

第5回 ACL European Chapter に参加して

Report on Fifth Conference of the European Chapter of the ACL

鈴木 雅実

Masami SUZUKI

ATR 自動翻訳電話研究所
ATR Interpreting Telephony
Research Laboratories

中村 順一

Jun-ichi NAKAMURA

九州工業大学
Kyushu Institute
of Technology

Abstract: The authors participated the 5th conference of the European Chapter of the ACL held in Berlin Apr. 11 to 13. In this article we report the trend of research activity in Europe and US. As a remarkable point we also introduce an initiative of "Reusable Linguistic Resources", on which a panel discussion took place in the conference. In addition, we briefly review the Second Bilateral Workshop for Computational Linguistics held in Manchester prior to the EACL.

1 はじめに

1991年4月9日から3日間にわたり、ドイツの旧東ベルリンで開催された、ACLのEuropean Chapter (EACL)の会議を聴講する機会を得たので、その概要を報告するとともに研究動向の分析を行なう。なお、これに先立つ4月4日、5日にイギリスのマンチェスター大学で開かれた、第2回日英計算言語学共同研究ワークショップ¹についても紹介する。

今回のEACLは、chairが旧東ドイツのJürgen KunzeとDorothee Reimannで、開催場所も旧東ベルリンの中心地である、アレキサンダー広場に面した会議場であった。当然のことながら、この会議は、90年10月の東西ドイツの統合以前から計画されていたものであり、会議の準備中に歴史的な変動があったことになる。この点については、Kunzeも、closing addressで、事務的な苦労話を含めて感慨深げにふれていた。会議とは直接関係しないが、自由に旧東西ベルリンの間を移動することにより感じられた、会議場周辺に色濃く残っている旧東ドイツらしさと繁華街などのある旧西ドイツとの対比が、筆者らやおそらく他の多くの参加者に、大きな印象を与えた。

以下では、まず、EACLの開催諸元、発表論文の分類と傾向、“Reusable Linguistic Resources”について、個々の発表で目についた論文の順に報告する。さらに、日英計算言語学共同研究ワークショップの内容を簡単に述べる。

¹これは、九州工業大学・情報工学部(代表・野村浩郷)とマンチェスター工科大学・計算言語学センタ(代表・Harold L. Somers, John McNaught)の共同研究プロジェクトの一環として開催されたものである。

2 ACL European Chapter の概要

2.1 開催諸元

- 開催期日 1991年4月9日 - 11日
- 開催場所 Kongresshalle, Alexanderplatz, Berlin
- 主催 Association for Computer Linguistics
- 資金協力 ヨーロッパ共同体委員会, ダイムラー・ベンツ財団, IBM ドイツ GmbH
- Chair Jürgen Kunze and Dorothee Reimann
Zentralinstitut für Sprachwissenschaft, Berlin
- 参加者 24ヶ国 から 265人
- 応募論文数 311 papers submitted,
25 late papers (10 days after),
- 採択論文数 286 papers reviewed,
50 papers accepted, 4 reserved
- 発表論文数 53 papers presented, 1 paper withdrawn

2.2 発表論文の分類と傾向

会議のプログラムは、初日の午前中以外、セッションAとセッションBに分かれて進行したため、筆者らが実際に聴講した発表件数は限られるが、論文全体を大まかに分類し、その傾向の分析を試みる。

- 言語現象の分析 1 件
Steve G. Pulman : Comparative and Ellipsis
- パーサ関係 7 件
David M. Magerman et.al. : A Probabilistic Chart Parser など
- ユニフィケーション・論理文法関係 6 件
Stephen J. Hegner : Horn Extended Feature Structure: Fast Unification with Negation and Limited Disjunction
- Tree Adjoining Grammar 2 件
Tilman Becker et.al. : Long Distance Scrambling and TAG など
- 語彙論・形態素論関係 4 件
Dan Tufis et.al. : A Unified Management and Processing of Word-Forms, Idioms and Analytical Compounds など
- 意味論関係 4 件
Elena V. Paducheva : Semantic Features and Selection Restrictions など
- 言語生成関係 6 件
Stephan Busemann : Structure-Driven Generation from Separate Semantic Representation など
- 学習・認知関係 3 件
Mori Rimon et.al. : The Recognition Capacity of Local Syntactic Constraints など
- 談話理論 3 件
Manfred Pinkal : On the Syntactic-Semantic Analysis of Bound Anaphora など
- 言語理解・知識モデル 4 件
Gudrun Klose et.al. : Modelling Knowledge for a Natural Language Understanding System など
- コーパスからの情報抽出 3 件
Sabine Berger : The Semantics of Collocational Patterns for Reporting Verbs など
- 機械翻訳関係 5 件
Bianka Buschbeck-Wolf et.al. : Limits of a Sentence Based Procedural Approach for Aspect Choice in German-Russian Machine Translation など
- 音声言語処理関係 4 件
Pete Whitelock : What kind of trees do we speak? – A Computational Model of the Syntax-Prosody Interface in Tokyo Japanese など
- 画像処理との接点 1 件
Wolfgang Wahlster et.al. : Designing Illustrated Texts: How Language Production is Influenced by Graphics Generation

