

TR-I-0218

言語データベースから抽出した
知識データの分布

Distribution of Knowledge Data
Extracted from a Language Data Base

江原暉将

Terumasa EHARA

1991.05.31

内容梗概

自然言語処理のための言語学的知識を構築する方法として、母国語話者による内省に基づく方法と、言語データベースから知識を抽出する方法がある。後者の方法は客観性に優れ、定量的な知識を得ることができる反面、内省では当然現われると考えられる知識が言語データベース中に全く存在しないか、存在してもその度数が小さいという問題がある。これは、データベースの過少性と呼ばれる。本報告では、知識をその知識に含まれる単語数で分類し、さらに、単語の分布にZipfの法則を仮定して、知識データの分布を求める。さらに、処理対象からの標本として言語データベースを構築した場合、処理対象中で未知知識データが出現する割合である未知データ率について考察し、データの過少性を具体的に示す。

Abstract

There are two methods to construct linguistic knowledge-base for natural language processing. One method is to construct the knowledge-base by the native speaker's introspective evaluation and the other method is to extract the knowledge from language database. The latter is more objective and quantitative, but has less-data problem. This is the fact that there are many knowledge data which is supposed to appear in the database by the native evaluation but actually they do not appear at all or appear unfrequently. Linguistic knowledges can be categorized by the number of words included in the knowledge. Using this categorization, the distribution of knowledge data extracted from language database is calculated assuming that the distribution of words is Zipf's distribution. The coverage of knowledge-base constructed from language database which is considered a sample from a population is computed.

1. はじめに

いかなる技術にしても、その技術を達成するには、対象にしようとしている自然現象を詳細に知ることが必要であり、自然言語処理技術の場合も例外ではない。自然言語処理が対象とするものは人間の用いる言語であるから、言語現象に関する知識を得ることが必要となる。従来、これらの知識は、主として、母国語話者の内省に頼って構築されてきたが、内省には揺れがあり、また定量性に欠ける欠点があった。一方、言語データベースを利用して、知識を構築する方法は、客観性に優れ、定量的な知識を得ることができる利点がある。しかし、言語データベースを用いた言語研究には、以下の問題点があることが指摘されている [Atkinson]。それは、データの過大性と過小性である。前者は、同一のデータがデータベースの中で繰り返し現れることを指し、後者は、内省によれば当然現れると思われる言語現象が全く現れないか、現れても少数であることを指す。この問題点は、データベースを用いた自然言語処理システムにも当てはまる。前者は、データベースの効率が悪さであり、後者はデータベースから抽出された知識の被覆率が小さいことである。

本報告では、言語データベースから知識を抽出する場合に、知識データがどのように分布するかを考察し、上記問題点を検討し、解決法を探る。

2. 言語処理における知識データの種類

本節では言語知識を単語数と言語数によって分類し、どのような知識がありうるかを示す。

知識は、まず言語数によって、1言語の知識と2言語の知識に分けられる。(多言語の知識も考えられるが、ここでは、最大2言語に限る。) 1言語の知識は、日本語なら日本語だけの知識であり、2言語の知識は、例えば日本語と英語の対応知識である。ATR対話データベース(ADD) [江原] は1言語と2言語の知識を含んでいる。次に、知識が幾つの単語によって構成されているかで分類する。単語数1の知識は、いわゆる単語辞書である。また、単語数2の知識は、例えば、単語間の係り受けの知識である [柿ヶ原] [江原2]。[Sumita] で利用している「名詞+の+名詞」の知識は、3個の単語から成るが、「の」の部分は固定しているので、これも単語数2の知識と見なせる。一般に機能語は各種の言語でその数が限られているため、ここでは、単語として、内容語のみを考え、機能語は単語数の計算には含めない。さらに複合語も全体で1語と見なす。そうすると、単語数と文節数は同じものになる。そこで、文節間の係り受け知識は、単語数2の知識となる。次に文に関する知識について考える。文は長さには長短があるので、単語数(文節数)を幾つと決めることが出来ない。しかし、述語文節が1つである単文に限れば、文節数は、格要素文節数+1となる。ここで言う格要素としては、自由格や副詞も含んでいる。格要素文節数は1ないし4ぐらいであるから、単文の知識は、単語数2ないし5ぐらいの知識と考えられる。平均的には単語数3程度ではないだろうか。次に文と文との係り受け関係の知識は、2文間の関係であるから、平均的には、単語数6程度の知識と考えられる。[工藤] で研究されている単文間の関係を用いた結束性の算出方式はこの種の知識の利用と考えられる。また、文節間の係り受けとその対訳を利用した翻訳方式 [幸山] は言語数2単語数2の知識を利用している。また、いくつかの翻訳単位を定めた、パターン方式の翻訳手法 [古瀬] [Witkam] は、最大、言語数2の文の対応知識を利用している。[佐藤] でも、単文を翻訳単位としているので、言語数2の文の対応知識である。以上をまとめると、表2. 1が得られる。

