

文章编号: 1003-0077(2022)01-0047-09

“细粒度英汉机器翻译错误分析语料库”的构建与思考

裴白莲^{1,2}, 王明文¹, 李茂西¹, 陈聪¹, 徐凡¹

(1. 江西师范大学 计算机信息工程学院, 江西 南昌 330022;

2. 华东交通大学 外国语学院, 江西 南昌 330013)

摘要: 机器翻译错误分析旨在找出机器译文中存在的错误, 包括错误类型、错误分布等, 它在机器翻译研究和应用中发挥着重要作用。该文将人工译后编辑与错误分析结合起来, 对译后编辑操作进行错误标注, 采用自动标注和人工标注相结合的方法, 构建了一个细粒度英汉机器翻译错误分析语料库, 其中每一个标注样本包括源语言句子、机器译文、人工参考译文、译后编辑译文、词错误率和错误类型标注; 标注的错误类型包括增词、漏词、错词、词序错误、未译和命名实体翻译错误等。标注的一致性检验表明了标注的有效性; 对标注语料的统计分析结果能有效地指导机器翻译系统的开发和人工译员的后编辑。

关键词: 机器翻译; 错误分析; 错误标注; 译后编辑

中图分类号: TP391

文献标识码: A

Construction of Fine-Grained Error Analysis Corpus of English-Chinese Machine Translation

QIU Bailian^{1,2}, WANG Mingwen¹, LI Maoxi¹, CHEN Cong¹, XU Fan¹

(1. School of Computer and Information Engineering, Jiangxi Normal University, Nanchang, Jiangxi 330022, China;

2. School of Foreign Languages, East China Jiaotong University, Nanchang, Jiangxi 330013, China)

Abstract: Machine translation error analysis, including error classes and error distribution etc. Error analysis of machine translation output, plays an important role in the research and application of machine translation. In this paper, post-editing is introduced into error analysis to annotate error labels. Automatic error annotation and manual annotation are applied to build a Fine-grained Error Analysis Corpus of English-Chinese Machine Translation (ErrAC), in which every annotated sample includes a source sentence, MT output, reference, post-edit, WER and error type. The annotated error types include addition, omission, lexical error, word order error, untranslated word, named entity translation error etc. Annotator agreement analysis shows the effectiveness of the annotation. The statistics and analysis based on the corpus provide effective guidance for the development of machine translation system and post-editing practice.

Keywords: machine translation; error analysis; error annotation; post-editing

0 引言

机器翻译质量评价是机器翻译研究的重要内容。机器翻译质量评价主要有人工评价和自动评价两种方式。由于人工评价成本较高, 周期较长, 不易获得, 目前机器翻译质量评价大多采用自动评价指标, 如 BLEU^[1], METEOR^[2] 和 TER^[3] 等。这些自动评价指标依据参考译文对机器译文给出整体得

分, 能够反映机器翻译质量整体情况, 但是无法反映机器译文具体存在哪些问题, 需要在哪些方面进行改进。为获取存在问题的具体信息, 就需要进行机器翻译错误分析。错误分析可以找出机器译文中具体存在的问题, 有助于了解机器翻译系统的不足, 找准改进的方向, 还可以为机器翻译质量估计、错误预测、自动译后编辑提供参考。近十几年来, 错误分析在国外机器翻译研究领域受到重视, 出现很多相关的研究, 例如, 使用错误分析评价机器翻译质量^[4]、

收稿日期: 2021-02-4 定稿日期: 2021-03-29

基金项目: 国家自然科学基金(61876074, 61662031, 61772246); 国家社会科学基金(19BYY121); 教育部人文社科基金(21YJC740040)

分析机器翻译的错误类型^[5]、通过错误分析对 NMT 和 PBMT 进行细粒度的人工评价^[6]。但国内相关研究还较少,仅有一些针对机器译文错误进行的语言学分析。例如,罗季美和李梅^[7]将机器译文错误分为词汇错译、句法错译、符号错译三大类并展开分析,或从短语和句子层面分析机器翻译的句法错误^[8]。这些研究仅使用独立的人工译文与机器译文做对比展开分析,而且针对的是传统的机器翻译系统如 RBMT,其错误分析的结果已不能反映当前机器翻译的水平。孙逸群^[9]对 5 篇海洋类论文摘要机辅翻译中的错误进行了剖析,其错误分析侧重实例分析和改错,而且语料规模小,不具代表性。据我们了解,目前还没有专门针对英汉机器翻译错误分析可公开获得的语料库。值得注意的是,随着神经机器翻译的发展,机器翻译质量极大提高,但是英汉翻译方向神经机器翻译质量究竟如何,还存在哪些具体问题,针对这些问题还鲜有专门的错误分析,本文尝试针对这些问题展开研究与探讨。

