

计算语言学  
与语言科技  
原文丛书

CAMBRIDGE

# THE SPOKEN LANGUAGE TRANSLATOR

# 口语机器翻译

编者 Manny Rayner, David Carter, Pierrette Bouillon,  
Vassilis Digalakis, Mats Wirén

导读 宗成庆



北京大学出版社  
PEKING UNIVERSITY PRESS

# The Spoken Language Translator

## 口语机器翻译

编者 Manny Rayner      David Carter  
Pierrette Bouillon      Vassilis Digalakis  
Mats Wirén

导读 宗成庆



北京大学出版社  
PEKING UNIVERSITY PRESS

著作权合同登记 图字 01-2009-4470

图书在版编目(CIP)数据

口语机器翻译: The Spoken Language Translator/瑞诺尔(Rayner, M.), 卡特尔(Carter, D.), 布伊隆(Bouillon, P.)等编. —北京:北京大学出版社, 2010. 8

(计算语言学与语言科技原文丛书)

ISBN 978-7-301-17156-1

I. 口… II. ①瑞… ②卡… ③布… III. 口语—机器翻译—英文  
IV. H085

中国版本图书馆 CIP 数据核字(2010)第 075822 号

*The Spoken Language Translator*, first edition (ISBN: 978-0-521-03882-0) by Rayner, Carter, Bouillon Digalakis and Wirén first published by Cambridge University Press 2000. All rights reserved.

This reprint edition for the People's Republic of China is published by arrangement with the Press Syndicate of the University of Cambridge, Cambridge, United Kingdom.  
© Cambridge University Press & Peking University Press 2010

This book is in copyright. No reproduction of any part may take place without the written permission of Cambridge University Press and Peking University Press.

This edition is for sale in the People's Republic of China (excluding Hong Kong SAR, Macau SAR and Taiwan Province) only.

此版本仅限在中华人民共和国境内(不包括香港、澳门特别行政区及台湾地区)销售。

书 名: 口语机器翻译

著作责任者: Manny Rayner David Carter Pierrette Bouillon  
Vassilis Digalakis Mats Wirén 编

责任编辑: 白雪李凌

标准书号: ISBN 978-7-301-17156-1/H·2493

出版发行: 北京大学出版社

地 址: 北京市海淀区成府路 205 号 100871

网 址: <http://www.pup.cn>

电子邮箱: [zpup@pup.pku.edu.cn](mailto:zpup@pup.pku.edu.cn)

电 话: 邮购部 62752015 发行部 62750672 编辑部 62753334  
出版部 62754962

印 刷 者: 世界知识印刷厂

经 销 者: 新华书店

787 毫米×1092 毫米 16 开本 23.5 印张 388 千字

2010 年 8 月第 1 版 2010 年 8 月第 1 次印刷

定 价: 46.00 元

未经许可,不得以任何方式复制或抄袭本书之部分或全部内容。

版权所有,侵权必究 举报电话: 010-62752024

电子邮箱: [fd@pup.pku.edu.cn](mailto:fd@pup.pku.edu.cn)

《计算语言学与语言科技原文丛书》由北京大学—香港理工大学汉语语言学研究、北京大学计算语言学研究所(由 973 课题“文本内容理解的数据基础”支持)和北京大学出版社合作推出

## 学术委员会 Academic Advisory Committee

主 任：

黄居仁(香港)

委 员：

陈克健(台北)

Chris Manning (Stanford)

董振东(北京)

Harold Somers (Dublin)

李宇明(北京)

陆俭明(北京)

Maarten de Rijke (Amsterdam)

沈 阳(北京)

石定栩(香港)

苏克毅(台北)

Suzanne Stevenson (Toronto)

王逢鑫(北京)

王厚峰(北京)

王士元(香港)

谢清俊(台北)

俞士汶(北京)

松本裕治(奈良)

郑锦全(Urbana-Champaign)

邹嘉彦(香港)

## 编委会

### Editorial Committee

主 编：

黄居仁教授(香港)

