

智能多轮对话改写研究综述

梁颖红, 桂文明, 吴秋玲

(金陵科技学院网络安全学院, 江苏 南京 211169)

摘 要:指代和省略是多轮对话改写技术中急需解决的关键问题之一,直接影响多轮对话的准确率。通过介绍以指代和省略为主的多轮对话改写的定义、基本思想、技术模型和方法、目前的数据集和评价指标、未来发展趋势等内容,给多轮对话改写研究形成完整的画像,为相关研究提供参考和借鉴。

关键词:多轮对话;指代;省略;改写;人工智能

中图分类号:TP391.1;TP18

文献标识码:A

文章编号:1672-755X(2021)03-0001-07

Survey of the Research on the Rewriting of Intelligent Multi Round Dialogue

LIANG Ying-hong, GUI Wen-ming, WU Qiu-ling

(Jinling Institute of Technology, Nanjing 211169, China)

Abstract: The pronoun reference and omission are the key problems in multi round dialogue technology, which directly affect the accuracy of multi round dialogue. The purpose of this paper is to provide a complete picture for the researchers by introducing the definition, the basic idea, the model and method, the current corpus and evaluation indicator, and the future development trend etc. of multi round dialogue rewriting of pronoun reference and omission, which could be used to provide reference for relevant research.

Key words: multi round dialogue; pronoun reference; omission; rewriting; artificial intelligence

近年来,随着人工智能(artificial intelligence, AI)技术的发展,多种 AI 产品应运而生,运用智能对话系统的对话机器人即其中之一。智能对话系统在各种智能算法的支撑下,使机器理解人类语言的意图并通过有效的人机交互执行特定任务或做出回答。智能对话系统大致可分为两种^[1]:1)任务导向型对话系统,旨在帮助用户完成实际具体的任务,例如帮助用户找寻商品、预订酒店餐厅等,典型代表有苹果语音助手 Siri 和亚马逊的 Echo。2)非任务导向型对话系统,也称为聊天机器人,典型代表是微软小冰。近年来集闲聊、娱乐、技能、任务为一体的天猫精灵、百度小度、腾讯叮当等也纷纷上市,开启了智能对话系统实用性的进程。

回顾智能对话系统的发展历程,大致可以分为 3 个阶段^[2]。第一阶段:通过关键词匹配和人工编写的一些规则来实现,典型代表是 1966 年 MIT 开发的模仿心理医生的机器人 Eliza。这种对话系统严重依赖于人工构建的知识,具有领域的局限性。第二阶段:采用统计机器学习的方法来实现,具有智能学习的能力,虽然摆脱了专家总结知识的弊端,但存在可解释性差、准确率不高等问题。第三阶段:通过深度学习,

收稿日期:2021-04-19

基金项目:江苏省“333”工程项目(BRA2015108);金陵科技学院高层次人才引进启动基金项目(jit-rcyj-201510)

作者简介:梁颖红(1970—),女,黑龙江齐齐哈尔人,教授,博士,江苏省“333”高层次人才,主要从事自然语言处理和对话系统研究。

利用大量的数据来学习特征表示和回复生成策略,提升了对话的智能性和准确率,是目前的主流方法。

智能对话分为单轮对话和多轮对话,多轮对话的难度要高于单轮对话,对话产品的推广应用需要突破多轮对话的技术壁垒。AI 对话技术中,语义理解(natural language understanding, NLU)是核心环节之一。多轮对话中,经常会发生指代和省略的情况^[3],对语义理解环节产生较大的影响。指代(pronoun reference)是指用代词、称谓和缩略语来指代前面提到的实体全称。省略(omission)也称为零指代(zero pronoun),是指在篇章中读者能够根据上下文的关系推断出来的部分经常被省略,被省略部分在句子中又承担相应的句法成分,并且回指前文中的某个语言单位。被省略的部分称为零指代项或者零代词,被指向的语言单位称为先行语^[4]。

下面是指代和缺省的例子:

指代:A:我读过《红楼梦》。

B:我在初中时也读过它。(这里的“它”指代前面的《红楼梦》。)

省略(零指代):A:你喜欢吃蛋糕吗?

