

単語の意味に関する概念ベースの高性能化

Nguyen Viet Ha^{†,1} 石川 勉[†] 帆苺 譲^{††} 笠原 要^{†††}

[†] 拓殖大学工学部情報工学科

^{††} NTTアドバンステクノロジー (株)

^{†††} NTTコミュニケーション科学基礎研究所

あらまし 単語間の意味的な類似性を判別するため、単語（以下、概念と呼ぶ）に関する知識ベース（概念ベースと呼ぶ）の構築を進めてきた。この概念ベースでは、各概念は、国語辞書の語義文から獲得される自立語を属性、その出現頻度を属性値として表現され、概念間の類似度は共通属性の比較により算出される。これまで、約4万の日常語からなる概念ベースを構築してきたが、具体的な応用においては概念数の不足や性能的な不十分さが問題となった。本論文では、語義文における属性の重要性およびシソーラス上の属性間の関係を考慮した属性値算出手法を提案し類似性判別能力の向上を図る。また、新語や造語を含めたあらゆる概念に対し類似性判別を可能にするため、一般的な単語に関する概念ベースおよび漢字概念ベースからなる大規模概念ベースの構築を行った。さらに、類語辞典を利用した類似性判別能力の評価法を提案した。この方法により評価した結果、従来の概念ベースに比べ大幅な性能向上が確認された。

1 はじめに

最近、コンピュータ上で柔らかな処理を実現しようとする研究が注目されている。実世界では、全てのデータが整った問題ばかりでなく、不完全なデータを許容する技術の重要性が高まってきたためである。このような柔軟な処理の実現を目指し、我々は、不完全な知識の下でも不完全さに応じた概略的な判断を行える推論法について研究を進めている [1]。具体的には、知識が欠落して厳密な解を得られない場合でも、それと類似した知識を補完することにより推論を進め、概略的な解を導くものである。すなわち、この方式では、知識間の類似性判定が必要となる。知識は基本的には単語を用いて表現されるので、知識間の類似性はそれを構成する単語間の類似性に基づいて決定されることになり、これを定量化する技術が必要となる。

単語に関する知識ベースとしては、日本ではEDR辞書 [2] やIPAL辞書 [3] が、米国ではCYC [4] が代表的であるが、これらは用途を限定せず、汎用性を意識し、人手で作成されている。これに対し、我々は用途を単語の意味の類似性判別に限定し、機械的に知識ベース（以下、概念ベースと呼ぶ）の構築を

進めている。概念ベースでは、各単語（以下、概念と呼ぶ）は複数の属性と属性値のペアで表現される。属性はその概念に関連した概念であり、属性値はその関連の度合いである。この場合、概念間の類似性の度合い（以下、類似度と呼ぶ）は、属性の共通性の度合いを利用して定義すれば単純に計算することが可能となる。

これまで、日常用語約4万語からなる概念ベースを構築し、概念間の類似性判別について一定の性能を確認した [5, 6]。しかし、新聞記事検索等への具体的な応用においては概念数が不足し、また、類似性判別能力も充分とはいえなかった [7]。

本論文では、これらに対処した大規模概念ベースの構築について述べる。本概念ベースは、一般語（約26万語）および漢字（約6千字）の各概念ベースにより構成され、前述した概念数不足に対しては、語数を拡大した一般語概念ベースと漢字概念ベースで対処する。ここで、漢字概念ベースは、概念が一般語概念ベースにない場合に、漢字単位の情報をもとにその概念を近似するために用いる。これにより、「激安」等といった漢字から構成される新語や造語に対しても、ある程度類似性判別が可能となる。また、類似性判別能力については、シソーラス上での属性間の関連性や語義文の各文の長さを考慮した属性値付

¹連絡先：拓殖大学工学部情報工学科
〒193-0985 東京都八王子市館町 815-1
E-mail: nguyen@cs.takushoku-u.ac.jp

与の適正化を行うとともに、各概念毎の属性数の均一化を行い、その向上を図る。

また、パーソナルコンピュータ等での利用を考慮して、概念ベースから主要概念を選定し、それに基づいた概念ベースのコンパクト化を行う。

さらに、類語辞典を利用した概念ベースの類似性判別能力の評価法を提案する。この方法を用いて、構築した大規模概念ベースを従来の概念ベースおよび共起辞書（新聞記事から評価用に作成）と比較し、定量的に評価する。