何らかの解析文法理論に関わる発表が、約 15 件で最も多い。ユニフィケーションだけでなく、各種の論理文法についての理論的考察や構文解析の高速化などを扱ったものが中心となっている。また、言語生成に関する発表が少しづつ増えてきているのも最近の傾向であろう。今回は、質問応答などのアプリケーションを目的とした生成の話題が多いように感じられた。談話理解に関する発表は数件見られるが、特に目立つほどではない。機械翻訳をテーマにした発表は、さすがにヨーロッパらしく 4 件あるが、全体からすれば一割に満たない。音声、特に韻律情報の処理に関連する発表は、若干増加しつつあるように思われる。

国別に見ると、英国 13 件、フランス 3 件、ドイツ 13 件、イタリア 2 件、オランダ 3 件、その他のヨーロッパ諸国 7 件であり、ヨーロッパ以外では、ソ連 1 件、米国 7 件、イスラエル・コロンビア各 1 件であった。

3 Reusable Linguistic Resources について

辞書・コーパスなどの大規模自然言語データは、計算言語学・自然言語処理のあらゆる分野で非常に重要なものになっており、研究者の財産として、再利用・共同利用でできることが望まれている。これらの言語データを扱うためのツールも、同様に共同で利用できれば、不要な重複開発が避けられるので、望ましい。そこで、1987 年ころから Bell Core の Donald Walker が中心になって Text Encoding Initiative (TEI) というテキスト・辞書データの国際的蓄積・交換のプロジェクトが行われている。これ以外にも、ACL Data Collection Initiative など各種のプロジェクトが行われている。

そこで、EACL では、イタリアの Antoni Zampolli や Donald Walker などを中心となり、パネル討論会が開催された。これは、2 回に分けて行われ、まず、opening の後に、プロジェクトの進捗状況、問題点などの報告があり、同日の最後のスロットに、その議論をふまえた討論が行われた。パネル討論の参加者は、次のような人々である。

1. Nicoletta Calzolari (Pisa Univ.)
2. Ulrich Heid (Univ. Stuttgart)
3. Remi Zajac (Univ. Stuttgart)
4. Nancy M. Ide (Vassar College)
5. Jasper Kamperman (EUROTRA)
6. Susan Warwick-Armstrong (Univ. de Geneve)

報告・議論は、1990 年 8 月の COLING90 の際に開催された「文章および辞書資源に関するワークショップ」を踏まえたものである。自然言語処理にとって”do it yourself”の時代は終わった、などとの意見も出されたが、残念ながら特に目新しい議論にはならなかった。Eurotra-7 (辞書・専門用語資料の再利用度調査プロジェクト)、Eureka Genelex (欧州語の一般辞書構築プロジェクト) などの進捗状況の報告もあったが、詳細な議論にはいたらなかった。

再利用, 特に, 共同利用は, 重要なテーマであるが, その具体的な実現は容易ではないことが感じられた。また, 現状では, 欧州語が中心であるので, 日本語などの多字種言語の扱いについても更に検討が必要であろう。なお, 言語データ (大量のコーパスなど) のを利用した研究に関する発表もいくつかあったが, これについては, 4.2 節で紹介する。

4 個々の論文の紹介

4.1 パーサ関係

- David M. Magerman & Mitchell P. Marcus
Pearl: A Probabilistic Chart Parser
- Tsuneko Nakazawa An Extended LR Parsing Algorithm for Grammars Using Feature-Based Syntactic Categories
- Gosse Bouma Prediction in Chart Parsing Algorithms for Categorical Unification Grammar

