



生成式人工智能时代的 语料库语言学研讨会

Symposium on Corpus Linguistics
in the Age of Generative AI

会议手册

主办：中国英汉语比较研究会语料库语言学专业委员会

北京外国语大学中国外语与教育研究中心

承办：外语教学与研究出版社

协办：《语料库语言学》杂志

2026年8月9日—8月11日

中国·北京

参会须知

各位老师，为方便您安排参会期间的学习和生活，会务组就有关事项作如下说明：

- ① 用餐:请凭餐券在大宴会厅用餐。餐券当日当餐有效，遗失不补。用餐时间：早餐07:00—08:00，午餐12:00—13:30，晚餐17:30—19:00。
- ② 请遵循用餐规范要求，厉行节约，会议期间严禁饮酒。严禁编造、散布、传播小道消息，不得通过自媒体等发布可能引发负面舆情的图片、言论等。
- ③ 严禁扰乱会议秩序，会议期间禁止摄像、摄影、录音，严禁泄露会议涉密内容、涉密文件，未经主办方同意，不得公开会议文件资料。
- ④ 参加会议应遵循会议议程和会议规定，按议程进行合理发言和讨论。
- ⑤ 会议期间请将手机保持静音状态。

如有违背上述情况，主办方自发现之日起有权取消相关人员参会资格，并有权要求相关人员赔偿对主办方造成的损失。

主办方享有本会议所有解释权。

感谢您的支持与配合，预祝您参会愉快！

生成式人工智能时代的语料库语言学研讨会会务组
2026年8月

目 录

CONTENTS

主办、承办及协办单位简介.....	2
会议总日程.....	7
主旨发言专家简介及发言摘要.....	9
分论坛日程.....	17
分组发言摘要（以第一作者姓氏音序排序）.....	27
邀请函.....	99

主办、承办及协办单位简介

中国英汉语比较研究会语料库语言学专业委员会简介

中国英汉语比较研究会语料库语言学专业委员会于2009年4月16日成立于上海交通大学。在2009年4月16日至2015年4月19日期间，学会名为“中国语料库语言学研究学会”，系中国修辞学会下属二级学会。2015年4月19日正式成为中国英汉语比较研究会（China Association for Comparative Studies of English and Chinese，简称CACSEC）下属二级学会。按照民政部规定，研究会现官方名称为“中国英汉语比较研究会语料库语言学专业委员会”，英文名称为Corpus Linguistics Society of China，简称CLSC。

专委会致力于开展各种学术研讨和交流活动，包括以下活动形式：每两年举办一次中国语料库语言学大会（Corpus Linguistics in China，简称CLIC），会议称为“第X届中国语料库语言学大会”，主办者为“中国英汉语比较研究会语料库语言学专业委员会”，承办单位和地点由专委会单位成员申请，通过理事会审定批准。专委会不定期举办语料库语言学圆桌论坛，为小规模专题论坛（Corpus Linguistics Open Colloquium，简称CLOC）。承办者及地点向理事会提出申请，并经理事会审定通过。

2026年3月29日，中国英汉语比较研究会语料库语言学专业委员会线上召开理事会会议。经全体理事投票表决，产生了新一届语料库语言学专业委员会理事会。第三届语料库语言学专业委员会组织架构如下：

顾问委员会：

名誉主任委员：杨惠中、卫乃兴、梁茂成

委 员：何安平、李文中、刘泽权、陈建生、张宝林

主任委员：许家金

副主任委员：邓耀臣、刘国兵、潘 璠、濮建忠、甄凤超

总 干 事：刘国兵

常 务 理 事：董 敏、胡春雨、雷 蕾、娄宝翠、陆 军、秦洪武、邵 斌、邢富坤

理 事：陈 功、戴光荣、邓 飞、高 霞、郭曙纶、黄立波、姜 峰、金碧希、金 檀、孔 蕾、李晶洁、李文超、林 铃、林维燕、刘鼎甲、刘运锋、钱毓芳、饶高琦、苏 杭、孙海燕、王德亮、王 丽、熊建国、熊文新、徐 昉、徐曼菲、杨林伟、杨 娜、张 超、张 懂、张瑞华、张 威、张 毓、邹艳丽

语料库语言学专业委员会官方网站：<https://clsc.bfsu.edu.cn>

北京外国语大学中国外语与教育研究中心简介

北京外国语大学中国外语与教育研究中心成立于 2000 年，是教育部人文社会科学重点研究基地，也是北京外国语大学国家重点学科“外国语言学及应用语言学”的主体单位。

中心以外语教育研究、语言本体研究、语料库研究、语言政策研究为主攻方向。主要任务是：1）承担重大科研项目，产出重要研究成果；2）培养高级研究人才，造就创新学术梯队；3）提供政策理论咨询，推动国内外学术交流。

中心成立以来，完成和在研的科研项目共计 100 余项，其中国家级项目 30 余项（含国家社科基金重大项目 4 项）。中心还荣获教育部和北京市人文社科优秀成果奖近 10 项，研究人员发表论文、提交咨询报告和出版专著超过 1000 篇 / 部，举办全国性、国际性学术会议百余场。

中心设有外国语言学及应用语言学、语言政策与规划学、外语教育学三个二级学科博士点、硕士点，是我国高级外语研究人才的重要培养基地。中心还常年接受培养国内外高校的访问学者、硕士 / 博士交流生和留学生。中心开设的专业和课程内容丰富新颖，紧跟国际前沿。各层次学生在导师指导下，视野开阔、思想活跃、勤奋钻研，并积极参与导师主持的科研项目。每年有多名毕业生获学校优秀博士 / 硕士论文奖，曾有多位博士生荣获北京市和全国优秀博士论文奖。

中心承办《外语教学与研究》《外语教育研究前沿》《语言政策与规划研究》三种 CSSCI 来源期刊和集刊。此外中心研究员还主编 *Chinese Journal of Applied Linguistics*、*Journal of World Languages*、《语料库语言学》《英语学习》及内部刊物《世界语言战略资讯》。

中国外语与教育研究中心官方网站：<https://sinotefl.bfsu.edu.cn>

外语教学与研究出版社简介

外语教学与研究出版社（简称“外研社”）由北京外国语大学于 1979 年创办，是以外语教育出版为特色，国内领先、国际知名的综合性文化教育出版机构。外研社以“记载人类文明，沟通世界文化”为使命，每年以 80 多种语言出版类别丰富的图书和期刊，并积极探索教育服务转型与数字化融合创新，构建了以教材出版为中心，覆盖教、学、测、评、研全流程的一站式教育服务解决方案。

服务国家战略。外研社出版《习近平总书记教育重要论述讲义》多语种版、《教育强国建设规划纲要（2024—2035 年）》多语种版，传播中国教育发展新思想新理念新观点；出版“理解当代中国”多语种系列教材，将习近平新时代中国特色社会主义思想系统融入外语类专业和大学外语课程教学；出版《大国语言战略》和“一带一路”国家文化教育大系等，促进世界文明交流互鉴。

引领教学改革。外研社教材涵盖基础教育、职业教育、高等教育领域，“以学生为中心”“产出导向法”“跨文化思辨”等理念首开先河，数百种教材不断再版，有力推动大中小学外语课程教学改革。外研社“四新”大学英语系列教材全新推出，聚焦“四新”各领域交叉学科以及国家紧缺专业方向，系统培养学生学科思维与科学素养，提升其在国际舞台讲好专业故事的能力。

推动数字赋能。外研社持续建设 U 校园 AI 版平台和 UMOOCs 慕课平台，不断优化 iTEST、iWrite、iTranslate 等智能教学工具，打造 iPublish 数字教材出版平台和 iTeach 教师数字素养平台，发布全球语言服务平台，建设“外语+”智能课程，推动人工智能赋能外语教学、测评、科研与实践，创新教学空间，切实赋能施教者、驱动学习者、提效管理者。

完善教学评价。外研社精研测评解决方案，研发人才评价标准，为高质量外语教育赋能增效。优诊学、iTEST 智能测试云平台等，以练促学、以测促教；国际人才英语考试、国际人才行业英语考试、VETS 实用英语交际职业技能等级考试等，对接语言测试与人才需求，以测评反拨院校育人，优化企业用人体系。

强化教师支持。外研社积极推动教师发展，通过“理解当代中国”全国大学生外语能力大赛和“教学之星”大赛等传递育人理念，展现育人成果；通过虚拟教研室、教学开放周、POA 云教研共同体、外语教育研究新星培育云共同体、AIGC 辅助教学创新提效实战工作坊、“行业+外语”双师双能型教师培训等为教师提供多元化、个性化服务。

推动教育研究。外研社发布中国外语教育基金项目、中国外语教材研究专项课题，通过《外语教学与研究》《外语教育研究前沿》《专门用途外语研究》等期刊设置专栏，推动理论与实践创新，构建产研融合、协同发展的新生态。出版外语学科核心话题前沿文库、“翻译中国”研究丛书、外国语言学及应用语言学学科史丛书等学术著作，展示最新研究成果，探索未来研究方向。

面对新形势与新要求，外研社将认真贯彻习近平新时代中国特色社会主义思想，全面落实立德树人根本任务，系统研究课程思政视角下的外语教育出版，还将积极探索建立科学评价体系和新型数字化内容与服务体系，全力支持外语教师课程改革与教学创新，推动高层次外语教育研究队伍建设，为新时代我国外语教育发展和外语人才培养作出新的更大贡献。

《语料库语言学》简介

为适应语言学与计算机科学以及人工智能之间交叉融合的发展趋势，北京外国语大学中国外语与教育研究中心于 2014 年创办《语料库语言学》杂志，并于 2017 年成为中国英汉语比较研究会语料库语言学专业委员会会刊。该刊物是亚洲地区最早创办的语料库语言学专业杂志。

创刊十年来，《语料库语言学》先后刊登 Tony McEnery、Douglas Biber、桂诗春、杨惠中、卫乃兴、梁茂成、秦洪武、许家金等国际、国内前沿语料库学者的科研成果。

杂志目前常设：研究论文、研制开发、书刊评介等。

该杂志被 CSSCI 来源期刊引用次数在数十本语言学集刊中位列前 10。

2024 年，《语料库语言学》荣获中国知网 2024 年度高影响力学术辑刊，位列获奖的 20 本语言类学术辑刊的第五名。

近年来，《语料库语言学》在办刊方面采取了诸多积极举措。例如，（1）全面采用通讯作者制；（2）详细标注多人合作论文的科研贡献；（3）2023 年起，明确杂志主编在任期间不在该杂志上发表成果；（4）组织学术沙龙和公益讲座。

刊物先后被中国知网、中国人文社会科学期刊综合评价指标体系 AMI、维普数据库、人大复印报刊资料收录入库。

欢迎学界同仁关注并积极投稿，支持《语料库语言学》的发展，共同推进人工智能时代语言学的创新发展，服务学科，服务社会，服务国家战略。

《语料库语言学》投稿网址：<https://ylyy.chinajournal.net.cn>

会议总日程

报到

日期	时间	地点
8月9日	14:00—20:00	北京外国语大学国际会议中心大堂
8月10日	8:30前	

2026 年 8 月 9 日（星期日）

时间	事项	地点
19:00—21:00	中国英汉语比较研究会语料库语言学 专业委员会理事会会议	培训楼101

2026 年 8 月 10 日（星期一）上午

开幕式 地点：第一多功能厅		
时间	议程	主持人
08:30—08:50	特邀嘉宾致辞 ——中国英汉语比较研究会语料库语言学专业委员会创会会长 卫乃兴教授	北京外国语大学许家金教授
	特邀嘉宾致辞 ——中国英汉语比较研究会语料库语言学专业委员会名誉会长 梁茂成教授	
08:50—09:10	合影	
主旨报告 地点：第一多功能厅		
时间	报告人及题目	主持人
09:10—09:45	AI-Assisted Corpus-Driven Linguistics: Technological Advances and Linguistic Innovations ——北京航空航天大学 卫乃兴教授	北京航空航天大学董敏教授
09:45—10:20	提示词优化策略探索 ——北京航空航天大学 梁茂成教授	华中科技大学潘璠教授
10:20—10:40	休息	
10:40—11:15	基于生成式大语言模型的术语自动抽取研究 ——大连外国语大学 邓耀臣教授	河南师范大学姜宝翠教授
11:15—11:50	语言复杂度的权衡与等复杂度假设 ——上海外国语大学 雷蕾教授	上海交通大学甄凤超教授
12:00—14:00	午餐午休	

