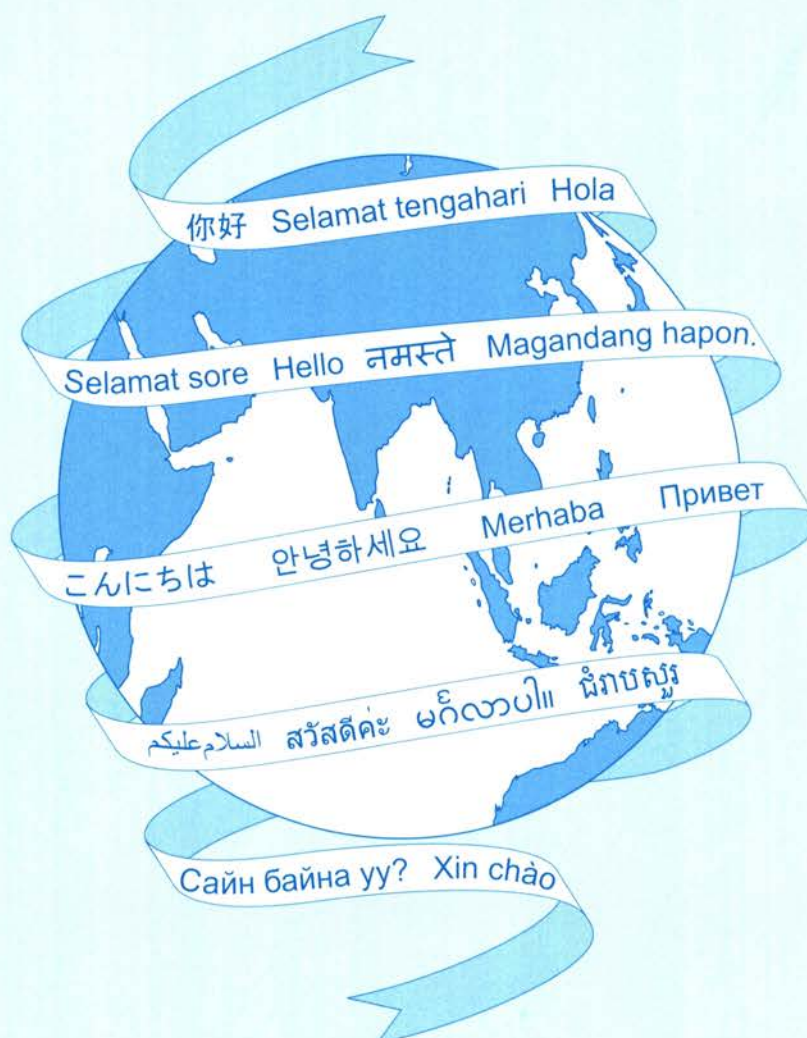


AAMT

Asia-Pacific Association for Machine Translation

Journal



October 2013 **No.54**

アジア太平洋機械翻訳協会

目 次

巻頭言：	翻訳関連団体のより積極的な協働・連携について 東 郁男 1
プロジェクト報告：	特許翻訳における統計翻訳とルールベース翻訳の比較および機械翻訳の有用性評価 後藤 功雄, 周嘉寶, 路斌, 隅田 英一郎, 鄒嘉彦 3
プロジェクト報告：	Building Chinese-English Machine Translation Systems for Patent Documents Zhongqiang Huang, Jacob Devlin, Jeff Ma, Rabih Zbib, and Rich Schwar 14
研究報告：	日英・英日に適した翻訳自動評価法 磯崎 秀樹 19
レポート：	知らなかったでは済まされない『翻訳の国際常識』 田嶋 奈々 21
レポート：	Translation before Demand Geert. BENOIT 25
レポート：	用途に合わせた機械翻訳のススメ 渡邊 麻呂 30
レポート：	コンピュータ処理用日本語表現辞書 JDMWE 首藤 公昭 34
シンポジウム参加報告：	第 14 回機械翻訳サミットおよび第 5 回特許翻訳ワークショップ参加報告 横山 晶一, 梶 博行, 二宮崇 39
シンポジウム案内：	「第 23 回 JTF 翻訳祭」のお知らせ 一般社団法人日本翻訳連盟翻訳祭企画実行委員会 ... 44
総会講演：	言葉の壁を越えたコミュニケーションの実現にむけて 那須 和徳 46
総会講演：	ソフトウェア ローカライゼーションにおける MT 利用と業界間断層 古田京子, 菊池邦明 49
AAMT 長尾賞受賞講演：	Crowd and Computer Translation Collaboration at Baobab 相良 美織 57
製品紹介：	翻訳支援ツール Transit ^{XT} のご紹介 株式会社シュタールジャパン ... 61
製品紹介：	PC-Transer 翻訳スタジオ V21 に搭載される新機能の紹介 株式会社クロスランゲージ ... 64
製品紹介：	「The 翻訳*エンタープライズ V16」 東芝ソリューション株式会社. 68
委員会活動報告：	AAMT 機械翻訳文テストセットに基づく日中機械翻訳の自動評価サイト作成報告 AAMT 機械翻訳課題調査委員会 WG1 70
AAMT 会員のひろば：	 YAMAGATA INTECH 株式会社 73
AAMT 長尾賞学生奨励賞：	「AAMT 長尾賞学生奨励賞」の新設について 中岩 浩巳 75
事務局からのお知らせ：	第 23 回通常総会および関連行事の報告 AAMT 事務局 77
	協会活動報告（2013 年 4 月～2013 年 8 月） AAMT 事務局 79
編集後記	 83

C O N T E N T

Foreword:	Toward a more proactive cooperation and partnership with translation industry organizationst..... <i>I. Higashi</i>1
Project Report:	Comparison of Statistical Machine Translation and Rule-Based Machine Translation and their Utility Evaluations in Patent Translations <i>I. Goto, K. P. Chow, B. Lu, E. Sumita, and B. K. Tsou</i>3
Project Report:	Building Chinese-English Machine Translation Systems for Patent Documents <i>Z. Huang, J. Devlin, J. Ma, R. Zbib, and R. Schwar</i>14
Report of research:	Automatic Evaluation Method for JE/EJ Translation..... <i>H. Isozaki</i>19
Report:	Global Translation Wisdom - What You Need to Know..... <i>N. Tajima</i>21
Report:	Translation before Demand <i>G. BENOIT</i>25
Report:	Encouraging Machine Translation According to Purpose <i>M. Watanabe</i>30
Report	Japanese Dictionary of Multiword Expressions: JDMWE: <i>K. Shudo</i>34
Symposium Report:	Report on the Machine Translation Summit XIV and the 5th Workshop on Patent Translation <i>S. Yokoyama, H. Kaji, and T. Ninomiya</i>39
Symposium Information:	Information for 23rd JTF TRANSLATION FESTIVAL TOKYO 2013 <i>Japan Translation Federation</i>44
General Meeting	Realization of communication to overcome the language barrier <i>K. Nasu</i>46
General Meeting	MT utilization in software localization and chasm between industries <i>K. Furuta, K. Kikuchi</i>49
AAMT Nagao Award	Crowd and Computer Translation Collaboration at Baobab..... <i>M. Sagara</i>57
Products	Introduction to Transit ^{NXT} <i>STAR Japan</i>61
Products	New Features of PC-Transer V21 <i>Cross Language</i>64
Products	The Honyaku Enterprise V16..... <i>Toshiba Solutions Corporation</i>68
Committee Report	A report on making a new website for automatic evaluation of Japanese-to-Chinese machine translation system based on AAMT testsets <i>WGI</i>70
AAMT Members:	 <i>YAMAGATA INTECH Corporation</i>73
AAMT Nagao Student Award	The establishment of the AAMT Nagao Student Award <i>H. Nakaiwa</i>75
AAMT Activities	General Meeting.....77
	AAMT Activities (From April to August, 2013)79
Editor's Note:	 <i>T. Utsuro</i>83

翻訳関連団体のより積極的な協働・連携について

一般社団法人 日本翻訳連盟

代表理事・会長 東 郁男

9月8日（日）早朝、IOCのロゲ会長の「TOKYO」という言葉で、2020年の夏季オリンピックの東京開催が決定しました。バブル崩壊、失われた20年、景気低迷など明るいニュースのなかった中で、久しぶりに日本中が湧き上がる瞬間でした。

オリンピック開催に向けてインフラ、観光、サービス関連だけではなく、さまざまな業界への波及効果が期待されています。いうまでもなく、これをきっかけに「人・物・資金」のグローバルな動きが加速することにより、翻訳・通訳関連需要の増加も期待されるところです。

（日本翻訳連盟の紹介）

まず、日本翻訳連盟（以下、JTF）について簡単に紹介させていただきます。JTFは1981年に任意団体として設立され、1990年には通商産業省（現経済産業省）認可の公益法人となりました。その後、公益法人制度改革に伴い、2012年4月に一般社団法人へ移行しております。JTFは産業翻訳の業界団体として会員間の情報交換・交流を深めながら、急速に進むグローバル化の中で会員のビジネスチャンスを拡大させ、業界の規模拡大・地位向上を目指しております。

具体的には、翻訳祭・翻訳セミナーなどの研修・講演事業、ほんやく検定などの資格能力審査事業、業界調査などの調査研究事業、ホームページ・JTFジャーナルによる情報提供事業などを行っております。

（翻訳需要と課題）

さて、オリンピック開催決定以前の業界の環境について述べますと、新政権によるアベノミクスの民間投資を喚起する成長戦略の中にも「製造業の支援」、「企業の海外展開支援」について言及されています。日本の国内市場が成熟化傾向にある中で、大企業だけではなく、中小企業、個人レベルでのグローバル展開も進んでいます。その過程においての翻訳・通訳需要についてはここに記載するまでもないと思いますが、翻訳・通訳事業はグローバル化にとって欠かせない事業だと考えます。

この点でも、今後さまざまな業界における翻訳・通訳のさらなる需要は期待できると考えますが、一方ではソースクライアント（翻訳を発注する企業）の「クオリティ・スピード・コスト」の要求基準は厳しくなりつつあります。

翻訳業界でも、従来から翻訳支援ソフトやマクロなどのツールを翻訳工程の一部に組み込むなど、いろいろな取り組みにより作業の効率化、品質維持・向上に努めている状況です。

しかしながら、「クオリティ・スピード・コスト」という観点から、翻訳業界がクライアントのニーズに応えられているか？と自問自答すると、残念ながら応えきれていないのが現状だと思います。実際、クライアントの中にも『よい品質の成果物を納品してくれれば、依頼したい案件は山のようにあるから…』、『もっと速くできれば…、もっと安くできれば…』と多くの方が話されます。

翻訳に関わる個人、会社を含む翻訳業界全体での作業効率が向上し、キャパシティが拡大すれば、今見えていない、顕在化していない翻訳ニーズはまだまだあると考えます。顕在化しているニーズ、顧客企業の要望に対してさえもまだ応えきれていない現状ですから、市場は未知数だと考えています。

（AAMTとの連携）

AAMTは、機械翻訳（MT）の研究開発者、製造販売者および利用者の3者で構成されておりますが、円滑なグローバルコミュニケーションが図られるように機械翻訳システムの開発・改良・啓蒙・普及を通じて機械翻訳の発展に努めておられます。

AAMTの方々、MTに関連する企業、団体の方々の長年の努力により機械翻訳の精度は向上し、また機械翻訳利用の可能性も拡大しています。

言語体系の違いはありますが、すでにヨーロッパにおいてはMTの利用を前提とした翻訳工程が普及しつつある状況です。クライアントにおけるビジネス上のニーズは非常に多岐にわたっています。そのニーズに対してMT関係者と翻訳業界関係者が積極的に連携して、さまざまな課題を解決していき、クライアントニーズを満足させる必要があるでしょう。

JTFなどの業界団体を含めて「翻訳」に関連する諸団体が連携していかなければならない事案ではないでしょうか？

なお、来る11月27日（水）に東京・市ヶ谷のアルカディア市ヶ谷（私学会館）において「第23回 JTF 翻訳祭」が開催されます（http://www.jtf.jp/jp/festival/festival_top.html）。今回は「大翻訳時代～Define Your Blue Ocean Strategy～」というテーマのもと、AAMTをはじめ多数の企業、翻訳者の方々にご協力いただき、講演・パネルディスカッションなど多数のイベントを準備しています。

「JTF 翻訳祭」への多くの皆さまのご来場を是非お待ちしております。

翻訳に関わる人々、団体のみならず、ひいては日本経済・社会の発展のため、翻訳に関連する諸団体がこれまで以上により緊密に協働・連携していけるよう尽力していきたいと思いをします。

特許翻訳における統計翻訳とルールベース翻訳の比較 および機械翻訳の有用性評価

—NTCIR-10 特許機械翻訳タスクの評価概要—

情報通信研究機構 後藤 功雄*

Hong Kong Institute of Education 周嘉寶** (Ka Po Chow)

City University of Hong Kong / Hong Kong Institute of Education 路斌*** (Bin Lu)

情報通信研究機構 隅田 英一郎

Hong Kong Institute of Education / City University of Hong Kong 鄒嘉彥**** (Benjamin K. Tsou)

1 はじめに

特許翻訳は、外国語で書かれた特許の理解や外国への特許出願のために、実用上において大きなニーズがある。そのために特許の機械翻訳の研究は重要であり意義がある。この研究の促進のために評価型ワークショップである NTCIR-9 および NTCIR-10 にて特許機械翻訳タスクを主催した[1,2]。NTCIR-9 でのランダムに選択した文の訳質評価から、現在の最高性能の機械翻訳システムは、特許文（説明文の部分のみでクレーム部分を除く）の翻訳において半数以上の文で原文の内容を理解できることが分かった。この結果から機械翻訳が実用上においても有用である可能性が高いと考え、特許翻訳が必要とされる状況において機械翻訳を利用した場合にどれだけ有用であるかという観点における評価の1つとして、NTCIR-10 特許機械翻訳タスクで特許審査での有用性に基づく評価（特許審査評価）を文の訳質評価に加えて実施した。

本稿では、NTCIR-10 特許機械翻訳タスクの概要、評価手法、評価結果について述べる。

2. 特許機械翻訳タスクの概要

特許機械翻訳タスクは、研究基盤を整備し、複数の翻訳手法に対して同じテストデータを用いた評価を実施して手法の有効性を明らかにすることで、特許翻訳技術の研究・開発を推進することを目的としている。以下にタスク実施の流れを示す。

- (1) 主催者が特許翻訳用の訓練データとテストデータを用意
- (2) 参加者が各自のシステムでテストデータを機械翻訳
- (3) 主催者が翻訳結果を評価
- (4) 参加者が研究成果をワークショップで発表

* 現在は NHK 放送技術研究所

** 現在は ChilLin Ltd, Hong Kong

*** 現在は Google Inc.

**** 現在は City University of Hong Kong / Hong Kong University of Science and Technology / ChiLin Ltd, Hong Kong

この活動を通して構築したデータ（訓練、テスト、評価結果、翻訳結果）は研究利用できるように管理される。NTCIR ワークショップは1年半単位で開催されている。第3回の NTCIR-3 から特許を対象としたタスクが実施され、第7回の NTCIR-7 から特許翻訳のタスクが実施されている。以下、NTCIR-10 特許機械翻訳タスクについて述べる。

実施した翻訳の言語対および翻訳方向は、日本語から英語（日英）、英語から日本語（英日）、中国語から英語（中英）である。

実施した評価は、文の訳質評価、特許審査評価、時系列評価およびマルチリンガル評価である。本稿では、これらのうちメインの評価である文の訳質評価と特許審査評価について述べる。

データは次の通りである。訓練データとして、日英・英日翻訳は約 320 万文対の日英対訳コーパス、中英翻訳は 100 万文対の中英対訳コーパス、翻訳先言語の単言語コーパスとして、日英・中英翻訳には英語の特許文 3 億文以上、英日翻訳には日本語の特許文 4 億文以上を参加者に提供した。開発データには 2,000 文対を用いた。これらのデータは NTCIR-9 で用いたデータと同一である。文の訳質評価のテストデータには新たに構築した 2,300 文、特許審査評価のテストデータには 29 文書を用いた。対訳コーパス、開発データおよびテストデータは、請求項の文を含まず、発明の詳細な説明部分の文からなる。

参加グループ数は全体で 21、日英翻訳は 13、英日翻訳は 11、中英翻訳は 12 であった。タスクに参加した翻訳エンジンの種類は、大きく分類すると統計翻訳（SMT）、ルールベース翻訳（RBMT）、用例翻訳（EBMT）、RBMT と SMT または EBMT との組み合わせ（HYBRID）の 4 種類である。

さらに、主催者がベースラインシステムの結果として 2 種類の SMT（フレーズベース SMT、階層フレーズベース SMT）、3 つの商用 RBMT^{1,2}、Google 翻訳による翻訳結果を追加した。利用したベースラインシステムを表 1 に示す。SMT ベースラインシステムは一般に公開されているソフトウェアで構築し、システムの構築と翻訳の手順は特許機械翻訳タスクの Web ページ³で公開している。ベースラインシステムの商用 RBMT のシステム ID は、本稿では匿名にしている。

表 1 ベースラインシステム

システム ID	システム	種類	日英	英日	中英
BASELINE1	Moses hierarchical phrase-based SMT system	SMT	✓	✓	✓
BASELINE2	Moses phrase-based SMT system		✓	✓	✓
RBMTx	The Honyaku 2009 premium patent edition（商用 RBMT）	RBMT	✓	✓	
RBMTx	ATLAS V14（商用 RBMT）		✓	✓	
RBMTx	PAT-Transer 2009（商用 RBMT）		✓	✓	
ONLINE1	Google online translation system（2012 年 10 月）	SMT	✓	✓	✓

¹ 3 つのうち人手で評価したシステムは NTCIR-9 で最も評価が高かったシステムのみである。

² 中英翻訳では、NTCIR-9 の評価で RBMT の訳質が SMT よりも低かったため、RBMT はベースラインシステムとして利用しなかった。

³ <http://ntcir.nii.ac.jp/PatentMT-2/>

3 評価手法

3. 1 文の訳質評価の手法

3. 1. 1 評価方法

評価者が文単位の訳質を評価した。各システムあたり 3 人の評価者がそれぞれ 100 文、合計 300 文を評価した。評価基準として、adequacy と acceptability の 2 つを用いた。これらは、主に翻訳結果を読んで内容を理解するという情報アクセスにおける有用性の評価基準である。これらの評価基準は NTCIR-9 で用いた評価基準と同じである。以下、本評価で用いた評価基準およびテストデータの構築方法について説明する。

3. 1. 2 評価基準

Adequacy

翻訳の適切さ (adequacy) の評価を 5 段階 (1-5) で実施した。この評価の目的は、システム間の比較である。本評価では、訳語選択が適切かどうかだけでなくフレーズや節の意味が正しいかも考慮して評価した。

Acceptability

図 1 に示す 5 段階の acceptability 評価を実施した。この評価の目的は、文レベルの意味が正しく伝わる文の割合を明らかにすることである。acceptability は入力文の意味が正しく伝わらない (たとえば、重要な情報が一つでも欠ける) と F になる。この評価は、adequacy と比べて、より実用に近い評価を目指している。例えば、システムへの要求水準が「入力文の意味が分かればよい」であれば、C 以上の評価となった訳文が有用であり、システムへの要求水準が「入力文の意味が分かり、かつ文法的に正しい」であれば、A 以上の評価となった訳文が有用である。このように、要求水準に応じた訳質の文の割合を明らかにすることができる。adequacy 評価の結果からは、このような文の割合は分らない。

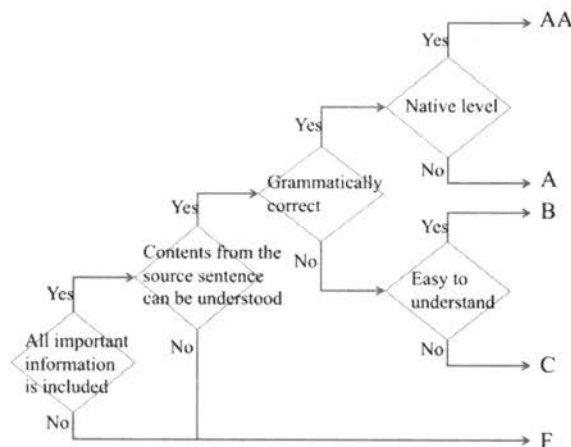


図 1: Acceptability

3. 1. 3 データの構築

訓練データは 2005 年以前の特許から構築しているため、テストデータは 2006～2007 年の特許から構築した。我々はテストデータを 2 つの方法で構築した。1 つ目の方法 (方法 1) は、自動的に抽出した対訳

文対からランダムに文対を選択し、さらにその中から人手で正しい対訳の文対を選択するというものである。2つ目の方法（方法2）は、原言語の特許文書からランダムに文を選択し、それらの参照訳は新たに翻訳して構築するというものである。2,300文のテストデータのうち2,000文は方法1で構築し、300文は方法2で構築した。方法1は、対訳文対の特許文書から自動的に抽出する際に実際の特許文の傾向が完全には反映されなくなる可能性がある。それに対して方法2は、前者の方法に比べて実際の特許文の傾向を反映できると考えられる。人手による評価は方法2で作成した300文に対して実施した。この構築方法はNTCIR-9からの改善点である。なお、2,300文のテストデータは自動評価値を計算する際に利用した。

3. 2 特許審査評価の手法

特許審査での有用性に基づく評価手法について説明する。特許審査官は、審査において既存の特許を調査して同じ技術が存在すれば、その既存の特許を引用して審査対象の特許を拒絶する。そのため、既存の特許に記載された技術的内容である事実を認定する必要がある。既存の特許が外国語で書かれていてその言語が分からない場合には、翻訳する必要がある。この翻訳を機械翻訳で行った場合に、翻訳結果から引用文書の特許を認定するために有用な事実をどれだけ認定できるかに基づいて、機械翻訳の有用性を評価する。本評価は日本知的財産翻訳協会（NIPTA）の協力により実施した。

3. 2. 1 全体の評価の流れ

ここで、全体の評価の流れについて説明する。

（1）準備 評価で利用するデータを構築する。まず審査の結論が拒絶の審決（審査の最終決定）とその審査で引用された特許を取得する。そして、審査において「引用文書から審査官が認定した事実」の根拠となる文を引用文書から抽出する。抽出した文をテストデータとする。なぜなら、この抽出した文は、「審査官が認定した事実」を表しているためである。そして、このテストデータが正しく翻訳されれば、翻訳結果から「審査官が認定した事実」を認定することができるためである。

（2）翻訳 テストデータを機械翻訳する。

（3）評価 翻訳結果から、「引用文書から審査官が認定した事実」をどれだけ認定できて、審査に有用であるかについて評価する。

3. 2. 2 評価方法

特許審査評価は日英翻訳と中英翻訳に対して行った。特許庁で審査官の経験があり、英語の能力が高い2人の評価者が評価した。1システムあたり各評価者が20文書、合計40文書（文書の重複を含む）の評価を行なった。評価したシステムは、adequacyの結果が上位のシステムから手法の多様性を考慮して日英、中英それぞれ3システムを選択した。

3. 2. 3 評価基準

特許審査評価で用いた評価基準を表2に示す。有用性の評価は、過去の審査で審査官が引用文書の特許から認定した事実を、引用文書を機械翻訳した結果からどれだけ認定できるかに基づいて行った。評価は引用文書単位で行なった。

3. 2. 4 データの構築

この評価に必要なデータを、審決を用いて以下のようにして構築した。

- (1) 結論が不成立（拒絶）の審決を取得する。
- (2) 審決中に記載されている「引用文書から審査官が認定した事実の説明」を抽出する。
- (3) 審査官が認定した事実を構成要素単位に分けて、それぞれの内容の根拠となる文を引用文書から抽出する。抽出した文を機械翻訳で翻訳するテストデータとする。

表 3 に、審査官が認定した事実と引用文書から抽出した文の例を示す。表 3 の左端の列は、審決中に記載されている審査官が引用文書から認定した事実である。表 3 の中央の列は、審査官が認定した事実を構成要素単位に分けたものである。表 3 の右端の列は、中央の列の事実を認定する根拠となった引用文書中の文を抽出したものであり、この文が翻訳するテストデータである。29 文書のテストデータを構築した。中英翻訳のテストデータは、日英翻訳のテストデータを日中翻訳して構築した。

表 2：特許審査評価の評価基準

VI	引用発明を認定するために有用な事実が全て認定できて、翻訳結果のみで審査可能
V	引用発明を認定するために有用な事実が半分以上認定できて、審査に有用
IV	引用発明を認定するために有用な事実が 1 つ以上認定できて、審査に有用
III	IVに至らないが、部分的に事実が認定できて、その文献が審査で無視できないことが分かる
II	一部の事実が認定できたが、審査に有用とはいえない
I	全く事実が認定できず、審査の役に立たない

(引用文書単位の評価)

表 3：審査官が認定した事実と引用文書から抽出した文の例

審査官が認定した事実	構成要素単位に分けた 審査官が認定した事実	引用文書から抽出した文 (テストデータ)
これらの記載事項によると、引用例には、 「内部において、先端側に良熱伝導金属部 43 が入り込んでいる中心電極 4 と、 中心電極 4 の先端部に溶接されている貴金属チップ 45 と、 中心電極 4 を電極先端部 41 が碍子先端部 31 から突出するように挿嵌保持する絶縁碍子 3 と、 絶縁碍子 3 を挿嵌保持する取付金具 2 と、 中心電極 4 の電極先端部 41 との間に火花放電ギャップ G を形成する接地電極 11 とを備えたスパークプラグにおいて、 中心電極 4 の直径は、1.2～2.2mm としたスパークプラグ。」 の発明が記載されていると認められる。	内部において、先端側に良熱伝導金属部 43 が入り込んでいる中心電極 4	また、図 3 に示すごとく、中心電極 4 の内部においては、上記露出開始部 431 よりも先端側にも良熱伝導金属部 43 が入り込んでいる。
	中心電極 4 の先端部に溶接されている貴金属チップ 45	また、中心電極 4 の先端部には、貴金属チップ 45 が溶接されている。
	中心電極 4 を電極先端部 41 が碍子先端部 31 から突出するように挿嵌保持する絶縁碍子 3	上記中心電極 4 は、電極先端部 41 が碍子先端部 31 から突出するように絶縁碍子 3 に挿嵌保持されている。
	絶縁碍子 3 を挿嵌保持する取付金具 2	上記絶縁碍子 3 は、碍子先端部 31 が突出するように取付金具 2 に挿嵌保持される。
	中心電極 4 の電極先端部 41 との間に火花放電ギャップ G を形成する接地電極 11	上記接地電極 11 は、図 2 に示すごとく、電極先端部 41 との間に火花放電ギャップ G を形成する。
	中心電極 4 の直径は、1.2～2.2mm	また、上記碍子固定部 22 の軸方向位置における中心電極 4 の直径は、例えば、1.2～2.2mm とすることができる。

4 評価結果

紙面の都合上、評価結果の要点のみ報告する。詳細な報告[2]および各グループの成果報告はオンラインで入手可能である⁴。システム名は、グループ ID（またはシステム ID）とプライオリティ番号の組で表示する。

4. 1 文の訳質評価の結果

評価結果を図 2 から 7 に示す。Adequacy の評価には平均値を用いた。

4. 1. 1 日英翻訳

図 2 と 3 が日英翻訳の評価結果である。JAPIO-1 と RBMT1-1 は RBMT, EIWA-1 と TORI-1 は HYBRID, KYOTO-1 は EBMT, それ以外は SMT である。RBMT の訳質が SMT より高いことが分かる。Acceptability が C 以上の割合は、RBMT のトップのシステム（JAPIO-1[3]）が 55%であったのに対し、SMT のトップのシステム（NTITI-1[4]）は 38%であった。C 以上の割合における RBMT に対する SMT の割合は NTCIR-9 では 39%（0.25/0.633）であったのに対し、NTCIR-10 では 69%（0.38/0.55）であり、割合が高くなっている。これは、RBMT に対して SMT の性能が向上していることを示している。NTCIR-10 でトップの RBMT の性能は NTCIR-9 でトップの RBMT と比べて同等以上と考えられるため、NTCIR-10 でトップの SMT の性能は NTCIR-9 でトップの SMT に対して向上しているといえる。