表 2. 1 知識データの言語数と単語数による分類

言語数	単語数	知識の例
1	1	単言語の単語辞書
	2	単語（文節）間係り受け辞書 [柿ヶ原] [江原 2]
	3	単文の辞書
	6	単文間の関係の辞書 [工藤]
2	1	対訳単語辞書
	2	係り受け対の対訳辞書 [Sumita] [幸山]
	3	単文の対訳辞書 [古瀬] [Witkam] [佐藤]
	6	単文の対の対訳辞書

3. 知識データの分布

ある言語データベースの中での単語データの分布はZipfの法則か、それを詳細化したMandelbrotの法則 [Mandelbrot] に、近似的に従うことが経験的に知られている。これらの法則は次のように表される。ある言語データベースの中の単語の異なり数をMとする。各異なりデータを度数の大きい順に並べたときの順位をrとすると、第r位の語のデータベース内での度数 $n(r)$ は次の式に従う。ここで、a、b、cは定数である。

$$n(r) = a / r \quad (3.1)$$

$$n(r) = a / (r + c)^b \quad (3.2)$$

(3.1) がZipfの法則であり、(3.2) がMandelbrotの法則である。順位Mのデータの度数は1であるから、(3.1) では、 $a = M$ となる。また、(3.2) では、 $a = (M + c)^b$ となる。単語数1の言語知識の分布の例として、図3.1にATR対話データベースから抽出した国際会議電話対話の単語データに対する順位と度数のグラフを示す。この時の異なり語数Mと延べ語数Nは

$$M = 4,797$$

$$N = 171,303$$

である。

では、単語数2の知識データはどのような分布に従うであろうか。単語数1の知識データの分布がZipfの法則に従うと仮定して、単語数2の知識データの分布を求めよう。単語数2のデータの例として、係り受けデータを考察し、片方の単語を係り元、他方の単語を係り先と呼ぶ。しかし、これは1例であり、単語数2のデータであれば、以下の議論は一般的に成立する。例えば、2語組（バイグラム）の分布にも当てはまる。さて、単語数1の知識データの分布はZipfの法則に従うとしたから、係り元および係り先の分布は共にZipfの法則に従う。この時、係り受けデータの分布を求めよう。係り元、係り先の異なりデータ数を共にMと仮定する。すると、順位1の係り元、係り先の度数 $n(1)$ は(3.1)式から共に

$$n(1) = M$$

となる。ここで、順位 r_1 位の係り元が順位 r_2 位の係り先に係る係り受けの度数を $n(r_1, r_2)$ と書き、

$$n(r_1, r_2) = \frac{a}{r_1 * r_2} \quad (3.6)$$

が成立すると仮定する。これは、係り元と係り先が独立に生起すると仮定したときの係り受けの度数である。

$$r = r_1 * r_2 \quad (3.7)$$

と置いて、この係り受けデータとしての順位を $R(r)$ とする。これを求めよう。まず、(3.6) の右辺分母の値が r となる r_1, r_2 の組合せの数 $k(r)$ を求める。N を自然数の集合として、ある自然数 r に対して、集合

$$D(r) = \{(r_1, r_2) \mid r_1 * r_2 = r, r_1 \in N, r_2 \in N\} \quad (3.8)$$

を定義すると、 $\#(A)$ で集合 A の要素数を表わすとして、

$$k(r) = \#(D(r)) \quad (3.9)$$

となる。 $k(r)$ は r を素因数分解

$$r = p_1^{n_1} * p_2^{n_2} * \dots * p_L^{n_L} \quad (3.10)$$

したとき、

$$k(r) = (n_1 + 1) * (n_2 + 1) * \dots * (n_L + 1) \quad (3.11)$$

となる。このとき、

$$\sum_{1 \leq i \leq r-1} k(i) < R(r) \leq \sum_{1 \leq i \leq r} k(i) \quad (3.12)$$

