

情報通信審議会 情報通信技術分科会
技術戦略委員会 次世代人工知能社会実装WG 事前会合
(自然言語処理と脳情報通信へのユーザーの期待)
議事概要

1. 開催日時

平成 29 年 1 月 16 日 (月) 14:00～16:00

2. 場所

総務省 10 階 第 1 会議室

3. 出席者 (敬称略)

主任：柳田 敏雄

構成員：東 博暢、池田 尚司、上田 修功、大岩 和弘、大竹 清敬、岡島 博司、加納 敏行、
栗本 雄太、鳥澤 健太郎、萩原 一平、原 裕貴 (代理：西野 文人)、森川 幸治、
八木 康史、山川 宏

ゲスト：上野山 勝也 (PKSHA Technology)、北出 大蔵 (トランスコスモス株式会社)、相良 美織
(株式会社バオバブ)、田口 隆久 (CiNet)、武田 秀樹 (株式会社 FRONTEO)、益子 信郎 (NICT)

オブザーバー：内閣府、文部科学省

事務局 (総務省)：技術政策課／野崎課長、山口企画官

研究推進室／越後室長、出葉推進官、中川補佐、皆川補佐

4. 議事概要

(1) 技術戦略委員会の今後の検討方向について

事務局より資料 1 に基づき説明が行われた。

(2) 自然言語処理技術及び脳情報通信技術に関する NICT の取り組み

鳥澤構成員から資料 2-1、CiNet 田口様から資料 2-2 に基づき、それぞれ説明が行われた。

(3) ユーザー・開発企業の取組み・課題、WG への期待

資料 3-1 から資料 3-8 により、参加いただいたユーザー・開発企業の方々より、それぞれ説明が行われた。

(4) 意見交換等

加納構成員：National Translation Bank の話が鳥澤構成員よりあったが、既に音声や言語に関しては Apple や Google がどんどん情報収集をしているため、弊社も色々なお客様とパートナーを組んで情報を集めてはいるものの、対抗できないレベルに来ている。やはりこれからは日本が結束して、守りではなく攻めの姿勢で、日本語だけでなく海外のデータも含め、感情データや、いわゆる音声の生データと、音声認識に使うような認識用のデータや、翻訳用のデータを組み合わせたような、新しい言語のデータベースを整備していかないといけないのではないか。その上で、どういう分析をしてどういう結果をだしていくかという

ところで、各企業や研究機関が競っていけばよいと思う。やはり、これからは国を挙げてひとつの共通のデータベース化が重要になってくると感じている。

鳥澤構成員：おっしゃる通りとは思いますが、一つ我々が違った認識をしている点としては、例えば siri 等においても正確な翻訳や、質問に対して正確な答えが返ってくるというような、入力と出力の組み合わせに価値がある。そのため、やはり入力と出力の組み合わせでちゃんとしたデータを作るということに意味があるのではないかと。アルゴリズムや分析の形態で勝負するのももちろんありだとは思いますが、一定の量や質は担保しつつ、新しいタイプの学習データとして、特色のあるデータを集めて提供していくことで勝負が優位になるのではないかと。Google にせよ IBM にせよ、人が手をかけて入力と出力のペアを組み合わせたデータについては未だそれほど圧倒的に保有しているわけでもなさそうである。そのところで勝負のしがいがあるのではないかと。総務省でご支援いただければ有難い。

萩原構成員：先ほども申し上げたが、CiNet には今、7 テスラと 3 テスラの計 4 台の MRI が入っている。日本は MRI の数がとても多く、大学にも次々に入っている。しかしそれを使うテクニシャン、データを解析するデータアナリストを拡充していかないと、研究者が全てやらなくてはならず、非常に時間の無駄である。特に産業応用や社会実装といった下流にいけばいくほど、先端でやっている研究者に関わってもらうより、テクニシャンやアナリストに関わってもらったほうが、どんどんデータを取ってどんどん色んなことをやって早くビジネスに結びつけることができる。ここは非常に大きな問題である。そもそも自分で MRI を保有している企業は日本では少なく、測るにはどこかに委託しなくてはならないが、どこにどう委託してどうやればいいのかも全くわからない状況。そのあたりの改善を是非お願いしたい。