⁴<http://research.nii.ac.jp/ntcir/workshop/OnlineProceedings10/>

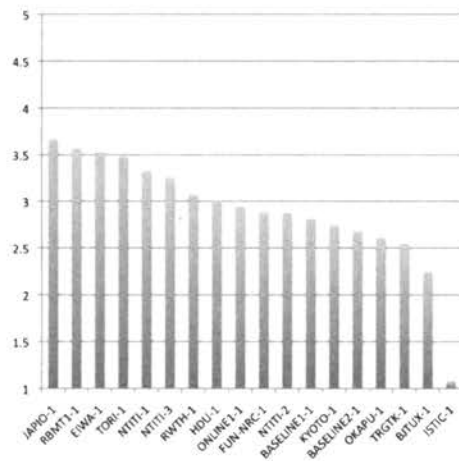


図 2: 日英 adequacy

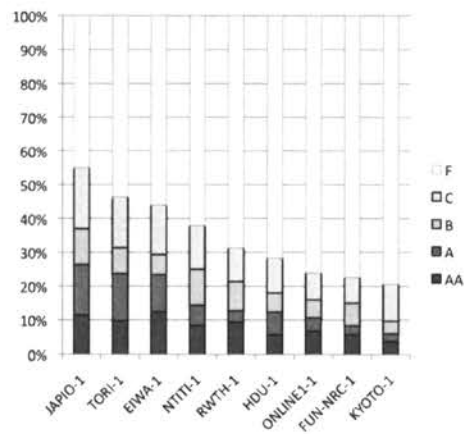


図 3: 日英 acceptability

4. 1. 2 英日翻訳

図 4 と 5 が英日翻訳の評価結果である。JAPIO-1 と RBMT6-1 は RBMT, KYOTO-1 は EBMT, それ以外は SMT である。トップの SMT (NTITI-2 および NTITI-1) の訳質がトップの RBMT よりも高いことが分かる。NTCIR-9 では、英日の特許翻訳で SMT がトップレベルの RBMT と同程度の訳質であったため、NTCIR-9 と比べてトップの SMT の性能が向上していることが分かる。NTITI のシステム[4]は、NTCIR-9 で NTT-UT システムが利用した「翻訳の前処理で英語の語順を日本語の語順に入れ替える手法[5]」を用い、さらに特許文の構文解析精度を向上させることによって、訳質を向上させた。トップの SMT のシステムはテスト文の 70%, トップの RBMT は 58% で acceptability が C 以上という結果を得た。

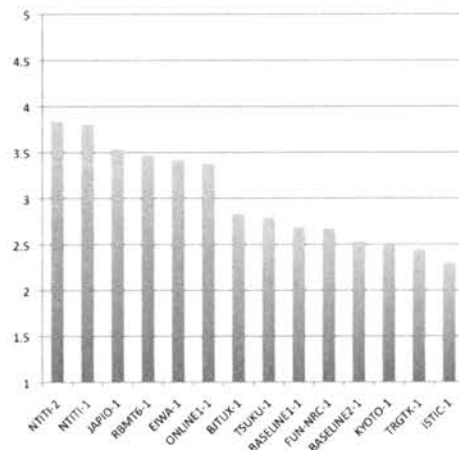


図 4: 英日 adequacy

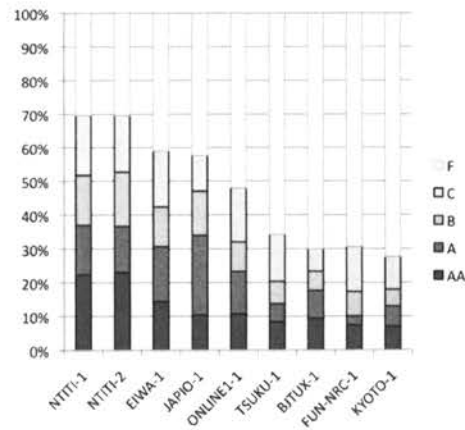


図 5: 英日 acceptability

4. 1. 3 中英翻訳

図 6 と 7 が中英翻訳の評価結果である。RWSYS-1 と EIWA-1 は HYBRID, BJTUX-2 は EBMT, それ以外は SMT である。中英翻訳では, NTCIR-9 で SMT が RBMT より訳質が高かったため, NTCIR-10 では RBMT の評価は行なわなかった。トップの BBN-1 システム[6]は, NTCIR-9 での BBN-1 システムに対して, 言語モデルや翻訳モデルなどをさらに改良し, テスト文の 67%で acceptability が C 以上という結果を得た。

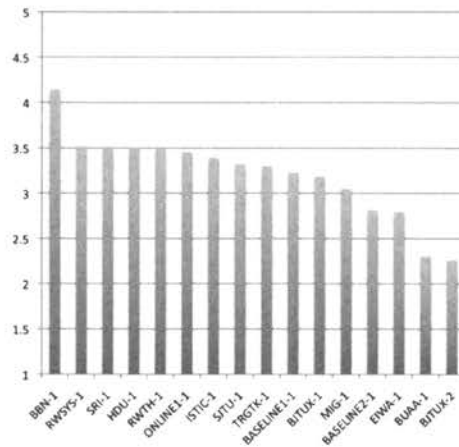


図 6: 中英 adequacy

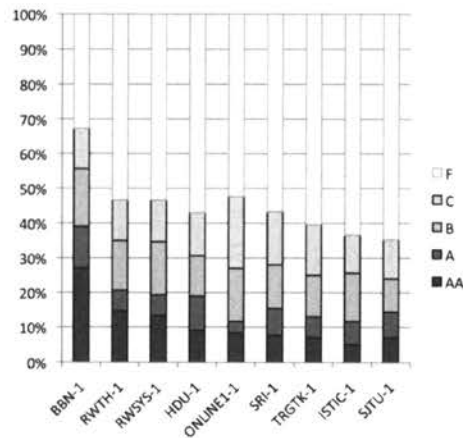


図 7: 中英 acceptability

4. 2 特許審査評価の結果

日英翻訳の評価結果を図 8 に，中英翻訳の評価結果を図 9 に示す．図 8 より，JAPIO-1 システム[3]は 6 割以上の文書で評価 VI を獲得し，全文書が評価 V 以上であった．これより，トップの日英翻訳システムの性能は特許審査で有用なレベルであることが分かった．SMT の NTITI-1 システム[4]も評価 V 以上が 6 割以上あり，ある程度有用であることが分かった．また，図 9 より，BBN-1 システム[6]は 2 割の文書で評価 VI を獲得し，9 割弱の文書で評価 V 以上であった．これより，トップの中英翻訳システムの性能も特許審査で有用なレベルであることが分かった．

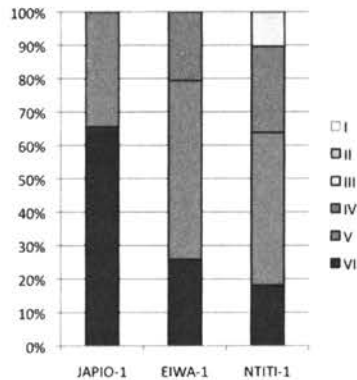


図 8: 日英 特許審査評価の結果

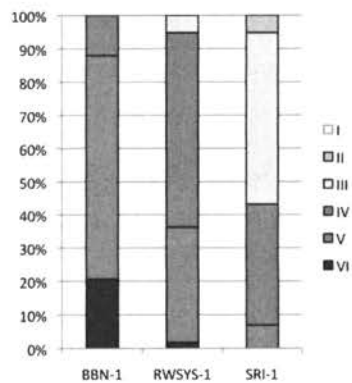


図 9: 中英 特許審査評価の結果

5 まとめ

NTCIR-10 特許機械翻訳タスクでの評価について述べた。この評価で実施した特許審査評価により、特許審査において、現在のトップレベルの機械翻訳システムが日英・中英翻訳で有用であることが明らかになった。また、英日翻訳ではトップレベルの SMT が RBMT よりも特許翻訳で高い訳質を達成したことが分かった。日英翻訳では、RBMT がトップレベルの SMT より依然として訳質が高かったが、SMT の性能が NTCIR-9 の時よりも向上していることが確認された。

謝辞

本評価に協力していただいた評価者およびデータ作成者に感謝します。また、特許審査評価のテストセット作成および評価に協力していただいた日本知的財産翻訳協会 (NIPTA) 石井理事長他に感謝します。

参考文献

[1] Isao Goto, Bin Lu, Ka Po Chow, Eiichiro Sumita, and Benjamin K. Tsou. Overview of the Patent Machine Translation Task at the NTCIR-9 Workshop. In Proceedings of NTCIR-9, pp. 559–578, 2011.

- [2] Isao Goto, Ka Po Chow, Bin Lu, Eiichiro Sumita, and Benjamin K. Tsou. Overview of the Patent Machine Translation Task at the NTCIR-10 Workshop. In Proceedings of NTCIR-10, pp. 260–286, 2013.
- [3] Tadaaki Oshio, Tomoharu Mitsuhashi and Tsuyoshi Kakita. Use of the Japio Technical Field Dictionaries and Commercial Rule-based Engine for NTCIR-PatentMT. In Proceedings of NTCIR-10, pp. 361–364, 2013.
- [4] Katsuhito Sudoh, Jun Suzuki, Hajime Tsukada, Masaaki Nagata, Sho Hoshino and Yusuke Miyao. NTT-NII Statistical Machine Translation for NTCIR-10 PatentMT. In Proceedings of NTCIR-10, pp. 294–300, 2013.
- [5] Hideki Isozaki, Tsutomu Hirao, Kevin Duh, Katsuhito Sudoh, and Hajime Tsukada. Automatic Evaluation of Translation Quality for Distant Language Pairs. In Proceedings of EMNLP, pp. 944–952, 2010.
- [6] Zhongqiang Huang, Jacob Devlin and Spyros Matsoukas. BBN's Systems for the Chinese-English Sub-task of the NTCIR-10 PatentMT Evaluation. In Proceedings of NTCIR-10, pp. 287–293, 2013.

Building Chinese-English Machine Translation Systems for Patent Documents

Zhongqiang Huang, Jacob Devlin, Jeff Ma, Rabih Zbib, and Rich Schwartz

1. INTRODUCTION

This paper provides a summary of the development of the Raytheon BBN Technologies systems [1] [2] for the Chinese-English sub-task of the NTCIR-10 Patent Machine Translation (PatentMT) evaluation [3]. The legal style of patent document, expressed in well-structured sentences and less ambiguous word meaning, makes patent documents easier to translate using Machine Translation (MT) systems. However, long and structurally complex sentences, as well as the frequent use of technical terminology and new terms, present unique challenges for MT. In this paper, we address some of the special challenges in patent documents and show how recent advancements in Statistical Machine Translation (SMT) improve patent translation. Starting with an adapted system that was originally designed for translating newswire data, we incrementally improved the translation accuracy on patent documents using various techniques. Our final system achieves state-of-the-art accuracy in both automatic and manual metrics in the NTCIR PatentMT evaluation.

2. TRANSLATION FRAMEWORK

Raytheon BBN Technologies has developed Statistical Machine Translation (SMT) systems for a variety of languages and genres, achieving top performance in official evaluations in DARPA sponsored projects, such as GALE and more recently BOLT. Our SMT systems are based on a string-to-dependency translation model [4], which employs hierarchical rules to translate source language strings to dependency trees in the target language. The systems utilize over a dozen features, including rule and lexical probabilities and a target language model (LM) score, in addition to about 50 thousand sparse features [5]. All feature weights are trained discriminatively to maximize expected BLEU [6]. We use a trigram LM to generate N-best hypotheses, which are then re-scored with a 5-gram LM. Word alignment models are trained with GIZA++.

3. EXPERIMENTAL SETUP

The data for the NTCIR PatentMT evaluation [3] consists of a parallel training corpus of one million sentence pairs (~45 million words), which we use to train the translation model. The data also includes a monolingual English corpus of US patent documents with 15 billion words. We train two separate language models. The first, which we refer to as the “45M LM”, is trained on the English side of the parallel corpus. The second language model, which we call the “14B LM”, is trained with the addition of the monolingual English corpus. We split the development set, consisting of two thousand sentence pairs, into a tuning set and a test set of roughly equal sizes.

4. BUILDING PATENT TRANSLATION SYSTEMS

For a system trained on the patent parallel data and the 45M LM, we get a mixed-case BLEU score of 34.01 on the test set. We improved on this result by addressing the special challenges of translating patent documents and

applying some recent advancements in SMT. Table 1 summarizes the resulting score improvements we obtained on the test set. We briefly describe these methods next.

4.1 Consistent Tokenization and Special Token Sharing

Chinese patent documents contain significantly more special ASCII strings (e.g. English words, patent numbers, mathematical expressions, and abbreviations) than Chinese newswire articles. Many of these strings were not aligned properly in the parallel corpus due to inconsistent tokenization on the source and target sides. We addressed this problem by tokenizing ASCII strings in Chinese sentences in the same way as in English, resulting in BLEU improvement from 34.01 to 34.56. We then replaced each type of special strings with a common token before training word alignment and language models, as well as before translating sentences, and replaced the original strings back in the translation output. This further improved BLEU to 34.97.

4.2 Re-training Casing Language Model

We originally performed true-casing of MT output based on a trigram casing LM that was trained on a collection of normal English text with mixed cases. We re-trained the casing LM using the mixed-case English sentences from the patent parallel corpus. This improved the BLEU score by 1.50 points to 36.47.

4.3 Optimizing Chinese Word Segmenter

Our Chinese word segmenter uses a simple left-to-right and longest-match-first algorithm based on a lexicon. The lexicon used in the previous experiments is a subset of a big Chinese word lexicon that was optimized for MT performance on newswire data. The optimization procedure starts with the big lexicon and iteratively removes words from the lexicon that are not aligned well in the GIZA++ alignments of the parallel training data. We stop removing words when the performance of the new MT system trained with the reduced lexicon stops improving. We re-optimized the lexicon on the patent parallel corpus, which improved the BLEU score to 36.95.

4.4 Miscellaneous Features

These include two sets of new features and various tweaks to the system. The first set of features is an extension of context-based lexical probabilities to model the joint likelihood of every target bigram given their aligned source words and the source context. The second set of features models MT hypotheses using traits [7], which are high-level characteristics of the output, such as the percentage of source context words that are not translated. The inclusion of the features and tweaks improved the BLEU score to 38.06.

4.5 CLIR-Based LM Adaptation

System	Test
Baseline with 45M LM	34.01
+ consistent tokenization	34.56
+ more token sharing	34.97
+ patent case-LM	36.47
+ optimized word segmenter	36.95
+ miscellaneous features	38.06
+ 14B LM	39.51
+ CLIR-based LM adaptation	40.95
+ robust context dependent translation	41.09
+ recurrent natural network LM	41.43
+ translation-based true caser	42.13

Table 1: Improvements in mixed-case BLEU on the test set resulting from various methods

Sentences in patent documents are well structured, and the word meanings are less ambiguous compared to newswire articles. Patent documents tend to re-use more n-grams than newswire articles. For example, we found that as many as 21% of the source 4-grams in the patent development set are also observed in the 45M patent parallel corpus, compared to only 5.4% for a newswire development set in a 227M word newswire training corpus.

We took advantage of this characteristic of patent documents by biasing the translation of a test sentence towards the passages in the 14B LM training data that are relevant to the test sentence, as measured by a cross-lingual information retrieval (CLIR) method [1]. This was achieved by scoring the translation hypotheses of a test sentence using an adapted LM, which is the linear interpolation of a general LM (the 14B LM in this case) and a biased LM trained on the relevant passages. In order to fairly evaluate the effectiveness of this approach, we re-trained the system to use the 14B LM, which improves BLEU from 38.06 to 39.51. The CLIR-based LM adaptation approach gives an additional improvement of 1.44 BLEU points, bringing the score to 40.95.

4.6 Robust Context-Dependent Translation

Our systems employed multiple context-based lexical translation and distortion models, which model probability distributions such as the distribution of target translations given the source word and its context. In order to alleviate the data sparsity issue that these models suffer from, we used a linear interpolation method:

$$p(e|c_1, \dots, c_m) = \sum_i w_i p(e|C_i)$$

where e is an event to be predicted based on the conditionals c_i in the context, C_i 's are ordered backoff contexts defined by heuristics, and each backoff model $p(e|C_i)$ is a maximum likelihood estimate (MLE) with fixed weight w_i . Unlike N-gram models in language modeling, where a clear backoff order exists, the optimal backoff order is not obvious in machine translation (e.g., whether the left context is more important than the right context). We addressed this problem by using a robust context-dependent translation method that considers all possible ways of backing off the context, and interpolates all of the backoff models together using optimized weights that depend on the reliability of the backoff components:

$$p(e|c_1, \dots, c_m) = \sum_i \frac{\delta_i + \gamma_i \log(N(C_i))}{\sum_j \delta_j + \gamma_j \log(N(C_j))} p(e|C_i)$$

where $N(C_i)$ is the count of component C_i in the training data, and parameters δ_i and γ_i are optimized on a held-out set. This model supports the intuition that more frequently observed components are considered more reliable. It provides an additional gain of 0.14 in BLEU to the system.

4.7 Recurrent Neural Network LM

We also investigated Recurrent Neural Network Language Models (RNNLM) [8]. A RNNLM uses the output of the hidden layer for word $(i-1)$ as part of the input layer for word i , effectively including the history of all previous words in the sentence. We trained a RNNLM on the 45M LM training corpus, and interpolated it with a 5-gram LM model trained on the same data. The scores computed from this model were used as an additional feature in n-best rescoring of MT hypotheses. The RNNLM contributes a gain of 0.43 BLEU points.

4.8 Translation-Based True Caser

We further improved on the trigram true caser by developing a new casing model similar to [9], which treats casing as translating from lower case sentences to upper case. The model uses several sparse features in addition to the probabilities of the rules and the LM. The improved true caser achieves a final mixed-case BLEU of 42.13.

5. EVALUATION RESULTS

We trained two systems for the evaluation. The primary system, labeled BBN-1, was trained on the 45M parallel training corpus and the 14B monolingual English patent corpus. The secondary system, labeled BBN-2, was trained on the 45M parallel training data and the LM data on the target side. Four types of evaluations were conducted for the NTCIR-10 PatentMT task: Intrinsic Evaluation (IE), Chronological Evaluation (ChE), Multilingual Evaluation (ME), and Patent Examination Evaluation (PEE). Our systems achieved the best performance among all submissions, by a significant margin, across all of the evaluation types and metrics.

System	IE	ChE	ME
Baseline1	32.52	30.74	17.96
Baseline2	31.34	29.34	18.05
BBN-1	42.68	39.44→41.09	27.62
BBN-2	39.98	36.69→38.93	N/A

Table 2: Automatic evaluation results in BLEU.

The arrow indicates improvement from NTCIR-9 to NTCIR-10 evaluation.

Table 2 shows the BLEU scores of our systems as well two baseline systems (Moses hierarchical and phrase-based SMT systems [10]) provided by the organizers, on multiple automatic evaluation types. Both BBN systems produced significantly better performance than the baseline systems in Intrinsic Evaluation. Compared to the NTCIR-9 results, our primary and secondary systems improved by 1.65 and 2.24 BLEU points respectively in the Chronological Evaluation. Our primary system also performed well in the Multilingual Evaluation.

In addition to automatic evaluations, the organizers also performed manual evaluation of the adequacy and acceptability of the translations of 300 sentences from the Intrinsic Evaluation set. Each translation was manually examined and assigned an adequacy score from 5 (best) to 1 (worse), as well as a 5-level acceptability score from AA (best) to F (worse). As shown in Figure 1, the BBN-1 system produced significantly better translations than the baseline systems in terms of adequacy, receiving an average adequacy score of 4.15. The distribution of acceptability score of the BBN-1 system is shown in Figure 2. The BBN-1 system also received a high pairwise acceptability score of 0.69, implying that its acceptability was much better than that of any other submission.

Patent Examination Evaluation (PEE) evaluates the utility of MT output for patent examination. Reference Japanese patents that had been used to reject real patent applications were first manually translated to Chinese, and then automatically translated to English using the participating systems. The translation outputs were rated by experienced patent examiners, based on the percentage of important facts that can be recognized from the translated text to reject the original patent application. As Figure 3 shows, a large portion of important facts was recognized from the translation output of the BBN-1 system. The translation results were considered useful in general for patent examination.

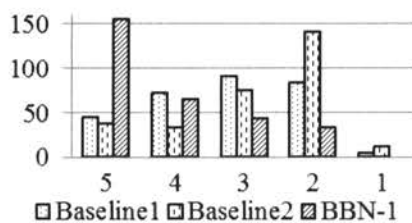


Figure 3: Adequacy

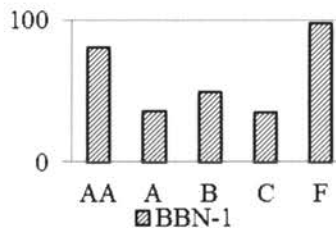


Figure 3: Acceptability

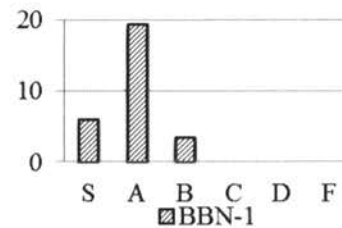


Figure 3: Patent examination

6. CONCLUSIONS

We have described the BBN systems for the Chinese-English sub-task of the NTCIR-10 PatentMT evaluation. We developed techniques to address the special challenges in translating patent documents and showed that some recent advancements in machine translation also improve patent translation. Our systems achieved the best results in both automatic and manual evaluations, and were judged useful for patent examination.

BIBLIOGRAPHY

- [1] J. Ma and S. Matsoukas, "Building a statistical machine translation system for translating patent documents," in *MT Summit XIII 4th Workshop on Patent Translation*, 2011.
- [2] Z. Huang, J. Devlin and S. Matsoukas, "BBN's Systems for the Chinese-English Sub-task of the NTCIR-10 PatentMT Evaluation," in *10th NTCIR Workshop Meeting on Evaluation and Information Access Technologies*, 2013.
- [3] I. Goto, K. P. Chow, B. Lu, E. Sumita and B. K. Tsou, "Overview of the Patent Machine Translation Task at the NTCIR-10 Workshop," in *10th NTCIR Workshop Meeting on Evaluation and Information Access Technologies*, 2013.
- [4] L. Shen, J. Xu and R. Weischedel, "String-to-dependency statistical machine translation," *Computational Linguistics*, 2010.
- [5] J. Devlin, "Lexical features for statistical machine translation," University of Maryland, College Park, 2009.
- [6] A.-V. I. Rosti, B. Zhang, S. Matsoukas and R. Schwartz, "BBN system description for WMT10 system combination task," in *Joint 5th Workshop on Statistical Machine Translation and MetricsMATR*, 2010.
- [7] J. Devlin and S. Matsoukas, "Trait-based hypothesis selection for machine translation," in *Conference of the North American Chapter of the Association for Computational Linguistics*, 2012.
- [8] T. Mikolov, M. Karafiat, L. Burget, J. Cernocky and S. Khudanpur, "Recurrent neural network based language model," in *INTERSPEECH*, 2010.
- [9] H. Hassan, Y. Ma and A. Way, "MaTrEx: the DCU Machine Translation System for IWSLT 2007," in *International Workshop on Spoken Language Translation*, 2006.
- [10] H. Hoang, P. Koehn and A. Lopez, "A unified framework for phrase-based, hierarchical, and syntax-based statistical machine translation," in *International Workshop on Spoken Language Translation*, 2009.

日英・英日翻訳に適した翻訳自動評価法

磯崎 秀樹

岡山県立大学 情報工学部

数年前、筆者は会社の研究所に所属していて、英語の医療関係の論文のアブストラクトを日本語で検索できる言語横断検索システムを作ってみようと思った。当時はテレビで医療関係の情報の捏造騒ぎがあり、筆者が前に作った、英語の情報を日本語で質問できる言語横断質問応答システムが好成績を収めていたからである。

医療用言語横断検索システムは、「臍がん」のような言葉を入力すると、それを辞書で英語に変換し、医療関係の英語論文のアブストラクトを検索する仕組みである。

同じグループで統計的機械翻訳(SMT)の研究もしていたので、その翻訳ソフトも組み込んでみたが、その翻訳結果には困ってしまった。因果関係がことごとく逆転しているのである。20年ほど前に機械翻訳を利用したときには、こんな間違いは見かけなかった。そこでインターネットのSMTのサービスを試してみると、やはり因果関係が逆転している。それどころか「メアリはジョンを殺した。」という短い文を「John killed Mary.」と訳し間違える。20年に比べて明らかに翻訳品質が悪い。

この原因は、SMTが欧米を中心に発展してきたためである。欧米の言語は語順が似ているので、単語を置き換えれば内容がかなりわかり、語順の問題を真剣に考える必要がない。SMTの研究者の話を聞いていると、語順の問題を解決するのに、過度の一般化をしているのが気になっていた。下手な一般化をすると、解ける問題も解けなくなるからである。

とりあえず、今必要な英日翻訳だけ、語順の問題を解決しようと考えた。すると日本語はhead finalなので、英語を解析して、headを後ろに持ってい

けばよい、ということに思い至り、英日翻訳の語順の問題はあっさり解決できた。

Headとは、ある表現の中で構文的な核となる部分であり、たとえば「美しい山」なら「山」、「山が美しい」なら「美しい」、「美しい山に登った」なら「登った」がheadである。つまり、修飾される言葉のことである。日本語では、修飾される言葉は修飾する言葉の後ろに来る。これがhead finalである。

英語をhead finalにすることで、どれくらい語順が近付いたかを計測することにした。語順の近さは順位相関係数「Kendallの τ 」で測定できる。しかし、Head finalにした英語(Head Final English, HFE)と日本語の参照訳の語順の近さは、すぐにわからない。どの英単語が日本語のどの単語と対応しているかわからないからである。

そこで、国立情報学研究所主催のNTCIR-7で作成された特許翻訳のデータにGIZA++というSMTのソフトを適用して単語の対応を求め、Kendallの τ を計算した。 τ は-1以上+1以下の値をとり、+1は完全に同じ順序、-1は完全に逆順序である。この測定により、HFEが日本語にかなり近いことが明らかになった。

NTCIR-7のレポートを見ていて、SMTの研究者が1ポイントを争う自動評価法BLEUが人間の評価とかけ離れていることに気付いた。ルールベースによる翻訳ソフトは、人間の評価では成績がよいのに、BLEUでは低く評価されている。こんな信用できないBLEUの1ポイントを争うことに意味があるのだろうか、と感じ始めた。

ここでふと思った。HFEと日本語参照訳の語順の近さを測定するには、単語対応を求めるために、

GIZA++を何時間か走らせる必要があった。でも、翻訳結果の日本語と参照訳の日本語の語順の近さを測定するには、GIZA++を起動する必要はない。同じ単語を見つけるだけである。すぐにプログラムを書いてみた。

実装上、いくつかの問題があった。同じ単語がなかったらどうするのか？同じ単語が複数あったらどうするのか？

両方の文に共通する単語だけで τ を計算し、一方にしか出現しない単語は無視することにした。同じ単語が複数ある場合は、前後の単語との組み合わせで判定することにした。そして τ を計測してみると、BLEU よりもずっと人間の評価に近いことがわかった。

しかし、共通する単語が2個しかなくてもいいのか？という疑問がわく。そこで共通する単語の少なさをペナルティとして扱うことにした。ここでペナルティの計算について、いくつかの選択肢がある。機械訳に現れる単語を基準とした単語適合率にするか、参照訳に現れる単語を基準とした単語再現率にするか、それとも、単語適合率と単語再現率の調和平均である単語F値にするか、という問題である。しかし、これらのペナルティをつけない τ そのもので十分人間の評価に近いので、これらのペナルティの影響はなるべく小さくしたかった。

試行錯誤の結果、 τ の値域を0以上1以下に変換して単語適合率の α 乗を掛ける、という最初のバージョンにたどり着いた。

τ に単語適合率を掛けると、 τ がマイナスの場合、単語適合率が高いほど成績が悪くなる、という問題があるので、 $(\tau + 1)/2$ という変換でマイナスにならないようにしてから単語適合率を掛けることにしたのである。 α の値をいくつに設定するかが問題であるが、これは実験によって決めた。

この自動評価法は RIBES (ライビーズ) という名前にした。Rank-based Intuitive Bilingual Evaluation Score の略であり、またカシスの仲間の植物を指す属名 (スグリ属) である。

RIBES は NTCIR-9 と NTCIR-10 で BLEU と並んで公式な自動評価法として採用されている。たとえば NTCIR-10 の結果によると、人手評価との相関係数は、RIBES が日英で 0.95、英日で 0.81 と非常に高いのに対し、BLEU は日英で 0.63、英日で 0.40 である。

参考文献

Hideki Isozaki et al. HPSG-based Preprocessing for English-to-Japanese Translation, ACM TALIP, Vol. 11, Issue 3, Article 8, 2012.