が成り立つ。ここで、 $\sum [p(i)] f(i)$ は述語 $p(i)$ が成立する i の範囲に亘る関数 $f(i)$ の和を意味する。ここで、(3.11) 式の L が r の対数に比例すると仮定し、各 $(n_i + 1)$ は 2 以上であるから、これを 2 で近似すると、 c, c' を定数として、

$$k(r) \doteq 2^{c' \cdot \log(r)} = c * r \quad (3.13)$$

となる。(3.13) が近似的に成立すると仮定する。すると、(3.12) は

$$c((r-1)^2/2 + (r-1)/2) < R(r) \leq c(r^2/2 + r/2) \quad (3.14)$$

となる。そこで、

$$R(r) \doteq c' \cdot r^2 \quad (3.15)$$

が成立する。これを逆に解くと、係り受けの順位 R のデータに対する、 r の値は、

$$r = c' \cdot R^{1/2} \quad (3.16)$$

となる。そこで、係り受けの順位 R のデータの度数は、

$$n(R) = a / R^{1/2} \quad (3.17)$$

となる。つまり、係り元、係り先の分布が Zipf の法則に従う場合、係り受けの分布は、 $b = 1/2$ 、 $c = 0$ とおいた Mandelbrot の分布に近似的に従うことが分かる。実際、図 3.2 に示す、係り受けデータの分布は上記近似の妥当性を示している。この議論を拡張すると、以下の性質が予測される。

性質 1 言語データベースから単語数 n の知識データを抽出したとき、この知識データベースの度数分布は

$$b = 1/n, \quad c = 0$$

の Mandelbrot 分布で近似される。

(性質終)

4. 知識データベースの被覆率

言語データベースから知識データを抽出して、知識データベースを作成した場合、このデータベースが処理対象に含まれる知識を十分覆っていることが必要である。例えば、言語データベースから単語を抽出して、単語辞書を作成した場合、処理対象に含まれる単語をどれだけ含むかが問題となる。もし、単語辞書に含まれない語が処理対象に存在すると、未知語となり、処理が不成功となることが予測される。本節では性質 1 を仮定して、この問題を考察する。

言語データベースを処理対象からの無作為標本とする。処理対象に含まれる知識の全体を DP とし、その(延べ)要素数を NP とする。また、言語データベースから抽出された知識の全体を D とし、その(延べ)要素数を N とする。 DP の異なり要素の分布が、 $c = 0$ の Mandelbrot 分布に従うとすると、異なり要素数を MP として、

$$NP = \sum [1 \leq r \leq MP] \{a / r^b\} \quad (4.1)$$

が成立しなければ成らない。また、 $r = MP$ において、 $n(r) = 1$ となると仮定すると、

$$1 = a / MP^b$$

であり、これから

$$a = MP^b$$

となる。そこで、

$$n(r) = (MP/r)^b \quad (4.4)$$

となる。一方、異なり要素の度数は整数値を取らなければならないから、順位 r が集合

$$R_m = \{r \mid m \leq (MP/r)^b < m+1\} \quad (4.5)$$

の要素である場合は、

$$n(r) = m \quad (4.6)$$

が成立すると仮定する。これは、図3.1、図3.2から妥当な仮定であることが分かる。 R_m の要素数 $\#(R_m)$ は、

$$\#(R_m) \doteq MP * \frac{(m+1)^{1/b} - m^{1/b}}{(m+1)^{1/b} * m^{1/b}} \quad (4.7)$$