パーサ関係の論文として, 筆者が興味を持った発表の内, これら 3 件について紹介する。最初のもは, 前節で紹介した言語データを活用する方法の一つとして興味深いものであり, 他の 2 つは, 素性構造ベースのパーサの高速化に関する話題である。

Magerman は, 文解析における品詞および構文的曖昧さの解消に統計的確率を使用する手法とその実験結果について発表した。解析アルゴリズムとしては, Earley-type の top-down 予測を用いる bottom-up chart parser である。統計的確率に基づく点数計算により, 品詞の決定を含め, 「最も良い」解から解析を進めるが, その他の解も平行して求められるようにしておき, 予測がはずれた場合には, 次善の解から計算を進めるものである。

興味深い点は, その点数計算にある。以下の 3 点が考慮されている。(1) これまでの実験の結果から, ルールの長さは点数に無関係である。(2) 訓練集合にあらわれなかった「まれな」事例を排除しないで, 一定の点数を与える。(3) Earley-type の top-down 予測が「まれな」事例であっても排除してはいけない。これらを踏まえ, 単語の品詞列の 3 組の頻度と top-down 予測を考慮した点数から点数を計算している。

この点数計算は, 特に, 対象分野を限定すれば効果的であるとの結果が示されている。たとえば, “Fruit flies like a banana.” という文に対して, 昆虫学の分野では, “flies /verb” の点数が “flies /noun” よりも低いだけでなく, “NP → noun noun” のような名詞連続の規則の方が単独の名詞よりも点数がよいという点から “fruit flies/NP” が選択できる。また, PP attachment については, 「方向判定」という特定の分野では, “from” や “to” は, ほぼ, 100% 動詞に係るというデータから非常に高い精度で判定可能である。

Nakazawa は, LR 法の構文解析アルゴリズムを素性構造に拡張する手法について発表した。LR 法では, あらかじめ規則から LR 表を作成するが, Tomita などの従来の手法では, この際に使用するのは, 単独のカテゴリシンボルであり, 解析で実際に使用する素性構造については, カテゴリレベルの解析が成功した場合に単一化を行うことにより扱っている。

これに対して Nakazawa は, LR 表を作成する段階から素性構造を考慮した表を作成するアルゴリズムを示した。カテゴリから素性構造に拡張する際に問題となる点の一つに, $V \rightarrow V NP$ のように素性構造が再帰的に作成される規則がある。実際の解析の際には, V の subcat により作成される構造が無限になるわけではないが, LR 表を作成する場合には, 具体的な subcat を与えることができないので, 対処が容易でない。

Shieber は, 再帰的に作成される構造の深さに制限を与えることにより, この問題に対処したが, これに対して, Nakazawa は, LR 表の項目の作成を工夫することにより対処した。これは, Shieber の方法とカテゴリシンボルだけを用いる方法との中間的なものになる。

現在の文法の表現方法が, カテゴリレベルの規則を詳細化するよりも, 語彙記述の素性構造を詳細化する方向である。このため, カテゴリシンボルだけによる top-down 予測は, 解析の効率化にあまり寄与しなくなっている。そこで, 素性構造を含めて予測する手法は興味深い。

従来のカテゴリレベルの解析の高速化手法が素性構造を基礎とする単一化文法では有効でない点に関して Bouma は, Categorical Unification Grammar (CUG) の場合について発表した。Categorical Unification Grammar では, ほぼすべての情報が語彙に記述されることになり, 規則としては, $X_0 \rightarrow X_1 X_2$ で, X_1 が head になるものと X_2 が head になるものの 2 つが中心である。このような規則に対して top-down 予測を行ってもほとんど有効でない。このため, より情報の多い語彙から出発する bottom-up 解析が一般に行われる。

しかし, Bouma は, その場合, 以下の 3 つの問題があることをまず指摘した。(1) 無駄な予測が行われることが多い。(2) 品詞の曖昧さが多いと, (潜在的に) 無駄な解析が行われる。(3) 先に示したの一般的な規則だけでなく, 文法規則に制約が書かれている場合, 効率がよくない。そこで, Bouma は, top-down 予測を導入する手法を提案し, 実際に実験結果を示した。

top-down 予測自体は, Left-corner 表により行う方法であるが, この表を作成する際に, ルールだけを調べていては, ルールが一般的であるため, ほとんど予測ができない。そこで, Bouma は, instantiated rule という考え方をを用いた。これは, 語彙カテゴリを一般的規則に代入することにより, 規則をより具体化しておくものである。直観的には, CUG の一般規則と語彙とを組み合わせ, 通常の CFG 規則に近いものを作成しておくことになる。実験結果によると, この方法は, 他の方法にくらべて良い結果が得られている。

