要素, 例えば格助詞の役割が, 適切に記述されないからである¹⁹.

これに対し, IDTM は日本語の構文現象も自然に表現する²⁰. ここでは, 例示のために, 次の (21) にある ‘壊す’ と ‘壊れる’ の自他形式の交替を考察する.

(21) a. X が (Z で) Y を壊す

は不要となる.

¹⁹これは, 英語の分析で文中の前置詞のプロファイルが不明なままであるという問題と本当はまったく同じ根をもつ問題であるが, 英語では項の位置が役割と強い相関をもっているので, 幸い深刻な問題として認識されていないにすぎないと思われる.

²⁰実際, IDTM が開発された最大の動機の一つは, 日本語の構文の特徴的意味を適切に記述したいという筆者の強い欲求である.

- b. Y が(Z で)壊れる (ただし, Z は原因/道具名詞句で, 場所名詞句ではない)

この自他交替は, 次の図式によって特徴づけることが可能である.

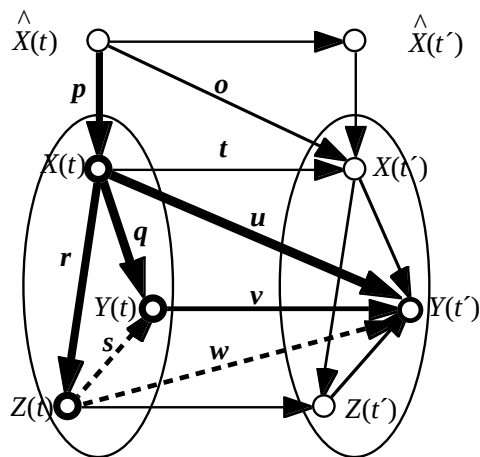


Fig. 18

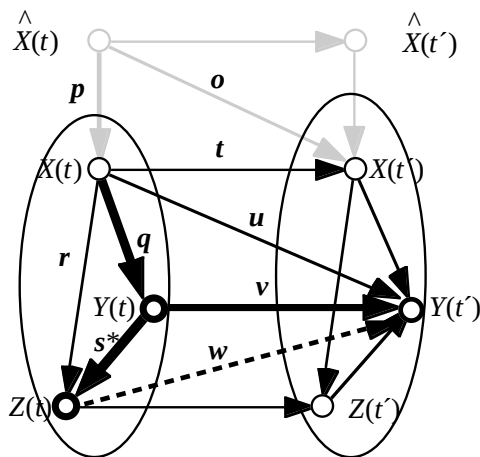


Fig. 19

Fig. 18 と Fig. 19 とともに格助詞は p, q, r, s で, 動的成分は t, u, v, w である. Fig. 18では p = ‘-が’, q = ‘-を’, r = ‘-で’, u = ‘壊す’, v = ‘壊れる’ である²¹. Fig. 19では q = ‘-が’, s = ‘-で’, v = ‘壊れる’ である²². ここでは, q が ‘-が’ と ‘-を’ の交替形式をもつことが特に興味深い.

主語を表す格助詞 ‘-が’ の役割 (p, q がコードする)を除けば, 英語の $X break Y$ の特徴的意味構造と日本語の ‘壊す’ の特徴的意味構造, ならびに $Y break$ の特徴的意味構造と日本語の ‘壊れる’ の特徴的意味構造とは, 驚くほど似ている. 両者の差は, 日本語では主語のマーキングのために p, q が現れ, $\hat{X}, \hat{Y}$ が介在することだけである. このような記述の収斂は, IDTM が十分に言語中立的な記述力をもった枠組みであることを強く示唆する.

4.4.1 超 ID 束の概念

$\hat{X}$ は特殊な要素であり, 正当化のために詳細な吟味を要するが, 詳しい説明は割愛する. 簡単に言うと, それは X の超 ID 束 (meta ID bundle) で, X がその値になるような超個体レベルでの変項である. それは, メンタル・スペース理論風に言うならば, X を値とする役割とも解釈できる.

²¹これには二つの解釈がある. (I) v = ‘こわ-’ として, 図18の u = ‘-す’, 図19の u = ‘-れる’ とするのと, (II) u = ‘こわ-’, o = ‘-す’, v = ‘-れる’ とするのとである. ただし, いずれの解釈もここでの分析の内容には影響しない.