以下、2章で類似性判別能力の向上という観点からみた概念ベース構築手法、3章で大規模化の観点からみた概念ベースの構成、4章で主要概念の選定と概念ベースのコンパクト化、5章で概念ベースの性能評価について述べる。

2 概念ベースの構築手法

本章では、従来の概念ベース構築法 [5] について述べた後、その構築法における類似性判別能力の問題点と改良手法について述べる。

2.1 従来の構築法

概念の知識表現法としては、その属性と属性値の集合を基本とする方法が典型的である [8]。例えば、「林檎」という概念は、{(形: 丸い), (色: 赤い), ...} のように表される。このような知識を獲得すれば、汎用的かつ強力な概念ベースが実現されるが、現在の技術ではこれを機械的に行うことは不可能といえる。

従って、我々は機械的な構築を第一義として、「tf-idf」手法 [9] に基づいて国語辞書から概念知識を獲得することとした。具体的には、単語（概念に相当）の語義文に含まれている自立語を抽出してそれを属性とし、その出現頻度を属性値とする。すなわち、この手法では概念は検索キーワードに相当し、その語義文は新聞記事等の検索対象に相当する。概念ベースの構築手順を以下に示す。

1) 辞書からの属性と属性値の獲得

単語の語義文に含まれる自立語をその単語と何らかの関係があるとみなして、属性とする。また、その自立語の出現頻度が高いほど関連性が高いと考え、「tf」の考えに基づいて出現頻度を属性値とする。ここで、語義文における自立語の重みは全て 1 とする。例えば、「馬」に関する辞書の語義文が、

『馬（うま）』家畜の一つ。たてがみが長い

草食の動物で ... 動物 ...

と書かれているとき、この語義文を元にして以下の属性と属性値の集合が獲得される。

「馬」 = {(家畜:1),(一つ:1), ... (動物:2), ...}.

2) 孫引き・逆引き参照

語義文からは十分な属性が獲得できない場合がある。これに対して、概念の孫引き及び逆引き参照で属性数を増やす。孫引きでは、属性自体も概念であることに着目し、その属性も元の概念の属性とする。また、逆引きでは、概念 g に属性 g' があるとき、 g と g' の間には関連性が存在するので、概念 g' にも属性 g を付加する。

3) 情報量による属性値の修正

概念ベースでは、全ての属性の重要度は一様ではない。例えば、辞書特有の言い回しに由来する属性「もの」や「こと」は多くの概念に含まれているため、概念に対する意味的な重要性は低いと考えられる。従って、「idf」の考えに基づいて出現頻度から属性の情報量を計算し、これを用いて属性値を修正する。

4) カテゴリーへの属性変換

この段階で概念ベースは、全ての概念が属性になりうるので、総概念数を N とすれば、 $N \times N$ の行列で表現されることになる。この場合、 N は辞書の総語彙数でもあり非常に大きいため、同義あるいは類似等互いに関連する属性が多く含まれることになる。すなわち、属性は互いに独立であるとはいえない。このため、ここでは意味の近い属性をシソーラス上の同一のカテゴリー（比較的意味的に独立といえる）にまとめている。なお、これには 2,715 種のカテゴリーが 12 段の木構造に分類されている NTT の日英翻訳システムのシソーラス [10] を用いている。

以上のように構築された概念ベースでは、概念 g_i は以下のように表現される。

$$g_i = \{(p_{i1}, q_{i1}), (p_{i2}, q_{i2}), \dots\}. \quad (1)$$

ここで、 p_{ij} は属性、 q_{ij} はその属性値である。シソーラスのカテゴリー（属性）の総数を $n (= 2715)$ とすると、 g_i は以下のように表現できる。

$$g_i = \{q_{i1}, q_{i2}, \dots, q_{in}\}. \quad (2)$$

ここで、2) までで獲得されなかった属性 p_j に対しては、属性値 q_{ij} は 0 とする。従って、概念は n 次元の意味空間で表現され、その空間でのベクトルであるとみなすことができる。

このようにベクトル表現される概念間の類似度は、各属性が独立（直交）と仮定すれば、単純にそのベクトルがなす角度 θ の余弦（正規化すると内積）として、以下のように定義することができる。

$$R(g_i, g_j) = \cos(\theta) = \sum_{k=1}^n q_{ik} q_{jk}. \quad (3)$$

従って、類似度は0～1の間の値をとることになる。なお、同式において属性値 q_{ij} はその二乗和が1になるように $(\sum_{j=1}^n q_{ij}^2 = 1)$ 正規化した値である。