4.2 コーパスを利用した研究

- **Sabine Berger** The Semantics of Collocational Patterns for Reporting Verbs
- **Michael R. Brent** Automatic Semantic Classification of Verbs from their Syntactic Contexts: An Implemented Classifier for Stativity
- **Jean Véronis and Nancy Ide** An Assessment of Semantic Information Automatically Extracted from Machine Readable Dictionaries

これらの論文は、いずれも機械可読なテキストリソースなり辞書から情報を抽出しようとする試みであり、その際に問題となる点が挙げられている。

Berger は新聞記事を題材に、いわゆる reporting verb (報告動詞) が、story と文脈を再構成する役割を果たしていることを示している。大きなコーパスにおける意味的な共起の分析から、語彙構造を提案する試みでもある。ニュース記事においては、純粋な情報 (pure information) とメタ的な情報を区別する必要があり、意味的な共起が辞書に記述されるべきであるとしている。

一例として、動詞 *insist* をとると、一連の Wall Street Journal の記事に出現する 477 回のうち、428 回には、*but* や *not* などマークされる対立表現が共起している。既存の辞書では、LDOCE の *insist* の語義説明に、“to declare firmly (when opposed)” と明記してあるが、一般に記述されることは少ない。従って、コーパス分析の結果から、A propositional opposition is implicit in the lexical semantic of *insist* という訳である。また、これらの動詞は、主語として、個人だけでなく、組織や団体など擬人的なものをとるが、コーパスを分析して、3 種類の素性により分類した結果が示されている。主語の NP が、1) single person 2) group people 3) institution のどれに相当するかを、各動詞ごとに分布を調べ、degree of animacy として特徴づけている。対象となった動詞は次のようにカテゴリ化された。

admit, denied, insist : -inst,

announce, claim : +person, *said, told* : other

さらに、この種の分析を基に、ニュースソースの名詞句用の文法を定義し、Lexical Conceptual Paradigm として実際の構文解析に役立てることを企てている点が興味深い。

Brent の研究は、コーパスのみを用いて、出現する動詞が、stative であるか、event を表すものであるかを判定しようとするものである。ここで使われる情報は、統語的な環境だけである。すなわち、進行形をとるか、程度を表す副詞によって修飾されるか、といった条件である。目指すところは、自動的に言語データベースを構築することである。結果は、かなり良い精度で判定に成功している。ただし、“think” のように、連接する要素によって意味が分かれるようなものもある。

think that (stative), *think of, think about* (event)

Véronis らの論文は、既存の機械可読な辞書から、階層関係情報を自動抽出する際の手法について考察している。使用された辞書は CED, OALD, COBUILD, LDOCE および W9 (Webster) の 5 種類の英語辞書である。個々の辞書を単独で使用した場合には、55 - 70 % の誤りが生じてしまうのに対して、5 種類の辞書からの情報をマージすると、誤り率は 6 % 低下した、と報告している。

これらの研究に見られる手法は、大量の言語データを利用して、可能な限り自動的に有用な情報を抽出しようとするものであるが、Berger の論文が示唆するように、その結果を各種の言語論的パラダイムに、どのように取り込んで行くか、また、その逆のフィードバックとして、何に着目した言語データの分析がなされるべきか、という点が今後の焦点となろう。

4.3 談話理解関係

- **Pete Whitelock** What sort of trees do we speak? -A Computational Model of the Syntax-Prosody Interface in Tokyo Japanese
- **Eric Bilange** A Task Independent Oral Dialogue Model

Whitelock の論文は、韻律情報を用いた、構文解析木の曖昧性解消のアプローチである。統語情報と韻律情報を素性構造上とともに扱う試みである。ここでは日本語 (標準東京) のアクセント、すなわちピッチレベルの高低による音韻的制約を直接、統語構造の決定に利用している点が興味深い。文節を韻律の単位 (prosodic unit) としている。一般に、統語構造と韻律構造は異なるので、左分枝の統語構造であっても、韻律的な構造がそれと一致するとは限らない。この研究では、2 文節・3 文節・4 文節の句について、統語的な分枝が成立するときの、ピッチの変化の違いを Categorical Unification Grammar の枠組を基本とする、素性の違いとして記述している。

