

WWWを丸ごとデータにすると

何がわかるか？

—— 格フレーム辞書を言語研究に利用する ——

特集
WWWを対象にした
日本語研究

黒田 航・李在鎬

1 はじめに

自然言語処理（NLP）の技術の進歩により、Webデータをコーパスとして扱う可能性が現実化し、その流れは一層堅固になって来ている。Webデータをコーパスとして使うことへの批判もないわけではないが、使い方次第では利点の方が多いように思う。(1)にWebコーパスの長所を、(2)に短所を幾つか挙げた（注1）。

- (1) a 電子化（e.g. 「書き起こし」）の手間が不要。
- b データの規模が（おそらく十分に）大きい。
- c 内容／用例が（おそらく十分に）多様である。
- (2) a 量と内容が共に常に変化し、動的である。

b 「ノイズ」が多い。
c 代表性に欠ける。

(1) bの利点と(1) cの利点は相関している（これは(1) cが(1) bを前提にするためである）。これは既存のコーパス（例えば新聞コーパス）に較べて明らかに優位である。Webコーパスとの比較を通じて新聞記事に使われる言語は文体の面（伊藤（文献5）でも語彙や語義の分布の面（竹内（文献10））でも偏っているという事実が明らかになりつつある。

当然、Webデータに欠点はある。(2) aが理由で、例えば正確な統計量（語数）を測定できない。(2) bの問題は、一定の割合で誤用、逸脱用法（e.g. 2chの言語）、道徳的

WWWを丸ごとデータにするとは何がわかるか？

に好ましくない内容を表わした文章が含まれる点にある（これらを自動的に除外することは難しい）。(2)cの問題は、Webコーパスは新聞コーパスに比べ比較的多様なデータを含んでいるとは言え（話しコトバを含まないの）均衡コーパスではないし、コンピュータで文章を書く習慣をもつ人々の言語実態しか対象にならない。

(1)の利点と(2)の弱点は表裏一体であり、(2)の弱点を容認することなしに(1)の利点を手に入れることは原理的に不可能である。となれば、Webデータをコーパスとして利用することは是非は、研究目的に合っているか否かで評価するしかない。研究の目的が日本語の正確な統計的実態の調査であれば、(2)aは本質的な難点となるだろうが、用例研究とその先にある意味研究の見地からは(1)cの効果は絶大である。

以上の問題意識の下、本稿では次の二つのことをする。
①長谷部（文献22）によるWikipedia日本語版をコーパス化する試みを2節で紹介し、②それに続く3節以下で河原・黒橋（文献19、20）によって開発され、<http://nlp.kuekyoto-u.ac.jp/nl-resource/caseframe.html>と一般公開されている「Webから自動構築された格フレーム辞書」（以下、単に「格フレーム辞書」）を使った言語研究を紹介する。4節は李ほか（文献24）の格フレーム辞書の

元になったWebコーパスを使った実証研究の紹介、5節は黒田ほか（文献11）の現行の格フレーム辞書の精度評価の研究である。

2 Wikipedia日本語版のコーパス化

日本語には著作権で保護されていない規模の大きなコーパスがない。これは日本語のコーパス基盤言語学がなかなか進展しない原因の一つである。この現状を改善する試みの一つとして、長谷部（文献22）は、Wikipediaの日本語版が比較的規模の大きな、著作権上の制約のない日本語コーパスとしての利用を可能にするWP2TXTとMconcというツールを開発し、<http://yohasebe.com/>で公開している。

WP2TXTは、Wikipediaのオリジナルデータをb22を展開しつつ、xmlデータをスキャンしてtitle要素およびテキスト要素内のデータを抽出する。これは生データから不要なwikiタグとhtmlタグを除去しテキストに変換する処理である（一般に構造化されたWebデータの平テキストへの整形の自動化は重要な処理である）。MconcはWP2TXTで処理したテキストデータを形態素解析し、ユーザが指定する文字列や形態素のパターンを抽出する

特集 WWWを対象にした日本語研究

機能をもつ。形態素解析は MeCab (注2) を用いて行っている。具体的には次の処理を行う。まずテキストデータを読み込んで文の単位に分解し、形態素解析を行う。そして出力結果を設定ファイルにあらかじめ記述された正規表現パターンに照らし合わせ、データの取捨選択を行い、抽出された文字列と詳細な形態素情報とをCSV (comma separated value) ファイルとして出力する。