また、類似度は文脈や見方（以下、観点と呼ぶ）によって変化する。例えば、観点「家畜」における「豚」と「馬」の類似度は、観点「乗る」におけるそれらの類似度より高いといえる。このような観点を考慮した類似度は、観点の主な属性に対応する概念の属性の属性値を増幅した後、式(3)を適用して算出する。

なお、以上の概念ベースの構築法および類似度計算法の詳細については文献[5]を参照されたい。

2.2 改良手法

前節に述べた構築手法では

- (1) 属性が意味的に独立でない
- (2) 概念あたりの属性数のばらつきが大きい
- (3) 語義文中の重要な属性が強調されていない

等の問題があり、適切な類似度が算出されない場合がある。これらに対し、それぞれ、シソーラスの情報を用いた上位属性の獲得、不要属性を削除した属性数の均一化、語義文の長さに応じた属性値の付与で対処する。以下、これらの手法について述べる。

2.2.1 上位カテゴリー付与

前述したように、概念間の類似度はベクトルの内積で求めている。しかし、このベクトルの基底である属性が完全に意味的に独立（直交）していないため、適切な類似度が算出されない場合がある。

例えば、以下のように概念「ダイニングテーブル」(g_a とする)が属性“机”と“材木”，概念「踏み台」(g_b とする)が属性“台”と“板”を持つと仮定しよう。

$$g_a = \{(\text{机}:0.7), (\text{材木}:0.7)\}.$$

$$g_b = \{(\text{台}:0.7), (\text{板}:0.7)\}.$$

この場合、類似度 $R(g_a, g_b)$ を求めると“0”となる。しかし、実際“机”と“台”，“材木”と“板”は、シソーラス的に表現すると、例えば図1のように関連する。従って、本来 g_a と g_b は意味的にある程度類似し、 $R(g_a, g_b) > 0$ となるべきである。このため、このような直交でない意味空間を考慮した類似度算出が必要となる。

直交する意味空間を構築する方法としては、行列の直交変換や主成分分析が利用されている[11, 12]。例えば、宮原らは辞書から獲得したデータに対して、固有値分解を行い、直交する意味空間を構築してい

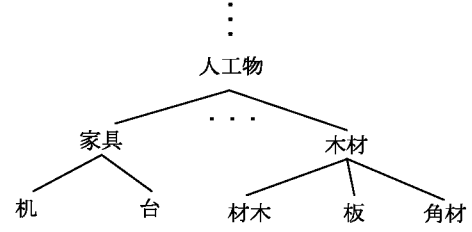


図 1: 属性間の関係

る[11]。しかし、これらの方法では構築された空間の基底はデータに依存して決定されるため、データが不完全な場合には得られた基底は真に直行しているとはいえない。また、たとえ完全なデータを獲得できたとしても、類似度を前述したベクトルの内積で求める場合、直交変換された後の空間での類似度は、変換前の空間での類似度と同じ値となる。

ここでは、シソーラス上の関係に着目して、属性の非直行性を補完する手法を提案する。概念はある属性を持ったとき、シソーラス上では下位属性と上位属性との間にはis-a関係等が存在するため、その属性の上位属性に対しても関連性を持つ[13]。この関係を利用し、もとの属性からその上位の属性を獲得し、概念ベクトルを修正する。例えば、先の例では上位属性「家具」や「木材」を以下のように概念 g_a , g_b の属性とする。

$$g_a = \{(\text{机}:0.7), (\text{家具}:q_1), (\text{材木}:0.7), (\text{木材}:q_2)\}.$$

$$g_b = \{(\text{台}:0.7), (\text{家具}:q_3), (\text{板}:0.7), (\text{木材}:q_4)\}.$$

これから分かるように、 g_a , g_b は共通な属性を持つことになり、類似度 $R(g_a, g_b) > 0$ となる。このような上位属性の付加により、実効的に属性の非直行性が補完されると考える。この場合、上位属性に対し、どの程度の値 $q_1 \sim q_4$ （付加値と呼ぶ）を付加するか問題となるが、付加値は下位属性の属性値とシソーラスの構造に依存すると考えられる。従って、ここでは、シソーラスの構造をもとに算出された属性の情報量を利用して付加値を決定する。すなわち、属性が上位の位置にあるほど下位属性の数が多くなり情報量が小さくなるので、付加値を小さくする。

