

階層化された意味フレームのネットワーク分析 (HFNA) は 単なるネットワーク分析ではない

FOCAL 流の言語分析/概念分析の理論と実践の背景説明

黒田 航*

中本 敬子†

Created: 2005/MM/DD; Modified: 2008/MM/DD, 2015/12/13

1 はじめに

このノートは主に認知言語学者 (の卵) 向けに、私たちが他のメンバーと一緒に実践している FOCAL 流 [1, 2, 3] の言語分析/概念分析の背景説明を目的とする。

2 第二期認知言語学のために

2.1 (日本の) 認知言語学の現状

今の認知言語学 — 特に日本の認知言語学 — では、研究全体が完全に言語学という「殻」の中に閉じこもって、その中で自己完結しているように思われる。実際、現在の認知言語学は、コーパス基盤のデータ主導型の言語研究との接点も少なく、また、非常に自由度の高い — つまり、何をしてもよい — 高次認知レベルで過度の一般化に耽り、コネクショニストモデル (connectionist models) のような低次認知レベルでのモデル化との整合性も気につけない状態にある¹⁾。そればかりか、今の認知言語学には認知心理学、認知科学との交流もほとんどなく、(計算的) 認知科学 ((Computational) Cognitive Science)、計算言語学 (Computational Linguistics) の成果を軽視、蔑視する傾向すら認められる。

その「殻」を破るためのキッカケは、次のよう

* (独) 情報通信研究機構 けいはんな情報通信融合研究センター

† 京都大学 教育学研究科

¹⁾ 私たちは例えば、Langacker [4, pp. 525–536] がコネクショニズムとの「互換性」を主張しても、それが用語上の互換性以上のものだとは思えないし、Lakoff と Johnson [5, pp. 569–583] が Neural Theory of Language をもち出しても、それがコネクショニズム全体の矮小化以上のものだとは思えない。

な態度であると私たちは考える：

- (1) コーパス言語学との接点を重視し、用法基盤 (usage-based) の言語分析を徹底する
- (2) コネクショニズム (connectionism) のような低次認知のモデル化との互換性を意識する²⁾
- (3) 言語現象の中途半端な「説明」ではなく、個々の言語表現の理解内容の詳細で正確な「記述」を通じて (言語の) 認知心理学、認知科学との接点を取り戻す

これらの意味することは軽微ではない。(1) は「認知能力」を想定した言語現象の「説明」を少なくとももいったん棚上げにする必要があることを意味するし、(2) は例えば (意味) 素性 (features) を言語分析に妥当な範囲で積極的に導入する必要性を肯定することを意味する³⁾。更に (3) は、Trajector/ Landmark や参照点 (reference point) のような認知言語学内部でしか通用しない「理論仮構物」群を、少なくとも心理実験や網羅的なコーパス分析によって検証する必要性を肯定することを意味する。

²⁾ 神経系との接点に重点を置きすぎることは現時点で好ましくないことは、私たちも認める。モデルや理論の価値を神経系との直接的な接点のみに求めるのは、様々な理由から好ましいことではない。神経系の活動を fMRI や NIRS を使って脳画像や何かで測ることができたとして、それを解釈するためのモデルが必要である。言語学、認知心理学が認知科学の脳神経科学的側面に貢献できる何かがあるとすれば、それは、中間的な解釈モデルとして現象理解の役に立つという点にあるはずだ。高次認知機能の低次認知機能への橋渡しは必要不可欠だが、そのための具体的方策に関して、私たちは特に強い主張をしているわけではなくて、単に「高次認知レベルでの一般化を自明のものとししないで、もっと低次認知のことも真面目に考えよう」と訴えるのみである。

³⁾ 素性を使わずにコネクショニストのシミュレーションを行なうのは、ほとんど無意味である。

2.2 認知言語学の「基盤」を再検討する必要性

認知言語学は、誕生当時、生成言語学のような他の流派の言語学と異なり、学際的であることが強調された。だが、そのような意気込みは、時の流れ、世代の交替と共に徐々に形骸化し、今となっては見えない影もない。

認知言語学が認知心理学から取り入れた最新の知見は、80年代のものである。[6]が書かれた後で認知心理学で研究成果として得られた知見の多くは、現在の認知言語学では取り入れられていない。大多数の認知言語学者は、認知心理学者がプロトタイプという説明概念を以前ほど真に受けていないという事情も知らない。認知言語学の認知心理学との交流は、事実上、80年代で終わっている⁴⁾。

だからこそ、認知言語学の学際性を本当に望むのであれば、次のことを理解するのは急務である：

- (4) 言語学者が言語データを直観に基づいて分析した結果のみから得られる一般化で、認知心理学的、認知科学的に無条件に妥当だと見なせるものは多くはない — 実際には皆無に等しい。

これが意味することは次のことである：

- (5) 認知言語学派の言語分析の「常套手段」、すなわち認知言語学の研究者コミュニティ内部で、従来ならば使うのが当然だと見なされてきた分析手法 (e.g., ネットワーク分析)、想定するのが当然だと見なされた説明概念 (e.g., Trajectory/Landmark, 参照点, メタファーリンク, メトニミーリンク, コネクターのような各種のリンク, 比喩写像, スキーマ, 各種の構文 (constructions) etc.) に対し、全面的な再検討が必要である。

すでに述べた理由から、この再検討の必要性は明白であるが、この論文では必要な再検討のすべ

⁴⁾ 認知心理学の知見を更新したいと思う人には、[7]のような入門書をお薦めする。

てを扱うことはできない。私たちが取り組むのは、認知言語学の「常套手段」の一つであるネットワーク分析やイメージスキーマによる解析の再検討である。

2.3 部外者から見た認知言語学の評価に学ぶ

§3での具体的な議論に先立って、ネットワーク分析やイメージスキーマによる意味の記述 (あるいは説明) が正確に何をやっていることなのかを読者に自覚してもらうことが必要だろう。このためには、認知言語学の「成果」が好意的な部外者にどう映っているかを客観的に紹介するのが効果的だと判断する。

この論文の著者のうちの一人 — 仮に *P* 氏としよう — は、認知言語学になじみがない認知心理学者で、部外者として枠組みを公平に評価できる立場にある。*P* 氏は、認知言語学関係の研究会に出ていてる際にしばしば、分析が何を対象としていて、何を仮定していて、分析結果が何の記述/説明になっているのか非常に不明確な印象を受けると言う。

ネットワークモデルやイメージスキーマによる記述/説明には、少なくとも下記の二通りの解釈がありえる。

- (6) a. 言語使用者が心の中に抱いている、刻々と変化する心的表示のモデル
- b. 意味の拡張や変化が起こるときの (マクロな) プロセス

この二つは — お互いに関連がないわけではないが — 決して同一ではない。*P* 氏は、認知言語学の文献を読んだり研究発表を聞いても、いったい分析者がどちらを意図しているのか分からないと訴える。

2.4 文脈効果を真剣に扱う必要性

分析の目的が (6a) なら、これは明らかにヒト x が文 s を読んだり聞いたりする際に x が実際に理解している「内容」の記述として不完全である。§4

で詳細を示すことになるが、ヒトの言語理解の内容は驚くほど具体的である。この点は、認知言語学に限らず、従来の言語学ではまったく自覚されていない。意味が重要だと声高に叫ぶ割には、表層形として現われていない内容の記述を怠っている—正確には、そのような対象を首尾一貫した形で扱う手法を言語学は今だに有していない。

実際、文意 (sentence meaning) の十分に妥当な記述、つまりヒトが文を聞いたり読んだりした際に何が理解されているのかを正確に、詳細に記述しようという試みは、言語学はおろか、認知言語学の内部にすら存在しない。

この点、文意に妥当な記述を与えようとする関心は生成語彙理論 (Generative Lexicon Theory) [8, 9, 10, 11] を動機づけているものであり、この方面では明らかに計算言語学の研究成果が認知言語学の研究成果に勝っている。

任意の文の文意の妥当な記述は重要な課題である。言語学分析は、語の多義性 (lexical polysemy) を扱うことができるだけ、語句の多義性の構造がどうやって生じるかが説明できるだけでは十分ではない。語句の多義性の構造の記述は正しい言語理論にとって必要不可欠だが、それで十分ではない。語句の多義性を解消する仕組みにも十分な説明を与えなければならない [12]。

注意深い言語研究者なら誰でも知っていることだが、語の意味は文脈—正確には共起している他の語の意味—から独立してない。語は一定の文脈の置かれたとき、一定の仕方で多義性が解消されるが、その脱曖昧化 (disambiguation) の妥当なモデル化ができなければならない。さもないと、語の意味の柔軟性に対応できない。

とすれば、文脈効果 (contextual effects)⁵⁾ の名で知られる効果を真剣に取り扱わない限り、認知言語学は不十分な理論体系である。

脱曖昧化の現象は、認知言語学内部で自覚されていないわけではない。例えば、[14, 4] は解釈の際に

⁵⁾ただし、関連性理論 (Relevance Theory) [13] の言う文脈効果ではなく、認知心理学で認定されている文脈効果である。

語の意味が文脈に対し意味の (相互) 調節 (semantic (mutual) accommodation) が示すことを認定している。だが、これは事実の認定にすぎず、突っ込んだ議論はない⁶⁾。これが意味することは、(認知) 言語学は真面目に文脈効果を扱うべきだということであるが、そのためには目的に見合った、特別な道具立てが必要である。

その理由を理解するのは簡単である: 文脈効果は複雑な現象である。それは明らかに、多体問題 (multi-body problem) である。文 s が n 個の語からなるとすると、該当する意味的調節は全部で $n!$ 通りある。 s の理解は、このような多数の相互作用の結果—おそらく自己組織化—の結果として生じる。この相互作用が (比喻) 写像やような規則ベースの手法で単純に記述できる現象でないのは明らかである⁷⁾。

このような目的のため、私たちはフレーム意味 (Frame Semantics: FS) [15, 16, 17, 18] の拡張である多層意味フレーム分析 (Multi-layered Semantic Frame Analysis)、並びに階層フレーム網分析 (Hierarchical Frame Network Analysis) という分析法が、その目的のために有効だということを論じる。詳細は §4 で展開する。

2.5 認知言語学は過剰般化の常習犯

ここで一旦、始めの論点に戻ろう。

ネットワークモデルやイメージスキーマによる記述/説明が (6b) だとすると、まず、拡張が起こったプロセスどおりに心的表現が構成されているという保証はどこにもないので、結果が概念分析になってない可能性が高い。

