

チャートパーザによる音声認識候補の効率的解析手法

田代 敏久 竹沢 寿幸
ATR 音声翻訳通信研究所
永田 昌明
NTT 情報通信網研究所

音声認識部と言語処理部のインターフェースとしてN-BESTインターフェースを採用した音声言語システムでは、言語処理部は音声認識候補のフィルタリングと入力の意味解釈の決定という2つの役割を持つ。本報告では、言語処理部にチャートパーザを用いて、複数の音声候補の共通部分の再計算を避け、もっともらしい候補を早期に発見し、入力の意味解釈の決定を行う方法を提案する。音声認識候補の再順序付けのための知識として、単一化文法で記述された統語的・意味的制約と文脈依存確率文法を用いて音声認識候補の解析実験を行い、共通部分の共有による効率向上と認識候補の再順序付けによる精度向上を確認した。

Efficient Chart Parsing of Speech Recognition Candidates

Toshihisa Tashiro Toshiyuki Takezawa
ATR Interpreting Telecommunications Research Laboratories
Masaaki Nagata
NTT Network Information Systems Laboratories

In a spoken language system that uses the N-best interface between the continuous-speech recognition(CSR) component and the natural language processing(NLP) component, the role of the NLP component is to filter CSR candidates and determine the semantic interpretation of the input. In this paper, we propose an efficient chart-parsing method of finding the most plausible candidate while avoiding the re-computation of the common substrings. In preliminary tests, our parser has been successful in reducing parsing steps by sharing common substrings and in selecting correct candidates first, using syntactic and semantic constraints described in the unification-based grammar and context-sensitive conditional probability CFG preferences.

1 はじめに

音声翻訳システムのような音声言語システムの構築には、音声認識部と言語処理部の効果的な統合手法の研究が重要である。音声認識部と言語処理部の統合手法は、ネットワークモデルのように音声認識と言語処理が一体化した密結合型と、階層モデルのように音声認識と言語処理が独立したモジュールを形成する疎結合型とに分けられる。密結合型のシステムは、すべての知識を統一的・集中的に利用できるという利点があるが、処理効率や各モジュールの独立性を考えると、疎結合型の方が優れている。

疎結合型のシステムの場合、音声認識部と言語処理部間で受け渡されるデータをどのように設定するかが問題となる。現在主流となっている方法は、音声認識部は複数の候補(上位 N 個)を出力し、最終的な解の決定を言語処理部に委ねるという方法である(N-BEST インターフェース)[11]。この方法は、

- 言語処理部の豊富な知識を、絞りこまれた候補に対してのみ適用するので処理効率がよい。
- 正解が N 候補の中に含まれていれば、たとえ第一候補が誤りであった場合でも、言語処理部で回復できる可能性がある。

等の利点がある。

音声翻訳システムのような音声言語システムを N-BEST インターフェースで設計した場合、言語処理部に求められる機能としては、以下の 2 つが考えられる。

- 音声認識候補をフィルタリングする機能。統語的・意味的に不適格な候補を排除したり、語用論的情報から判断してよりもっともらしい候補を優先的に選択したりする必要がある。
- 入力正しい(もっとももっともらしい)意味解釈を決定する機能。一般的に、一つの入力に対して可能な意味解釈は複数存在するが、言語処理部では最終的にはもっとももっともらしい解釈一つに決定する必要がある。

これらの機能を実現するために、言語処理部は高度な統語的・意味的な知識や語用論的知識、あるいはタスクに関する深い知識等を利用して処理を進める必要がある。一般にこれらの知識の適用は計算コストが大きいので、音声認識部と言語処理部のインターフェースの設計の際には、処理効率について十分に考慮し、できる限り小さい計算量で適切な音声認識候補を選択し、さらに選択した音声認識候補のもっともらしい意味解釈を決定できるようにする必要がある。

