

第58回全国議事記録議事運営事務研修会

音声認識技術の現状と会議録作成への適用可能性(2)

京都大学教授 学術情報メディアセンター 河原達也氏

・音声認識の原理

音声認識の原理

- 入力音声Xに対して、以下の確率が最大となる単語列Wを見つける

$$p(W|X) = \frac{p(W) \cdot p(X|W)}{p(X)}$$

$p(W)$: (その状況において)単語列Wが発せられる確率

- 日本語における単語連鎖
- 話題・状況に起因する語彙
- スタイルに起因する言い回し

言語モデル

外国語では、状況や話題が特定できていると聞き取りやすい

$p(X|W)$: 単語列(音素列)Wに対する音声Xの確からしさ

- 日本語の標準的な発声パターン
- 個人の発音の特徴
- 周囲環境や入力装置の音響特性

音響モデル

母国語では、聴覚だけで聞き取れる

いきなり式がついて恐縮なんです、音声認識の原理を一言で言うところの式になってきます。数学的にというか、記号を使っていますと、入力音声(X)に対して、この確率が最大となる単語(W)を見つける。音声が入ってきたときに、「はい」と言っているのか、いろいろ考えながら聞いていると思うのですが、その確率が一番高いと思うもの、こう言っているのだらうという確信の高いものを最終的に認識結果とします。

言語モデル、音響モデルという二つの確率の項があります。音響モデルは単語列、あるいは「あ」「い」という音素に対して、その音声とはこういうものであるということ。子供のときから言葉が聞いていると、自然に耳ができてくるのと

ム・ジョニール」といった単語が出てくるということが何となく予測できずから、聞き取りやすくなってくるわけです。外国語を聞く場合には、非常に重要なかぎになっているのです。外国語ではこちらの言語モデルがないと聞きにくいし、あると聞きやすくなります。

・音声認識の基本要素

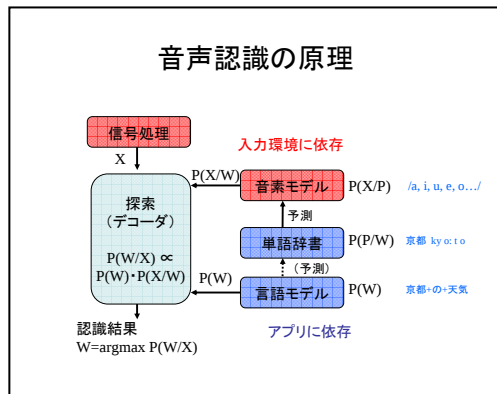
このシステム、機械をつくるときは、さつき申しましたように、日本語の言語学的な知見、「私」の次には「は」「の」が来やすいとか、話題とか状況がある程度モデル化するとか、会議でしたら、「何々に関しては何々でございます」といった一定の表現、言い回しを機械に学習させます。こういう表現がよく来るんだというスタイルを機械に覚えさせるわけです。

会議と電話の会話では、話題や言い回しが全然違いますから、それぞれタスクドメインごとにつくらないといけないということになります。

一方、音響モデルの方は、純粋に「あ」とか「い」という声をたくさん集めてきて、性別、環境ごとにパターンをつくり

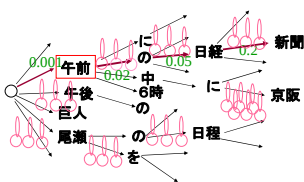
ます。カーナビと電話では当然環境が違いますから、それぞれ音のパターンを覚えさせてつくりまします。

音声認識の原理



これがブロック図にしたものです。音声に周波数分析みたいなことをします。それに対して、先ほど申しました音響・音素モデルと言語モデルは、全部確率的な枠組みでやっていますので、全部尤度に落ちます。天気の内容をするアプリケ

音声認識エンジン(デコーダ)による単語列仮説の探索



ーションを考えるなら、言語モデルでは「京都の天気」「あしたの空模様」といった表現が来やすいというのを覚えさせて、「京都」は「KYOTO」というスペルも覚えさせ、「K」「T」はこういう音のパターンだということを覚えさせて、それぞれを総合して認識します。最終的に「京都の天気」と言った確率が一番高ければ、それを認識結果にします。