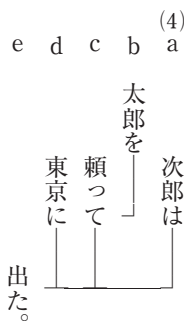
3 「格フレーム辞書」とは？

Web から自動構築した大規模格フレーム辞書とは、①5億文のWeb データの一文一文を形態素解析し、②それに「係り受け解析」と呼ばれる構文解析を行なって名詞句と述語の間の述語／項関係⇨依存関係を同定し、③こうして得られた項と述語の対の集合を、項の意味の類似度に従って意味グループ化したデータ (注3) である。この処理の際に具体的に行われることを(3)を例にして解説しておく。

(3) 次郎は太郎を頼って東京に出た。

(3) の文を形態素解析プログラム JUMAN (注4) で形態素解析し、その結果に係り受け解析プログラム KNP (注5) に渡して係り受け解析して得られる解析結果は

次の(4)である。



この出力は①「次郎は」は「出た」に係り、②「太郎を」は「頼って」に係り、③「頼って」が「出た」に係り、④「東京に」が「出た」に係っていることを表わしている。多くの例に対して同様の処理を行なうことで、S⇨「次郎は、出た」、「東京に、出た」(注6)、S⇨「太郎を、頼って」のような項と述語の対 [Arg, Pred] の集合の集合 S⇨ S⇨ S⇨ が得られる。S の部分集合のグループ化の結果が格フレーム辞書である。

[1] 格フレーム辞書でどんな検索ができるか

(5) a 基本的な検索…述語 (i.e., 動詞、形容詞、形容

動詞) を入力することで、それを主要部とする格フレームが一覧で表示される。それぞれの格フレームには格要素として現れる語の一覧が表示される (頻度2以上、上位20位まで)。

b 原文の参照…格フレームの検索結果で格要素

WWWを丸ごとデータにすると何がわかるか？

[2] 係り受け解析に関する幾つかの注意

- 係り受け解析は広い意味での依存(構造)解析(dependency (structure) parsing)の一種であり、その解析の出力となる係り受け関係の集合は基本的には依存関係⇨述語／項関係の集合であるが、次の点には注意が必要である(注7)。
- c 格要素との組み合わせでの検索…「裏目に出る」のように格要素と述語の組み合わせを(スペースを空けて)入力すれば、特定の格要素を含む格フレームが検索可能。
 - d 受身形／使役形での検索…検索語の後に「ing」をつけることで受身形の検索が、ingをつけることで使役形の検索が可能。
 - e 名詞の共起関係の検索…「名詞から検索」チェックボックスをオンにして名詞を項に取る述語が検索可能(表示されるのは頻度2以上現れる格フレームの一覧)。リンクをクリックすることで、その格フレームの詳細が表示される。

- (6) 助詞(例えばニやデ)の曖昧性は解消されていない(このため、述語の二格とデ格の分類精度はラ

に比べて非常に低い5[1]②で後述)。

- (7) 「次郎は、上京した」のような主題マーカーのハで格助詞が非明示化されている係り受け関係は、格フレーム辞書構築の候補集合から除かれている(ガ格の生起数は一般に期待より少なくなる傾向があり、§4で後述する事実の一部の[13]bに影響が出ていると思われる)(注8)。

- (8) 格フレーム辞書で名詞と述語との共起も調べることが出来るが(複合名詞の分析基準に問題がある)、検索には些か不合理な制約がある(例えば「犯人」や「防人」は検索できるが「料理人」や「使用人」は検索できない。同じ理由で「御者」や「患者」は検索できるが「逮捕者」や「追跡者」は検索できない)(注9)。

4 格フレーム辞書を使った連体節の使用実態調査

李ほか(文献24)は、日本語の連体節の実態調査を行うため、11個の述語(「恐れる」「積む」「冷やす」「寄与する」「固まる」「散る」「心掛ける」「賛成する」「逃げる」「逃れる」「躍進する」)の連体表現五七〇五例を格フレーム辞書の元コーパスから機械的に抽出し、人手で構造を分析