错误分析和译后编辑是高度相关的工作,错误分析是找出机器译文的错误,译后编辑是改正机器译文的错误。错误分析和译后编辑都可以用来评价机器翻译的质量,但以往的研究大多把错误分析和译后编辑单独使用或单独作为研究对象,较少有把两者结合起来的研究。我们将译后编辑和错误分析结合起来,先对机器译文进行译后编辑,然后以译后编辑译文(PE 译文)作为参照,对机器译文进行错误标注。在此基础上,构建了一个细粒度英汉机器翻译错误分析语料库(Fine-grained Error Analysis Corpus of English-Chinese Machine Translation, ErrAC)。PE 译文比参考译文更适合作为错误标注参照的原因在于,翻译本来就存在一文多译的现象,同一个源语言句子可以有多种不同的正确译文,而在机器译文的基础上进行译后编辑,力求 PE 译文是最接近机器译文的正确译文,其编辑距离最短。因此,以 PE 译文来衡量机器翻译的质量相对而言更客观,更能准确地找出机器翻译真正存在的问题。文献^[3]表明,使用人工译后编辑译文得到的 HTER 值,相比最接近机器译文的参考译文的 TER,更能准确地衡量机器翻译的质量,而且,HTER 与人工评价的相关性比 BLEU 与人工评价的相关性更高。下面给出了 WMT19 新闻机器翻译测试集上的两个实例,它们表明以人工参考译文和 PE 译文作为错误分析参照的区别。

例 1

源语言句子: It would be extremely ill advised to venture out into the desert on foot with the threat of tropical rainfall.

机器译文: 在热带降雨的威胁下,徒步冒险进入沙漠是极不明智的。

PE 译文: 在热带降雨的威胁下,徒步冒险进入沙漠是极不明智的。

参考译文: 由于热带降雨的威胁,沙漠冒险活动将十分危险。

PE 译文 WER 0.00 参考译文 WER 76.92

例 2

源语言句子: Do you think he's telling the truth to the country?

机器译文: 你认为他对国家说的是真话吗?

PE 译文: 你认为他对国人说的是真话吗?

参考译文: 你觉得他对国人所说的是事实吗?

PE 译文 WER 8.33 参考译文 WER 35.71

从例 1 可见,机器译文是正确的译文,达到了对翻译的忠实、通顺的要求,但是与参考译文有很大的差别。如果按照参考译文来标注错误,那么会得出这一机器译文质量低劣的结果,其 WER 值^①高达 76.92,这样显然无法准确、有效地衡量机器译文质量。例 2 中,译后编辑实际上只需要一次替换的编辑操作,即修改一处错词,就可以达到忠实、通顺的要求,PE 译文 WER 为 8.33,但是机器译文与参考译文的差别较大,WER 为 35.71,把机器译文修改成参考译文需要三次替换操作和一次插入操作。由此可见,使用 PE 译文作为参照对机器译文进行错误标注,比直接使用参考译文更客观,更能准确地反映机器翻译系统的问题。

本文工作的意义体现在以下三个方面: 1) 获得对神经机器翻译质量更客观、更准确的评价; 2) 为机器翻译系统开发、译后编辑工作提供参考; 3) 可以为机器翻译质量估计、错误预测、自动译后编辑提供数据和参考。

本文组织结构如下: 第 1 节介绍错误分析和译后编辑相关研究和相关语料库建设情况; 第 2 节介绍语料来源和语料库构建过程; 第 3 节对错误标注结果进行统计与分析; 第 4 节总结全文。

1 相关工作

错误分析可以以人工和自动两种方式进行。

^① Word Error Rate, 词错误率。

Vilar 等人^[10]建立了人工错误分析的框架,定义了错误类型,根据错误分类对机器译文进行错误标注。Popovic 等人^[11]提出基于屈折变化和句法信息的自动错误分析框架,自动获得错误的细节信息。机器翻译错误分析主要有以下几种应用:第一,用于评价某一机器翻译系统的质量^[12],或比较几种不同的机器翻译系统,通常是比较 SMT 和 NMT 等不同系统^[13-15];第二,考察不同错误类型对机器翻译质量的影响^[16-17];第三,用于译后编辑的相关研究,考察不同错误类型对译后编辑工作量不同方面的影响^[18-19]。但是,这些研究或是在机器译文上进行错误分析,或以参考译文为参照进行错误分析,而不是以 PE 译文为参考,这会导致错误分析与实际情况存在偏差。

随着译后编辑在翻译行业越来越普遍,逐渐出现了一些可公开获得的译后编辑语料^[20]。WMT 从 2012 年开始质量估计子任务,从 2015 年开始自动译后编辑子任务,这两个子任务都提供了译后编辑译文语料。部分语料还有错误标注,包括基本的编辑距离操作,如替换、删除、插入和移位,或者“好”“差”二元标签,其语言对涉及英德、英俄等。CWMT 从 2018 年开始翻译质量估计任务,提供英汉语言对机器翻译译后编辑译文,部分语料有“好”“差”二元标签,部分语料有每个句子的 HTER 值。这些语料对机器翻译质量估计、错误预测、自动译后编辑、译后编辑人员培训都非常有用。

同时,还出现了一些做了错误标注的译后编辑语料库。例如,TRACE 语料库^[21]包含法英、英法译后编辑译文,其中有基本编辑距离错误类型的标注。Koponen^[22]使用英西机器翻译语料,提供译后编辑译文,对语料进行错误标注,研究错误类型与估计的译后编辑工作量、实际编辑操作之间的关系,但是其语料不能公开获得。