具体的には、概念 g_i において、属性 j に属性値が存在した場合には、その上位属性 k に対する付加値 q_{ik}^j を以下のように求める。

$$q_{ik}^j = q_{ij} \left(\frac{I_k}{I_j} \right). \quad (4)$$

ここで、 I_j , I_k は、それぞれ属性 j , k の情報量であり、 q_{ij} は属性 j の属性値である。また、その上位

属性 k には複数の下位属性が存在したり、それ自体に初めから属性値が存在する場合もある．これらを考慮し、ここでは、最終的な属性 k の属性値 $\hat{q}_{ik}$ を以下のように求めることとする．

$$\hat{q}_{ik} = \max \{q_{ik}, q_{ik}^{j_1}, q_{ik}^{j_2}, \dots\}. \quad (5)$$

ここで、 q_{ik} は初めからの属性 k の属性値であり、 $q_{ik}^{j_1}, q_{ik}^{j_2}, \dots$ は k の下位属性から算出された付加値である．例えば、前述の例で概念「踏み台」が属性「角材」と「木材」も持つとする（図 1）．この場合、属性「木材」の付加後の属性値は、「木材」の属性値と「板」、「角材」の属性値から算出された付加値の最大値となる．

また、属性の情報量はその属性に属する下位属性の数をもとに算出することとし、以下のように設定する．

$$I_k = -\log_2\left(\frac{n_k + 1}{n}\right). \quad (6)$$

ここで、 n_k はシソーラスにおける属性 k の下位属性の数であり、 n は総属性数である．

2.2.2 属性数均一化

概念ベースの応用の一つである概略推論法 [1] では、二つの概念が類似するか否かの判断を行う．すなわち、この判断のための固定したしきい値が必要となる．また、 g_a, g_b が類似し、かつ g_c, g_d が類似しないとすれば、 $R(g_a, g_b) > R(g_c, g_d)$ が成り立つ必要がある．

一方、国語辞書では語義文の記述量が各概念で同じのではなく、その量が多い概念には獲得される属性の数が多くなる．概念間の類似度は平均的には、共通属性の比較に基づいているので、属性数の多い概念のペアではその値が高くなり、属性数の少ない概念のペアではそれが小さくなる傾向がある．属性が平均的に分布する概念ベースモデル [6]（総属性数 M 、平均属性数 r ）を用いると、全ての概念間の平均的な類似度 $\bar{R}$ は以下で与えられる（属性値は均一で $1/\sqrt{r}$ と仮定）．

$$\bar{R} = \sum_{k=0}^r \frac{M C_k \times M - k \ C_{r-k} \times M - r \ C_{r-k}}{(M C_r)^2} \times \frac{k}{r}. \quad (7)$$

従って、 r が大きければ大きいほど $\bar{R}$ が大きくなる．例えば、上式では $M = 12$ 、 $r = 3, 4, 5$ の場合、 $\bar{R}$ はそれぞれ、0.25、0.33、0.42 となる．このため、属性数が多ければ実際には類似しない概念間の類似度が、類似する概念間の類似度より大きくなる場合が生じ得る．

一方、前述した孫引きや逆引き参照で属性を増やすため、概念に関連しないノイズとなる属性が付与されることがある．属性数が多ければ多いほどノイズも多くなると考えられる．しかし、ノイズとなる属性は、本来その概念と関連しないので、語義文中の出現頻度が低く、属性値が小さいと考えられる．

以上を考慮して、ここでは属性値の小さい属性をノイズとみなし、これらを削除して属性数の均一化を行う．なお、このときの属性数は実験的に決定する．

2.2.3 語義文解析手法

従来の手法では、属性値を単純に語義文における属性（自立語）の出現頻度で決定している．すなわち、全ての自立語の重みは同じとしている．しかし、語義文中の個々の文においては、その文が短いほど自立語の持つ意味が大きいと考えられる．見出し語（概念）を少ない語数で説明しているため、自立語と見出し語の関連性が高いといえるからである．従って、ここでは文の長さ（自立語の数）の逆数を重みとし、属性値を出現頻度とその重みで算出する．具体的には、一文の中の自立語の数を m とすると、自立語の重要度 d を以下のように求める．