Hideki Isozaki et al. Automatic Evaluation of Translation Quality for Distant Language Pairs, in Proc. of EMNLP-2010, pp.944-952, 2010.

Hideki Isozaki et al. Head Finalization: A Simple Reordering Rule for SOV Languages, Proc. of WMT-2010, pp.244-251, 2010.

Isao Goto et al. Overview of the Patent Machine Translation Task at the NTCIR-10 Workshop, Proc. of NTCIR-10, pp.260-286, 2013.

知らなかったでは済まされない『翻訳の国際常識』

株式会社翻訳センター

品質管理推進部 田畠 奈々

国際ガイドランスは既に発行されていた

2013年6月、南アフリカ・プレトリアにて ISO TC37 の総会が開催された。

ISO は皆さんご存知のとおり。TC37 とは「言語、内容および情報資産」の標準化を扱うために ISO 内に設置された専門委員会 (Technical Committee) のことだ。この TC37 では翻訳の国際規格を制定するための分科委員会 (Subcommittee) を起ち上げ、2 年前から原案の作成を行ってきた。今年の総会では、この規格原案がついに DIS (Draft International Standard) と呼ばれるステージに突入したことが報告された。DIS ステージとは、それまで専門委員会の枠組み内だけで検討されてきた規格原案がついに表舞台に姿を現すステージのことである。このステージに突入すると業界関係者はこの規格原案を入手し、コメントをすることができる。日本としてこの原案に賛成か反対か、ひとつの答えを出さなければならない。5 か月の投票期限が設けられているのは、業界関係者で意見を統一する必要があるからだ。この投票で賛成可決されると、早ければ 2014 年には正式な国際規格として発行される。翻訳の国際規格の誕生というわけだ。

規格の名前は ISO 17100 - Requirements for translation services。概要については後ほどご紹介するが、その前に重要なことがもうひとつある。この規格、一体何をもとに作成されたのか、ということである。そう、実はこの規格のベースとなった国際ガイドランスなるものが存在するのだ。

この国際ガイドランスとは、ISO/TS 11669 Translation projects - General guidance のこと。2012 年 5 月に ISO から発行された。

もちろん、この原案を作成したのも TC37 であるが、プロモーションが不十分だったせいか、この存在を知る人は少ない。しかし、先ほど述べた翻訳規格の誕生がすぐそこに迫ってきた今、この国際ガイドランスの重要性が増している。これを知らないと世界から置いてきぼりを食らってしまう。

秘かに誕生していた国際ガイドランス、誕生の背景

ヨーロッパでは 1996 年以降、品質マネジメント規格である ISO 9001 は翻訳業界に馴染まないとの理由から翻訳に特化した規格を独自に制定する動きが活発化した。2006 年には、各国独自に制定した規格を統合する形で欧州規格 EN 15038 が誕生した。それをきっかけに、ヨーロッパ以外でもこの EN 15038 を参考に翻訳に特化した規格を制定する国が出てきた。つまり EN 15038 が徐々に浸透を始めたのである。これに焦りを感じた ISO が翻訳の規格制定に遅ればせながら乗り出したのは 2011 年のこと。全世界で適用できる規格を制定する、それが目的だった。しかし、いきなり規格を制定するのは現実的ではない。しかも、既に浸透を始めている EN 15038 を無視するわけにはいかない。そんな理由から、EN 15038 をベースにしつつ、世界各国が翻訳に対して共通理解が得られるようなガイドランス文書の作成に着手した。

国際ガイドランス、その内容とは

この文書は名前のとおり、ガイドランスである。翻訳者や翻訳会社が必ず遵守すべき要求事項を定めたものではない。あくまで、翻訳に対する世界の共通認識を明文化したものである。

とはいえ、侮るなかれ。このガイダンス文書、1～7章の構成で全35ページとなかなか中身が濃い。

1章の「適用範囲」にも記載されているが、このガイダンス文書が面白いのは、単に翻訳者や翻訳会社だけを対象に書かれたものではなく、翻訳を依頼する側、すなわちクライアントやエンドユーザー向けのガイダンスも兼ねている点である。当たり前だが、あらゆる関係者が共通の認識を持って翻訳プロジェクトに関わることが重要である、ということなのだ。

2章の「用語の定義」も実に面白い。最も基本的な translate という用語から、聞きなれない overt/covert translation という用語まで幅広くカバーしている。A-language, B-language, C-language という新たな概念も登場している。また、翻訳業界にとっては兼ねてより課題であった各プロセスの名称の不統一問題を解決する手がかりにもなりえる。つまり、check, revise, review, proofread などのプロセス名が明確に定義されているのだ。これで、翻訳会社によって review の内容が異なるなんて問題は解消されるかもしれない。これだけでも相当な内容が盛り込まれていることを分かって頂けるのではないだろうか。

3章では肝心の翻訳者や校正者の能力について記載されている。要求事項ではなく努力義務という扱いではあるが、翻訳の学位を持っている方が望ましいなど厳しい条件も含まれている。その中でも、特に注目すべきなのは校正者の能力である。

翻訳者と同等以上の能力が求められているのだ。翻訳者の卵が校正を担当することも多い日本にとっては、最も影響が大きい部分といえよう。とはいえ、翻訳者の翻訳文を修正する校正者が翻訳者よりも上位の位置づけであるべき、というのは納得性が高い。他国の委員から聞いた話によると、翻訳者よりも高い能力を持つ校正者を確保するのが難しいという状況は日本だけではないようなので、業界として「あるべき姿」を追求するということのようだ。

4～5章は翻訳前の準備から納品後のフォローアップまで細かな翻訳プロセスについての記載となっている。翻訳者や翻訳会社にとっては当たり前の内容であるが、翻訳業界に不慣れなクライアントには是非とも目を通してもらいたい部分。短納期や予算の制限など、翻訳プロジェクトに重大な影響を与える要素が列挙されている点も注目だ。

6～7章ではプロジェクト仕様書について触れている。なぜ仕様書が必要なのか、誰が作成すべきなのか、そしてどんな項目を含めるべきなのか。驚くなかれ、仕様書に含めるべきものとしてなんと21に及ぶ項目が紹介されている。このガイダンスで最も価値のある章である。例えば、項目の一つに「origin」というのがあるが、ここではソース言語を確認すべし、という単純な注意書きにとどまらず、「クライアントからフランス語から英語への翻訳依頼があった場合、実はその文書のオリジナルは英語であり、それをクライアントが知らないだけかもしれないので、既にその言語で文書が発行されていないか確認しよう」とか「原稿を書いた人のネイティブ言語は何なのかを確認しよう」など非常に細やかな例が示されている。翻訳会社の営業社員やプロジェクトマネージャーの教育資料としても使えるそう。

翻訳業界に与える影響

繰り返しになるがこの文書はガイダンスという位置づけである。今すぐに翻訳者や翻訳会社に影響を与えるものではない。しかし、冒頭で説明したとおり、このガイダンスをベースにした国際規格が早ければ2014年に発行される。

このガイダンス文書ではあくまで努力義務だったことも、規格となれば要求事項になることも考えられる。翻訳者や校正者であれば求められる能力について、翻訳会社であれば必須のプロセスについて早めに確認しておいて損はない。

要求事項でないにしても、一旦こういった形で指針が示された以上、翻訳者や翻訳会社が独自のプロセスをあたかも標準であるかのようにクライアントに押し付けるなんてことはできなくなる。逆もまた然り。そういった意味では、「本来はこうあるべき」という明確な指針が示されたことは翻訳者、翻訳会社、クライアントすべてにとって嬉しいことだ。

翻訳の国際規格、まもなく誕生

さて、本題の国際規格の方に話を移したい。既にご紹介したとおり、規格の名前は ISO 17100 - Requirements for translation services。業界関係者に情報が開示される DIS ステージに突入していることから、希望者は日本規格協会から原案を有料で入手することができる。名前からも分かるが、こちらは先ほどのガイダンスとは異なり、要求事項を示すもの、いわゆる認証規格である。1～6 章、全体で約 20 枚の構成。ガイダンス文書では丁寧な説明が多かったが、こちらはどちらかというと要求事項が淡々と記載されている、という印象。

翻訳者・校正者に求められるものとは

この規格で着目すべき点は大きくいうと 2 つ。1 つめは、翻訳者と校正者に対する要求事項である。この規格によると、翻訳者は「翻訳の学位」、「翻訳以外の学位+翻訳実務経験 2 年」、「翻訳実務経験 5 年」あるいは「政府認定の翻訳資格」のいずれかを有していなければならない。日本には「翻訳の学位」や「政府認定の翻訳資格」が存在しないことから、新規参入者にとっては厳しい条件となっている。しかし、この規格は製品規格のような位置づけになる可能性が高く、翻訳会社およびクライアントはプロジェクトに応じて ISO 準拠品かそうでないかを選択することができる。

つまり、実務経験のない翻訳者は ISO に準拠しないプロジェクトで経験を積んでいくという構造ができあがる。

これまで駆け出しの翻訳者や素人の翻訳者と同じような扱いしかしてもらえなかったベテラン翻訳者にとってはありがたい話だ。当然のことであるが、ISO 準拠品を担当する翻訳者の単価は高くなり、翻訳会社の売り価格もそれ相応のものとなると予想される。

校正者に対する要求事項にも着目しておきたい。校正者は翻訳者と同等以上の能力を有することが求められる。ガイダンスの方でも説明したが、翻訳者の卵が校正を担当することも多い日本にとっては見逃せないポイントである。実際、このような要求事項を満たせる校正者はどれくらいいるのだろうか。翻訳会社にとっては頭の痛い話になりそう。規格が発行されていない今の段階で予想するのは難しいが、この規格に準拠するとなると、ある翻訳者が翻訳したものを別の翻訳者が校正をするという方法しか考えられない。「校正者＝翻訳者の卵」という形態は徐々に減少し、「校正者＝翻訳者」「校正者＝翻訳者より上位」という構造が定着していくものと思われる。これによって、校正者の待遇は自動的に引き上げられることになると予想されるが、これまであまりの待遇の悪さに校正者を目指す人が少なかったこの業界の問題がいよいよ解決されるきっかけになれば、と不安よりも期待が高まる。これまで誰が校正を行っているのか不透明だっただけに価格の引き上げには納得がいかなかったクライアントも、この要求事項を見れば納得をせざるを得ないだろう。

翻訳会社が遵守すべきプロセスとは

では、2 つめを紹介しよう。この規格では標準プロセスとして、「翻訳(translate)」→「翻訳者による自己点検(check)」→「校正(revise)」→「レビュー(review)」→「最終校閲(proofread)」→「最終検品(final verification)」というフローが示されている。

校正(revise)とは翻訳者以外の第三者がソース言語とターゲット言語の両方を見るバイリンガルチェック、レビュー(review)とはその分野の専門家がターゲット言語のみを見るモノリンガルチェックといった具合にひとつひとつのプロセスが定義づけられている。また、レビュー(review)と最終校閲(proofread)は、クライアントの要望があった場合に実施するオプションのプロセスとなっているが、その他は必須のプロセスだ。翻訳会社が ISO 準拠と謳うからには、校正(revise)というプロセスを省略してはならない、ということになる。

どのように導入されていくのか

この規格がどのように業界に導入されていくのか、今の時点では明らかにされていない。実際、TC37 内にはプロモーションのための委員会というものがあるが設置されており、今後その委員会の中で方針が詰められていくものと予想される。しかし、既に発行されているガイドンスとは明らかに位置づけが異なってくることは間違いない。ISO の他の規格と同様に認証機関が選定されるのか、はたまた自己宣言という形になるのか。実際にこの規格を読んで頂ければわかるが、すべてのプロジェクトをこのガイドンスに準拠して進めるのは決して容易ではなく、また現実的でもない。仮にすべてのプロジェクトでこの規格を遵守できたとしても、場合によってオーバースペックとなる可能性があるし、オーバースペックの製品に対して高額を支払わなければならないクライアントが納得するわけがない。おそらく、導入後は ISO に準拠した翻訳プロジェクトとそうでないプロジェクトを明確に区別していくことになるだろう。

翻訳者や翻訳会社であれば、過去に一度は「こんな安い価格と短い納期で完璧な翻訳物を求められても困る！」という経験をしたことがあるのではないだろうか。

一方、「かなり高額の見積もりだったので当然翻訳にはネイティブチェックが含まれていると思いきや、ろくにチェックも入っていない翻訳を受け取ってしまった…」というクライアントも少なくないはず。適正なプロセスを明確にし、それに対して適正な価格を設定する。翻訳者、翻訳会社、クライアントすべてにとってフェアな時代がやってきたといえる。翻訳者や翻訳会社が、万が一この規格を「クライアントに要求されたから仕方なくこの準拠するもの」と捉えてしまったとすれば、それはあまりにも勿体ない、というのが筆者の個人的な考えだ。翻訳業界全体がこの規格をむしろ前向きに受け止め、クライアントやエンドユーザーを啓蒙し、業界全体のさらなる発展に繋げていきたいところである。

Translation before Demand

YAMAGATA EUROPE
Geert BENOIT

Manga

Yamagata Europe is a Belgian company with Japanese roots, a perfect fit: Belgians and Japanese both passionately love manga. The most famous Belgian manga figure is undoubtedly "Tintin". Tintin represents focus on detail and quality, but because we want to focus on speed today, I would like to introduce another beloved Belgian manga figure: Lucky Luke. Lucky Luke is a character that was created by Morris. Lucky Luke is known as the cowboy who can shoot faster than his own shadow. This level of speediness should be the ultimate target for our translation services.



Speed is a mentality

The environment in the Far West has changed: we are surrounded by Big Data, and even more with stories about Big Data. It indeed looks like the Big Data hype is growing bigger than the Data itself.

Big Data stories usually talk about the 3V's: Volume, Velocity and Variety.

In this case study we will only focus on the Velocity aspect. We will not focus on the Velocity that enables us to process big volume jobs, nor on the speed and power offered by MT and other tools to tackle the massive data volumes. No, we want to micro-focus on the speed aspect in order to handle the small chunks of data that together make the big data big. We will show how MT can be an enabler to bring data to the user immediately.

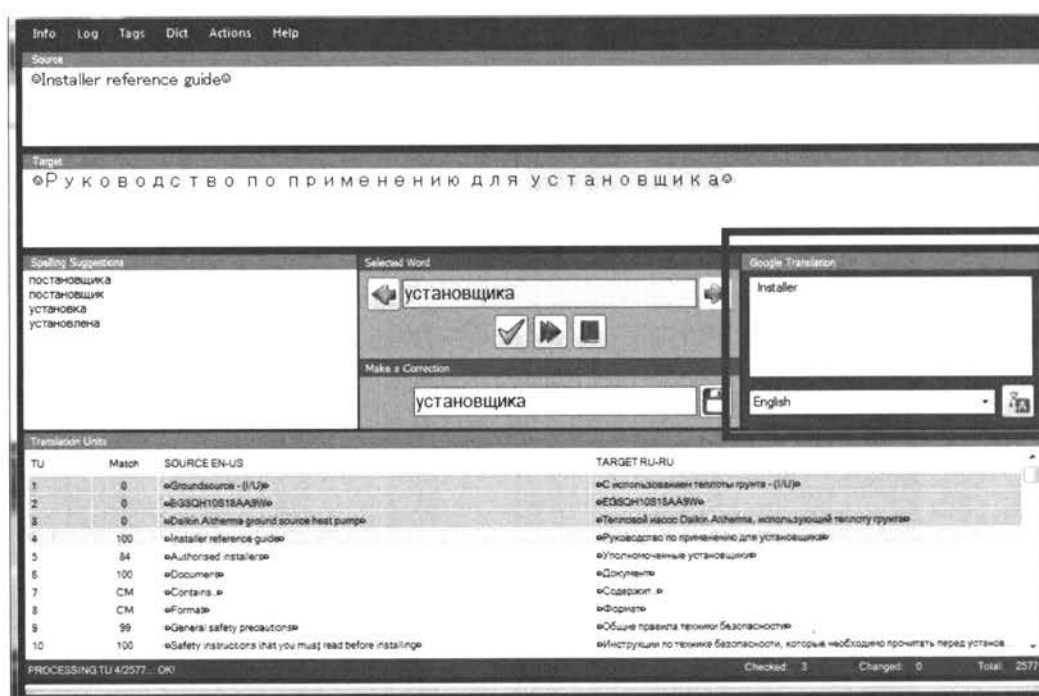
Speed should be a mentality, and MT is just a tool to help us to realise our targets, MT is not a target in itself. So this article is not about MT, it is about speed. Personally I feel that MT is

too often seen as the holy grail, the ultimate technological target in itself. MT should be considered as just another tool to plug into your process.

In our organisation we have seen that MT – when available at hand for different users - quickly find its way where speed is needed. Therefore, we will look at three applications where MT is embedded for speed reasons. The love and need for speed brought us to MT.

SnellSpell:

A first tool where we integrated MT is in our spellchecker called “SnellSpell”. We developed this spellchecker a few years ago because we needed a strong and fast tool to batch-process multiple files (sometimes +100 files) in one spellcheck action. Also we needed a tool that created a log file as evidence that the spell check was actually done. The translator who works for Yamagata uses SnellSpell to do all spellchecking after translation. At Yamagata, we use SnellSpell to double check whether the spellcheck was actually done. We use the log file that tracks all movements as evidence in case of discussion afterwards. However because we obviously do not master all languages at Yamagata Europe, we build in an MT function to double check the meaning of highlighted potential spelling problems. The tool gives us an MT translation (using the Google API) of the highlighted words in languages we do not fully master.



Strawberries on demand

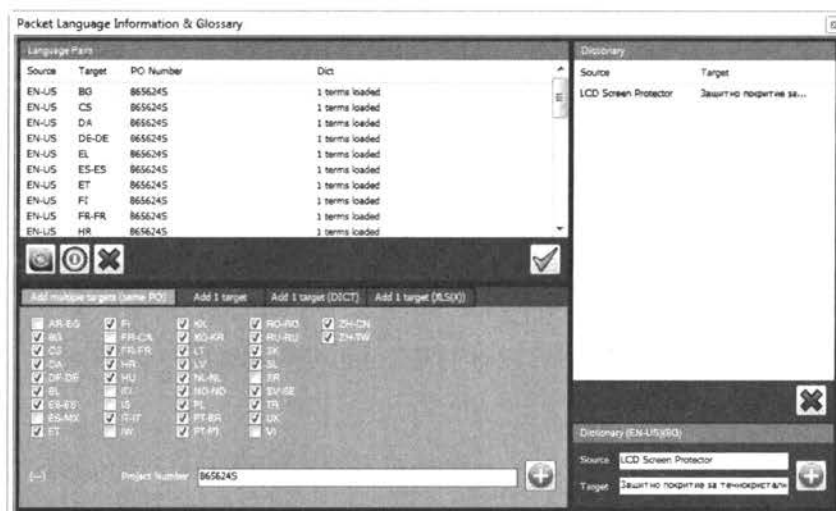
A second tool where we embedded MT is “Ichigo”. Ichigo means “strawberry” in Japanese, but it can also mean “one (ichi) sentence/word (gyo/go)”. Ichigo was developed because we

wanted to radically simplify and speed-up the administrative side of the process for really small and additional translations (one sentence or one word).

In a complex organisation like Yamagata where we handle many different translation projects, output formats, glossaries, ... we noticed that the administrative systems got far too complicated for small and fast jobs. The typical Japanese love for detail and fear for failure had resulted in unnecessary complicated instructions to be submitted with each tiny translation job. Comes in the need for simplicity. The answer is Ichigo: a handy tool that offers everything in one box: ordering, instructions, translation interface, QA, invoicing, ... One toolbox for all parties involved: end client, MLV, translator, accounting.



With Ichigo successfully running for one year, we found it to be a small step to add MT to this process. Today MT allows us to immediately feedback translation to the client as raw MT output, or to pre-translate and have the translator post-edit the MT output. We plug in Systran, KantanMT or Google depending on the language combination and/or customer needs.



Translation before demand:

The third project is born out of the Honda project that was described in the September Issue of the AAMT Journal by Heidi Van Hiel (page 20). For this case, we start from two basic assumptions:

1. The client has multilingual user generated content stored in databases.
2. The data owner knows to identify the content fields in the database that potentially need to be understood in another language.

Because the translation industry has been focusing on the translation and processing of "on purpose created documentation", we tend to believe that all translatable data is stored in well managed Content Management Systems. Not true of course.

Honda provided us with a project where we started to translate warranty claim information and user generated content about technical issues of the Honda products. This data was available in the local language of the end user (car/bike dealer) in the Honda warranty system. After analysis of the market and the market shares, Honda decided to translate only three languages into English to cover a very broad scope of the warranty issues. With this system, and together with Honda we started to use MT as an integrated translation component to handle content in content management systems with un-structured content.

In this first phase, Honda looked at the major languages and the statistically minimal option in order to capture most warranty issues. Today Honda has a different approach and asks us to translate more than 10 languages including even languages where the translation need is not yet proven. Adding a language and supplying the translation is considered cheap enough to just do it. When the translation would eventually be needed no time is lost and immediate action can be taken based on the available translated data. We call this: translation "before" demand.

Fast data = translated knowledge

The MT integrated translation workflows can bring assumedly interesting data in a speedy way to the user at virtually no cost, as such bringing the translated content to the user before he sees a need to get it translated. The MT output is content in the language of the user, content that can be post-edited to a perfect translation if needed.

We see similar MT requests from other clients often including projects where content is being made available in other languages to speed up processes, without specific requests from the end users.

This is what we call “Fast Data” at Yamagata Europe: translation gets done faster than our shadow.... Thanks Lucky Luke, for setting a great example!

「用途に合わせた機械翻訳のススメ」

株式会社十印 代表取締役社長

渡邊 麻呂

本稿では、産業翻訳会社（Language Service Provider、以下、LSP と記す）として、数多くのユーザーの様々なドキュメントを翻訳した経験から、LSP にとっての機械翻訳（Machine Translation、以下、MT と記す）の現状を述べたうえで、様々な考察をおこない、最後に MT の未来像を提示してみたい。

なお、ここに記述することは、筆者が TAUS のセミナーで発表した内容に沿ったものである。

1. LSP における MT の現状

現在の産業翻訳において、ユーザーから対価を取って納品する翻訳物は、専門の翻訳者の手による翻訳作業により生み出されている。実際のワークフローにおいては、翻訳者の翻訳で完成物にはならず、専門のチェッカーによる QA（Quality Assurance）工程を経た後に完成物、いわゆる納品物となるものである。

そこで、この翻訳というワークフローにおける、現在の MT は以下のように位置付けられるのである。

- ・平均的な MT の使用率は、約 10%
- ・MT を適用する言語の組み合わせは、英文多言語、英文和訳、和文英訳の順
- ・MT を適用するドキュメントの分野、種類は限られている
- ・MT で生成された翻訳文のみでは完成物にならず、ポストエディットが必要となる

このポストエディットは QA 工程に比べ労力を要する

- ・MT を使用することで、翻訳者の翻訳に比べ、リードタイム、工数は減少するが、劇的な削減効果

は得られていない

MT エンジン、および、MT ソフトウェアを開発しない多くの LSP は、既存の MT ソフトウェアを自社のワークフローに組み込み使用している。ワークフローにおいて、MT の翻訳精度の向上、適用可能と思われるドキュメントの分野、種類の拡大 等、向上を望むところはあるが、現状ではワークフローの工夫により補っている。この事は、MT が翻訳者による翻訳に敵わないまでも、翻訳者を補完する役割を果たしているものではないと考えられる。

2. MT に対するユーザーの評価

LSP に翻訳を依頼するユーザーは、納期を守り、正確で間違いのない翻訳を、適切な価格で納品されることを LSP に望んでいる。一見すると、「納期」、「品質」、「価格」の三要素に収斂されるように見えるが、それぞれの要素には様々な変数をあてはめることで、ユーザーが LSP に求める納品物が決まるのである。ここにおいて、ユーザーが「納期を早く」、「価格を安価に」と考える際に、LSP に MT の使用を打診することが見受けられる。このことは、ユーザーは翻訳者よりも「早く」、「安価」に翻訳ができるソフトウェアが MT だと評価しているものである。しかし、多くのユーザーが「早く」、「安価」だけでなく、「品質」も望むのが現実である。ここには、MT の翻訳精度への期待もあると捉えるべきである。

3. LSP が考えるあるべき姿

ここで、LSP として、現在の MT をより積極的に使用することを考察したい。

近年、ユーザーが MT ソフトウェアを導入して翻訳

をおこない、LSP に MT で生成された翻訳文のエディットを依頼する事例が増えている。この場合、ユーザーは MT の翻訳結果に満足せず、LSP は後工程のエディットのみをおこない、そこで出される MT への要望は MT 開発者に届かないという図式になっている。

そこで、ユーザー、LSP、MT 開発者のあらたな枠組みを考えてみたい。

ユーザーは翻訳に求める結果のみを LSP に伝え、LSP はユーザーが求める結果を出すワークフローを考える、そして、LSP は言語の専門家として MT への様々な要望を MT 開発者につたえ、MT 開発者は MT の向上をはかる、という枠組みである。これを、三者の協業体制として提案するものである。

筆者が、TAUS のセミナーで講演した際、この協業体制を述べたことに対して、自社内で MT を使うソフトウェアベンダーから、「そのような三者協業を目指さなくても、MT のポストエディットを LSP に発注することで、良い関係が築かれているのではないか。」と指摘があった。

そこで、筆者は会場の出席者に向かって「この中で、今のようなポストエディットをおこなうことがハッピーですか？」と問うたところ、同意をしたのは一人だけで、多くの出席者は MT のポストエディットをおこなう LSP という構造に、疑問を感じていることが分かった。

このことから、ユーザー、LSP、そして MT 開発者の従属関係ではなく、皆がハッピーになる三者協業体制の必要性を感じた次第である。現在、弊社では三者協業を実現しており、良い環境を築いている。

弊社の実績から、LSP が先導して用途やユーザーによって必要とされる品質基準を変えることにより、今の MT を使用方法を見出してみたい。その過程で、MT のほうが人間翻訳よりも優れている

ポイントも提示できると考えるものである。

この前提のもとで、「ユーザーが翻訳に求める結果」を「用途と品質」に置き換えて、「用途と品質」別にいくつかのドキュメントを MT で翻訳した例と、その評価、翻訳者による翻訳との比較を示してみたい。

この事例は、MT を組み込んだワークフローによるものである。

〔事例 1〕

用途：社内向け会報

品質：なんとなく意味がわかれば OK

（原文）

It was very helpful to see all those concerned with the project last Friday. I'd like to summarize the major points of the meeting. We agreed that this is the best time to enter the mobile phone market.

（訳文）

それらがすべてこの前の金曜日、プロジェクトに関係があったのを見ることは非常に有用でした。私は、会合の主な主張の要点を要約したい。私たちは、これが携帯電話市場に参入する最高の時間であることに合意しました。

メリット：低コストで迅速に社員へ情報を伝えることができる

コスト：50～90%削減 納期：50～90%短縮

〔事例 2〕

用途：エンドユーザー向け技術文書

品質：わかっている人が読むので読めれば OK

（原文）

You can place a menupopup inside the button to cause a menu to drop down when the button is pressed, much like the menulist. However, in this case you must set the type attribute to the value menu.