本報告では、効率的な音声認識部と言語処理部のインターフェースを実現するための手段として、アク

ティブチャートパーザによる N-BEST 音声認識候補の解析手法を提案する。この手法は、

- 複数の音声認識候補の共通部分を共有した初期チャート(ラティス構造)を作成し、同一語句の再計算を回避することにより効率的な解析を行う。
- 経験則的・統計的ヒューリスティックを用いたアジェンダ制御により、複数の音声認識候補を再順序付けし、もっともらしい意味解釈を早期に発見する。

という特徴を持っている。

2 N-BEST インターフェース

最近の音声言語システムは、音声認識部と言語処理部のインターフェースとして N-BEST インターフェースを採用しているものが多い。N-BEST インターフェースとは、

- まず、計算コストが小さい知識(音響モデル、単純な統語モデル等)を用いて候補の絞り込みを行い、上位 N 個の候補を出力する。
- その後、より計算コストがかかる知識(複雑な syntax、semantics 等)を用いて、候補を再順序付け(reorder)する。

というテクニックである。

N-BEST インターフェースの音声言語システムを作成する場合には、以下の 4 つの点について研究する必要がある。

- 音声認識部での N-BEST 探索手法:

正確な N-BEST を出力するのは難しい問題である。また N の数をどれ位に設定すればよいのか、ということも実証的に検討する必要がある。

- N-BEST 候補のデータ構造:

連続音声認識で N-BEST の候補を出力した場合、各候補には共通部分が含まれることが多い。この場合、単語グラフ(単語ラティス)を受け渡しのデータ構造とすれば、言語処理部の処理の負担を軽減できる[2, 7]。また、受け渡しのデータ構造を単なる文字列とするか、形態素単位とするか、ということも重要な問題である[10]。

- 各フェーズで用いる知識:

ディクテーションのような音声言語システムでは、入力の深い意味解釈等はあまり必要ではない。そのため、音声認識部では単語の bigram を用い、言語処理部では trigram を用いるというように、比較的浅い知識を用いる場合が多い。この場合、 N の数は比較的大きくすることが可能かつ必要で

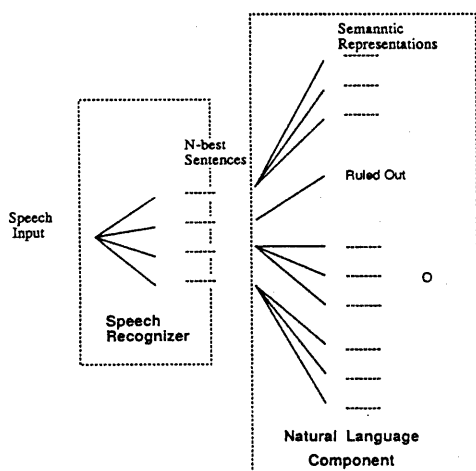


図 1: 音声認識システムの探索空間(木)

ある。一方、音声翻訳システムのような音声言語システムでは、音声認識部でも CFG のような比較的高度な統語モデルを用い、言語処理部では単一化文法のような高度な言語知識を利用して意味解釈を行う場合が多い。この場合には、N の数は比較的小さくすることが可能かつ必要である。

● 言語処理部での N-BEST 候補の処理方法:

言語処理部では、入力された N-BEST 候補に対して、より高度な知識を用いて再順序付けを行う。この場合、複数の候補をどのように探索するかが問題となる。特に、言語処理部で深い意味解釈を行う場合には、できるかぎり効率のよい探索手法を用いる必要がある。

従来の N-BEST インターフェースに関する研究は、主として音声認識研究者によって行われてきたため、言語処理部での N-BEST 候補処理方法に関する研究は少ない。また発表された研究においても、

- N-BEST 候補処理に有効なヒューリスティックな選好をほとんど用いていない [2, 9]。
- 実際の音声認識候補の処理手法や実験結果が示されていない [5]。

等の問題点があった。

本研究の目的は、ヒューリスティックな選好をチャートパーザのアジェンダ制御に用いることにより、実際の音声認識候補の効果的な解析が可能なることを示すことである。

3 言語処理部に必要な知識

3.1 統語的・意味的制約

入力の統語的・意味的な適格性を判断するために用いられる知識である。一般の言語処理システム(文字入力前提となっているシステム)では、入力の統語構造や意味解釈を決定する際に、不適格な構造や解釈を棄却するために用いられる。音声言語システムにおいては、不適格な認識候補を棄却(rule out)するためにも用いることが可能である。