$$d = 1/m. \quad (8)$$

属性 p の属性値 q は、それが語義文に n 回出現し、重要度がそれぞれ、 $d_1, \dots, d_n$ のとき、下式で求める．

$$q = \sum_{i=1}^n d_i. \quad (9)$$

例えば、「空（そら）」の語義文が「地上に高くそびえる空間．天．」であったとする．この場合、「空」の各属性値は以下のように算出される．

$$\text{「空」} = \{(\text{地上:0.25}), (\text{高く:0.25}), (\text{そびえる:0.25}), (\text{空間:0.25}), (\text{天:1})\}.$$

なお、従来の手法を用いると、この例では属性値は全て 1 となる．

3 大規模概念ベースの構成

2 章に述べた手法を用いて概念ベースの構築を行う．しかし、あらゆる概念に対して、この手法で概念ベクトルを生成することは、新語や造語の存在を考えると不可能である．従って、ここでは、大規模な一般語の概念ベースとともに漢字概念ベースを構築し、一般語概念ベースに存在しない概念に対し、漢字概念で近似することとする [14]．

3.1 一般語概念ベース

三つの国語辞典 [15, 16, 17] を用い、2 章で示した方法により総数 257,905 の一般語概念ベース（一般

語 GB) を構築した。また、属性数を 100 で均一化した。5 章に述べる評価手法で評価した結果、均一化する前と均一化した後の評価値は、それぞれ 0.609 と 0.611 となり、均一化することにより性能が若干向上した。評価値の差は誤差範囲内であるが、ノイズが削除され性能が向上したとも考えられる。なお、均一化前は属性数が 10 から 1,619 までばらつき、その平均は 229 であった。また、均一化後は、属性数 100 未満の概念数は 47,125 であり、これらの概念の平均属性数は 70 であった。

3.2 漢字概念ベース

漢和辞典 [18] と国語辞典を用いて総字数 5,651 の漢字概念ベース (漢字 GB) を構築した。国語辞典の場合、漢字一字の見出し語が存在するため、その語義文も利用することとした。この概念ベースでは、語義文中の自立語 (属性) は、基本的には漢字一字ではないため漢字概念とはならない。従って、属性も概念であるという特徴を利用した「孫引き・逆引き」参照は行わない。また、前節と同様に属性数を 100 で均一化した。均一化前は、属性数は 4 ~ 1,499 までばらつき、その平均は 467 であった。また、均一化後は、属性数 100 未満の概念数は 443 であり、これらの概念の平均属性数は 59 であった。

漢字 GB を利用した概念ベクトルの合成は以下のように行う。

- 1) 概念 (単語) を漢字に分解し、漢字ごとの属性ベクトルを抽出する。
- 2) 漢字のベクトルを線形結合し、属性値の二乗和が 1 になるように正規化を行い、概念ベクトルとする。

概念ベクトルの合成例を図 2 に示す。このように、「激安」等の新語や造語に対して概念ベクトルを合成することができる。

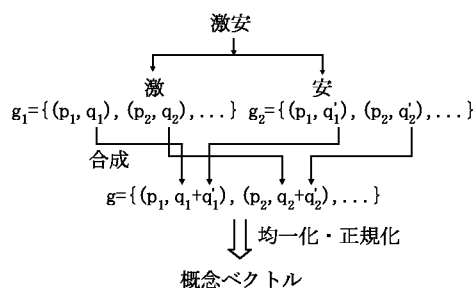


図 2: 漢字 GB による概念ベクトル合成例

4 主要語選定と概念ベースのコンパクト化

前章で述べた概念ベースの容量はテキスト形式で 500MB 以上になり、パーソナルコンピュータ等の主記憶上での利用は困難である。従って、このような利用を想定した、コンパクト版概念ベースの構築を行った。コンパクト化は、主要な概念を選定し、主要でない概念を主要概念で近似することにより実現した。以下、これらについて述べる。