(訳文)

ボタンが押されたときに、メニューをドロップダウンさせるために menulist のように、ボタンの内側に menupopup を置くことができます。ただし、この場合、型属性を値メニューにセットする必要があります。

メリット: 品質はそれなりだが短納期&低コストで翻訳できる

コスト: 10~50%削減 納期: 20~60%短縮

[事例 3]

用途: ソフトウェアユーザーガイド

品質: 量が多く、急いでいる。品質もそれなりに。

(原文)

Procedure 5-2: To alter the automated nightly data collection

1. Select [Tasks & Logs]! [View All Scheduled Tasks...] from the top menu bar.
2. From the list of current tasks, select [Collect ALL Capacity Advisor Data Nightly....]
3. If you would like to modify the timing, click [Edit]. The [Task Confirmation] screen is displayed.
4. Click [Schedule].
5. Select [Periodically] from the list labeled [When would you like this task to run?].

(訳文)

手順 5-2: 自動的な夜間のデータ収集を変更するには

1. トップメニュー・バーから [タスク&ログ] → [すべてのスケジュールされたタスクを表示...] を選択します。
2. 現在のタスクのリストから、[夜間にすべてのキャパシティー・アドバイザー・データを収集] を選択します。
3. タイミングを変更したい場合は、[編集] をクリックします。[タスク確認] スクリーンが表示されます。

4. [スケジュール] をクリックします。

5. [いつこのタスクを実行しますか?] というラベルのリストから [定期的] を選択します。

メリット: 品質もある程度担保しつつ短期間でリリースできる

コスト: 10~30%削減 納期: 10~40%短縮

[事例 4]

用途: ソフトウェア仕様書

品質: 人間翻訳と変わらないレベルで

(原文)

The main program also instantiate the visualisation manager (SSAVisManager) before transferring control over to the G4UI, so that the user can change the parameters associated with the SSAT execution, including the number of events to be run.

(訳文)

また、メインプログラムは、実行するイベントの数など、SSAT 実行に関連付けられているパラメータをユーザーが変更できるように、G4UI に制御を転送する前に、可視化マネージャー (SSAVisManager) をインスタンス化します。

メリット: 低コストは難しいが、高品質&短納期でできる

コスト: 0~20%削減 納期: 10~30%短縮

このように、あらかじめ個々のドキュメントの翻訳に対して「用途と品質」を設定したうえで MT の翻訳をおこなえば、その「用途と品質」は充分満たされると考えられる。

そのうえで、使用する MT のクセを分析して、MT が解釈しやすいように原文に手を加える「原文プレエディット」をおこなうことで、訳文の向上が認められた。

また、対象ドキュメントの翻訳において MT を継続することで以下のようなメリットも生み出されることが分かっている。

- ・作成した辞書が無駄にならない

（一回のみの MT の使用では、作成した辞書は一回限り）

- ・MT のクセを習得し、希望の翻訳が出てくるような編集能力が習熟できる
- ・統計ベースの場合はコーパスが増えていくことにより翻訳精度が向上する

また、以下のように MT が翻訳者による翻訳よりも優れている点も挙げられる。

- ・数字を間違えない
- ・用語を間違えない
- ・常に同じ訳文が出る
- ・訳文の統一がとれる
- ・大量に処理できる
- ・同じ間違いをする（対策を立てられる）

このように、「用途と品質」を設定しうえて、継続使用することともに、「機械翻訳が優れている点」を考慮すれば、MT は翻訳現場で使えるものとするものである。

参考までに、機械翻訳に向いているドキュメントの種類を以下に提示する。

規格文書、仕様書、ソフトウェアユーザズガイド、リリースノート、機械関係の捜査説明書、そして、大量のドキュメント

また、マーケティング系文書においても、短文のコピーライトは難しいが、長文のコピーライトなら十分役立つという結果を得ている。

4. MT の将来像、そして未来

三者協業体制により、MT の精度向上、MT に対応するドキュメントの「用途と品質」は拡大するものと思われる。その先には、ユーザー、LSP、MT 開発者が連携し翻訳工程のバックグラウンドを意識せず、クラウド上に展開するグローバルで新たなワーク

フローが実現できると確信するものである。

コンピュータ処理用日本語表現辞書：JDMWE

首藤 公昭

背景と概要

日常の広範な文を対象とする機械翻訳を始めとする自然言語処理(NLP)では単語と接辞の辞書だけでは語彙資源として十分ではなく、必要に応じてフレーズを処理単位と捉えることが不可欠である事はよく知られており、近年、フレーズベース統計翻訳などの研究が盛んに行なわれています。(文献2、5、6など) しかし、統計的手法においては、相当大規模なコーパスでも言語表現の生起に希薄性が残るロングテール問題や考慮すべきフレーズの種類や特徴をどのように規定するかなどの難しい問題があり、実用レベルの十分な成果はまだ見られないようです。特に、フレーズを単語のようにカプセル化して扱うのは柔軟性を欠くため、必要に応じてギャップを含む不連続フレーズを扱うことが不可欠ですが、このことを考慮すると統計翻訳はさらに難しさが増すと思われます。筆者はフレーズ対フレーズを基本にした機械翻訳のための日本語フレーズ辞書(JDMWE: Japanese Dictionary of Multiword Expressions)を人手で構築する作業を古くから行なってきました。他言語のフレーズと対応を取るパラレル化はこれからの課題ですが、日本語側の辞書としては一定のレベルに達したので、公開に踏み切りました。詳細は文献8、9に譲り、本稿ではその現状を簡単にご紹介いたします。

NLPの世界では2002年のI. A. Sag氏らの論文、“Multiword Expressions: A Pain in the Neck for NLP”(文献7)がきっかけとなって、特異表現=複単語表現(MWE: Multiword Expression)の重要性が再認識され、この種の表現辞書の構築に向けて、人手による方法、統計処理を援用する方法等が世界中で試みられていますが、まだ有効かつ包括的な成果

は得られていません。

慣用句の様な意味の判定が難しい表現や塊りとして扱った方が得な決まり文句的表現を出来るだけ網羅するというのが当初からのJDMWE開発の基本方針で、不特定多数の出版物(新聞、雑誌、小説、辞書等)やテレビ、ラジオの放送文から内省によって約111,000表現(見出し)が採集されています。JDMWEの特徴としては、収録表現の網羅性のほか、基本的に不連続フレーズを扱っていること、記載情報の多様性、処理への即応性などが挙げられます。

JDMWEの編成・規模

JDMWEは次の様な部分辞書から構成されています。(ただし、これらは必ずしも排他的ではありません。)

1. 日本語 MWE 辞書__慣用句編

「油を売る」、「顔を洗って出直す」、「真っ赤なウソ」などの慣用句約4,400表現収録。

2. 日本語 MWE 辞書__動詞性表現(I類)編

「手を結ぶ」、「思いが募る」、「沽券に関わる」のような『名詞+を+動詞』、『名詞+が+動詞』、『名詞+に+動詞』の形式の表現、約35,000を収録。

3. 日本語 MWE 辞書__動詞性表現(II類)編

「骨の髄までしゃぶる」、「猿も木から落ちる」、「ゼロからやり直す」など動詞が支配するI類、III類以外の表現約13,000を収録。

4. 日本語 MWE 辞書__動詞性表現(III類)編

「放り出す」、「飲んだくれる」など、単語性や構成性に疑問が残る、複合動詞を中心とした比較的単純な構造の表現約3,500を収録。

5. 日本語 MWE 辞書__形容詞性表現編
「頭が痛い」、「傷跡が生々しい」など、形容詞が支配する約 4,800 表現を収録。
6. 日本語 MWE 辞書__形容動詞性(様態)表現編
「痘痕も笑窪」、「願ったり叶ったり」、「驚く程」など、広義の形容動詞性ともいえる様態表現約 10,000 を収録。
7. 日本語 MWE 辞書__副詞性表現編
「思いがけず」、「心を鬼にして」など、連用修飾句としてよく用いられる約 16,000 表現を収録。
8. 日本語 MWE 辞書__連体詞性表現編
「気のきいた」、「世に云う所の」、「トントン拍子の」など、連体修飾句として良く使われる約 15,000 表現を収録。
9. 日本語 MWE 辞書__名詞性表現編
「天下の一大事」、「無二の親友」など、名詞性の約 18,000 表現を収録。
10. 日本語 MWE 辞書__連結詞・文副詞性表現編
「そうは言うものの」、「このような状況で」など、文頭で用いられる連結詞・文副詞性表現、1,100 を収録。
11. 日本語 MWE 辞書__オノマトペ表現編
「グラグラ」、「シュルシュルと」、「ガッツリ食う」など、擬声・擬音・擬態語およびそれらを用いた典型表現約 12,000 を収録。
12. 日本語 MWE 辞書__格言・諺・成句・決まり文句編
「急がば回れ」、「義を見てせざるは勇無きなり」などの約 3,900 表現を収録。
13. 日本語 MWE 辞書__四字熟語編
「四面楚歌」、「切磋琢磨」などの約 500 表現を収録。
14. 日本語 MWE 辞書__挨拶・呼びかけ・応答・独言表現編
「何という事だ」、「御苦労さま」など良く使われる挨拶・呼びかけ・応答・独言表現約 450 を

収録。

15. 日本語 MWE 辞書__クランベリー型表現編
「しがみつく」、「後ろめたい」などのクランベリー型表現候補約 40 を収録。
16. 日本語 MWE 辞書__機能語(助詞、助動詞)性表現編
「を理由に」、「において」、「における」、「ほうがいい」、「てください」など、助詞、助動詞相当の機能を持つ表現約 4,000 を収録。

記載情報

上記の JDMWE 各辞書には、原則として以下の a～f の情報が記載されています。

a. 区切り、表記情報

辞書見出しは「ばかをみる」のように平仮名ベタ表記で与え、「ばか/を/みる」と要素単語に区切られること、そのうち、「ばか」は「バカ」、「馬鹿」、「莫迦」、「みる」は「見る」、「観る」と表記可能であることが記載されています。従って、「ばかをみる」、「バカをみる」、「馬鹿をみる」、「莫迦をみる」、「ばかを見る」、「バカを見る」、「馬鹿を見る」、「莫迦を見る」、「ばかを観る」、「バカを観る」、「馬鹿を観る」、「莫迦を観る」の 12 種の異表記が与えられていることになります。

b. 文法機能情報

例えば、「墓穴を掘る」は全体として動詞句(VP)、「命の洗濯」は名詞句(NP)、「年がら年中」は副詞句(AdvP)の働きをする、など、表現全体の相当カテゴリーが記載されています。

c. 文法構造情報

例えば、「目が点になる」は、目→が→なる、点→に→なる という依存構造を持つことを自立語は品詞記号、付属語はローマ字綴りを使い、カッコ [] で [[Nga]][[Nni]V₃₀]] と 2 項の句表示をしています。(N は名詞、V₃₀ は動詞終止形の品詞記号です。) また、並列構造はカッコ<>、あるいは《 》で、並

列要素はカッコ()で表示しています。例えば、「泣く子と地頭」の構造は<([V₄₀N])to(N)>と表わします。(V₄₀は動詞連体形の品詞記号です。)

d. 内部修飾可否情報

例えば、慣用句「油を売る」は「油をいつもの店で売る」のように内部修飾句をとり、ギャップが生じることがあります。このことを c. の構造記述中にアスタリスクで [[Nwo]*V₃₀] のように記載します。この情報は、慣用句などをいつも単語化して扱う弊害を避け、より柔軟なギャップ付きフレーズ(不連続フレーズ)として扱うための枠組みです。

e. 文脈情報

例えば、軽動詞構文「顔をする」は「困った顔をする」の様に文頭側に連体修飾句を要求します。また、副詞的表現「一つたりとも」は後方に「与えない」の様な否定句を要求します。この種の必須(あるいは選好)文脈が記載されています。

f. 連体、連用、動詞化情報

本辞書は「独りよがり」、「針で突いた程」、「子供だまし」のような、物事の様態を表わす形容動詞的と言える表現を含んでおり、これらが連体、連用修飾句として用いられ、動詞化して用いられ、あるいは後続要素を記載しています。例えば、擬態語「フラフラ」は「フラフラの」、「フラフラした」、「フラフラとした」で連体修飾、「フラフラ」、「フラフラと」、「フラフラして」、「フラフラとして」で連用修飾、「フラフラする」、「フラフラとする」と動詞化すること、これに対して「グングン」は、「グングン」、「グングンと」の連用修飾形のみ存在することなどが記載されています。

JDMWE の性質

JDMWE の統計的性質を調べる 2 種の調査を行いました。詳細は文献 8、9 に譲りますが、結果は概略、次の通りです。

(1) ランダムに選んだ日本経済新聞の 1 ページと

最終ページの記事 1,459 文中、1,928 か所に JDMWE の収録表現が何らかの形態(活用変化形など)で使われていました。10 文中に 13 か所程度という比率です。(このうち、6 か所は上記辞書 16. の助詞・助動詞性表現でした。) このように新聞記事における JDMWE の比較的高いトークン・カバー率が見られます。

(2) Google 社が公開した Web 上の 200 億日本語文に現れる n- 単語連鎖の出現頻度データ LDC2009T08 (文献 4) と上記辞書 2. に収録した動詞性(I 類)表現とを比較しました。その結果、『名詞+助詞(が、を、に)』から『動詞』に移る遷移確率が高いものほど JDMWE に多く採録されているという傾向が確認され、表現採録の妥当性がある程度推定できました。また、同調査から JDMWE 収録表現の 10% 程度は Google のデータに表れていないことが分り、コーパスにおける MWE のロングテール分布が覗われました。さらに、Web 上に生起する動詞性表現(I 類)の約 2.5% にすぎない JDMWE 収録の動詞性表現(I 類)が、トークンレベルでは Web 上に生起する同形式表現の約 14% をカバーしていることが分かりました。

このように JDMWE には比較的高頻度生起の表現が選ばれていることが分ります。

JDMWE の応用

JDMWE が与える意味的な纏まりをもつ表現ブロックを処理単位とする手法は構文・意味解析、機械翻訳において威力を発揮します。例えば、JDMWE には「手に付かず」、「散歩に出る」、「事にする」がそれぞれ、副詞性表現(AdvP)、動詞性表現(VP)、助動詞性表現(Aux)として収録され、それぞれに前記 c. の構造 [[Nni][V₁₂zu]]、[[Nni]V₃₀]、[[Nni]suru] が記載されています。また、「手に付かず」には前記 e. の文頭側必須文脈条件として「が」格の後置詞句が指定されています。これらの情報によって入力

文:「彼は仕事が手に付かず、散歩に出る事にした」の構文解析を行った場合の解析例を次ページの図1に示します。(処理の細部は省きます。) 入力文は8文節からなっていますが、JDMWEを優先的に採用すれば、5文節からなる文として取り扱うことができ、処理が簡素化されています。同時に、相当数の不正解(曖昧さ)が排除できていることは明らかです。また、この例の場合、JDMWEの各表現に英訳情報、例えば、“as *SUB* is unable to get down to doing *SUB*’s *N*”、“go out for a walk”、“decide to”などを与えたと仮定すれば、図2のように上記の解析にほぼ並行した形で自然な(意味に忠実な)訳出を行うことが出来る可能性が在ります。(処理の細部は省きます。)

また、前述の統計的性質から、JDMWEは入力単語を予測しながら読み進み、解析木を効率よく生成していく予測的構文解析を実現するためにも有効な資源となります。むしろ、誤ってMWEを認定する可能性もあるでしょうが、巧く設計すれば、妥当な解析に早期に導かれる確率が高いと思われます。

収録表現数 68,000 程度であった JDMWE の旧バージョンを当時市販されていたワープロソフトに組み込み、見出し、および、前記 a. の単語間区切り情報と異表記情報を用いることにより、カナ漢字変換の初回正変換率を 7 ポイント程度向上させたことが報告されています。(文献 3) 現状の JDMWE では収録表現、記載情報がより充実しているため、カナ漢字変換のさらなる精度向上にも寄与できるだろうと思われます

おわりに

2011 年の全米計算言語学会(ACL)年次大会における MWE ワークショップで K. Church 氏は “How Many Multiword Expressions do People Know?”と題する招待講演を行ないました。(文献 1) JDMWE はこの疑問に対する日本語における一つの暫定的

な解答案となっています。(ただし、大半の専門用語、固有表現、時間空間表現は JDMWE の対象外です。)

なお、JDMWE の公開情報については筆者のサイト <http://jefi.info/> をご覧ください。

参考文献

- (1) Church, K., 2011. How Many Multiword Expressions do People Know? Proceedings of the 49th Annual Meeting of the ACL, workshop on MWE,
- (2) Koehn, P. et al., 2007. Moses: Open Source Toolkit for Statistical Machine Translation, Proceedings of the 45th Annual Meeting of the ACL.
- (3) 小山、檜原、吉村、首藤, 1998. 連語データを利用したカナ漢字変換, 情報処理学会論文誌, 39-11
- (4) Kudo, T., Kazawa, H., 2009. Japanese Web N-gram Version 1. Linguistic Data Consortium.
- (5) Galley, M. et al., 2010. Accurate Non-Hierarchical Phrase-Based Translation, NAACL.
- (6) Och, F. J. et al. 2000. Improved Statistical Alignment Models, Proceedings of the 38th Annual Meeting of the ACL.
- (7) Sag, I. A., et al. D., 2002. Multiword Expressions: A Pain in the Neck for NLP. Proceedings of the 3rd International Conference on Intelligent Text Processing and Computational Linguistics (CICLing 2002).
- (8) Shudo, K., Kurahone, A., Tanabe, T., 2011. A Comprehensive Dictionary of Multiword Expressions. Proceedings of the 49th Annual Meeting of the ACL.
- (9) 首藤, 田辺, 2010. 日本語複単語表現辞書 JDMWE, 自然言語処理, 17-5.

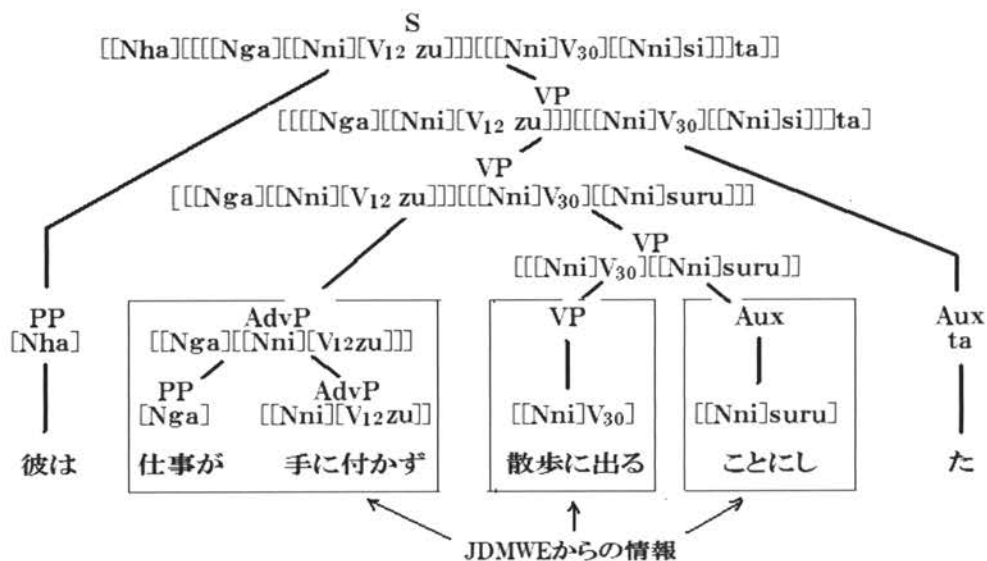


図1 JDMWEによる構文解析例

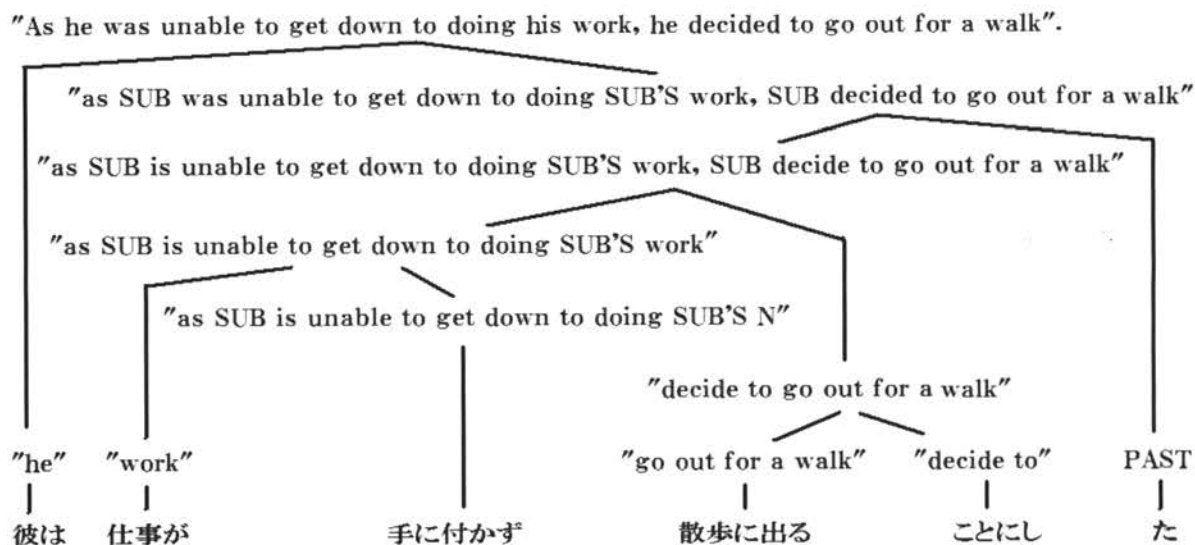


図2 JDMWEによる英訳例

第 14 回機械翻訳サミット(Machine Translation Summit XIV) および
第 5 回特許翻訳ワークショップ(The 5th Workshop on Patent Translation)参加報告

横山晶一(山形大学)、梶博行(静岡大学)、二宮崇(愛媛大学)

1. 会議概要

第 14 回機械翻訳サミット (Machine Translation Summit XIV、以下 MT Summit XIV と略称) は、2013 年 9 月 2~6 日、フランス、ニース(Nice)の Acropolis Conference Centre で開催された。会場は、ニースの中心街や旧市街に近い静かな場所で、アクセスも便利な場所にある。

本会議は、4 日~6 日の 3 日間で、参加者は主催者側発表で 229 名であった。本会議に先立って 2 日と 3 日に 6 つの Tutorial と 4 つの Workshop が行われた。Tutorial の参加者は計 71 名、Workshop の参加者は計 135 名であった。なお、3 日に歓迎レセプション、5 日にパンケットが開かれた。

AAMT/Japio 特許翻訳研究会関係では、筆者ら 3 名の他に、宮澤信一郎 (秀明大学)、宇津呂武仁 (筑波大学) の各委員、Japio から、松田成正部長、星山直人主任が参加した。

本会議では、2 つの招待講演、4 つのユーザ講演と 1 つのパネル討論が行われた。研究論文は 73 編の応募のうち 37 編 (約 50%) が採択され、ユーザ論文は、24 編の応募のうち 17 編 (約 71%) が採択された。これらの論文は 3 並列のオーラルセッション (研究 2、ユーザ 1) とポスターセッション (研究のみ) で発表された。また、FP7 プロジェクトのポスターセッションで 20 あまりの EU プロジェクトの発表があった。

以下に筆者らがオーガナイズした特許翻訳ワークショップと本会議の概要を報告する。

2. 第 5 回特許翻訳ワークショップ

このワークショップは、AAMT/Japio 特許翻訳

研究会 (委員長: 辻井潤一) における種々の議論をもとに、MT Summit でワークショップを開いて、この分野の研究者の意見交換ができればよいという趣旨の下で開催された。過去には第 1 回が Phuket (Chair: 横山晶一)、第 2 回が Copenhagen (Co-Chair: 辻井潤一、横山晶一)、第 3 回が Ottawa (Co-Chair: 江原暉将、横山晶一) 第 4 回が厦門 (Co-Chair: 横山晶一、江原暉将、王丹(CIPO)) で開催されている。

今回は横山晶一、梶博行が Co-Chair として開催した。本会議の 2 日前で、併設の Tutorial や Workshop があったため、参加者は 20 名ほどであった。

以下、プログラムの発表順に、概要を簡単に述べる。以下の時間は、そのセッショントータルの時間を示す。間で空いている時間は、休憩や昼食に当てられている。

(1) 招待講演(9:00-10:40) (司会: 梶博行)

Chair の横山晶一による開会のあいさつの後、2 つの招待講演が行われた。

・ Paul Schwander (European Patent Office):
Machine Translation at the EPO – Removing Language Barriers from Patent Documentation

EPO の機械翻訳への戦略的な取り組みについて EPO の情報管理国際プログラムの責任者であるポール・シュワンダー氏が講演した。内容は以下のとおり。特許情報は膨大で機械翻訳が必須。審査官は機械翻訳の結果をみて人手翻訳すべき特許を選択している。中国特許庁の出願の伸びが著しく、日本語や韓国語を含むアジア言語の特許情報へのアクセスが EPO でも大きな課題となっている。Google と協力して特許情報の機械翻訳サービス Patent Translate を 2012 年 2 月に開始した。現在、日本語、

中国語を含む 21 の言語と英語の間の翻訳がサポートされ、3 億 500 万件の翻訳特許文書にアクセスすることができる。訳文の質を 5 段階で評価するとともに質の向上をはかっている。2014 年末までに 32 言語に拡大する計画で、その中には韓国語やロシア語が含まれる。

・ Hitoshi HONDA (Japan Patent Office): Current Status of Machine Translation System in the JPO

日本の特許庁における機械翻訳システムの利用について特許庁・特許情報企画室・特許情報企画調査係長の本田仁氏が講演した。内容は以下のとおり。審査官の業務の 1 ステップである先行技術のサーチにおいて外国語文書へのアクセスを容易にする手段として機械翻訳が果たす役割は大きい。現在、米国や欧州の特許については日本語抄録（人手翻訳）と英語の全文テキスト（原文）を利用することができるが、韓国や中国の特許については英文抄録（人手翻訳）しか利用できない。2014 年度中には韓国、中国特許の機械翻訳による日本語全文テキストを提供する予定である。外国の特許庁との間の審査情報の相互利用にも力を入れており、日本特許の審査情報の日英機械翻訳結果を EPO や USPTO（米国特許庁）に提供している。EPO や USPTO のユーザ（審査官）からはフィードバックがあり、システムの改良に役立っている。また、中国特許庁の出願件数の著しい増加に鑑みて、中日機械翻訳システムの構築に力を入れている。2012 年には、中国と日本の対応する特許のテキストから対訳関係を抽出し、100 万語規模の日中専門用語対訳辞書を構築した。

(2) 一般講演 1(11:10-12:10) (司会：二宮崇)

ここでは 3 件の発表が行われた。

・ Svetlana Sheremetyeva: On Integrating Hybrid And Rule-Based Components For Patent MT With Several Levels Of Output

Sheremetyeva による発表。この論文の提案手法では、意味役割付与まで行う構文解析を用い、原言

語の意味表現から翻訳先言語の意味表現に変換し、翻訳先言語の文を生成することで機械翻訳を行っている。システムは、NP チャンカーや VP チャンカーを用いる浅いモジュールと意味役割まで扱う深いモジュールから構成され、深いモジュールでは、意味役割のついた述語項構造が用いられている。また、解析結果を修正するインタラクティブモードや、単文の解析を出力するモードを備えている。まだ大規模な実験は行われていないが、より多くの知識やルールを記述することで実用的な機械翻訳システムの実現を目指している。

・ Itsuki Toyota, Zi Long, Lijuan Dong, Takehito Utsuro, and Mikio Yamamoto: Compositional Translation of Technical Terms by Integrating Patent Families as a Parallel Corpus and a Comparable Corpus