現在のところ、韻律情報の正確な検出は難しく、音声認識上の重要な研究課題となっているが、この発表が示唆するように、韻律情報を統語解析に役立ようとする研究は、今後一層活発になるものと思われる。

Bilange の論文は、タスク指向型の人間-機械の対話モデルにおいて、対話知識とタスクの知識を明確に分離し、タスクによらない対話モデルの構築が可能であることを示そうとするものである。この研究は SUNDIAL ESPRIT プロジェクトの中で、システムの談話管理モジュールの開発として行なわれている。ここで採用されている対話モデルでは、4 つの decision layer が定義されている。

1) Rules of conversation 2) System dialogue act computation 3) User dialogue act interpretation 4) Strategic decision level

英語・フランス語・ドイツ語によるプロトタイプを作成し、フライトの予約、フライトに関する問合せ、列車の時刻表の問合せの 3 種類のタスクについて、対話モデルの有効性を検証中である。

5 日英計算言語学共同研究ワークショップ

概要

このワークショップは、九州工業大学情報工学部とマンチェスター工科大学 (UMIST) の計算言語学研究センター (CCL) とのイニシアティブによって始められた、日英共同の自然言語処理に関する研究会であり、第一回は昨年秋に九工大で開催されている。今回の参加者は約 30 名、そのうち発表者は 13 名であった。以下では、これらを分類整理し、幾つかの発表内容について、その特徴を述べる。

5.1 翻訳関係

- **Danny Jones** Example-Based Machine Translation

最近注目され始めている用例主導型の機械翻訳のアプローチの利点を述べたものである。京都大学の佐藤やオランダの Sadler, ATR の隅田の論文を多く引用し、用例ベースの手法の重要性を説いている。[3]

- **Jun-ichi Tsujii and Kimikazu Fujita** Lexical Transfer based on Bilingual Signs: Towards interaction during transfer

言語間の距離が大きい場合に問題となる、構造変換の問題を取り上げ、これを Bilingual Sign という論理的な述語を用いて、見通し良く記述しようとする試みである。例えば、

(1) The teacher runs the program.

(2) The teacher runs the company.

という例文に対して、2 つの Bilingual Sign

[RUN:JIKKOUSURU] と [RUN:UN'EISURU] を定義する。もう少し複雑な例では、

(私が) なんとか 論文を 仕上げた。

←→ I managed to complete [the/a] paper.

この場合、[MANAGE:NANTOKA] という Bilingual Sign を定義することになるが、品詞が異なるため、その内容は次のようになる。

```
(Def-Pred [MANAGE:NABTOKA]
  arg1:=.
  arg2:=[event:dekigoto],
  eng :=head :=e-lex := manage,
      agt := <! arg1>,
      evt := <! arg2>,
  jpn :=<! arg2>,
      arg := <! arg1>,
      /adv := head := j-lex := nantoka)
```

なお、同テーマの論文が、EACL でも発表された。

- **John Phillips** Linguistic Context and Machine Translation

機械翻訳における compositionality の問題に深い関わりのある、文脈の言語内の側面を取り上げている。

- **Masami Suzuki** Lexical Choice in Dialogue Translation

ATR で音声言語翻訳の対象領域としている、「国際会議への参加申し込み」の日英対訳コーパスにおける、動詞の訳語の分析をもとに、対話翻訳の問題点を考察したものである。通訳者によって対訳がつけられた対話コーパスでは、一つの日本語動詞に対して、実に多様な英語表現が与えられている。その分布を調査することによって、個々の動詞の特性を把握するとともに、会話文における訳語選択のストラテジーを見い出すことをねらいとしている。特に、訳語のタイプに関して、意味的な階層的な関係、状況依存性、視点の変化などによる同義関係の分類を行ない、会話文における柔軟な訳語選択への課題を考察している。

5.2 辞書関係

- **Iris Arad** Automatic generation of bilingual dictionaries

コーパスを利用して、特定の領域と部分言語を対象とした機械翻訳のための、2ヶ国語辞書を自動作成する手法についての発表である。

- **Masasuke Tominaga, Seiji Miike, Hiroshi Uchida and Toshio Yokoi** Development of the EDR Concept Dictionary