野崎課長：一点補足する。この WG を開始することを先週金曜日に情報通信技術分科会へ報告をした。そこでの質疑として、最近話題になっている Amazon の Echo は、一気に 500 万台もアメリカで普及していて、値段も 5 千円程度という。このように各家庭などのあらゆる情報の入り口をおさえられてしまうと、対話領域において日本はどうなってしまうのかという質問があった。これについては、Amazon や Google 等が取り組んでいる中で、まさに今 NICT が実用化している Wisdom-X のような本当に高度な日本語対話プラットフォームが必要であり、個々の企業ではなかなか開発できないのではないかと、そういうところは、やはり国の役割というのが期待されるのではないかと。今後議論していきたい。

柳田主任：ベンチャーをやっている方は、将来 Google や Facebook に勝てるのか、ご意見をお聞かせ頂きたい。

相良様（バオバブ）：11 月 11 日、機械翻訳の世界では、Google Neural Network ショックといわれ、大変な騒ぎになっているが、当然すぐに各研究者及び弊社でもその評価を行っている。世間で騒がれるほど、Google の精度が上回っているわけではない。特に分野を特化すると、如実に結果が現れている。そのため、既に研究者の中では大体どのあたりが弱いとか、こういう風にすれば Google の精度が上回るし、こういったミスを犯しがち、というのは出ている頃。Google 脅威、Google が全てをおさえってしまうと世間では言われているが、絶対に勝てると思って私たちはやっている。

柳田主任：どの部分は勝てる、というのをおっしゃっていただくと、その部分に投資しやすい。

相良様（バオバブ）：意気込みではなくて結果として精度が出ている。

上野山様（PKSHA Technology）：弊社はまさに対話応答のニューラルネットのエンジンを作っている会社である。先ほど Echo の話等がでたが、ここの市場構造が今後どういうレイヤー構造になっていくのかというところをどう見立てるかが重要だと思っている。Echo は色んな家電やデバイスに載ってくるわけだが、Echo 自体が実装している対話応答があらゆる会話をいきなり実現できるかというところはおそらくそうではない。その上側に色んなアプリケーションが載ってくるという形になるだろう。何かを買う、配送するというところはおそらく垂直統合されるところで、その上側にどういう構造でニューラルネットが何層構造になるのか。汎用 AI と専門 AI というのはよく混同されて議論されるが、どの時間軸でどのプレイヤーをとっていくのかをどう見立てるのがすごく重要。我々は専門型からいくつかのファンクションであがっていくというやり方を取っている。一方で総務省の方々や NICT の方とはうまく日本企業がやっていけるような形で連携できると良い。

北出様（トランスコスモス）：実際に自動応答等の領域を普及させるにあたって大きく 3 つポイントがあると感じている。一つはやはり音声認識エンジンを含めた実装というのは、実際のビジネスで広く使うには、価格であれ認識精度の技術面であれ、それをどういう領域で活用するのかというビジネスケースの具体化という部分であれ、非常に一般的な企業で使うにはハードルが高すぎる。その結果学習データがそもそも作られない。そこで 8 割 9 割欠いているという認識でいる。ここのところで何かしらの取組み、共有できるようなエンジン等になるのかもしれないが、そういったものがあれば良いと思う。ただ、それでいたずらに音声認識用のテキスト化したデータを集めても意味があるかというところはどう思っている。鳥澤構成員もおっしゃっていたが、そこにアノテーションをして意味のあるデータをつけていくということが重要。ただ、例えば弊社のコールセンターも 1 日 1 万人以上の人間が日々コールセンターで応対をしている。そこで意味のある応対をしているはずだが、1 年～3 年程でそのデータが消えていっている。すなわちみなさんが貴重と言っているアノテーションデータが、ある瞬間にはあるが、時がたつとどこかに消え去ってしまっている状況。ストックであるという意識が全く現場に無いので、そういったような社会理解を広めることが必要。また、全てをデータとしてアノテーションして残しておくわけにいかないで、どういうデータをアノテーションして残しておくか、それがどういう出力に活かされるのかというガイドラインをもう少し示すことが重要ではないかと思っている。三点目としては、さきほど上野山様の話にもあったが、アノテーションしたデータは実際には学習済みのモデルとして消化して保有することになると思うが、学習済みモデルが知財的にどのように扱われるか、またそれを共同利用していく時やビジネスで提供していく時にどういう商形態があるかというところが見えない中で、企業が取り組むのはリスクがあるとも感じている。そういった法制度の整備も必要ではないかと思っている。