それに加えて、ネットワークを構成するメトニミーリンク、メタファーリンク、コネクターのよう

⁶⁾記述の明示性に関して言うと、生成辞書理論 (GLT) [8, 9, 10, 11] の方が同様の事実をずっと明示的に取り扱っている。GLT の難点を挙げれば、認知心理学的な妥当性が考慮されていない点であろうか。

⁷⁾そのことを直観できないならば、始めから事実を認識していないという意味で観察力が欠けているか、あるいは複雑系の科学者として現象の複雑さを正しく評価するための感受性が欠けていると思われる。

な各種のリンクを心内プロセスと見なすなら、直観のみに基づいた言語データの直接分析から判断できることは極めて限られており、そこから一般化して言えることは、非常に限定されるはずである。この観点からすれば、概念比喩理論 [19, 20, 6, 21, 5], ブレンド理論 [22, 23, 24, 25, 26, 27, ?, 28] を擁する認知言語学者は、「不十分証拠に基づく過剰な一般化」の常習犯ということになる。

2.6 経験基盤主義は単なる言い逃れ

P 氏は続いて、(6a), (6b) に共通の問題として、次の点に関する説明が皆無であることを指摘する：

- (7) a. 各種のリンクは何から生じているのか、
- b. イメージスキーマは何から獲得されるのか

これは驚くべきことではない。事実上、リンクは事実上、認知的基本要素 (cognitive primitives) として扱われているし、イメージスキーマは認知的必然性 (cognitive necessities) と見なされている。

これらの「実証」の要請に対して、単に「経験基盤」説 (experientialism) をもちだし、それらに「身体基盤」 (embodiment) があると抽象するだけでは、十分ではない。それは「神様が決めたから」というのと同じ程度にしか意味がない。これらの理論仮構物のすべてに対し、すでに §2.2 で述べた理由で再検討が必要である。

このようなこと背景に、P 氏が認知心理学者として部外者の観点から認知言語学に対してもった印象は、結局「スローガンや方向づけはいい感じ、でもねえ…」なのである。

P 氏が指摘するのは、用法基盤とか身体経験とか言っている割に、その内実に対する直観も考察も、認知心理学の見地から見れば不十分だという点である。特に、実際に私たちが得ることのできるのは個別的で具体的な個々の経験でしかないということが忘れられていると P 氏は指摘する。

生成言語学の思弁を廃すると言いながら実は、生成言語学の言語能力 (competence) と同じくらい所

与のもとになっている「認知機能」が役割が大きい。これでは「神様」を「言語」能力から「認知」能力に取り換えているだけである⁸⁾。

2.7 結論

このような現状は明らかに好ましいものではなく、打破する必要がある。だが、どうやって？

具体的な解決法は §4 に示すことにして、§3 ではまず何が問題なのかを明らかにする。

3 ネットワーク分析を越えて

認知意味論 (Cognitive Semantics: CS) [6], 認知文法 (Cognitive Grammar: CG) [14, 4, ?, 29] のような認知言語学の枠組みでは、ネットワークモデル (Network Model), あるいはネットワーク分析 (Network Analysis: NA) が概念体系の適切な分析モデルであると論じられ、興味深い実例も数多く存在する。手軽にアクセスできる例を挙げれば、例えば、Lakoff [6, p. 436, Fig. 27] は *over* の多義性を ICM のネットワークとして分析しているし、Langacker [14, p. 74, Fig. 2.3a] は [APPLE], [BANANA], [PEAR], [TOMATO] を [FRUIT] のスキーマのネットワーク (schematic network) の拡張として記述し、[29, p. 123] は、構文交替現象を捉えるため、[send NP1 NP2], [send NP2 to NP1], [give NP1 NP2], [give NP2 to NP1] を部分的にネットワークで分析し、[29, p. 134] は、Rubba [30] を紹介しながら、現代アラム語 (modern Aramic) の非線型形態論の断片をネットワークで分析している。また、その根拠が明示されることは多くないが、構文文法 (Construction Grammar) の分析 [31, 32, 33] もネットワーク分析を基本的に踏襲している。

だが、ネットワーク分析には問題がないわけではない。第一にネットワーク分析に明らかな欠点が存在するという点でそうであり、第二に「そもそもネットワークとは何か？」という問題に明確な答

⁸⁾この種のご都合主義は、認知言語学に限らず、日本への米国産の言語理論が「侵入」を開始した時から一貫して認められる傾向であるけれど。

えがあるわけではないという理由からそうである．
これらを明らかにすることから論を始めよう．

3.1 階層性の不在：ネットワーク分析の欠点 1

ネットワーク分析は確かに有用な分析手法だが，一見して明らかな限界もある．その一つは，ネットワーク N 自体が N を構成する構造のあいだの階層性 (hierarchy) を表わすための制約を内在していないことである⁹⁾．

Langacker [14, p. 74] は，階層性をもっと自然に表わしそうな彼のスキーマのネットワークに関して，次のように規定する：

- (8) I refer to such an assembly of categorizing units as a **schematic network**. Schematic networks are conveniently represented in diagrams like Fig. 2.3(a), which often resemble standard taxonomic hierarchies. However, *these are no intrinsic restrictions on the configuration of these networks*, and diagrams like Fig. 2.3(a) can be regarded as summary abbreviations for sets of individuals categorizing relationships, as shown in Fig. 2.3(b). Categorizations like these define a **schematic plane** of relationships, which are fundamental to linguistic organization. Structures in this plane serve three crucial functions, which are often considered distinct but receive a unified account in cognitive grammar: (1) categorization; (2) the capture of generalizations (expressed by schemas); and (3) the sanction of novel structures (the categorization of nonunits). [著者によるイタリック体による強調]

イタリックで強調した部分をなぜ強調するのは，私たちにはサッパリわからない．ネットワークが制約されていないということは，それが何も説明しないということではないのか？

ネットワーク N 自体が N を構成する構造のあいだの階層性を表わすための制約を内在していないことは，次の図 1 にある三つの構造記述モデルのあいだの対応関係を検討すれば，直ちに明らかである．

図 1 は，構造の階層モデル (hierarchical model)，集合モデル (set model)，ネットワークモデルの三つのモデルのあいだの対応関係を表わしたもので

⁹⁾これは事実，認知言語学の研究成果が — 例えば WordNet [34] のような — 非認知系の研究結果とうまく折り合いがつかないことの一因となっている．

ある．最上部にあるのが階層モデル，最下部にあるのがネットワークモデルである¹⁰⁾．

上部二つの階層モデル，集合モデルを統合すると，オブジェクト指向プログラミング (Object-Oriented Programming: OOP) を通じて普及したクラス/インスタンス分析 (Class/Instance Analysis) [35, 36] となる¹¹⁾．

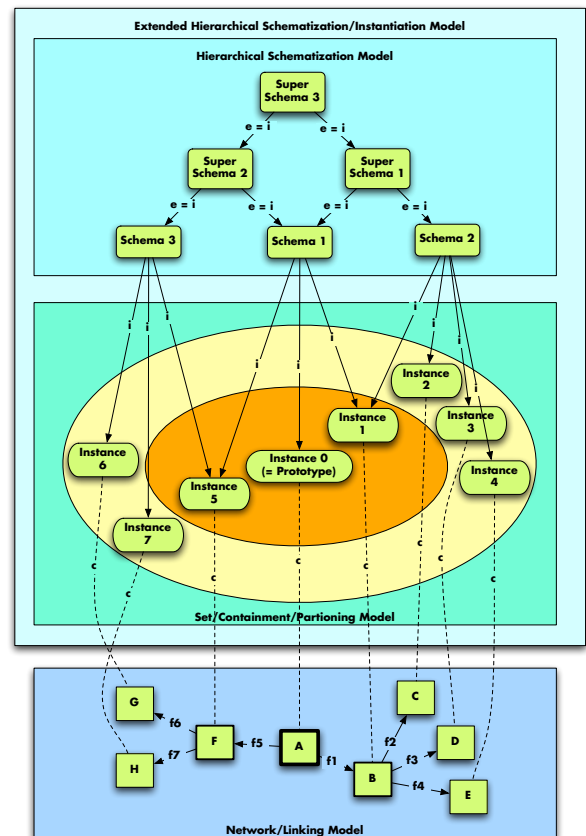


Figure 1: 階層モデル，集合モデル，ネットワークモデルのあいだの対応関係 (e は詳細化，i は具体化，c は対応関係を表わし，f_i は各種のリンクを表わす)

¹⁰⁾この点に関して言うと，Langacker のスキーマのネットワークは最下層にあるネットワークとは別種の構造であり，明らかに何を称してネットワークと呼ぶのかに関して，認知言語学者のあいだで概念的に混乱があることになる．

また，数学的にはネットワークは単なるグラフ構造 (graph structure) なので，下層にあるもののみをネットワークと呼ぶ必然性はないのだが，これはこれでネットワークの概念が形骸化して，好ましくない．実際，この定義では，ツリー構造もネットワークの一種なのである．

¹¹⁾[37] に認知文法と OOP/OOA の対応関係に関する面白い考察がある．

3.1.1 モデルの組織化の原理

三つのモデルはおおの、組織化の原理が異なる．具体的には，

- (9) 階層モデルの組織化の原理は，具体/具現化 (instantiation) = 詳細化 (elaboration) である (図では e が詳細化を， i が具体化の関係を表わしている)．図には明示されていないが，具現化の逆は(上位)スキーマ化 ((super-)schematization) である．これは抽象化 (abstraction) と同一視できる．このような特徴があるために，このモデルは WordNet [34] や NTT 日本語語彙大系 [38] のようなシソーラス (thesaurus) として体系化された概念の階層性をうまく記述する，
- (10) 集合モデルの組織化の原理は包含関係 (inclusion) である．これはベン図によって表わされている．
- (11) (狭義の) ネットワークモデルの組織化の原理は，様々な種類の結びつけ (linking) (の関数) である．この関数はリンク (links) と呼ばれ，図では f によって表わされている．メタファーリンク (metaphor link) やメトニミーリンク (metonymy link) [6] がリンクの代表例であるが，それには明らかにコネクター (connectors) [23, 24, 25] も含まれる (図 1 には，メトニミーリンク，メタファーリンク，コネクターの種類の区別は表わしていない)．

なお，詳細化の関係 e と具体化の関係 i の区別は (“ $e = i$ ” と明記してあるように) 本質的ではない．実際， e は i の特殊な場合である¹²⁾． e, i には (必然的に) 方向性があるが，対応関係 c には方向性はない．