同じです。個人の発音の特徴も、ある程度入ってきますし、周囲の環境を多少考慮してつくりまします。

それに対して言語モデル。ネイティブの母国語話者は、音響モデルだけでほとんど聞き取れるようになるのですけれども、特定の状況において、この単語列が発せられる確率、日本語でしたら例えば「私」の次には「は」が来るとか「の」が来るといった文法的な制約もございまして、もうちょっと言いますと、例えば政治の話をしていたら政治の単語が出てくる、経済の話をしていたら「株」が替わるといった単語が出てくる語彙とか、話題とか状況に応じて出てくる語彙とか表現がある程度決まってくることを考慮してつくりまします。

それから、スタイル。議会とか講演でしたら、「ございます」といった丁寧な言い回しがされるであろうといったことです。このような状況を想定して、そこでどういう表現が使われているかを予測すると、聞き取りやすいということです。英語を聞く場合は非常にわかりやすいと思うのですが、北朝鮮のニュースをしていたら「アトミック」「キ

これは、その過程をイラストにしたものです。外国語や英語を聞いているときに、ああ言っているのか、こう言っているのか考えながら聞くことがあります、それを機械でもいろいろやっているのです。

これは我々の宣伝になるのですが、私どもの京都大学が中心になって、フリーの音声認識ソフト「Julius」を十年ぐらい前からつくっております。このウェブ 사이트は、「音声認識」で検索すると五番目以内には出てくると思うのですが、多分一番に出てくると思います。この特徴は、商用のソフトと違って、フリーソフトということ。タダということもあるのですが、ソースコードも開示しております。パソコンで言うとういんドウズとリナックスというのがありますが、リナックスの方は中身が全部見えるようになっていて、それに近いものでもありまして、一般の方にとって必ずしも使いやすいわけではありませんが、システムを開発する方には非常によく使っていただいています。

それから、市販のものと違って、リナ

ックスでもウインドウズでも動くようになっていきますし、最近はあるメーカーの方のマイコンにも移植してもらっています。また、日本語だけではなくて、英語やいろんな言語に持っていけます。

ただ、公開しておりますバージョンは、基本的にはデイクテーションのモデルでありまして、かなりゆっくり丁寧に発声しないと認識できません。

では、なぜそれしか公開しないのかと言いますと、先ほども言っていますように、それ以外のものは全部オーダーメイドになってきますので、議会なら議会用にコッソリデータを集めてつくらないといけないという過程がどうしても要ります。万能で何でも認識できるというレベルでは、丁寧に発声することを前提にしなければいけません。

オーダーメイドで開発するときにかぎとなりますのがコーパスデータです。さっき音響モデルと言語モデルという二つのモデルがあると言いましたけれども、それぞれ統計量を学習させます。人間の場合でも、たくさん子供のころから言葉を聞いていると、だんだんわかるようになってきたり、英語でも、何十時

間、何百時間とヒアリングの経験を積むと、英語のパターンが身についてきますけれども、それに相当することをやっております。

具体的には、音ですと少なくとも数十時間、できれば数百時間聞く。当然、一人の人の声だけ聞いていてもいけません。いろいろな話し方の人の声を聞かないと、聞き取りはできませんので、目安としては数百人が必要と言われています。

言語モデルの方は、数千万語ぐらい。数千万と言われてもピンと来ないと思いますけれども、新聞記事で言うところの十年分。あるいは今、日本じゅうに数億あると言われているウェブページを全部集めてきてつくるということをやっています。

ただ、それらは基本的には書き言葉です。どうしても議会のような話し言葉にはそのままでは適用できません。ですから、議会の会議録といったものをベースにつくっています。

・話し言葉音声認識の問題
では、具体的にどうつくっているかというお話です。

これは、話し言葉と書き言葉、読み上げ音声と実際の自然な音声でどういう違いか、これも体感していただくために音を出してみます。今度は日本語です。サンプル一 また、政治に金がかかるということと体がわかない。