2026 年 8 月 10 日（星期一）下午

分论坛发言		
时间：2026 年 8 月 10 日 14:00-17:30		
分论坛主题	主持人	地点
分论坛一：大语言模型与智能语料标注（一）	邹艳丽、刘运锋	培训楼201
分论坛二：大语言模型与智能语料标注（二）	熊建国、李文超	培训楼203
分论坛三：大模型语言能力与人机语言对比研究	孙海燕、陈功	培训楼205
分论坛四：智能翻译、翻译质量与译者风格研究	张威、王丽	培训楼207
分论坛五：多模态语料库与智能传播研究	孔蕾、张超	培训楼208
分论坛六：智能学术语篇与语言教学研究	徐昉、林铃	培训楼301
分论坛七：智能构式、局部语法与词汇语义研究	姜峰、张懂	培训楼303
分论坛八：智能媒体话语、国际传播与社会议题（一）	钱毓芳、金碧希	培训楼305
分论坛九：智能媒体话语、国际传播与社会议题（二）	邓飞、杨娜	培训楼307
分论坛十：智能语料资源、语言变异与跨学科应用	王德亮、杨林伟	培训楼309
17:30—19:00	晚餐	

2026 年 8 月 11 日（星期二）上午

主旨报告		
地点：第一多功能厅		
时间	题目	主持人
09:00—09:35	从多语种文体计量到大模型语用识别的探索与实践 ——国防科技大学 原伟教授	浙江外国语学院邢富坤教授
09:35—10:10	局部语法视角下大语言模型言语行为能力评估 ——上海外国语大学 苏杭教授	浙江大学邵斌教授
10:10—10:30	休息	
10:30—11:05	大语言模型的资源、主体和媒介角色 ——北京语言大学 饶高琦副研究员	广东外语外贸大学胡春雨教授
11:05—11:40	大语言模型辅助的跨模态图文意义建构量化研究 ——北京外国语大学 刘鼎甲副教授	扬州大学陆军教授

闭幕式

时间	议程	主持人
11:40—11:50	大会总结发言 ——中国英汉语比较研究会语料库语言学专业委员会总干事 刘国兵教授	北京外国语大学黄鼎博士
12:00—14:00	午餐	

主旨发言专家简介及发言摘要

卫乃兴 教授

北京航空航天大学

亚太语料库语言学协会创会成员、中国英汉语比较研究会语料库语言学专业委员会创会会长。现任北京航空航天大学外国语学院教授、学术委员会主任。主要研究兴趣包括语料库语言学、短语学、局部语法、话语分析等。先后主持国家社会科学基金项目 5 项；出版《词语学要义》等专著 7 部。在国内外重要期刊发表论文百余篇。



AI-Assisted Corpus-Driven Linguistics: Technological Advances and Linguistic Innovations

The sweeping developments of generative AI have mounted serious challenges for linguistics while also opening up new possibilities for its innovation. In this talk, I set out to address the main features of AI-assisted Corpus-driven Linguistics, which I propose and intend to regard as a potentially promising approach to corpus linguistic studies in future years. I start with a brief discussion of the convergences and divergences of AI research and corpus linguistics in the past fifty years, trying to spell out the sameness and differences between the two neighbouring disciplines. I argue that while AI research and corpus linguistics share a host of analytic techniques and instruments, it is their fundamental differences in first principle thinking that set them apart and drive them to different orientations. Secondly I move on to a discussion of the major technological features of the current generative AI enterprise, such as the nearly total population corpora and the algorithms of so-called next token prediction, among others. Thirdly and finally, I focus on the important features of the AI-assisted corpus-driven approach to linguistics, which broadly include huge authentic corpora data for research, minimal assumptions about language, alignment of AI data and linguistics and cross-border falsification, among other things.

梁茂成 教授

北京航空航天大学



北京航空航天大学外国语学院教授、博士生导师、院长，国家“万人计划”哲学社会科学领军人才，国务院政府特殊津贴专家，现任国务院学位委员会外国语言文学学科评议组成员，中国英汉语比较研究会语料库翻译学专业委员会副会长等职。主要研究兴趣涉及语料库语言学、数据科学、应用语言学等，主持省部级以上项目十余项，著有《中国学生英语作文自动评分模型的构建》《大规模考试英语作文自动评分系统的研制》《语料库应用教程》（合著）、《什么是语料库语言学》等，发表各类学术论文近百篇；开发各类语言数据处理和语料库分析工具 30 余款，是 iWrite 英语作文智能评阅系统（外研社）和韩素音国际翻译大赛自动评分系统（中国翻译协会）的设计人。

提示词优化策略探索

大语言模型的输出质量高度依赖提示词设计，但编写高质量提示词本身需要任务所涉及领域的专业知识、合理设置大语言模型的温度（temperature），并熟练掌握提示词工程技巧。因此，设计高质量提示词对普通用户完成复杂任务构成较大挑战。本研究介绍一种提示词优化策略。研究中利用大语言模型，由用户提供的提示词推导任务涉及领域所需的人设（persona），并通过设置温度范围搜索最佳温度取值，以此对用户提示词进行量化诊断和全面优化。实验表明，研究中提出的提示词优化策略能够系统提升提示词质量和下游任务的完成度。

邓耀臣 教授

大连外国语大学



大连外国语大学教授、博士生导师,《外语与外语教学》主编,享受国务院政府特殊津贴。研究方向为语料库语言学理论及应用、计算术语学。主持完成国家社科基金一般项目 2 项,国家社科基金重大项目子课题 1 项,参与国家社科基金重大项目,国家社科基金、教育部社科基金一般项目多项。在 *Applied Linguistics*、*International Journal of Corpus Linguistics*、*Journal of Quantitative Linguistics* 以及《外国语》《中国外语》《上海翻译》《外语教学》《外语与外语教学》等国内外外语类权威期刊发表论文 60 余篇,主编出版学术专著 3 部。

基于生成式大语言模型的术语自动抽取研究

生成式人工智能技术的发展为术语自动抽取研究带来了新的契机,但针对基于大语言模型的术语自动抽取的可行性及优势,尚缺乏系统的实证探讨。本研究通过术语自动抽取实验,从实证角度系统评估大语言模型在术语自动抽取中的优势与潜力。结果表明,基于生成式大语言模型的方法在术语自动抽取性能上显著优于传统方法。研究同时发现,相对于大语言模型的架构差异,参数规模和提示设计的差异是影响基于生成式大语言模型术语自动抽取效率的重要因素。具体而言,当大语言模型参数规模达到 2000 亿时,术语自动抽取达到最佳效果;而在提示设计中包含专业领域信息能够有效提高术语自动抽取效率。本研究结果有望为人工智能赋能的术语自动抽取理论探索与实际应用提供重要参考。

雷蕾 教授

上海外国语大学

博士，上海外国语大学教授、博士生导师。研究兴趣涉及基于语料库和计量方法的现代汉语、古代汉语、英语的词汇 / 句法描写及历时研究、学习者语言研究、语言数字人文研究等领域。在剑桥大学出版社等出版专著 5 部，在 *Applied Linguistics*、*Journal of Second Language Writing*、*Corpus Linguistics and Linguistic Theory* 等 SSCI 期刊发表研究性论文 60 余篇、书评 10 余篇，其中两篇论文入选 ESI 高被引论文；在 CSSCI 期刊发表论文或书评 10 余篇。主持国家社科基金项目 2 项。兼任 *Journal of English for Academic Purposes* (SSCI) 等国内外期刊编委、*Corpus-based Studies across Humanities* (De Gruyter) 副主编。入选爱思唯尔中国高被引学者名单。

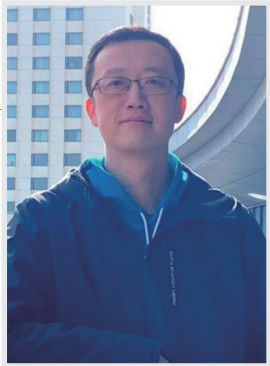


语言复杂度的权衡与等复杂度假设

复杂度权衡假设 (complexity trade-off hypothesis) 认为，某一语言领域 (如词汇层面) 难度的增加，往往会通过另一领域 (如句法层面) 的简化来抵消。这一概念植根于等复杂度假说 (equi-complexity hypothesis)。该假说认为，尽管人类语言呈现多样性，但所有语言在总复杂度上均维持相近水平。这一原则对于理解语言如何受认知约束及高效交流需求的影响具有重要意义。本次发言将首先回顾语言复杂度权衡的相关研究，然后展示团队的最新研究成果，并期望这些成果能为理解语言复杂度权衡及其研究提供新视角。

原伟 教授

国防科技大学



国防科技大学教授、博士生导师、博士后合作导师、外国语学院俄罗斯中亚西亚语言文化系主任，人才工程青年学者，中国英汉语比较研究会语言服务专业委员会理事、外语教育技术专业委员会理事，教育部国家语委“国防语言能力发展研究中心”专职研究员，江苏省“国防语言智能处理”军民融合创新平台专职研究员。精通乌兹别克语、俄语和英语，主要从事计算语言学、计量语言学和数字人文研究，主持国家社科基金项目 3 项、国家社科基金重大项目子课题 2 项、中国博士后基金一等及特别资助项目 2 项、省部级课题 2 项及其他各类课题 30 余项。近年来出版专著 3 部，主编出版词典 4 部及教材 5 部，获得江苏省优秀成果 1 项，获得软件著作权 6 项，在 SSCI、A&HCI、EI、CSSCI、CSCD 等各类期刊发表论文 50 余篇。

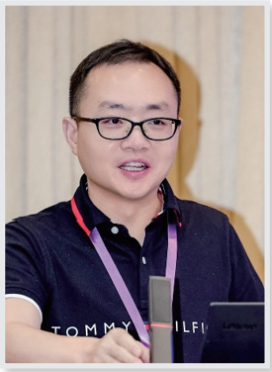
从多语种文体计量到大模型语用识别的探索与实践

在全球化与数字化背景下，网络空间的信息失序（如虚假新闻与仇恨言论）给多语种自然语言处理带来严峻挑战，而构建高质量的多语种标注语料库是驱动技术突破的核心。依托文体计量学与传统机器学习，本研究对英俄双语网络失序信息进行了系统量化与自动识别，并在特定数据集上取得良好成效。然而，传统计算方法高度依赖人工特征工程，跨语言迁移成本高，且难以捕捉反讽、隐喻等深层且隐蔽的语用特征。大语言模型凭借强大的跨语种泛化与深度上下文理解能力，为突破传统方法瓶颈带来了革命性转机。但在引入大模型作为自动化标注器时，其推理过程的“黑盒”属性导致生成结果具有随机性与不可解释性。因此，探索如何将大模型主观的语言生成过程转化为客观、可复现的标准化标注流程，已成为当前语言科学与多语智能交叉领域亟待解决的关键命题。

苏杭 教授

上海外国语大学

伯明翰大学博士、北京航空航天大学博士后，现为上海外国语大学语言科学研究院教授、博士生导师。研究兴趣包括语料库语言学、系统功能语言学、语用学、学术英语等；主持完成国家社科基金、中国博士后基金等科研项目三项，出版著作两部，在 *Applied Linguistics* 以及《外语教学与研究》等期刊发表论文 50 余篇。学术兼职包括中国英汉语比较研究会功能语言学专业委员会常务理事、中国英汉语比较研究会语料库语言学专业委员会理事、*Journal of English for Academic Purposes* 编委等。曾入选“重庆英才·青年拔尖人才”支持计划、第四批重庆市学术技术带头人后备人选，获得重庆市社会科学优秀成果奖二等奖（2024）、高等教育（研究生）国家级教学成果奖二等奖（2023）以及重庆市高等教育教学成果奖三等奖（2022）等奖项。



局部语法视角下大语言模型言语行为能力评估

大语言模型在语言研究与教学中的应用受到广泛关注，相关研究涉及功能语篇分析、礼貌知识生成、反语识别以及外语教学等领域。然而，当前研究较少考察大语言模型的语用能力，评估大语言模型言语行为能力的研究更是少见。基于此，本研究采用局部语法路径，依托其“形式—功能”对齐及透明化功能标注的分析优势，尝试评估大语言模型言语行为能力。研究表明，局部语法有助于将评估从大语言模型“能否实施”某类言语行为，推进至其“如何实施”“能否恰当实施”以及“功能结构是否完整”的机制层面，进而为提升大语言模型语用能力提供语言学依据。

饶高琦 副研究员

北京语言大学

北京语言大学语言科学院副研究员，语言资源高精尖创新中心兼职研究员，硕士生导师，中国标准化委员会语言与术语技术委员会 SAC TC62 委员、中国中文信息学会青年工作委员会委员。主要研究领域涉及语言资源建设、语言规划和数字人文等。在国内外期刊、会议发表论文 70 余篇，起草多项团体标准和国家标准。创办并主编语言学公众号“汉语堂”。