3.2 体系性の由来の不透明さ: ネットワーク分析の欠点 2

ネットワーク分析の背後にある組織化の原理が各種のリンクによる構造の結びつけであると分かっ

たので，ネットワーク分析のもう一つの欠点を述べるができる．それは，ネットワーク分析が事例同士の (あるいは構造同士の) 関係のタイプを特定しないという点である．この欠点は，各種リンクの存在論が「説明」されない限り，決して解消されない¹³⁾．この理由により，ネットワークモデルによって記述された体系は，その由来が不透明となる．

具体例を挙げて論じよう．図では， A が中核(事例) (core member(s))，別名プロトタイプ (prototypes) で， B, F は A の拡張(事例) (extensions)，あるいは周辺(事例) (peripheral members) である． C, D, E は A の， F, G は B の，拡張(事例)例である． B, F は A に f_1, f_2 のリンクで， C, D, E は B に f_2, f_3, f_4 のリンクで結びつけられているが， $f_1, f_2, \dots, f_7$ のタイプはどれも同じだとは限らない．そのうちの幾つかはメタファーリンクであり，また幾つかはメトニミーリンクであるだろう．だが，それを図にあるネットワークの構造のみから知ることはできない．そのことを明示するためには $f_1, \dots, f_7$ にいちいちラベルづけする必要がある．ネットワークモデルは事例同士の関係のタイプを特定しないとは，このことを言う．

このような問題はクラス/インスタンス分析では生じない．なぜなら，可能な結びつけはスキーマの具現化 ($e = i$) と，その逆である事例のスキーマ化のみに限定されているからである¹⁴⁾．

明らかにネットワーク分析とクラス/インスタンス分析には一長一短がある．ネットワーク分析ではメタファーやメトニミーがネットワークのどこに起こっているのか明示することは簡単である．それはリンクを分類すれば良いからである．これに対し，クラス/インスタンス分析ではメタファーや

¹³⁾ 実際，認知の仕組みの解明の目的に照らして見ると各種のリンクは，文法記述の目的に照らして見たときに，60 年代，70 年代に生成文法家が言っていた文法規則 (rules of grammar) がもっている説明力と同じ資格しかもたない．それらはいずれも同じく，現象に関する記述的一般化以上のものではない．従って，文法規則が文法を説明しないのと同じく，リンクはヒトの認知の仕組みを説明するわけではない．

¹⁴⁾ クラス/インスタンス分析で問題となるのは，単に表記の面での煩雑さである．

¹²⁾ あるいは反対に， i が e の特殊な場合であるとも言える．

メトニミーが体系のどこで起こっているのかを明示することは困難である。それは相当に複雑なスキーマの具現化のプロセスの副産物だからである。

これに対して、クラス/インスタンス分析ではメタファーやメトニミーがどのように生じているのを知るのは比較的容易である。これに対して、ネットワーク分析ではメタファーやメトニミーがどのように生じているのを知ることは困難である。それを明示しようと思ったら、(比喩)写像 ((metaphorical) mapping) という形で記述的一般化を述べるしかない。

だが、比喩写像という形での記述的一般化は、明らかに表面的な説明である。従って、認知言語学者が認知科学的な意味での説明的妥当性を追及するならば、メタファーリンク、メトニミーリンク、各種コネクタのような (ICM 間の) リンクの実体性を真に受けるべきではないのである。そのようなリンクは、それ自体がより“根本的”な説明項 (例えばニューロンの活動パターン) によって記述され、説明される必要がある¹⁵⁾。

認知の根本的な説明項は、言語や思考のような「高次」認知現象の記述に関して妥当であるばかりでなく、知覚のような「低次」認知の記述と互換性をもつものでなければならない。高次認知と低次認知のあいだには非常に大きな隔たりがあり、高次認知レベルで得られた一般化がそのままの形で低次認知に妥当すると考えることは—認知言語学では往々にして認められることだとは言え—無理を通り越して無謀である。

高次認知レベルの分析が低次認知レベルで妥当性をもつことを保証するためには、それがニューロンの構造、ニューロン体系の挙動と互換性をもつものであることを保証しなければならない¹⁶⁾。個々

¹⁵⁾これが意味することの一つは、[20, 21, 6, 5] の概念比喩理論 (Conceptual Metaphor Theory) は、実際には認知の仕組みに関する記述的一般化の羅列であって、正しい意味での説明ではないということである。

¹⁶⁾ここで言う互換性は、単なる用語上の互換性 (terminological compatibility) ではない。このことを明記するのは、ある言語学の理論、理論的枠組が例えばコネクショニズム (connectionism) と互換性があると主張するだけでは十分ではないという点を強調するためである。例えば、Langacker の認知文法

のニューロンは、素性 (feature) として記述できるような何らかの特徴 (characteristic) の有無をコードする以上のことはしていないという事実を理解するのは、この点で本質的に重要である。

3.3 ネットワーク分析の総合評価

以上の考察に基づいて認知言語学で主流となっているネットワーク分析の今日的評価として、ここで結論的に言えることは、次のようなことだろう:

- (12) 各種のリンクは高次認知レベルの記述的一般化としては妥当であるが、それが知覚のような低次レベルの認知の仕組みにどれぐらい互換なのかは、まったく未知数である。

とりわけリンクを用いた記述的一般化の、神経細胞レベルの低次認知機構との互換性を確保するためにしなければならないことは、例えば、

- (13) 特に各種リンクが素性表示 (feature representation) と互換性をもつことを示すこと
- (14) 各種リンクによって結びつけられている表象 (representations) の「内容」を十分に詳細に特定し、記述すること

の二つである。

だが、これを達成するのは簡単なことではない。少なくとも、これまで認知言語学がなし遂げたたと一般に考えられている成果の一部を放棄しなければならないのは明らかである。

(13) が意味することは、認知言語学者はリンクのような高次認知レベルの理論仮構物を認知的実体として真に受けてはいけないうに気をつけ、低次認知機構と高次認知機構とに「橋渡し」があるようなモデル化に留意する必要があることである。これに関してはすでにある程度詳しく論じたので、細部は繰り返さない。

の枠組みがコネクショニスト・モデルと単なる用語上の互換性以上のものをもつかどうか、Lakoff-Johnson の Neural Theory of Language がコネクショニズム全体の手前勝手な矮小化に基づかないものであるかどうか、極めて疑わしいと私たちは考える。

(14) が意味することは、認知言語学の目指すべきことは、細部を犠牲にして中途半端な一般化を述べるという「手抜き」を行なってはならないことである。これは生成言語学以前の記述主義の再評価と再興につながる姿勢である。これは、§5.3.1 で用法基盤の文法モデル (usage-based model of grammar) の詳細を論じる際に指針となる考えである。

3.3.1 いい加減な説明より妥当な記述を

これに関連して特に強調しておきたいのは、言語の認知科学への貢献を目指す語学に必要なのは、徹底した現実主義、記述主義だということである。言語学者は言語という自然的な現象を科学的解析する研究者として、機能の面でも構造の面でも、言語の細部は、しばしばグロテスクなくらい錯綜しているという現実を直視するべきだということである。私たちは、生成言語学が長らくそうしてきたように「認知」や「心」を隠れ蓑にするべきではない。

その理由は簡単で、認知という自然的な現象は、機能の面でも構造の面でも、しばしばグロテスクなくらい錯綜した細部をもっているという現実があるからである、何より重要なことは、そのような細部は、まだまだ十分によく知られておらず、詳細に記述するに値するという事実である。

明らかに、妥当な記述なしに妥当な説明などありえない。認知言語学がこれまでどれほど妥当な説明を行ってきたかは、偏にそれがこれまでどれだけ詳しい記述をどれだけ十分に行ってきたかにかかっている。

3.3.2 認知言語学の「成果」の公正な評価

この点に関して言うと、認知言語学の達成した「成果」には大いに問題ありとしなければならない。言語構造の記述に「ありもしない体系性」をもちんで平気な顔をしているのは、チョムスキー革命以来の悪癖であるが、これは残念ながら認知言語学に移行して改められるどころか、むしろ悪化している感すらある。

今の認知言語学の枠組みに何よりも必要なのは、みかけの説明力の向上ではなくて、実質的な現象の記述力を向上である。この際に重要になるのは、リンク自体の一般化ではなく、互いにリンクされている表象内容の、十分に精密で妥当な記述である¹⁷⁾。

3.3.3 意味フレーム分析

以下で私たちが取り組もうと思っているのは、互いにリンクされ、全体としてネットワークを構成している個々の表象の具体的「内容」を意味フレーム (semantic frames) の概念を用いて詳細に記述することである。そのための理論的枠組みは言語のフレーム指向概念分析 (Frame-Oriented Concept Analysis of Language: FOCAL) と呼ばれる。

FOCAL は Fillmore のフレーム意味論 (Frame Semantics: FS) [15, 16, 17, 18]、並びに **Berkeley FrameNet** (BFN) [39, 40] の枠組みを概念的に拡張したものである。

BFN/FS と FOCAL には明白な差違が幾らか存在するが、この論文では詳細には立ち入らない。関心のある方は、[?, 41] のようなオンライン論文を参照されたい。

3.3.4 イメージスキーマへの還元は必要か？

すでに明言したように、私たちが望んでいるのはネットワークを構成する表象の「内容」の具体的な記述である。その際、私たちは具体的な記述をイメージスキーマ (image schemas) [6, 20] へ帰着させることは考えない。少なくとも認知科学の成果の現状を踏まえる限り、それは、控え目に言っても未熟すぎる方向づけであり、おそらく自家撞着的な説明以上の成果はもたらさないだろう。

繰り返しになるが、私たちは未熟な説明項による、空虚な説明は求めていない。私たちが求めているのは、あれやこれやの事例にあてはまるよう

¹⁷⁾ 認知言語学者は概して言うと、表象という説明概念に対して、生成言語学への反感に端を発して反感をもつことが多いけれど、そのような反感は忘れてもらう必要がある。

に見えるが、正確にどれくらい当てはまっているのか正しく評価できないような空虚な一般論ではない。

私たちが求めているのは、言語に関する現象の詳細な記述を通じた、言語の認知科学への実質的な貢献である。私たちは、過去 10 年間で認知言語学が学派内部での「内輪ウケ」に明け暮れる以上のことをなしたのか非常に怪しく思うし、そのような現状をまったく肯定しない。

4 意味フレーム分析の理論と実践

[この節の内容は 01/10/2005 の時点で未完]

4.1 簡単な実例: x が y を襲う

イメージスキーマへの還元の本質的難点は、それが特定する概念構造が抽象的すぎ、ヒトの発話理解/文章理解の具体的内容の重要な部分をまったくと言ってよいほど反映しないという点である。

ヒトの言語理解の内容が驚くほど具体的であり、信じられないほど詳細であるのを見るには、次の (15) のような簡単な例を検討するだけで十分である:

(15) 大型の台風が九州 (の人々) を襲った

(15) には

- (16) a. “九州で暮している人々が、台風の接近により損害を受けた”
b. “九州で暮している人々の生活が、台風の接近により害された”

とは明示的に書かれていないが、(15) の理解には、明らかに (16) に明示したような内容が伴う。

このような理解内容の大部分はイメージスキーマでは説明できないし、また、このような理解がメトニミーだと言ったところで、それはせいぜい現象を正しく分類しているだけで、なぜそれが可能なのかに関して、実質的に何かを説明しているわけではない。このような効果は、特別の方法で記述しな

いと、存在を特定することすら難しい。以下、そのために技法を多層意味フレーム分析 (Multi-layered Semantic Frame Analysis: MSFA) という名の下に紹介する。

4.2 多層意味フレーム分析 (MSFA)

(15) の理解内容を意味フレームを用いて解析したのが表 2 に示す MSFA である。

表の列には意味フレーム $F_1, \dots, F_{13}$ が、行には (15) の形態素 $M = \{ \text{大型, の, 台風, が, 九州, (の人々), を, 襲っ, た} \}$ を配置している。フレーム F と形態素 m の交点にあるセルには、 m が実現する F の意味役割名が指定されている。セルに指定がない場合、意味役割が実現されないということである。

例えば、[台風] は

- (17) a. F_1 : 〈発生〉フレームの意味役割の一つである〈発生体〉,
b. F_2 : 〈経路移動〉フレームの意味役割の一つである〈経路移動体〉,
c. F_3 : 〈(経路不定) 移動〉フレームの意味役割の一つである〈移動体〉,
d. F_5 : 〈加害〉フレームの意味役割の一つである〈加害体〉,
e. F_6 : 〈被災〉フレームの意味役割の一つである〈原因〉,
f. F_7 : 〈経験〉フレームの意味役割の一つである〈内容〉,
g. F_8 : 〈働きかけ〉 (= 使役) フレームの意味役割の一つである〈影響源〉

という複数の意味役割を同時に実現している。

このような実現関係のうちの一部は具現化である。例えば、 F_2 : 〈経路移動〉は F_3 : 〈(経路不定) 移動〉の特殊な場合であり、移動体としての性質は F_3 から継承 (inherit) されている。

		F1	F2	F3	F4	F5	F6	F7	F8	F9	F10	F11	F12	F13
Index	形態素	発生 [+inferr ed]	経路移動 [+inferr ed]	移動 [+inferr ed]	経路 [+inferr ed]	加害	被災	経験 [+inferr ed]	働きかけ	ヒトの生 活 [+indire ct]	生物の棲 息 [+inferr ed]	国土	収容	分割
1	大型	規模				規模?	規模?	規模?						
2	の	MARKER												
3	*	GOV												
4	台風	発生体	移動 体:EVO	移動体		加害体	原因	内容	影響源					
5	が					MARKER								
6	*	発生地	起点	起点	起点				被影響体					外部
7	*		着点	着点	着点									
8	*		通過点	経過 点:EVO	経過点: EVO							国境	収容 器:EVO	境界
9	*											領土		内部:EVO
10	九州					場所	被災地	場所		生活地域	棲息地域	地 域:EVO	収容物	
11	*									GOV				
12	(の人々)					被害者	被災者	経験者		生活者	棲息生物			
13	を					MARKER								
14	襲っ					GOV	EVO							
15	た													

Figure 2: (15) の MSEFA

4.2.1 意味役割に関する概念上の注意

意味役割は主題役割 (thematic roles/ θ -roles) ではない。主題役割は意味役割の特殊な場合である。主題役割が特殊な意味役割だと言うのは、それらが統語構造に反映される特殊な意味役割だという意味においてである。実際、意味役割の特徴の大部分は統語には反映されない¹⁸⁾。

4.2.2 効果としてのメトニミーはどこから来るか

これがメトニミーだというのは、控え目に言っても現象の矮小化である。(15) の内容理解には少なくとも次のようなメトニミー関係の解消が係わっている:

(18) R1. [九州] → [九州で暮す人々]

R2. [九州で暮す人々] → [九州で暮す人々の生活]

R3. [九州で暮す人々の生活] → [九州で暮す人々の生活空間に関係する様々な物事]

強風で家が壊れたり、電気や水道が止まったり、

¹⁸⁾これが理論的に含意することは決して軽微ではない。例えば、統語に反映されるかどうかを意味役割の認定基準してはけない。それは実際、意味役割の定義に反する。意味役割の認定基準は、純粋に概念的なものでなければならない。この理由から、統語に反映されないことを理由に意味役割でないことを認定してはいけない。

R1, R2, R3 は〈人々〉が、〈特定地域〉で、生活する〉フレーム (F9) を基盤にして解消されている。

4.2.3 文脈「効果」としてのメタファーはどこから来るか

「台風が y を襲う」という表現は、明らかにメタファーである。台風は生物ではないので、実際にはヒトを襲ったりはしない。これは明らかに次のような例とはちがう:

(19) 空腹のライオンがインパラの群れを襲った。

(20) 覆面の男が銀行を襲った。

だが、(15) のような比喩はどこから来るのか? それを説明するのに、例えば (21) のような概念比喩を想定すれば十分なのだろうか?

(21) [自然災害は捕食者である]

([NATURAL DISASTER IS A PREDATOR])

そんなことはまったくない。[42] で断片的に示したことだが、(15) の内容理解を可能にしているのは、実際にはイメージスキーマのような抽象的な概念構造ではなくて、もっと高次の複合的概念構造、例えば〈ヒト y が台風 x に襲われる〉ことのイメージ構造であり、それは、

- (22) a. 〈 x の接近によって発生する強風で y の住んでいる家 z が部分的に壊れた〉り、
 b. 〈 z に供給されている電気や水道が一時的に止まった〉り、
 c. 〈 y が日常的に利用している交通機関が一時的に利用不能になった〉り、
 d. 〈 y の知り合い y^* が事故に遭った〉り

すること — そういう一連の具体的なイメージ群が一つのまとまりをなした全体、あるいは状況のクラスター (situation cluster) が喚起 (evoke) され、イメージのネットワークが想起できることである。

このようなイメージのネットワークは〈ヒト y が台風 x に襲われ、直接、間接に被害を被る〉ことの ICM だと言える。だが、本質的な問題は、その ICM を構成する具体的なイメージの一つ一つを、例えば、

- (23) a. SOURCE-PATH-GOAL スキーマ
 b. CONTAINER スキーマ
 c. LINK スキーマ

のような概念的“根源”に還元することにどれほどの価値があるか、という点にある。

私たちが以下で示そうと思っているのは、一般に認知言語学で行われている分析の反対である。具体的には、次のことを示すことになる：

- (24) 例えば (15) の内容理解に (23) に列挙したような抽象的なイメージスキーマが貢献してる程度は一般に信じられいてるよりも少ないこと、
 (25) それ故、一般に理解内容のイメージスキーマへの還元は — 言語学が理解内容の妥当な記述を目指す限り — 有害であっても有益ではない。

4.2.4 抽象的イメージスキーマは理解にどう貢献するか

(15) の理解内容に関与する意味フレームでイメージスキーマでうまく特徴づけられるものは、F1: 〈台

風の発生〉、F2: 〈台風の経路移動〉、F12: 〈収容〉である。F1 は出現スキーマで、F2 は経路 (移動) スキーマで、F12 は容器スキーマで特徴づけることができるだろう。

確かに台風は経路移動すると理解される。その際、移動の起点 (source) は台風が発生した地点 (おそらくフィリピン沖) で、台風の (通過) 経路 (TRAJECTORY = PATH) 上の通過点 (TRANSITORY POINT) が九州地域であったことは理解される。

だが、それ以外の意味フレームはどうだろうか？例えば、F5: 〈加害〉フレーム、F7: 〈経験〉フレームは、それらにうまく対応するイメージスキーマはあるのだろうか？

比喩的例を考慮に入れると、〈打撃〉フレームが一部の〈被害〉や〈被災〉の比喩的理解の背後にあるのがわかる。だが、〈打撃〉スキーマというものはあるのだろうか？それが存在するとしても、それは (23) にある抽象レベルのイメージスキーマと同じ性質のものだろうか？

そうだとはいえられない。無理に説明に固執するのでなければ、「抽象的イメージスキーマには個々の文の理解内容を十分に詳しく記述する効力が無い」と素直に認める方が簡単だし、なすべきことをなすための道も開ける。これに対し、意味フレームは文の理解内容を十分に詳しく記述する効力がある。

FOCAL は用法基盤を真剣に考える。その結果、認知言語学の主流とは違い、記述内容の具体性を美德とし、過剰般化に基づく中途半端な説明を狙わない。これが私たちの分析手法が理解内容の記述に有効だと考えられる最大の理由の一つである。このような目的にとって、意味フレームは好ましい記述装置である。それは状況を基盤にしているので、特定する構造記述の内容が具体性である。

もちろん、(23) にある抽象的イメージスキーマの (15) の内容理解における貢献は無ではない。それは無ではないのだが、認知言語学で考えられているのとは働き方が異なっている。抽象的イメー

ジスキーマの貢献は、具体的理解内容の特定に対してより、適切な統語構造の選択に対して働く。これが (23) にあるようなイメージスキーマ群が言語分析で目立つ理由である。

だが、これは概念構造の構成原理として、これらの抽象的イメージスキーマ群が、[20, 6] が主張しているような仕方で概念体系の形成、保持に強力に機能していることは意味しない。従って、[20, 6] が試みているように、抽象的イメージスキーマ群が概念体系の形成、保持に強力に機能していることを示したいなら、それは言語分析とは別の仕方で達成されなければならない。言語が—正確には統語構造が—一部の抽象的イメージスキーマ群と結びついていることが示されたからといって、それから概念体系が一部の抽象的イメージスキーマ群と結びついていることは示されていない。

4.3 階層的フレーム網分析 HFNA

以上の注意点を明確に表わしてくれるのが、図 3 に示した (15) の階層的フレーム網分析 (Hierarchical Frame Network Analysis: HFNA) である。

左端のあるのは (15) の形態素解析で、それは F1: 〈発生〉, F2: 〈経路移動〉フレーム, F5: 〈加害〉フレームに対応づけられている。

全体として、図の左にある構成物から右にある構成物に対して喚起 (evocation) の関係、あるいは活性化 (activation) の流れがある。ただし、矢印で表われている具現化の関係は意味フレーム単位ではなく、意味役割単位で示されている。