サンプル二 このお話、幹事の方からいただいたときも、そのやりとりの中でも実はあったのですが……

こんな感じで、最初のように丁寧に発声すると九五%行くのですが、後のように、本当にそのまましゃべっているのを生で入れると六五%になります。

二つを聞き比べていただきますと、どういう違いがあるかといいますと、話し方の、特にスピードが後の方が速いということもありますし、話し言葉の特徴として、速くなったりゆっくりになったりします。考えているとゆっくりになりますし、スラスラ出てくれば速くなるといった大きな変動があります。音声認識システムにとっては結構クリティカルで、一定速度で発声しないと認識できないという特徴があります。

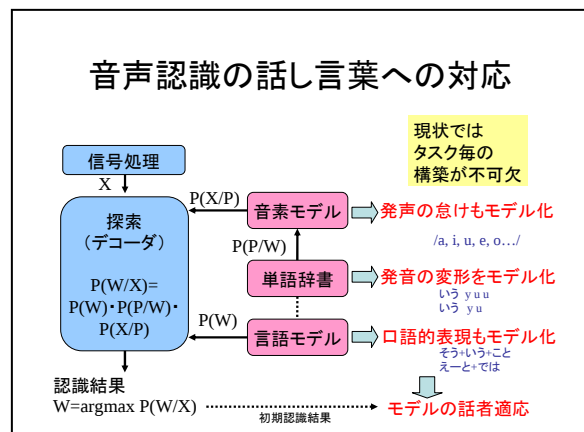
あとは、発声の怠けとか口語表現であるとか、言いよどみといった現象があり

ます。

また、話し言葉の特徴として、「何とかですが、何とかでして……」と、一分でも二分でも話す方がおられますけれども、発話の区切りがよくわからないといった現象もあります。

もう一つ、一番の特徴として、話し言葉の場合、個人差が大きいということがあります。つまり、表現にしても、発音にしても、クセというか方言も含めまして、標準語を意識してしゃべればそうなのですが、地言葉になると差が大きくなりますので、できることならば、人ごとに適応していくことが望ましいと言われています。基本的には、音響モデル、言語モデルと全部話し言葉用にカスタマイズしてつくります。

この図は、さっきと同じブロックダイアグラムです。発声の怠けや発音の変形、話し言葉の口語的な表現をモデル化します。それから、話者ごとに適応します。現状では、これらをタスクごとにつくる必要があります。



・議会向け音声認識モジュールの構成

衆議院向け
音声認識モジュールの構成(京大)

	現状	将来(研究開発中)
音素モデル	CSJ	CSJ+審議コーパス
単語辞書	会議録+CSJ	会議録+審議コーパス
言語モデル	会議録+CSJ	会議録+審議コーパス
認識率	80%(予算委)	85-90%??

CSJ: 日本語話し言葉コーパス(学会講演:228時間)
審議コーパス: 衆議院審議の忠実な書き起こし:150時間

これは、今、京都大学で研究開発しています衆議院向けの音声認識モジュールの構成です。現状はCSJと書いておりますが、学会講演を二百時間ぐらい集めたものと、衆議院の会議録をベースにつくっています。現在、衆議院さんの方から百五十時間ぐらいの実際の方々のデータの提供していただいております。そして、それをベースにこれらをつくつ

ています。現状では、予算委員会で八〇%ぐらいの認識率ですが、目標としては、一年以内にこれを八五〜九〇にしたい。これは確たる成算があるわけではありませんが、現状よりは幾らかよくなるだろうという見通しです。

・衆議院審議コーパスの例

これは実際の衆議院のコーパスの例です。黒い字で書いてありますのが会議録でありまして、これは皆さんの方が身近でよくおわかりかと思えます。

実際の発話を忠実に書き起こしますと、「えー」「そのー」が随分入っております。赤い部分がフイラーです。「ですねー」も会議録では落とされています。

それから、「だけど」が「だけれども」になったり、「なるんで」が「なるので」という文書体に整形することもなされています。

「ただいてる」を「ただいている」のように、「い」を補完するなど、実際の忠実な話し言葉と会議録でも随分ギャップがあります。衆議院の予算委員会の例で数えてみますと、単語のうち一二%ぐらいが違っていると観測されて