Terra 语料库^[23]是可以公开获得的人工错误标注语料库,用于自动错误分类工具 Addicter^[24]和 Hjerson^[25]的评估。这个语料库由不同研究小组独立标注,标注策略各不相同,有的小组不使用参考译文,有的小组使用参考译文,这样会导致错误标注一致性不高,因为标注策略不同,标注的结果会有较大差异。而且,这项工作中的人工错误分类和自动错误分类是完全独立进行的。TARAXü 语料库^[26]也能够公开获得,该语料库包含译后编辑译文和机器翻译错误标注数据,但是这两项工作是完全独立进行的,并且不是在同一个数据集上进行的。PE2rr

语料库^[27]在译后编辑译文的基础上进行错误标注,更准确地反映了机器译文的错误情况,可以公开获得,但是该语料库只包含英语、塞尔维亚语、德语、西班牙语语料。这些语言均属于印欧语系,语言之间的差别相对较小,其错误分析的数据可能无法一般化。英语和汉语分属于不同的语系,差别较大。一些印欧语系语言之间机器翻译常见的错误,如屈折错误,在英汉语言方向上并不存在,而有的错误类型则可能比较突出,那么英汉机器翻译与其它相同语系语言之间机器翻译的错误情况、错误分布有没有差异、有什么差异,就需要专门进行英汉机器翻译错误分析,而目前英汉语言对机器翻译质量评价和错误分析还缺少类似的语料库。

2 语料库构建

本节介绍语料来源和语料库构建过程。语料库构建过程分为两个阶段:译后编辑和错误标注。首先由专业人士进行译后编辑,然后采用自动错误标注加人工标注的方式进行错误标注。

2.1 语料来源

我们的语料来源为 WMT2019 新闻机器翻译测试集英中翻译方向,该测试集包括源语言句子、机器译文和人工翻译的参考译文。我们将测试集按照新闻内容分为六类:政治、经济、社会、体育、科教和文艺。我们使用的机器译文是 KSAI 组(金山 AI)提交的机器译文,该小组在英中机器翻译任务中人工打分排名第一。KSAI 组提交的机器译文是基于各种神经机器翻译模型,以 Transformer 作为基线系统,使用了几种数据过滤和回译作为数据清洁和数据增强的方法。最终模型是经过多模型集成、重排序、后处理的系统组合^[28]。语料库的统计信息见表 1,句子数为 1 997,源语言句子词数为 42 034,机器译文词数为 51 693,编辑词数为 8 380。编辑词数百分比是按照编辑词数与机器译文词数加漏词数量的百分比来计算的。

表 1 源语言句子词数、机器译文词数及编辑词数

新闻类别	句子数	源语言句子词数	机器译文词数	编辑词数 (编辑词数百分比/%)
政治	700	15 425	18 622	2 365(12.5)
经济	112	2 473	3 019	473(15.4)
社会	494	9 293	10 958	1 532(13.7)

续表

新闻类别	句子数	源语言句子词数	机器译文词数	编辑词数 (编辑词数百分比/%)
体育	305	6 269	8 154	2 144(25.7)
科教	174	3 791	4 670	651(13.6)
文艺	212	4 783	6 270	1 215(18.9)
总数	1 997	42 034	51 693	8 380(16.2)

2.2 译后编辑

进行本次译后编辑工作的译后编辑人员为两名翻译专业教师,均精通英汉两种语言,具有丰富的翻译和译审经验。为保证译后编辑质量,在进行译后编辑之前,译后编辑人员经过多次讨论和修改,知晓此次译后编辑的目标和原则。译后编辑的目标是修改机器译文的错误,使译文达到忠实源语言句子、语句通顺的要求,质量适中即可。本次译后编辑采取轻度译后编辑的原则,即只进行最少量的必要的编辑操作以达到译文质量可接受的效果,不考虑风格、文采问题,也不考虑译后编辑人员在用词习惯、语法结构等方面的个人喜好问题。针对本次译后编辑制定五条具体指南:(1)力求译文意思正确、语句通顺;(2)确保没有信息增加或遗漏;(3)尽可能多地使用机器译文;(4)除非影响语义,否则不修改句子结构;(5)单纯的风格问题无需修改。

表 1 表明了语料库句子数、源语言句子词数、机器译文词数,以及机器译文经过译后编辑的编辑词数。整体编辑词数百分比为 16.2%,需要进行编辑修改的比例不是很大,这表明在新闻翻译对质量要求适中的应用场景中,在领域语料比较丰富的基础上,神经机器翻译质量可接受程度高,机器译文在很大程度上可用。在各种新闻类别中,编辑词数百分比最大的是体育新闻,达到 25.7%,是出现错误最多、需要译后编辑量最大的新闻类别。而编辑词数百分比最小的是政治新闻,为 12.5%,是需要译后编辑量最小的新闻类别。可见,不同新闻类别之间机器翻译质量的差别较大,其可能的原因在于相关领域训练语料规模的大小。