たとえば、次のような音声認識候補の例では、格助詞「は」が欠けており、日本語の統語的制約に違反しているため、言語処理部によって棄却できる。

こちら会議事務局です
(正解: こちらは会議事務局です)

また、次のような例は、統語的制約は満たしているものの動詞「持つ」の「が格」(主格)に名詞「登録用紙」を取ることは意味的制約に違反しているため棄却できる¹。

登録用紙が既にお持ちでしょうか
(正解: 登録用紙は既にお持ちでしょうか)

我々は、このような統語的・意味的な制約を記述するための枠組みとして単一化文法を用いている。単一化文法は、制約が宣言的に記述されているため

- 文法の保守や拡張が容易である。
- 制約の適用方法が自由であり、様々な戦略を用いることが可能である。

等の利点がある。

3.2 ヒューリスティックな選好 (preference)

入力のもっともらしさを判断するために用いられる知識である。一般の言語処理システムでは、入力の統語構造や意味解釈の曖昧性の解消に用いたり、もっともらしい解釈を早期に決定するために用いることができる。音声言語システムにおいては、複数音声認識候補からもっともらしい候補を選択するためにも用いることが可能である。

これは、さらに次の2つに分類できる。

1. 内省 (insight) に基づく選好:

言語学的・心理学的な見地から、文の読み(理解)の選好(偏好)のモデルを明らかにできれば、それ

¹このような選択制限が常に制約として記述しきれるのか、という疑問は当然存在する。本来はこれも制約ではなく、選好として扱うべきであろう。特に、自由発話を扱うことが可能な頑健な解析システムを作成するためには、制約の段階的緩和等の処理を行なう必要があるため [12]、選択制限を制約として使うのはさらに難しくなるだろう。

を文の解釈の決定に利用することが可能である。
この種の選好としては、以下のようなものが挙げられる。

- Close Attachment:
ある語句の修飾先は一般に複数ある(曖昧性がある)が、一般的にはもっとも近い語句を修飾することが多い、というヒューリスティックである。
- ゼロ代名詞最小:
一般にゼロ代名詞(必須要素の欠落)は少ないほど、文としてのまとまりがよい、というヒューリスティックである。

2. コーパスから統計的に得られる選好:

コーパスから自動的に抽出可能な知識である。この種の選好として用いることが可能な統計情報としては、以下のようなものが挙げられる。

- 単語(形態素)の統計情報:
単語(形態素)の出現頻度、単語(形態素)の連鎖情報(マルコフモデル)等がある。
- 共起関係の統計情報:
単語や句(文節)の係り受け情報等がある。
- 文法規則(CFG 規則)の統計情報
確率文法(書き換え規則が使われる条件を考慮しない・するにより大きく2つに区分できる)等がある。

以下のような例では、どちらの文も統語的・意味的制約違反は起こしていないが、2番目の候補がよりもっともらしいと判断するのが妥当である。

- 1 待つ登録用紙で手続きをしてください。
- 2 まず登録用紙で手続きをしてください。

内省に基づく場合には、一番目の文は、「待つ」という語の必須格が埋まっていないことが不自然である、等と判断することができる。一方、統計情報による場合には、「待つ」という語がこういった文脈で出現することが稀だということにより、2番目の文がよりもっともらしいと判断するわけである。

4 音声認識候補の再順序付け(reorder)手法

N-BEST インターフェースを採用した音声処理システムは、図1に示すような探索空間(木)の探索問題として規定することができる。よって、言語処理部によるN-BESTの音声認識候補を処理する方法としては、次のような方法が考えられる。

- 統語的・意味的制約のみを用いた再順序付け(深さ優先探索)

もっとも単純な手法である。言語処理部は、音声認識部から渡されたN-best候補の上位の候補から処理を試みる。もし候補が統語的・意味的制約に違反していなければ、その候補の意味解釈を決定し、言語処理部の最終的な解とし処理を終了する。もし制約に違反している場合には、次候補を選択し処理を続行する。