特集 WWWを対象にした日本語研究

した。これから Keenan & Comrie (文献3) や井上 (文献23) の提唱した関係節化の階層が、全体の傾向としては記述的に妥当であることが示唆されたが、その一方で語彙項目 (かそのタイプ) によって異なる分布を示していることも明らかになった。これは文法レベルでの一般化とは別に語彙項目の特性を踏まえた、構文的なレベルでの一般化も必要なことを示唆する結果である。

[1] 先行研究と問題の所在

従来、日本語の連体 (修飾) 表現 (広義の関係節) をめぐっては記述的・理論的観点から様々な考察がなされてきた (例えば久野 (文献4)、奥津 (文献7)、井上 (文献23)、寺村 (文献15)、三原 (文献9)、加藤 (文献16、17) など)。本節では、先行研究の指摘を踏まえながら日本語の連体 (修飾) 表現を記述する上で重要な視点について述べると同時に、本研究の問題提起を行う。

日本語の連体 (修飾) 表現の初めての包括的研究として寺村 (文献15) (元の論文は一九七五、一九七八) が挙げられる。その記述的・一般化の中心軸として(9) aの「内の関係」と(9) bの「外の関係」による対比はよく知られている。

- (9) a 秋刀魚を焼く男
b 秋刀魚を焼くにおい

(9) aは連体先の名詞が連体節内の用言に対して補語の関係にあったと考えられることから「内の関係」と呼び、一方の(9) bはそのような関係が成立しないことから「外の関係」と呼ばれている。

内の関係に関する先行研究の観察として、寺村 (文献14) では連体 (修飾) 節の可否をめぐり、表1の観察を示した (以下では特定の名詞を連体 (修飾) 節化することを装定化と呼び、その逆を述定化と呼ぶ)。

表1の観察をより精緻的に捉えた研究として Keenan & Comrie (文献3) と井上 (文献23) が挙げられる。前

表1 格関係に基づく装定化の可否

区分	装定化の可否
ガ格	可能
ヲ格	可能
ニ格	一部不可能 (原因はやや難, 変化の結果や判断基準は不可)
ヘ格	可能
デ格	ほとんど可能 (範囲 [クラスで一番背が高い] は不可)
カラ格	部分的に可能 (出発点, 原因, 起点は無理)
ト格	部分的に可能 (動詞の項であれば可能だが, 非項は無理)
ノ	可能

WWWを丸ごとデータにすると何がわかるか？

者は五〇の言語の類型的調査から装定化の容易さの順序を表わす階層として(10) aを提案している(注10)。後者は日本語の形態格の曖昧性を考慮した階層として(10) bを提案している。

(10) a 主語V直接目的語V間接目的語V前置詞の目的語V所有格V比較級の目的語

b 主格V直接目的格V間接目的格V位置格ニV位置格ヲV目標格ニまたはへV位置格デV助格デV基準格デV奪格V所有格V基点格V随格V理由格V比較格

(10)の階層性をめぐっては英語を中心とした一部の言語では使用頻度の調査(例えば Keenan(文献2)や Fox(文献1))が行われていたが、日本語に関しては丸元・乾(文献18)による内の関係に対する計量的調査はあるが、外の関係も含めた網羅的実態調査は行われていない。

寺村(文献14)では外の関係に関する分類を試み、構造的相違が明確なタイプとして(11)と(12)の対比を指摘している。

(11) a 選挙に出る考え

b 赤ちゃんが笑っている写真

(12) a タバコを買ったおつり

b 文子が座ったうしろ

(11)は被修飾名詞の内容を説明するものであることから「内容節」と呼ばれている。(12)は「逆補充」と呼ばれ、被修飾名詞が持つ相対的概念を利用した連体(修飾)表現の一つとされ、日本語特有の構造だと指摘されている。

これまでの日本語学・言語学の知見はいずれも日本語の構造を理解する上では重要であり、多くの研究成果が上がっているが、問題がないわけではない。このような問題が意識される一方で、個々の一般化が日本語の使用例に対してどの程度まで妥当かを評価した研究はない。その理由の一つは、評価のための有用な資料が存在しなかったからである。この難点は格フレーム辞書データを使うことで解消されると考えられる。