²²格助詞 ‘-が’ の特徴づけとしては, それを $t = R(Y, Y')$ とする案もある. しかし, 完全(非能格)自動詞 (e.g., ‘泣く’) の主語をマークする ‘-が’ の記述に困難を生じるため, ここにある図法に与えた特徴づけを採用した. だが, これは日本語の自動詞が内在的に再帰性を有する他動詞の一種として定義されれば本質的な困難とはならないので, 今後, この方向に修正される可能性はある.

興味深いのは、超 ID 束と ID との $\langle \hat{X}, X \rangle$ の関係は、ID と属性との関係 $\langle X, \tilde{X} \rangle$ と相同 homologous だという点である。これは重要な意味をもつ事実であり、例えば日本語の二重主語構文 (e.g., “ゾウは(その)鼻が長い”, “魚は(そのうち({の, で}))鯛がいい”) の分析や、ウナギ文 (e.g., “ボク({が頼んだ, が飼っている, ...})のはウナギだ”) の分析に重要な役割を演じるが、ここでは追及しない。§4.3.2 で言及した属性地図と参照点構造との類似がその理由になっていることのみを指摘しておく。

4.4.2 IDTM は語彙分類のために、新たな認知的な動機づけを提供する

IDTM は、関係的な要素であれば、それが構文中で動詞であるか、形態素であるかは区別しない。実際、動的な成分は、動詞の場合もあるし、前置詞、後置詞の場合もある。具体例としては、Fig. 14において、 t は Y become Z と Y into Z ととも実現可能な意味成分である。

IDTM が提供するのとは、そのような「純言語的」な区別を捨象した抽象的なレベルで成立している心内現象の記述のための道具であり、実際、それに基づいて新たに品詞とは別のレベルで成立する語彙クラス分類も可能となる。実際、この性質が IDTM に望ましい言語非依存性を与えていると考えられる。

4.5 同時並行事象の記述: GIVE の場合

IDTM の最大の利点の一つは、次の *give* が示す同時並行事象 (parallel events) の概念構造を正しく記述することにある。

- (22) a. X give Y to Z e.g., *She gave her money to her son.*
 b. X give Z Y e.g., *She gave him the information.*

様々な研究者の努力 (e.g., 中村 2001, Newman 1996) にも関わらず、 x GIVE y TO z のタイプの構文がいわゆる玉突きモデルではうまく扱えない。

まず、玉突きモデルでは、(22)a, b の二つの構文のちがいをうまく区別できない。中村 2001 は両者の違いはプロファイルの違いであると分析しているが、彼のプロファイルの定義はラネカーが (Langacker 1987) で与えたもの、あるいは、それが本来もつべき表層形との対応関係を損なっている。また、これとは別に矢印の向きを逆にすることはできるだろう。だが、それはエネルギーの流れが反対になっていることを意味する。

仮に、このような問題が解決できたとしても、最後に y TO z の特徴づけが残っている。中村 (2001) を含めて、典型的なラネカー流の分析では、 x GIVE y と y TO z のプロファイルが区別されない。これが原因で、*to* が付加語を導入する x LEAD y TO z と x GIVE y TO z との有意味な区

別は, ラネカー流の図法では不可能であるように思われる.

もっと重要なことに, 玉突きモデルはいずれのタイプでも, 二つの異なる動作 y の (抽象的) 地点 z への「位置の移動」と, y の x から z への「帰属の変化」とが, 同時に進行する並行事象だという点を, 首尾一貫した形で記述しているとはいいがたい. 所有の変化と, 位置の変化の間には時間的順序はない. だから, それらの関係は玉突きモデルでは, うまく特徴づけられない.

中村 2001 は確かに, この平行性を捉えるように図法を拡張しているが, それは対症療法的であり, 認知文法全体の枠組みの中で一貫した扱いなのか, 些か疑問である. この点に関しては, §4.5.3 でもう一度詳しく見る.

これに対し, IDTM は GIVE の意味構造を, 並行的に進行する三つの局面 α, β, γ の複合として記述することを可能にする.