大语言模型的资源、主体和媒介角色

进入数智时代，大语言模型成为人机共生的语言生活形成和发展的重要部分。在这一前提下，大语言模型扮演了语言资源、语言主体和语言媒介三重角色。大模型因为“占有”海量语料，从而成为现实语言生活的“压缩包”（荀恩东教授比喻），自身也具有了语言资源的属性（有用性、稀缺性、可交换性）。大模型因为其生成能力、理解能力而直接进入语言生活；互联网中甚至一半以上的内容直接出自大模型，人机协同内容生产成为行业主流，其主体性已成为现实。在此基础上，伴随大模型及其驱动的智能体渗入生活，大模型具有了语言媒介的属性。训练语料中，现实语言生活无法接触的语言现象被大模型投送到更广泛的用户端。而其作为中介也勾连起了更多人、机生产的内容。三种角色互相影响，彼此促进，值得今天的语料库研究者和数据科学家深入思考。

刘鼎甲 副教授

北京外国语大学



博士，北京外国语大学中国外语与教育研究中心、国家语言能力发展研究中心副教授、专职研究员，《语料库语言学》杂志副主编，兼任中国英汉语比较研究会语料库语言学专业委员会、语料库翻译学专业委员会理事。主持国家社科基金、国家语委重点课题等多项科研项目，出版专著 1 部，发表论文近 30 篇。主要研究兴趣为语料库语言学、语料库翻译学、计算语言学与语言数据库建设，代表作有《大数据视角下大规模英汉平行语料库的加工、检索与应用》。

大语言模型辅助的跨模态图文意义建构量化研究

语境共选观强调意义产生于语言单位与词汇、语法、语义及语用环境之间反复出现的选择关系。在多模态语篇中，图像与文字的意义建构同样体现为符号资源在特定语境中的系统共现。本文以药品包装正面为研究对象，基于社会符号学与系统功能语言学元功能理论，从概念再现、人际互动和语篇组织三个维度考察图文协同构义。研究构建了涵盖中国、美国、欧洲和日本四个市场、四类非处方药的多模态语料库，共收录 512 件包装和 3106 项文字单元，并借助大语言模型完成初步识别与分类，经人工校正后分析视觉过程、情态、构图、显著性、框架与文字功能、位置及显著性，进而揭示模态内和模态间的意义共现格局。结果显示，药品包装主要采用叙事性图像、中性情态、“上下”式信息结构、品牌凸显和清晰框架。文字模态则集中承担产品识别、适应症说明、规格标识和监管认证等功能。不同地区与品类在信息密度、情态取向、框架强度和文字配置方面呈现差异，展现出有明显区别的跨模态意义组合。研究表明，图文意义源于视觉与语言资源的系统共选和功能互补，为多模态语篇的语料库量化研究提供了可验证的分析路径。

分论坛日程

分论坛一：大语言模型与智能语料标注（一）

序号	作者姓名	发言题目	发言时间
1	孙文青、董 敏	Exploring a Large Language Model-Assisted Approach to Annotation of Appraisal as Co-Selection	14:00-14:15
2	余丹妮、喻儒辰、Marina Bondi、Ken Hyland	大语言模型辅助的语步标注优化：聚焦模型间的不一致标注	14:15-14:30
3	陈馨怡	Evaluating LLM-Assisted Appraisal Annotation and Identity Construction in Oscar Acceptance Speeches	14:30-14:45
4	王 婕	大模型辅助下的美媒移民话语趋近化自动标注与人机协同分析	14:45-15:00
5	杨传鸣、孔祥瑞、郭云洁、刘语轩	大语言模型辅助化学期刊英文学术论文摘要语步结构识别的研究	15:00-15:15
6	李志慧	基于短语学视角的大语言模型隐喻自动标注方法研究——以TED商务文本为例	15:15-15:30
7	陆娇娇	社交媒体文本的自动化叙事标注 ——一项使用大语言模型的探索性研究	15:30-15:45
8	张可人	大模型辅助语料库标注框架的构建与应用：以车载语音英中葡三语质检语料为例	15:45-16:00
9	敖乌日力格	大语言模型辅助低资源语言语料标注的效能与局限——以蒙古文为例	16:00-16:15
10	刘育含	生成式人工智能辅助科技语体“的”字功能标注研究——基于小型历时语料的初步考察	16:15-16:30
11	郑宇熹	生成式人工智能辅助的上古汉语语法树标注生产系统	16:30-16:45
12	仲子懿	人机协同视角下语料库语步标注研究：基于本硕英语论文摘要的IMRD分析	16:45-17:00
13	黎顺苗	大语言模型辅助澳门中文词义标注研究	17:00-17:15

分论坛二：大语言模型与智能语料标注（二）

序号	作者姓名	发言题目	发言时间
1	张 懂、蔡拂晓	大语言模型的构式知识表征能力研究	14:00-14:15
2	李文超	A Pre-Trained-Model-Based Tool for Diagnosing Non-Native Japanese Learners' Oral Development	14:15-14:30
3	郭欣月	基于生成式AI的中亚留学生作文语义偏误自动识别与分类研究	14:30-14:45
4	李鹤林	大语言模型驱动下的构式语法知识建模与英语构式识别研究	14:45-15:00
5	耿文殊	Bridging Bottom-Up Usage and Top-Down Dictionary Inventories Through Generative AI	15:00-15:15
6	周穗芳	基于少样本学习的隐喻自动识别研究	15:15-15:30
7	赵 奕	人机协同的庭审问答语料库构建与语用功能标注研究	15:30-15:45
8	赵泽栋	From Statistical Filtering to Semantic Retrieval: An Automatic Pipeline for Metadiscourse Phrase Extraction	15:45-16:00
9	刘苗苗	生成式人工智能赋能语用现象识别：海外中文餐厅评论中编码负面情感的检测、分类与解码	16:00-16:15
10	徐铤伟	生成式人工智能驱动的人智协同双语语料智能分析系统构建研究——基于RAG与多智能体框架的实践探索	16:15-16:30
11	杨 茗	Can LLMs Parse Classical Chinese? A Prompting-Based Evaluation	16:30-16:45
12	廖 元	A DeBERTa-Based Approach to Automatic Metaphor Identification Under the MIPVU Framework	16:45-17:00

分论坛三：大模型语言能力与人机语言对比研究

序号	作者姓名	发言题目	发言时间
1	李佳蕾	Simplified by Algorithms, Explicated by Humans: The Paradox of Syntactic Complexity in Machine and Human Translation Across Genres	14:00-14:15
2	杨明托	系统功能语言学与人工神经网络的生成机制	14:15-14:30
3	陈新炜、郝美玲	AI时代，人类的语言表达多样性降低了吗？——来自豆瓣《霸王别姬》历时影评语料的证据	14:30-14:45
4	谢 敏、张 滢	Readability of Human-Written and LLM-Generated Summaries of Maritime Accident Reports: A Corpus-Based Comparative Study	14:45-15:00
5	李 琳	AI生成对话与汉语自然口语对话的多维度对比分析	15:00-15:15
6	尤 美	情感问答语域中大语言模型与人类语言特征对比研究	15:15-15:30
7	李颖慧	Can AI Construct Research Space? A Corpus-Based Study of Engagement in Human- and AI-Generated Research Article Introductions	15:30-15:45
8	何 荷	ChatGPT译文和人工译文的词汇丰富度和句法复杂度研究：以贾平凹《高兴》为例	15:45-16:00
9	刘彩晴、马 蓉	Exemplification in GenAI-Generated and Expert-Written Research Articles: A Local Grammar-Based Investigation	16:00-16:15
10	王 媛	Affordances and Constraints of Qualitative Corpus Analysis in the GenAI Age: Several Epistemological and Methodological Complexities from Simulated Intersubjectivity and Participatory Sense-Making in Second Language Education	16:15-16:30
11	黄雅诗	Constructing Interpersonal Meaning in Psychotherapy Discourse: A Corpus-Assisted Comparison of Human Therapist and ChatGPT-Generated Counseling Responses	16:30-16:45
12	郭嘉雯	评价理论视角下China Daily与DeepSeek生成京剧新闻报道的对比研究	16:45-17:00
13	华美淇	An Evaluation Study of Doubao’s Chinese Conversational Implicature Reasoning Ability Based on the SwordsmanImp Corpus	17:00-17:15

分论坛四：智能翻译、翻译质量与译者风格研究

序号	作者姓名	发言题目	发言时间
1	董 敏、许嘉杨眉、 刘 颖	基于大语言模型的文学翻译自动评估探索：语义—情感二维相似度计算	14:00-14:15
2	曹 婷、高晚晚、 陆 军	人工智能辅助的英汉被动式对比短语学研究	14:15-14:30
3	王 丽	Can Machine Translation Keep the Speaker’s Stance? Appraisal and Interpersonal Meaning in English-Chinese Renderings of TED Talks	14:30-14:45
4	龚集伟	基于BERT模型的《红楼梦》英译本海外传播比较研究	14:45-15:00
5	梁妍钰	生成式人工智能与人工口译的语料库多维特征对比研究——以联合国演讲文本为例	15:00-15:15
6	兰 博	基于语料库的《西游记》海外传播形象建构研究——以“三打白骨精”为例	15:15-15:30
7	徐道冉、翁义明	基于语料库的人机翻译对比研究 ——以中华学术外译项目文化类文本为例	15:30-15:45
8	田文杰	基于依存句法的大语言模型翻译质量优化研究	15:45-16:00
9	苏玥玥	How Do Student Interpreters of Different Proficiency Levels Handle Mode Change? A Corpus-Based Entropy Study on Consecutive and Simultaneous Interpreting	16:00-16:15
10	柳 玥	中华学术外译中语言学著作英译本的短语化模式及其规范趋同与差异——一项大语言模型辅助的语料库研究	16:15-16:30
11	王昭鉴	“一译三体”：大语言模型辅助的华兹生 <i>Basic Writings of Mo Tzu, Hsiün Tzu, and Han Fei Tzu</i> 译者风格研究	16:30-16:45
12	陈 霞	霸权话语的语料对冲：中非气候政策术语的多语平行语料库构建方法	16:45-17:00

分论坛五：多模态语料库与智能传播研究

序号	作者姓名	发言题目	发言时间
1	朱周晔	ChatGPT看图作文叙事特征研究	14:00-14:15
2	么 娟、王一漫	A Quantitative Analysis of Image-Text Dynamics and Strategic Choices in Multimodal Advertising	14:15-14:30
3	韩智巍	多模态语料库人机协同标注模式探索——以创意字幕翻译为例	14:30-14:45
4	赵明圆、张继光	“China Travel” 视频评论的话语建构：基于BERTopic与批评话语分析的研究	14:45-15:00
5	齐兆贤	局部语法框架下韩国媒体“朝鲜族”多文化报道的图文共选机制与多模态话语表征分析（2009 – 2023）：基于BERTopic主题建模与Vertex AI图片标注	15:00-15:15
6	吴 悠	多模态大语言模型视觉语法标注能力评估——以药品包装再现结构识别为例	15:15-15:30
7	相贺延	A Study on Human-Machine Collaborative Annotation of a Multimodal Discourse Corpus of Classroom Teaching Videos	15:30-15:45
8	朱宇吟、汪雄飞	生成式人工智能辅助下的多模态符际翻译与叙事节奏重构研究——以电影《银翼杀手》跨媒介转译为例	15:45-16:00
9	潘晓琪、李 琳	冰雪旅游专题多模态语料库的设计与建设	16:00-16:15
10	任卓璇	From Visual-Textual Relations to Linguistic Choices: A Probe into Chinese Print Advertising	16:15-16:30
11	吕兴克	Human-Machine Collaborative Annotation and Validation of Multimodal Corpora: Generative-Model-Driven Image-Text Logico-Semantic Relations in Digital News	16:30-16:45
12	汪雄飞、朱宇吟	生成式AI辅助的小说—译本—电影多模态隐喻标注研究——以《了不起的盖茨比》为例	16:45-17:00

分论坛六：智能学术语篇与语言教学研究

序号	作者姓名	发言题目	发言时间
1	徐迎春、姜宝翠	A Corpus-Based Comparative Study of Hedges in English and Chinese Research Articles	14:00-14:15
2	宋畏林、林 铃	Exemplification in University Student Assignments: A Local Grammar Analysis Across Levels Of Study	14:15-14:30
3	许晨琛、徐 昉	“Such Comparison Fails To...”: The Neglected Open-Ended Negatives in the Literature Review Section of Social Science Research Articles	14:30-14:45
4	于银燕、张伶俐	Rhetorical Moves and Lexico-Grammatical Features in Applied Economics Abstracts: An LLM-Assisted Corpus-Based Analysis	14:45-15:00
5	宋瑛明	Expressing Opinions in Spoken English across L1 Backgrounds and Task Types: A Human-LLM Collaborative Exploration	15:00-15:15
6	王冉然	“语料库+大模型”双轮驱动的垂直领域语言智能教学范式构建——以轮机英语为例	15:15-15:30
7	康 卉	Comparing the LLM-Assisted and Teacher-Annotator Feedback Modes: Impacts on Syntactic Complexity Development in EFL Writing	15:30-15:45
8	邓海龙	生成式人工智能兴起背景下学术摘要中的立场表达与作者身份建构：一项历时语料库研究	15:45-16:00
9	王 悦	Register Differences in the Syntactic Hierarchical Structure of L2 Chinese Written Production: A Synergetic Linguistic Perspective	16:00-16:15
10	王雅慧	基于句向量的英语二语学习者作文语篇衔接评估	16:15-16:30
11	王一凡	How Academic Writers Summarize: Towards a Local Grammar of Summarisation	16:30-16:45
12	闫鹏飞	大语言模型学术论文结论生成的语体特征对比及预测研究	16:45-17:00