4.3.1 HFNA はクラス/インスタンス分析の一例

HFNA はクラス/インスタンス分析の一例であり、概念体系の階層性をうまく表現する。左端が最下位のクラス、右端が最上位のクラスである。図 1 では階層関係が上下関係で表わされていたが、図 3 では左右関係で表わされている。

5 FOCAL の概要

この節では FOCAL の枠組みを概説するが、すでに述べたように、FOCAL は Berkeley FrameNet¹⁹⁾ に啓発された枠組みである。従って、まず、背景知識として、BFN の情報を簡単に提供しておくことは、FOCAL の理解の助けになると思われる。

5.1 Berkeley FrameNet とはどんな研究企画か

BFN は意味フレームの大規模なデータベースを開発する研究企画である。開発は第二期目に入っており、現時点で数百程度の意味フレームがデータベース化されている²⁰⁾。BFN の日本語版は、日本語フレームネット (Japanese FrameNet: JFN)²¹⁾ [44] という名称で進行している。

BFN は Fillmore の FS [16] の発展的応用であるが、FS から BFN への移行は単調ではなく、重要な概念的変更も含まれる²²⁾。例えば FS 初期の [16] の時点で想定されていた (意味) フレームは、解釈フレーム (interpretive frames) と呼ばれ、相当広い意味で理解の背景となる知識構造を指すもので、Fillmore 自身の表現を借りれば、“unified frameworks of knowledge, or coherent schematizations of experience” [16, p. 232] である。

これに対し、BFN では同じく (意味) フレームという名称が使われていながら、記述対象が事実上、状況 (内の行為) のような場合に限定されている。このような変更は理論的なものなのか作業の都合によるものなのか定かではない。これから示唆されるように「フレーム」という語が何を規定しているのかは、フレーム意味論の枠組みの中ですら必ずしも一定でも明白でもない。BFN 流の語義の分析の具体例を一つ、A に挙げた。

¹⁹⁾ ホームページは <http://www.icsi.berkeley.edu/~framenet/> で、数多くの良質なオンライン論文が入手できる。

²⁰⁾ 詳細は [43, 40]などを参照されたい。

²¹⁾ ホームページは <http://www.nak.ics.keio.ac.jp/jfn/index.html>

²²⁾ FS と BFN との違いの詳細は [40]などを参照されたい。

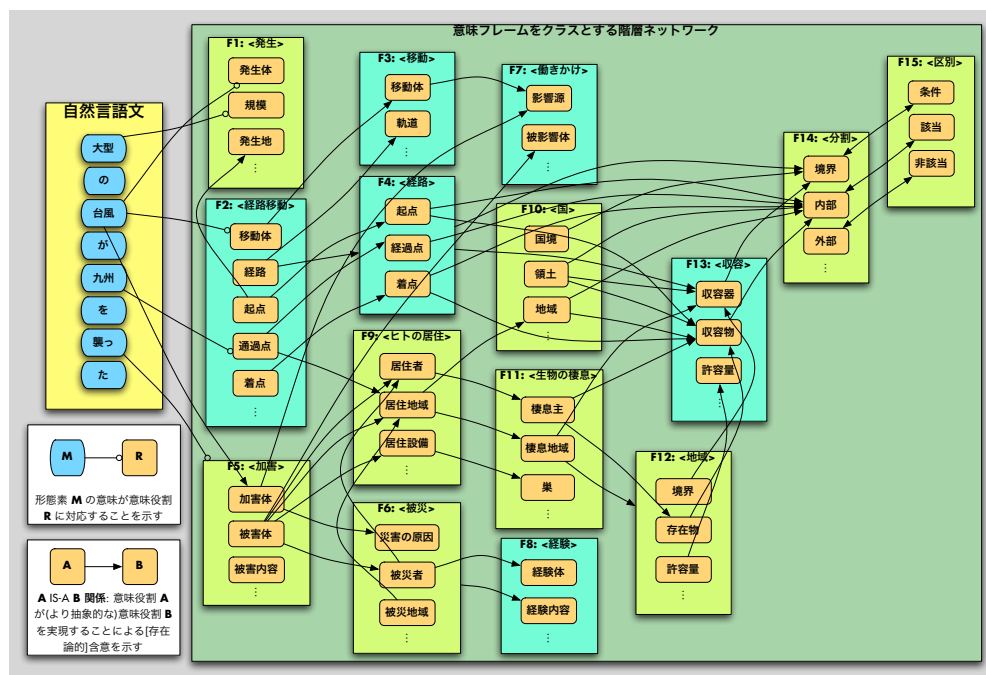


Figure 3: (15) の HFNA

変更の理由はともかく、私たちが(意味)フレームと呼んでいるものは、初期 FS の解釈フレームとしての(意味)フレームではなく、BFN の意味での、より限定された(意味)フレームの概念に相当する。FS に関する既得知識がある場合、この違いは誤解の元になりがちなので、留意して頂きたい。

5.1.1 本家 FS からの「逸脱」に関する注意

説明の便宜のため、Fillmore が定義した FS [15, 16, 17, 18] を Berkeley FS (BFS) と呼ぶ。私たちの研究は BFS の拡張だが、現時点での研究方向づけ、特に作業仮説は BFS のそれと同一というわけではない。私たちはフレームの定義に関して、BFS に比べて限定的な立場を取っている。

具体的には、BFS は具象名詞 (e.g., { 本, 机, メガネ, ... }) がフレームをもつことを許す—正確には排除できない—が、その定義は寛容すぎて効力がなくなる危険があると私たちは考える。これら具象名詞群が何らかのフレームに関係していること—特にフレームを喚起する効果があること—は明らかであるが、このこと自体は、これらの概念がフレーム構造をもつことは意味しない。(フレームと

いうデータ構造を使った)モノの内部構造の記述には、それらが使われる状況の記述が含まれるべきではない。そうではなくて、そのような情報は(フレームというデータ構造を使った)状況 σ の内部構造 $F(\sigma)$ の記述に、 F を構成する意味役割の実現値として記述されるべきである。そうでないとモノのフレームの記述内容に歯止めが利かなくなる。

このような判断を動機づけているのは、BFS の枠組みに許されている自由度の過剰に対する懸念である。²³⁾ 自由度が過剰であると、説明は恣意的になりがちである。

5.1.2 フレーム意味論を制約する必要性

私たちが想定する「制約された FS の枠組み」は、BFS との区別のために FOCAL (Frame-Oriented Concept Analysis of Language) と呼ぶ。詳細は、黒田ほか [1]、中本ほか [45] を参照されたい。以下、これら二つの枠組みの共通点と相違点について、紙面が許す限りで解説する。特に、意味フレームと

²³⁾ 基本文献 [15, 16] ではフレームの例が列挙されているばかりで、ある概念構造がフレームと呼べるための判定条件に関して明確な規定がないことが、この自由度の過剰の原因である。

いう用語が特定する概念内容が基本文献で多かれ少なかれ不定であることを確認し、現時点で私たちの研究方向にとって重要となる側面に焦点を当て、紹介する。

5.1.3 意味フレームは(状況)理解の単位である

フレーム F を限定的に定義した場合、 F は状況の理想化で、典型的には $\langle\langle\text{何が}\rangle\langle\text{いつ}\rangle\langle\text{どこで}\rangle\langle\text{何のために}\rangle\langle\text{何を}\rangle\langle\text{どうする}\rangle\rangle$ のような形で表現できる。この意味でのフレームは、外延的には自然言語処理で(表層)格フレーム ((surface) case frames) [46, 47, 48] と呼ばれているものと実質的に同一だと考えられる。この予想は [49] で部分的に確かめられている。

状況の理想化としてのフレームは、ヒトが区別可能な状況の一つ一つをコードしている非言語的な単位²⁴⁾で、有限個しか存在しないが、その数は少ないとは言えない(人手コーディングの経験に基づいて、少な目に見積もっても意味フレームの数は、文化ごとに数千から数万はあると推測される)。この集合がヒトが理解できる状況の全体を定義すると考えられる。

5.1.4 意味フレームは意味役割の源泉となる

同一のモノ (e.g., “本”) は(そのアフォーダンス [51])に基づいて、異なる状況 $S_1, \dots, S_n$ で異なる現われ $R_1, \dots, R_n$ (e.g., $\langle\text{出版物}\rangle, \langle\text{内容}\rangle, \langle\text{表現手段}\rangle, \dots$) をもつ。状況 F での x の現われ $F.R(x)$ が x の F での意味役割である。この意味での意味役割を BFN はフレーム要素 (Frame Elements: FEs) と呼ぶ [40]。

5.2 FOCAL 流のフレーム意味論の拡張と制約

私たちは以上のような BFN の洞察を取り入れつつ、BFS を認知科学的観点から独自の拡張と制約を加えた。この枠組みが FOCAL であり、具体的には以下のように仮定する²⁵⁾

²⁴⁾[50] によって部分的に実在性が確認されている。

²⁵⁾基本文献 [16, 39, 40] では重要な概念に正確な定義がないことが多いため、どれが BFN にない、私たち独自の拡張なの

5.2.1 勝ち残り方式のフレームの特定

BFN/BFS は、人が語の意味を理解するときに、具体的にどういう処理過程を通じてフレームが特定しているのか考察していない。これは認知科学的な観点からすると不満が残る。この問題を解消するために、私たちは次のように仮定する:

語の意義の曖昧性が解消されるとは、ほかの語群との共起によって意味フレームが特定され、フレーム内でのその語の意味役割が定まることである。この際、複数のフレームの競合が起こり、勝ち残ったフレームが解釈を決定する。これは文脈(効果)の明示的モデル化である。

その品詞に係わらず、語は様々なフレームを喚起するが、喚起の仕方は語種によって異なる。動詞はフレームの特定に大きく貢献し、この性質故に BFN で支配項 (governor) と呼ばれ、重要視される [52] が、FOCAL では文脈効果を捉えるため、次の仮定を置く:

(26) どんな語も単独では意味フレームを特定する力はない。

(27) 語は幾つもの異なるフレームを同時に喚起する場合があり、この場合、両立しないフレーム群が競合する。この競合は「最適者の勝ち残り」方式で解消される。

これはつまり、フレームの特定は、並列的、分散的だということである。この意味で、FOCAL は—語の意味という“比較的マクロ”なレベルで—並列分散意味論 (Parallel Distributed Semantics) を指向する枠組みである。

5.2.2 意味フレームの基本レベル存在仮説

私たちは、意味フレームのネットワーク構造には「基本レベル」と、その「上位、下位レベル」の区別があると考えた。これを(意味フレームに関する)基本レベルの存在仮説と呼ぶ。これは概念の階層に基本レベルの概念 (e.g., イヌ)、その上位概念 {動物、かを正確に言うのは難しいと断っておきたい。