となる。 $m=1$ つまり度数1の異なり要素数は、

$$\#(R_1) \doteq MP * \left(1 - \frac{1}{2^{1/b}}\right) \quad (4.8)$$

となる。 $b=1$ の時は、総異なり数の $1/2$ が度数1であり、 $b=1/2$ の時は、 $3/4$ が度数1である。このように、言語データベースから抽出した知識データには度数の少ないデータが大量にあることが示された。これは、データの過少性を意味しており、この点に対する対策が必要である。

言語データベース D は母集団 DP の無作為標本であるから、 DP のある要素が D に存在する確率は、

$$p = N/NP \quad (4.9)$$

である。そこで、 DP での度数が m の異なり要素が D にサンプルされない確率は

$$p(m) = (1-p)^m \quad (4.10)$$

となる。 D にサンプルされない要素は未知データとなる。 DP での度数が m の要素のうち、未知データとなる要素の延べ度数の平均を $f(m)$ と書くと、

$$f(m) = m * \#(R_m) * p(m) \quad (4.11)$$

となる。そこで、未知データの総延べ度数の平均 f は

$$f = \sum [1 \leq m \leq MP] f(m) \quad (4.12)$$

となる。 DP の総要素数に占める未知データの割合を未知データ率と呼び、 u と

替くと、

$$u = f / NP \quad (4.13)$$

となる。p(m) は m の巾のオーダーであり、#(Rm) は m の多項式のオーダーであるから、f(m) は m の巾のオーダーで小さくなる。そこで、MP' を MP より小さい適当な数として、

$$f \doteq \sum [1 \leq m \leq MP'] f(m) \quad (4.14)$$

と近似できる。

N/NP = 1/2 とした時の、b = 1 と 1/2 に対する f(m) の値を表 4.1 に示す。ここで、MP' = 5 と仮定した。ただし、#(Rm) と f(m) は表中の数値に MP を乗じた値である。

表 4.1 未知データ数の値

b	m	#(Rm)	p(m)	f(m)
1	1	1/2	1/2	0.250
	2	1/6	1/4	0.083
	3	1/12	1/8	0.031
	4	1/20	1/16	0.013
	5	1/30	1/32	0.005
1/2	1	3/4	1/2	0.375
	2	5/36	1/4	0.081
	3	7/144	1/8	0.018
	4	9/400	1/16	0.006
	5	11/900	1/32	0.002

表 4.1 より、f の値は、

$$f = 0.382 * MP \quad (b = 1) \quad (4.15)$$

$$f = 0.482 * MP \quad (b = 1/2) \quad (4.16)$$

となる。一方 NP は次の様にして求まる。(4.1)(4.4) 式より

$$NP = \sum [1 \leq r \leq MP] (MP/r)^b \quad (4.17)$$

である。この式の右辺を度数 m が MP' より小さい部分と大きい部分に分ける。

$$R' = \sum [1 \leq m \leq MP'] \#(Rm) \quad (4.18)$$

とおくと、

$$NP = \sum [1 \leq r < MP - R'] (MP/r)^b + \sum [1 \leq m \leq MP'] m * \#(Rm) \quad (4.19)$$

となる。(4. 19) 式の右辺第1項を NP' と第2項を NP'' とおく。 NP' を積分で近似すると、

$$NP' \doteq \int [1 \leq r \leq MP - R'] (MP/r)^b dr$$

$$= MP^b \frac{1}{1-b} \{ (MP - R')^{1-b} - 1 \} \quad (b \neq 1)$$

$$MP * \log (MP - R') \quad (b = 1)$$

(4. 20)

となる。 $b = 1$ と $1/2$ について、計算する。

$$MP' = 5 \quad (4. 21)$$

と置く。表4. 1より

$$R' = MP * 0.833 \quad (b = 1)$$

$$MP * 0.972 \quad (b = 1/2) \quad (4. 22)$$

となる。そこで、(4. 20) 式より

$$NP' = MP * \log (0.167 * MP) \quad (b = 1)$$

$$MP^{1/2} * \frac{1}{2} \{ (0.028 * MP)^{1/2} - 1 \} \quad (b = 1/2)$$

(4. 23)