我们以 WER 表示机器译文每个句子的编辑距离,按照编辑距离的大小将所需的译后编辑工作量(体现为所进行的实际编辑操作)分为四个等级,结果见表 2。从表 2 可以看出,有 26%的句子已经可接受,无需任何编辑操作,47.8%的句子只需要少量编辑操作即可达到质量适中的要求,17.2%的句子需要中等编辑工作量,只有 9%的句子需要进行大

量修改。ErrAC 语料库中也给出了每个句子的 WER 值。在不同新闻类别中,体育类需要的译后编辑工作量相对较高,不需要编辑操作的比例为 14.4%,低于其他所有类别,而需要进行大量译后编辑操作的比例达到 21%。

表 2 译后编辑工作量等级分布

新闻类别	译后编辑操作(以 WER 作为编辑距离)			
	无 (0%)	低 (0%~25%)	中 (25%~50%)	高 (>50%)
政治	29.7	51.9	12.1	6.3
经济	17.0	53.6	20.5	8.9
社会	33.8	44.1	16.0	6.1
体育	14.4	37.1	27.5	21.0
科教	23.6	57.5	16.1	2.9
文艺	19.3	47.2	20.8	12.7
总数	26.0	47.8	17.2	9.0

2.3 错误标注

错误标注工作分两个阶段进行。首先,以 PE 译文为参考,使用 Hjerson 自动错误标注工具进行错误标注;然后,将自动错误标注的结果一一进行人工核对和修改,并细化和扩展错误类型。

进行错误标注之前先对中文语料进行预处理,采用清华 THULAC 分词工具进行分词。Hjerson 工具以机器译文和 PE 译文作为输入,以词为单位进行错误标注,输出错误标注结果。Hjerson 工具可以识别和标注五种类型的错误,即增词、漏词、错词、词序错误和屈折错误(动词时态/人称/情态/格/性/数)。Hjerson 工具主要是针对英语、德语等印欧语系语言开发的,其中的屈折错误常出现于印欧语系语言之间的翻译中,而英汉机器翻译的目的语为汉语,汉语不是屈折语言,没有屈折错误,因此 Hjerson 实际上标注出来的错误有四种,即增词、漏词、错词和词序错误,在 ErrAC 语料库中分别以 ext、miss、lex 和 reord 表示。除漏词错误,所有其他错误均针对机器译文做标注。其中,在机器译文中出现了而在 PE 译文中没有出现的词标注为增词。在 PE 译文中出现了而在机器译文中没有出现的词标注为漏词。漏词错误需要针对 PE 译文做标注,因为漏词是机器译文中没有的词,无法在机器译文的标注中体现,在 PE 译文上做标注,才能体现漏词错误及漏词的位置。

自动标注之后进行人工标注,标注者为本文作

者之一,知晓标注规则和方法。在人工标注阶段,除核对和修改自动错误标注,还对错误类型进行了细化和扩展。细化针对增词错误类型,细化的标注有两种,一是数词加量词,二是人称代词加结构助词。由于英汉语言习惯的差别,这两种增词错误是英汉翻译中经常出现的问题,在机器翻译中更为明显。英文中的冠词 a 或 an,在机器翻译中常被译为一个、一种、一名等,而很多情况下按照汉语的习惯用法这些是应该省略的,如例 3 所示。数词和量词的增词分别标注为 ext-num 和 ext-cla,其出现次数分别为 81 次和 83 次,占增词总数的 4.76%和 4.87%。

例 3

源语言句子: Thomas Bjorn, the European captain, knows from experience that a sizeable lead heading into the last-day singles in the Ryder Cup can easily turn into an uncomfortable ride.
机器译文: 欧洲 队长 托马斯 · 比约恩 (Thomas Bjorn) 从 经验 中 知道, 在 莱德杯 最后 一天 的 单打 比赛 中, 一个 相当 大 的 领先 优势 很 容易 演 变 成 一 场 不 舒 服 的 比 赛。
PE 译文: 欧洲 队长 托马斯 · 比约恩 (Thomas Bjorn) 根据 经验 知 道, 在 莱德杯 最后 一 天 的 单 打 比 赛 中, 大 比 分 的 领 先 优 势 也 很 容 易 变 成 不 利 局 面。
机器译文标注: x x x x x x x x lex x lex x x x x x x x x x x ext-num ext-cla lex lex x x x x x lex x ext-num ext-cla lex lex lex lex x
PE 译文标注: x x x x x x x x lex x x x x x x x x x x x x x lex x x x miss x x lex x lex lex x

此外,英语中的人称代词 we/he/she/they 等及其相应物主代词 our/his/her/their 等,在机器翻译中基本都按原本译出,但是根据汉语使用习惯,很多时候在译文中都应该省略,否则译文不自然、不通顺,如例 4 所示。人称代词和结构助词增词分别标注为 ext-pro 和 ext-aux,分别出现 114 次和 84 次,分别占增词总数的 6.69%和 4.93%。

例 4

源语言句子: We've transformed the look and feel of our beauty aisles to enhance the environment for our customers.
机器译文: 我们 已经 改变 了 我们 美容 通道 的 外 观 和 感 觉, 为 我 们 的 客 户 改 善 了 环 境。
PE 译文: 我们 已经 改变 了 美容 通道 的 外 观 和 氛 围, 为 客 户 改 善 环 境。
机器译文标注: x x x x x ext-pro x x x x x lex x x ext-pro ext-aux x x ext x x
PE 译文标注: x x x x x x x x x lex x x x x x x