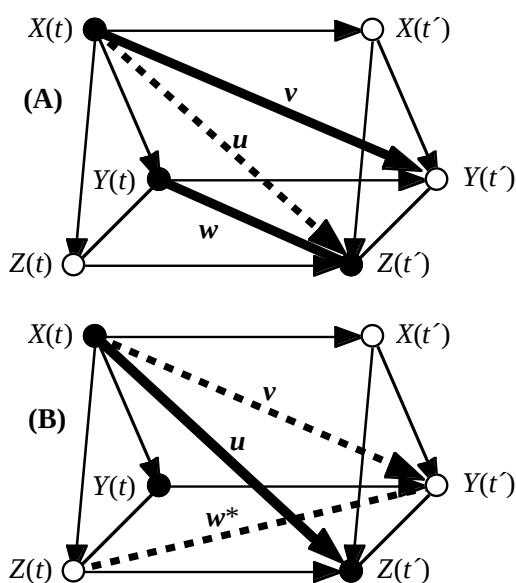


Fig. 20

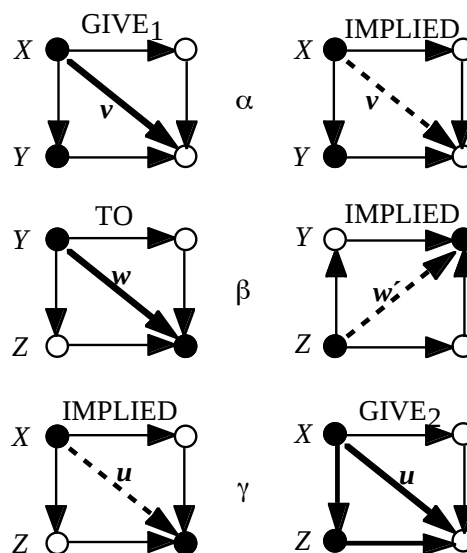


Fig. 21

Fig. 20-A は $F = x \text{ GIVE } y \text{ TO } z (\text{GIVE}_1)$ を, Fig. 20-B は $F' = x \text{ GIVE } z y (\text{GIVE}_2)$ の意味構造を特徴づける. 破線は非明示的な関係を示す. Fig. 21-A, B は二つの構文の比較のために, 局面 α, β, γ に何が現れるかを示す.

局面 α は $F (= \text{GIVE}_1)$ で X から Y への働きかけ (この場合は譲渡) をコードするが, $F' (= \text{GIVE}_2)$ ではコードは暗黙的である (これは破線で示してある). 局面 β は $\text{GIVE}_1, \text{GIVE}_2$ で意味が異なる. GIVE_1 では, 局面 β は Y から Z への移動をコードするのに対し, GIVE_2 では, 局面 β は Z の Y への働きかけ (この場合は受容) をコードする. 局面 γ は, X から Z への対人的相互作用をコードする. これは, GIVE_1 では間接的, GIVE_2 では直接的にコードされている.

重要なのは、与格構文で暗示されていた影響関係の成分 u が二重目的構文では明示されていること、 β に現れる働きかけの方向 w, w^* が逆になっていることである。Fig. 21-A- β の構造は $Y\{\text{MOVE, GO, ...}\}$ to Z の構造と共通で、Fig. 21-B- β の構造は $Z\{\text{HAVE, TAKE, ...}\}$ Y の構造と共通である。

4.5.1 IDTM での構文「効果」の説明

これが, Goldberg 1995 流の与格および二重目的語構文の「効果」の説明になるなら, 構文というのは語彙の組みあわせから派生的に生じる現象として再定義できる可能性が示唆される。ほとんどの場合, 派生的, いわゆる「構文的」な意味は, 語彙的ベクトルの合成として表現されている。これは, 構文的な意味が始めから心内辞書に「登録」されているようなものではなくて, 語彙的意味の相互作用の産物だという主張とうまく折りあう。これが正しい主張であるならば, 構文というのは「創発」的 emergent なもので, 単なる用法の目録化やネットワーク関係の記述によっては, その全貌を捉えられないことになる²³。

実際, 構文「効果」が確認できるからといって, それから構文が実際に存在することにはならない。実際, 英語においてですら, Goldberg の研究が構文効果の目録化以上のことをしたかどうかは, かなり怪しいと思われる。

4.5.2 ラネカー1987 の分析との比較

ここで IDTM で提案された分析をラネカーの分析と比較する。

Langacker 1987: 327 では, 次の Fig. 22 が $F = X \text{ give } Y \text{ to } Z$ の意味構造を記述すると主張されている²⁴。ここでは, X : AG [= Agent], Y : TH [= Theme], Z : EXPER [= Experiencer] と解釈する。