Toyota による発表。パラレルコーパスとコンパラブルコーパスを用いた構成的な対訳辞書生成手法の提案をしている。特許翻訳の学習および評価でよく用いられているパラレルコーパスは、コンパラブルコーパスから対訳文対を抽出することにより生成されており、そのため元のコンパラブルコーパスのうち 30%程度しか用いられていない。従って、パラレルコーパスからの対訳辞書抽出では残りの 70%が有効に利用されていない。この論文の提案手法では、コンパラブルコーパスから対訳専門用語の候補を抽出し、英辞郎およびパラレルコーパスから学習した翻訳テーブルを用い、構成的に対訳候補のスコアを計算し、対訳専門用語を決定する。この手法によりコンパラブルコーパスの残り 70%から対訳専門用語対が得られる。この手法により比較的高い精度でより多くの対訳語対が得られることが実験により示された。

・ Shoichi Yokoyama: Analysis of Parallel Structures in Patent Sentences, Focusing on the Head Words

Yokoyama による発表。特許データベースから、並列構造を形成していると思われる主辞(head

word)を抽出し、頻度の高い語を用いて、並列構造の誤りを修正した結果、58%程度の誤りを訂正できたという結果を述べた。

(3) 基調講演(14:30-15:30) (司会：横山晶一)

・ Philipp Koehn (University of Edinburgh, Great Britain):Advances in Machine Translation for Patent Translation and Beyond

午前中に自らの Tutorial をこなし、わざわざこの基調講演を引き受けてくれたことをまず感謝したい。内容は、本会議の招待講演や、この会議を通じて常に話題になっていたトピックである Domain Adaptation や異なるモデルの combination、最近の Operation Sequence Model などを使いやすく説明し、最後に big data に対する挑戦について述べるといった、近年の MT の進展を解説するものであった。

(4) 一般講演 2(16:00-16:40) (司会：Svetlana Sheremetyeva)

2 件の発表があった。

・ Yun Jin and Oh-Woog Kwon, Seung-Hoon Na, and Young-Gil Kim: Patent Translations as Technical Document Translation: Customizing a Chinese-Korean MT System to Patent Domain

韓国のグループによる発表。特許文書と技術文書の類似性に着目して、Domain Adaptation に技術文書をもちいることにより、中韓 MT システムの性能向上を図った結果について述べた。

・ Rahma Sellami and Fatiha Sadat, and Lamia Hadrich Belguith: Exploiting Multiple Resources for Japanese to English Patent Translation

チュニジアとカナダのグループによる発表。NTCIR-10 の日英特許翻訳システムの翻訳性能を改善するために Domain Adaptation や resource の拡張を行うのに、日英で同じと思われる Wikipedia を使用したという内容であった。

(5) まとめ

横山による締めくくりの挨拶。次回も行いたい

という趣旨のことを述べた。

3. 本会議

3.1. 招待講演、ユーザプレゼンテーション

招待講演は、4 日 (今回の主催者である Andy Way の開会挨拶に続いて)、5 日にそれぞれ 1 つずつ行われた。

・ Hinrich Schuetze: The operation sequence model: integrating translation and reordering operations in a single left-to-right model (4 日)

ミュンヘン大学教授による講演。これまでの PB-SMT など、種々のモデルの利点と欠点を分かりやすく解説するとともに、それらを統合、克服するモデルとしての operation sequence model について解説し、その現状と将来展望について述べた。

・ Harold Elsen: Demystifying Machine Translation: Learning from the real world (5 日)

ドイツの DELTA International CITS GmbH の managing director による講演。SMT と RMT とのさまざまな比較や、MT と TM の結合、術語、post-edit など、リアルワールドにおけるさまざまな問題点をあげて、MT の将来像はこうあるべきだという持論を述べた。

ユーザプレゼンテーションは、Hans Uszkoreit による "Including the User in Research and Application Improvement"(ただし DCU のメンバーによる代読)のほか、LionBridge, IBM Germany, Microsoft の 3 社が発表した。

3.2. 一般講演、ポスターセッション

一般講演は、論文採録のところでも述べたように、Research and Development 2 つと User Track 1 つに分かれて、3 つのセッションが並行して進められるという形式で行われた。

話題は、今回は比較的狭い範囲に絞られたような感想を持っている。

まず一つは Domain Adaptation である。これは、たとえば初日の招待講演の直後に行われた Patrick Lehnert et al.: Conditional Random Field Using

Intermediate Classes for Statistical Machine Translation や、George Foster et al.: Simulating Discriminative Training for Linear Mixture Adaptation in Statistical Machine Translation (ベストペーパー賞を獲得)、Katharina Waeschle et al.: Generative Discriminative Methods for Online Adaptation in SMT などがそれにあたる。

Terminology も主要なテーマである。岡山大の竹内孔一氏らのグループの発表、Koichi Satoh et al.: Terminology-driven Augmentation of Bilingual Terminologies などである。

また、pre-editing, post-editing も話題の一つで、Guillaume Wisniewski et al.: Analysis of a Large Corpus of Post-Edited Translations: Quality Estimation, Failure Analysis and the Variability of Post-Editon や、Michel Simard et al.: PEPr: Post-Edit Propagation Using Phrase-based Statistical Machine Translation、ポスターで発表された Lucia Morado Vazquez et al.: Comparing Forum Data Post-Editing Performance Using Translation Memory and Machine Translation Output: A Pilot Study などがある。

Estimation もテーマの一つで、Kahif Shah et al.: An Investigation on the Effectiveness of Features for Translation Quality Estimation, Sudip Kumar Naskar et al.: Meta- Evaluation of a Diagnostic Quality Metric for Machine Translation, Ahmed Hamdi et al.: The Effect of Factorizing Root and Pattern Mapping in Bidirectional Tunisian- Standard Arabic Machine Translation, Luis Formigo et al.: Real-Life Translation Quality Estimation for MT System Selection, Pratyush Banerjee et al.: Quality Estimation-guided Data Selection for Domain Adaptation of SMT, Maaxi Li et al.: Listwise Approach to Learning to Rank for Automatic Evaluation of Machine Translation などがある。

3.3. FP7 プロジェクトセッション

EU の The 7th Framework Programme for Research and Technological Development (第7次研究・技術開

発フレームワークプログラム, 2007~2013 年) の支援を受けたプロジェクトのセッションで、以下のプロジェクトのポスター発表があった。ヨーロッパでは機械翻訳技術の実用化、応用がさまざまな形で進められていることがわかる。

- ACCEPT (ジュネーブ大学ほか): インターネットコミュニティコンテンツの機械翻訳
- BOLOGNA (CrossLang ほか): 統計的機械翻訳とポストエディット、翻訳メモリを統合したシラバス翻訳サービス
- CASMACAT (エディンバラ大学ほか): 翻訳プロセスの認知モデルに基づく次世代翻訳ワークベンチ
- CNGL (ダブリンシティ大学): 多言語・マルチモーダル・マルチメディアコンテンツの生成、処理、統合
- EUBRIDGE (カールスルーエ工科大学ほか): テレビのキャプション、大学の講義、EU 議会などの翻訳
- EXCITEMENT (Bar Ilan 大学ほか): 多言語テキスト推論プラットフォーム
- EXPERT (Pangeanic ほか): ハイブリッド言語翻訳技術
- FAUST (ケンブリッジ大学ほか): リアルタイムでユーザーフィードバックに適応する機械翻訳システム
- Let'sMT (コペンハーゲン大学ほか): 個人に合わせてカスタマイズされた翻訳システムの構築ツール
- MateCat (FBK ほか): 機械翻訳のコンピュータ支援翻訳への統合
- METANET (DFKI ほか): 翻訳品質評価 QTLanchPad、Web コンテンツのためのメタデータ MultilingualWeb-LT など
- MOLTO (イエーテボリ大学ほか): 制限言語に基づく多言語オンライン翻訳
- MONNET (DELI ほか): オントロジーの翻訳

- MosesCore (エディンバラ大学ほか) : オープンソース機械翻訳の開発と利用
- PANACEA (Pompeu Fabra 大学ほか) : 言語資源工場 (機械翻訳のための言語資源の獲得、生成、メンテナンス)
- QT (ダブリンシティ大学ほか) : 翻訳品質評価技術
- SIGNSPEAK (アーヘン工科大学ほか) : 手話認識・翻訳
- SUMAT (Deluxe Digital Studios ほか) : AVコンテンツのサブタイトルの機械翻訳
- TaaS (Tilde) : 多言語ターミノロジーの獲得、処理、再利用
- transLectures (パレンシア技術大学ほか) : ビデオ講義の転写・翻訳
- TTC (ナント大学) : コンパラブルコーパスからのターム抽出
- XLIKE (ザグレブ大学ほか) : クロスリンガル知識抽出

3.4. パネルディスカッション

パネルは、The MT Tower of Babel という題で、司会者と 6 名のパネリストが壇上に上がった。MT システムはすでに 50 年以上も前から作られているのに、Google や Bing が出てくるまではほとんど一般の関心をひかなかったのはどういうわけかというやや刺激的な問いかけで、研究者や翻訳会社からのパネラーが討論した。ただ、この種のパネルの常として、時間が足りなかったのが残念であった。

4. まとめと考察

本会議も、ワークショップも、現在の機械翻訳のかかえる問題点が明確になり、大変有意義であった。機械翻訳関係の研究者、ユーザ、メーカー等が世界レベルで一同に会する機会はこの会議以外には非常に少ない。そのためにも今後ともこの会議の継続と発展を祈っている。今回の会議では、free の

翻訳者も結構参加していたのは、日本では余り見られない現象であった。

なお、会議の最後に IAMT Award of Honor が、SMT の貢献などに対する業績で、Hermann Ney 教授に授与された。また、この分野で長年活躍してきた John Hutchins 氏に特別表彰が行われ、会場は standing ovation で、氏をたたえた。次回の会議は、EAMT 会長の Andy Way 氏から、AMTA 顧問 (前会長) の Alon Lavie 氏に IAMT 会長が引き継がれ、アメリカ Miami で 2015 年 10 月下旬に ATA National Conference と連動して開催の予定である。なお、AAMT の中岩浩巳会長が IAMT 副会長に就任した。また、EAMT-2014 の会議は、クロアチアの世界遺産都市 Dubrovnik で、2014 年 6 月 16~18 日に行われることも合わせて発表された。

今回のアレンジや論文査読などに委員会の多くのメンバーの助力があったことに感謝する。第 6 回のワークショップも、今後の議論の中で詰めてゆきたい。

「第 23 回 JTF 翻訳祭」のお知らせ

一般社団法人日本翻訳連盟

翻訳祭企画実行委員会

<「第 23 回 JTF 翻訳祭」開催のご案内>

日本翻訳連盟では、来る 11 月 27 日（水）、東京のアルカディア市ヶ谷（私学会館）を会場として、恒例の「JTF 翻訳祭」を開催いたします。今年で 23 回目を迎えます。AAMT 様はじめ各関連団体より後援をいただいております。

今年掲げるテーマは、「大翻訳時代 ～Define Your Blue Ocean Strategy～」です。分科会形式による「講演・パネルディスカッション」及び「プレゼン・製品説明コーナー」（全 30 セッション）、「翻訳プラザ（展示会）」、「交流パーティー」などの催し物を用意しております。交流パーティーでは参加者同士の交流促進や情報交換の場としてご活用いただければ幸いです。

翻訳祭は翻訳者、翻訳会社、クライアント、翻訳支援ツールメーカー、翻訳スクール等の業界関係者が一堂に会するイベントです。翻訳祭を通して、参加者の皆様に新たな人脈形成やビジネス機会が生まれることを願っております。

翻訳祭に是非ともご出席賜りますよう、よろしくお願い申し上げます。

★★★ 第 23 回 JTF 翻訳祭 ★★★

http://www.jtf.jp/jp/festival/festival_top.html

【日時】

2013 年 11 月 27 日（水）9:30～20:30（開場・展示会開始 9:00）

【場所】

「アルカディア市ヶ谷（私学会館）」

※JR・地下鉄「市ヶ谷駅」より徒歩 2 分

【全体テーマ】

「大翻訳時代 ～Define Your Blue Ocean Strategy～」

【概要】

この数年間、製造業の不調や際限のない円高など先の見えない閉塞感と向き合ってきた翻訳産業界。でもちょっと背伸びして見渡してみれば、とても近い距離にある新興国と手を携えて成長に転じていける可能性があったのです。B to B も B to C も、ビジネスはまず『ことばありき』。リーマンショックから 5 年、そして東日本大震災から 3 年目の今年を、『大翻訳時代』の到来年と位置付けました。サブテーマは、欧州経営大学院の W・チャン・キム教授らの経営論「ブルー・オーシャン戦略」。競争の激しい既存市場「レッド・オーシャン（赤い海、血で血を洗う競争の激しい領域）」を離れ、参入者のいない未開拓市場「ブルー・オーシャン（青い海、競合相手のいない領域）」を切り拓くべきだと説いています。

これらのテーマのもと、ゲームや映像などのエンターテイメントから、特許、IT などの技術分野、そしてメディカルや文学までの多彩なジャンルを対象に、翻訳産業の起爆剤となり得る可能性を秘める機械翻訳の最新動向や東南アジアの翻訳市場動向も交え、多様な方向から Blue Ocean Strategy を探ります。どんな革新的なチャレンジができるのか、一緒に考えていただけたら幸いです。

今年の翻訳祭では、AAMT 様にパネルディスカッションを企画していただきました。

11:30～13:00（90 分）

「機械翻訳による競争戦略

—機械翻訳で差別化できるか—」

<パネリスト>

○秋元圭氏 (㈱クロスランゲージアルアンドディ)

○目次由美子氏 (㈱シュタール ジャパン)

○山本ゆうじ氏 (秋桜舎・代表)

<モデレーター>

長瀬友樹氏 (㈱富士通研究所)

【申込締切】

2013 年 11 月 20 日 (水) まで

【主催・運営】

一般社団法人 日本翻訳連盟

第 23 回 JTF 翻訳祭企画実行委員会

【後援・特別協力】

日本翻訳者協会 (JAT)

【後援】

経済産業省・アジア太平洋機械翻訳協会(AAMT) 他

【プログラム】

・開場/受付開始 9:00~

・トラック 1~6 [6 会場]9:30~18:00

全 24 セッション

・プレゼン・製品説明コーナー

9:30~17:30 全 5 セッション

・翻訳プラザ (展示会)

9:00~17:30

・交流パーティー [会場]3 階 富士の間

18:30~20:30 (120 分)

→詳細はウェブサイトをご覧ください。

●翻訳プラザ (展示会) 出展 35 社

昨年に続き翻訳プラザは、例年よりも広い会場で開催します。より多くの方にお立ち寄りいただけるよう、講演間やお昼の休憩時間も長く設定します。翻訳プラザは、翻訳業界の企業が参加する展示会です。来場者と出展企業が交流し、情報交換する場で

す。AAMT 様も出展されます。どうぞ 3 階の展示会会場にお立ち寄りください。

<展示内容>

・翻訳支援ツールメーカーの製品デモ

・翻訳会社との翻訳相談 (翻訳発注・翻訳者登録等)

・翻訳学校の翻訳講座紹介

・出版社の翻訳書籍販売

●プレゼン・製品説明コーナー 参加 5 社

今年も展示会に隣接して「プレゼン・製品説明コーナー」を併設します。プロジェクター・スクリーンを備えた特設会場となっています。

・1 セッション 60 分/5 セッション

・翻訳プラザに隣接した特設会場にて、各社のプレゼン発表を受講!

・プレゼンコーナーを増設し、65 名収容可能!

※翻訳プラザ、プレゼン・製品説明コーナーへの参加は無料です (事前登録不要)。

プレゼン発表企業一覧 (5 社)

●交流パーティー

情報交換を行い、これからの業界の展望などを語り合う機会となります。翻訳祭の講演者やパネリストの方々も参加される予定です。また昨年に続いて「ほんやく検定 1 級合格者の表彰式」を行います。お仕事のための人脈作りや最新の業界の動向をつかむためにぜひご参加ください。

<ボランティアスタッフ募集中! (11/7 まで)>

募集内容は、「受付」「会場司会」「会場アシスタント」「書記」「写真」です。ボランティアスタッフにご協力いただく方の「講演・パネルディスカッション」参加料は無料です (パーティーは有料)。

【お問い合わせ】

一般社団法人 日本翻訳連盟 事務局

<http://www.jtf.jp/> E-mail:info@jtf.jp

言葉の壁を越えたコミュニケーションの実現にむけて

株式会社NTTドコモ サービス&ソリューション開発部

那須 和徳

1. はじめに

最近10年間で日本の移動通信市場を見てみると、NTTドコモが1999年に開始したi-modeサービスに代表されるモバイル（本稿では「モバイル」を第三世代以降の携帯電話を中心に以降記述）インターネット市場の発展と成熟の10年間であったと言い換えることも出来る。モバイルインターネット上では従来の電話や電子メールといったコミュニケーションサービスとは異なる新たなコミュニケーションツールとして、LINEやTwitter、Facebook、といったover the topプレイヤー（以下、OTT）によるクラウドサービスが展開されている。

このような事業環境の中、NTTドコモでは2012年に、「ドコモクラウド」というブランド名を付与し、当該ブランドサービスの一つとして自動音声翻訳サービス「はなして翻訳」を開始した。本稿では、「はなして翻訳」のサービス状況からみた、ユーザーが翻訳機能に求めるものは何か？についての考察を述べる。

2. 「はなして翻訳」サービス概要

はなして翻訳は、スマートフォンを通じて、ユーザーの言葉を相手の言葉に翻訳する。日本人と外国人のコミュニケーションにおいて、その場にあたかも通訳者がいて、コミュニケーションをサポートする、そんな風景を再現したサービスである。対応言語は、英語、中国語、韓国語、ドイツ語、フランス語、イタリア語、スペイン語、ポルトガル語、タイ語、インドネシア語の10ヶ国語に対応している。

サービスの特徴は次の三つだと考える。一つ目は高いレスポンス性の実現、二つ目は安心機能の充実（翻訳結果を再度母国語に翻訳し正誤を確認）である。これら二つは、本サービスを翻訳辞書ではなく、

コミュニケーション利用のために使うことを想定した工夫箇所である。三つ目は、離れた場所にいる二人の電話中の会話を翻訳する遠隔翻訳である。肉声（母国語）と合成音（外国語）と文字（母国語・外国語）でコミュニケーションをとることができるサービスである。

3. 「はなして翻訳」の仕組み

はなして翻訳は幾つかの機能を組み合わせることで実現している。発話された音声は、音声認識処理を実施し音声から文字へと変換する。その後、機械翻訳処理を実施し母国語の文字から外国語の文字へ変換する。そして音声合成処理にて外国語の文字から合成音を作成し音声挿入する。“音声録音”“音声認識”“機械翻訳”“音声合成”“音声挿入”という5つの機能を組み合わせて実現している。

これらを実現するうえでサービスシナリオと呼ぶ機能が重要な役割を果たしている。サービスシナリオは二つの機能を有している。一点目は、機能部品を組み合わせることでサービスを柔軟に組み立てる機能。複数の機能を組み合わせや順序制御を簡単な記述で変更でき、サービスバリエーションの拡大が短期間で可能である。二点目は、クライアント／サーバ型通信（以降、C/S）とピア to ピア型通信（以降、P2P）の2種類の通信処理を融合する機能である。現在、多くのWEBサービスに利用されている通信はC/S型通信である。端末からのリクエストに対し、サーバで処理を実行し、実行結果の応答を端末に返すという通信である。これに対し、ネットワークを介して端末と端末がデータ送受信する電話のような通信がP2P型通信である。これらの異なる2つの通信処理を融合することで、今までにはないサービス実現が可能である。

このような特徴的な仕組みにより実現した高いレスポンス性能や翻訳精度などを評価頂き、はなして翻訳は、世界の著名な展示会にて表彰いただいた。（「CTEC JAPAN2012 米国メディアパネルイノベーションアワード 2012 グランプリ」、「MobileWorldCongress2013 グローバルモバイル賞 消費者向けベストネットワーク商品 ソリューション部門賞」、「CTIA2013_E-Tech 賞モバイルアプリ・コンテンツ、ソーシャル、メディア&エンターテインメント部門賞」を受賞）

4. ユーザーの声に基づく考察

はなして翻訳は、約一年間のトライアルを経て2012年11月よりサービスを開始した。2013年8月末時点で、アプリインストール数230万、翻訳回数一日平均20万回程度の利用状況である。これまでの期間SNSなどを通して多くのユーザーの声を収集した。トライアル期間においては、法人利用として約80ヶ所の店舗など外国人観光客の接客で、個人利用は400名のモニターに利用頂いた。ここでは、開発者が直接ユーザーに会うことで実感を伴う生の声（改善要望など）を収集した。

モニター期間の本サービスへの改善箇所についてアンケート集計を行った。結果、断トツ一位は、“操作性”であった。アンケート総数のうち操作性向上の要望件数が45%を占め、ついで、音声認識向上が10%、翻訳精度向上が10%であった。以外に思われる方も多いのではないだろうか。また、モニターユーザーから直接、「今より、少しでも意味が伝わるようになれば良い。完璧なものを求める必要はない。」とアドバイス頂いたことも数件あり、精度について寛容であることが分かった。（もちろん、一定の精度は必要であるが。）

精度よりも操作性が重視されるという結果であるが、これら二つに因果関係があるとも考えている。開発段階においてのエピソードである。音声認識エンジンの評価の一つに、認識精度は同等であるがレスポンス性の異なる二パターンで評価した際、文字正答率などの定量評価は変わらないが、ユーザー感覚としてはレスポンスが早いエンジンが高精度と感じるのである。待った挙句に間違えたので

は、待たされたことへの不満と間違ったことへの不満がダブルで効いてくるのである。つまり操作性の向上はユーザー感覚としては精度向上にもつながるということである。

また、商用サービス開始後の利用状況に、ひとつ面白い結果が現れた。スマートフォンに表示された音声認識結果や翻訳結果が誤っている場合、母国語の入力文を正しく修正する機能があり、この機能により母国語の表現方法を変更することで正しい翻訳に改善される場合も少なくない。この機能の利用率（修正回数／翻訳回数）が、なんと20%となっているのである。事前の想定では、ユーザーは誤った翻訳結果が出ると、「このサービスは使えない。」と早々に判断し、その後の利用を止めるのだと考えていた。結果からみると想定とは異なり、このサービスを完全ではないとの理解のうえで、使い方の工夫しているユーザーが多いという結果であった。「今より、少しでも意味が伝わるようになれば良い。」を裏付ける結果である。

これらの結果をまとめると、ユーザーは、翻訳精度が不十分なことは理解のうえ補完ツールとして翻訳サービスを利用したい意向を持っている。ただし、操作性の要求レベルは厳しく、少しでも使いづらかったり、レスポンスが遅く待たされたりすることへの許容度は低いと考える。

5. おわりに

プロジェクトを開始した段階では自然言語処理技術者からは、「実用レベルのものは、まだ無理。」といった意見が多かった。故に、大規模なユーザーに触れることは、コーパス獲得による精度向上が期待される側面と、悪評価となるリスクから臆病になる側面、そんな両面が感じられた。はなして翻訳の開発者は通信技術者であり、機械翻訳や音声認識などの自然言語処理技術の素人のため、ある種、怖いもの知らずの状況が幸いし、積極的に市場に出せたことが複数の表彰を受賞できた要因であると思う。

「理想を追い求めるだけではなく、今できる範囲で最高のものをユーザーに提供する。」この言葉は、

弊社の幹部より本プロジェクト開始においての開
発者への叱咤激励の言葉である。意図は、「怖がら
ずに早い段階で世の中の評価を受けることが重要
である。」ということであり、引き続き実践してい
きたいと考える。

先日、2020 年オリンピック・パラリンピックの
開催地が東京に決定した。7 年後には、世界各国か
ら日本に集まる外国人とのコミュニケーションを
不自由なくとれる世界の実現を目指して、機械翻訳
に携わる多くのプレイヤーの皆様と協業してい
きたいと考える。今後も、言葉の壁を越えたコミュニ
ケーションの実現に向けて、多くの力を結集するた
めの役割を果たしていく所存である。

ソフトウェア ローカライゼーションにおける MT 利用と業界間断層

CA Technologies、Globalization Translation

古田京子（マネージャ）、菊池邦明（シニア トランスレータ）

第 1 部：今後期待される翻訳ツール像

1. 弊社における MT の利用

私達の所属する部署では、弊社ソフトウェア製品のユーザ インターフェイス文字列およびドキュメントの翻訳（ソフトウェア ローカライゼーション）を担当している。他社と同様に弊社でも翻訳メモリ（TM）を活用する翻訳支援（CAT）ツールを使用してきたが、更なる効率化のために 2009 年から機械翻訳（MT）ツールをローカライゼーションに利用し始めた。MT の適用対象は、TM の適用では 100% マッチやファジー マッチにならない、新規の翻訳文字列である。

MT の適用ではプリプロセスとポストプロセスが重要な役割を果たしており、これらの処理は実装方法によって 2 つの種類がある。1 つは弊社のツールにソフトウェア的に組み込まれたプリ/ポストプロセスで、これらはツールのオプションによってそれぞれオン/オフを切り替えることができる。たとえば、弊社のソフトウェア リソースに出現する変数文字列が機械翻訳によって翻訳されないようにする機能などがこれには含まれる。

もう 1 つは翻訳者が任意にカスタマイズすることが可能な文字列置換テーブルによる文字列処理である。たとえば弊社のスタイルでは、「～しなければならない」という表現の使用は避けているが、文字列置換テーブルを使用して「～する必要がある」という表現に置換されるようにしている。文字列置換テーブルでは正規表現が使用できるため、多くのケースに柔軟に対応できる。

これらの工夫や改善を経て、今では MT は弊社のローカライゼーションではなくてはならないツ

ルとなり、翻訳の効率化およびコスト削減に貢献している。

2. 翻訳ツールの現状

CAT ツールおよび MT ツールの両方を使用していると、各ツールの特徴がより明らかになると共に、両者の狭間に何か不足している点がより強く意識されるようになった。現場で翻訳作業を行う者の立場から見ると、現在存在する技術をうまく組み合わせることで、両者の長所の合計以上のことが可能な、ある意味で理想的な翻訳ツールを作ることも不可能ではないように思われた。

CAT ツールおよび MT ツールの比較という観点から各ツールの特徴を述べてみたい。

CAT ツールの特徴

- ◆ TM の活用
 - 同一/類似の文を検索し、再利用する
 - 詳細なレバレッジ解析機能（マッチ率ごとのワード数）
- ◆ 文中に現れるタグ要素（書式、変数）への対応
- ◆ 多くのファイル形式への対応
 - 翻訳先言語で元のファイル形式に戻しても問題が比較的少ない

MT ツールの特徴

- ◆ 新規の翻訳元テキストから翻訳先テキストを生成できる
- ◆ TM 機能またはタグ要素への対応がないか限られている
 - TM のエクスポート/インポートで文字種が変更
 - Unicode の未サポート
 - レバレッジ解析機能の未サポート
 - タグ要素（書式、変数）の破壊や消失。MT 出力

にも悪影響

- ◆ 対応ファイル形式が限られている

- 対応ファイル形式に MT を適用した場合もレイアウトなどの問題が発生しやすい

あたかも MT ツールに短所が多いように見えてしまうが、翻訳元テキストから翻訳先テキストを生成することは CAT ツールにはできないことであり、最大の長所である。近年、CAT ツールおよび MT ツール両方の長所を活かそうとする試みが多く見られるようになってきた。以下に例を示す。

- ◆ CAT ツール内部からの MT 利用

- CAT ツールが対応している MT エンジンのみ利用可能

SDL Trados Studio、Atril Déjà Vu X2、MultiCorpora MultiTrans Prism など

- ◆ CAT ツールの進化 (TM の統計的、EBMT 的处理)

- Advanced Leveraging 技術 (文よりも細かいレベルでの翻訳再利用)

Déjà Vu X2 の DeepMiner、MultiTrans Prism の ALTM など

- ◆ CAT および MT ツールの統合

- CAT/MT ツール統合環境

富士通 Cliché、秋桜舎 SATILA など

3. 期待される翻訳ツール

現状を踏まえ、今後期待される翻訳ツールがどのようなものになるか考察してみた。これらはあくまでも一人の翻訳者の立場から希望するものであり、厳密性に欠け、夢想であるかのような感は否めないが、ご容赦いただければ幸いである。今後期待され