分论坛七：智能构式、局部语法与词汇语义研究

序号	作者姓名	发言题目	发言时间
1	明志龙、董 敏	Revisiting the Extended Unit of Meaning Based on Dependency-Relation: A Distributional Semantic Approach to Semantic Prosody	14:00-14:15
2	张 懂、欧阳霖	Alternation Among Chinese Passive Constructions: A Multivariate Diachronic Analysis	14:15-14:30
3	林晨希	次短语的计算建模：基于非典型宾语结构的自然语言生成探究	14:30-14:45
4	马金辉	基于语义向量空间模型的语义空间演变历时研究——以“V the crap out of NP”构式为例	14:45-15:00
5	汤慧桃	Headedness, Overlap, and Salience: The Multi-dimensional Tension in English Blends	15:00-15:15
6	吴术燕、周 卫	基于语料库标注分析的汉语完句性问题研究	15:15-15:30
7	卢明远	“He’s Indeed Not as Handsome as Me”: Local Grammar Patterns and Semantic Prosody of Negative Comparison in Putonghua Chinese	15:30-15:45
8	王诚文	数智时代背景下汉语构式搭配分析的智能路径探索	15:45-16:00
9	郭俊杉	“勇”人文基因语料库构建研究——以《论语》中“勇”相关论述为例	16:00-16:15
10	陈静垚	Local Grammar and Construction Grammar Approaches to Discourse Acts in Academic Writing: The Case of Classification	16:15-16:30
11	邱音竹、张芸洁	Local Grammar of Cause and Effect in Economic Discourse: A Corpus Analysis of the Chinese Report on the Work of the Government	16:30-16:45
12	邱辛新	A Local Grammar of Relevance Statements Across Engineering and Medical Sciences Research Articles	16:45-17:00

分论坛八：智能媒体话语、国际传播与社会议题（一）

序号	作者姓名	发言题目	发言时间
1	胡春雨、肖龙鸿	“语料库一大模型”协同的企业原型叙事建构研究	14:00-14:15
2	钱毓芳	Constructing the “Other”: Representations of Traditional Chinese Medicine in the British Press (1800-1980)	14:15-14:30
3	王悦	数字时代的算法规训与语言生活治理——基于大语言模型春节交际互动的话语分析	14:30-14:45
4	刘文欣、边琳杰、杜易霖、 骆屹霄、鲁正豪、金碧希	中美主流媒体生成式AI伦理风险话语图景建构和中国应对研究	14:45-15:00
5	彭伊寒	算法深描：基于Vibe Coding的欧盟对华产业政策批评性话语分析	15:00-15:15
6	任俊强	语境中的“中国”：美国国情咨文涉华话语的语义聚类与框架分析（1972年—2026年）	15:15-15:30
7	林敏奋、谢钦	Narrating the Corporate Self Across Institutional Divides: A Corpus-Based Comparison of Huawei’s and Apple’s 2025 Annual Reports	15:30-15:45
8	林博泽、林敏奋	中美新能源车企ESG话语的搭配词对比研究——以比亚迪与特斯拉为例	15:45-16:00
9	林树义	非洲媒体报道中的中国AI形象研究——基于语料库的及物性分析	16:00-16:15
10	李海萍、马倩、周方媛	计算语料库辅助下新加坡智库涉华文本的话语历史研究	16:15-16:30
11	向珮源	Combining CDA and Topic Modeling: Analyzing the Evolution of Green Economy Media Discourse in China Daily (2001-2024)	16:30-16:45
12	王子玥	隐喻场景理论视域下的中美能源企业身份建构研究	16:45-17:00

分论坛九：智能媒体话语、国际传播与社会议题（二）

序号	作者姓名	发言题目	发言时间
1	张 岚、董 敏	A Local Grammar Approach to Evaluative Language in the Register of Fake News	14:00-14:15
2	卢婷婷	LLM辅助的RT巴以冲突报道标题关键形态研究	14:15-14:30
3	晋洋榕	RT以伊冲突报道标题的语法格关键形态分析	14:30-14:45
4	陈永丽	英国主流报刊关于网络安全的话语建构研究	14:45-15:00
5	柴昭甲	面向大学生心理健康的微博文本语言表达特征分析——基于社交媒体语料库	15:00-15:15
6	王旭佳	大模型辅助情感对话语料隐喻标注的人机对比研究——以中文线上心理咨询文本为例	15:15-15:30
7	何思卓	从婚姻话语看奥斯汀女性观的演变——基于AntConc的三期作品语料解读	15:30-15:45
8	任纹漩	A UFA-Approach to Diachronic Investigation into the Identity Construction of Veganism in China Daily	15:45-16:00
9	谷孟雪	基于语料库的美国媒体关于“犹太资本”的批评话语分析	16:00-16:15
10	刘宇欣	Negative Semantic Prosody in Weibo Discourse: A Corpus-Based Study of a Viral Public Controversy	16:15-16:30
11	韩小溪、张 媛	联合国场域下中美外交话语的对比研究（1971年—2025年）	16:30-16:45
12	姜 萱	美国主流报刊的光伏叙事：一项语料库话语建构研究	16:45-17:00
13	刘诗睿	Evaluative Collocational Patterns of Menstruation-Related Node Words in Chinese Social Media Weibo and Xiaohongshu: A Corpus-Based Study	17:00-17:15

分论坛十：智能语料资源、语言变异与跨学科应用

序号	作者姓名	发言题目	发言时间
1	梁佳益、林铃	Patterns of Thanking in Asian Englishes: A Local Grammar-Based Investigation	14:00-14:15
2	王德亮、孙钰涵	从语言资源库到互动知识库：生成式人工智能时代汉语对话语料库建设的探索与实践	14:15-14:30
3	金 檀、徐曼菲、李元科	利用“一针三库”平台开展语料库智能研究	14:30-14:45
4	王宇婷	真实生活融入导向的网络语言筛选——以流行语榜单（2013年—2025年）及《人民日报》用例为依据	14:45-15:00
5	杨鑫悦	帕金森病语言及言语障碍研究现状与发展趋势——基于Bibliometrix的可视化分析	15:00-15:15
6	郭梦溪、荀恩东	语料库共享开放视角下的数据、技术与服务：利用BCC建设个性化语料库	15:15-15:30
7	刘风光、柴宜林	商务语境指责回应话语策略及其社会语用理据	15:30-15:45
8	周潇然	藏东南区域语言关系计算研究——基于核心词的编辑距离算法与历史比较法的互补验证	15:45-16:00
9	马彩霞	人机协同语料库模式下多元评论话语特征探究——以“京师心理大学堂”公众号为例	16:00-16:15
10	宋婧婧	媒体高频词语频序轨迹的生命周期模式及话题变化机制——基于2010年—2024年词表的软聚类研究	16:15-16:30
11	陈苏姗	专题语料库的人机协同研究模式构建——兼论区域国别研究的方法优化	16:30-16:45
12	王 珊、王少茗	大语言模型在放置动词教学中的应用探索	16:45-17:00
13	陈伊扬、缪青青	中日韩高中英语教材文化呈现的比较研究——基于内容分析法的多国教材分析	17:00-17:15

分组发言摘要

(以第一作者姓氏音序排序)

大语言模型辅助低资源语言语料标注的效能与局限
——以蒙古文为例

敖乌日力格 内蒙古大学

低资源语言普遍面临语料规模有限、人工标注成本较高和专业标注力量不足等问题。大语言模型在上下文理解与少样本学习方面的能力，为低资源语言语料标注提供了新的可能，但其实际效能仍会受到具体语言结构和语料特征的影响。本文以蒙古文命名实体标注为例，选取具有典型蒙古文语言特征且易引发标注偏差的语料实例，重点考察大语言模型在零样本和少样本条件下对人名、地名和机构名等实体的识别表现。研究将从标注准确性、实体边界、类别判定和结果稳定性等方面进行比较，并结合蒙古文词形变化、格附加成分、固有名词与普通名词区分以及上下文歧义等现象，分析模型出现实体遗漏、类别混淆、边界识别偏差和格式不一致等问题的原因。本文旨在从语言学角度揭示大语言模型辅助蒙古文语料标注的有效环节与可靠性边界，并为低资源语言语料标注中的样本选择、结果校验和质量控制提供参考。

人工智能辅助的英汉被动式对比短语学研究

曹 婷、高晚晚、陆 军 扬州大学

英语被动语态与汉语“被”字结构等往往被视为对等结构，但实际对应关系错综复杂，其间的对应规律难以系统揭示。本研究以扩展意义单位模型和语篇衔接理论为框架，通过人工智能技术辅助的汉英平行语料库和英汉平行语料库数据对比分析，分别从形式、意义和功能维度考察相关被动结构译文的共选特征，尝试揭示英汉被动式的对应规律。数据显示，汉语高频被动结构（如“被”字句与“受”字句）英译时趋于高频共享 Be+V-ed 等 8 种主要类联接型式，“被”字句与“受”字句的英译类联接分布差异显著，中性/积极语义韵和消极语义韵分布因英文翻译结构不同而异；英语被动式汉译有“被”字结构等 12 种主要型式。其中，定语短语和主动式型式呈现出明显的中性/积极语义韵，“被”字句译文的消极语义韵比例高于中性/积极语义韵，其他型式的中性/积极语义韵仅略高于消极语义韵。分析表明，英汉被动结构的对应关系由核心动词、类联接型式、态度意义（语义韵）以及跨句语境信息等要素共同界定。其中，跨句语境信息影响因英汉语言的衔接机制不同而异。被动式与主动式的交叉对应受到语义韵和 / 或语篇衔接功能等因素制约。

面向大学生心理健康的微博文本语言表达特征分析 ——基于社交媒体语料库

柴昭甲 上海外国语大学

Web 2.0 的发展促进了以用户交互为主的在线社交媒体的产生, 以微博为主流的短文本社交平台中富含大量情绪信息, 用户通过发表观点、转发博文、评论内容等方式进行信息交互。由于字数的限制, 社交文本的情感内容更加集中, 其个体使用语言的表达方式在一定程度上可以反映用户的性格与情绪特点。通过分析大学生用户的微博文本语言表达特点, 可以为后续的智能化心理健康评估提供支持。因此, 本研究爬取了患有抑郁症的心理异常学生用户与心理健康学生用户的微博文本数据, 构建了社交媒体语料库, 分别使用不同方法从文字语言特征和非文字语言特征两个方面进行了分析。文字语言特征分为词汇层级和句子层级, 每个层级细分为不同的类别进行进一步考察。在非文字语言特征方面, 通过对比用户关于表情符号、标点符号和数字语言的使用, 分析用户语言表达特征。

Local Grammar and Construction Grammar Approaches to Discourse Acts in Academic Writing: The Case of Classification

陈静垚 四川外国语大学

This study integrates local grammar with construction grammar to offer a principled and function-oriented framework for describing realisations of discourse acts in academic writing. According to Goldberg (1995, 2006), the knowledge of a language is the knowledge of constructions. However, how to identify constructions associated with discourse acts has received limited systematic exploration. Classification is a fundamental discourse act through which academic writers organise knowledge, establish category membership, and represent hierarchical relations. Although previous studies have explored linguistic realisations of classification in academic discourse, a more systematic approach to teaching classification is needed. Such an approach should not only explain what classification means, but also show how classificatory meanings are realised through recurrent linguistic patterns in authentic academic writing. In this study, a corpus of applied linguistics research articles was compiled and instances of classificatory markers (e.g., classify, belong to) were retrieved. Local Grammar was then employed to map discourse-semantic elements onto their corresponding formal realisations, yielding recurrent form-meaning pairings that characterise the discourse act of classification.

These local grammar patterns were subsequently interpreted as candidate constructions and generalised into mid-level classification constructions. Finally, the identified constructions were organised into a hierarchical constructional network. Methodologically, the study demonstrates how local grammar can function as an analytical bridge between discourse acts and constructions, providing a principled and replicable procedure for identifying constructions of discourse acts in academic writing. The resulting constructional network extends construction-based approaches to EAP research and offers empirically grounded pedagogical resources for teaching classification as a discourse act.