ペット, ... }, 下位概念 (e.g., { 柴犬, チワワ, ... }) があるのと同じである。

私たちの「襲う」の分析ではフレームの基本レベルを特定するには到っていないが, 具体性のレベルの違いは意味フレームの階層化に反映されている。[53, 54]。概念の基本レベルとは, 概念階層の他の階層よりも認知的に優位な階層のことである。同様の階層性は, 対象概念だけでなく, 行為や事象の概念にも見られると指摘されている [55, 56]。

5.2.3 フレームの具体性と理解の深さは相関する

基本レベルのフレームの存在から基本レベルの理解の存在が予想される。語群に最適な意味フレームが決定されたとき, あるレベルの理解が達成されるが, 理解のレベルはフレームの具体性の度合いによって決まる。フレームの抽象性が低ければ「深い理解」が生じ, 抽象性が高ければ「浅い理解」が生じる。

5.3 意味フレーム解析の作業手順

ここではコーパス基盤の階層的フレーム網 (Hierarchical Frame Network: HFN) の構築法を論じる。実例として“ x が y を襲う”の意味フレームの階層的ネットワークを特定するための作業手順を説明する。

5.3.1 コーパス事例の収集とデータ編集: 用法基盤主義の実践

日英対訳コーパス²⁶⁾から KWIC (Key Word In Context) ツールを使って「襲 { わ, い, う, え, お, って }」の全用例を収集した。これは参照したコーパスで観察可能な「襲う」の全用例の分析である²⁷⁾。これに BFN の手法 [52] にならって意味役割を割り当

²⁶⁾<http://www2.nict.go.jp/x/x161/members/mutiyama/align/index.html> で公開。

²⁷⁾この際, 用例が比喩的であるか否かの区別は, 意図的に行わなかった。比喩的な表現と非比喩的な表現の区別は, データの解析自体から現われてくるべきもので, 分析者が恣意的に導入すべきものではないと考えたからである。実際, 比喩と非比喩の区別を分析の前提としないフレーム基盤アプローチが, その区別を前提とする比喩写像理論 [5] 同等以上に比喩性の問題をうまく扱えるのを示すことが, [57] の論点であった。

て。データベース化した。最終的に得られた事例は 416 例であった。

次に KWIC 形式の言語データを加工ツールで編集する。この作業は Microsoft 社の Excel で可能であるほど容易である。Excel での作業環境を想定して話を進めると, 加工を始める前, 一つ一つの事例は (コーパスの名前, 文のコーパス内での事例の生起位置を示す情報を除けば) (i) L(ef t Context), (ii) K(eyword), (iii) R(ight Context) の三列のみからなっている。人手で意味フレームを特定しコーディングするためには, 次の [a; b1, 2, 3; c1,2,3; d] の情報を指定する必要がある:

- (28) (a) 文 S が実現しているフレーム名; (b1) S の主語句 s と (b2) s の意味型, 並びに (b3) s の意味役割 (= FE); (c1) S の目的語句 o と (c2) o の意味型, 並びに (c3) o の意味役割; (d) S の意味フレーム。

表 1 に「襲う」のコーディング例をあげる。

BFN が企画された理由の一つはそういう作業の資源としての意味フレームのデータベースを提供することにあった。実際, 意味フレームのデータベース $D(F)$ (と注釈支援ツール) が利用可能な状態であれば, 表 1 に示したようなコーディング作業は比較的容易だと思われる。

5.3.2 “ x が y を襲う”のフレームの階層的ネットワーク

このデータの意味素性表示に基づいて, 意味フレームの階層的ネットワークが同定された。このような構造のことを階層フレームネットワーク (Hierarchical Frame Network: HFN) と言う。「襲う」の HFN の最終形を図 4 に示した (階層が上のフレームは抽象的, 下のフレームは具体的なフレームである)。実例を除く最下位ノードは (29) の 15 個のフレームに相当する:

- (29) **F01:** 武力抗争; **F02:** 軍事侵略; **F03:** (強盗などの) 資源強奪; **F04:** 強姦; **F05:** 虐待; **F06:** 動物の攻撃 (捕食系); **F07:** 動物の攻撃 (非捕食系);

Table 1: S1, S2 の主語句, 目的語句の意味型, 意味役割, フレーム名のコーディング例

文 ID	L(K)	K	R(K)	Sの主語句の文字列	Sの主語句の意味タイプ名	Sの主語句のFE名	Sの目的語句の文字列	Sの目的語句の意味タイプ名	Sの目的語句のFE名	Sの実現しているFrame名
S1	二人組が銀行を	襲った	。	二人組	人間 [+grouped]	強盗	銀行	施設 OR 機関	金融機関	強盗
S2	サメが傷ついたイルカを	襲った	。	サメ	肉食魚類 [+animate]	捕食動物	(傷ついた)イルカ	哺乳動物 [+animate]	獲物	捕食

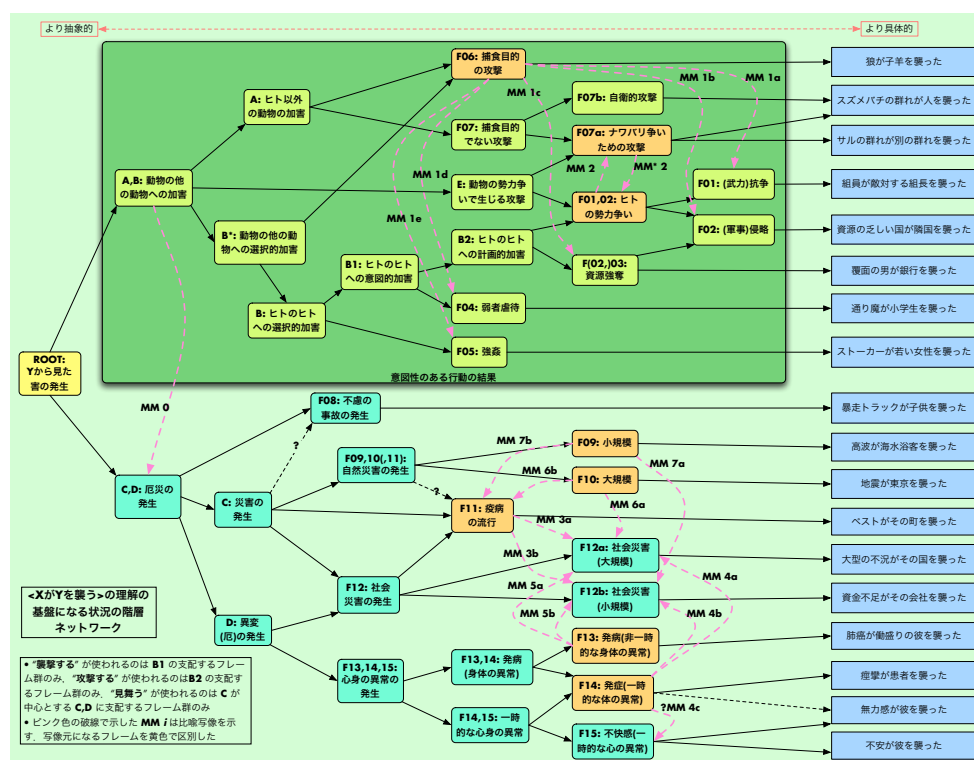


Figure 4: 「襲う」の HFNA: 矢印は (部分的) 具体化 (elaboration) の関係を表す

F08: 人為災害の発生; F09: (高波などの) 小規模な自然災害; F10: (地震などの) 大規模な自然災害; F11: 疫病の流行; F12: 活動への打撃の発生; F13: 発病 (非一時的な心身の異常); F14: 発症 (一時的な体の異常); F15: 悪感情の発生 (一時的な心の異常)

これらの意味フレームは「襲う」という語の選択制限に反映される限り, なるべく細かく区別した. 例えば「{地震, 台風, ...} が太郎 [+human, -grouped] を襲った」が奇妙であるのに対し「{高波, 突風, ...} が太郎 [+human, -grouped] を襲った」は自然である. これは F08, F09 の区別の根拠となる.

5.3.3 HFN の上位ノード

図 5 には「襲う」の意味フレームのネットワーク (つまり「襲う」の HFN) の上部構造のみを示した. HFN のルートは「被害の発生」で, それを除く大規模の意味フレーム群が 9 個存在することを表している. 図 5 の B, E, I, J は図 4 の A, B, C, D に対応する²⁸⁾.

私たちは, ネットワーク構造の成形に素性表示を用いた. 実際, 図 5 にあるように「襲う」の FN の上位 10 つのノード (上位フレーム群: Root, A, ...,

²⁸⁾図 5 の区別は素性ラティスによる論理的な区別で, 図 4 の区別より細かく, すべてに認知的な意味があるとはいえない.

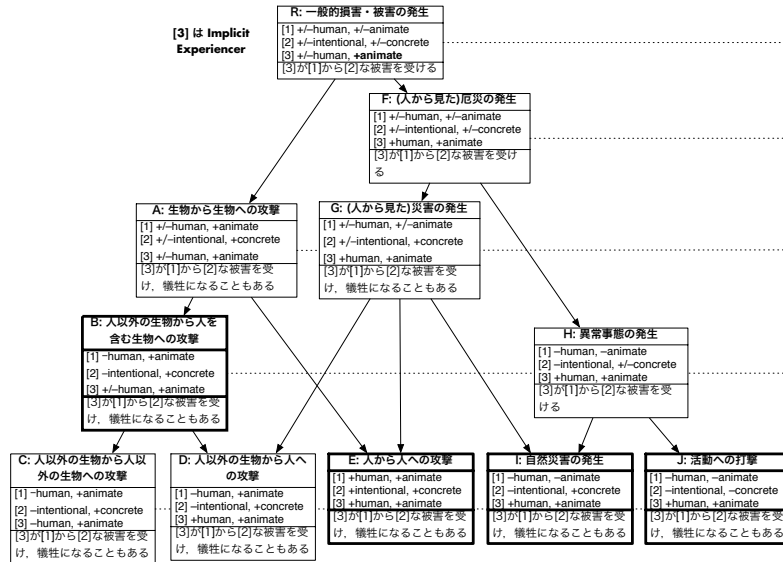


Figure 5: 「襲う」の HFN (上位 10 ノードのみ)

J) は、次のような意味素性の組み合わせによって記述されている:

- (30) a. 主語名詞の指示対象 s について: $[\pm\text{human}(s)], [\pm\text{animate}(s)], [\pm\text{intentional}(s)]$;
 b. 効果 e について: $[\pm\text{concrete}(e)]$;
 c. 目的語の指示対象 o について: $[\pm\text{animate}(o)], [\pm\text{human}(o)]$