上田構成員：理研は、原則として基礎研究なので、目的志向的な部分でも基礎をやっている。対話で汎用的な部分はやはりみなさんおっしゃったようにデータが全てになるだろう。対話でいえば、例えばうつや認知症のようなところを改善し、間接的に日本の産業、労働力を高めるといった研究は目指している。

武田様 (FRONTEO) : うちがエキスパートの認識と判断の代替というところをメインにやっているが、我々のようなドメインでやっている、対話や音声認識というのは前段階の技術として非常に接続性が良い。そのような技術に関しては国・民間を問わず積極的にパートナーシップを結んでいきたいと思っているので、そういうところを是非簡単に使える状態にしていきたい。また、データの蓄積については、例えば訴訟の対応をされていて企業のインターナルな情報、なかなか外にでてこない情報に対して大量に触れる機会があったため、そこを学習として取組むことができた。やはりパブリックな情報に関しては既に先行している Google のようなところが大量に貯めこんでいるため、なかなか差異をだすのは難しい。例えば企業の中にあって取り出しにくい情報や、そういう希少性の高い情報の集積という点でないとなかなか差異がでないのではないか。もう一点、何か作るというところはもちろんよいのだが、実際にユーザー企業が使うというところが進まないと、どうしても社会実装を含めなかなか進んでいかない。そのため企業が人工知能技術を使うことによって何らかのインセンティブがでるような、そういう方法があると我々として非常にやりやすい。

柳田主任 : NICT がラベルしたデータを作って、そのデータを公であるから、研究で得たデータであるから開放しなさいといわれると、こちらのインセンティブがなくなってしまう。やはり企業の方と、一緒に会話しながら大きなプロジェクトでデータを取るというシステムを作っていないと、研究者と話をしながら少しデータをもらうという形ではたいしたデータ量にならない。企業も無料とは思わずに、データをつくることを一緒にやり、その代わりかなりの資金は提供するというようなシステムを国には作って頂きたい。データがあり、それが応用できるというところまで道筋が見えていけば、学生は育つと思う。これに関し、八木先生に最後に人材育成についてご意見を頂いて終わりにしたい。

八木構成員 : やはり、実データで学んでいくようにしないとなかなか人が育たないと思う。社会実装の中で実際に大切なデータが大学における教育にも活用できると、大学からも企業で活躍してもらえる人材をたくさん排出できると思う。ぜひ企業にはそういう観点でとらえてもらえるとうれしい。

柳田主任 : おそらく、少しラベルしたようなデータをしっかり取ってそれが応用に結びつくというシステムができてしまえば、学生は非常に興味をもって大勢来ると思っている。是非そういうシステムをまず作ることが大事だと思う。1960 年代は半導体と IC では、アメリカには絶対勝てないといわれていたが、20-30 年したら Japan as No.1 になった。車もそうだと思うが、この分野もそうなるように頑張りたい。我々が大学を出た頃は、絶対に半導体と IC でアメリカに勝てるなんて誰も言わなかったし、今より深刻な状況だった。しかしちゃんと勝ったのだから、また Japan as No.1 になるように皆さんで頑張りたい。

(5) その他

第 1 回次世代人工知能社会実装 WG は 1 月 30 日 (月) に開催予定。