专题语料库的人机协同研究模式构建 ——兼论区域国别研究的方法优化

陈苏姗 北京外国语大学

生成式人工智能的发展正深刻重塑人文社科的知识生产方式。面对多语种、多平台、多模态信息快速增长带来的挑战，区域国别研究传统依赖人工，在获取、组织与分析资料等方面已难以满足效率和规模等需求。专题语料库作为区域国别研究的重要知识基础，逐渐成为连接人工智能与区域国别研究的重要载体。大语言模型能够承担语料获取、清洗、分类、关联及初步分析等专题语料库建设工作，但其生成幻觉、文化偏差、事实可靠性不足以及价值判断缺失等局限，决定其难以替代研究者完成研究设计、知识解释与理论建构。因此，本文构建面向区域国别研究的“四环节人机协同语料库研究模式”，即“问题发现—智能处理—人工释义—协同建构”，形成以研究者主导、人工智能辅助的专题语料库协同工作机制，推动人工智能由工具应用转向服务区域国别研究的方法优化，为生成式人工智能时代专题语料库建设与区域国别研究提供研究思路参考。

霸权话语的语料对冲： 中非气候政策术语的多语平行语料库构建方法

陈霞 浙江越秀外国语学院

在英语霸权主导的气候治理话语体系中，中非技术合作面临多层级转译失效的困境。本文提出“殖民性隐喻筛查—多语效能校准”双重语料库构建框架，通过系统收集中国、欧洲、美国及非洲四方气候政策文本，建立首个聚焦中非气候技术合作的中—英—法—斯瓦希里语平行语料库（2010年—2024年）。研究突破传统单向对齐模式，设计三级权力标注体系：术语殖民性残留指数（CRI），用

于识别“技术转移”等概念的隐性二元对立隐喻；向量偏移值测算，利用嵌入空间量化“气候正义”等核心术语在不同语言间的语义扭曲度；本土接受度预标注，整合 UNEP 非洲生态谚语数据库实现文化语境编码。方法论层面形成四步处理链：多源政策文本抓取、斯瓦希里语长尾术语处理（借助 BERTurk 模型）、权力权重分配算法（综合制度性权力与传播渗透率）、可视化对比界面生成隐喻网络差异图谱。实证分析揭示三类典型转译断裂：中方“生态文明”在英译中集体主义语义显著流失，非方“Hali ya Hewa”在技术文本中被降维为天气现象，“碳中和”等概念在法语转译中遭基督教救赎叙事强加。本研究理论首创“语料对冲”概念，将批评话语分析转化为可计算参数；方法上提供低成本多语种政策术语处理方案，效率较传统人工标注大幅度提升；实践层面产出《中非气候技术传播敏感术语清单（初级版）》，为后续智能翻译系统开发奠定基础。未来计划扩展至基隆迪语等更多非洲本土语言，并探索区块链技术实现语料库动态权力校准。

AI 时代，人类的语言表达多样性降低了吗？ ——来自豆瓣《霸王别姬》历时影评语料的证据

陈新炜、郝美玲 北京外国语大学

AI 的普及正在改变人类写作与语言表达。已有实验研究主要关注 AI 辅助写作的短期即时效应，即比较个体在使用或不使用 AI 工具时生成文本的差异，并发现 AI 辅助写作可能增加不同作者文本之间的相似性、降低文本词汇多样性。然而，这类研究尚难揭示 AI 广泛进入社会生活后，人类语言表达本身是否发生变化。对此进行考察，有助于将 AI 写作研究从特定工具使用的即时效应，推进到 AI 时代语言表达生态变化的讨论。基于此，本研究以豆瓣电影《霸王别姬》的用户影评为语料，按发布时间划分为基线期（2013 年—2015 年）、前 AI 期（2019 年—2021 年）和后 AI 期（2024 年—2026 年），每年随机抽取 100 条影评，经预处理后最终纳入 899 条文本。研究从个体与群体两个层面考察语言表达变化：个体层面采用 MATTR50 衡量词汇多样性，群体层面采用多语言 Sentence-BERT 模型计算组内文本相似性，并通过非参数检验及 Bonferroni 校正后的事后比较检验时期差异。结果显示，基线期与前 AI 期在两项指标上均无显著差异，而后 AI 期的个体词汇多样性和组内文本相似度均显著高于前两个时期。结果表明，后 AI 期影评语言并未表现为个体表达贫乏，而是呈现出个体词汇多样性提升但群体表达趋同增强的趋势。本研究为理解 AI 时代语言表达多样性的变化提供了初步语料证据。

Evaluating LLM-Assisted Appraisal Annotation and Identity Construction in Oscar Acceptance Speeches

陈馨怡 北京航空航天大学

This study investigates how Oscar winners use evaluative language to perform public identity work, drawing on a corpus of acceptance speeches delivered between 2020 and 2025. Guided by Appraisal Theory and Bucholtz and Hall's (2005) identity tactics framework, the study uses LLM-assisted annotation with manual verification, and reflects on this method along the way. The LLM performs reliably on broad Attitude type classification, but its accuracy drops noticeably when the labels are theoretically motivated and identity oriented, since these require inference beyond lexical semantics. This pattern fits existing evidence that LLM annotation degrades as task complexity and theoretical granularity increase. Turning to the findings, Affect emerges as the dominant Attitude type, and it often supports authentication by helping to construct a credible and relatable winner persona. A diachronic tendency also appears. Earlier speeches draw more heavily on Judgement, while later speeches show a gradual shift toward Appreciation, a change that may reflect broader shifts in award speech conventions after the disruptions of the pandemic. Identity tactics surface through recurrent links between evaluative choices and target selection, with adequation and authentication as the most prominent orientations. Taken together, the findings show how award recipients navigate the tension between institutional recognition and communal belonging through patterned evaluative strategies.

中日韩高中英语教材文化呈现的比较研究 ——基于内容分析法的多国教材分析

陈伊扬 暨南大学 缪青青 西北师范大学

在全球化与区域一体化背景下，高中英语教材承载着培养学生跨文化能力的重要使命。本研究基于文化要素构成观（Moran, 2004），借鉴四维文化呈现分析框架（国别归属—文化主题—文化层面—呈现路径），采用内容分析法，对中日韩三国高中英语教材中的文化信息点进行系统提取、编码与量化比较。研究语料包括中国人民教育出版社的《普通高中教科书·英语》（2019版）必修1-3册、选择性必修1-4册（共7册），该版本在全国范围内广泛使用，在中国高中英语教材中较具权威性和代表性；韩国NE Neungyule出版社的《High School English》I - II册及《High School Common English》I - II册（共4册）——NE Neungyule是韩国英语教材出版领域最具影响力的出版社之一，该套教材

依据韩国教育部 2015 年修订的教育课程编制，广泛用于韩国高中英语教学；日本三省堂的《CROWN English Communication》共 5 册，三省堂是日本历史最悠久的教科书出版社之一，该系列经日本文部科学省审定，是日本高中英语教科书的代表性版本。国别归属分为本土文化、目标语文化、国际文化、无具体指向的文化及跨文化五类；文化主题参照 45 个三级主题词表进行分类；文化层面涵盖文化实践、文化产品、文化观念、文化个体、文化社群及综合类；呈现路径按教材板块（课文正文、语言练习、文化练习等）进行统计。研究发现：（1）在国别分布上，三国教材均呈现文化多样性，但本土文化占比差异显著——中国教材本土文化占比最高，日本教材次之，韩国教材最低；三国教材均以英美为目标语文化代表，但呈现频次不均衡。（2）在文化主题上，“跨文化交际”这一普遍文化主题在三国教材中均占主导，具体国别文化主题的分布则折射出各国不同的文化战略与教育导向。（3）在文化层面上，三国教材均以“文化实践”和“文化产品”为主，“文化观念”占比普遍偏低，反映出高中英语教材在深层文化价值观挖掘上的普遍不足。（4）在呈现路径上，课文正文是三国教材承载文化信息的最主要载体，但文化练习板块的文化承载量差异显著，其中中国教材的设计占比明显高于日韩。本研究揭示了东亚三国高中英语教材在文化呈现上的共性与差异，为外语教材的跨文化编写与国别比较研究提供了实证依据，并对高中英语教材如何优化文化配置、深化文化观念呈现提出了具体建议。

英国主流报刊关于网络安全的话语建构研究

陈永丽 天津科技大学

随着网络安全风险不断扩展至社会治理、企业运行、国家安全和国际关系，新闻媒体对网络安全议题的建构值得关注。本文以 2023 年—2025 年英国四份全国性主流报刊《泰晤士报》《每日电讯报》《独立报》和《金融时报》中 144 篇网络安全报道为研究对象，基于框架理论，结合内容分析与话语分析，考察其议题类型、原因归因、对策建议和信息来源。研究发现，英国媒体中的网络安全报道以社会议题和外交议题为主，技术议题占比较低，说明网络安全已被建构为涉及公共服务、企业责任、关键基础设施和国际冲突的复合性安全议题。报道在原因归因上强调黑客团伙和国家行为体，在对策上突出企业责任、技术改进、政府治理和国际合作。话语分析表明，英国媒体主要通过犯罪化框架、治理框架、军事化与国家化框架以及 AI 风险框架建构网络安全意义。本文补充了现有研究中对英国主流媒体网络安全话语关注不足的缺憾，也为理解数字化依赖、大国竞争和人工智能发展背景下网络安全议题的媒体建构提供了参考。

生成式人工智能兴起背景下学术摘要中的立场表达与作者身份建构：一项历时语料库研究

邓海龙 赣南师范大学

已有研究多讨论人工智能生成文本的特征、识别方法以及人工智能辅助写作的教学利弊，但对于生成式人工智能兴起前后，已发表学术论文摘要的写作方式是否发生变化，仍缺乏基于大规模语料的系统考察。本文以 2015 年 - 2025 年 Web of Science 收录的英文研究论文摘要为研究对象，构建了包含两千余万篇摘要的大规模历时语料库，考察学术摘要中立场表达和作者身份建构方式的变化。结果显示，近十年来英文论文摘要中的立场表达呈现出较为复杂的变化趋势。其中，模糊限制语明显减少，增强语仅略有下降，说明摘要写作中直接限定和保留判断的表达有所弱化，但并未简单走向更强的确定性表达。与此同时，被动结构使用减少，而以“本研究”“本文”等为核心的研究框架表达增多，表明学术摘要正在从传统的非人称表述逐渐转向以研究本身为中心的表达方式。部分变化在时间上与 2022 年以后生成式人工智能的广泛使用相重合，但本文并不简单认定二者存在直接因果关系，而是将其视为学术写作规范、发表实践和技术环境共同变化的结果。研究表明，在人工智能时代，学术英语写作教学不应只关注人工智能工具使用和文本检测问题，还应引导学习者理解学术写作中的立场调控、表达分寸和作者身份建构。

基于大语言模型的文学翻译自动评估探索： 语义—情感二维相似度计算

董 敏、许嘉杨眉、刘 颖 北京航空航天大学

本研究以刘慈欣《三体》（第一、二部）为语料，利用大语言模型及微调技术，对中文原文、ChatGPT 生成译文与专家译文的语义相似度与情感相似度进行两两比较，开展文学翻译自动评估研究。不同于 BLEU、COMET 等将译文质量简化为表层接近度或命题意义传递的做法，本文区分了“说了什么”（语义相似度）与“如何表达情感”（情感相似度）两个维度。在语义维度，我们采用多语言语义表示模型 BAAI/bge-m3，将文本映射至统一向量空间，并计算余弦相似度；在情感维度，我们使用经 BRIGHTER 多语言情感强度数据集微调的 XLM-RoBERTa-large 模型，构建了六维情感强度向量，并计算跨语言情感相似度，同时引入 Qwen3-14B 本地模型进行开放式情感标签提取，以补充固定情感类别可能无法覆盖的文学性情感线索。研究结果显示，AI 译文与原文的语义相似度与情感相似度（分别为 0.77 和 83.8）均高于专家译文（分别为 0.75 和 82.8）。这一结果表明，AI 译文在信息对齐、句法对应和表达

规范性方面通常表现更稳定 (Bahdanau *et al.*, 2015; Vanmassenhove *et al.*, 2021), 因此自动评估往往更易获得高分。专家译文则通常更倾向于灵活调整翻译策略, 如结构重组、语义显化与语气调节等, 并通过创造性重构实现情感与风格的再现, 因而相似度较低。本研究的发现也在一定程度上为译者主体性提供了更多的实证支撑。此外, 该结果可能更多反映了基于相似度计算的自动评估方法的偏向性 (Liu *et al.*, 2023), 而非文学翻译质量的真实呈现。本研究认为, 自动评估方法难以有效衡量文学翻译的创造性表达、风格迁移和文学性再现等复杂维度 (Karpinska & Iyyer, 2025)。在文学翻译评估研究中, 相似度计算指标仍需与人工评估框架 (如 BWS、MQM 等) 相互参照, 并拓展更多的评估维度。

Bridging Bottom-Up Usage and Top-Down Dictionary Inventories Through Generative AI

耿文殊 北京第二外国语学院

Word meaning in academic discourse is rarely uniform. A single polysemous word may carry different operative meanings depending on the conceptual orientation of the discourse community that uses it. Corpus-based and lexicographic traditions have approached this property of language from separate vantage points, the former describing usage inductively, with categories emerging from the data, the latter drawing on usage evidence to illustrate senses without a systematic account of how each sense's weight shifts across contexts. This study asks whether dictionary-defined sense categories, the most systematic account of a word's semantic potential, can be systematically connected to large-scale corpus evidence. If such a connection can be established, it becomes possible to ask how that semantic potential is selectively activated across academic disciplines. We address this question by developing SenseMapper, a framework that fine-tunes the pretrained language model ModernBERT through contrastive learning on dictionary-derived sense examples, producing embeddings that assign each corpus instance of a word to its most probable dictionary sense. Because dictionary entries supply only a handful of examples per sense, GPT-4 expands this training data, raising classification accuracy from 0.68 to 0.88 and macro F1 from 0.40 to 0.85 on held-out examples. Applied to 600,000 academic abstracts from six disciplines, the framework generates discipline-specific sense profiles. Most nouns maintain stable sense distributions across fields, while a subset exhibit systematic discipline-specific preferences. Because these profiles remain anchored to named dictionary senses rather than unsupervised clusters, a disciplinary preference can be named rather than merely detected as undifferentiated variation, as when family shifts from a social sense in the soft sciences to a taxonomic sense in the hard sciences. Generative AI supplies the connective tissue between a bottom-up, usage-based tradition and a top-down, dictionary-based one, without itself deciding which sense any instance receives.