る翻訳ツールで重要になるであろうと思われる機能に以下の 3 点がある。

- ◆ 依存構造木情報を持った TM

- ◆ ポストエディット結果の活用

- ◆ 信頼性スコア

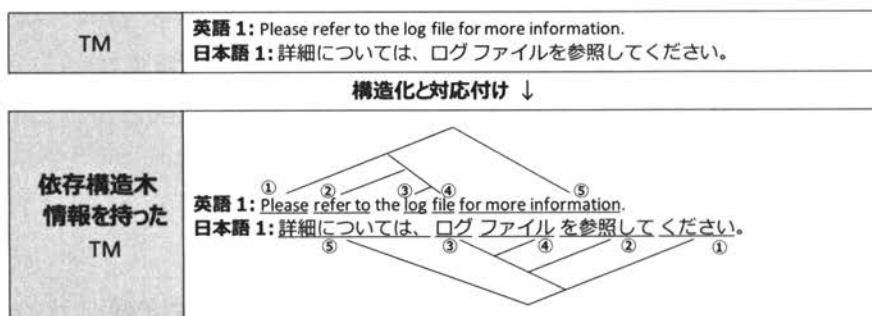
依存構造木情報を持った TM

TM を活用する CAT ツールがソフトウェアローカライゼーション業界で使用されるようになって 15 年以上が経過しているが、TM の基本構造はいまだに翻訳元および翻訳先テキストの対データである。TM に依存構造木情報を付加することによってデータの有用性が高まるように思われる。

依存構造木情報を付加する方法としては、翻訳元および翻訳先テキストの両方をルールベース機械翻訳 (RBMT) 的に辞書と文法情報から解析する方法、統計機械翻訳 (SMT) または用例ベース機械翻訳 (EBMT) が行うような統計的解析を利用する方法、などが考えられる。翻訳元および翻訳先テキストの対データに依存構造木情報を付加した場合の例を示す (図 1)。翻訳元および翻訳先テキストが厳密に単語レベルまでは一致しないケースも多く発生すると思われるが、そのような場合は解析の細かさを下げることで可能な範囲で対応付けが行われることが望まれる。

TM に依存構造木情報を付加することで、TM に含まれている既存訳データをより有効にレバレッジ (再利用) できるようになる。図 2 の例では、翻訳対象英語と TM が比較され、③の箇所のみが異なることが検出されている。次に、この該当箇所のみ MT が適用され、日本語側で差異が補われて翻訳

図 1：依存構造木情報を持った TM



先テキストが完成されている。

図 3 の例では、②～④を含む枝全体が翻訳対象英語で動詞句として再利用できることが検出されている。次に、この動詞句の日本語訳を再利用しながら残りの部分に MT が適用され、翻訳先テキストが完成されている。

ポストエディット結果の活用

MT の出力結果をポストエディットしていると、何度も似たような修正作業を行っていることに気付くことがある。このような場合、翻訳者が行う文

字列の追加/置換/削除/移動の情報がポストエディットを行う編集環境によって自動的に取得され、以降の作業における翻訳先テキストに自動的に反映されることが望まれる。

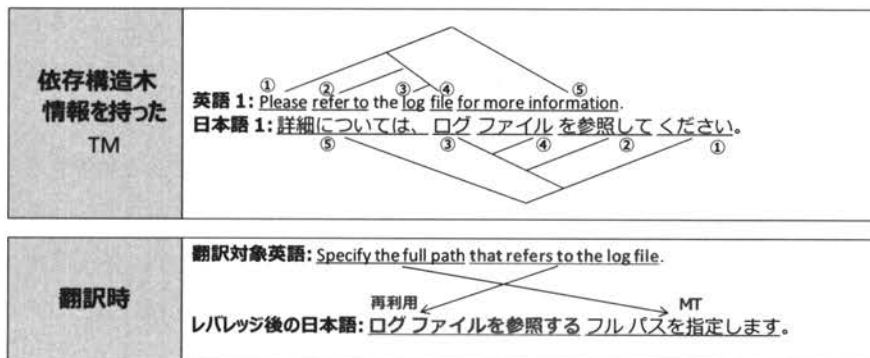
図 4 の例では、「operating system(s)」から翻訳された「オペレーティング システム」が「OS」に書き換えられるべきことが学習され、後に出現した翻訳先テキストで自動的に置換（自動ポストエディット）が行われている。実際の動作としては、ある特定の編集が 1 回行われただけで以降の翻訳

図 2： Advanced Leveraging の例 (差異のみの補完)



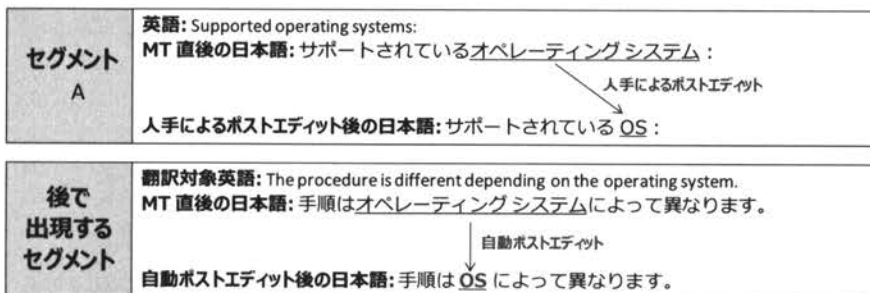
③のみが異なることを検出。該当箇所のみ MT 適用後置換し、残りの部分は再利用。

図 3： Advanced Leveraging の例 (サブセグメントレベルでの再利用)



②～④を含む枝全体が動詞句として再利用できることを検出。残りの部分は MT 適用して組み合わせる。

図 4： ポストエディット結果の活用



「operating system(s)」から翻訳された「オペレーティングシステム」が、「OS」に書き換えられるべきことを学習。

先テキストで自動ポストエディットが行われてしまうと過剰な対応になってしまうと思われるため、複数回（たとえば 3 回など）同じ編集が行われたら自動ポストエディットの対象にするなど、柔軟な調整が必要になるとと思われる。

信頼性スコア

上の「依存構造木情報を持った TM」で述べたように TM および MT の両者を相補的かつ効果的に活用できるようになると、翻訳ツールによって最も適切であろうと判断される形で翻訳先テキストが生成されるようになる。このため、翻訳元テキストの類似度が高い翻訳先テキストを再利用する、CAT ツールで言うところのファジー マッチが存在しなくなる。そのため、従来のファジー マッチの一致率に相当するものとして、信頼性スコアのようなものが必要になってくると考えられる。

TM および MT によって組み立てられた翻訳先テキストをレビューおよびポストエディットする場で、そのテキストにどの程度の信頼性があるのか何の指標もないとすると、翻訳者は結局のところ翻訳元および翻訳先テキストの整合性を一字一句確認する必要に迫られ、レビューに要する時間が増大してしまう可能性がある。翻訳先テキストがどのように生成されたかに基づき、0 ～ 100 などの信頼性を示すスコアが提供されるようになると、翻訳者

はそれを参考に重点的にレビューを遂行できるようになる。また、CAT ツールではファジー マッチの翻訳元テキストの差異が確認できるのと同様に、どの部分が TM からの再利用部分で、どの部分が MT の適用部分であるのか確認できるようになるのが望ましい。

信頼性スコアのその他の用途としては、たとえばスコアが 95 以上である場合にはレビューなしでそのまま利用するなどの自動的な処理にも応用することができる。

信頼性スコアの算出に考慮できる要素としては、たとえば以下のものが考えられる。

- TM 由来の箇所と MT 適用箇所の比率
- TM 由来の箇所に関し、TM 内の該当数と一貫性
- テキストのワード数や文字数
- RBMT 的、統計的に翻訳元および翻訳先テキストを解析し比較

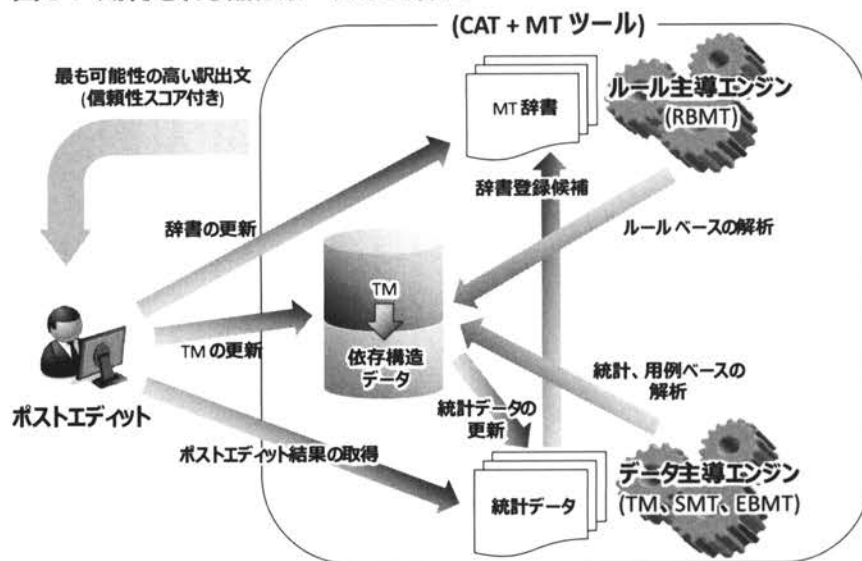
4. 期待される翻訳ツールの懸念事項

今後期待される翻訳ツールで懸念されると思われる主な事項を以下に列挙してみる。

ユーザ インターフェース

現状の CAT ツールでも、すでに翻訳元および翻訳先テキスト、TM の参照、用語集の参照などで画面上に多くの情報が満載されている。これらに加

図 5：期待される翻訳ツールの動作図



え、翻訳先テキストがどのように生成されたかを表示する機能や、RBMT 辞書を編集する機能などが必要になると、ユーザを混乱させることなくそれらの機能をどのように効果的に画面上に配置するかが課題になる可能性がある。他方で、たとえばソフトウェアの統合開発環境では多数の情報や機能が効率的に利用できるようになっていていることを考えると、工夫によってかなり改善できる可能性もある。

ソフトウェアのパフォーマンス

期待される翻訳ツールでは、依存構造の解析やユーザが加える編集情報の取得がリアルタイムで行われる。これらの最新の情報に加え、既存の大量の TM、統計的情報、および RBMT 的情報に基づき、以降の翻訳先テキストが生成される。これらの中で特に統計的 (SMT および EBMT 的) 処理の部分で処理の負荷が増大し、実用時に求められるパフォーマンスに至らない可能性が考えられる。

機能統合によるブラックボックス化

期待される翻訳ツールでは CAT 機能と MT 機能が密接に融合しているため、個別に分離独立させて置き換えることが困難である。また、蓄積される依存構造データや統計的データは、おそらくこのツール上でしか利用できないものになる。他の翻訳ツールに乗り換える場合には、これらのデータの喪失を覚悟する必要がある。機能の統合により、オープン化、モジュール化、標準化の流れに逆行するように見える。

ツールの使用方法の習得

期待される翻訳ツールには、従来の CAT ツールにあるような機能に加え、RBMT エンジンの辞書編集機能なども搭載される。RBMT エンジン辞書の編集は、悪影響の可能性なども考慮しながら慎重に行う必要がある場面もあり、多少の慣れが要求される。予期しない結果を招くことなくツールを効果的に使用できるようになるまでには、従来の翻訳ツールよりも長い習熟期間が必要になると考えられる。

以上、今後期待される翻訳ツールについて自由に述べてみたが、実際の問題としてこれらの機能の技術的な実現方法については何も述べられていない。また、RBMT 的処理と統計的処理との間で競合が起こった場合の解決方法、複数の翻訳者が同時に作業する場合の動作など、詳細な部分で考慮および調整すべき点が多数あると思われる。しかしながら、翻訳者が使用する将来の翻訳ツールの方向性を考えるにあたり、MT/CAT ツールの研究者/開発者の方々の何かの参考になれば幸いである。

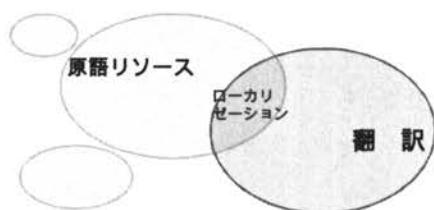
5. おわりに

第 2 部：MT と翻訳（業界）の変遷

第 1 部では、将来的に期待される MT 像について、かなり詳細に説明させていただいた。第 2 部では、より俯瞰的な視点から、翻訳業務における機械翻訳（MT）のポジショニングについて考えてみたい。

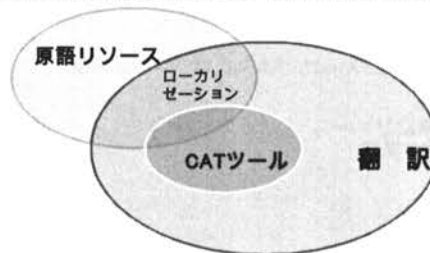
翻訳業界の基本は「原語ソースと翻訳需要ありき」。様々な原語のオリジナルソースに対して、ローカル言語への翻訳需要/サービスが発生する。図 A に示す、原語リソースを表す円に、翻訳（サービス）を表す円が徐々に近づいて重なっていく単純なイメージがそれである。

図 A：まず翻訳ニーズありき（IT業界のケース）



やがて、単に「需要あり×サービスあり」だった世界にさまざまな要求（ボリューム、スピード、等）が追加されていくと、今度はプロセスの効率化を助けるツールが現れる。CAT ツール（Computer Aided Translation Tool）である。既存の翻訳資産を翻訳メモリとして生かしつつ、大量/反復データ処理はコンピュータに任せることで翻訳作業の効率化と標準化に貢献してきた。現実の翻訳作業の渦中から、そのニーズに答えるべく生まれてきたツールである。図 B は CAT ツールが台頭してきた頃のイメージ図で、現在よりも少し前の状況を表している。CAT ツールのユーザもまだ一部の翻訳者のみに限られていた時代である。

図 B：翻訳需要が増え、CATツールが翻訳業界に参入



やがて、CAT ツールは産業翻訳界に浸透していき、図 C のように勢力を伸ばしてきた。台頭からこちら、このツールの隆盛には実に目を見張るものがあり、一部のエリアではマストのツールになったといっても過言ではないと思われるが、図 C はそれでもまだ現在の状況を示しているとは言えない。なぜなら、ここにはまだ足りない要素があるからだ。

図 C：翻訳需要はますます増え、CATツールが翻訳業界に定着

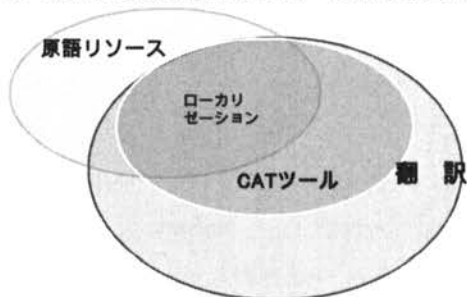
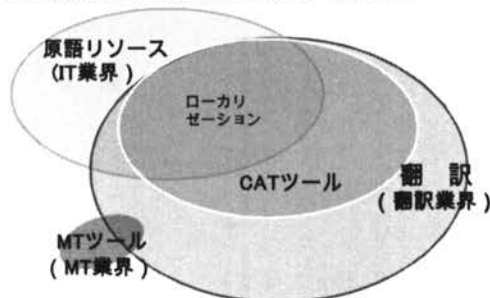


図 C に含まれていない要素とは、ずばり、MT である。その要素を加えたものが図 D で、90 年代の第一次 MT ブームの頃のイメージとお考え頂きたい。図 D で特にご注目頂きたいのは、MT を表す円が翻訳サービスを表す円の外側からアプローチしている点だ。CAT ツールが図 B で翻訳サービスの円の内側に存在しているのとは対照的である。MT は、CAT ツールと同様に翻訳を支援するツールであるものの、翻訳業界からは少し離れた場所で生まれた技術を基礎としていることが、CAT ツールとの大きな違いである。

通常、翻訳は人手によるサービスなので、当然ながら人間の言語処理能力をそのまま利用する。すでにある能力を利用するだけなので、あらたに学究的な処理理論は必要としない。反対に、MT では自然言語処理/言語解析にまつわる部分こそが生命線。そのため、翻訳業界とは隣り合っているものの、や

や離れたところで発展してきた研究結果が、その進化に大きく寄与してきたのである。

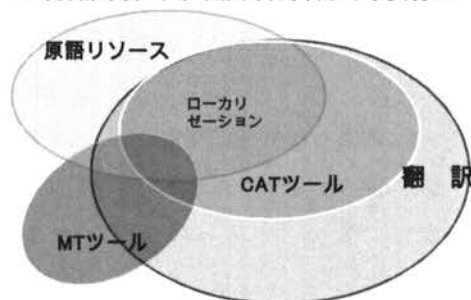
図 D：翻訳にコスト減が求められ、MTツールが台頭



90年代の第一次 MT ブームでは、さまざまな MT エンジンが華々しく市場に紹介されたこともあって、世間一般的にも翻訳業界的にもかなり大きなインパクトを与えたが、事実上の成果は実験段階にとどまり、本格的に翻訳業界に参入したとはいえなかった。弊社でもこの時代に MT を導入したことがあるものの、その MT エンジンでは残念ながら大量翻訳プロセスを作り上げるには至っていない。

しかしながら、その後 21 世紀に入って、MT の性能は 90 年代とは一線を画すほどの飛躍を遂げ、一部の産業翻訳界にはかなりの食い込みを見せるようになってきた。その健闘ぶりは図 E をご覧頂きたい。これが現在のイメージ図である。CAT ツールが翻訳業界ですでに不動の地位を築いているのに比べればまだまだ MT には開拓の余地がありそうに見える。

図 E：両者かなり歩み寄ったが、まだギャップのある現状？



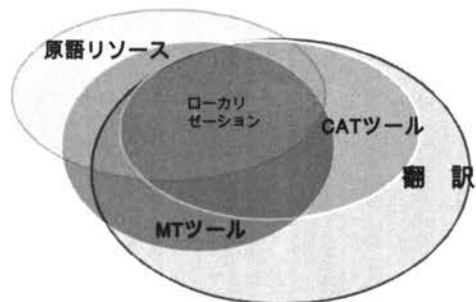
CAT ツールと MT のポジショニングに開きができている状況については、もちろん、それぞれに成り立ちの歴史なり、開発元各社の事情なりがあるわけだが、エンドユーザにとってはいずれも産業翻訳に欠かせないツールであることに変わりはない。

両者がお互いの存在意義を認め合い、同じ翻訳

サービスの中でユーザを支えるツール“同志”として、ツールの垣根を超えてコラボレーションの可能性を模索してほしいと願ってやまない。

図 E では特にご注目頂きたい点がある。MT を示す円が、原語リソースを示す円とのみ重なっている部分（つまり翻訳の円とは重なっていない箇所）なのだが、これは、翻訳プロセスにおけるリバースエンジニアリングとも呼べる新しい動きである。弊社での実例を紹介させていただきたい。弊社では MT をローカライゼーションプロセスの中心に据えており、MT による大幅な効率化への期待も高い。そこで、MT で処理しにくいパターンの原語リソースは原語側から変えてしまおう、という動きが出てきている。これまで翻訳業界の常識では、「原語リソースは変更できないもの」だったと思われるが、もし、この弊社で見られるようになったリバースエンジニアリング的な流れが定着していくなら、MT をプロセスに組み込むことが、これまでの常識的な翻訳プロセスとは逆の流れ（＝ローカル言語から発信して翻訳の標準化を促す）を生み出すことにもつながる。原語からローカル言語へ、というベクトルが変わることなく続いてきた翻訳の世界において、MT だからこそ新たな可能性をもたらし、新風を吹き込むこともできるのでは、と感じる。

図 F：ユーザの思い描く理想像
“翻訳”の中で両者が共存共栄してほしい



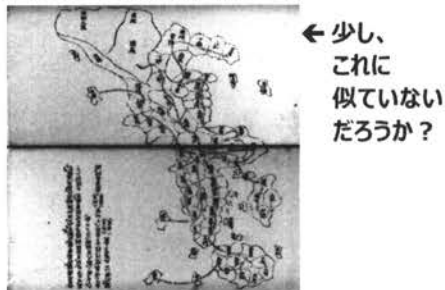
ただ現時点で多くの方は、いまだ MT に対して旧態依然としたイメージしかお持ちではないのではないだろうか。前にお見せした、90 年代の第一次 MT ブームの頃の図 D の時代から、印象が刷新されていないのではないか、と危惧する。

“なんだかよく知らないが MT というものもある

らしいね…、聞いただけで自分もこれまで使ったことはない…、使う機会もなさそうだなあ…、知人から使用感をきいたこともないし…、ウワサによるとまだ使いづらいらしいよ、単にウワサだけだね…”、とまあ、こんな具合なのではないかと。

この状況は、なにやら奈良時代の僧侶、行基の作ったとされる日本地図を想起させる。この地図（図 G）に描かれた東北および北海道エリアの茫漠たる描かれぶりと、現在の翻訳業界における MT のポジションは、どことなく似ていると思われるのだ。誰も行ったことがないからわからないエリアも誰も使ったことがないからわからない MT。

図 G：MTが普及しない状況分析 (1/3)



人間は、自己保存の本能から未知のものには近づきたがらない。うかつに近寄って怪我をするくらいなら近づかない。君子危うきに近寄らず（図 H）である。とはいえ、対象を身近に感じられるようになれば、この考え方も変わっていくはずだ。

図 H：MTが普及しない状況分析 (2/3)

➤ 何事によらず、“未ユーザ”である理由は

「よく知らないから」
「使ったことが無いから」
「使いたくないから」

➤ 要するに…

プロバガンダが進んでいない（便利だと知られてない）
使う機会がない（簡単なら使いたい人もいるはず）
ユーザの反応を知らない（生の声を聞けば気も変わる）

もうひとつ、実際に使っているエンドユーザとして申し上げておきたいのは、現実の翻訳作業と並行して MT を併用しづらいことも普及を妨げている理由のひとつではないかということだ（図 I）。この点は第 1 部でも言及させていただいた通りである。

翻訳ファイルの処理と MT 処理の両方を別々に作業しなければならないケースはまだ多数派だと思われるが、このような状態では、労力/時間（つま

りコスト）の問題から MT に手を出せないと判断するユーザは多いだろう。

だが、もしワンストップで全ての作業が片付いて、しかも効率化を実現できるとしたら？ これなら、労力/時間（つまりコスト）の問題にアタマをかかえるエンドユーザにも魅力的に映るはずだ。

図 I：MTが普及しない状況分析 (3/3)

➤ 既ユーザからの意見

翻訳作業時に並行して使いづらい
→ メンテナンスに労力が必要
→ 専任者が必要

➤ ということは…

■ メンテナンスに労力が必要ならば、
→ つい、使いたくなる、かも？
■ 専任の MT 担当者が必要ならば、
→ 小規模組織では使いにくい？

最後に、おまけで産業翻訳における七つ道具の新旧版を作ってみた（図 J）。

図 J：産業翻訳の七つ道具（旧 vs. 新）

旧	新
言語知識／一般辞書	言語知識／一般辞書
業界知識	業界知識
業界に特化した用語集	翻訳メモリ
電話/FAX/モデム	ネットワーク環境
ワープロ	PC（コンピュータ/タラシ）
PC、CATツール	CATツール
翻訳メモリ	MTツール

左側は少し前の時代の「翻訳七つ道具」で、右側が我々の理想をいれた「翻訳七つ道具」。右側の新版は、目下弊社で日々の業務に活躍中である。

Crowd and Computer Translation Collaboration at Baobab

株式会社バオバブ

相良 美織

この度、2013 年 AAMT 長尾賞という過分な賞を頂戴しました。驚き感激すると同時に恐縮して身のすくむ思いでございますが「もっと日本の機械翻訳発展の為に頑張りたい。」という叱咤と受けとめております。

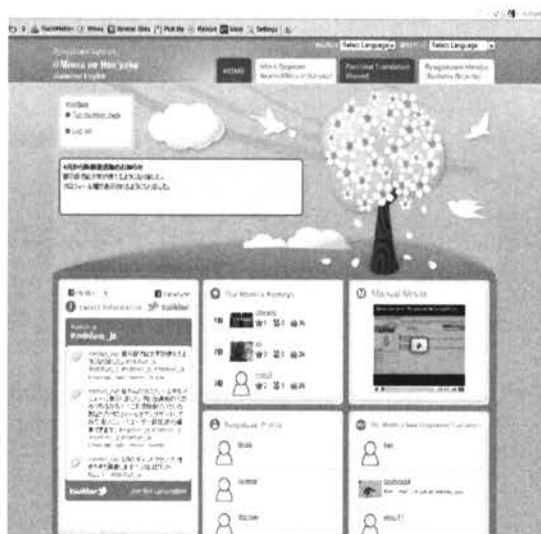
心からこれまでご指導頂いた皆様に厚く感謝を申し上げますとともに、紙面を借りて私どもの活動をご紹介しますと思います。

1. はじめに

株式会社バオバブは 2010 年に設立されました。2008 年に某省による国内の E コマース海外展開の後押しの為、ボトルネックとなっている翻訳の解決策を検証する研究会に私が携わったのがきっかけです。それまで全く翻訳のバックグラウンドはありませんでした。次年の 2009 年あいにく当該プロジェクトは諸々の事情により継続されませんでした

図 1 「留学生ネットワーク@みんなの翻訳」

トップページ



が、E コマース運営会社からは海外展開の意向に伴う高精度な自動翻訳構築へのご要望、そして株式会社国際電気通信基礎情報研究所 (ATR) と京都大学の先生方からは「社会的意義が非常に高く研究的見地からも挑戦したい分野であるので是非何らかの形で検討を継続してほしい。」という後押しを受けて、E コマースに特化した機械翻訳エンジン構築への試みが始まりました。しかしながら、その過程において早々に「機械翻訳構築の為に大量の対訳をいかに安価に構築するか」という壁にあたることになります。この高く厚い壁を打破する為に独立行政法人情報通信研究機構 (以下、NICT と記述) による「みんなの翻訳」という翻訳支援ツール (参照: AAMT ジャーナル 2009 年 6 月号) を基盤とした「留学生ネットワーク@みんなの翻訳」を 2010 年 3 月にリリースしました。つまり、もともとは「留学生ネットワーク@みんなの翻訳」はコーパスベース機械翻訳の為に対訳構築の為に構築されたプラットフォームでした。

当初は国内の留学生を中心に翻訳者を募集しておりました。(現在は 40% 程度がプロフェッショナルの翻訳者の方々になります) 言語は、英語・中国語・韓国語になります。

2. 「留学生ネットワーク@みんなの翻訳」の特徴

前述の通り、「留学生ネットワーク@みんなの翻訳」(以下、RMH と記述) は NICT の「みんなの翻訳」を基にいくつかツールを追加し、システムをあらたに再構築したクラウドトランスレーションプラットフォームです。主な特徴は以下になります。

- ① 機械翻訳を下訳として提示することで効率とスピードを向上
- ② 用語登録・類似文章の提示・辞書のワンクリックサーチ等翻訳支援ツールの充実
- ③ チャットや掲示板など翻訳者同士のつながりを促進するソーシャルなツールを実装
- ④ ポイントなどゲーミフィケーションの要素を採用し、楽しく仕事ができる様に工夫
- ⑤ “翻訳が難しい”（例：原文に誤字がある場合等）文章をスキップ出来る「スキップ機能」により、翻訳作業のスタックによるストレスから翻訳者を解放
- ⑥ 翻訳文字数カウント及び給与計算・抜き打ちテスト自動生成・一斉メール送信等、管理者の負担を軽減させ運営間接費を削減