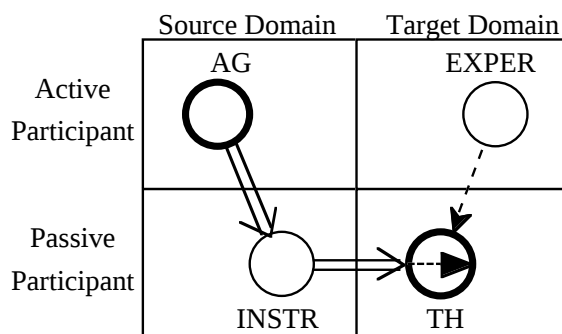


Fig. 22: Langacker 1991: 427, Fig. 7.5

²³この創発主義の主張は, Kuroda (1997) に詳しく展開されている。ここで言う創発とは、複雑系 complex systems の科学の鍵概念である。これに関する簡単な情報は Waldrop (1992) で得られる。Prigogine and Stengers (1984) もう少し専門的な内容も扱っている。

²⁴プロフィールはこのままでは不適切だと思うが、この問題は今は深く追求しない。

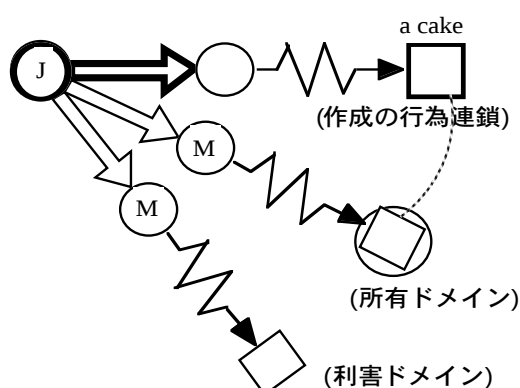


Fig. 25

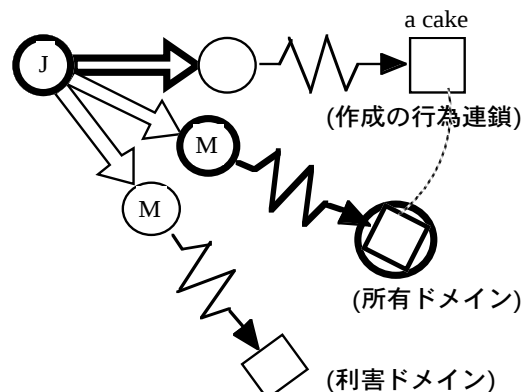


Fig. 26

Fig. 25, Fig. 26は, 中村 2001: 68 (4a), (4b) の図を少しばかり修正したものである.

IDTM の分析と中村 (2001) の認知文法の枠組みでの分析には, 幾つかの興味深い共通点もあるが, 同時に際立った違いもある. 以下ではそれを検討するが, その前に, 中村の図法に内在する幾つかの問題を指摘しておく.

まず, *send* と *bake* との意味構造の違いは作成 (= D1), 所有 (= D2), 利害 (= D3) の三つのドメインの分化に帰せられているが, それが何に動機づけられているのが明確ではない. 特に (23)b に所有ドメインは本当に関与しないのだろうか?

次に, $X \text{ send } Y \text{ to } Z \sim X \text{ send } Z \text{ to } Y$ の分析では $Y (= a \text{ book})$ の到達状態は丸で表され, $X \text{ bake } Y \text{ for } Z \sim X \text{ bake } Z \text{ for } Y$ の分析では $Y (= a \text{ cake})$ の到達状態が四角で表されている. これは, 明らかに *bake* 作成の効果を *send* の移動の効果と区別するためにであるが, これはいかにもアドホックである. これは意味役割の数だけ形の区別が必要だということを示唆する. とすれば, 最終的に, いったいどれぐらいの数の区別が必要なのだろうか? これは明らかに優雅な図の利用法ではない.

更に, Fig. 25, Fig. 26では, 作成, 所有, 利害の三つのドメインが主語 *John* を共有しながら放射状に広がる構造として描かれている. なぜドメインの数は正確に三つなのか, その理由は与えられていない. また, この構造化が「正確に」何を表わしているのか, 些か茫漠としている. なぜ始めから三つのドメインの間の対応として表さないのだろうか? その理由はこれが参照点構造の拡張だからなのかも知れない. だが, それは明示的に言われていない. 実際, 主語を共有しない図は, ずっと IDTM が与える図に類似する.