编 委：

冯志伟教授(北京)

顾曰国教授(北京)

黄伟道教授(Singapore)

黄莹菁教授(上海)

姬东鸿教授(武汉)

陆 勤教授(香港)

蒙美玲教授(香港)

苏新春教授(厦门)

孙茂松教授(北京)

陶红印教授(Los Angeles)

徐飞玉教授(Saarbrücken)

薛念文教授(Waltham)

杨立范编审(北京)

俞士汶教授(北京)

曾淑娟副研究员(台北)

詹卫东副教授(北京)

赵铁军教授(哈尔滨)

周 明研究员(北京)

宗成庆研究员(北京)

常宝宝副教授(执行秘书)(北京)

## 丛书前言

计算语言学(Computational Linguistics, CL)在语言科学与信息科学的研究领域扮演关键性的角色。语言学理论寻求对语言现象规律性的揭示与完整的解释。计算语言学正好提供了验证与应用这些规律与解释的大好机会。作为语言学、信息科学乃至心理学与认知科学结合的交叉学科,计算语言学更提供了语言学基础研究与应用研究的绝佳界面。事实上,计算语言学与人类语言科技(Human Language Technology, HLT)可以视为一体两面,不可分割。

计算语言学研究滥觞于上世纪五六十年代的机器翻译研究。中文的相关研究也几乎同步开始,1960年起在柏克莱加州大学研究室,王士元、邹嘉彦、C. Y. Dougherty 等人已开始研究中英、中俄机器翻译。他们的中文计算语言学研究,可说是与世界最尖端科技同步的。中国国内中俄翻译研究也不遑多让,大约在上世纪50年代中期便已开始。可惜的是,这些中文相关早期机器翻译研究,由于硬件与软件的限制,没能延续下来。中文计算语言学研究比较有系统的进展,还要等到1986年;海峡两岸在同一年成立了两个致力于中文计算语言学基础架构建立的研究群。北京大学的计算语言学研究所在朱德熙先生倡导下成立,随后一段时间由陆俭明、俞士汶主持。而台湾“中研院”的中文词知识库小组,由谢清俊创立,陈克健主持,黄居仁1987年返台后加入。

中文计算语言学的研究,20余年来已累积了相当可观的成绩,重要研究领域与议题中都有可观的研究成果,华人计算语言学者也渐渐在国际学术界崭露头角。随着世界经济转向知识密集产业,跨语言跨文化沟通与知识整合成为知识产业的关键,语言科技的发展日渐成为国际主流。在这个有利发展的大环境下,我们相信,中文计算语言学与华人计算语言学学者的成绩,将会百尺竿头更进一步,进入计算语言学学术核心,并产生把握学科动态、引领学术走向的大师。

回顾计算语言学研究在过去二十年的蓬勃发展,统计模式的引入应该是最主要的原因之一。但二十年后学界也开始看到了统计模式的局限,因此最近几届 ACL 终身成就奖得主,不约而同地大力提倡结合语言学理论与概率模型的研究,来提升计算语言学研究的层次,以寻求新的突破。

回顾中国国内的计算语言学发展,来自计算机科学的贡献多于语言学的贡献。这在理论与概率模型整合研究的大趋势下,不免令人忧心。这也许可以部分归咎于英文研究专著获得不易。国内较易取得期刊或会议论文,但由于篇幅的限制,往往无法对理论做深入完整的阐述,因此也导致国内年轻学者,长于运算而拙于理据。因此,藉由英文专书来弥补不足,巩固研究理据,进而开拓研究视野,是非常重要的步骤。

剑桥大学计算语言学原版书系列的引进,就是在上述背景下产生的。本人忝为 Cambridge University Press 所出版的 Studies in Natural Language Processing 系列编辑委员之一,并将于 2010 接任主编。能够将此系列中较重要的几部著作引进国内,责无旁贷。引进原版,不是难事;要真正搭建知识的桥梁,使国内学者与学生开拓研究视野,将原文著作的理论精髓,更多应用于中文研究,则需另加努力。因此,本丛书的特色,是在保留原版的基础上,每本书都邀请一位专家撰写中文导读,其着力点有三:

其一,全书内容简介。导读作者长年浸淫于该领域,对原著能提纲挈领,切中肯綮,并提供相关研究背景。可助读者更准确地掌握并吸收该书的内容。其二,中文相关研究。原作不一定会提到相关的中文研究。由导读专家补充介绍,能搭起理论与中文相关应用的桥梁,从而能够使读者掌握在这个议题进入中文研究的最佳切入点,让相关中文研究的开拓者获得理论的参照和指导。其三,补充原书出版后该领域研究的新发展。现代科技发展迅速,任何经典著作出版后,几乎马上有新的相关研究。因此,在理论架构的脉络中,加上新近发展,使读者能更贴切地掌握研究脉动。全书摘要通常采用文字叙述。而中文相关研究及最新研究发展则分别以文字叙述及延伸阅读书目的方式呈现。延伸阅读书目,使读者可以很快上手,进入相关研究领域,也是本丛书策划者的苦心所在。可以说导读是本丛书的亮点,不特为原书增色,亦且增加了不少附加价值。

本丛书的出版,是多方协作的结果。在规划出版的漫长过程中,北大计算语言研究所俞士汶老师及常宝宝老师提供了无私无悔的支持。香港理工大学,特别是北大一理大汉语语言学研究中心与陈瑞端、石定栩、沈阳几位在关键时刻的挹注,也起到了关键作用。当然,整个系列能够顺利出版,离不开有学术眼光和胸襟的北大出版社的支持,而剑桥出版社主管编辑 Helen Barton 从中斡旋,使合约能顺利签订,是必不可少的一环。最后,我要感谢本丛书的国内编委,特别是此次担任导读的各位主笔的辛勤付出,他们为读者搭建了进入学术殿堂的台阶。本丛书的出版,适逢 2010 COLING 国际计算语言学会议在北京举办之际,正象征着国内计算语言学研究与国际的接轨;国内学者风云际会,大展身手,跻身计算语言学的国际舞台,将指日可待。

丛书主编

黄居仁

谨志于香港红磡

二零一零年元月

# 导 读

宗成庆

## 1. 背景概述

口语翻译(Spoken Language Translation, SLT<sup>①</sup>)是指让计算机实现从一种语言的语音到另一种语言的语音自动翻译的过程。其理想目标是,让计算机像人一样充当持不同语言的说话人之间的翻译角色。会议演讲、交谈(通过电话、网络或面对面)、广播等场景下的话语翻译都是口语翻译应用的重要领域。由于多数情况下说话人的话语都以口语风格为主,人们尤其希望翻译系统可以接受并实现任意口语化的、自由交谈式的对话语音直接翻译,因此,口语翻译又称对话翻译(dialogue translation),或者直观地称为语音翻译(speech-to-speech translation, S2ST)。

口语翻译涉及计算语言学、计算机科学、语音和通讯技术等多种学科与技术,因此,开展这项研究具有非常重要的理论意义。同时,该技术一旦获得突破,可以广泛地应用于人类社会生活的各个方面,例如,商贸会谈、民航信息咨询、国际会议(包括体育运动会)信息综合服务、旅游信息服务等等,因此,该技术蕴涵着潜在的巨大社会效益和经济利益。近年来,随着全球化信息社会的到来,人们在经济、商贸、体育、文化、旅游等各个领域的国际交流日益广泛,口语翻译成为人们迫切需要的技术而倍受关注。

口语翻译研究最早起始于上个世纪 80 年代中后期。1989 年美国卡内基—梅隆大学(Carnegie Mellon University, CMU)演示了世界上第一

---

① 请注意,此处的缩写“SLT”含意不同于后面的“口语翻译系统(Spoken Language Translator, SLT)”项目。本导读中“SLT”除了在此处指“口语翻译”以外,其余地方均指“SLT 项目”。