一方、 NP'' は(4. 21) 式の下で

$$NP'' = 1.450 * MP \quad (b = 1)$$

$$1.325 * MP \quad (b = 1/2) \quad (4. 24)$$

となる。そこで、

$$NP = MP * \log (0.167 * MP) + 1.450 * MP \quad (b = 1)$$

$$MP^{1/2} * \frac{1}{2} \{ (0.028 * MP)^{1/2} - 1 \} + 1.325 * MP$$

(4. 25)

となる。 MP と NP の関係を $b = 1$ の場合について、表4. 3に示す。 $b = 1/2$ の場合は、表4. 3の範囲で、未知データ率は0.7程度であまり変化しない。

表 4. 3 異なりデータ数と延べデータ数、未知データ率の関係
($b = 1$)

M P	N P	u
1 * 1 0 ³	6. 6 * 1 0 ³	0. 1 2 6
1 * 1 0 ⁴	8. 8 * 1 0 ⁴	0. 0 9 4
1 * 1 0 ⁵	1. 1 * 1 0 ⁶	0. 0 7 5
1 * 1 0 ⁶	1. 3 * 1 0 ⁷	0. 0 6 2
1 * 1 0 ⁷	1. 6 * 1 0 ⁸	0. 0 5 3

実際には、M Pが増えるに従って、 b の値も大きくなるので、未知データ率はこれらの値より小さくなる。表 4. 4 に実際のデータによる $N / N P = 1 / 2$ の時の未知データ率を示す。ただし、上記議論では、処理対象を母集団とし、言語データベースを標本としたが、表 4. 4 では、言語データベース自体を母集団とし、そこから、抽出率 $1 / 2$ の標本を抽出する場合を仮定した。A D DとA Pのデータは単語データであるから、 $b = 1$ でモデル化でき、係り受けのデータは $b = 1 / 2$ でモデル化できる。しかし、いずれも、表 4. 3 やそのすぐ上の記述の値よりも実際のデータの方が小さくなっている。ここで、A P電のデータは[浦谷]によった。また、係り受けのデータは新聞データベース[江原2]によるものである。

表 4. 4 実際の言語データに対する未知データ率

データ名	M P	N P	f (1)	f (2)	f (3)	u
A D D	2,875	67,000	1,146	420	222	0.027
A P 電	75,000	3,500,000	25,000	?	?	0.007
係り受け 「を」格	101,666	167,281	76,166	16,100	3,600	0.576

5. おわりに

言語データベースから言語知識を抽出して、知識データベースを作成するときのデータの分布について考察し、度数の少ないデータが大量に存在する、いわゆるデータの過少性を示した。

この過少性に対処するためには、なんらかの単語のカテゴリー化が必要であろう。その1方法は[江原2]で提案されている。

参考文献

- [Atkinson] Martin Atkinson et al.: Foundations of General Linguistics, P.33, George Allen & Unwin, London, 1982.
- [Mandelbrot] Benoit Mandelbrot: On the Theory of Word Frequencies and on Related Markovian Models of Discourse, Proc. of the Symposium on Applied Mathematics, Vol.12, pp.190-219, 1961.
- [Sumita] Eiichiro Sumita et al: Example-Based Approach in Machine Translation, Proc. of Info Japan'90, Oct., 1990.
- [Witkam] Toon Witkam: Analogy-Based MT, Presentation for ATR, July, 1990.
- [浦谷] 浦谷則好: A P 電の単語分布、私信、1990年。
- [江原] 江原暉将ほか: A T R 対話データベースの内容、T R - I - 0 1 8 6、1990年10月。
- [江原2] 江原暉将: 単語のカテゴリーを用いた係り受け整合度の平滑化、T R - I - 0 2 1 5、1991年5月。
- [柿ヶ原] 柿ヶ原康二: 係り受け関係を用いた文節候補選択処理、T R - I - 0 1 0 5、1989年9月。
- [工藤] 工藤育男: 文と文の結束性を捕らえるための知識、T R - I - 0 1 3 4、1990年2月。
- [幸山] 幸山秀雄ほか: 用例を用いた依存関係単位の翻訳、41回情全大、4S-5、1990年9月。
- [佐藤] 佐藤理史ほか: 実例に基づいた翻訳、情研資、N L - 7 0 - 9、1989年1月。
- [古瀬] 古瀬蔵ほか: 変換主導型機械翻訳の実現手法、情研資、N L - 8 0 - 8、1990年11月。