人工标注阶段扩展的三种错误类型为未译、命名实体翻译错误和标点符号错误。机器译文中出现了一些未经翻译的英文单词,标注为 untr。机器译文中还出现了一些命名实体翻译错误或命名实体翻译前后不一致的问题,包括人名、地名、组织机构名称等。未译的大多都是命名实体,但因为错误形式不同,所以做了区分。命名实体翻译错误标注为 nen。此外,还有标点符号错误、多余或遗漏的问题,这类问题全部归类为标点符号错误,标注为 punc。

除了细化和扩展错误类型,在人工标注阶段还进行了多标签错误标注。因为有的词存在多种错误,如错词、未译、命名实体翻译错误也可能出现在错误的位置上,即同时也是词序错误。这种情况自动错误标注工具无法标注,在人工阶段做了补充,针对叠加的词序错误标注了多错误标签,在语料库中表示为+reord。

错误标注完成之后,为检验标注质量,我们进行了标注者一致性分析。采用取样的方法,取数据集中前 100 个句子,分别由 A1 和 A2 两位标注者独立进行标注,两位标注者均经过培训,知晓标注规则和方法。错误标注不是简单地打分或排序,它涉及所标注的错误数量、错误类型和标注的位置,标注者一致性不容易计算。我们采用关于错误标注不同标注者一致性的计算方法^[29],该计算方法关注所标注错误的共现情况,如式(1)所示。

Agreement = (2 * A^agree) / (A1^all + A2^all) (1)

其中,上标 all 表示每位标注者标注的总数,上标 agree 表示两位标注者标注错误类型相同的数量。不同标注者一致性详见表 3,整体一致性达 90.6%。可见,在自动标注工具的基础上进行人工修改,不仅提高了错误标注效率,也有助于提高标注者一致性。

表 3 不同标注者一致性 (单位: %)

错词	88.8	未译	100
增词	74.9	命名实体	100
漏词	97.6	标点符号	100
词序	99.5	总数	90.6

该计算方法关注所标注错误的类型和数量,没有考虑标注错误的位置。在 ErrAC 语料库中,我们经过观察发现,不同标注者出现标注位置不一致的主要是词序错误,即 reord 的标注位置会有差异,其他错误类型的标注位置基本上差异不大。在各种错

误类型中,增词的标注者一致性相对较低,这是因为在英汉翻译中,词与词并不是一一对应的,词一对多、多对一的情况很常见,会造成标注者对于某个词是属于增词还是错词的标注产生差异。例如,源语言句子中“holiday homes”,机器译文为“度假之家”,PE 译文为“度假屋”,标注者 A1 标注为“lex lex lex”,标注者 A2 标注为“lex ext lex”。两者对“之”字的错误类型标注不一致,分歧的原因在于标注者 A1 将“度假之家”三个词理解为对应源语言句子“holiday homes”两个词,而标注者 A2 的理解是“度假”对应源语言句子“holiday”,“家”对应源语言句子“home”,那么“之”就理解为增词。

采用同样的计算方法,我们还计算了同一标注者一致性。在标注者 A1 完成第一次标注之后,间隔两

个月的时间,随机取数据集中 100 个句子再次进行标注。经过计算得出,同一标注者一致性为 93.6%。

3 统计与分析

3.1 结果统计

我们对错误标注结果做了统计,每种错误类型的数量和错误率见表 4 和表 5。错误率是错误数量与文本总词数的百分比,这样方便对不同的机器译文进行错误分析时相互比较。从表 4 可见,数量最多的错误类型是错词,即在机器翻译中选择了错误的词汇进行翻译,错词数量为 3 861,约占编辑词数的 46%。其次是增词,数量为 1 703,约占编辑词数的 20%。词序错误和漏词分别约占 16%和 13%。

表 4 错误类型数量

新闻类别	各种错误类型数量									
	增词	错词	+词序错误	漏词	词序错误	未译	+词序错误	命名实体	+词序错误	标点符号
政治	500	1087	+62	326	414	56	+0	38	+1	1
经济	104	226	+16	53	56	47	+2	3	+0	7
社会	287	662	+38	237	241	98	+0	42	+0	9
体育	425	1131	+64	202	303	25	+3	124	+5	21
科教	155	262	+16	101	96	40	+2	12	+1	6
文艺	232	493	+41	165	205	88	+3	65	+1	25
总数	1703	3861	+239	1084	1315	354	+10	284	+8	87

表 5 错误率

(单位: %)

新闻类别	各种错误类型错误率									
	增词	错词	+词序错误	漏词	词序错误	未译	+词序错误	命名实体	+词序错误	标点符号
政治	2.68	5.84	+0.33	1.75	2.22	0.30	+0.00	0.20	+0.01	0.01
经济	3.44	7.49	+0.53	1.76	1.85	1.56	+0.07	0.10	+0.00	0.23
社会	2.62	6.04	+0.35	2.16	2.20	0.89	+0.00	0.38	+0.00	0.08
体育	5.21	13.87	+0.78	2.48	3.72	0.31	+0.03	1.52	+0.06	0.26
科教	3.32	5.61	+0.34	2.16	2.06	0.86	+0.04	0.26	+0.02	0.13
文艺	3.70	7.86	+0.65	2.63	3.27	1.40	+0.05	1.04	+0.02	0.40
总数	3.29	7.47	+0.46	2.10	2.54	0.68	+0.02	0.55	+0.02	0.17