要するに, 中村のモデルは, 何がどういう風に制約されているのか (つまり, 図で何が描けて, 何が描けないのか) 明確ではない. 従って, 中村の図法は (6) の PSR を満足しておらず, 読み手の読みこみに依存している分だけ, 恣意的である.

この理由の一つに, 中村ではプロフィールの解釈に関して整合性が欠けているという点がある. 中村 (2001: 67) は (23)a を記述するスキーマ (= Fig. 23) で *Mary* がプロフィールされている

ないとしているが、それは *Mary* という形で表層形に現れているという事実を矛盾するのでないだろうか？ 彼の想定するプロファイルの定義が円弧のプロファイルの場合と矛盾しないのか、判然としない(円弧では、円(周)は言語表現として実現しないことに注意されたい)。同様の問題は他の図にも当てはまる。プロファイルがある(あるいはない)ための基準が明確に与えられていないし、それが一貫しているとも思われない²⁵。

共通点としては、まず、中村の言う所有ドメイン、移動/作成ドメイン、利害ドメインが、IDTM で言う局面 α, β, γ に、おのおの正確に対応し、どの局面が脱プロファイルされているかに関して、二つの分析は見事に見解が一致している点である。二つ目の共通点は、中村の分析でも与格構文や二重目的語構文の「意味」が語彙的意味の延長として捉えられているという点である。この点に関しては、IDTM はより明示的に、構文「効果」があくまで語彙的な情報の統合の副作用として記述されている。

異なっている点としては、IDTM には中村の言う力動ドメインが明示的に現われない。それは中村が基盤として玉突きモデルが力と運動のメタファーに基づく動力的なモデルであるのに対し、IDTM にはそのような力動的な側面は個々の ID の属性の変化の中に捨象されているからである。議論の分けれるところであろうが、筆者はこれを好ましいことだと判断する。

このような明白な違いが幾つかあるにせよ、二つの分析のあいだにある共通点の幾つかは、大いに強調されてよい。特に、GIVE が関与する二つの構文の意味が異なるドメインに分散しているという論点は重要な二つの分析の重要な一致点である。確かに、中村の分析ではドメインの区別が天下りの的に与えられ、その指定に独立の動機づけが欠けているが、その欠落を IDTM が自然に補っていると解釈するのがもっとも適切であろう。

結論として言えるのは、中村 (2001) は全体として正しい洞察を得ているが、それを図で適切に表現することに失敗している。そして、それは彼が依拠している視覚化の技法それ自体の内在的な困難から来ていると考えられる。IDTM は、そのような正しい洞察を正しく視覚化するための道具立てを提供する可能性が高い。

4.5.4 前置詞に内在する文法的「主語」

IDTM で理論的に予測されていることの一つに、前置詞が関係的であるならば、(i) それは二つ以上の項をもち、(ii) その一つは目的語であり、もう一つは主語である、ということがある。

実際、前置詞が固有の主語をもつという性質が成り立つので、中村 2001 が三つのドメインの並列性を主語を共有しながら放射状に展開するものとして記述しているのと実質的に同じことが、まったく独立の前提から出発している IDTM 分析でも同じく得られている。 $F = X \text{ give}$

²⁵中村の言うプロファイルの「有無」が、実際にはプロファイルの当たり具合の「相対差」だと好意的に解釈することはできるだろう。しかし、これは IDTM がそうしているように、プロファイルの当たり具合は全/無ではなく、段階性が認められていることが前提となる。

Y to Z の意味は、実際、 $G = X$ give Y という動詞成分と $H = Y$ to Z という前置詞成分とに分散的に表現されており、何らかの解釈域での Y の移動は H 成分によってコードされていると考えられるのが、もっともデータとの整合性が高い分析であろう。実際、 $F' = X$ lead Y to Z や $F'' = X$ move Y (on) to Z のような「複合的」パターンが F と区別されるのは、必ずしも自明ではない。実際、その差は、記述的には、 $F = X V Y$ to Z の Z の名詞句のタイプの変動に相関して、 $H = Y$ to Z の解釈領域が所有になるかならないかという、たったそれだけのちがいによるものである。