B:喜欢! 你呢?(补全后:你喜欢吃蛋糕吗?)

A:我也喜欢。(补全后:我也喜欢吃蛋糕。)

1 多轮对话改写的定义

在多轮对话中,往往要获取上下文的信息,找到指代和省略的对应部分,以便做出正确的理解。如果把所有的对话信息全部作为输入,会增加模型处理的难度和时间,而且对一些长对话来说,处理起来更为复杂。

图 1 是多轮语音对话的流程图。语音识别模块产生语音识别结果(用户话语);语义理解模块产生用户对话行为;对话管理模块选择需要执行的系统行为;如果这个系统行为需要和用户交互,那么语言生成模块会被触发,生成自然语言或者系统话语;最后,生成的语言由语言合成模块朗读出来。

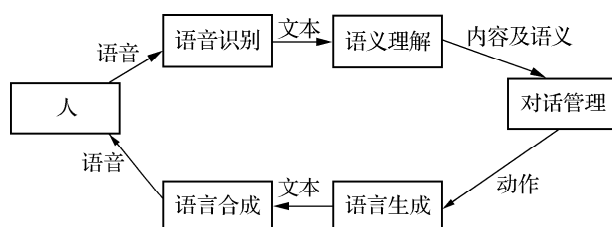


图 1 多轮语音对话流程

多轮对话改写,属于图 1 中语义理解模块需要完成的主要任务,是根据历史轮和当前轮的信息将当前轮语句中的指代和省略补齐,改写后的句子能够表达完整的语义信息。多轮改写的形式化定义如下:

以历史轮和当前轮的对话作为输入(H, U_n),其中 $H = \{U_1, U_2, \dots, U_{n-1}\}$, 目标是根据历史对话信息,将当前轮的语句 U_n 改写成 R , 即 $(H, U_n \rightarrow R)$ 。

2 多轮对话改写的模型及其比较

多轮对话改写的核心其实在于对 U_n 进行指代消解和省略补全,希望改写之后的语句 R 能够表达更完整的语义信息,使多轮对话变成单轮对话。

2.1 多轮对话改写的模型和算法

通常把多轮对话改写看成以下 3 种情况。

2.1.1 把多轮对话改写任务分成指代和省略的检测以及指代消解

采用 pipeline 模型,先解决指代和省略的检测,然后再处理指代消解,是一种顺序处理模型。

指代和省略的检测,就是发现 U_n 中代词的位置以及可能出现省略的位置^[5-6]。方法有两种:1)基于

规则的方式。通过字符串匹配的方式找到代词,基于规则模板匹配省略的位置。2)基于序列标注的方式。用符号(如0表示常规的单词,1表示代词的边界,2表示代词)标出代词的开始边界和位置,从而找出代词。

指代消解^[7-8],是根据历史轮和当前轮的对话,从候选填充词集合中找到最适合对应位置的答案。采用分词以及N-gram组合的方式选出候选填充词。

2.1.2 把多轮对话改写任务看成句子生成问题

随着深度学习的应用,把多轮对话改写任务看成句子生成问题是目前的主流策略。基于该策略的模型大多运用神经网络模型,大致可以分为两类:基于长短期记忆(long short-term memory, LSTM)的模型和基于转换(transformer)的模型。

1)基于LSTM的模型

A. 带有注意力(attention)机制的序列到序列(seq2seq)模型

seq2seq模型使用编码器-解码器架构^[9-11],将源序列映射到目标序列。注意力机制避免为每个句子学习单一的向量表示,而根据注意力权值来关注输入序列的特定输入向量。在每一解码步骤中,解码器将被告知需要使用一组注意力权重对每个输入单词给予“关注”。这些注意力权重为解码器提供上下文语义信息。

B. 带有复制(copy)机制的seq2seq模型

该模型同样由两部分组成:编码器和解码器^[12]。编码器将自然语言文本转换成定长的语义向量,解码器根据该语义向量生成三元组。解码器根据数量不同可以分为OneDecoder和MultiDecoder。OneDecoder就是用一个解码器来生成所有三元组,MultiDecoder是由多个解码器组成,一个解码器生成一个三元组。通过模型提供位置信息,能够在一定程度上解决未登录词(out-of-vocabulary, OOV)问题。