この方法の欠点として、上位の候補が制約を満たしておりかつ下位の候補が正解の場合、下位の候補を選択する基準が音声認識候補の尤度以外に存在しない、ということが挙げられる。逆に利点としては、音声認識部の一位正解率が高い場合、この方法は無駄な探索をしない、ということが挙げられる。

● 選好を用いた再順序付け-1(全数探索)

選好に基づくスコアを意味解釈に付与することにより、制約だけでは区別できない音声認識候補も順序付けが可能となる。これを実現するもっとも簡単な方法は、1)すべての候補を処理し、可能な意味解釈をすべて出力する、2)処理途中に、経験則的・統計的選好に基づくスコアも同時に計算しておき、最終的に最も高いスコアを得たものを解とする、という方法である。

この方法の欠点として、言語処理部(特に単一化文法等を用いた高度な言語処理)は、計算コストが大きいプロセスなので、すべての候補を処理する(すべての探索空間を辿る)のはメモリ量、計算量の点で困難である、ということが挙げられる。逆に利点としては、スコア付けに用いた評価基準による最適解が保証されていることが挙げられる。

● 選好に基づく再順序付け-2(ヒューリスティックな探索)

選好に基づくスコアを言語処理過程(通常は構文解析部)でヒューリスティックとして利用し、すべての探索空間を辿ることなく適切な解を発見しようとする方法である。

この方法は、選好に基づくスコアが良好なヒューリスティックとして機能するならば、探索空間を減少させることにより効率的な処理が可能となる。半面、もしスコアが良好なヒューリスティックとして機能しない場合には、求める解の信頼性は低くなってしまふ恐れがある。

いずれの探索方法を採用するにせよ、N-BEST 音声認識候補にはかなりの共通の部分文字列が含まれているので、共通部分の再計算を避ける仕組みが必要である。また、経験則的・統計的選好に基づくスコアを探索時のヒューリスティックとして利用する場合には、N-BEST 候補が一つのデータ構造としてまとまってお

り、かつ言語処理(解析処理)の過程を制御できる必要がある。

5 アクティブチャートパーザを用いた音声認識候補解析

本節では、アクティブチャートパーザによる音声認識候補解析手法の実際について説明する。

5.1 共通部分文字列を共有した初期チャートの作成

図2に示すように、N-BEST 音声認識候補には多くの共通部分文字列が含まれることが多いので、効率的な処理を行なうためには共通部分の再計算を防ぐ仕組みが必要である。

チャートパーザはもともと解析過程で同一句(well formed substring)の再計算を避けるために、チャートというデータ構造を採用したパーザなので、複数の候補を共有した初期チャートを作成すれば、共通部分の再計算を自然に避けることが可能である。

そこで、我々は図2に示すような音声認識候補をもとに、図3に示すような初期チャートを作成するプログラムを作成した。

初期チャートを構成する各弧には、入力のシンボルと文の番号のリストが付与されている。文番号のリストを用いることにより、入力のN-BEST 候補に含まれない語句の生成(図2の例では「こちらに会議事務局でした」)を防ぐことができる。また、この文番号により、音声認識候補のスコアをアクセスすることも可能となっている。

5.2 言語モデル

チャートパーザのアジェンダを、選好に基づくスコアにより制御し、もっともらしい解を早期に発見しようとする研究は数多く行なわれている[5, 6, 1]。アジェンダ制御に用いられる知識には様々なものがあるが、確率文法、特に書き換え規則が適用される条件(文脈)を考慮した(文脈依存の)確率文法は、

- コーパスから自動的に抽出できる。
- 自然言語の持つ文脈依存性を捕えることが可能である。

等の点で有望と考えられる。

本研究では、北[4]により提案された言語モデル(rule-bigram)をアジェンダ制御のための知識として採用している。この言語モデルでは、ある木の確率を、その木の最右導出の逆過程(図4参照)を条件とした条件付確率の積、

$$P(T) = \prod_{i=2}^n P(r_i | r_1, r_2, \dots, r_{i-1}) \quad (1)$$

と定義し、これを次のような単純マルコフ過程(bigram)で近似している。

$$P(T) \simeq \prod_{i=2}^n P(r_i | r_{i-1}) \quad (2)$$

北の研究[4]では、この言語モデルをLRパーザによる音声認識に用いている。この言語モデル自体は特定のパーザのアルゴリズムとは無関係なので、チャートパーザや他のCFGパーザに応用することも可能である。