前置詞句 (e.g., *to Brazil*) が自動詞 (e.g., *come*) と結合し複合的に他動詞句 (e.g., *come to Brazil*) を構成しうるのは、動詞 X come と前置詞 X to *Brazil* とが主語 X を共有するからである。一般に、表層分布から見て前置詞の主語が存在しないように見えるのは、それが動詞の項 (e.g., 主語, 目的語) と共有されているためであろう²⁶。

4.5.5 IDTM は二種類の前置詞の区別を予測する

理路論的な要請として、IDTM は二種類の異なった前置詞があることを要請する。動的なものと静的なもの(あるいは単に非動的なもの)である。動的な前置詞の代表例は *to*, *into*, *from*, *across*, *over* であり、非動的な前置詞の代表例は *in*, *on*, *at*, *of* である。(22)a の *to* 与格で移動を表せるのは、*to* が動的な前置詞であることの一つの現われであろう。

このような区別は、はじめは奇異に見えるかも知れないが、§4.4.2で説明した Y into Z と Y become Z との交替を説明するものでもあり、IDTM では重要な説明的役割を演じる。

4.6 IDTM の現時点での問題点

IDTM の分析に問題がないわけではない。例えば、次のような場合が存在しないか、しないならば、それは何故かに関する自明な説明はない。

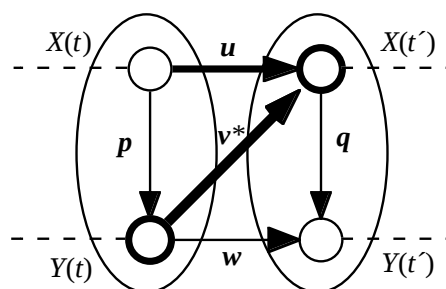


Fig. 27

これは、恣意性を減少させるために、矢印の向きがうまく制約される必要があることを示す事例だが、これをうまく制約する原則は現時点では見つかっていない。今後の研究から、このよう

²⁶この点は、Kuroda (2000) で詳細に論じられている。

なタイプの現象が存在するか,それが存在しないならば,それを制約する原理の存在が明らかになると期待される.

5 まとめ

IDTM は,正確には言語構造のモデルではなく,その背景にある**概念化**のモデルである. IDTM が提供する概念化の表示は玉突きモデルのそれより複雑であり,そのため,表層形との写像関係は些か複雑になる. 実際, IDTM が表示する内容は,より抽象的であり,玉突きモデルほど単純明快ではない.

だが, IDTM は,それによって玉突きモデルよりも詳細で,かつ恣意性の低い(つまり,読者の読み込みに依存する度合いの少ない)記述を提供しており(事前の適合性の選り分けなしに選ばれた)任意の言語表現の正しい記述を得るという目的のためには,複雑化によって失うものよりも得るものの方が大きいと考えられる. 特に, IDTM はその表示の抽象性ゆえに,結果として個別言語に特有のメタファー的バイアス(英語に構文おける力動モデルの支配力)に依存しない高次のレベルで意味構造の記述を提供する枠組みであると考えられる.

参考文献

- Chomsky, Noam (1965) *Aspects of the Theory of Syntax*. Cambridge, MA: MIT Press.
- _____ (1982) *Some Concepts and Consequences of the Theory of Government and Binding*. Cambridge, MA: MIT Press.
- _____ (1995) *The Minimalist Program*. Cambridge, MA: MIT Press.
- Croft, William (1991) *Syntactic Categories and Grammatical Relations: The Cognitive Organization of Information*. Chicago, IL: University of Chicago Press.
- Clements, G. N. and Elizabeth V. Hume (1995) The internal organization of speech sounds.
In: John Goldsmith (ed.) *The Handbook of Phonological Theory*, 245-306. Oxford, UK: Basil Blackwell.
- Deane, Paul D. (1992) *Grammar in Mind and Brain: Explorations in Cognitive Syntax*. Berlin/New York: Mouton de Gruyter.
- Fauconnier, Gilles R. (1985) *Mental Spaces: Aspects of Meaning Construction in Natural Language*. Cambridge, MA, MIT Press.
- _____ (1997) *Mappings in Thought and Language*. Cambridge: Cambridge University Press.
- Goldberg, Adele E. (1995) *Constructions: A Construction Grammar Approach to Argument*