个语音翻译实验系统——SpeechTrans,该系统因此被称为语音翻译研究的里程碑。几乎在同一时期,日本政府和企业投入巨资成立了国际先端通信技术研究所以(Advanced Telecommunications Research Institute International, ATR),并专门设立了电话语音翻译实验室(Interpreting Telephony Laboratories, ITL),从事电话语音翻译技术的研究。进入90年代以后,欧美各国政府纷纷设立重大项目开展语音翻译研究,其中,当时的德国联邦教育研究部(Federal Ministry of Education and Research, BMBF)资助的 VERBMOBIL 语音翻译研究项目(1993年至2000年)是最具代表性的大项目,同时,美国国防部(DARPA)也资助了一系列语音翻译研究项目。

本书介绍的内容是由瑞典斯德哥尔摩 Telia Research 公司资助的首批语音翻译项目之一——“口语翻译系统(Spoken Language Translator, SLT)”的研究工作。SLT项目起始于1992年下半年,期间除了短暂的间断以外,项目一直持续到1999年中期结束。确切地说,SLT项目包含三个子项目:SLT-1项目从1992年到1993年中期,由瑞典 Telia Research、美国 SRI International 和瑞典计算机科学研究所(Swedish Institute of Computer Science, SICS)联合承担;SLT-2项目从1994年中期到1996年底结束,SICS退出了该期项目,取而代之的是日内瓦大学(University of Geneva)和克利特技术大学(Technical University of Crete)的相关机构;SLT-3在SLT-2结束之后很快上马,由SLT-2的同批人马执行,直到1999年中期结束。SLT翻译系统限定在航空旅游信息系统领域(Air Travel Information System, ATIS),实现了英语和瑞典语双向口语翻译和英语、瑞典语到法语的口语翻译,从法语到英语的口语翻译只是一个初始版本。

本书全面、系统地介绍了SLT项目研究的主要成果,内容包括语言处理与语料收集、语言覆盖性、语音处理和系统评估等各个方面,其中语言处理部分是本书的主要内容。

## 2. 内容介绍

全书内容包括四个部分:引言、语言处理与语料收集、语言覆盖性、语

音处理、系统评估及全书的结束语。

## 2.1 引言部分

引言部分包括 7 个小节,其中,1.1 节简要介绍本书的背景和关于口语翻译的基本概念,包括:为什么要做口语翻译(其应用价值是什么)? 口语翻译的基本问题是什么? 目前(本书写作时)口语翻译实现的可行性在哪里? 已经实现了哪些技术?

1.2 节给出了 SLT 的系统框架。1.3 节则通过一个英语句子被翻译成法语句子的实现过程简要解释 SLT 系统的翻译方法。1.4 节对 SLT 项目中使用手工编码的语法方法进行语言处理的理由给予了简要解释,并对如何将一个效率较低的通用句法分析器转换成特定领域的高效分析器的基本思想进行了介绍。1.5 节通过英法翻译的例子说明 SLT 采用的混合转换方法(hybrid transfer)中知识源划分的基本思路 and 不同知识源之间的相互作用。1.6 节简要介绍了 SLT 系统中语音识别模块的基本要点。1.7 节简要说明了关于 ATIS 语料的基本情况。

## 2.2 第一部分:语言处理与语料收集

第一部分介绍 SLT 系统的语言处理方法和语料收集情况,包括第 2 章至第 8 章的内容。

其中,第 2 章主要介绍利用 SRI 核心语言引擎(core language engine, CLE)实现机器翻译的基本方法。SLT 的翻译模块采用多引擎翻译机制。所谓多引擎实际上只有两个翻译引擎,一个是基于表层转换算法的词对词(word-to-word)翻译引擎,2.2 节简要介绍了该引擎的基本思想;另一个是基于准逻辑形式(quasi logical form, QLF)的翻译引擎。2.3 节对 QLF 的基本思想、结构形式、QLF 的形式化表示以及 QLF 的中心词关系给予了简要介绍。CLE 利用合一文法(unification grammar)规则定义字符串(句子)和 QLF 之间的关系。2.4 节面向 CLE 实现语言分析的问题(将表层字符串转换成 QLF 表示和基于 QLF 生成自然语言句子),介绍了 CLE 合一文法的形式化表示,并分别给出了利用合一文法表示法语名词短语、瑞典语从句和瑞典语关系从句的例子。2.5 节介绍的是词语形态分析方法。2.6 介绍翻译转换规则的基本形式。在