3.課題-安定期な品質保持と生産性の最大化

RMH は前述の通り様々なツールの実装をしています。これらは単に翻訳作業の効率やスピードを改善させるだけでなく、以下に効果がある事が認められています。

- ・ 翻訳達の定着を保持
- ・ 時間あたりの生産性向上
- ・ 品質の安定化

当初はあまり想定していませんでしたが、このようなツールの実装がいわば **human management** の一部代替を担っています。

翻訳単価を上げる事は、翻訳者にとって大きなインセンティブになることは間違い有りません。一方、翻訳単価を上げ続ける事は顧客に提供する単価に影響を与えます。

「人間は何をインセンティブとして動くのか」は我々の永遠の課題であります。

「いかに楽しんで翻訳の仕事をしてもらうか」このソリューションにも私たちはシステムや多様なツールを採用しています。

沢山の失敗を経て、現在は

「自分の家族もこのアルバイトをやりたいと言っている。」

「友人からすすめられて、このアルバイトを見つけた。」

図2 「留学生ネットワーク@みんなの翻訳」編集画面



「母国に帰国する事になったがこのアルバイトは続けたい。」という口コミにより自然に翻訳者が集まってくるようになり、また翻訳者達の定着率も飛躍的に上がりました。

翻訳者同士が、「この単語、どうやって訳するべきだと思う？」というようなディスカッションを自発的に活発にできることはプラットフォーム全体の成長を促しながら、翻訳者個人に満足感と帰属意識を持たせる事を可能にしています。

4.更なる多様なサービスの提供とプラットフォームの充実を目指して

RMH は 11 月に大きなリニューアルを予定し、「BAOBAB」サイトとして生まれ変わります。特にソーシャルな翻訳者同士のつながりやゲーミフィケーションの要素を強化し、より作業効率を向上させことを狙います。

また、翻訳以外のタスクも本格的にご提供して参ります。

- ・タグ付け
- ・書き起こし
- ・文書分類
- ・順位付け
- ・機械翻訳の人力評価

これにより、画像や音声を書き起こし→翻訳などもワンストップでご提供する事が可能になります。

機械翻訳の評価は、「人力評価は費用が高いが、BLEU は信頼性に欠ける。」というニーズに基づき、機械翻訳研究機関にヒアリングし 5 段階の基準を設定し、評価者育成のドリルを作成しました。

今後は、多言語資源データを構築するプラットフォームに大きく脱皮して参ります。

5.人力翻訳マーケットシェア縮小への懸念と Google の脅威について

今回の長尾賞受賞式の後の懇親会にて、「機械翻訳の品質が向上したら、人力翻訳業は今後一体どうなってしまうのか?」「Google 翻訳のシェアがますます

拡大していると、日本の機械翻訳の業界は衰退していくのではないか?」という声を多くうかがいました。

この 2 点について以下に私見を述べたいと思います。

①人力翻訳のマーケットについて

機械翻訳の実務化に携われば携わるほど、人力翻訳の品質にどうしても及ばない分野に無数にぶつかります。

「村上春樹氏の小説をなんとか機械翻訳出来ないか?」と問い合わせを受けたことも有りますが、まさにそのような翻訳は村上ワールドを理解されている翻訳者の方の付加価値がいかに発揮される場であり、当分機械翻訳には難しいと思われます。

もともと、バオバブが始めた「E コマースの大量商品説明文章の翻訳」は、それまで人力翻訳では莫大な費用と時間がかさむため、市場はほぼ 0 でした。機械翻訳の採用により、商品説明文章以外の箇所（ヘルプページや個人情報の扱いなど）の人力翻訳や機械翻訳訳出結果への人力による修正業務が発生しました。

- ・高品質・誤訳の許されない箇所は人力翻訳
- ・商品の回転が早く「そこそこの品質でよい」膨大な商品説明文章の翻訳には機械翻訳

というように適材適所に使い分ける事で、顧客に未開拓市場参入、そして新たな人力翻訳への供給を生むと信じています。

②Google は脅威か

もちろん無料の Google 翻訳の精度向上は我々にとっても好ましい事象ではありません。しかしながら、ご提供する顧客の目的やニーズにきめ細かく対応し、ある一定のセグメント（分野）に限定すれば、彼らの精度を上回るのはそれほど難しくないと考えています。

最後に

機械翻訳の翻訳精度はまだまだ大きく改善の余地があります。ルールベースによる翻訳は確かに成熟期を迎えているかもしれませんが、アカデミックの世界では日夜新しい研究や試みが行われており格段の進歩を遂げています。

日々、クライアントと接してニーズやレビューを伺い、実用に耐える新しい機械翻訳の時代が到来している事を確信しています。

今後もますます日本の機械翻訳を発展すべく邁進して参ります。どうぞ宜しくお願い致します。



翻訳支援ツール Transit^{NXT} のご紹介

株式会社 シュタール ジャパン 目次 由美子

はじめに

株式会社 シュタール ジャパンは、技術翻訳、DTP、情報処理を業務とし、スイスに本社を置く 31ヶ国 47 拠点を擁するシュタール グループの一員として活動しています。現地在住のネイティブスピーカーが翻訳を担当する「翻訳現地主義」を採り、多言語翻訳でも豊富な実績を上げてきました。25 年以上に渡って培われてきた翻訳会社としての経験と知識は、ソフトウェアの自社開発にも活用されており、高品質の翻訳と最先端のテクノロジーの両方をお客様にご提供できる唯一のソリューションプロバイダーとして、グローバルに活躍する翻訳会社の中でもユニークな地位を確立しています。

今回ご紹介する Transit^{NXT} は、スイス本社とドイツの専門開発チームによって作成された翻訳支援ツールであり、シュタール ジャパンにて日本語化し、販売とサポートを行っています。

Transit のあゆみ

Transit は 1980 年代末にシュタール社内でインハウス使用が開始されました。日本語に対応し、日本で販売を開始したのは 1996 年末のことです。2009 年にリリースされた最新版の Transit^{NXT} (トランジット ネクスト) は、機能性と効率性を飛躍的に高め、かつ優れた操作性を備えたシュタール社の技術力の結実ともいえます。

シュタール ジャパンで受注する仕事の中でもっとも多いのは、ドイツ系メーカーのサービスマニュアルの日本語へのローカライズです。ドイツ語から日本語に翻訳し、データ処理、版下作成、印刷までを行います。

過去にまったく同じ文章があったとしても、新しいマニュアルを作成するためには一から翻訳をし

なければならないというのはドキュメント担当者にとって悩みの種でした。たとえば、モデルチェンジに際して改訂が発生する自動車メーカーにとっては、新旧マニュアル間の相違箇所を把握するだけでも大変手間のかかる作業でした。

Transit は新旧バージョンを照合し、旧版にはなかった部分のみを翻訳者が翻訳できるように処理を実行します。翻訳者は「ファジーマッチ」機能を使用して旧版に含まれる類似する文章を検索し、過去の訳を提示することで、新版の翻訳に流用することもできます。多言語用語管理ツールの TermStar が連動するため、Transit を使用する翻訳作業では専門用語の訳の統一が容易に実現できます。また、「グローバルオープン」機能により Transit では複数の翻訳ファイル(ランゲージペア)を仮想的に一つのファイルとして一括して開くことができるため、品質チェックを効率良く実行することが可能です。さらに、Transit^{NXT} に新たに装備されたプレビュー機能を利用して、レイアウトを確認しながら翻訳作業を進めることもできます。

翻訳が完了した後に実行するエクスポートによって原文のレイアウトが再現されるため、翻訳済みファイルの DTP 作業も大幅に軽減されます。

Transit での作業は、I. [インポート]、II. [翻訳]、III. [エクスポート] と大別することができます。Transit を使用してドキュメントを英語から日本語へ翻訳する場合の作業手順を簡単にご紹介いたします。



I. [インポート]

ステップ 1: プロジェクトの作成

Transit^{NXT} で新たに翻訳を行うためには、プロジェクトの作成から始めます。ある翻訳ジョブに対して Transit で必要とされる設定は総称して「プロジェクト」と呼ばれます。翻訳対象の原文ファイルの格納場所や翻訳言語などをプロジェクトに設定します。Transit を使用する翻訳は、プロジェクト単位で管理されます。



ステップ 2: フィルタリングとセグメンテーション
新しく作成されたプロジェクトの設定内容に基づき、翻訳対象のファイルが Transit 形式に変換されます。まず「フィルタリング」によって翻訳対象ファイルのテキスト情報とレイアウト情報が分離され、ランゲージペア（原文用と訳文用テキストファイル）と COD（レイアウト情報保存ファイル）が生成されます。次に、「セグメンテーション」によって文または段落単位で区切りが挿入され、新旧のテキストを比較するための単位となるセグメントに分割されます。



ステップ 3: プリトランスレーション

プリトランスレーションでは、原文用ファイルと翻訳メモリのそれぞれの原文がセグメント単位で比較されます。完全に一致するセグメントが検出された場合、訳文用ファイルの該当セグメントは、対応する翻訳メモリの訳文で自動的に置き換えられます。つまり完全一致が見つからなかったセグメントは原文のまま残るため、訳文用ファイルには訳文（日本語）と原文（英語）が混在することになります。

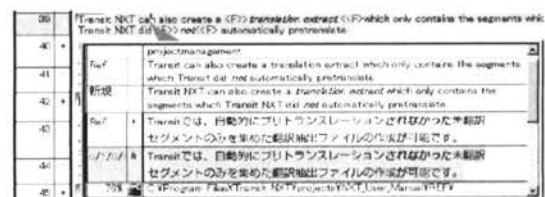


II. [翻訳]

ステップ 4: ファジーマッチ、辞書、プレビュー機能の利用

自動的に実行されるプリトランスレーションを含むインポートの工程を経て、原文のまま残っているセグメントを人手により翻訳していきます。訳文ファイル上ではプリトランスレーション済みのセグメントはスキップし、未翻訳のセグメントへジャンプするショートカットキーも利用することができます。

ファジーマッチ（あいまい検索）を利用することにより、翻訳メモリ内の 99%以下の類似率を持つセグメントを検索して提示させ、類似する過去の訳例を新規翻訳に流用することができます。



Transit と連動する TermStar 辞書には、技術用語を登録します。Transit での翻訳またはチェック作業を行いながら、TermStar 辞書へ用語を登録することもできます。また、Excel で作成した用語集などがあれば TermStar 辞書に変換することも可能です。TermStar 辞書に登録済みの用語は、Transit エディタ上の文中にハイライト表示され、辞書ウィンドウにその訳語が表示されます。訳語は訳文用ファイルへショートカットキーで簡単に取り込むこともできます。翻訳者が自ら辞書を引くまでもなく登録された用語が提示されるので、用語統一に効果を発揮します。

さらに、Word、InDesign などの DTP データには「PDF プレビュー」で原文のレイアウトや文脈を確認しながら、HTML や XML などのマークアップ言語には原文のレイアウトに加え訳文のレイアウトまでも「HTML プレビュー」で確認しながら翻訳を進めることができます。DLL や EXE ファイルについては「リソースエディタ」を使い、長くなってしまった訳文テキストに合わせ、ダイアログのリサイズを行うことも可能です。



III. [エクスポート]

ステップ5: セグメントマーカの削除とレイアウト結合

インポートのステップ2で挿入されたセグメントマーカを自動的に除去し、分離した訳文ファイルとレイアウト情報を再結合させます。こうして原文のレイアウトが再現された翻訳済みのドキュメントができあがります。

エクスポート後、ファイル形式によってはレイアウト調整やコンパイルなどの作業も必要となります。また、将来の改訂に備え、翻訳メモリや辞書を更新管理することも翻訳支援ツールを継続使用の上での重要な作業の一つです。

翻訳支援ツールの導入初期は翻訳メモリや辞書に用語が十分に蓄積されていないため、効果を充分には実感できない場合もありますが、使用を継続することによって効果は増幅し、その恩恵を享受することができます。

Transit のこれから

欧米では20年以上に渡り、日本では特にこの10年、翻訳支援ツールは翻訳業界に浸透してきました。膨大なテキストを含む自動車業界やソフトウェア業界をはじめ、さまざまな分野でのマニュアル翻訳において、翻訳支援ツールは絶大なる効果を発揮し、受け入れられてきました。フリーウェアを含め多数の製品がありますが、翻訳支援ツールに求められる機能は標準化されつつあります。

昨今、産業翻訳およびマニュアル作成に関連するセミナーなどでは、機械翻訳ツールの台頭が目覚ましく見受けられます。たとえば用語辞書の充実や原文の洗練によってより適切な訳を引き出すなど、機械翻訳ユーザーによる工夫や努力も多く紹介されています。同時に、ポストエディティング（機械翻訳後の人手による改訂）の重要性は依然として高いままです。

機械翻訳で生成された結果に対するポストエディティングを翻訳支援ツールで行うのが当然であるというような将来も遠くないのではないかと、個人的には考えています。Transit^{NXT}は翻訳作業だけでなく品質管理・チェック作業にも優れた機能を兼ね備えています。機械翻訳済みのデータに対しても、多様な機能をあますところなく活用できると信じています。

Transit^{NXT}は機械翻訳との統合が可能です。現状として、この機能はまだ多用されていませんが、機械翻訳の需要が増大していく中、翻訳支援ツールもその機能を積極的に取り入れていくべきと考えます。産業翻訳におけるこの2大テクノロジーの融合がもたらす相乗効果は計り知ることができません。さらに効率の良い翻訳作業を実現するためのツールとして、Transit^{NXT}が先駆として大いに活用されることを切望しています。そして、Transit^{NXT}およびその他翻訳支援ツールがアジア太平洋機械翻訳協会の活動および機械翻訳のさらなる発展の一役を担うことを希望しています。

PC-Transer 翻訳スタジオ V21 に搭載される新機能の紹介

株式会社クロスランゲージ

株式会社クロスランゲージは、2013 年 11 月 1 日に翻訳ソフトウェアの主力製品である「PC-Transer 翻訳スタジオ V21」を発売する。

「PC-Transer」は、1991 年に国内初のパーソナルコンピュータ用「自動翻訳ソフト」として誕生した「Transer（トランサー）」の改定第 21 版となる製品である。PC-Transer シリーズは、専門語辞書搭載による「専門性の高い分野の機械翻訳」、機械翻訳と人間の手により完成させた訳文を反映する翻訳メモリとの「ハイブリッド方式による翻訳」など、常に時代をリードするさまざまな翻訳ソリューションを提供し続けてきた翻訳ソフトである。



今回の改訂では、翻訳者に求められつつある時代のニーズに応えるべく、2つの新機能を搭載した。

1. 一般 XML ファイル対応機能
2. 対訳アライン機能

1. 一般 XML 対応機能

最近では公的機関やデータベース系の翻訳の仕事で、XML 形式で納品するケースが増えてきているため、特許、論文や大型ローカライズ案件など、タグ付き

XML 文章を翻訳者のレベルでもある程度対応できるように機能を追加した。V20 で Microsoft Open XML に対応したのにつき、今回は一般 XML ファイルについても対応した。(XML 1.1 対応)

翻訳エディタでは、XML ファイルを読み込んで文字データを取得し、原文エリアに表示して翻訳を行うことができる。翻訳結果は、タグ情報を保持したまま XML ファイルとして保存することも可能である。

【処理の流れ】

ファイルの読み込み時に、ファイルの種類で XML ファイル形式を選択する。



本ソフトでは、読み込んだ原文の XML ファイルのタグ情報を保持したまま、文字データが翻訳された XML ファイルとして保存することができるだけでなく、通常の対訳翻訳ファイルや HTML ファイルとして保存することも可能である。読み込んだ XML ファイルの構造を以下のようにツリー表示することもできる。

翻訳メモリに登録するための原文とそれに対する訳文が用意されている時、原文とそれに対応する訳文が1対1になるように整列(アライン)させる必要があり、この読み込みから整列までの一連の流れをアライン処理という。

これまでの PC-Transer シリーズ・翻訳エディタでアライン処理を行うには、原文訳文読み込みの機能を使用して原文と訳文を翻訳エディタに読み込ませて、原文の右隣に対応する訳文を手作業により移動させる必要があった。

【処理の流れ】

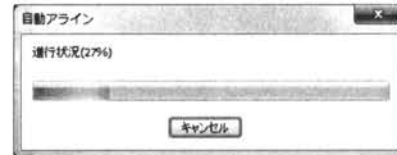
メニューから[翻訳メモリ]-[自動アラインの実行]を選択する。ファイルオープンダイアログが表示されるので、原文となるファイルを選択する。



再度ファイルオープンダイアログが表示されるので、対応する訳文ファイルを選択する。



訳文ファイルを選択した後、自動アライン処理が開始される。

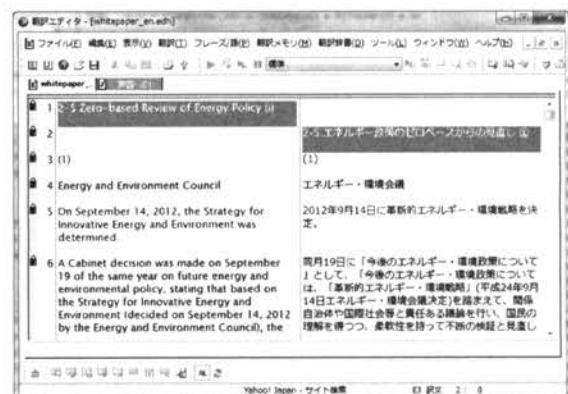


処理が終了すると、自動アラインにより、原文と訳文の対応がとれたものが並んで表示される。これはいままでの編集画面とは異なり、「アラインモード」というアライン作業専用の画面モードである。



対応する原文と訳文が欠けているものは、片側が空白となつて表示される。例えば、下記の例では冒頭の1行目は、訳文側が空白行となつて表示されている。

この自動アライン処理が終わった時点で編集画面はアラインモードになっているので、うまく対応がとれていない場合はアラインの調整を行う。上記例では、左側の1行目と右側の2行目の対応が取れていないが、内容としては対応しているの、それぞれのセル内の、文字のない空白部分をクリックすると、文の色が反転する。



ここで選択した文を右クリックすると、以下のよう
なメニューが表示される。



ここで [対訳設定 (T)] を選択すると、選択した
原文の隣に訳文が移動して、その文は対訳が確定し
たとみなして黄色い色がつく。同時に、アライン処
理における編集へのロックがかかった状態になる。

現実的には、過去の翻訳資産を読み込んだ場合、
原文が 1 文に対して訳文が 2 文であることもあり、
その逆の場合もある。そこで、アラインモードには、
効率的に対訳のアラインを実行するため、下記のよ
うな機能を搭載した。

○ブロックの入れ替え

選択した 2 つの文ブロックの内容を入れ替える
ことが可能。原文側同士、又は訳文側同士で、入れ
替えたい文ブロックを 2 つ選択して、[入れ替え]
を選択して実行する。

○文章の連結

選択した文ブロックの文章を 1 つに連結する。原
文側同士、又は訳文側同士で、連結したい文ブロ
ックを 2 つ以上選択して、[文章の連結] から、用途に
応じて以下の種類の項目から選択して実行する。

- ・ 上の文章の末尾に連結
- ・ 下の文章の行頭に連結

○文章の分割

文章内のキャレットの位置から文章を 2 つに分
割する。文章内の分割したい位置をクリックしてキ

ャレットを移動させ、[文章の分割] を選択して実行
する。

○空行挿入

選択したセルの、上または下の行に空行を挿入す
ることにより、アライン状態を調整する。

○自動アラインの再実行

[自動アラインの再実行] をクリックし、自動アラ
インを再実行すると、[対訳設定] で確定した対訳間
の未確定の原文と訳文を自動的にアラインするこ
とができる。



これにより、うまく自動アラインできない複雑な
文書でも、数か所を対訳設定してやるだけで、残り
は自動的に対応をとってくれる。

これらの新機能とは別に、既存の原文訳文読み込
み機能（自動アラインなし）もあるので、ユーザー
は使いやすい方を選択すること可能である。

これらの機能は PC-Transer 製品シリーズには今
回初搭載であるが、今後もユーザーの使い勝手と効
率性を上げるべく、研究開発を継続し、成果を製品
に反映してゆく予定である。

【本製品に関してのお問合わせ先】

株式会社クロスランゲージ マーケティング 2 部
電話：03-5215-7631 / info@crosslanguage.co.jp

中国語・韓国語の最新翻訳エンジンを搭載した企業内翻訳サーバ

「The 翻訳®エンタープライズ V16」

東芝ソリューション株式会社

東芝ソリューション株式会社は、翻訳精度を向上した中国語・韓国語の最新翻訳エンジンを搭載し、さらに翻訳の業務効率を高める機能も充実したサーバ型翻訳ソフトウェア「The 翻訳エンタープライズ V16」を発売しました。

1. The 翻訳エンタープライズとは

The 翻訳エンタープライズは、クライアントからブラウザを介して翻訳を実行することができる機械翻訳システムです。

○翻訳環境の共有で均質な翻訳を提供

辞書データや翻訳に使用される設定をすべてサーバ上で一元管理し、利用者全体に均質な翻訳結果を提供します。

○ブラウザ、メールで翻訳要求

クライアントにプログラムインストールが不要で、ライセンス管理も容易です。

○情報漏洩リスクを軽減

イントラネット内にサーバを構築し、情報漏洩リスクを軽減します。

○大量文書の高速翻訳

マルチ CPU を使用した並列処理が可能です。

<翻訳対象言語>

- ・英語⇄日本語
- ・中国語（簡体字、繁体字）⇄日本語
- ・韓国語⇄日本語

<主な機能>

- ・ホームページ翻訳
- ・テキスト翻訳
- ・メール翻訳

・ファイル翻訳：

テキスト、Microsoft® Word/Microsoft® Excel/PowerPoint®、PDF、HTML、XML、SGML

<翻訳方法>

- ・ブラウザ、Microsoft® Office に組み込んだツールボタンからの翻訳
- ・ブラウザ画面から翻訳
- ・メールにより翻訳原文送信
- ・CGI を利用した外部プログラムからの翻訳
- ・Ajax インターフェースを用いた外部プログラムからの翻訳
- ・API による外部プログラムからの翻訳

2. 新バージョンの強化点

The 翻訳エンタープライズ V16 は、当社従来商品と比較して、中日/日中翻訳、韓日/日韓翻訳の精度を向上しました（中日翻訳 12%向上、韓日翻訳 25%向上（*1））。専門用語辞書の分野増強などにより、辞書知識を強化（中日/日中 173 万語、韓日/日韓 155 万語）し、韓国語翻訳では、当社独自の技術である統計的訳語選択の知識を強化しました。

翻訳エンジンの高速化も行ない、翻訳サーバとしてのパフォーマンスを向上しました。前バージョンでは、韓日/日韓翻訳の機能に一部制限がありましたが、ホームページ翻訳やファイル翻訳（*2）も行えるようにしました。

3. 韓日翻訳エンジンの統計的訳語選択技術

韓国語と日本語は語順が非常に類似しているため、英日翻訳や中日翻訳のような原文から訳文への大きな構造変換は必要ありません。韓国語の語順を可能な限り利用して日本語に変換します。一方、韓

国語を表記するハングルは表音文字であり、表意文字の日本語に翻訳する場合、同音（同綴）異義語が非常に多いという特徴があります。したがって、訳語選択の技術が韓日翻訳の精度を左右します。

統計的訳語選択技術は、翻訳対象の文書に関連した文書集合から、単語の並びに関する統計情報である言語モデルを学習し、これに基づいて、複数の候補から最適な訳語を選択するものです。学習対象の文書集合を特許文書とすることで、特許文書の翻訳に適した訳語選択も可能になりました。

4. その他の特長

4.1 翻訳支援ツールのデファクトスタンダード「SDL Trados Studio」(*3)との連携

The 翻訳エンタープライズのプラグインを追加すると、SDL Trados Studio の画面上に、The 翻訳の訳文を自動で挿入することができ、訳文作成の手間をさらに軽減できます。SDL MultiTerm（用語DB）と The 翻訳エンタープライズのユーザ辞書は、データの相互利用が可能です。また、SDL Trados Studio の翻訳メモリを The 翻訳エンタープライズで利用することもできます。

なお、SDL Trados Studio と The 翻訳エンタープライズとのデータ相互利用は、英日/日英翻訳のみで可能です。

4.2 他システムへの翻訳機能の組み込み

①API を用いた連携

The 翻訳のライブラリ（C/Java）を使用して、業務プログラムから直接翻訳機能をコールし、翻訳結果を得ることができます。

- ・対象はテキスト翻訳とファイル翻訳
- ・C や Java で作成されるプログラムでの翻訳実行に有効

②Web サービスを用いた連携

Web サービスインターフェースにデータ送信することで、翻訳結果を得ることができます。

- ・対象はテキスト翻訳
- ・Web サービスクライアントでの翻訳実行に有効

<ライセンス体系>

- ・CPU ライセンス（利用者数無制限）
- ・ユーザライセンス（最小 5 ユーザ）

<動作環境>

HW：PC サーバ

（The 翻訳エンタープライズ使用分として）

メモリ：384MB 以上（韓日/日韓 1.5GB 以上）

ディスク：2.0GB 以上

（他の作業領域として 2GB 以上を推奨）

OS：【Windows 版】

日本語 Windows® Server 2003（32 ビット版）
/2008（32 ビット版・64 ビット版）

【Linux 版】

Red Hat Enterprise Linux 5.7～5.9、6.1～
6.4、CentOS 5.7～5.9、6.1～6.4

- *1 特許文書の翻訳精度を BLEU スコアで機械評価したもの（当社調べ）
- *2 Linux 環境でのファイル翻訳は行えません。
- *3 SDL Trados Studio は SDL 社の翻訳ソフトウェアです。

■製品情報

「The 翻訳エンタープライズホームページ」

<http://mt-server.toshiba-sol.co.jp/>

- ・本内容は予告なく変更する場合があります。
- ・The 翻訳は、東芝ソリューション株式会社の商標です。
- ・Microsoft、Windows、PowerPoint は、米国 Microsoft Corporation の米国及びその他の国における登録商標または商標です。
- ・Java は、Oracle Corporation 及びその子会社、関連会社の米国及びその他の国における登録商標です。
- ・本文中の製品名称はそれぞれ各社が商標として使用している場合があります。

AAMT 機械翻訳文テストセットに基づく 日中機械翻訳の自動評価サイト作成報告

機械翻訳課題調査委員会 ワーキンググループ 1

1. はじめに

アジア太平洋機会翻訳協会の機械翻訳課題調査委員会（以下、AAMT 調査委員会）では、機械翻訳システムの翻訳品質を評価する方法として「テストセット評価」に着目し、機械翻訳システムの評価のためのテストセットを作成してきた。2011 年度には日中機械翻訳システムを対象とした基本文テストセットを作成し、AAMT ホームページからのダウンロードを可能とした。2012 年度以降はテストセット評価の自動化について検討を進め、この度、Web ブラウザ上で手軽にテストセット評価を実行できる自動評価サイトを AAMT ホームページで公開できる運びとなった。本稿では、テストセット評価の概要と他の評価手法との関係について述べた後、テストセット評価に基づく自動評価サイトの構成、使い方等について紹介を行う。

2. テストセット評価とは

テストセットとは、翻訳言語対毎に整備した原文、参照訳（正解）の対と、それぞれの文対に付与した設問からなるデータの集合を指す。

表 1 テストセットのサンプル

原文	The trash can was thrown away.
参照訳(正解)	ごみのカンが捨てられた。
設問	“can”が「カン/缶」のように名詞で訳されていますか？

機械翻訳文の品質評価については、これまでも多くの方法が提案されてきた。実際に用いられている方法としては、正確性や流暢性の観点で人間が評価する主観評価と、参照訳（正解）と機械翻訳文の類似性を数値化することで評価を行う自動評価に大別されるが、どちらの方法も評価結果は文単位あるいはドキュメントの単位で品質が数値化されるだけのものである。これ

に対して「テストセット評価」では、あるエンジンの翻訳結果が他と比べてどこが優れて（劣って）いるのかという種類の定性的な情報も得ることができる。

たとえば表 1 の設問の解答が“no”であった場合、そのシステムは助動詞に関わる品詞選択が適切にされない可能性があることを評価者が知ることができる。この種の情報は、特に翻訳システムの研究開発者にとって有益といえる。