C. 指针网络(pointer network)结构模型

指针网络是针对编码器-解码器中的注意力机制进行改进的^[13],是一个全新的神经网络模型,相较于利用许多经典的RNN方法处理自回归问题,pointer network的最大贡献是解决了输出字典依赖输入序列长度的问题。与经典RNN中seq2seq的方法相比,pointer network不再依赖解码器的输出词典表,直接输出的是输入序列各值在当前解码位置的概率分布,并以最大值方式输出可能的输入序列中的值向量。

D. 混合指针生成网络模型

seq2seq模型有2个缺点:复制的细节不准确和倾向于重复自己。而混合指针生成网络从两方面做了改进^[14-15]:1)使用指针生成器网络通过指向从源文本中复制单词,这有助于准确复制信息,同时保留generator的生成能力;2)使用coverage跟踪内容,不断更新注意力,从而阻止文本不断重复。

2)基于转换(transformer)的模型

transformer是由一个个transformer块堆叠而成的多层架构^[14-16],深层transformer堆叠+海量无监督语料预训练的模式受到很多研究者的关注。目前神经语言程序学在各业务应用的语言模型如GPT(generative pre-training)、BERT(bidirectional encoder representations from transformers)等,都是基于transformer的模型。基于transformer的模型的主要特征是多头注意力(multi-head attention)机制和位置编码^[15]。

multi-head attention的本质是,在参数总量保持不变的情况下,将同样的query, key, value映射到原来高维空间的不同子空间中进行attention计算,在最后一步合并不同子空间中的attention信息。这样降低了计算每个head的attention时的向量维度,在某种意义上防止了过拟合。由于attention在不同子空间中有不同的分布, multi-head attention实际上是寻找了序列之间不同角度的关联关系,并在最后一步中,将不同子空间中捕获到的关联关系综合起来。

位置编码是transformer框架中特有的组成部分,保存单词在序列中的相对或绝对位置,补充了attention机制本身不能捕捉位置信息的缺陷,位置信息直接叠加于语义信息embedding之上,使得每个token的位置信息和它的embedding充分融合,并被传递到后续所有经过复杂变换的序列表达中去。

2.1.3 把多轮对话改写任务看成语义分割问题

这种方法将多轮对话改写看作是语义分割问题^[17],将改写之后的结果看作对不完整的输入进行一系列编辑之后的结果。具体来说,拼接所有上下文语句得到长度为 M 的序列 $C = \{c_1, c_2, \dots, c_M\}$ 。同时,将不完整的输入表示为 $X = \{x_1, x_2, \dots, x_N\}$ 。于是,可以根据不完整的语句 X 以及上下文的序列 C ,编辑得到重写之后的语句 X^* 。

如图 2 所示:用 c_1, c_2 替换 x_3 ,把 c_5, c_6 插入到 x_6 之前,这样通过替换与插入等编辑操作完成了句子的改写。

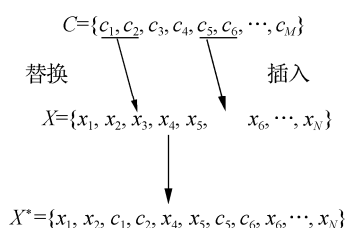


图 2 编辑操作示意图

2.2 基于 LSTM 的模型和基于 transformer 的模型的比较

基于 LSTM 的模型和基于 transformer 的模型是目前多轮对话改写的主流模型,下面对两者进行比较:

基于 transformer 的模型相较于基于 LSTM 的模型出现晚,在诞生之初就针对基于 LSTM 的模型的缺陷进行了改进,因此,在目前的应用中,模型的效果大多优于 LSTM 模型。具体分析如下:基于 LSTM 的模型无法并行训练和推理,不利于大规模快速训练;基于 LSTM 的模型会有长程依赖信息丢失问题,基于 transformer 的模型可以动态建立输入序列的长程依赖关系;基于 transformer 的模型采用多头注意力机制,可以实现多个转换块的叠加,各个注意头可以执行不同的任务;虽然基于 transformer 的模型中,自注意力(self-attention)没有循环结构,但采用位置嵌入(position embedding)来表示序列中元素的相对或绝对位置关系;基于 transformer 的模型缺乏对时间维度的建模,即使有位置嵌入也和 LSTM 这种天然的时序网络有差距。