います。

衆議院審議コーパスの例

- ・(えー)それでは少し、今(その一)最初に大臣からも、(その一)貯蓄から投資へという流れの中に(ま)資するんじゃないだろうかとかいうような話もありましたけれども、(だけど/だけれども)、(まあ)あなたが言うとうとうそらしくなる(んで/ので)ですね、えー)もう少し(ですね、あの一)これは(あー)財務大臣に(えー)お尋ねをしたいんです(が)。
- ・(ま)その(あの)見通しはどうかということでありますけれども、これについては、(あの一)委員御承知の(その)「改革と展望」の中で(ですね)、我々の今(あの一)予測可能な範囲で(えー)見通せるものについてはかなりはつきりと書かせていただいて(い)るつもりでございます。

ですから、会議録は実際たくさんあるのですけれども、本当に話している言葉をモデル化するには不十分ということとで、こういった忠実な書き起こしをつくつていきます。

それから、これはうちの特徴ですけれども、先ほど申しましたように、話者によつて発音は随分変わりますから、一回普通の音声認識システムで認識して、その結果に基づいて、話者ごとにデータを

区切ります。だれそれさん、だれそれさんと、話者ごとにモデルをチューンしてからも一回やるので、二回認識するのです。これは当然リアルタイムな処理はできませんが、一回目よりは時間は短くなりますので、実用的にはそんなに問題はなからうということで、実際には五%ぐらいのゲイン(向上)原因があります。もっと望ましいのは、総理とか大臣とか、よく発言される方は、それぞれその人ごとにデータを集めてつくることなのですけれども、現実的に可能なかどうかは、まだ検討しているところであり

・発音モデル

辞書の方も、話し言葉に特有な発音を入れていかないとけません。

この例での「音声に起因」が、実際の発音の「オンセインキーン」が「キーン」と長音化することがあります。こういう現象をとらえておくと、「閑静な寺院」という単語についても同様に発音がわかってきます。

これは実際に日本語でどういう現象が多いかということですが、「本当に」

が「ほんとに」とか、「それから」が「そえから」、「けれども」は「けども」など、長音化、促音化が多いのですけれども、こういう現象があります。もう少しおもしろいのは、例えば「用いる・用いれ」のときは「もちーる」と長音化して発音しますが、「用いれ」のときは余り長音化しない。「用いよう」になるとほとんど長音化しないといった現象もとらえております。

・言語モデル

言語モデルとは、言い回し、単語の連鎖がどういふことかということです。

「えー」のようなフイラーも、どこでもかしこでも入り得るわけではなくて、文頭に入りやすいとか、(ん)「これは」という後に入りやすいなど、ある程度の特徴がありますので、そういうものもある程度モデル化します。

「このー」の前に「えー」が入ることが多いというのがわかりますと、そのほかにも入りやすいだろうということが予測できます。

これも、今ある百五十時間ぐらいのデータで分析しますと、「えー」が一番多

くて、そのほかに「ですね」が多い。置換の例としては、「ている」が「てる」、「という」が「ていう」になることがわかっていきます。これも全部統計量を求めて、音声認識システムに反映させています。認識するときは「えー」を予測して、後で捨てることをやっています。

こういうモデルをつくるときも、いろいろ検討事項はあるのですが、まずは議会ごとにつくらないといけないということと、東京で話される話題とか語彙と北海道で話される話題や語彙が随分違いますから、それらを反映する必要があります。

もっといいますと、衆議院でしたら、財務委員会とか法務委員会があります。委員会ごとに話題とか語彙も違ってくるので、やってみないとわからないのですけれども、それぞれつくる必要があるのかなのかという検討はしないといけないと思います。

また、きょうの会議ではこういう質問がなされるということが事前にわかっていましたら、質問書、答弁書から予測もできます。NHKのニュースは、当日のアナウンサーの原稿を機械に入れて