図5にrule-bigramの例を示す。この例は、「...ている」等、文のアスペクトを表現する助動詞を含む動詞句が成立した後($V \rightarrow V\text{-ASPECT}$)、

- 完全な文として成立する($\text{SENT} \rightarrow V$, 「論文を書いている」) 確率が0.3
- 前方の文副詞とより大きな動詞句を作る($V \rightarrow \text{ADV-SENT } V$, 「もちろん書いている」) 確率が0.3
- 形式名詞「の」を伴う($\text{FN} \rightarrow \text{の}$, 「書いているのは」) 確率が0.2、
- 助動詞「ます」を伴う($\text{AUXV} \rightarrow \text{ます}$, 「書いています」) 確率が0.2

であることを示している。

5.3 スコア計算

解析途中の弧は、上記の言語モデルによりスコア付けされる。スコア計算時には、Magerman[5]が提案した確率の幾何平均を採用し、弧の支配する部分文字列長を正規化し、低頻度現象による弧の過小評価を防いでいる。

この言語モデルは、解析途中の弧(部分木)をスコア付けするためのものなので、予測弧(確定した部分木を含まない弧)にはスコアは与えられない。現在の実装では、アジェンダ内に予測弧がある場合は、予測弧を優先して処理を進め、そうでない場合には、もっとも高いスコアを持つ弧を優先して処理を進める、というbest-first探索を用いている。これは、

- 予測弧を先に処理する方が、自然な文の解釈を早期に決定できる[13]。
- 予測弧を優先せずに、best-first探索を行なうと、文頭近くでの探索誤りの影響が大きくなってしまふ。

という理由による。

- 1: こちら-に-会議事務局-です (score = -40.596821)
 - 2: こちら-や-会議事務局-です (score = -40.442726)
 - 3: こちら-は-会議事務局-です (score = -40.245636)
 - 4: こちら-は-会議事務局-でし-た (score = -40.162292)
 - 5: こちら-も-会議事務局-です (score = -40.126244)
- (正解: こちらは会議事務局です)

図 2: N-BEST 音声認識候補 (N=5) の例

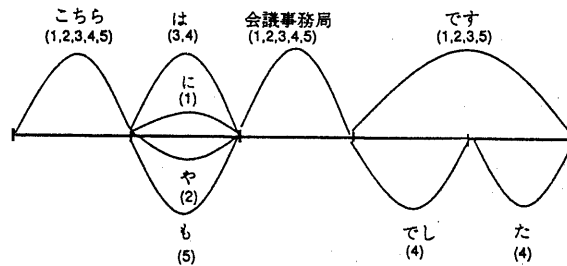


図 3: 初期チャートの例

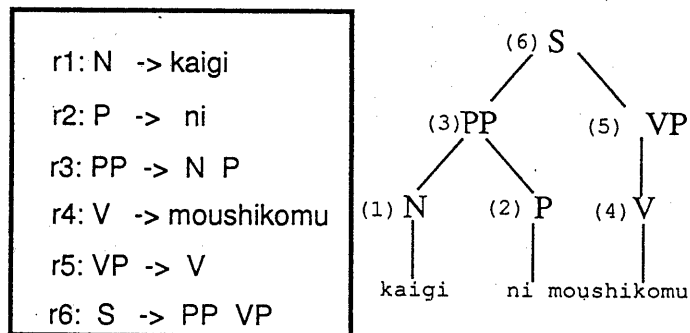


図 4: 最右導出の逆過程

V -> V-ASPECT		SENT -> V		0.3
V -> V-ASPECT		V -> ADV-SENT V		0.3
V -> V-ASPECT		FN -> の		0.2
V -> V-ASPECT		AUXV -> ます		0.2