在基于 LSTM 的模型中,带有 copy 机制的 seq2seq 模型与指针结构网络模型之间的比较如下:带有 copy 机制的 seq2seq 模型实现有些复杂,指针结构网络模型的效果更好,且更容易实现;指针结构网络模型使源文本生成单词更加容易,甚至可以复制源文本中的非正式单词,处理未登录词问题;指针结构网络模型允许使用更小规模的词汇集,需要较少的计算资源和存储空间,训练速度更快。

3 多轮对话的数据集和评价指标

3.1 多轮对话的数据集

从事多轮对话改写研究的人员大多使用多轮对话的数据集,数据集的具体情况如下:

3.1.1 中文数据集

A. 清华大学 SU H, SHEN X Y 等^[17]从几个流行的中文社交媒体平台上抓取了 20 万个候选的多轮会话数据,针对多轮对话改写任务进行了人工标注,形成了专门针对多轮对话改写任务的数据集。

B. CrossWOZ 大规模跨领域任务导向多轮对话数据集^[18], <https://github.com/thu-coai/Cross-WOZ>。CrossWOZ 是第一个大规模中文跨领域任务导向数据集,包含 6 000 个对话,10.2 万个句子,涉及 5 个领域(景点、酒店、餐馆、地铁、出租)。平均每个对话涉及 3.2 个领域,远超之前的多领域对话数据集,增添了对话管理的难度。

C. OpenViDial 多模态对话数据集^[19], <https://github.com/ShannonAI/Op>。从电影和电视剧中抽取句子、图片对,经过数据处理与清洗,最终得到 100 万余条句子,及其对应的图片信息。

D. NaturalConv 数据集^[20]。这是一个主题驱动的中文多轮对话数据集,它更接近于人类对话,具有

自然属性,包括场景假设、自由话题扩展、问候等。包含约 40 万个语句和 1.99 万个对话,涉及多个领域(包括但不限于体育、娱乐、科技)。平均轮数为 20,明显长于其他语料库。

E. 中文电影台词数据集。这是电影对白语料集, https://github.com/rustch3n/dgk_lost_conv。

3.1.2 英文数据集

A. 斯坦福大学 NLP 小组多轮对话数据集, <http://nlp.stanford.edu/projects/kvret>。这组数据集包含了 3 031 条多轮对话数据,内容主要包含日程安排、天气信息检索和兴趣点导航等。这个对话集是通过知识库建立的,确保系统对自然语言处理得灵活流利。

B. MuTual 是基于人标签推理的多轮对话数据集, <https://github.com/Nealcly/MuTual>。该数据集用于针对性地评测模型在多轮对话中的推理能力,有助于将来进行多轮对话推理问题的研究。

C. OpenSubtitles 电影字幕集, <http://opus.lingfil.uu.se/OpenSubtitles.php>。

D. Twitter 数据集, https://github.com/Marsan-Ma/twitter_scraper。

E. Papaya Conversational Data Set,由 Cornell、Reddit 等数据集重新整理而来, <https://github.com/bshao001/ChatLearner>。

3.2 多轮对话的评价指标

1) BLEU(bilingual evaluation understudy)通过累积 N-gram BLEU 的得分,评估重写后的句子与标签句子的相似程度。

首先计算 N-gram 的概率^[14]:

$$p_n = \frac{\sum_{C \in Candidates} \sum_{N\text{-gram} \in C} Count_{clip}(N\text{-gram})}{\sum_{C' \in Candidates} \sum_{N\text{-gram}' \in C'} Count(N\text{-gram}')} \quad (1)$$

其中, *Candidates* 表示模型输出的结果, *N-gram* 表示每一个词组。首先计算在所有参考答案中 *N-gram* 出现的次数,取其最大值 *Max*,将 *Max* 和候选答案中出现的次数取最小值得到 *Min*,将所有 *N-gram* 对应的 *Min* 相加得到式子的分子。分母为每一个 *N-gram* 在候选答案中出现的次数的和。