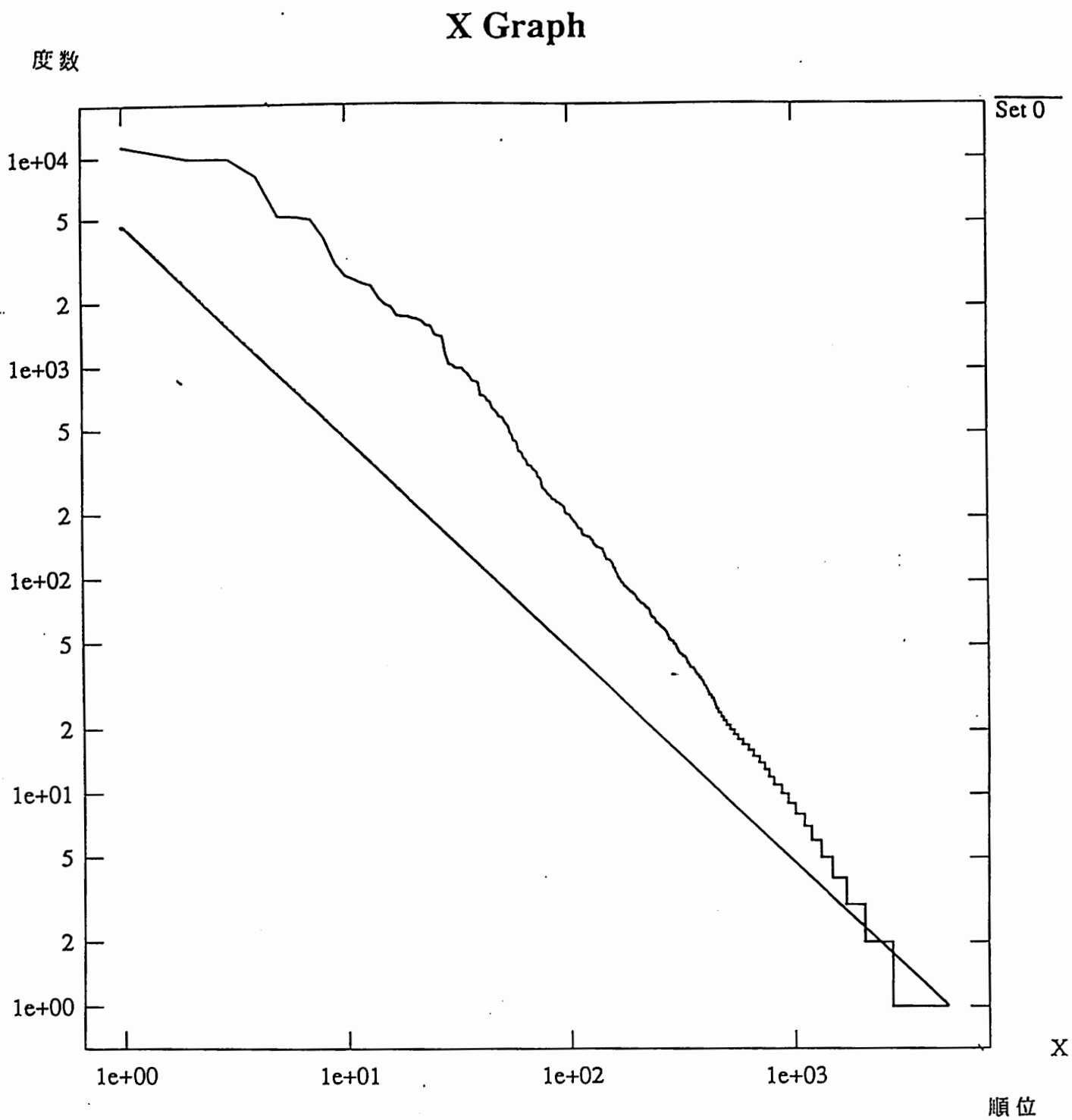


図 3. 1 単語数 1 の知識データ (単語データベース) の分布
 A D D 国際会議 電話対話
 直線は $b = 1$ の Mandelbrot 分布