[2] データの収集法と解析法

李ほか(文献24)は次の方法でデータを収集した。① 格フレーム辞書の元コーパスから、11個の述語の連体(修飾)表現を抽出し、その中から「くしたAのB」のように係り先に曖昧性がある表現を除くことによって、分析対象とするデータを作成した(注11)。② 人手で修飾構造の特定を行った。分析の際には、先行研究の指摘に従って構造的タイプとして、内の関係、内容節、逆補充の3タイプに分類すると同時に、連体節内での格関係を

Category	①内的関係	②内容節	③逆補充	④内的関係-内容節	⑤逆補充-内容節
寄与する	188	322	15	35	22
恐れる	638	157	130	112	73
固まる	161	99	132	20	61
散る	178	18	49	16	28
賛成する	139	5	29	12	12
心掛ける	53	63	16	10	36
積む	865	195	153	42	107
逃げる	125	73	183	7	7
逃れる	158	19	266	49	19
躍進する	58	5	1	6	6
冷やす	295	174	75	47	73
合計	2858	1130	1048	231	538

[illegible]

特定した。本研究では、加藤（文献16、17）の示唆に従い、格関係間の排他性は仮定していない。

[3] 結果と考察

調査の結果、連体節のタイプとして図1の分布が得られた。全体の傾向としては、内の関係がもっとも多く、それに内容節と逆補充のタイプが続いている。その一方で非常に興味深いのは、語彙項目によって分布の内訳が異なっていることである。「寄与する」や「心がける」では内容節が主に用いられている。「逃げる」や「逃れる」では逆補充が主に用いられている。すべての語彙項目ごとに均質に分布しているわけではない（この分布の差は分散分析で有意差も出ているので、誤差ではない）。これは説明に値する事実の発見だと言える。

続いて内の関係における格関係の集計を行なった。格関係と動詞のクロス集計を行った結果を図2に示す。これから次の4点が示唆される。

(13) a 全体における使用頻度および生産性の面から、その傾向を捉えた場合、「ガ格Vヲ格Vニ格Vヘ格Vデ格Vニヨッテ格Vヲ通ジテ格」という傾向が観察される。

b 「積む」を除く他動詞「恐れる」「心掛ける」「冷

やす」の場合、ガ格よりヲ格の方がより生産的である。

c 「散る」「積む」「逃げる」のように移動を表す動詞群においては、ガ格に次いで生産的なものとしてニ格がある。

d 「固まる」「賛成する」「躍進する」のような変化を表す動詞群においては、9割近い用例でガ格が使用されている。

(13) bでは係り受け解析の段階で八句に隠れたガ格を取りこぼしていることが原因になっている可能性がある（注12）。

[4] まとめ

本研究の調査によって全体の傾向についての先行研究の指摘の妥当性が示唆される一方、語彙項目によって一般的な傾向から外れた分布が認められる可能性が示唆された。これはある集団の平均的な挙動とその部分の挙動は必ずしも同じにならないということ、統計解析では一般に「データの層別化の問題」として知られる問題の一種が生じている可能性が示唆される。これは文法レベルでの一般化とは別に語彙項目の特性を踏まえた、構文的なレベルでの一般化も必要なことを示唆する結果だと言

える。

本研究の妥当性で問題となるのは二点である。①11個の述語の選択基準が明確ではなく、得られた結果の代表性が明らかではない。②それに問題がないとしても人為的分類ミスが排除できていない。この種の問題が残っているものの、調査結果の妥当性以上に私たちが重要だと考えているのは、本研究がWebコーパスを基盤として構築された大規模格フレーム辞書が言語研究用の有益な研究資源になる可能性を提示している点である。

5 格フレームの精度の人手評価と改良のための提案

4節から格フレーム辞書の言語研究への利用可能性が示せたと思う。ただ、現時点で格フレーム辞書がどれほど信頼できる資源なのか不明である。黒田ほか(文献11)はこの問題意識から、格フレーム辞書の精度を人手で評価した。この節ではその結果を簡単に報告し、言語学の観点から精度向上のための提案を行なう。これは言語学からNLPへのフィードバックである。そうするのは、現行のデータベースの精度が不十分であるという事実の根本的原因の一部が、言語学がそれらについて正しい理