为了避免出现候选句越短,参考句越长,分数反而越高的情况,需要增加一个惩罚因子:

$$BP = \begin{cases} 1, & c > r \\ e^{(1-\frac{c}{r})}, & c \leq r \end{cases} \quad (2)$$

其中, *c* 表示候选答案文本单词数, *r* 表示参考答案文本单词数。

$$BLEU = BP \cdot \exp\left(\sum_{n=1}^N w_n \log_2 p_n\right) \quad (3)$$

其中, w_n 是一个恒定的权重, p_n 表示 *N-gram* 的最大值,一般取 4。

2) 用 ROUGE-N 度量重写的句子和标签句子的 *N-gram* 重叠^[16]。

$$ROUGE-N = \frac{\sum_{S \in GroundTruth} \sum_{N\text{-gram} \in S} Count_{match}(N\text{-gram})}{\sum_{S \in GroundTruth} \sum_{N\text{-gram} \in S} Count(N\text{-gram})} \quad (4)$$

其中,分子表示参考文本和生成文本共有的 *N-gram* 的个数。分母表示参考文本中 *N-gram* 的个数(也可以是生成文本中 *N-gram* 的数量,只不过这里更关心召回率)。对于存在多条 references 的情况,让 candidate 针对每个 reference 进行计算,取其中的最大值。

3) 用 ROUGE-L 度量重写的句子和标签句子之间的最长匹配序列,同时考虑 Precision 和 Recall^[14]。

$$\begin{cases} R_{lcs} = \frac{LCS(X,Y)}{m} \\ P_{lcs} = \frac{LCS(X,Y)}{n} \\ F_{lcs} = \frac{(1+\beta^2)R_{lcs}P_{lcs}}{R_{lcs} + \beta^2 P_{lcs}} \end{cases} \quad (5)$$

其中, $LCS(X,Y)$ 表示 *X* 和 *Y* 的最长公共子序列的长度, *m* 和 *n* 分别表示参考答案和改写结果的长度。

4) EM代表精确匹配准确度^[15],这是最严格的评估指标,表示完全匹配的改写结果占有测试集的比例,即:

$$EM = \frac{Count_{\text{exact-match}}}{Count_{\text{test}}} \quad (6)$$

4 多轮对话改写的发展趋势

目前,单轮对话技术相对比较成熟,但很多应用中单轮对话无法满足需要,因此多轮对话技术急需发展。多轮对话改写直接影响多轮对话最后的结果,在今后一段时间将会引起广泛关注。

传统对话系统由于缺乏大量的领域知识及基础常识的支撑,目前还不能执行深度的推理与决策。下一代对话系统应该是同时具备很强的任务完成能力和社交连接能力,满足用户的信息需求和社交需求^[21-22]。对话系统下一步必然走向全方位多模态的交互方式,通过视觉、语音、语言、知识等交互融合,使得人与机器的交流变成无限制的交流。另外,目前广泛应用的对话系统服务于越来越多的人,通过互动、理解和推理的学习过程,对话助手可以无意中隐蔽地存储一些较为敏感的信息,因此隐私保护也是对话系统需要考虑的问题。

从技术上,当前的大模型和大数据肯定是一个研究趋势和潮流,多轮对话改写研究采用深度学习的神经网络模型应该是主要趋势。融入常识和知识,精准理解对话历史,强化记忆机制,从而提高上下文的一致性,将是进一步的研究目标。对话模型将有能力从与人的交互中主动去学习,持续探索的技术方向有生成模型、强化模型、迁移学习、机器阅读理解、情感分析等。

参考文献:

- [1] 赵阳洋,王振宇,王佩,等.任务型对话系统研究综述[J].计算机学报,2020,43(10):1862-1896
- [2] 姚冬,李舟军,陈舒玮,等.面向任务的基于深度学习的多轮对话系统与技术[J].计算机科学,2021,48(5):232-238
- [3] YIN Q Y,ZHANG Y,ZHANG W N,et al. Deep reinforcement learning for Chinese zero pronoun resolution[C]. Melbourne:Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics,2018:569-578
- [4] YIN Q Y,ZHANG W N,ZHANG Y,et al. A deep neural network for Chinese zero pronoun resolution[C]. Melbourne:Proceedings of the Twenty-Sixth International Joint Conference on Artificial Intelligence,2017:3322-3328
- [5] YIN Q,YU Z,ZHANG W,et al. Zero pronoun resolution with attention-based neural network[C]. New Mexico:Proceedings of the 27th International Conference on Computational Linguistics,2018:13-23
- [6] LIU T,CUI Y M,YIN Q Y,et al. Generating and exploiting large-scale pseudo training data for zero pronoun resolution [C]. Vancouver:Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics,2017:102-111
- [7] ZHANG H M,SONG Y,SONG Y Q,et al. Knowledge-aware pronoun coreference resolution[C]. Florence:Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics,2019:867-876
- [8] YANG W,QIAO R,QIN H C,et al. End-to-end neural context reconstruction in Chinese dialogue[C]. Florence:Proceedings of the First Workshop on NLP for Conversational AI,2019:68-76
- [9] MALMI E,KRAUSE S,ROTHER S,et al. Encode,tag,realize:high-precision text editing[C]. Hong Kong:Conference on Empirical Methods in Natural Language Processing(EMNLP),2019:5054-5065
- [10] GU J,LU Z,LI H,et al. Incorporating copying mechanism in sequence-to-sequence learning[C]. Berlin:Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics,2016:1631-1640
- [11] VINYALS O,FORTUNATO M,JAITLEY N. Pointer networks[J]. Computer Science,2015(1):28-35
- [12] SEE A,LIU P J. Get to the point;summarization with pointer-generator networks[C]. Vancouver:Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics,2017:1073-1083
- [13] 祝凯华. Transformer 多轮对话改写实践[EB/OL]. (2020-04-29)[2021-03-27]. <https://zhuanlan.zhihu.com/p/137127209>
- [14] HUANG M Z,LI F,ZOU W H,et al. SARG:A novel semi autoregressive generator for multi-turn incomplete utterance

- restoration[C]. Beijing: Proceedings of the Association for the Advance of Artificial Intelligence(AAAI), 2021
- [15] LIU Q, CHEN B, LOU J G, et al. Incomplete utterance rewriting as semantic segmentation[C]. Pontacana: Proceedings of the Empirical Methods in Natural Language Processing(EMNLP), 2020: 1 – 12
- [16] GU S H, FENG Y, LIU Q. Improving domain adaptation translation with domain invariant and specific information[C]. Minneapolis: 2019 Annual Conference of the North American Chapter of the Association for Computational Linguistics (NAACL-HLT 2019), 2019: 3081 – 3091
- [17] SU H, SHEN X Y, ZHANG R Z, et al. Improving multi-turn dialogue modelling with utterance rewriter[C]. Florence: Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics, 2019: 1 – 11
- [18] ZHU Q, HUANG K L, ZHANG Z, et al. CrossWOZ: A large-scale Chinese cross-domain task-oriented dialogue dataset [J]. Transactions of the Association for Computational Linguistics, 2020, 2: 1 – 14
- [19] MENG Y X, WANG S H, HAN Q, et al. OpenViDial: a large-scale, open-domain dialogue dataset with visual contexts [J]. arXiv: 2012.15015v1
- [20] WANG X Y, LI C, ZHAO J Q, et al. NaturalConv: a Chinese dialogue dataset towards multi-turn topic-driven conversation[C]. Beijing: Proceedings of the Association for the Advance of Artificial Intelligence(AAAI), 2021: 1 – 9
- [21] 肖鹏, 于丹, 王建超, 等. 情感型对话机器人技术的研究综述[J]. 软件工程, 2021, 24(2): 14 – 18
- [22] 庄寅, 刘箴, 刘婷婷, 等. 文本情感对话系统研究综述[J]. 计算机科学与探索, 2021, 15(5): 825 – 837

(责任编辑: 湛江)

声 明

本刊已被中国知网、万方、维普、超星等数据库收录, 如无特殊申明, 即视为投稿者同意授权本刊及本刊合作媒体以数字化方式复制, 并通过信息网络传播和发行。本刊支付的稿酬已包括上述所有使用方式的报酬。

作者向本刊提交文章发表的行为视为同意我刊上述声明。

本刊编辑部