4.1 主要概念の選定

主要概念は国語辞書 [17] の 10 万語の中から手作業で選定した。選定の基準としては、日常的な基本単語を取り出すことを目指して以下を設定した。

- i) 複合語 (例えば「立体図形」) を除く。ただし、極めてよく使うもの (例えば、「株式会社」) は選択する。
- ii) 慣用的な用語 (例えば、「舌を巻く」) を除く。
- iii) 接頭辞、接尾辞のついた語 (例えば、「不確実」や「歴史的」) を除く。ただし、極めてよく使うもの (例えば、「不真面目」や「合理的」) は選択する。
- iv) 複数の動詞が結合した単語 (例えば、「叩き出す」) を除く。ただし、よく使うもの (例えば、「組み合わせる」) は選定する。

これらの基準は極めて曖昧であるともいえるが、厳密な基準は、自然言語を対象としているためその設定自体が困難である。また、その設定ができたとしても、それを用いて人間が判定を行うため、適切にその基準に合致するとは考えづらい。さらに、判定に要する時間も膨大となる。従って、ここではこのような基準を用いることとした。選定にあたっては、この曖昧性を補完し主要概念を取りこぼさないために、4 人の実験者で行い、それぞれ選定された概念の和集合を主要概念集合とした。その結果、47,441 の概念が選定された。なお、各実験者が選定した主要概念の数はそれぞれ、31,731, 26,055, 18,962, 22,266 であった。また、4 人、3 人以上、二人以上、一人以上に選ばれた概念の数はそれぞれ、7,155, 16,183, 28,235, 47,441 であった。

なお、主要概念集合として、従来 GB の概念集合をそのまま利用することも可能であったが、機械的に構築しているため、一部重要な概念が欠落していた。従って、前述したように新たに主要概念の選定を行った。

4.2 概念の代表化

概念ベースを圧縮する一つの手法として、主要でない概念をそれと類似した主要概念で代替（以下、代表化と呼ぶ）することが考えられる．ここでは主要概念でない概念 g に対し、主要概念の中からそれと一番類似する概念 g' （式（3）で求めた類似度が一番高いもの）を検索し、これを g の代表概念とした．

前節の主要概念を用いて代表化を行った結果、本 GB の容量は 100MB となり、約 1/5 に圧縮された．なお、もとの概念 (g) と代表概念 (g') との類似度の平均と標準偏差の実測値は、それぞれ 0.85 と 0.13 であった．この結果から、このような単純な代表化でも、もとの概念ベースをある程度近似できているといえる．

5 評価

概念ベースの類似性判別能力を定量的に評価する必要がある．ここでは、類語辞典を利用した自動評価法を提案する．この方法を用いて、作成した大規模概念ベースと従来の概念ベースを評価した．また、相対的に性能を比較するために、新聞記事より共起概念ベースを構築し、同様の手法で評価を行った．以下、評価手法とその結果について示す．

5.1 評価手法

概念ベースの特性としては、

- (1) 類似する概念との間の類似度と全く類似しない概念との間の類似度の差が大きい
- (2) 類似する概念との間の類似度は関連する概念（類似しないが何らかの関連性を持つ）との間の類似度より大きい

が必要といえる．これを考慮した評価法には類語辞典を用いた評価手法 [6] があるが、ここではこれを評価値が最大 1 になるように改良する．具体的には、まず類語辞典のなかからランダムに n 個のサンプル概念 (g_a) を抽出する．各サンプル概念 g_a の類語として挙げられている単語の一つを類似する概念 g_b 、類語より一つ上のレベルの分類に含まれる単語の一つを関連する概念 g_c とする．また、最大分類の異なった分類の中からランダムに一つの単語を選択し無関係の概念 g_d とする．これらの n 組の $g_a \sim g_d$ を用いて、上記の (1) の評価指数として以下を設定する．

$$f_1 = (r_1 - r_3)/(1 + \sigma_1 + \sigma_3). \quad (10)$$

ここで、 r_1 、 r_3 はすべての組の $R(g_a, g_b)$ 、 $R(g_a, g_d)$ の平均値であり、 σ_1 、 σ_3 はその標準偏差である．同式の分母は、二つの類似度分布の重なり具合を表わし

ており、これが小さいほど判別能力が高いと考える．

一方、(2) の評価指数としては、単純に正しい判断 ($R(g_a, g_b) > R(g_a, g_c)$) の割合が大きいほどよいとし、以下を設定する．

$$f_2 = m/n. \quad (11)$$

ここで、 m は正しい判断の数を表す．

これらを用い、概念ベースの類似性判別能力の評価指数を以下のように設定する．

$$F = f_1 \times f_2. \quad (12)$$

この総合評価指数 F は、類似概念間の類似度が 1、無関係概念間の類似度が 0、また類似概念と関連概念との類似度の逆転がない、理想的な概念ベースの場合で 1 となる．