- Structure*. Chicago/London: University of Chicago Press.
- Kosslyn, Stephen M. (1980) *Image and Mind*. Cambridge, MA: Harvard University Press.
- Kuroda, Kow (1997) Where do constructional meanings come from? 『言語科学論集』 3: 17-44. 京都大学基礎科学科, 京都.
- Kuroda, Kow (2000) *Foundations of Pattern Matching Analysis: A New Framework Proposed for a Cognitively Realistic Description of Natural Language Syntax*. Unpublished doctoral dissertation, Kyoto University, Japan. [http://clsl.hi.h.kyoto-u.ac.jp/~kkuroda/papers/kuroda2000/*.pdf]
- Kuroda, Kow (2001) Presenting the *Pattern Matching Analysis*, A framework proposed for the realistic description of natural language syntax. *Journal of English Linguistic Society* 17: 71-80. English Linguistic Society of Japan.
- 黒田 航 (2003a). 「認知形態論」 『認知音韻・形態・語彙論』 (シリーズ認知言語学第二巻): . 東京: 大修館.
- 黒田 航 (2003b). 「概念化の ID 追跡モデルの」 の提案. 第4回日本認知言語学会大会予稿集: 23-26. 日本認知言語学会.
- 黒宮 公彦 (2003). 概念化の ID 追跡モデルの妥当性. 第4回日本認知言語学会大会予稿集: 27-30. 日本認知言語学会.
- Newman, John (1996) *Give: A Cognitive Linguistic Study*. Berlin and New York: Mouton de Gruyter.
- Lakoff, George (1987) *Women, Fire, and Dangerous Things: What Categories Reveal about the Mind*. Chicago, IL: University of Chicago Press.
- Langacker, Ronald W. (1987) *Foundations of Cognitive Grammar, Vol. 1: Theoretical Prerequisites*. Stanford, CA: Stanford University Press.
- _____ (1988) A usage-based model. In: B. Rudzka-Östyn (ed.) *Topics in Cognitive Linguistics*, 127-161. Amsterdam/Philadelphia: John Benjamins.
- _____ (1991a) *Concept, Image, and Symbol: The Cognitive Basis of Grammar*. Mouton de Gruyter.
- _____ (1991b) *Foundations of Cognitive Grammar, Vol. 2: Descriptive Application*. Stanford, CA: Stanford University Press.
- _____ (1993) Reference-point constructions. *Cognitive Linguistics* 4(1): 1-38. [Reprinted in Langacker (2000)]
- _____ (1997) Constituency, dependency, and conceptual grouping. *Cognitive Linguistics*, 8(1), 1-32. [Reprinted in Langacker (2000) in revised form].
- _____ (2000) *Grammar and Conceptualization*. Berlin/New York: Mouton de Gruyter.

- McCawley, James D. (1971) Prelexical syntax. In: Peter A. M. Seuren (ed.) *Semantic Syntax*, 29-42. London, Oxford University Press.
- 中村 芳久 (2001) 「二重目的語構文の認知構造: 構文内ネットワークと構文間ネットワークの症例」 山梨正明 (編) 『認知言語学論考』 (第一巻): 59-110. 東京: ひつじ書房.
- Prigogine, Iriya and Isabelle Stengers (1984) *Order out of Chaos: Man's New Dialogue with Nature*. New York, Bantam Books. [I.プリゴジン・I.スタンジェール. 伏見康治ほか(訳) 『渾沌からの秩序』 東京: みすず書房]
- Pullum, Geoffrey K. (1996) Nostalgic views from Building 20. *Journal of Linguistics* 32: 137-147.
- Pustejovsky, James. (1995) *The Generative Lexicon*. Cambridge, MA, MIT Press.
- 定延利幸 (1993) 「事態認知モデル構成要素としての状態の必要性」 『日本認知科学会第10論文集』: 76-77. 日本認知科学会.
- Sadanobu, Toshiyuki (1995) "Two types of event models: billiard-ball model and mold-growth model" 『国際文化学研究』 第四号: 57-110. 神戸: 神戸大学国際文化学部.
- Shepard, Roger N. (1978) The mental image. *American Psychologist* 33: 125-137.
- Waldrop, M. Mitchel (1992) *Complexity: The Emerging Science at the Edge of Order and Chaos*. Penguin Books. [邦訳: M・ミッチェル・ワールドロップ [田中三彦・遠山峻征 訳] 『複雑系—科学革命の震源地・サンタフェ研究所の天才たち』 東京: 新潮文庫]