論を用意していないことが原因になっている可能性は高いと考えるからである。

[1] 評価の結果と個別の傾向の分析

① 格フレーム辞書が表現すべき情報の定義

格フレーム辞書は有用なデータベースだが、用法クラスターの一つ一つが表現している情報は何なのかは明確ではない。文献8の示唆に従うならば、うまく行った述語の用例クラスターとは、「C」を総合的に判断して、文献8の言う意味での特定の状況を表わすと判断できるクラスターである。

格フレーム辞書の用例クラスターCが理想的な状態にあるのは、Cの格要素のどれにも、Cが状況Sを表わすなら入っているべきではない要素の混じっていない状態のことである(どの述語を見ても格フレーム辞書の現状はこの理想からはほど遠い)。

述語の用例クラスターは格要素集合ごとに適当な得点を与え、その合計点から次のように分類した。極上品(√150)、美品(150√125)、美良品(125√100)、良品(100√70)、並品(70√50)、難あり(50√25)、評価不可(25√0)。

16個の述語の精度評価の結果を図3と図4に個数グラフと割合グラフにして示す(個数グラフは述語の使用頻度

WWWを丸ごとデータにすると何がわかるか？

図3 横軸 = 述語ごとのクラスターの個数 (品質別)

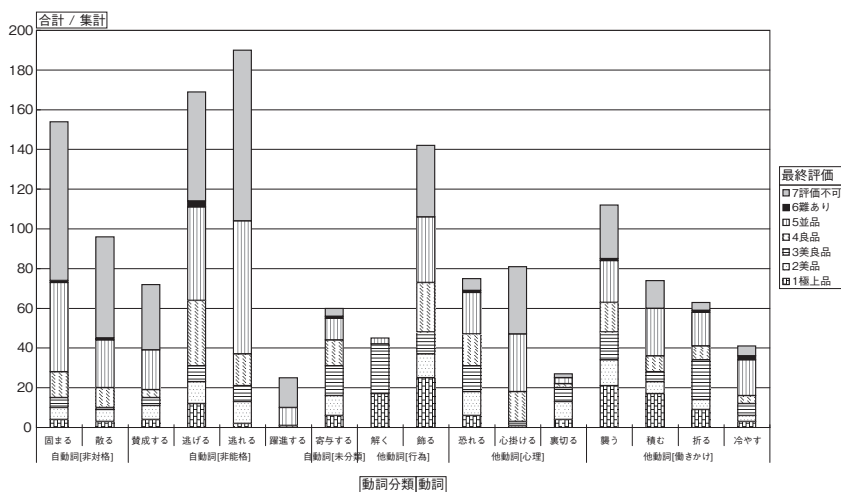
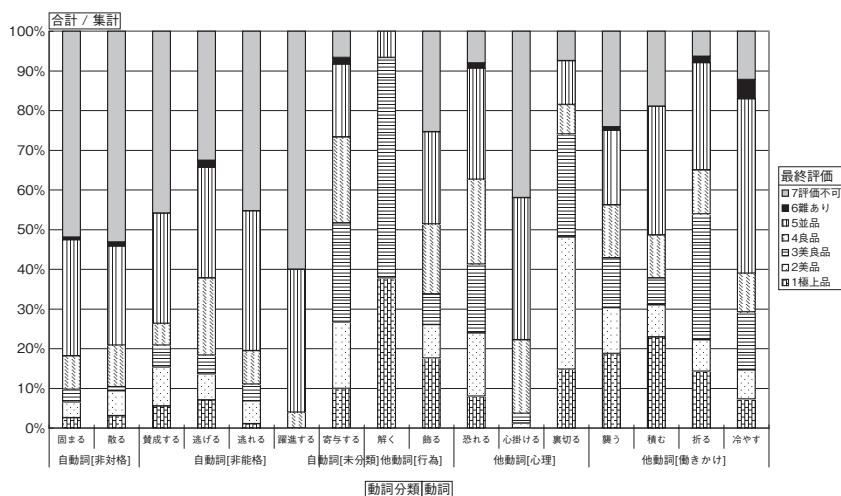


図4 横軸 = 述語ごとのクラスターの相対度数 (品質別)