基于 BERT 模型的《红楼梦》英译本海外传播比较研究

龚集伟 兰州交通大学

本研究基于 BERT 的情感、情绪分析与主题聚类方法，对《红楼梦》杨宪益译本与霍克斯译本的海外读者评论进行比较分析。研究发现：在情感和情绪分布上，两个译本的评价均偏向正面，具体情绪以愉悦为主。但霍译本整体评价更为均衡和积极，杨译本则呈现更集中的极端好评与差评。在主题分布上，两译本评论均关注情节、译文与图书体验，但杨译本评论还涉及出版形态相关内容，且读者购书物流体验更差，霍译本评论则更聚焦故事人物与背景等叙事要素。基于此，研究提出：中国经典文学的外译与传播更应注重归化与异化译本的协同流通，以满足不同读者群体的需求，并强化版本管理，完善国际物流支撑，通过涵盖翻译、出版与流通等多维度的系统工程，统筹推进，方能促进中国文学的国际接纳与长远影响。

基于语料库的美国媒体关于“犹太资本”的批评话语分析

谷孟雪 天津科技大学

近年来，随着美国政治极化和身份政治加剧，“犹太资本”相关话题在美国社会政治话语中获得了新的生命力。因此，本研究以 LexisNexis 数据库作为语料来源，基于框架理论与语料库语言学方法，考察美国媒体关于“犹太资本”的话语建构。研究发现，美国媒体话语中，“犹太资本”主要与政治权力、反犹主义阴谋论两类名词呈现显著共现关系，话语主要围绕“操控美国政治的工具”与“反犹主义阴谋论”两种话语框架；媒体的政治倾向及背后资本对其话语建构具有显著影响：右翼媒体主要围绕“操控美国政治的工具”框架展开，左翼媒体则集中于“反犹主义阴谋论”框架，而政治立场相对温和的媒体其话语框架呈多元化特征。本研究揭示了美国媒体对“犹太资本”的话语建构以及其与背后资本的关系，为相关研究提供了方法论参考。

评价理论视角下 China Daily 与 DeepSeek 生成京剧新闻报道的对比研究

郭嘉雯 中国矿业大学（北京）

本研究以评价理论为分析框架，对比 China Daily 真实京剧英文新闻报道（CD）与 DeepSeek 模拟

生成京剧英文新闻报道（DS）在评价性话语策略上的异同。研究以 CD132 篇真实新闻报道的标题为提示词，调用 DeepSeek-V3 API 生成对应的模拟报道，借助 Claude Sonnet 模型完成情感、判断、鉴赏三类态度资源的大规模标注，旨在聚焦于同一话题下人机生成新闻的评价性话语差异。数据结果显示，CD 与 DS 的态度资源在结构上较为一致，但均衡性存在显著差异。CD 新闻中情感、判断、鉴赏资源分布较均衡，多依托京剧演员、传承人等第一人称的方式呈现，整体呈现“以人为本”的人格化倾向；DS 高度集中于鉴赏资源，倾向于借助权威转述、群体反应及艺术细节描写实现正面评价，呈现“以物为本”的审美化倾向。这表明，大语言模型在模拟官方媒体话语时一定程度上可以把握情感价值取向，但尚难以模拟真实记者借助经验所实现的情感共鸣与文化认同建构。研究为大语言模型辅助对外文化传播写作及评价理论自动标注方法提供了实证参照。

“勇”人文基因语料库构建研究 ——以《论语》中“勇”相关论述为例

郭俊杉 首都师范大学文学院

人文基因是贮存人文功能的语言单位序列，主要包括显性人文基因和隐性人文基因。“勇”或“勇气”一直是值得探讨的内涵丰富的道德概念。为解决大模型理解汉语文言文中“勇”这一概念的问题，构建“勇”人文基因语料库，主要包括“勇”人文基因词库和句库。在“勇”人文基因词库方面，依据常用的古汉语辞典，以“勇”“敢”“胆”为语素构建“勇”人文基因正向候选词库，以“怕”“惧”“怯”为语素构建“勇”人文基因负向候选词库。在“勇”人文基因句库方面，以《论语》中“勇”相关表述为例，依据“勇的等级性”，把“勇”划分为：为己之勇和利他之勇两种人文表现。结合局部语法理论，以“勇”的意义为单位，对“勇”相关论述标注。具体有三个方面：一、显性人文基因“勇”表现为显性人格基因、显性行为基因和显性心理基因；二、隐性人文基因“敢”为代表表现为隐性行为基因和隐性心理基因；三、人文基因结构整合为“勇”表现为显性意指基因与隐性挑战基因的合一。此外，结合“勇”相关的细化标注维度，提出“勇气的局部语法”标注框架，构建“勇”人文基因语料库。

语料库共享开放视角下的数据、技术与服务： 利用 BCC 建设个性化语料库

郭梦溪、荀恩东 北京语言大学

语料库共享与开放的关键，不仅在于扩大语料规模或推动资源集中汇聚，还在于把分散在研究者

和专业群体手中的语言材料组织为可持续利用的语料资源，并进一步形成服务能力。本文以数据、技术与服务三个层面为分析框架，讨论 BCC（北京语言大学语料库中心）个性化语料库的建设逻辑与实践路径。文章认为，BCC 个性化语料库的核心并非为个人建立封闭的小型语料库，而是使语料拥有者能够在统一技术规范下自主组织语料、持续维护资源，并在条件成熟时进入更大范围的共享与服务体系。数据层强调语料类型、场景知识和元数据组织的差异；技术层强调通用格式、自动建库、索引生成和全过程支撑；服务层强调检索分析、接口调用和分层供给能力。BCC 个性化语料库通过 BCC 语料格式、LangSC 工具体系、索引机制和统一检索语言，将建库、索引和检索能力下放给语料拥有者，使分散材料先在本地图形成可检索、可维护、可调用的语料服务，并为后续共享开放预留接口。该实践证明，个性化语料库建设是连接私域语料自主组织与公共语料服务体系的重要路径。

基于生成式 AI 的中亚留学生作文语义偏误自动识别与分类研究

郭欣月 新疆师范大学

生成式人工智能在第二语言教学领域的应用潜力正受到广泛关注，然而现有研究多聚焦于语法偏误的自动识别，对语义层面偏误的探讨相对不足。本研究以中亚留学生汉语作文为对象，尝试探索生成式 AI 在语义偏误自动识别与分类任务中的表现。首先，构建一个小型中亚留学生作文语料库，涵盖不同学习时长的学生写作文本。其次，设计特定提示词，引导生成式 AI 扮演“汉语教师”角色，重点识别语法正确但语义不通的偏误句子，并对其进行分类。再次，将 AI 识别结果与人工标注结果进行系统对比，计算准确率与召回率等量化指标。最后，对 AI 识别出的典型语义偏误类型进行归纳，并从词汇概念义混淆、母语语义迁移、文化背景差异等角度分析其成因。研究发现，生成式 AI 在语义偏误识别方面具有一定潜力，但在边界判断和细粒度分类上仍存在局限。本研究将偏误分析的重点从语法层面延伸至更为复杂的语义层面，为探索 AI 在国际中文教育中的应用边界提供了实证参考。

联合国场域下中美外交话语的对比研究（1971 年—2025 年）

韩小溪、张 媛 山东师范大学

联合国大会一般性辩论是各国阐述核心立场、进行“话语斗争”的国际舞台。中美作为国际社会中的核心行为体，其外交话语演变反映了两国身份定位与战略叙事变迁，进而牵动着国际话语权力的消长。本研究基于 1971 年—2025 年中美两国在联合国大会一般性辩论的发言，结合语料库辅助批评话语分析方法，从关键词、搭配、核心概念语义和互文性切入，量化双方外交话语的历时变化。研究

发现，中方话语从反霸叙事转向发展合作，再推进至全球治理倡议；美方变化集中在“自由”这一核心概念的用法变迁及议题的局部调整。中方发言持续锚定《联合国宪章》，美方则更多援引宣言等多元文本。相关发现揭示了中美外交话语演化路径和话语策略的差异，展现了话语与权力、社会现实的互动关系，对理解当代国际秩序下的话语斗争、构建对外话语体系以提升国际话语权具有参照价值。

多模态语料库人机协同标注模式探索——以创意字幕翻译为例

韩智巍 北京外国语大学

近年来，多模态语料库在语言研究中的价值已得到学界广泛认可。然而，多模态语料的标注长期依赖人工完成，建库成本高、耗费时间长，成为制约多模态语料库规模化发展的重要因素。随着生成式人工智能多模态理解能力的迅速提升，多模态语料库的自动化标注正在成为可能。在此背景下，本文以创意字幕翻译的多模态分析为例，旨在探索人机协同的语料库标注模式。研究构建三层标注框架：首先由多模态大模型依据既有编码体系完成全样本自动标注；然后引入另一多模态大模型进行独立复编，通过跨模型一致性来评估标注结果的可靠性；最后将模型分歧样本提交人工评判。研究语料来自 YouTube 某频道 26 集视频，共包含 962 个编码对象。结果表明，大模型标注能够较好复现原研究的人工编码分布及主要统计结论，模型分歧主要集中于人类编码员同样容易产生争议的部分。研究认为，基于三层框架的人机协同标注模式有助于在保证信度的前提下显著提升标注效率，可为生成式人工智能时代多模态语料库的自动化标注提供借鉴。

ChatGPT 译文和人工译文的词汇丰富度与句法复杂度研究： 以贾平凹《高兴》为例

何 荷 北京外国语大学

生成式人工智能的快速发展给翻译研究带来了机遇和挑战。其中之一是挖掘出 ChatGPT 译文与人工译文的语言特征差异，以便反思利用大语言模型提升译文质量和译后编辑的效率。本文以汉学家韩斌英译的贾平凹《高兴》为例，借助 AntWordProfiler、AntConc、WordSmith 8.0 等工具考察 ChatGPT 译文和人工译文的词汇丰富度，采用 L2SCA 考察其句法复杂度。研究发现，在词汇丰富度层面，ChatGPT 译文的词汇数量、词汇密度和低频词的比率更胜一筹，文本信息含量更丰富；在句法层面，ChatGPT 译文在单位产出长度、从属结构数量、并列结构数量和短语复杂度方面优于人工译文，但是译文的可读性却遭到了削弱，一定程度上降低了译文的可接受性。在文学翻译的译后编辑中，译文审

阅者应更关注平衡译文的可读性和可接受性，增强译入语规范意识，弥补大语言模型对社会文化因素制约的钝化反应，更好发挥其在文学翻译中的“守门员”作用。

从婚姻话语看奥斯汀女性观的演变 ——基于 AntConc 的三期作品语料解读

何思卓 大连外国语大学

本研究选取简·奥斯汀不同阶段的 5 部代表性小说，选取其中涉及婚姻问题的部分篇章语料，采用定性分析的方法搜集与归纳奥斯汀小说中的相关婚姻话语，进而探索奥斯汀婚姻话语及女性观的发展脉络。研究表明：奥斯汀的早期婚姻话语由爱情、个性两方面组成，体现了奥斯汀女性自主意识的初步觉醒；中期婚姻话语体现了更多的道德性、责任性和规范性，意味着奥斯汀女性观进一步的成熟；晚期婚姻话语用平等、尊重、自我价值等词语贯穿全文，女性主体意识高度升华。AntConc 软件作为精准检索工具，很好地解决了人工文本分析中的片面性问题。奥斯汀女性观发展的脉络是由感性觉醒到理性完善的历程，这不仅是奥斯汀本人的创作认识，也体现了 19 世纪英国女性家庭觉醒的时代背景。