注: 错误率为错误数量与文本总词数的百分比。

3.2 问题与建议

错误分析对机器翻译系统开发具有很好的参考价值,其主要意义在于,有助于了解机器翻译系统存在的具体问题,了解系统的不足和短板,明确改进的方向,为机器翻译系统开发提供参考。我们对神经机器翻译译文进行错误分析,根据所发现的主要问

题,对机器翻译系统开发提出如下建议。

第一,针对一词多义问题。通过错误分析可知,错词问题是神经机器翻译的主要问题。机器译文中错词问题大多是因为源语言句子中一词多义,而目前的神经机器翻译技术没有对句子进行真正的理解,无法根据领域和上下文信息来选择正确的义项,导致翻译时选词错误。建议机器翻译系统开发时,

一方面,通过引入外部的领域知识库或知识图谱,充分利用外部知识;另一方面,通过大型单语语料库训练准确的语境词向量进行词义消歧,充分利用上下文信息,来缓解一词多义导致的错词问题。

第二,针对增词错误。在 ErrAC 语料库中,代词加结构助词、数词加量词这两种类型的增词占增词总数的 21.25%。在机器翻译系统开发时,可以考虑对这些词类的翻译设置一定的约束。此外,还需要提高训练语料的质量。如果训练语料在这些词类的翻译上处理得比较好,神经机器翻译在这方面也会有更好的表现。

第三,针对术语翻译错误。以体育类新闻为例,体育类新闻中错词的数量多达 1 131 处,占语料库中错词总数(3 861)的 29.3%,其错误率为 13.87%。原因在于,体育类新闻中很多词是专业术语,这些术语在译文中也需要对等地翻译成专业术语,而机器翻译往往把这些词按照常用义项译出,没有根据领域来选择合适的义项,导致翻译错误。例如,The attempt sailed high above the box,句中的“box”,机器翻译为“盒子”,而在足球术语中应为“禁区”。建议开发机器翻译系统时,引入相关领域的术语词典资源,并使系统在译文本输入时可以识别其所属领域,即时调用相关领域术语资源,以缓解术语翻译错误的问题。

第四,针对代词引起的翻译错误问题。机器翻译中由于对代词指代对象不明,导致出现翻译错误的情况很多,有时甚至引起整个句子的意思出现偏差。代词指代不明有多种原因,比如句中代词可指代的对象有多个,导致代词指代模糊;或者代词的指代对象距离代词很远,跨越了单个句子。目前神经机器翻译模型大多是句子级别的,无法很好地利用篇章上下文信息解决跨越句子的指代问题。建议开发和改善以段落、篇章为输入单元的翻译模型,开发基于篇章级别的神经机器翻译系统,这样的系统还可以获取句子之间的依赖关系,更连贯地翻译整个篇章文本。

第五,针对缺乏训练语料问题。领域相关语料稀缺会直接影响翻译质量,比如体育类新闻中命名实体翻译错误多达 124 处,占语料库中命名实体翻译错误总数(284)的 43.7%。原因在于,体育类新闻中人名、球队名、俱乐部名称等出现的频率比其他类新闻更高。在机器译文中,这些命名实体翻译出现译错以及翻译前后不一致的情况很多。这些命名实体不能正确翻译的直接原因是相关领域的训练语料较少。针对这一问题,一方面是尽可能增加语料的数量,扩大训练语料的覆盖度;另一方面是提高训练语料的质量。应当避免直接从网上爬取双语语料作

为训练语料,而要仔细甄别双语语料的质量,使用高质量的双语语料。获得大量高质量的双语语料对于提高神经机器翻译质量具有决定性作用。此外,针对命名实体翻译问题,建议在机器翻译系统中加入命名实体翻译检查机制,检查并改正命名实体翻译前后不一致的情况。

3.3 经验与建议

从 ErrAC 语料库的数据中可以总结出一些经验教训供译后编辑人员参考。

第一,关注一词多义引起的错词问题。在各种类型的错误中,错词数量最多,达到 3 861 次,可见一词多义仍然是机器翻译的一个障碍,目前神经机器翻译系统还无法根据领域和上下文选择正确的词义进行翻译。因此,在译后编辑过程中,需要关注一个词在不同领域、不同上下文中表达的不同意义,应关注词义选择的问题,以提高译后编辑的准确率和效率。

第二,善于发现和修改词序错误能有效提高译后编辑效率。词序错误占编辑词数的 16%。研究发现,词序错误是机器翻译使用者最不喜欢的错误类型^[30]。其原因可能在于词序错误更难发现和修改,特别是长距离词序错误。Popovic 等人^[31]发现,错词和词序错误所需要的认知努力最大。如果是错词叠加词序错误,需要的译后编辑认知努力更大,需要的译后编辑时间更多。因此,词序错误所需要的译后编辑工作量可能相对较大,在译后编辑过程中需要予以关注。译后编辑人员应该熟悉中英文在词序方面的差异,增强对翻译中词序问题的敏感性。