実際にやっていますけれども、それに近いことをやるとよくなる。よくなるとは思いますが、こういうことをやると手間がふえますので、どのくらいよくなつて、どのくらい手間がふえるかというトレードオフも考えなければいけません。

こういったことを反映して開発しておりまして、現時点では衆議院の予算委員会の総括では、八〇%ぐらいの認識精度になつていくということで、現在これの改善を目指して研究開発を行つております。

私は、予算委員会がほかと比べて難しいかどうかはよくわかつておりませんが、けれども、本会議よりは予算委員会の方が難しいのは自明でして、ほかの外務委員会とか法務委員会に比べてどうかというところは、まだ検討はしておりませんが、そんなに大きく違いはないだろうと考えております。

・デモンストレーション

ここで、どのぐらいのものかというのを体感していただきたいと思います。七〇%、八〇%、九〇%という順番にお見

せします。最初が七〇%の例です。

この議論は、ちよつと精査しなければならぬと思いますが、例の米朝の合意ができましたね、一九九三年か一九九四年かなあの時点から、アメリカは北に優しくなりましただね。韓国も、北に優しくなりましたね。日本もコメなんか支援したりして、優しくなりましたね。結局あの時点で、北朝鮮の核は大体おさまったなど、みんな安堵の胸をなでおろして、その後は、それ以上激しくすると、また彼らが機嫌を損なうという意味で、結構太陽政策みたいなものが始まつたのは、そのあたりからだと思うんです。

しかし、そういうことをやりながらつき合つてきたら、ふたをあげたら核という問題が出てくる。そういう意味で、これは交渉次第だと思ひますけれども、特に経済協力等については、あるいは人道主義的なものについては、変に施すみたいな気持ちにならずに、核は核として検証して、そのことが本当に北朝鮮の核を放棄するような方向に動くのかという、そこだけは本当に精査した上で心を、気配りをしてもらいたいというのが、私の心からの願いでございます。

これはえらくタイムリーな話題なんです、実際には二年前の委員会のデータです。七〇%ではなかなか使いものにならないのではないかと思います。多いのではないかと思います。

次に八〇%の例です。

○答弁 E T F に関しましては、これは元本保証のない商品でございます。そのことは事実として重要な点であるかと思ひます。

不適切な発言に不適切な部分があつたということ踏まえ、以後、注意をいたしたいと思ひます。

○議長 海江田万里君。

○海江田万里君 これは大いに反省をしていただかなければいけないわけですが、そういう形で、一日も早く削除した方がいいですよ、本当に、これはね。それはやっていただけのわけですね。それはぜひ、本当に一日も早くお願いをしたいと思ひます。

それでは今、最初に竹中大臣からも貯蓄から投資へという流れの中に資するんじゃないだろうとかいうような話もありましたけれども、ただどなたが言うつと、本当にうそらしくなるので、もう少し、こ

れは塩川財務大臣にお尋ねをしたいんですが、もう片一方で、個人国債、個人向けの国債が大変売れたということですが、その個人向けの国債が売れたということが、この貯蓄から投資へという方向に役立っているのかどうかということ、ご意見を伺いたしたいわけですが、これについてはどうでしょうか。

できているところとできていないところが随分はつきりしていると思うのですが、けれども、やっぱり早口のところは難しいという感じはするかと思ひます。最後に、九〇%の例です。

○赤嶺君 査察を成功させるという立場に、外務大臣、立つておられるのか、あるいは査察について、今どんなふう考えているのか極めて不明なんです、査察を成功させようという立場ではないんですか。

○議長 川口外務大臣。

○川口外務大臣 私は、ぜひ査察を成功させたいと思ひています。そして、その成功させたいという国際社会の気持ちをイラクが理解をして、成功させるための能動的な協力をする必要がある。それがなければ査察は成功しない、そういうことだと思ひます。

こんな感じで、みんな川口さんのようにしゃべつてくれると非常にありがたいのですが、この九〇%の例ですと結構使いものになるかなという印象をお持ちいただけたかと思ひます。