gokei

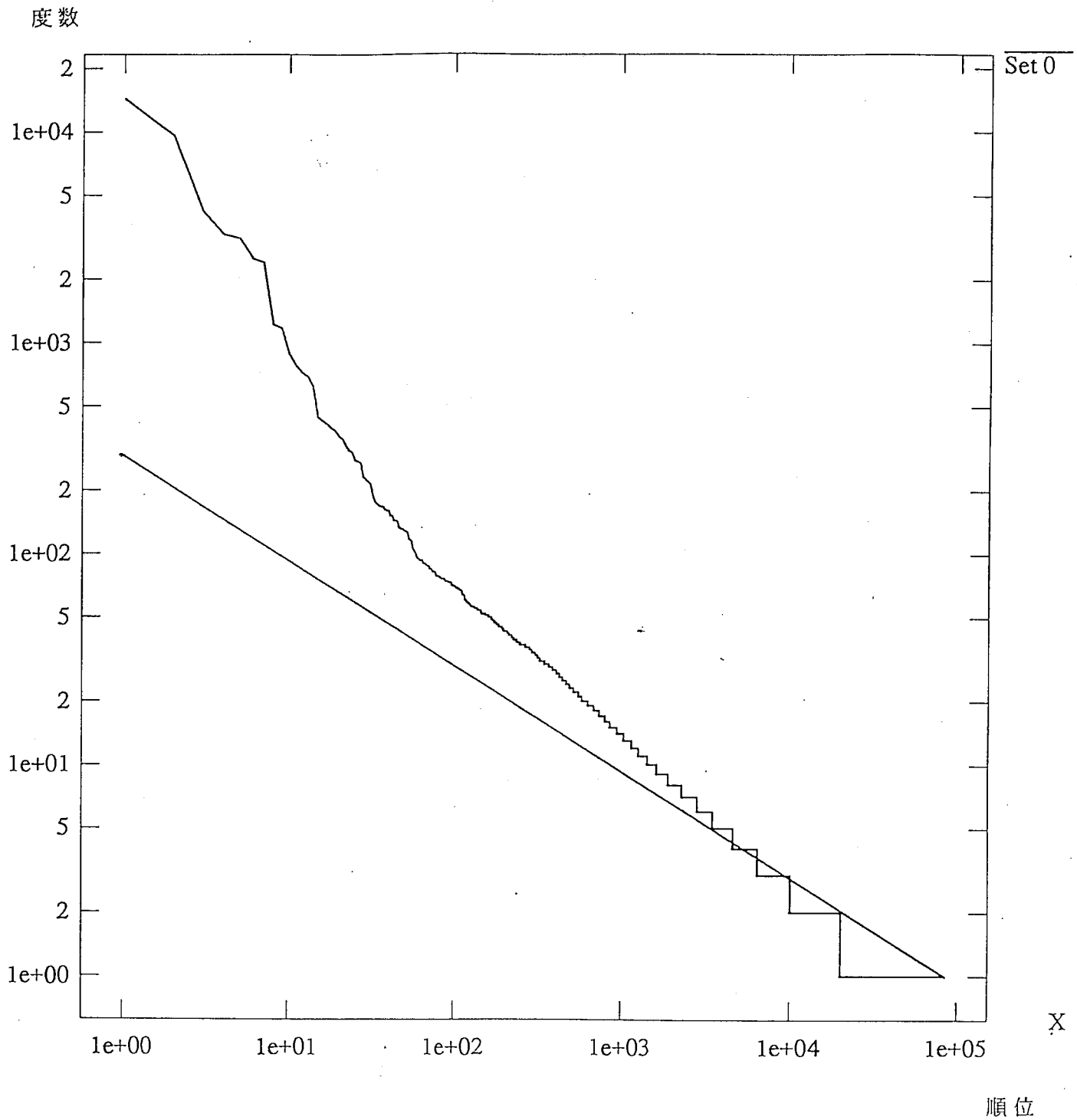


図 3. 2 単語数 2 の知識データ (係り受けデータベース) の分布
「が」格の係り受けデータ
直線は $b = 1/2$ の Mandelbrot 分布