An Evaluation Study of Doubao's Chinese Conversational Implicature Reasoning Ability Based on the SwordsmanImp Corpus

华美淇 香港城市大学

Understanding the implicit meaning in Chinese social dialogue is crucial to realizing human-computer interaction. This study evaluates Doubao, a leading Chinese Large Language Model (LLM), using SwordsmanImp—the first multi-turn conversational implicature corpus derived from the classic Chinese sitcom *My Own Swordsman* (Yue *et al.*, 2024). This study also constructs a dual-paradigm evaluation framework that combines an open-ended free generation task with a Five-Alternative Forced-Choice (5AFC) selection task to systematically evaluate the performance of the model in terms of pragmatic generation, multi-turn understanding, selection bias and reliance on training data distribution, and compares it with the human baseline. The two experiments were devised to compare the implicature comprehension and response capabilities between humans and Doubao. Experiment 1 uses a 5-point Likert scale to score overall response quality along five

dimensions: character consistency, implicature accuracy, colloquial naturalness, contextual appropriateness and fluency, and compares the answer performance of Doubao and humans respectively. Experiment 2 presents five types of options (high-scoring human answers, literal answers, original scripts, model self-generated answers and contextual distractors), and measures the rate of consensus between Doubao and 12 native Chinese speakers. The results show that: (1) Doubao's average generation score is 2.86 (out of 5), which is significantly lower than the human score of 3.36 ($t(118) = 2.72$, $p = 0.008$, Cohen's $d = 0.62$), with major weaknesses in character consistency, implicature accuracy, and colloquial naturalness. (2) In the 5AFC task, native speakers reached a high degree of agreement on 68.2% of the items, while Doubao's agreement rate on these items is only 40.0% ($\chi^2 = 5.50$, $p = 0.019$, Cramér's $V = 0.61$). The model also shows obvious dependence on the original script (selection rate 20.8%) and self-preference (12.5%). (3) Typological analysis reveals that Doubao performs worst in high-context-dependent types (such as evasive utterances and polite refusals, score gap exceeding 0.7 points relative to humans), while the gap narrows for explicit types like humour and sarcasm. Therefore, this study concludes that although Doubao possesses basic pragmatic reasoning ability, there are still substantial qualitative and quantitative gaps compared with human native speakers in terms of cultural implicit meanings, character dynamics, and multi-turn strategic interaction.

“语料库—大模型”协同的企业原型叙事建构研究

胡春雨、肖龙鸿 广东外语外贸大学

学界对结合语料库语言学与叙事分析的研究兴趣日益增长，但传统方法普遍难以在海量文本中依托数据驱动实现原型叙事的建构。本文借鉴并整合“基于语料库的叙事分析 (Corpus-Based Narrative Analysis, CBNA)”等研究成果，创新性地提出“语料库—大模型”协同的企业原型叙事建构分析框架，并以《中国日报》(2016年—2024年)关于华为的英文报道为个案探讨该框架的适切性。研究发现，语料库的海量数据发现与大模型的深度语义理解等方面的协同赋能，可有效揭示该媒体报道高度契合“发生—矛盾激化—评价—结果”的经典叙事结构。在这一历时轨迹中，报道分别构建了“前沿科技的全球赋能者”“遭遇地缘政治霸凌的受害者”“利用规则维权的抗争者”与“实现底层技术重构的破局者”四大原型叙事。其连贯叙述了全球科技巨头华为如何在极限施压下一步步破茧成蝶、逆境重生，这不仅是对企业实践的语境化重构，更是在反霸权语境下重塑中国科技企业韧性与国家形象的成功话语实践。本研究不仅为“语料库—大模型”协同的叙事分析提供了具有可操作性的方法论借鉴，也为中国媒体提升国际叙事传播效能与话语权提供了重要参照。

Constructing Interpersonal Meaning in Psychotherapy Discourse: A Corpus-Assisted Comparison of Human Therapist and ChatGPT-Generated Counseling Responses

黄雅诗 西南大学

With the increasing use of large language models (LLMs) in mental health support, it is important to examine how counseling responses generated by LLMs differ from those produced by human therapists at the discourse level. Drawing on the interpersonal metafunction in Systemic Functional Linguistics, this study constructs a parallel corpus consisting of client turns, human therapist responses, and ChatGPT-generated responses based on 20 real psychotherapy sessions focusing on depressive emotions and related difficulties. Using corpus-assisted discourse analysis, the study compares the two types of responses from both macro and micro perspectives. At the macro level, it examines the distribution and combination of major counseling strategies, including questions, reflections, empathy, affirmations, advice, and information giving. At the micro level, it further analyzes interpersonal linguistic resources such as question types, reflection complexity, the realization of empathy and affirmation, mitigating devices in advice, and modality in information giving. The findings show that human therapists tend to demonstrate stronger sequential sensitivity and closer attention to clients' specific experiences. Their responses are often shorter and more closely tied to clients' immediate discourse, relying on questions, reflections, and brief acknowledgements to facilitate the unfolding of therapeutic interaction. In contrast, ChatGPT-generated responses are generally longer and more strategy-dense, frequently combining empathy, affirmation, explanation, and advice within a single turn. Although this makes ChatGPT responses appear structured and supportive, it may also lead to more generalized and less context-sensitive responses. These differences reveal contrasting ways of organizing counseling strategies and constructing interpersonal meaning. The study contributes to research on psychotherapy discourse and LLM-mediated mental health support, and provides linguistic implications for the responsible design and use of LLM-based counseling tools.

美国主流报刊的光伏叙事：一项语料库话语建构研究

姜 萱 天津科技大学

气候变化是全球治理的重要议题。2015年12月12日,《巴黎协定》的通过标志着全球气候治理

进入绿色低碳转型的新阶段，光伏产业在全球快速扩张。美国在光伏发展史中具有特殊地位，因此探寻其媒体关于光伏的话语建构具有重要意义。本研究以 NOW 语料库为数据来源，提取该库中 2015 年 12 月 12 日至 2025 年 12 月 31 日期间所有包含“光伏”的词语索引，自建“光伏词语索引语料库”，并依据该时段内全球光伏产业的关键政策与市场节点，将语料划分为三个历史子时期。研究采用语料库辅助的话语分析法，结合框架理论和批评话语分析理论，旨在考察不同时期美国媒体对光伏的话语建构及其嬗变。研究发现，美国媒体对光伏的建构经历了由技术话语到经济话语，最后到能源安全与竞争话语的演变。媒体在第一阶段主要通过技术突破、成本优化、发展前景框架等，将光伏建构为新兴能源技术；第二阶段通过能源转型、市场、中国双重形象框架等，将光伏建构为全球新兴产业；第三阶段通过储能协同、供应链风险框架将光伏建构为全球战略性产业，反映出美国能源政策调整、利益诉求及中美战略竞争加剧背景下的话语转向。

利用“一针三库”平台开展语料库智能研究

金 檀、徐曼菲、李元科 华南师范大学

本次发言介绍金檀教授带领的华南师范大学语料库智能研究团队研制的“一针三库”智能教研平台，汇报本团队近年来在“阅读难度调控研究”（金檀教授主研），“教材语篇整合研究”（徐曼菲教授主研）和“写作习得追踪研究”（李元科教授主研）等方面取得的成果。

RT 以伊冲突报道标题的语法格关键形态分析

晋洋榕 北京外国语大学

新闻话语不仅是对客观事件的镜像反映，更是意识形态与地缘政治立场的隐性建构场域。传统批评话语分析多聚焦于词汇与宏观叙事，忽视了屈折语中强制性形态特征的认知功能。本研究基于关键形态分析（Keymorph Analysis, KMA）与 Janda 的俄语语法格认知语义理论，以“今日俄罗斯”（RT）关于 2025 年 6 月以伊冲突的 1129 条新闻标题为语料，通过标准化残差分析，量化考察了核心政治实体的语法格分布模式。研究发现，RT 对不同阵营采取了差异化的语法格配置策略：在冲突当事方层面，两国军事机构（“ЦАХАЛ”与“КСИР”）均表现为主格显著超用，被塑造成冲突的显性主体；但在国家层面，以色列（Израиль）多通过属格隐匿为动作来源与关系背景，而伊朗（Иран）则通过与格和前置格被“被动化”与“场域化”。在域外国家层面，美国（США）的属格凸显揭示了其作为隐性结构力量的特征，俄罗斯（Россия）则借由主格的显著超用被塑造为积极斡旋的外交主体。

Comparing the LLM-Assisted and Teacher-Annotator Feedback Modes: Impacts on Syntactic Complexity Development in EFL Writing

康 卉 大连外国语大学

This two-year longitudinal study examines the differential effects of LLM-assisted versus teacher-annotator feedback on syntactic complexity development in EFL writing. Analysis of 35 complexity metrics across 397 student texts reveals two distinct developmental profiles. The LLM-assisted approach fosters a narrow, vertical, and robust pattern, which is marked by sustained growth in complex noun phrases and dependent clauses, with stability across task difficulties. The teacher-annotator approach promotes a broad, horizontal, yet less consistent pattern, which is characterized by balanced development across clauses and noun phrases, increased coordination, and greater susceptibility to task difficulty. These trajectories suggest a trade-off: LLM-assisted feedback concentrates and stabilizes gains in specific syntactic domains, while teacher-annotator feedback elicits broader but more variable development. The study provides empirical guidance for integrating LLMs into L2 writing pedagogy.

基于语料库的《西游记》海外传播形象建构研究 ——以“三打白骨精”为例

兰 博 西安外国语大学

在《西游记》海外传播研究中，“唐僧”这一形象长期被孙悟空所遮蔽。本研究以“三打白骨精”这一关键情节为分析样本，基于语料库方法，对比考察蓝诗玲、余国藩、詹纳尔三个当代英译本对唐僧形象的建构差异。通过构建态度表征、关系表征、伦理身份表征的三维分析框架，结合词频、搭配、话语特征等量化指标与定性解读。三个译本分别受学术忠实、文化外宣、市场导向的驱动，共同塑造了“唐僧”在英语世界“慈悲却蒙昧、权威而脆弱”的混杂形象。本研究拟为翻译形象学提供可操作的实证路径，对中国神话人物的海外传播具有方法借鉴意义。

计算语料库辅助下新加坡智库涉华文本的话语历史研究

李海萍、马 倩、周方媛 昆明理工大学

在中国—东盟高层战略沟通不断加深、区域地缘安全变局加剧的背景下，东盟核心智库的涉华态度是观察区域对华认知与政策动向的风向标。本研究尝试融合大模型时代的语义聚类技术（BERTopic）与传统语料库语言学方法，选取2013年1月至2026年4月期间，以“China”为核心词检索新加坡RSIS、ISEAS、SIIA三大权威智库，自建涉华专题语料库，共1868篇文本。目前，本研究已完成第一阶段的宏观量化分析。结果表明，新加坡智库涉华研究热点经历了从“区域综合安全”向“经贸机制”及“印太地缘博弈”的历时演化。涉华话语逻辑呈现出从早期“嵌套于东盟多边机制”向近期“聚焦大国双边博弈”的显性化转向，安全议题边界在四个历时阶段中持续扩容。作为正在推进中的研究，本研究第二阶段计划聚焦浮现的核心“安全”语义场，依托话语历史分析法，借助AntConc工具，开展微观的话语策略解码与历史互文透视。

大语言模型驱动下的构式语法知识建模与英语构式识别研究

李鹤林 北京外国语大学

构式语法将语言知识视为由形式—意义配对构成的复杂网络，为语言的结构化描述提供了重要理论基础，但构式之间的多类型链接（继承、相似、多义、隐喻）在现有构式库（如English Constructicon、CASACon等）中缺乏系统化的显式表达与计算表征手段。针对上述问题，本文采用OWL本体语言构建面向英语程度副词+形容词构式的知识图谱表征框架。在知识建模层面，定义垂直继承链接（inherits from/inherited by）与水平相似链接（same form、similar function等）两类核心边类型，以具体构式（如强化构式、过度构式、感叹构式）为节点，形成层级与关联的网络化表示。在知识应用层面，利用OWL推理机的传递性规则自动推导构式间的递归继承关系，并通过SPARQL查询实现同功能构式的聚类检索。以程度副词+形容词构式家族（如very good, too big, so cool）为案例进行实例化与可视化展示，验证了该框架对构式层级抽象与横向对比的表征有效性。本研究为构式知识可计算化提供了一条低门槛、可复用的技术路径，其成果可服务于智能语言服务与语言教学等领域。

Simplified by Algorithms, Explicated by Humans: The Paradox of Syntactic Complexity in Machine and Human Translation Across Genres

李佳蕾 北京邮电大学

Although machine translation (MT) has become popular in cross-cultural communication, some people criticize that MT is lexically polished but syntactically rigid. To examine why the syntactic patterns of MT are different from those produced by human translators, we compared 14 syntactic complexity measures in three corpora of MT, human translation (HT), and native texts (NT) based on translation universals. Our research shows opposing trends of syntactic simplification and explication in MT and HT, to varying degrees across news, general prose, academic texts and fiction. MT simplification is most prominent in general prose, followed by news, fiction, and academic texts. HT explication stands out in news and fiction. This challenges the traditional model of translation universals by revealing that simplification and explication are not mutually exclusive but can coexist in MT and HT. Simplification in MT is attributed to its embedded algorithms, whereas explication in HT constitutes a conscious cognitive strategy.