主観的に申しまして、九〇%ぐらいだったらいかなと思ひましたけれども、七〇%ぐらいではちよつとつらいのではないか。八〇%ぐらいだと、微妙だなと思われた方が多いのではないかと思うのですが、先ほど申しましたように平均で八割ぐらいで、目標としては九割に近づけるべく努力をしているという状況であります。

認識は書き起こすだけではありまして、幾つか関連するトピックとしまして、雑音とか残響対策、特に反響、エコーが入るところの検証も必要かもしれない。音響機材をよくするかどうかで認識率が随分変わつてくるとも聞いております。それから、話者の特定をする。実際に会議録を作成される際には、発言者も入れないといけませんので、話者のラベルが自動で入るといい。

先ほどの例では、議長の指名の発話のチャンネルは落としていましたけれど

も、それを音声認識できるかもしれせんし、話し手の特徴をとらえてラベルをつけるということも、技術的にはそんなに難しいことはありません。

それから、「えー」とか「あー」といったフイラーが随分多いので、それも自動的に除去することも必要かと思ひます。

・音声認識結果の情報

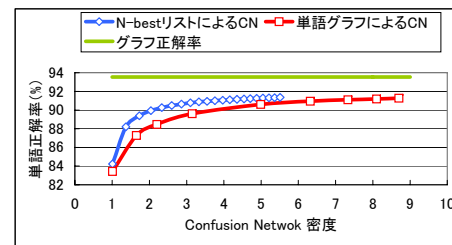
以上が音声認識の技術的な話ですけれども、最後にこれを会議録作成にどのように適用していくかにつきまして、幾つか衆議院の方とも何回か議論を持たせていただいたことも踏まえまして、現在考えていますことを御紹介いたします。

まず音声認識システムの課題としましては、今ビデオで見ていただいたような認識結果の第一候補だけではございませんで、第二候補、第三候補も出ております。

それから、認識には一応信頼度がある。この辺は自信を持って出しているけれども、この辺は怪しいということもある程度ついております。

一番いいのは、単語ごとに対応する音声の区間がわかっておりますから、認識が怪しいところ、あるいは認識が間違っていると判断されたところのみをピンポイントで再生することができます。あるいは、そこだけゆくり再生することもできます。

音声認識の候補数と認識精度(正解率)



平均的に、第2候補までで90%の正解が含まれる

このグラフは、複数候補を出したときにどのぐらいの認識率になるかというものの。例えば青い方を見ていただきますと、認識率の数え方が若干違うのですけれども、第一候補で八二%ぐらいのところを第五候補ぐらいまでとりますと九〇%をちよつと超える。第三候補ぐらいでも九〇%を超えることがわかります。ただけるかと思えます。

・音声認識結果の利用法

こういう情報を踏まえまして、どのように使えるかということについて、幾つか紹介いたします。

まず一番の方法としまして、ワープロに先立って人手修正を介するテキスト化を行ってしまふ。これは北海道議会さんを含めまして、幾つか入れられているところはこういうやり方でやっておりますと思います。とにかく最初に完全なテキストをつくってから、後で会議録としての体裁を整えていく。このメリットとしては、ワープロと音声認識修正ソフトを分けられるということがあります。テキスト化する処理と、後で会議録をワ

ードプロで作成する段階と、二段階のフェーズが入ります。もう少し具体的にいうと二通りありまして、北海道議会さんとかで使われていますアドバンストメディアの商品では、第一候補のみを表示して、だめならタイプで修正します。

それに対して、京都大学で試しにつくってみたものでは、せつかく第三候補まであれば九〇%まで行くのですから、それを表示して選択できるようにすればいいのではないかとすることにしました。ワープロの仮名漢字変換で第一候補がだめで、何回かキーをたたけば第二第三候補で出てくると同じようなインターフェースになっています。

アドバンストメディアさんの「リライタール」という商品の画面をいただいています。特に宣伝ではありませんで、参考のために見ていただければと思います。

音声波形がございまして、今、編集しているところの認識結果が出てきます。例えばここがおかしいというところをクリックすると、その音声が入るようになっていて、画面で音声聞きながら、間違ったところを逐一直してい