を表わすわけではないことに注意)。

② 精度低下の主要因…二格、デ格の曖昧性

一見して述語ごとの精度のばらつきが大きい。細かく見てゆくと次のようなことがわかった。

(14) a 述語の頻度のクラスター精度への影響は顕著ではない。

b ヲ格が強い述語はうまく用例がクラスター化されているが、ヲ格を取らない述語のクラスター化精度が低い。

(14) a を考慮に入れて (14) b の結果を見ると、クラスター化精度はおそらく (格) 助詞の意味の広がりにより反比例していると言える (注 13)。曖昧性が著しいのは二格、デ格、ノ格であり、比較的に曖昧性が低いのがヲ格、ガ格である。ガ格は曖昧性は低い、おそらくハへの吸収率が高いために獲得される実例数が少なく、用例クラスターの弁別に貢献していない。

二格が関係する述語の精度が低い最大の理由は二の用法の多様性にある (用法によって二格が項に相当する要素だったり、しなかったりする) (注 14)。改善策は (17) で示す。二格要素と同じことがデ格要素の集合に関しても言えるが、デ格はそれでもまだ二格よりは曖昧性が少なく、ま

た、二格と違って (準) 項をマークする率も低いので、デで終わっている文節の下位分類が効果的ではないことは、深刻な精度の低下には繋がっていないように思える。

[2] 言語学者による対策の提案

以上の問題は次の対応を取ることで改善できるのではないだろうか？

(15) 用例結合のための閾値は現在は一律に設定されているが、述語の意味クラスごとに別の閾値を最適化したり、統計量を反映させることには精度向上の効果があるのではないかと思う。図 3 の結果を見る限り、結合の閾値をヲ格をもつ述語とともたない述語に区別するだけでも、それなりの効果があるだろう。

(16) 用例収集・結合の活用形ごとの層別化…大きな効果は期待できないだろうが、李ほか (文献 13) の結果を考えると述語をタ／テイタ／ル／テイル形ごとに採集し、活用形ごとに用例クラスター化した後に一つに統合することで、より良い結果が得られる可能性もある。一部の述語のクラスターングはこれで精度が向上するかも知れない。

(17) 二格要素と非二格要素の識別…二で終わってい

る文節（＝名詞句あるいは後置詞句）から二格を識別する必要については5節[1]②ですでに述べた。これには次の方法に即効性のある効果が期待できると思う（注15）。

a 「Xに」の要素 X について他の格（ ϕ ） X が、 X を、 X から、 X の、……）の要素になるかを判定し、

b 他の格に現れる（か期待される格分布に従っている）と判断できる）なら、二格要素と、

c 他の格に現れない（か期待される格分布から明らかに偏っていると判断できる）なら、非二格要素とする。

これによって統計量を使つて二格と非二格を効果的に区別できる可能性は高いと思われる（名詞句に期待される格分布の正確な姿は一般にはわかっていないので、これは独立に調べておく必要がある）。

話を更に一般化すると、次の可能性に繋がる。

- (18) 格分布クラスを使つた名詞の下位分類の推定…標準化された格列 A ＝「―ガ、―ガ2、―ヲ、―ニ、―デ、―カラ、―マデ、……」がある時（注16）、 $A(N) = [p_1, \dots, p_n]$ ($p_i = \frac{n_i}{\sum_{j=1}^n n_j}$) によつて名詞 N の格頻度分布を定義し、 $A(N_1), \dots, A(N_n)$ を同値類

に分類（例えばクラスタリング）して $N_1, \dots, N_n$ の下位分類を得て、それで N のクラスを推定する。

[3] 現行の辞書構築の主な問題点

評価を通じて明らかになった現行の格フレーム辞書の（言語学者から見る限り）問題だと思われる点は次の通りである。

- (19) 共起情報の損失…現行の格フレーム辞書では名詞（句）同士の共起情報を利用されていない。このため「 x が y を襲つた」から二格名詞として y が採られた時のガ格名詞が何であつたかが明示されない。

- (20) 同表記異語の扱い…「避ける」には「さける」と「よける」が混在していた。読みの推定と語義の推定は表裏一体なので、語義推定の手法の進歩が問題の根本的解決には不可欠だろう。

- (21) より大きな単位での検索の必要性…意味を単位とした場合、形態素よりも大きな単位の分割が必要である。この意味での過分割の発生頻度は高く、影響は無視できない。「者」のような接尾辞、「上」のような接尾辞化した形式名詞が「見かけの被覆率を向上させつゝクラスター結合の精度の低下