図 5: rule-bigram の例

	全候補	一発話平均
非共有	326083	7953
共有	199559	4847
改善度	61%	-

表 1: 共有の効果 (ステップ数)

6 実験及び考察

6.1 実験の概要

チャートパーザによる N-BEST 音声認識候補の効率的分析手法の有効性を調べるために、実際の音声認識候補に対して小規模な解析実験を行なった。実験対象は、国際会議の申し込みをタスクとする会話 (男性話者 1 名、計 41 発話 1 発話につき最大 5best 候補、計 187 文) とした。41 発話中、29 発話 (71%) については 1 位候補が正解であり、39 発話 (95%) は 5 位以内に正解を含んでいた。なお、音声認識には SSS-LR 音声認識システム [8] を利用した。

使用した文法は約 1700 ルール (内語彙規則約 1500) の規模であり、単一文法の枠組みで記述されている。この文法は「目的指向型」電話会話における日本語の基本表現の約 90% の表現をカバーしている [14]。言語モデルとして使用した rule-bigram は、実験対象を含まない 800 文を解析した後、人手により正解を選んだ解析木をもとに学習した。また sparse data に対応するため、deleted interpolation [3] の技法を用いてスムージングを行なった。

6.2 共有の効果

まず、共通部分の共有の効果を確認するために、単に N-BEST 候補を一文毎に解析する場合のステップ数と、共通部分を共有し発話単位毎に解析する場合のステップ数を測定した。表 1 に結果を示す。

共通部分の共有により、約 40% ステップ数を減少させることができた。

6.3 音声認識候補の再順序付け能力

統語的・意味的制約とヒューリスティックな選好による音声認識候補の再順序付け能力を測定した。動作例を図 5 に、結果を表 2 に示す。

統語的・意味的制約のみを用いた深さ優先探索では、一位に正しく認識された音声認識候補をすべて正しく解析した上で、統語的・意味的に不適格な候補を棄却することにより、さらに 3 発話の正しい音声認識候補を選ぶことができた。統語的・意味的制約に加え、ヒューリスティックな選好をアジェンダ制御に用いた場合には、33 発話について正しい音声認識候補を最初に選択した。また全数探索を行なった後、もっと

	一位正解発話 (率)	副作用
音声認識	29(71%)	-
深さ優先探索	32(78%)	-
全数探索	34(83%)	3
ヒューリスティック探索	33(80%)	4

表 2: 音声認識候補の再順序付け (音声認識スコアなし)

> (test-acp test-2-7)

((1 来月お申し込みになりますと読まないんです)
(2 来月お申し込みになりますと読まれるんです)
(3 来月お申し込みになりますと四万円です)
(4 来月お申し込みになりますと読まないのんです)
(5 来月お申し込みになりますと読まないでるんです))

A New Result has been found! [2439 steps:11sec]

score = -2.868712199567113

sentence NO.=(SENTENCE 3)

[[[RELN だ-IDENTICAL]

[ASPT STAT]

[COND [[[PARM !X02[]]

[RESTR [[[RELN と-CONDITIONAL]

[IDEN [[[RELN 申し込み-1]

[AGEN []]

[TLOC [[[PARM !X01[]]

[RESTR

[[[RELN 来月-1]

[ENTITY !X01]]]]]

[ASPT UNRL]

[OBJE []]]]

[ENTITY !X02]]]]]

[OBJE []]

[IDEN [[[PARM !X03[]]

[RESTR [[[RELN NUMBER]

[UNIT []]

[IDEN [[[PLACE-10000

四-1]

[PLACE-1000 []]

[PLACE-100 []]

[PLACE-10 []]

[PLACE-1 []]]]

[ENTITY !X03]]]]]]]

図 6: パーザの動作例

	一位正解発話 (率)	副作用
全数探索	35(85%)	2
ヒューリスティック探索	34(83%)	3

表 3: 音声認識候補の再順序付け能力 (音声認識スコア: 言語モデルスコア = 3:1)