心理実験との整合性を保証するため、素性 [F] の値は連続で区間 $[0, 1]$ のどんな値でも取りうるとする。この解釈の下では、 $\{0, 1\}$ は妥当性の評価関数の最小値と最大値であり、値が $1/2$ の場合、真偽性は「中和されている」と考える。これは $\{-1, 0, +1\}$ の二極分化体系を考えるのと実質的には同じことである。これは古典的な素性理論の自然な拡張で、これによってプロトタイプ効果も正しく捉えることができる。

6 意味フレームによる言語現象の説明

この節では、以上の説明の下で、幾つかの言語現象を意味フレームの直接、間接の効果として記述することを試みる。

6.1 選択制限は素性共起制限に由来する

この節での議論の実験的証拠は [58] に提供されている。

$[F(x)]$ を $x (x = \{\text{主語名詞 } s, \text{ 目的語名詞 } o\})$ についての素性 [F] の妥当性だとすると「襲う」のネットワークでは例えば、次の素性の共起制限が成立している:

- (31) a. $*[\dots, +\text{intentional}(x), \dots, -\text{animate}(x), \dots]$ が含意関係 R_1 から帰結

$$R_1: [+ \text{intentional}(x)] \rightarrow [+ \text{animate}(x)]$$

- b. $*[\dots, +\text{human}(x), \dots, -\text{animate}(x), \dots]$ が含意関係 R_2 から帰結

$$R_2: [+ \text{human}(x)] \rightarrow [+ \text{animate}(x)]$$

- c. $*[\dots, +\text{human}(x), \dots, -\text{intentional}(x), \dots]$ が含意関係 R_3 から帰結

$$R_3: [+ \text{intentional}(x)] \rightarrow [+ \text{human}(x)]$$

- d. $*[\dots, +\text{animate}(s), \dots, -\text{animate}(o), \dots]$ が含意関係 R_4 から帰結

$$R_4: [+ \text{animate}(s)] \rightarrow [+ \text{animate}(o)]$$

ここで、含意関係 $R: [+F] \rightarrow [+G]$ は「F の値が + ならば、G の値も + である (が、その逆は必ずしも真ではない)」ことを表わす。これにより、例え

ば R_2 が $[\dots, +human(x), \dots, -animate(x), \dots]$ という素性共起が不可能である理由になる．これに対し $[\dots, -human(x), \dots, +animate(x), \dots]$ には問題がない．

6.1.1 素性の共起制限には身体性が反映する

素性に共起制限があるということは、素性の値が相互に独立していないということである．つまり自由度が押さえ込まれている．これは行動制御で有名な Bernstein 問題 (Bernstein Problem) [59, 60] の意味論版の自然な解消であると見なすことができる．

これは近年重要視されている「身体性」の問題に直結する．「意味が身体化されている」という昨今の主張の正確な理解は、結局は「身体が素性の共起制限の発生源となって、恣意性を押さえこんでいる」ということであろう．これが含意するのは「意味が生じるのは、私たちが身体をもっているからだ」ということである．

6.1.2 意味フレームへの解釈の“引き込み”

(31a) は含意関係 R_1 に起因する個体属性を記述し、(31d) は含意関係 R_2 に起因する「生物 s が襲うのは生物 o である」という生体間 (s, o) の相互作用を記述している．これは意味フレームの存在に起因する、解釈上のフレームへの引き込み効果 (attraction-to-frame effect)、あるいは単にフレーム効果の一例である．

6.1.3 選択制限の由来

明らかに、選択制限 (selectional restrictions) はフレーム効果の一例である．「襲う」に限らず、一般に他動詞の s, o の選択制限が独立していないのは、フレーム効果が存在するからである．(32c, d) が奇妙なのは、フレーム効果の一例である．

- (32) a. 強盗が銀行を襲った (F03)
 b. 通り魔が小学生を襲った (F05)
 c. ???通り魔が銀行を襲った (F03, F05 Mix)

- d. ??強盗が小学生を襲った (F03, F05 Mix)

外界を情報状態 I として一般化すれば、属性は I の個体情報を、フレームは I の環境情報をコードしていると言える．従って、「フレームとは何か?」「認知モデルとは何か?」のような根源的な問いの答えは「ヒトの(脳内で)環境情報をコードする構造体」となる．

フレーム効果は例えば、(33) のような文の解釈で発生するメトニミー効果もうまく説明する:

- (33) 疫病がその { 地方, 動物園, ... } を襲った

(33) での概念〈地方〉や〈動物園〉は $[+animate]$ の素性をもたず、図 5 の R の $[+animate(o)]$ の指定に反しているが、(33) で理解される内容は「地方」や「動物園」という(名の)抽象的な対象が襲われたということではなくて、(34)に近い．

- (34) 疫病がその { 地方に住む人々, 動物園で飼育されている動物, ... } を襲った

これは、襲われるものが $[+animate]$ なものとして再構成されるということである．

(33) $\Rightarrow$ (34) の“再解釈”の効果をメトニミーと言うのは簡単であるが、そう名づけること、そう分類することは何かの説明であるわけではない．

だが、これをフレーム効果の一例だと見なせば、(33) の目的語の箇所にメトニミーが現れる必然性が説明される．

6.2 メトニミーの一部は解釈の状況への引きこみ効果である

この節での提案する意味解釈の状況への引きこみ効果の実験的証拠の一部は [61] で提示されている．

6.2.1 道具主語文に道具使用者を読み込む

道具主語の文 (35) の主語名詞「凶刃」にも同様のフレーム効果を認めることが可能である．

- (35) 凶刃が通行人を襲った

これに係わっているのは〈道具の使用〉フレームである．道具名の言及によって道具の使用者が喚起されるのは，何かが〈道具〉であることは必然的にその〈使用者〉がいるというオントロジー上の前提あるからである．そのような知識は使用という状況を認定することで効率的に記述できるはずである．

6.2.2 制御不能になった道具は自律性が読み込まれる

これに対し，相対的に自律性の高く，それ故に道具性の低い[(暴走)車]が主語となった(36)の場合，(37)にあるように暴走車が運転者の身体の延長，つまり道具として解釈できると同時に，何らかの理由(例えばブレーキの故障など)で運転手の制御から離れた車自体に(創発的な)[+animate(s)]を認める解釈も可能である．

(36) 暴走車が通行人を襲った

(37) 運転手が{運転ミスで, 不注意で, 錯乱して}, 運転していた車で通行人を殺傷した

この違いは〈加害〉フレームの要求である[+animate(s)]という特性を，どのレベルの存在体(entity)に帰着するかの差によって説明可能である．

7 おわりに

どれほどうまく行ったかは別にして，私たちが試みたのは，認知言語学の「基盤」の見直しである．

A Berkeley FrameNet の実例

BFN で(加熱)調理(cook)フレーム²⁹⁾を見ると，(38)のような定義が与えられている³⁰⁾

(38) Cooking_creation

Definition: This frame describes food and meal preparation. A COOK creates a PRODUCED FOOD from (raw) INGREDIENTS. The

²⁹⁾フレームには文化差が反映する．英語の〈COOKING〉では加熱が前提となるが，日本語の〈料理〉，〈調理〉はちがう．

³⁰⁾これは FrameSQL 2.1 の定義である．

HEATING INSTRUMENT and/or the CONTAINER may also be specified: [Cook Caitlin] [Gov baked] [Food some cookies] [Ingr from the pre-packaged dough]³¹⁾

FEs:

Core:

- a. CONTAINER [CONTAINER]: This FE identifies the CONTAINER that holds the food being produced.
 - b. COOK [COOK]: The COOK prepares the PRODUCED FOOD.
 - c. HEATING INSTRUMENT [HEAT INSTR]: This FE identifies the HEATING INSTRUMENT with which the COOK prepares the FOOD.
 - d. INGREDIENTS [INGR]: The INGREDIENTS which are altered by the COOK alters INGREDIENTS to create the PRODUCED FOOD.
 - e. PRODUCED FOOD [FOOD]:
 - f. RECIPIENT [REC]: RECIPIENT identifies the person for whom the food has been prepared.
- (39) Relations: { Inherits From: —, Is Inherited By: —, Subframe of: —, Has Subframes: —, Uses: Apply_heat, Is Used By: —, Is Inchoative of: —, Is Causative of: —, See Also: Apply_heat }
- (40) Lexical Units: *bake.v, concoct.v, cook.v, cook up.v, make.v, prepare.v, put together.v, whip up.v*
Created by cota on 2002:02-12 15:25:43.0

Lexical Units (LU's) とはフレームを実現する語彙項目である．

BFN のデータベースでは，ここにあるようなフレーム群が(39)にあるような継承や合成などの関係によって結びつけられ，構造化されている[43]．

References

- [1] 黒田 航, 中本 敬子, 金丸 敏幸, 龍岡 昌弘, 野澤 元. フレーム指向概念分析 (FOCAL) の目標と手法: Berkeley FrameNet を超えて. [未発表論文: <http://clsl.hi.h.kyoto-u.ac.jp/~kkuroda/papers/focal-manifesto.pdf>], 2004.
- [2] 黒田 航, 中本 敬子, 金丸 敏幸, 龍岡 昌弘, 野澤 元. 「意味フレーム」に基づく概念分析の射程: Berkeley FrameNet and Beyond. 日本認知言語学会第5回大会 Conference Handbook, pp. 133–153. 日本認知言語学会 (JCLA), 2004.