5.2 評価結果

類語辞典 [19] を用い、200 組のサンプル概念を選択し、各概念ベースについて評価を行った．結果を表 1 に示す．同表から分かるように、本 GB の評価値は従来 GB の評価値を大きく上回った．すなわち、2.2 節に示した改良手法は有効であったといえる．評価値 $F = 0.611$ は、前述 $F = 1$ の理想的な場合に対し、かなりよい値といえる．また、漢字 GB のみを用いた場合でも、類似性判別がある程度可能であることが明らかになった．なお、サンプル概念の一部を表 2 に示す．

表 1: 評価結果

	本 GB (一般語)	本 GB (漢字)	従来 GB	共起 GB
f_1	0.650	0.345	0.456	0.114
f_2	0.940	0.770	0.880	0.624
F	0.611	0.266	0.401	0.071

表 2: 評価データの例

サンプル (g_a)	類似 (g_b)	関連 (g_c)	無関係 (g_d)
全力	総力	人力	鉾物
視野	視界	視線	戦死
悪臭	異臭	臭い	船
絶食	断食	試食	本
歓声	歓呼	叫ぶ	貯蓄
跳躍	飛躍	快速	唾
熟睡	熟眠	仮眠	番号
近道	早道	通行	学友
行為	行動	言動	橋

また、このようなベクトル表現の概念ベースの一つとして、新聞記事データ [20] から共起概念ベース

を作成した．ここでは，文書中の共起情報に基づいて属性と属性値を決定する．例えば，「予防接種の普及で発生数は減っている」という文があったとき，自立語である「予防」，「接種」，「普及」，「発生」，「数」，「減る」を抽出し，これらは互いに関連すると考える．具体的には，例えば，「予防」の概念に対し，他の概念をこの概念の属性とする．また，属性値は，文のなかの距離が短いほど関連性が高いとし，距離の逆数とする．例えば，「発生」は「予防」から数え3番目の単語であるため，距離を3とし，「発生」の属性値を $1/3$ とする．すなわち，「予防」の属性と属性値は以下のように算出される．

$$\text{「予防」} = \{(\text{接種:1.0}), (\text{普及:0.5}), (\text{発生:0.33}), (\text{数:0.25}), (\text{減る:0.2})\}.$$

このように属性と属性値を新聞記事の全ての文で獲得した後，2章で述べた他の手順を用いて共起GBを構築した．なお，出現頻度の低い（3以下）単語（概念）は，十分な属性数を確保することができないと考え，共起GBから外すこととした．

こうして構築した共起GBに対し，同様に評価を行った結果，評価値 $F = 0.071$ が得られた（表1）．なお，共起GBは一年分の記事データしか用いていないため，前述の200組のサンプル全てが含まれてはいなかった．従って，評価はこれに含まれていた157組について行った．また，この157組を用いて，本GB（一般語）を再評価した結果， $F = 0.611$ が得られた．これらの結果から，類似性判別等の用途においては国語辞書を知識源とした概念ベースは，新聞記事等を知識源とした共起ベースより圧倒的に有効であるといえる．

5.3 連想検索による評価

前述した機械的評価では感覚的な性能を把握しづらいため，一例として，概念「台風」に対して，もっとも類似度が高い概念の上位10個を各概念ベースで抽出した．結果を表3に示す．同表から本GBでは，従来GBや共起GBに比べ，適切な類似概念が多く検出されていることが分かる．また，本GBの主要概念を用いた場合，日常的に多用される概念が多く検出されることが分かる．すなわち，類語検索等にはこれを利用した方がよいといえる．

6 むすび

本論文では，単語の意味の類似性判別のための大規模概念ベースの構築手法について提案した．本手法では，シソーラス上での属性間の関連性や語義文長を考慮した属性値付与，さらに属性数の均一化等

表 3: 「台風」の上位類似概念

本 GB (一般語)	本 GB (主要概念)	従来 GB	共起 GB
颱風	嵐	野分	低気圧
室戸台風	夕風	来襲	曇り
枕崎台風	野分	具風	南々西
野分	朝風	竜巻	前線
嵐	暴風	爽涼	雨
夕風	大嵐	不揃い	降り
野分	一荒れ	秋涼	東南東
朝風	風	春一番	高波
夕風	暴風雨	秋日和	見舞う
暴風	強風	秒速	洪水