電子化辞書研究所における、概念辞書の開発について述べたものである。[4]

- **Blaise Nkwentz Azeh and Kerry Maxwell** The Eurotra (UK) Monolingual Dictionary

EUROTRA プロジェクトにおける、辞書 (というよりは、基礎データ) の内容についての発表である。対象としている分野は、衛星通信に関する文書である。

5.3 言語解析関係

- **Vincent Tanda** Coordination in Bafut ... and Japanese

等位構造に関する、Bafut 語と日本語の類似性 (英語に対して) を述べ、各種の等位構造の構文解析手法を提案している。

- **James Au-Yeung** Parsing Mixed Left- and Right-branching Structures

英語に代表される右分岐型の言語や、日本語などの左分岐型言語のパーズ手法に対して、両方の言語類型に適用可能な一般的なパーズの枠組が研究されているが、中国語のような混合型の言語の場合、問題点が多い。この発表では、人間の言語使用における心理言語学的制約に基づくプリファレンス (Minimal Attachment や Late Closure) が、英語には適用できても、中国語には向かないことを示し、これを解決するための加算的な構文解釈アルゴリズム “Fine Grain Incremental Interpretation” を提案している。

- **Jeremy Lindop** The Analysis of English in Eurotra 多言語間の翻訳を目的とした、EUROTRA プロジェクトの中での、英語の解析について報告している。解析結果は ECS(Eurotra CONSTITUENT Structure) から、IS(INTERFACE Structure) を経由して、他の言語へ変換される。

5.4 対話理解その他

- **Chris Bowerman** Writing and the computer: the nature of the problem and an intelligent Tutoring System solution

外国語（この場合はドイツ語）学習者のための一種の CAI システムにおける、計算機 = Tutor と学習者の対話の制御に関する問題点をテーマとした発表である。扱っている主題は、文を書く行為のモデル化である。ものを書く行為は、目標に基づく一連の制約を満足することと見なし、Prewriting/Planning, Generating, Revising の3つの過程と定義している。

- **Jun-ichi Nakamura, Yuko Den, Yasuharu Den and Sho Yoshida** Towards a Sentence Generation based on the Point of View

視点と言語表現との関係に着目し、この観点からみて自然な文だけを生成する日本語文生成システムについての発表である。視点と言語表現の関係を定式化するため、久野の視点の理論を参考にしてている。これは、同じ意味内容を持つ様々な形式の文を自然なものとは不自然なものに分ける原則として記述する。例えば、受動文では二格よりもガ格の方に視点がおかれるので、本来視点がおかれるべき話者を二格に持つ受動文は不自然である。受動文・授与動詞・使役文などの言語表現と視点の関係に関する原則について考察し、これらの原則を Lexical Functional Grammar を用いた文生成システムに適用した。原則の記述は、語彙に関する制約として表現されている。

- **Hirosato Nomura and Hideki Iwamoto** Linguistic Model of Law Sentences and Computer Processing

法律文の言語モデルに基づき、言語情報と論理情報を同時にコード化するための表現手法を提案している。言語モデルは、いわゆる制限言語として定義されている。論理情報は法律上の事象を推論するために用いられる。この枠組を基に、法律知識ベースの構築を目指している。[5]

謝辞

九州工業大学の吉田 将 情報工学部長ならびに野村 浩郷 教授、ATR 自動翻訳電話研究所の樽松 明社長、森元 逞 データ処理研究室長には、今回の国際会議参加の機会を与えて頂いたことに感謝します。また、EACL の registration desk で親切に対応してくれたベルリンの人々、マンチェスターでお世話になった方々に、この場を借りてお礼を述べたいと思います。

参考文献

- [1] Proceeding of the European Chapter of the ACL, 1991
- [2] Computational Linguistics Research at KIT and UMIST, 1991
- [3] E. Sumita et.al. Translating with Examples: A New Approach to Machine Translation, *The Third MT Conference*, 1990
- [4] H. Uchida, Electronic Dictionary, *International AI Symposium 90 Nagoya*, JSAI, 1990
- [5] H. Iwamoto and H. Nomura, Linguistic Model of Law Sentences and its Application to the Natural Language Processing, *Japan-Australia Joint Symposium on NLP*, to appear 1991