³¹⁾Web インターフェイスでは FE は色分けされている．

- [3] 黒田 航, 中本 敬子, 野澤 元. 意味フレームに基づく概念分析の理論と実践. 正明山梨, 他(編), 認知言語学論考第4巻, pp. 133–269. ひつじ書房, 2005. [増補改訂版: <http://cls1.hi.h.kyoto-u.ac.jp/~kkuroda/papers/roles-and-frames.pdf>].
- [4] R. W. Langacker. *Foundations of Cognitive Grammar, Vol. 2: Descriptive Applications*. Stanford University Press, 1991.
- [5] G. Lakoff and M. Johnson. *The Philosophy in the Flesh*. Basic Books, 1999.
- [6] G. Lakoff. *Women, Fire, and Dangerous Things*. University of Chicago Press, 1987. [邦訳: 『認知意味論』(池上 嘉彦・河上 誓作 訳). 紀伊国屋書店.].
- [7] 道又爾, 北崎充晃, 大久保街亜, 今井久登, 山川恵子, 黒沢学. 認知心理学: 知のアーキテクチャを探る. 有斐閣, 2003.
- [8] J. Pustejovsky. The generative lexicon. *Computational Linguistics*, Vol. 17, No. 4, pp. 409–440, 1991.
- [9] J. Pustejovsky. *The Generative Lexicon*. MIT Press, 1995.
- [10] James Pustejovsky and B. Boguraev. A richer characterization of dictionary entries: The role of knowledge representation. In B. T. S. Atkins and A. Zampoli, editors, *Computational Approaches to the Lexicon*, pp. 295–311. Oxford University Press, Oxford, 1994.
- [11] James Pustejovsky and P. Bouillon. Aspectual coercion and logical polysemy. In James Pustejovsky and B. Boguraev, editors, *Lexical Semantics: The Problem of Polysemy*, pp. 133–162. Oxford University Press, Oxford, UK, 1996.
- [12] Y. Ravin and C. Leacock. Polysemy: Overview. In Y. Ravin and C. Leacock, editors, *Polysemy: Theoretical and Computational Approaches*, pp. 1–29. Oxford University Press, Oxford, UK, 1999.
- [13] D. Sperber and D. Wilson. *Relevance: Communication and Cognition*. Blackwell, 2nd edition, 1995.
- [14] R. W. Langacker. *Foundations of Cognitive Grammar, Vol. 1: Theoretical Prerequisites*. Stanford University Press, 1987.
- [15] Charles J. Fillmore. Frame semantics. In Linguistic Society of Korea, editor, *Linguistics in the Morning Calm*, pp. 111–137. Hanshin Publishing, Seoul, 1982.
- [16] Charles J. Fillmore. Frames and the semantics of understanding. *Quaderni di Semantica*, Vol. 6, No. 2, pp. 222–254, 1985.
- [17] Charles J. Fillmore and B. T. S. Atkins. Towards a frame-based lexicon: The semantics of RISK and its neighbors. In A. Lehrer and Eva F. Kittay, editors, *Frames, Fields and Contrasts: New Essays in Semantic and Lexical Organization*, pp. 75–102. Lawrence Erlbaum Associates, 1992.
- [18] Charles J. Fillmore and B. T. S. Atkins. Starting where the dictionaries stop: The challenge for computational lexicography. In B. T. S. Atkins and A. Zampoli, editors, *Computational Approaches to the Lexicon*, pp. 349–393. Clarendon Press, Oxford, UK, 1994.
- [19] J. Grady. Theories are buildings revisited. *Cognitive Linguistics*, Vol. 8, No. 4, pp. 267–290, 1997.
- [20] M. Johnson. *Body in the Mind*. University of Chicago Press, 1987.
- [21] G. Lakoff and M. Johnson. *Metaphors We Live By*. University of Chicago Press, 1980. [邦訳: 『レトリックと人生』(渡部昇一ほか 訳). 大修館.].
- [22] S. Coulson. *Semantic Leaps: Frame-Shifting and Conceptual Blending in Meaning Construction*. Cambridge University Press, 2001.
- [23] Gilles R. Fauconnier. *Mental Spaces: Aspects of Meaning Construction in Natural Language*. Cambridge, MA: MIT Press, 1985.
- [24] Gilles R. Fauconnier. *Mappings in Thought and Language*. Cambridge, MA: Cambridge University Press, 1997.
- [25] Gilles R. Fauconnier and Mark Turner. Conceptual projections and middle spaces. Cognitive Science Technical Report (TR-9401), Cognitive Science Department, UCSD, 1994.
- [26] Gilles R. Fauconnier and Mark Turner. Blending as a central process of grammar. In Adele G. Goldberg, editor, *Conceptual Structure, Discourse, and Language*. CSLI Publications, 1996.
- [27] Gilles R. Fauconnier and M. Turner. Conceptual integration networks. *Cognitive Science*, Vol. 22, pp. 133–187, 1998.
- [28] Mark Turner. *Cognitive Dimensions of Social Sciences*. Oxford University Press, Oxford, 2001.
- [29] R. W. Langacker. *Grammar and Conceptualization*. Mouton de Gruyter, 2000.
- [30] J. Rubba. *Discontinuous Morphology in Modern Aramic*. PhD thesis, Department of Linguistics, University of California, San Diego, 1993.
- [31] W. Croft. *Radical Construction Grammar*. Oxford University Press, Oxford, 2000.
- [32] W. Croft. Lexical rules vs. Constructions: A false dichotomy. In H. Cuykens, Th. Berg, R. Dirven, and K.-U. Panther, editors, *Motivation in Language: Studies in Honor of Günter Radden*. John Benjamins, Amsterdam, 2003.
- [33] A. D. Goldberg. *Constructions: A Construction Grammar Approach to Argument Structure*. University of Chicago Press, Chicago, IL, 1995.
- [34] Christiane Fellbaum, editor. *WordNet: An Electronic Lexical Database*. MIT Press, 1998.

- [35] P. Coad and E. Yourdon. *Objected-Oriented Analysis*. Prentice-Hall, 2nd edition, 1990.
- [36] P. Coad and E. Yourdon. *Objected-Oriented Design*. Prentice-Hall, 1991.
- [37] 長谷部 陽一郎. コンピュータ・アナロジー再考: 認知言語学とオブジェクト指向の観点から. 日本認知言語学会第5回記念大会 Conference Handbook, pp. 53–56. 日本認知言語学会 (JCLA), 2004.
- [38] NTT コミュニケーション科学研究所 (監修). 日本語彙大系. 東京: 岩波書店, 1997.
- [39] Charles J. Fillmore, C. Wooters, and C. F. Baker. Building a large lexical databank which provides deep semantics. In *Proceedings of the Pacific Asian Conference on Language, Information and Computation, Hong Kong*, 2001.
- [40] Charles J. Fillmore, C. R. Johnson, and M. R. L. Petruck. Background to FrameNet. *International Journal of Lexicography*, Vol. 16, No. 3, pp. 235–250, 2003.
- [41] 黒田 航, 中本 敬子, 金丸 敏幸, 龍岡 昌弘, 野澤 元. 「意味フレーム」に基づく概念分析の射程: Berkeley FrameNet and Beyond. 日本認知言語学会第5回大会 Conference Handbook, pp. 133–153. 日本認知言語学会 (JCLA), 2004.
- [42] 黒田 航. 概念メタファーの体系性, 生産性はどの程度か? 日本語学, Vol. 24, No. 6, pp. 38–57, 2005.
- [43] C. F. Baker, Charles J. Fillmore, and B. Cronin. The structure of FrameNet database. *International Journal of Lexicography*, Vol. 16, No. 3, pp. 281–295, 2003.
- [44] K. H. Ohara, Seiko Fujii, Hiroaki Sato, Shun Ishizaki, Toshio Ohori, and Ryoko Suzuki. The Japanese FrameNet project: A preliminary report. In *Proceedings of PACLING '03*, pp. 249–254, 2003.
- [45] 中本 敬子, 黒田 航, 野澤 元, 金丸 敏幸, 龍岡 昌弘. FOCAL/PDS 入門: フレーム指向概念分析/並列分散意味論の具体的紹介. [未発表論文: <http://clsl.hi.h.kyoto-u.ac.jp/~kkuroda/papers/introduction-to-focal.pdf>], 2004.
- [46] 河原 大輔, 黒橋 禎夫. 用言と直前の格要素の組を単位とする格フレームの自動獲得. 自然言語処理, Vol. 9, No. 1, pp. 1–16, 2002.
- [47] 河原 大輔, 黒橋 禎夫. 格フレーム辞書の漸次的自動構築. 自然言語処理, Vol. 12, No. 2, pp. 109–131, 2005.
- [48] 荻野 孝野, 小林正博, 井佐原 均. 日本語動詞の結合価. 東京: 三省堂, 2003.
- [49] 中本 敬子, 黒田 航. 「逃れる」の階層的意味フレーム分析とその意義: 「言語学・心理学からの理論的, 実証的裏づけ」のある言語資源開発の可能性. 言語処理学会第12回大会発表論文集, pp. 592–595, 2006. 発表 P4-1.
- [50] 中本 敬子, 野澤 元, 黒田 航. 動詞「襲う」の多義性: カード分類課題と意味素性評価課題による検討. 認知心理学会第二回大会口頭発表論文集, p. 38, 2004. [<http://clsl.hi.h.kyoto-u.ac.jp/~kkuroda/papers/Nakamoto-et-al-CogPsy2004-Original.pdf>].
- [51] 佐々木正人. アフォーダンス: 新しい認知の理論. 岩波科学ライブラリー, 1994.
- [52] B. T. S. Atkins, M. Rundell, and H. Sato. The contribution of FrameNet to practical lexicography. *International Journal of Lexicography*, Vol. 16, No. 3, pp. 333–357, 2003.
- [53] A. Pansky and A. Koriati. The basic level convergence effect in memory distortions. *Psychological Science*, Vol. 15, pp. 52–59, 2004.
- [54] E. Rosch, C. B. Mervis, W. Gray, D. Johnson, and P. Boyes-Braem. Basic objects in natural categories. *Cognitive Psychology*, Vol. 8, pp. 382–439, 1976.
- [55] M. W. Morris and Gregory L. Murphy. Converging operations on a basic level in event taxonomies. *Memory and Cognition*, Vol. 18, pp. 407–418, 1990.
- [56] J. M. Zacks and B. Tversky. Event structure in perception and conception. *Psychological Bulletin*, Vol. 127, pp. 3–21, 2001.
- [57] 黒田 航, 野澤 元. 比喩理解におけるフレーム的知識の重要性: FrameNet との接点. [COE21 ワークショップ「メタファーへの認知的アプローチ」のための研究論文 <http://clsl.hi.h.kyoto-u.ac.jp/~kkuroda/papers/metaphor-and-frames.pdf>], 2004.
- [58] 中本 敬子, 黒田 航. 意味フレームに基づく選択制限の表現: 動詞「襲う」を例にした心理実験による検討. 言語科学会第7回大会ハンドブック, pp. 75–78, 2005. (June 25–26, 2005, 上智大学).
- [59] N. A. Bernstein. *The Coordination and Regulation of Movements*. Oxford University Press, 1967.
- [60] N. A. Bernstein. *On Dexterity and its Development*. Lawrence Earlbaum, 1996. edited by M. Turvey, translated from Russian by M. L. Latash. [邦訳: 『デクステリィー: 巧みさとその発達』. 工藤和俊 (訳). 佐々木正人 (監訳). 金子書房. 2003.].
- [61] 黒田 航, 中本 敬子, 野澤 元. 状況理解の単位としての意味フレームの実在性に関する研究. 日本認知科学会 第21回大会 発表論文集, pp. 190–191, 2004.