を行い，より適切な概念ベクトルの生成を実現している．構築した概念ベースは，総数約26万語の一般語概念ベースと約6千字の漢字概念ベースからなる．漢字概念ベースは，新語や造語等，一般語概念ベースにない概念についての概念ベクトル生成に用いる．

構築した概念ベースの類似性判別能力について評価した結果，従来の概念ベースに対し，大幅な性能向上が実現された．また，主要概念を選定して概念ベースのコンパクト化を行い，小規模のシステムでも容易に利用可能とした．

謝辞

本研究の一部は文部省科学研究費補助金（課題番号11650412）及び同志社大学の学術フロンティア研究プロジェクト「知能情報科学とその応用」の研究助成によって実施された．

参考文献

- [1] Nguyen V. H., 石川 勉, 阿部明典: 知識の類似性を利用した概略推論法, 電子情報通信学会論文誌 D-I, Vol.J84-D-I, No.4, pp.389–400 (2001).
- [2] 横井俊夫, 仲尾山雄, 荻野孝野, 田中裕一: 概念レベルにおける電子化辞書の情報構造, 情報処理学会論文誌, Vol.38, No.1, pp.32–43 (1997).
- [3] 青山文啓, 橋本三奈子: 名詞の辞書記述, 情報処理学会自然言語処理研究, Vol.94–104, pp.9–16 (1991).
- [4] Guha R. V. and Lenat D. B.: Cyc: A midterm report, AI Magazine, Vol.11, No.3, pp.32–59 (1990).
- [5] 笠原 要, 松沢和光, 石川 勉: 国語辞書を利用した日常語の類似性判別, 情報処理学会論文誌, Vol.38, No.7, pp.1272–1283 (1997).

- [6] 石川 勉, 井澤潤一郎, Nguyen V. H., 笠原 要 : 単語の意味に関する概念ベースの類似性判別能力からの最適構成, 人工知能学会誌, Vol.13, No.3, pp.470–479 (1998).
- [7] 熊本 睦, 島田茂夫, 加藤恒昭 : 概念ベースの情報検索への適用, 情処研報, Vol.99, No.1, pp.9–16 (1999).
- [8] 安西祐一郎 : 認識と学習, 岩波書店 (1989).
- [9] Salton G. and Allen J.: Text Retrieval Using the Vector Processing Model, 3rd Annual Symposium on Document Analysis and Information Retrieval, pp.9–22 (1994).
- [10] 池原 悟, 宮崎正弘, 横尾昭男 : 日英機械翻訳のための意味解析辞書, 情処学会自然言語処理研究, Vol.84–13, pp.95–102 (1991).
- [11] 宮原隆行, 清木 康, 北川高嗣 : 意味の数学モデルによる意味的連想検索の高速化アルゴリズムとその実現方式, 情処学会論文誌, Vol.38, No.7, pp.1399–1411 (1997).
- [12] 小島秀樹, 伊藤 昭 : 文脈依存的に単語間の意味距離を計算する一手法, 情処学会論文誌, Vol.38, No.3, pp.482–489 (1997).
- [13] Nguyen V. H., 石川 勉, 笠原 要 : ベクトル表現された概念に対する類似度計算法, AI 学会ことば工学研究, SIG–LSE–A001, pp.49–55 (2000).
- [14] 帆苅 讓, 石川 勉, 笠原 要 : 言葉の意味に関する階層型大規模概念ベースの構築, 情処研報, Vol.99, No.1, pp.25–32 (1999).
- [15] 松村 明 (編) : 大辞林第二版, 三省堂 (1995).
- [16] 新村 出 (編) : 広辞苑第四版, 岩波書店 (1991).
- [17] 金田一春彦, 池田弥三郎 (編) : 学研国語大辞典第二版, 学習研究社 (1988).
- [18] 藤堂明保, 松本 昭, 竹田 晃 (編) : 新版漢字源, 学習研究社 (1994).
- [19] 遠藤織枝, 他 (編) : 類語例解辞典, 小学館 (1994).
- [20] 毎日新聞社 : CD 毎日新聞 94 データ集 (1994).