も高いスコアの解を選びだしたところ、34 発話について正しい音声認識候補を選択した。

表中の「副作用」の欄は、音声認識で正しく 1 位に認識された発話を、言語知識及び探索手法により評価を下げてしまった発話の数である。これらは、

- 統計量が十分ではない。
- 音声認識のスコアを考慮していない。
- best first 探索なので、最適解が最初に見つかる保証がない。

こと等に起因していると考えられる。

そこで、音声認識スコアと言語モデルのスコアを 3:1 (実験的に決定した) の重み付けで再順序付けを行なってみた。結果を表 3 に示す。全数探索では、元の音声認識結果に比べ、約 14% の精度向上が実現できた。前に述べたように、39 発話 (95%) については候補中に正解を含んでいるので、さらにより言語モデルや探索手法を用いれば、より精度を上げることが可能だろう。

7 おわりに

本稿では、まず音声認識部と言語処理部の統合手法としての N-BEST インターフェースについて述べ、言語処理部に必要な言語知識、言語処理部での N-BEST 候補の探索手法について考察した。さらに、効果的な音声認識候補の解析処理をチャートパーザで実現する方法を提案し、実験結果を報告した。

今後は、さらに大規模なデータを用いて評価実験を続けながら、

- より正確な言語モデル
- より正確で効率的な探索手法

について検討し、高精度で効率的な音声言語解析手法の研究を進める予定である。

謝辞 本研究の機会を与えて下さいました ATR 音声翻訳通信研究所 山崎 泰弘社長に感謝致します。また、適切な助言を頂いた第 4 研究室 森元 逞室長に感謝致します。

参考文献

- [1] Ralph Grishman, Catherine Macleod and John Sterling: "Evaluating Parsing Strategies Using Standardized Parse Files," ANLP-92, pp.156-161, 1992
- [2] Mary P. Harper et al.: "Semantics and Constraint Parsing of Word Graphs", ICASSP-93, pp.63-66, 1993
- [3] Jelinek, F. and Mercer, R.L.: "Interpolated estimation of Markov Source Parameters from Sparse Data," Workshop Pattern Recognition in Practice, pp.381-397, 1980
- [4] Kita, K et al.: "Continuously Spoken Sentence Recognition by HMM-LR," ICSLP-92, pp.305-308, 1992
- [5] Magerman, D.M. and Marcus, M. P.: "Pearl: A Probabilistic Chart Parser," EACL-91, pp.193-199, 1991
- [6] Magerman, D.M. and Weir C.: "Efficiency, Robustness and Accuracy in Picky Chart Parsing," ACL-92, pp.40-47, 1992
- [7] Martin Oerder and Hermann Ney: "Word Graphs: An Efficient Interface between Continuous-Speech Recognition and Language Understanding," ICASSP-93, pp.119-122, 1993
- [8] 永井明人、鷹見淳一、嵯峨山茂樹: "逐次状態分割法 (SSS) と音素コンテキスト依存 LR パーザを統合した SSS-LR 連続音声認識システム," 信学技報, SP92-33, (1992-06)
- [9] Nagata, M. and Morimoto, T.: "A Unification-Based Japanese Parser for Speech-to-Speech Translation," IEICE Trans. Inf. & Syst., Vol.E76-D, No.1, pp.51-61 (1993).
- [10] 永田昌明ほか: "疎結合かつ階層的な音声言語インターフェース: 音声認識用文法と言語処理用文法の段階的統合を目指して," 情報処理学会第 43 回全国大会, 6V-01, 1991
- [11] Schwartz, R. M. and Chow, Y.: "The N-Best Algorithm: Efficient Procedure for Finding Top N sentence Hypotheses," ICASSP-90, pp.81-84, 1990
- [12] Stephanie Seneff: "Robust Parsing for Spoken Language Systems," ICASSP-92, pp.189-192, 1992
- [13] 島津明: "文の読みの偏好と解析の付加との関係 - Mental OS の観点から -," 情報処理学会大 4 3 回全国大会, 3G-5, 1991
- [14] 浦谷則好ほか: "話し言葉の日英翻訳システムの評価法," 情報処理学会第 4 6 回全国大会, 6B-4, 1993