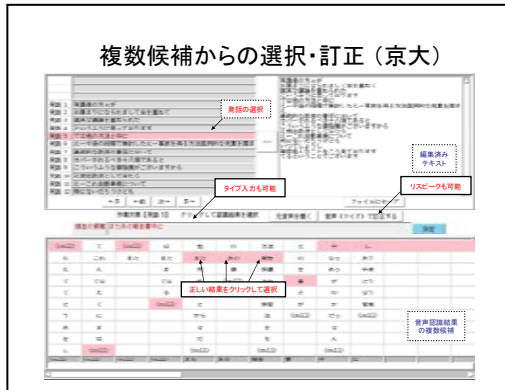
きます。この原稿をお出ししたのが一カ月ぐらい前だったのですけれども、十月になつてから、NECさんの方が同じような「同じような」と言うと怒られますけれども、議会向けのソリューションシステムを発表されたので、それをちよつといただいています。

見た目は、似た感じですがけれども、テキストの編集機能がかなり充実しているという違いがあります。例えば発言者の自動推定もやっていますし、ラベルも入れられるようになっていて、あと、会議録に必要な「(拍手)」、「(起立)」といったコメントも、この画面で入れられます。

また、これは商品ではないのですが、うちで試しにつくってみたインターフェイスです。さつきと同じで、最初の変換が間違っていたら第二候補、第三候補を出すやり方です。それを一タクリックしていただいたら手間がかかりますから、最初から複数個を出して、そこから選んでいけるようになっていて、さつきのグラフで申しましたように、全部合わせると九割ぐらい正解が含ま

れているということ、九割というのはさつきの川口さんの例でありましたように、結構イケているという意味で、こういう方法もあるのかなと考えています。

複数候補からの選択・訂正(京大)



実際に、何人かの方とお話ししますと、こういうやり方は余りよろしくなくて、先ほどのやり方の方が効率がいいと聞いております。タイプ歴にもよるので

すけれども、速記者のようにタイプがなれている方でしたら、キーボードだけで全部仕事が終わる方が速いということ、私もそれはそうかなと思います。

今お見せしたのが、世の中にある商品というシステムですが、世の中にある商品というやり方の可能性として、例えば音声認識結果をいきなり出さないで、人間がタイプをしていくと、音声認識結果を引っ張ってきて、合っている候補を逐次出すやり方もあるかと思えます。

携帯電話で入力をするときに、「おは」と入れば「おはよう」と出るとか、予測入力に近い考え方です。ワープロソフトなら一回で済みますから、作業手順をふやすことなく音声認識結果もそれなりに効果的に利用できるのではないかと考えていますが、ソフトをつくるのが面倒くさい、あるいはワープロソフトのバージョンが変わったら使えなくなるのではないかとといった安定性、将来性に関する懸念があります。

もう一つの方法としては、全部音声認識結果を使うのではなくて、先ほどの海江田さんの例でござんいただきましたように、できているところはほとんど合

・まとめ
話は以上です。これは私の願望に近いところもあるのですが、品質面で、例え

つていますし、特に早口になったところはほとんどできませんので、できているところだけ音声認識結果を出して、できていないところはタイプすると、ストレスも少なくなるのではないかと考えられます。個人の好みに応じて閾値を変え

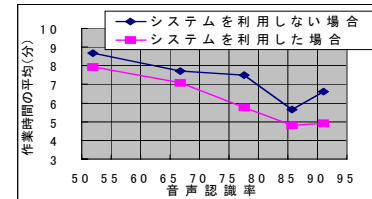
このグラフは、認識率と被験者の方の主観的な評価の関係をプロットしたものです。一応七段階評価で、使い勝手がいかどうかを評定してもらいました。四が真ん中なので、五以上がいいと思われるところで、五を超えるところは大体七五%ぐらいです。

これで見えるのは、使っている側としては、認識率が高い方が使い勝手がいい。さっきのグラフとは違う結論ですけれども、認識率の高い方が、このシステムの使い勝手がいいと感じられる傾向にあります。その閾値が七五%ぐらいのところにあります。