第三,在 ErrAC 语料库中,增词错误数量较多,但相对比较容易修改。删除增词的编辑操作所需要的译后编辑认知努力和时间最少^[31]。而且,关于增词错误,还可以关注代词加结构助词、数词加量词这样的增词,在本语料库中,这几种类型的增词占增词总数的 21.25%。这样有针对性地进行译后编辑,有助于提高译后编辑的速度和效率。

第四,具备全局意识,从篇章整体的角度修改错误。在机器译文中,经常出现命名实体翻译前后不一致的问题,影响篇章的连贯性,导致译文读者理解困难。虽然译后编辑人员在篇章全局的理解和把握上有优势,但有时容易忽略篇章信息,更多关注单个句子的细节。因此,在译后编辑过程中需要对该问题予以注意,修改译名不一致的问题,保证命名实体翻译前后一致,加强译文篇章的连贯性和可读性。

第五,适当关注标点符号,根据中文习惯来修改。在英汉翻译中,受英文句子结构的影响,机器译

文常出现中文长句。在译后编辑过程中,需要根据中文习惯合理断句,插入标点符号,尤其是逗号。在 ErrAC 语料库中,插入标点符号的译后编辑操作达 165 次,其中大多数是插入逗号。

最后,加强对机器翻译的了解。译后编辑人员除了需要具备扎实的双语能力和翻译能力,还需要对机器翻译有较好的了解。他们需要了解机器翻译系统的不足和问题,熟悉机器译文中常出现的错误,尝试摸索总结其错误模式,并掌握有针对性的纠错方法。只有在译后编辑实践中不断积累经验,才能不断提高译后编辑的质量和效率。译后编辑人员可以充分利用机器翻译提供的便利,同时发挥人工编辑的优势,促进人机融合翻译模式的发展。

4 总结

我们构建了一个可公开获得的细粒度英汉机器翻译错误分析语料库,语料库中每一个标注样本包括源语言句子、机器译文、参考译文、PE 译文、词错误率,以及基于 PE 译文所进行的错误标注。错误分析语料库可以准确、有效地评价机器翻译质量,获得关于机器译文错误类型、错误分布的数据,有助于了解目前神经机器翻译存在的具体问题,为机器翻译系统开发提供参考。我们将译后编辑与错误分析结合起来,对所进行的译后编辑操作进行错误标注,这比使用参考译文作为参照进行错误标注,更能准确地反映机器译文的具体问题,更符合人对机器译文错误的认知。错误分析对机器翻译系统的开发和译后编辑工作都有很好的参考作用,还可以为机器翻译质量估计、错误预测、自动译后编辑和译后编辑教学提供数据基础和参考作用。错误分析语料库中的数据可用于机器翻译质量估计实验,可用于对机器译文进行错误预测,结合错误分析的译后编辑数据在自动译后编辑研究中也非常有用。由于人工的限制,目前数据库规模还比较有限,而且只针对神经机器翻译做了错误分析,没有涉及 SMT 等其他系统的错误分析和相互比较。未来的工作除扩大语料库规模以涵盖更多领域和不同机器翻译系统的语料之外,还将基于该语料库构建初步的计算模型,用于机器翻译质量估计和自动译后编辑实验。

参考文献

[1] Papineni K, Roukos S, Ward T, et al. BLEU: a method for automatic evaluation of machine translation [C]//Proceedings of the 40th Annual Meeting of the

- Association for Computational Linguistics, 2002: 311-318.
- [2] Banerjee S, Lavie A. METEOR: an automatic metric for MT evaluation with improved correlation with human judgments [C]//Proceedings of the ACL Workshop on Intrinsic and Extrinsic Evaluation Measures for Machine Translation and/or Summarization, 2005: 65-72.
- [3] Snover M, Dorr B, Schwartz R, et al. A study of translation edit rate with targeted human annotation [C]//Proceedings of the 7th Conference of the Association for Machine Translation in the Americas, 2006: 223-231.
- [4] Koponen M. Assessing machine translation quality with error analysis [C]//Proceedings of the KåTu Symposium on Translation and Interpreting Studies, 2010: 1-12.
- [5] Bojar O. Analyzing error types in English-Czech machine translation [J]. The Prague Bulletin of Mathematical Linguistics, 2011, 95: 63-76.
- [6] Klubicka F, Toral A, Sánchez Cartagena VM. Fine-grained human evaluation of neural versus phrase-based machine translation [J]. The Prague Bulletin of Mathematical Linguistics, 2017, 108: 121-132.
- [7] 罗季美,李梅.机器翻译译文错误分析[J].中国翻译, 2012, 30(05): 84-89.
- [8] 罗季美.机器翻译句法错误分析[J].同济大学学报(社会科学版), 2014, 25(01): 111-118,124.
- [9] 孙逸群.海洋类论文摘要机辅翻译错误剖析[J].中国科技翻译, 2019, 32(02): 31-33,30.
- [10] Vilar D, Xu J, D'Haro L F, et al. Error analysis of statistical machine translation output [C]//Proceedings of the 5th International Conference on Language Resources and Evaluation, 2006: 697-702.
- [11] Popovic M, Ney H. Morpho-syntactic information for automatic error analysis of statistical machine translation output [C]//Proceedings of the Workshop on Statistical Machine Translation, 2006: 1-6.
- [12] Lommel A, Burchardt A, Popovic M. Using a new analytic measure for the annotation and analysis of MT errors on real data [C]//Proceedings of the 17th Annual Conference of the European Association for Machine Translation, 2014: 165-172.
- [13] Bentivogli L, Bisazza A, Cettolo M, et al. Neural versus phrased-based machine translation quality: a case study [C]//Proceedings of the Conference on Empirical Methods in Natural Language Processing, 2016: 257-267.
- [14] Toral A, Sánchez-Cartagena V M. A multifaceted evaluation of neural versus statistical machine translation for 9