SLT 系统中转换规则包括两类,一类为简单转换规则,另一类为递归转换规则。简单转换规则用于将源语言的 QLF 片断映射到目标语言,而递归转换规则用转换变量表示一个 QLF 表达式的转换,包括一个或多个子表达式的转换。2.7 节简要介绍了基于 QLF 的语言处理机制,包括源语言句子到 QLF 的转换、源语言 QLF 到目标语言 QLF 的转换以及基于 QLF 的目标语言句子生成。

第 3 章和第 4 章介绍如何将一个通用语法通过自适应方法调整后用于特定领域的语句解析。其中,第三章侧重介绍基于解释学习方法(explanation-based learning, EBL)的语法专用化处理方法,而第四章则主要介绍句法分析中的剪枝策略和候选结果筛选方法。

第 5 章介绍的是一个树库工具(TreeBanker)。TreeBanker 是一个数据库管理系统,用于对 CLE 模块进行有监督地训练以获得面向特定用户的判别式。即对于一个用户来说(用户未必是系统专家),TreeBanker 利用 CLE 为语料库中的每一个语句提供一组可能的判别式,由用户选择正确的 QLF,从而在较短的时间内为特定用户训练一个实用的 CLE。

对于 CLE 来说,一个重要的问题是当系统被移植到一个新的领域时,如何确保足够的词汇覆盖率,即如何将新领域的词汇补充到系统已有的词汇表里,同时又保证这些词汇的语义能够得到正确的解释。第 6 章就是专门介绍在 SLT 项目中如何实现词条获取的,包括如何利用一批语料判断每一个词汇是否要被添加到系统的词汇表里,然后利用系统可读的格式描述这些词汇的词形、形态、句法和语义等信息。

第 7 章介绍 CLE 模块形态分析规则的编译器和实现环境。需要说明的是,尽管 SLT 系统中使用了 CLE 的形态分析功能,但是,这部分工作并非专门为 SLT 系统所做,而是在其他项目中已经完成的。

第 8 章介绍 SLT 项目中语料和数据收集的相关工作。为了使 SLT 系统在某个特定的领域中达到最佳性能,许多模块需要利用数据驱动的方法在领域语料上进行训练。这一章专门介绍如何收集和处理领域语料。

## 2.3 第二部分:语言覆盖问题

第二部分介绍 SLT 系统的语言覆盖情况,包括第 9 章至第 13 章的

内容。

第9章比较详细地描述了英语语法,包括所有规则和词汇在 ATIS 领域中的覆盖情况,而第10章和第11章分别介绍法语和瑞典语的语言覆盖情况。由于从时间顺序上来讲,英语语法被开发的最早(实际上,英语语法的开发原本是专门研究 CLE 的项目中的一项任务),而且法语和瑞典语的语法与英语语法仅有很小的差异,甚至有些地方没有差异,因此,法语和瑞典语的语法都是在英语语法的基础上改编过来的。为了避免这三章内容的重复,第10章和第11章仅对法语和瑞典语两种语言中与英语语法有差别,且读者可能感兴趣的部分给予了简要的介绍。总起来看,第9至11章介绍的是英语、法语和瑞典语三种语言各自语法的覆盖情况。

第12章介绍的则是用于基于 QLF 翻译的双语转换规则的覆盖情况。这一章较完整地描述了三种手工编制的转换规则集的覆盖情况:英语→法语、英语→瑞典语、瑞典语→英语。该章的作者(Manny Rayner, Pierrette Bouilon 和 Ivan Bretan)考虑到可能大多数读者对法语的了解程度多于对瑞典语的了解,因此,这一章的大多数例子都是以英法转换规则为例,但是,其基本思想在这三个语言对中都是实用的。