一応タイピングスキルとも関係がありますので、それとの関係を見たところ、パソコン使用歴の浅い人の方が、このシステムを使い勝手がいいという傾向にあります。そういう意味で、プロの方が使うとどうなるかというのは、今後評価検証していかないといけないところになります。

京大・龍谷大でのシステムの評価

- 被験者：大学生・社会人18名
- 比較対象：タイプ入力
- 認識率に関係なく、音声認識利用の方が作業時間が短い
← 認識率が低い箇所(話者)は、人間でも書き起こしが困難



先ほど申しました今できているもので、実際に京大と龍谷大学の学生さんに試しにやってみたら例であります。大学生が多いので、皆さんのようにプロの速記者と、タイプ入力のスキルが随分違うと思うのですけれども、一応参考にはなるかと思っています。

課題としては、先ほどお見せした衆議院の予算委員会の音声認識結果をもとに、会議録を作成してもらいました。横軸が認識率で、五〇%から九五%までランダムに選んでいます。

赤い方は音声認識を使った場合で、黒い方は一から聞いた場合。聞くところはさっきと同じような画面を使っていますので、違いは音声認識結果からスタートしているかいないかだけですが、これは僕らにとってはグッドニュースで、認識率には関係なく音声認識結果を使った方が、作業時間は短くなっています。興味深いのは、認識率が高いと、全部人手でやっても、機械でやっても速くなりますし、認識率が低い人は、機械でやっても、全部人手でやっても時間はかかる。認識率が低い箇所、あるいは低い人

ば単語とか固有名詞を音声認識の辞書に登録しておきますと、特に固有名詞の揺れがなくなるといふ副次的な効果もあります。

それから、さっきの繰り返しですが、熟練した人が集中した場面において作業時間が短くなるかというのは今後評価検証していく必要があるのですけれども、少なくともタイプし続ける状況はなくなりやすから、長時間作業の負荷は少なくなるかと思っています。

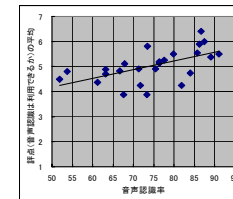
具体的に言いますと、今、例えば衆議院では五分ぐらいのデータが限界ですけれども、こういうシステムを使うと、もうちょっと長くやってもそんなに疲れないのではないかと期待しています。もう少し明らかなのは、このシステムを使う訓練が今の速記を学習するのに比べて少なくて済むであろうということです。

以上をまとめますと、話し言葉の音声認識は一般に容易ではありませんし、それができたかできないかと言われると、私は「できない」と答えるを得ないと思います。議会のような特定の分野のオーダーメイドにすると、それなりにで

は、人間が書き起こしても難しいということが言えるかと思っています。

これは以前、衆議院の方とお話したことがあるのですけれども、議員さんごとに認識率のグラフを出して、それをお見せしたところ、「この先生は難しいんですよね」ということをおっしゃっていました。機械がやっても難しいところは人間がやっても難しいということが言えます。

システムの主観的評価と音声認識率との関係



- 認識率との間に一定の相関
- 認識率が75%を超えると、「音声認識はテキスト化作業に利用可能か」について、比較的高い評定値(7段階評価で5以上)

きつつあるということで、特に議会の音声認識は、実は比較的那中でも容易な分類になります。それは、膨大な会議録などの学習データが存在するからです。実際に衆議院の会議録というのは、年間に千八百万〜二千万単語ぐらいの量がありまして、それは数えると、大体新聞記事一年分と同じです。衆議院の速記者の方がつくられているテキスト、あるいはその議員さんが話されている内容は、新聞と匹敵するぐらいの量があるのです。

もう一つは、収録条件が比較的良好ということ。マイクも音響設備も、またたのものに比べれば、ちゃんとセッティングされています。話者も、基本的には話したれた方が丁寧に発声することで、八〇〜八五%ぐらいの認識精度は可能と考えられます。しかしながら、それが使えないものになるかということに関しては、まだ幾つかの事例はありますけれども、今後一層の検討が必要になります。

どうも長らくおつき合いました。ありがとうございます。以上で私の話は終わらせていただきます。(拍手)