を招いている可能性は高い。これは係り受けの前に形態素の結合処理を入れることで対処できるだろう。

(22) 複合動詞の扱い…複合動詞の格支配がうまく処理されていない。 V_1 (e.g. 「襲う」「解く」) が V_2 (e.g. 「襲い+かかる」「解き+放つ」) という形で複合動詞をなす場合に、 V_2 の格要素となつてゐる形が「」の格要素として抽出されてゐる。

(23) ハに吸収されてゐる格の扱いの改善…係助詞ハの要素は収集から外されてゐる。これはハに化けて現れやすい格要素の事例が十分に収拾できないという結果を生んでゐる。評価の結果を見ると、影響を受けやすいのはガ格である。ガ格は一般に省略されやすいと考えられてゐるが、ハに吸収されてゐるために取りこぼし率も高いと思われる。ガ格の被覆率を上げるにはハに隠された格の推定が必要かつ有効だろう。

6 終わりに代えて…加工されたデータを使う理由

Webデータを言語研究に使えるようにするには加工

「前処理が必要である。その加工は長谷部（文献22）の実装したマークアップ除去処理には限らない。表層情報よりも先に進み、意味情報にアクセスしようと思えば、形態素解析、構文解析の自動化の重要性が増す。格フレーム辞書は、人手で解析することは不可能なほど大規模なデータから抽出され、それなりの精度で一般化された述語／項関係の大規模データベースである。その被覆率の高さは圧倒的である。これをうまく利用することで意味研究のための前処理を省ける。

一方、NLP研究者は言語学者からのフィードバックに期待してゐる。係り受けの精度向上は時間が自然に解決するタイプのものではない。言語学者が格フレーム辞書を使い、その改良のためにNLPにフィードバックをすることは、言語学の将来のための投資だとは考えられないだろうか？

注

1 より詳細な検討は前川（文献6）などを参照のこと。

2 <http://mecab.sourceforge.net/>

3 現状では類似度の計算にシソーラスの情報を使つてゐるが、それを使わない方法への移行を進めてゐることである。現行の構築法でグループ化に使つてゐる類似度の指標は、述語の直前にある格要素（直前格）の名詞の類似度である。

4 <http://nlp.keekyoto-u.ac.jp/~nlp-resource/juman.html>

- 5 <http://nlp.kueekyoto-u.ac.jp/nl-resource/knp.html>
- 6 依存構造の定義を尊重すれば厳密にはsは「頼って」出た」も含むべきだが、現行の格フレーム辞書ではVの連用形+テを項／述語対の集合に含めていない。
- 7 係り受け解析と依存構造解析の同一視には問題がある。詳細は黒田・飯田（文献12）を参照。
- 8 「次郎は」が「頼って」の（ガ格の）項であることは係り受け解析では明示的には表現されず、(V₁、V₂ のような)並列構造の処理などを経て復元する必要がある。
- 9 これは現在、開発者の河原大輔（NICT）氏が対応策を実装している最中である。
- 10 この階層性を形成する要因として認知的な理解しやすさなどが挙げられており、使用頻度もこの階層に沿うと考えられている。
- 11 この抽出作業には格フレーム辞書の検索のwebインターフェイス <http://reed.kueekyoto-u.ac.jp/cf-search/> は使っていない。ただし、格フレーム辞書の検索結果から原文を参照できるので、同じことをwebインターフェイスを使って行なうことは（手間はかかるが）可能である。
- 12 ただ、これが「積む」「固まる」「賛成する」「躍進する」の用法クラスターとの対比で説明として成立するためには、ハがガを隠しやすい述語（か環境）とそうでない述語（か環境）の区別があることを示す必要があり、そのための独立の調査が必要である。
- 13 この可能性は黒田ほか（文献11）の研究に先立って藤井ほか（文献21）で示唆されている。
- 14 ヲ格を項に取らず、ニ格を項に取る述語の中では「寄与する」が例外的に精度がよい。これはこの動詞と共起するニ格の曖昧性が例外的に低く、程度や様態を表わす副詞類が共起

- しないことが理由だった。
- 15 この方法は加藤鉦三（信州大学）の提案による。
- 16 ただし、自動詞のガ格と他動詞のガ格は区別する必要があるかも知れない。