第13章介绍语言数据的重用问题。这一章探讨的基本思想非常简单,而且无可争议:从某种意义上讲,所有的自然语言都是相似的,因此,处理某一种语言的软件在一定程度上应该适用于处理其他语言。也就是说,如果语言 L1 和 L2 足够相似,那么,直接利用处理 L1 的系统处理 L2 一定比专门为 L2 从头编写一个处理系统更容易。这就是软件的重用问题。对于语言数据的利用来说存在同样的问题。该章介绍了 SLT 项目中语言数据的重用方法,包括两部分:一是语法和词汇的重用方法;二是转换规则的重用方法。

## 2.4 第三部分:语音处理

第三部分介绍 SLT 项目的语音处理技术,包括第14章至第19章的内容。

SLT 项目中语音识别模块的早期版本是 SRI 开发的 DECIPHER<sup>TM</sup> 系统,后来使用的是 DECIPHER<sup>TM</sup> 的改进版——Nuance 语音识别系统。

DECIPHER<sup>TM</sup> 和 Nuance 都是基于隐马尔可夫模型 (Hidden Markov Model, HMM) 的非特定人、大词汇量、连续语音识别系统。

第 14 章对 HMM 和基于统计方法的语音识别技术给予了简要介绍。当时本书编写时 HMM 还没有被广泛地应用于自然语言处理,编者增加这一章的目的就是希望 HMM 能够很快地被用于自然语言理解和语言建模中。当然,今天 HMM 早已不再是自然语言处理领域陌生的工具。

第 15 章介绍了 SLT 系统中语音识别模块的声学模型——连续密度的 HMM (continuous density HMMs, CDHMMs) 及其相关技术。SLT 原型系统的第一阶段 (SLT-1) 中使用了 CDHMMs 声学模型,但是其复杂性太高,难以达到实时性要求,因此,在第二阶段 (SLT-2) 系统中增加了很多提高速度指标的技术,包括提高概率计算和假设搜索速度等。本章对 CDHMMs 和其改进技术及最终的系统性能给予了简要介绍。

第 16 章介绍语音识别中的语言模型,主要介绍语言模型中的数据稀疏问题处理方法和面向多语言语音识别器开发的语言模型实现技术。

第 17 章介绍对新语言的语音识别器移植技术。SLT 项目的第二阶段 (SLT-2) 实现了第一个连续语音的瑞典语语音识别系统,该系统基于第 15 章介绍的 CDHMMs 模型及其相关技术。因此,这里存在一个原有系统向新的识别语言 (目标语言) 移植的问题。本章首先介绍了用于声学模型和语言模型训练的语音语料收集情况,然后介绍了目标语言单词的发音数据收集方法,最后介绍了利用这些语音语料训练声学模型和语言模型的基本方法。

说话人的口音适应问题始终是实现一个说话人无关的语音识别系统必须面对的基本问题,而对于多语言口语翻译系统来说,输入语音的语种识别技术也是系统不可缺少的。因此,第 18 章专门介绍 SLT 系统中语音识别器对不同口音的适应性处理方法和多语言语音识别器的实现方法。

第 19 章关注 SLT 系统中与语音和语言处理部分同时相关的一些问题和方法,包括语音识别模块 DECIPHER<sup>TM</sup> 与语言处理模块 CLE 的接口问题、由于某些语言 (如瑞典语) 允许使用多种形式表达复合名词而引起的识别问题以及在语言翻译中尝试利用输入信号中音素信息的方法等。