- language directions [C]//Proceedings of the 15th Conference of the European Chapter of the Association for Computational Linguistics, 2017: 1063-1073.
- [15] Bentivogli L, Bisazza A, Cettolo M, et al. Neural versus phrase-based MT quality: an in-depth analysis on English-German and English-French [J]. Computer Speech and Language, 2018, 49: 52-70.
- [16] Popovic M, Avramidis E, Burchardt A, et al. Learning from human judgments of machine translation output [C]// Proceedings of the XIV MT Summit, 2013: 231-238.
- [17] Federico M, Negri M, Bentivogli L, et al. Assessing the impact of translation errors on machine translation quality with mixed-effects models [C]//Proceedings of the Conference on Empirical Methods in Natural Language Processing, 2014: 1643-1653.
- [18] Krings H. Repairing texts: Empirical investigations of machine translation post-editing process [M]. Kent, OH: The Kent State University Press, 2001.
- [19] Zaretskay A, Vela M, Pastor G C, et al. Measuring post-editing time and effort for different types of machine translation errors [J]. New Voices in Translation Studies, 2016, 15: 63-92.
- [20] Potet M, Esperanca Rodier E, Besacier L, et al. Collection of a large database of French-English SMT output corrections [C]//Proceedings of the 8th International Conference on Language Resources and Evaluation, 2012: 4043-4048.
- [21] Wisniewski G, Singh A K, Segal N, et al. Design and analysis of a large corpus of post-edited translation: quality estimation, failure analysis and the variability of post-edition [C]//Proceedings of the MT Summit XIV, 2013: 117-124.
- [22] Koponen M. Comparing human perceptions of post-editing effort with post-editing operations [C]//Proceedings of the 7th Workshop on Statistical Machine Translation, 2012: 181-190.
- [23] Fishel M, Bojar O, Popovic M. Terra: a collection of translation error-annotated corpora [C]// Proceedings of the 8th International Conference on Language Resources and Evaluation, 2012: 7-14.
- [24] Zeman D, Fishel M, Berka J, et al. Addicter: what is wrong with my translation? [J]. The Prague Bulletin of Mathematical Linguistics, 2011, 96: 79-88.
- [25] Popovic M. Hjerson: an open source tool for automatic error classification of machine translation output [J]. The Prague Bulletin of Mathematical Linguistics, 2011, 96: 59-68.
- [26] Avramidis E, Burchardt A, Hunsicker S, et al. The TARAXü corpus of human-annotated machine translations [C]//Proceedings of the 9th International Conference on Language Resources and Evaluation, 2014: 2679-2682.
- [27] Popovic M, Arcan M. PE2rr Corpus: manual error annotation of automatically pre-annotated MT post-edits [C]//Proceedings of the 10th International Conference on Language Resources and Evaluation, 2016: 27-32.
- [28] Barrault L, Bojar O, Costa-jussa M R, et al. Findings of the 2019 conference on machine translation [C]//Proceedings of the 4th Conference on Machine Translation, 2019: 1-61.
- [29] Symne S, Ahrenberg L. On the practice of error analysis for machine translation evaluation [C]//Proceedings of the 8th International Conference on Language Resources and Evaluation, 2012: 1785-1790.
- [30] Kirchhoff K, Capurro D, Turner A. Evaluating user preferences in machine translation using conjoint analysis [C]//Proceedings of the 16th EAMT Conference, 2012: 119-126.
- [31] Popovic M, Lommel A, Burchardt A, et al. Relations between different types of post-editing operations, cognitive effort and temporal effort [C]//Proceedings of the 17th Annual Conference of the European Association for Machine Translation, 2014: 191-198.



袁白莲(1981—),博士研究生,讲师,主要研究领域为计算语言学、机器翻译。
Email: qiubl@ecjtu.edu.cn



李茂西(1977—),博士,教授,主要研究领域为自然语言处理和机器翻译。
E-mail: mosesli@jxnu.edu.cn



王明文(1964—),通信作者,博士,教授,博士生导师,主要研究领域为自然语言处理、信息检索、数据挖掘、机器学习。
Email: mwwang@jxnu.edu.cn