参考文献

- 1 B. A. Fox. The noun phrase accessibility hierarchy reinterpreted: Subject primacy or the absolutive hypothesis? *Language*, 63:856-870, 1987.
- 2 E. Keenan. Variation in universal grammar. In R. Fasold & R. Shuy, editors, *Analyzing Variation in Language*, pp. 136-148. Georgetown University Press, Washington, D. C., 1975.
- 3 E. Keenan & B. Comrie. Noun phrase accessibility and Universal Grammar. *Linguistic Inquiry*, 8:639, 1977.
- 4 久野暉（一九七三）『日本文法研究』（大修館書店）
- 5 伊藤雅光（二〇〇五）「計量言語学とコーパス言語学」、『計量言語学』25(2)、八九〜九七頁
- 6 前川喜久雄（二〇〇七）「大規模均衡コーパスが開く可能性」、『日本語言語学会第134回大会予稿集』公開シンポジウム、二四〜二九頁、日本語学会
- 7 奥津敬一郎（一九七四）『生成日本文法論』（大修館書店）
- 8 中本敬子・黒田航（二〇〇六）「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』五九二〜五九五頁、発表四（一頁）
- 9 三原健一（一九九四）『日本語の統語構造』（松柏社）
- 10 竹内孔一（二〇〇七）「グラフ構造に基づく同時クラスターリングを利用した動詞の属性クラスの抽出」、『情報処理学会

特集 WWWを対象にした日本語研究

- 研究報告』2007-NL-182 (9)
- 11 黒田航・李在鎬・渋谷良方・河原大輔・井佐原均 (二〇〇七)
「自動獲得された大規模格フレーム辞書の精度向上を見込んだ人手評価」(『言語処理学会第13回年次大会発表論文集』一〇五四～一〇五七頁)
- 12 黒田航・飯田龍 (二〇〇六)「文中の複数の語の(共)項構造の同時的、並列的表現法: Pattern Matching Analysis (Simplified) の観点からの「係り受け」概念の拡張」(『信学技法』106 (191)・1・5)
- 13 李在鎬・鈴木幸平・永田由香・黒田航・井佐原均 (二〇〇七)
「動詞「流れる」の語形と意味の問題をめぐって」(『計量国語学』26 (2)、六四～七四頁)
- 14 寺村秀夫 8一九八一「日本語の文法」(下) (国立国語研究所)
- 15 寺村秀夫 (一九九三)『寺村秀夫論文集Ⅰ…日本語文法編』(くろしお出版)
- 16 加藤重広 (一九九九)「日本語関係節の成立要件(2)…先行研究の整理とその問題点」(『富山大学人文学部紀要』30・65、111)
- 17 加藤重広 (二〇〇〇)「日本語関係節の成立要件(2)…文法論的要因と語用論的要因」(『富山大学人文学部紀要』31・71、156)
- 18 丸元聡子・乾裕子 (二〇〇〇)「連体修飾を受ける体言の格構造の復元…コーパスに基づく「内の関係」の分析」(『言語処理学会第6回年次大会発表論文集』一六～一九頁、言語処理学会)
- 19 河原大輔・黒橋禎夫 (二〇〇五)「格フレーム辞書の漸次的自動構築」(『自然言語処理』12 (2)、一〇九～一三二頁)
- 20 河原大輔・黒橋禎夫 (二〇〇六)「高性能計算環境を用いたWebからの大規模格フレーム構築」(『情報処理学会研究報告』2006-NL-171・六七～七三頁)
- 21 藤井敦・秋山典丈・徳永健伸・田中穂積 (一九九五)「動詞の多義性解消における格の弁別能力と集中度の有効性について」(『言語処理学会第一回年次大会発表論文集』一一七～一二〇頁)
- 22 長谷部陽一郎 (二〇〇六)「Wikipedia 日本語版を利用した言語研究の手法」(『言語文化』9 (2)、三七四～四〇三頁)
- 23 井上和子 (二〇〇七)『変形文法と日本語』(上)(大修館書店)
- 24 李在鎬・黒田航・渋谷良方・河原大輔・井佐原均 (二〇〇七)
「コーパス分析に基づく連体表現の使用調査」(『第134回日本言語学会大会予稿集』四二二～四二七)
- (くろだ・こう 独立行政法人 情報通信研究機構専攻研究員)
(リ・ジェホ 独立行政法人 情報通信研究機構専攻研究員)