## 2.5 第四部分:系统评估与结语

第四部分包括第 20 章和第 21 章两章的内容。

第 20 章详细介绍了整个 SLT 语音翻译系统的实验评估方法。本章的重点是介绍多个语音识别系统在未知语音数据上测试的识别性能。

本书的最后一章第 21 章作为结束语再次归纳了 SLT 项目所取得的主要成果和技术进展。

另外,本书的两个附录分别是区别式打分方法的数学描述和基于 QLF 处理方法中的一些具体标识符号及其关系定义。

## 3. 对读者的建议与预期

口语翻译是集口语语音识别、口语理解和翻译以及语音合成技术于一体的综合应用系统,既有语音识别、机器翻译和语音合成三项核心技术本身的问题,又有语料(包括语音和语言数据)收集与加工、系统集成和人机界面设计等方面的问题。本书对 SLT 口语翻译项目中这些问题的处理方法给予了较全面的介绍,内容丰富,通俗易懂,相信读者能够通过阅读本书对口语翻译的基本原理、系统实现中存在的pecific问题和相应的处理方法有较全面的了解,同时希望 SLT 项目研究中的一些经验能够为从事口语翻译理论研究和系统开发的读者提供有益的参考。需要提醒读者注意的是,本书的首版时间是 2000 年,书中所介绍的内容是 1992 年至 1999 年大约 7 年间 SLT 口语翻译项目的研究工作,而这一时段恰好是语音技术(包括语音识别和合成)、机器翻译和自然语言处理相关技术迅速发展的重要时期。在这一时期,隐马尔可夫模型(HMM)、语言模型、最大熵等一批统计模型和机器学习方法在语音识别和自然语言处理领域得到了广泛研究和实践,尤其在自然语言处理领域,1990 年代初期基于理性主义的分析方法与基于经验主义的统计方法同庭抗争的纷杂局面逐渐转变成为以统计方法为主、分析方法和统计方法相融合的和谐发展阶段,这一根本性的转变极大地促进了自然语言处理技术的发展。伴随语音识别技术和统计机器翻译技术飞跃式的发展,口语翻译研究取得了显著进步。因此,从某种角度讲本书介绍的内容不能代表目前口语翻译研究的

最新技术状况,尤其对于致力于口语翻译理论研究和技術开发的读者来说,延伸阅读近几年来的相关文献是非常必要的。

#### 4. 延伸阅读指南

自上个世纪80年代中后期人们致力于口语翻译研究以来,口语翻译无论在理论方法上还是在系统实现上,都取得了长足的进步。概括地讲,这些进展大致可以归纳为如下几个方面:

(1) 系统识别和翻译的词汇量逐步扩大,由研究初期系统只能识别和翻译几百个源语言词汇增长到目前的几千个或上万个(特定领域),甚至基本没有词汇量限制(非特定领域)。

(2) 系统可处理的句型和语音已基本口语化。在语音翻译研究的初期,系统对输入语句的结构和发音方式一般都有较严格的限制和约束,系统只能识别相对标准的发音和处理相对规范的句型,而目前的系统一般都是针对日常口语会话的,系统对输入语句的句型和说话人的语音都不需要有严格的限制,口语识别和理解的鲁棒性(Robustness)、正确率都有了较大的提高。

(3) 翻译方法开始走向多元化和集成化,译文质量有了较大的提高。目前的口语翻译系统往往不是采用单一的翻译方法或单个翻译引擎,而是采用多种翻译方法和多个翻译引擎进行融合的实现策略。尤其基于大规模真实语料的统计翻译方法的迅速崛起和发展,口语翻译系统的鲁棒性和可移植性有了较大的提高,同时,大大缩短了系统开发的周期。

(4) 多模态交互式口语翻译系统正在被研究。为了提高口语理解和翻译的正确率,人们已尝试将对话情景知识,诸如说话人的手势、表情等信息引入到口语翻译系统中,并采用说话人与系统交互式方法以达到帮助系统理解说话者意图的目的,从而提高对话翻译的可理解率和正确率。

无论如何,口语翻译是一项极具挑战性的高难度、多学科交叉技术,面临许多难以克服的困难和障碍,尽管目前已经取得了令人可喜的成果,但仍有大量的难题有待于进一步研究。归纳起来,这些问题主要包括如下几个方面:

(1) 口语的声学特性分析有待于进一步加强,以提高语音识别的鲁