3. 他の評価方法との関係

図 1 は人間による主観評価と従来型自動評価との比較において、テストセット評価の位置づけを示したものである。

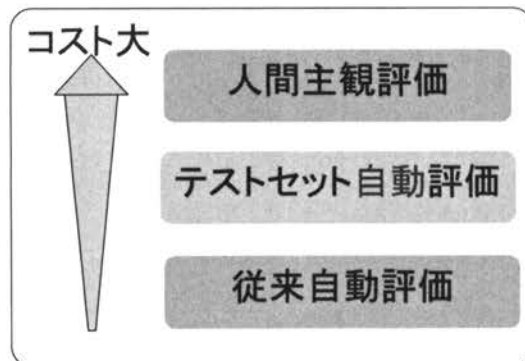


図 1 評価手法の比較

人間主観評価は、人間が一文ずつ機械翻訳文を読んで評価を行うことから非常に大きなコストがかかる。これに対して BLEU 等に代表される従来自動評価は、参照訳（正解）を用意するだけで評価は自動的にやってくれる。テストセット評価は、参照訳に加えて設問作成を行う必要があり、コスト的には人間主観評価と従来自動評価の中間に位置すると考えられる。

テストセットの各設問回答について “yes” を 1 点、

“no”を0点とみなし、テストセット全体の平均を計算することによって、テストセット評価の結果を数値化することができる。このようにして算出した評価値は、人間による主観評価値との間に高い相関があり、従来自動評価値と主観評価値との相関と比較しても遜色のないものである。

4. 自動評価サイトの作成

AAMT 調査委員会では、日中機械翻訳システムの評価用に AAMT 機械翻訳文テストセットおよび自動評価プログラムをサーバ上に配置し、Web ブラウザ上で中日機械翻訳評価が行えるサイトを作成した。評価用テストセットは 251 文である。

設問評価の自動化は、図 2 に示すように機械翻訳文との文字列マッチングに基づいて“yes”, “no”を判定している。同義語についても考慮されている。

```

read a line;
if the line match /严|严厉|严格|厉害/ then
    print "Yes";
else
    print "No";
endif

read a line;
if the line match /去买东西|去购物|去逛街|去逛商店/ then
    print "Yes";
else

```

図 2 設問評価の自動化プログラム (pseudo code)

今回作成した評価用サイトは下記の特徴を持っている。

- ・世界中のどこからでも Web ブラウザ上でテストセットに基づく日中翻訳の評価が可能
- ・評価結果は 10 の文法項目(図 4 の各軸に対応)別に点数表示される
- ・代表的な日中システム(Web で公開中の 6 種)の平均値と比較できる

評価サイトの構成と利用手順を図 3 にまとめる。利用者は、まず、原文(日本語)をサーバからダウンロードし、

評価対象の日中エンジンを用いて原文をすべて中国語に翻訳する。次に、翻訳結果を Web ページ上の領域にペーストしてサーバに送信すると、評価結果がブラウザ上に表示される。テストセットは全部で 251 文であるがダウンロードされる原文には大量のダミー文を加えている。これは、ベンダがテストセットの原文に絞って品質を上げることを困難にするための対策である。(テストセット単体でのダウンロードは、評価サイトの公開に伴い中止とした。)

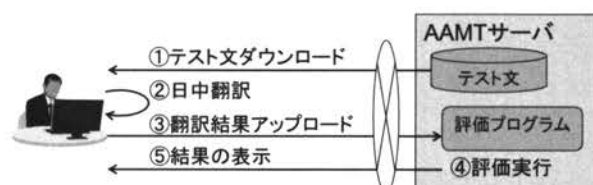


図 3 評価サイトの構成と利用手順

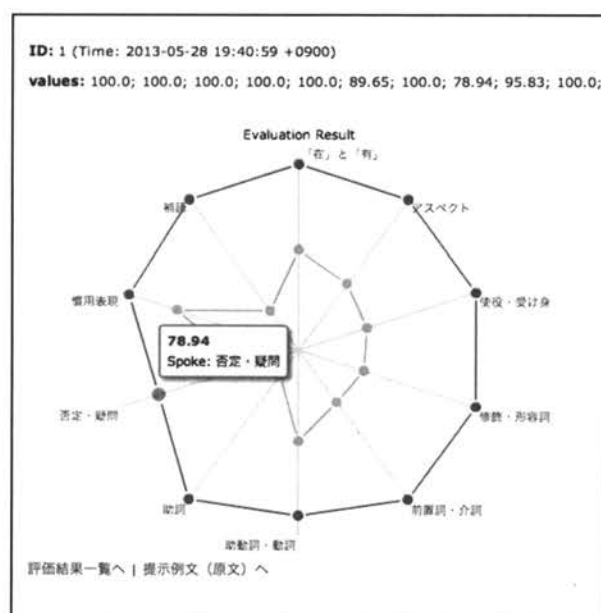


図 4 評価結果表示例

評価結果は図 4 のように 10 の文法項目別に点数がレーダチャートで表示される。評価対象システムと代表システムの平均とを同時に確認することができる。マウスカーソルを評価点のプロット上に持っていくと小窓が開いて数値が表示される。

5. おわりに

AAMT 調査委員会では、現在、中日基本文テストセットと中日特許文用テストセットを作成中であり、2013 年度末までに今回作成した評価サイトの仕組みを活用して公開を予定している。

AAMT会員のひろば

AAMT 会員の新たな交流の場を AAMT Journal 誌面上で提供するべくスタートいたしました「AAMT 会員のひろば」、会員の皆さまのご助力をいただきまして、第一回の No.41 のスタートから第十一回を迎えることができました。今号では、法人会員一社からのご寄稿をいただいております。

独自のお取り組みのご紹介、機械翻訳研究への提言、AAMT の活動へのご要望など、今回も貴重なご意見をお寄せいただきました。

AAMT Journal では今後も引き続き、会員の皆さまからのご寄稿を心よりお待ちしております。

ご寄稿・お問い合わせは AAMT 事務局(E-mail: AAMT-info@AAMT.info)まで宜しく願いいたします。

法人会員（敬称略・50 音順）

会員名

YAMAGATA INTECH 株式会社／YAMAGATA INTECH Corporation

自己紹介

YAMAGATA INTECH 株式会社は、YAMAGATA 株式会社（創業 1906 年）のドキュメント制作部門として 1983 年に発足し、今年 30 周年を迎えました。

「人にものをわかりやすくつたえること」をモットーに、マニュアル・仕様書といったテクニカルドキュメントの制作及び翻訳に携わってまいりました。

昨今ではマニュアルも印刷物に限らず、スマートフォンといった電子媒体向けの HTML マニュアルやアプリの電子化も急速に広がっており、電子コンテンツの企画からソフトウェア開発までもが制作業務の一部となっております。

また、ソリューション部門ではドキュメントの制作支援システムやシステム開発（CMS）、デザイン部門では携帯電話・スマートフォンをはじめ、各種電子機器の GUI 制作など、事業領域は拡大しております。

MT/翻訳とのかかわり

MT および翻訳業界に期待すること

私たちは機械翻訳をとて有望なソリューションの一つと考えています。

しかし、現時点で人力翻訳と比較した場合、機械翻訳はその正確さにおいて一歩遅れを取っているといえましょう。そこで弊社では機械翻訳の弱点を補い、利点を最大限に活かすために人力翻訳の下訳として利

用するのではなく、正確さよりもスピード感を優先するプロジェクトのソリューションとして利用しました。今回はその事例を紹介させていただきます。

現在弊社のヨーロッパ法人である Yamagata Europe（在ベルギー王国）は、あるお客様からの依頼を受けて、2 種類の機械翻訳プロジェクトに取り組んでおります。一つ目のプロジェクトはヨーロッパ各国の代理店から送られた多言語のテキストを 12 時間以内に英語に翻訳するというものです。二つ目のプロジェクトは多言語での同時チャットサービスです。いずれも代理店を含めたお客様社内データベースに組み込まれたものであり、社内での使用に限定されています。

12 時間以内の多言語→英語のスピード翻訳を実現するためには、毎日現地代理店のデータベースに入ってくる保証要求に関する情報の流れを Yamagata Europe で管理する必要がありました。入ってくる情報は Yamagata Europe で英語に機械翻訳され、ウェブ・サービスを使ってお客様に転送されます。その結果、QA エンジニアは有意義なフィードバックを最適なタイミングで受けられることになります。

二つ目のプロジェクト、同時チャットサービスの実現は、まさに機械翻訳に最適な案件でした。各国代理店と技術サポートセンター本部間の技術的なチャットセッションの即時翻訳が不可欠であり、11ヶ国語から英語への翻訳を 2 秒以内に終えなければならないという要求の達成は人力翻訳では不可能だったためです。

このプロジェクトにおいて、お客様からは「90%のわかりやすさ」を要求されていました。しかし、この数字を機械翻訳だけで達成するのは困難だったため、機械翻訳の精度を上げるためには原文の簡潔さが重要であることをエンジニアたちに説明し、トレーニングを行いました。その結果、エンジニアたちは機械翻訳に適したライティングを徐々に身につけていき、翻訳品質も向上していったのです。

機械翻訳の普及によって、翻訳業界にも新たなビジネスチャンスが生まれてくるはずです。我々は機械翻訳を「翻訳」ということだけでとらえるのではなく、この技術を活かして新たな「ソリューション」を生み出していこうと考えております。

YAMAGATA INTECH 株式会社 翻訳ビジネス部 部長 古河 師武

AAMT への要望

弊社は本年度より本協会へ加盟させて頂きました新米でございます。宜しくお願い申し上げます。

協会の皆様から学ばせて頂くことが多々あると思います。機械翻訳のみならず、翻訳品質向上に向けて色々と情報共有させて頂ける機会が増えることを常に願っております。今後とも宜しくお願い申し上げます。

「AAMT 長尾賞学生奨励賞」の新設について

アジア太平洋機械翻訳協会会長

中岩浩巳

アジア太平洋機械翻訳協会(AAMT)の初代会長の長尾真先生が 2005 年日本国際賞を受賞されましたが、同賞の賞金の一部を AAMT の活動に資するようにと AAMT にご寄付くださいました。これを受けて、AAMT では、機械翻訳システムの実用化の促進および実用化のための研究開発に貢献した個人あるいはグループを表彰の対象とした AAMT 長尾賞を設け、2006 年より 11 のグループ・個人が受賞されてきました。機械翻訳の発展に長年取り組んでこられた長尾真先生の機械翻訳への熱い思いを、次世代を担う学生諸君に伝え、将来の機械翻訳の発展に資する研究を見出し、もって学生諸君の奮起を促すことを趣旨として、長尾真先生から頂いた寄付を原資として運営する AAMT 長尾賞学生奨励賞を新設いたすこととし、AAMT 第 23 回通常総会において承認されました。AAMT 長尾賞学生奨励賞の応募要領等の詳細につきましては、AAMT Forum 等を通じてご案内いたす予定です。賞の趣旨をご理解いただき、多くの方の応募をお待ちしております。

名称： AAMT 長尾賞学生奨励賞

審査体制： 「AAMT 長尾賞学生奨励賞審査委員会」により審査を行います。

応募期間： 総会の 2 ヶ月以上前までとします。

賞状及び賞品： 受賞者には賞状と賞品（図書カード 1 万円分）を贈呈します。

応募条件：

- ・受賞対象者は応募論文執筆時点で学生であること。
- ・受賞対象者あるいは受賞対象者の指導教員は AAMT 会員であること。
(会員期間は応募時期（前年度）と受賞時期とします)
- ・受賞対象者は過去に AAMT 長尾賞学生奨励賞を受賞していないこと。
- ・応募者は受賞対象者あるいは受賞対象者の指導教官であること。
- ・応募者は受賞対象者が主筆者として執筆した論文を提出すること。

応募方法： 応募条件を満たす受賞対象者に関して、以下を AAMT 事務局に提出すること。
「応募用紙」、「応募論文」

応募論文： 未発表、もしくは募集開始日から過去 3 年以内に発表した論文（卒業論文、学位論文を含む）かつ、AAMT 長尾賞学生奨励賞に未応募の論文

受賞論文の取り扱い： AAMT 長尾賞学生奨励賞を受賞した論文の概要文を AAMT ジャーナルに掲載します。受賞者には AAMT ジャーナル掲載用の概要文を作成して頂きます。

AAMT長尾賞学生奨励賞規約

第1条

長尾博士による「2005年日本国際賞」受賞を記念し、アジア太平洋機械翻訳協会（以下「AAMT」と略す）は、機械翻訳研究に携わる優秀な研究者と判断される当協会個人会員もしくは当協会個人会員の研究室に所属する個人に対して、年次総会において賞を授与する。賞の名称は「AAMT長尾賞学生奨励賞」とする。

第2条

AAMTは、受賞者の公正な選考のために「AAMT長尾賞学生奨励賞審査委員会」を設置することができる。AAMT会長は3~5名の委員を任命し、そのうち1名が委員長を務めるものとする。

第3条

応募者は総会の2ヶ月以上前までに長尾賞学生奨励賞委員会に提出することができる。

第4条

委員長は、本規約の運用にあたって細則を定めることができる。

「AAMT長尾賞学生奨励賞」規約運用細則

授賞対象について次の通り定める。

- (1) 賞状と賞品（図書カード1万円分）を授与する。
- (2) 受賞対象者は応募論文執筆時点で学生であること。
 2. 受賞対象者あるいは受賞対象者の指導教員は AAMT 会員であること。
（会員期間は応募時期（前年度）と受賞時期とする）
 3. 受賞対象者は過去に AAMT 長尾賞学生奨励賞を受賞していないこと。
 4. 応募者は受賞対象者あるいは受賞対象者の指導教官であること。
 5. 応募者は受賞対象者が主筆者として執筆した論文を提出すること。
- (3) 応募論文は、未発表、もしくは募集開始日から過去3年以内に発表した論文（卒業論文、学位論文を含む）かつ、AAMT長尾賞学生奨励賞に未応募の論文とする。
2. 受賞者は受賞論文の概要文を作成し、AAMT ジャーナルへ寄稿すること。

第 23 回通常総会および関連行事の報告

AAMT 事務局

当協会の第 23 回通常総会が 2013 年 6 月 19 日（水）13 時より・ホテルアジュール竹芝にて開催されました。総会後、各委員会からの報告会、講演会、そして第 8 回 AAMT 長尾賞授与式と受賞者による記念講演会が盛況のうちに行われました。

第 23 回通常総会

1. 開会の辞
2. 会長挨拶 NTT メディアインテリジェンス研究所・中岩浩巳
3. ご来賓挨拶
4. 出席会員の確認
5. 議案
 - 第 1 号議案 2012 年度事業報告（案） 第 2 号議案 2012 年度決算報告（案）
 - 第 3 号議案 2013 年度事業計画（案） 第 4 号議案 2013 年度収支予算（案）
 - 第 5 号議案 長尾賞の規約の改定について（案）について（案）
 - 第 6 号議案 長尾賞学生奨励賞の新設について（案）
 - その他・会員提案事項
6. 閉会の辞

報告会

- 開会挨拶 会長 中岩浩巳（NTT メディアインテリジェンス研究所）
1. 機械翻訳課題調査委員会 委員長 長瀬 友樹（(株)富士通研究所）
 2. AAMT/Japio 特許翻訳研究会 副委員長 梶 博行（静岡大学）
 3. インターネットワーキンググループ リーダー 富士 秀（(株)富士通研究所）
 4. 編集委員会 委員 小谷克則（関西外国語大学）

講演会

- ・言葉の壁を越えたコミュニケーションの実現に向けて
那須和徳氏（NTT ドコモ サービス&ソリューション開発部）
- ・ソフトウェアローカリゼーションにおける MT 活用と業界間断層
古田京子・菊池邦明氏（CA Technologies）

第8回 AAMT 長尾賞授与式・記念講演会

受賞者：株式会社バオバブ代表取締役社長 相良美織

受賞理由：

対訳コーパスを使う翻訳支援技術の下、「クラウドソーシング翻訳」のためのプラットフォームを構築し、その上に外国人留学生を中心とした翻訳コミュニティを形成し、廉価な人手翻訳を提供するビジネスを立ち上げ、ネット通販会社、IT 関連会社、大学などで利用実績を上げていることが認められることから、機械翻訳技術を活かしたインターネット時代の新たなMTビジネスを切り開いた点で長尾賞にふさわしい。

選考委員長：飯田 仁（東京工科大学）

選考委員：Key-Sun Choi（韓国 KAIST） 宇津呂 武仁（筑波大学）
永田昌明（日本電信電話株式会社）

推薦者：安達 久博（株式会社サン・フレア）

同意人：袋谷丈夫（株式会社ATR-Trek）



授賞式

懇親会

本会後の懇親会は、多数の参加者にお集りいただき、学識経験者、研究者、翻訳家等、幅広い分野の方々が活発な意見交換を行い有意義な交流の場となりました。ご参加誠に有難うございました。

協会活動報告

(2013 年 4 月～2013 年 8 月)

第 23 回通常総会

2013 年 6 月 19 日

- | | | | |
|------------|---------------------|---------|-----------------|
| 第 1 号議案 | 2012 年度事業報告 (案) | 第 2 号議案 | 2012 年度決算報告 (案) |
| 第 3 号議案 | 2013 年度事業計画 (案) | 第 4 号議案 | 2013 年度収支予算 (案) |
| 第 5 号議案 | 長尾賞の規約の改定について (案) | | |
| 第 6 号議案 | 長尾賞学生奨励賞の新設について (案) | | |
| その他・会員提案事項 | | | |

報告会

2013 年 6 月 19 日

- | | |
|--------------|---------------------|
| ①機械翻訳課題調査委員会 | ②AAMT/Japio 特許翻訳研究会 |
| ③インターネット WG | ④編集委員会 |

講演会

○講演：

- ・言葉の壁を越えたコミュニケーションの実現に向けて
那須和徳氏 (NTT ドコモ サービス&ソリューション開発部)
- ・ソフトウェアローカリゼーションにおける MT 活用と業界間断層
古田京子・菊池邦明氏 (CA Technologies)

第 8 回 AAMT 長尾賞受賞式・記念講演会

受賞者：株式会社バオバブ 代表取締役社長 相良美織

受賞理由：

対訳コーパスを使う翻訳支援技術の下、「クラウドソーシング翻訳」のためのプラットフォームを構築し、その上に外国人留学生を中心とした翻訳コミュニティを形成し、廉価な人手翻訳を提供するビジネスを立ち上げ、ネット通販会社、IT 関連会社、大学などで利用実績を上げていることが認められることから、機械翻訳技術を活かしたインターネット時代の新たな MT ビジネスを切り開いた点で長尾賞にふさわしい。

懇親会

2013 年 6 月 19 日 ホテルアジュール竹芝 21 階フレンチレストラン『ベイサイド』

決算理事会

2013 年 6 月 19 日

- | | | | |
|------------|--------------------|---------|----------------|
| 第 1 号議案 | 2012 年度事業報告（案） | 第 2 号議案 | 2012 年度決算報告（案） |
| 第 3 号議案 | 2013 年度事業計画（案） | 第 4 号議案 | 2013 年度収支予算（案） |
| 第 5 号議案 | 長尾賞の規約の改定について（案） | | |
| 第 6 号議案 | 長尾賞学生奨励賞の新設について（案） | | |
| その他・会員提案事項 | | | |

機械翻訳課題調査委員会

2013 年 4 月 12 日（2013 年度 第 1 回）

- ① 前回委員会の議事録の確認
- ② 各 WG の活動について（各 WG に分かれて議論）
- ③ 活動内容の報告（各 WG から）
- ④ 活動内容についての議論
- ⑤ まとめと次回委員会について

2013 年 5 月 17 日（2013 年度 第 2 回）

- ① 前回委員会の議事録の確認
- ② 各 WG の活動について（各 WG に分かれて議論）
- ③ 活動内容の報告（各 WG から）
- ④ 活動内容についての議論
- ⑤ まとめと次回委員会について

2013 年 6 月 28 日（2013 年度 第 3 回）

- ① 前回委員会の議事録の確認
- ② 各 WG の活動について（各 WG に分かれて議論）
- ③ 活動内容の報告（各 WG から）
- ④ 活動内容についての議論
- ⑤ まとめと次回委員会について

2013 年 8 月 2 日（2013 年度 第 4 回）

- ① 前回委員会の議事録の確認
- ② 各 WG の活動について（各 WG に分かれて議論）
- ③ 活動内容の報告（各 WG から）
- ④ 活動内容についての議論
- ⑤ まとめと次回委員会について

インターネット WG

- ① AAMT Forum メールマガジンのコンテンツ作成・配信業務の、課題調査委員会への移管作業
- ② 新 AAMT ホームページ公開に向けた更新体制の立ち上げ
- ③ 他委員会の AAMT ホームページ利用に関するサポート
- ④ AAMT ホームページの更新
- ⑤ その他、事務局ネットワークのインフラ管理全般

編集委員会

2013 年 6 月 25 日（2013 年度 第 1 回）

- ① AAMT Journal No.54 の記事の執筆依頼計画について
- ② AAMT Journal No.54 のスケジュールについて

AAMT/Japio 特許翻訳研究会

AAMT/Japio 特許翻訳研究会 拡大評価部会

2013 年 5 月 10 日（金）（2013 年度 第 1 回）

1. 新年度のご挨拶・紹介（委員長、委員、Japio）
2. 新事務局からのご挨拶
3. 平成 24 年度納品報告（事務局）
4. 報告書寄稿内容説明（各著者）
5. 拡大評価部会関連（江原先生）
6. Workshop 関連および MT Summit 関連
7. AAMT 総会関連
8. 次回の開催について
 - 開催の日時（場所）
 - 主な議題

2013 年 6 月 7 日（金）（2013 年度 第 2 回）

1. 前回議事録の確認
2. WS について（JPO, EPO, KIPO, KIPI などの状況）
3. 投稿状況（綱川先生）
4. AAMT 総会関連

5. 次回の開催について

- 開催の日時案（7月26日（金曜日）18：00～20：00）
- 主な議題

2013年7月19日（金）（2013年度 第3回）

1. 前回議事録の確認
2. 各機関／委員の25年度研究計画（各委員）
3. 拡大評価部会活動計画（江原委員）
4. 第5回特許翻訳ワークショッププログラム（横山副委員長）
5. NTCIR10 Patent Translation Task 報告（後藤委員）
6. 「分野別ユーザ辞書と市販ルールベース機械翻訳」について（Japio）
7. 事務局から銀行口座開設のお願い
8. 次回の開催について
 - 開催の日時案（10月4日（金曜日）18：00～20：00）
 - 主な議題

AAMT/Japio 特許翻訳研究会 拡大評価部会

2013年5月10日（金）（2013年度 第1回）

1. 新年度のご挨拶・紹介（委員長、委員、Japio）
2. 新事務局からのご挨拶
3. 自動評価
4. 人手評価
5. テストセット
6. その他

- 5th Workshop on Patent Translation（研究会主催）

2013年9月2日(月) Acropolis Conference Centre, Nice, France

（MT Summit XIV 併設 Workshop として開催）

AAMT ジャーナル編集委員会委員長
筑波大学 システム情報系 知能機能工学域
宇津呂 武仁

AAMT ジャーナル 54 号をお送りします。

今号の巻頭言は、日本翻訳連盟(JTF) 東郁男 代表理事・会長より、御寄稿を頂きました。

また、今号が会員の皆様のお手元に届くその少し後、来る 11 月 27 日(水)に、AAMT も後援をします第 23 回 JTF 翻訳祭が開催されます。日本翻訳連盟(JTF) 翻訳祭企画実行委員会の担当の方からは、この第 23 回 JTF 翻訳祭のお知らせをご寄稿いただきました。

一方、AAMT の活動としては、今号に先立ちまして、2013 年 6 月に開催されました総会におきまして、NTT ドコモの那須和徳氏ならびに CA Technologies の菊池邦明氏、古田京子氏より貴重な御講演を賜りましたが、今号におきましては、御講演内容についての貴重な御寄稿を頂きました。

あわせて、第 8 回の AAMT 長尾賞の選考結果を受けまして、受賞者である株式会社バオバブの相良美織社長より、受賞理由となった「Crowd and Computer Translation Collaboration at Baobab」についての紹介記事をご寄稿していただきました。

また、今年の総会におきまして、AAMT 長尾賞学生奨励賞の新設が承認されたことを受けまして、中岩浩巳会長より同賞の紹介記事をご寄稿いただきました。

プロジェクト報告として、NTCIR-10 特許翻訳タスクの国内主催者である NICT の後藤功雄氏、隅田英一郎氏、ならびに、海外からの主催者メンバーの方々より、NTCIR-10 特許翻訳タスクにおける評価結果概要についてご寄稿頂くとともに、海外参加チームから 1 件のご寄稿をいただきました。また、山形大学横山晶一先生、静岡大学梶博行先生、愛媛大学二宮崇先生より、第 14 回機械翻訳サミットおよび第 5 回特許翻訳ワークショップへの参加報告をご寄稿いただきました。その他、海外からのレポートを 1 件、国内からの研究報告、レポートを計 4 件、それぞれご寄稿いただきました。

AAMT 内の活動報告として、機械翻訳課題調査委員会から、「AAMT 機械翻訳文テストセットに基づく日中機械翻訳の自動評価サイト」の作成報告をご寄稿頂きました。

その他、「AAMT 会員のひろば」の企画におきましては、法人会員の紹介文 1 件を掲載しました。

AAMT

Asia-Pacific Association for Machine Translation

AAMT 入会のご案内

AAMT は、機械翻訳の発展を目的として、機械翻訳の研究者、開発者、製造者、利用者が集まった任意の組織です。委員会による定期的な調査研究をはじめ、機関誌の発行、シンポジウム、セミナー等各イベントの開催など幅広く活動を行っています。

機械翻訳にご関心のあるすべての方にご入会をお勧めします。

* * AAMT 会員の特典 * *

1.AAMT Journal の購読ができます。

会員には、機関誌である AAMT Journal（年 2～3 回発刊予定）が送付されます。購読料は年会費に含まれています。

2.機械翻訳関連の最新情報をメールでお届け

会員専用メーリングリストで、最新の機械翻訳関連の情報をお届けします。

MT 新製品、新サービスの紹介、国際会議、シンポジウムのお知らせ、WEB での MT 関連記事の紹介など盛りだくさんです。

3. AAMT が組織する委員会や調査活動に参加し、機械翻訳や翻訳に関心のある方との交流を深め、知見を広めることができます。

機械翻訳に関する言語資料の調査、広報、標準化活動に参加したり、AAMT Journal や会員専用メーリングリストで、自社製品、サービスの紹介を行うことができます。

4.関連機関の主催する国際会議に参加できます。

IAMT の主催で隔年開催される MT Summitをはじめ、AAMT、AMTA*、EAMT**の主催する会議やワークショップに参加できます。

AMTA* : Association for Machine Translation in the Americas

EAMT** : European Association for Machine Translation

年会費は以下の通りです。

法人会員：入会金 1 口 10,000 円 年会費 1 口 50,000 円

個人会員：入会金 1,000 円 年会費 5,000 円（学生は学生会費 1,000 円）

ご関心のある方は、事務局までお問い合わせください。

アジア太平洋機械翻訳協会（AAMT）

ホームページ：<http://www.aamt.info>

電子メール：aamt-info@aamt.info

AAMT

Asia-Pacific Association for Machine Translation

入会申込書

以下の通り、アジア太平洋機械翻訳協会の会員申し込みを致します。

申込日 20 年 月 日

氏名（ローマ字）	
氏名（漢字）	
電話番号	
メールアドレス	
所属先	
所属先住所	〒
種別	☐ ユーザ ☐ 研究開発者 ☐ その他
機械翻訳に関するお知らせメールの配信	☐ 希望する ☐ 希望しない

コメント

--

AAMT



AAMTジャーナル No.54

発行：アジア太平洋機械翻訳協会（AAMT）

ホームページ：<http://www.aamt.info>

住所：〒171-0014 東京都豊島区池袋2-55-2鈴木ビル3階

（株）日本システムアプリケーション内

phone：03-5951-3961 fax：03-5951-3966

編集委員会：宇津呂 武仁 小谷 克則 大倉 清司

鈴木 博和 阿部 さつき

表紙(図部分)デザイン：阿部 さつき

事務局：神崎 享子 荻野 孝野

印刷所：株式会社ユリクリエイト

Asia-Pacific Association for Machine Translation (AAMT)

c/o Japan System Application Co., Ltd.

Suzuki Building 3F 2-55-2, Ikebukuro, Toshima-ku Tokyo 171-0014, JAPAN