AI 生成对话与汉语自然口语对话的多维度对比分析

李琳 浙江外国语学院

基于大语言模型的生成式人工智能展现出的超强语言理解和文本生成能力引发了语言学界的广泛关注与深入讨论，越来越多的学者开始探讨大语言模型在语言研究与教学中的应用潜力。然而，现有研究多聚焦学术论文等书面语体分析，针对 AI 生成对话与人类自然口语对话的对比分析稍显不足。本研究基于一个 24 万字的汉语自然口语语料库，以 DeepSeek、豆包、Grok 和 ChatGPT 生成的对话文本为例，采用多维度分析法，通过统计分析选取了 20 个语言特征，从“即时交际性与信息具体化”“互动性”“叙事性”“口语表达的多样性与情感表达”四个维度对比分析了 AI 生成对话与汉语自然口语对话的共性和差异。研究表明，AI 生成对话文本较好地掌握了汉语口语交际的互动性和叙事性语言特征，但整体上倾向于使用相对单一的程式化表达方式，在反映口语交际的即时性和语言表达的多样性等方面与汉语口语对话尚有显著差异。本研究认为在应用 AI 生成语料时，要明确 AI 生成语料的优点和缺点，避免盲目地将 AI 生成语言当作人类语言的替代品来开展语言研究与教学。对于语言学者来讲，AI 时代更需要做好语言事实的挖掘与描写，尤其是对口语中语言规律细节的深入挖掘与描写。

大语言模型辅助澳门中文词义标注研究

黎顺苗 暨南大学

词义消歧是自然语言处理的基础任务。澳门中文作为典型的区域中文变体，其多义词兼具粤语、普通话、葡语语言接触形成的区域语义特征，传统人工标注存在标注成本高、标注质量参差不齐等问题。大语言模型的发展为区域中文的词义标注语料库建设提供了新的方法。本文采集 2009 年—2023 年《澳门日报》《澳门时报》《华侨报》等语料，涵盖新闻、体育、经济、娱乐、科技、金融等主题，选取高频多义词“飞”，探索大语言模型辅助澳门中文词义标注的可行性。“飞”在澳门中文兼具通用基本义、区域引申义，因此选取其作为词义标注实验的个案。研究参考 Cantonese WordNet 的义项划分，并结合语料的实际情况，构建目标词“飞”的人工词义标注基准，并设计不同提示策略，比较 GPT-5.5、Qwen 等主流大语言模型在澳门中文词义自动标注上的表现，评估大语言模型辅助澳门中文的标注能力。本研究旨在回答：（1）大语言模型能否辅助构建澳门中文词义标注语料库？（2）不同提示策略能否提升大语言模型对澳门中文区域词义的自动标注性能？（3）不同大语言模型在澳门中文词义标注中的表现是否存在差异？研究通过探索面向澳门中文的人机协同词义标注流程，评估大语言模型辅助低资源区域中文词义标注语料库建设的可行性。

A Pre-Trained Model-Based Tool for Diagnosing Non-Native Japanese Learners' Oral Development

李文超 浙江大学

This study develops an automated tool that combines a fine-tuned pre-trained speech model with natural language processing techniques to assess and track the oral development of non-native Japanese learners. Drawing on the International Corpus of Japanese as a Second Language, the study uses 1,755 recordings (46.5 hours; approximately 560,000 words) from learners across 16 countries and 12 first-language backgrounds, of which 1,000 were used to fine-tune the Transformer-based ASR model OpenAI Whisper for accent adaptation. The tool recognizes accented non-native Japanese and automatically transcribes and annotates pauses, fillers, repetitions, and reformulations. The tool re-operationalizes the Analysis-of-Speech unit on Japanese typological grounds—covering the te-form, topicalization, and clefts—and proposes fourteen computable measures across six criteria: fluency, pronunciation, content, grammar, and lexical and syntactic complexity. It provides learners with both a holistic score and detailed, criterion-by-criterion feedback, making it suitable for self-directed

learning as well as for teaching and assessment. Multi-faceted Rasch analysis yields a measured separation of 3.77 and a separation reliability of .93, with thirteen of the fourteen measures varying linearly across proficiency levels (A1–C2). Tool–human annotation agreement reaches an F-score of up to .996, and tool–human scoring agreement reaches a quadratically weighted kappa of up to .781, surpassing existing semi-automated tools. The tool reveals fine-grained developmental trajectories—a phonological-to-lexical shift in fillers, a U-shaped trajectory in lexical density, and a linear rise in syntactic complexity—together with L1-transfer patterns across learners of different backgrounds. By making these trajectories visible, the tool not only measures where a learner currently stands but also pinpoints what to practise next, offering a basis for diagnostic feedback, self-regulated learning, and the design of localized teaching materials.

Can AI Construct Research Space? A Corpus-Based Study of Engagement in Human –and AI–Generated Research Article Introductions

李颖慧 浙江工商大学

The integration of generative artificial intelligence (GenAI) into academic writing has intensified scholarly debate regarding its capacity to produce discourse aligned with established conventions of academic argumentation. Despite a growing body of research on lexical, syntactic, and discursal features of GenAI-generated texts, relatively little attention has been given to its role in the rhetorical construction of research space in research article (RA) introductions. In particular, the deployment of engagement resources in stance construction remains underexplored. Drawing on Swales' Create-a-Research-Space (CARS) model and Appraisal Theory, this study investigates how engagement strategies are deployed in human-written and GenAI-generated RA introductions, with particular attention to the construction of research space. Adopting a corpus-based comparative design, the study compiles two specialized corpora of Applied Linguistics RA introductions comprising human-authored texts and ChatGPT-generated counterparts produced under controlled prompting conditions. The analysis combines quantitative corpus techniques with qualitative discourse analysis of engagement resources, including Entertain, Attribute, Disclaim, Counter, Pronounce, and Concur. Particular attention is given to: (1) the distribution and rhetorical functions of engagement resources across different CARS moves; (2) variation in dialogic positioning and stance construction between human- and AI-generated texts. The findings are expected to provide empirical insights into the strengths and limitations of GenAI in academic discourse and offer implications for developing discipline-specific writing pedagogy that integrates AI tools effectively.

基于短语学视角的大语言模型隐喻自动标注方法研究 ——以 TED 商务文本为例

李志慧 北京航空航天大学

随着大语言模型 (LLMs) 在语言研究中的广泛应用, 其在隐喻识别中的潜力日益受到关注。然而, 现有研究主要聚焦于词汇层面的隐喻识别, 对篇章层面的多词隐喻表达及概念隐喻源域—目标域映射关注不足, 也缺乏与语言学理论深度融合的方法框架。基于此, 本研究以 TED 商务文本为对象, 从短语学视角出发, 将 MIPVU 隐喻识别程序转化为链式思维 (CoT) 提示语逻辑, 并结合 MSDIP 框架构建源域—目标域标注流程, 探索大语言模型辅助隐喻识别与概念域标注的方法体系。研究采用 GPT-5 进行零样本与少样本实验, 并以人工标注结果为参照, 对隐喻识别及概念域标注效果进行评估。初步实验结果表明, 融合语言学理论的设计能够显著提升大语言模型在篇章层面隐喻识别和概念域标注中的表现。研究旨在构建一种兼具语言学可解释性、可复制性与可扩展性的隐喻自动标注方法, 为语料库隐喻研究及大语言模型在语言学研究中的应用提供新的方法论路径。

Patterns of Thanking in Asian Englishes: A Local Grammar-Based Investigation

梁佳益、林 铃 上海交通大学

In recent decades, the study of varieties of English has gained prominence, especially as English evolves in diverse postcolonial contexts. Speech Acts such as thanking reveal culturally embedded politeness norms, yet cross-regional comparisons in Asian Englishes are limited. Drawing on data from the four components of the International Corpus of English (ICE), the study presents a local grammar-based approach to thanking patterns across the four varieties of English. It aims to advance our understanding of thanking and its associated socio-cultural norms and to enrich the toolkit for researching speech acts in world Englishes. Local Grammar is an alternative approach to linguistic analysis that focuses on one specific discourse function (Hunston and Su 2019). Instead of using grammatical terms such as Subject and Object, local grammar maps context-specific functional terminologies onto corresponding formal elements. Drawing on Su (2018), the data were annotated using functional labels such as Thanking, Benefactor, Beneficiary, Hinge, Intensifier and Specifier. In total, 22 local grammar patterns of thanking were identified and compared across varieties, with further analysis of patterns featuring intensification, specification and explicit benefactor realization. The present study reveals

both shared and variety-preferential features in thanking across the four varieties of English. IndE displays the most distinctive profile, favouring benefactor naming and suggesting a subtle shift that longer utterances are not necessarily perceived as more polite in this variety. HKE and SgE show a stronger alignment with BrE norms while retaining local preferences: HKE favours specifying reasons for thanking, whereas SgE shows minimal use of intensifiers. These findings suggest that thanking is shaped by both local politeness norms and target-language conventions, and highlight local grammar as an innovative corpus-analytic approach for quantifying and comparing speech act realisations across varieties of English.

生成式人工智能与人工口译的语料库多维特征对比研究 ——以联合国演讲文本为例

梁妍钰 长春理工大学

以大语言模型为代表的生成式人工智能技术迅猛发展，为语料库语言学与翻译研究带来了全新机遇。然而，针对 AI 生成语言与人类口译产出在政治话语中的量化对比研究仍较为有限。本研究聚焦大模型时代的机器笔译与人工口译的异同，旨在回答两个核心问题：（1）两者在宏观词汇与句法特征上有何显著差异？（2）在微观层面处理复杂外交修辞时，两者呈现何种策略偏好与局限？本研究依托计算机辅助翻译（CAT）课程教学实践，以联合国演讲为源语，收集 25 名外语专业学生的口译产出语料，选取“豆包”生成对应译文，构建“源语—学生口译—AI 笔译”复合型平行微型语料库。本研究借助 Gemini 作为智能分析工具，宏观量化测算两者的词汇密度与句法复杂度；微观层面则重点考察政治等效词及长难句切分的典型偏误。研究发现：（1）在宏观特征上，AI 译文在句法复杂度和词汇多样性上显著高于学生口译产出，呈现出更强的书面语体特征；（2）在微观策略上，学生口译在应对高认知负荷时多采用“显化”与“简化”策略，而“豆包”在处理深层外交隐喻与语境时，仍存在机械直译与语用流失的局限。本研究不仅客观评估了生成式大模型在特定政论文体中的表现，也为大模型辅助语料库分析及智能化翻译教学提供了实证参考。

A DeBERTa-Based Approach to Automatic Metaphor Identification Under the MIPVU Framework

廖 元 广东外语外贸大学

Metaphor is pervasive in everyday discourse: Steen *et al.*'s (2010) corpus-based study of roughly 190,000

words found that metaphor-related words (MRWs) account for 13.6% of running text. As generative large language models (LLMs) become widely available, corpus linguistics faces a methodological question distinctive to the GenAI era: should LLMs simply replace discriminative sequence-tagging models as the tool of choice for metaphor annotation, or does a more principled division of labor exist? This study builds an automatic metaphor identification system grounded in the MIPVU procedure, comprising two complementary pipelines for direct and indirect metaphors. For direct metaphors, gold annotations are scarce (fewer than 200 instances in VUAMC), so DeepSeek-V4-Flash generates zero-shot hard labels over the BE06 and OANC corpora, which are then used to train a DeBERTa-v3-large discriminative tagger (three-way mFlag/mrw_lit classification). For indirect metaphors, VUAMC's roughly 18,000 gold-annotated tokens support a two-stage knowledge distillation pipeline: a DeBERTa teacher first produces soft labels over BE06, which pre-train a student model before supervised fine-tuning on VUAMC. On the VUAMC held-out test set, the system achieved a sentence-level F1 of 88.03% for direct metaphor, marginally above the LLM zero-shot baseline (87.50%) under the same optimized prompt, while leading by 6–9 points on word-level metrics (mFlag F1, mrw_lit F1). For indirect metaphor, five independent random-seed runs yielded a stable All-POS F1 of $82.26\% \pm 0.25\%$, outperforming the best system from the 2018 VUA shared task by roughly 17 points. Domain-control contrastive experiments and a multi-layer contamination analysis are reported to make the system's evaluation boundaries and limitations transparent. These results suggest that, in the GenAI era, the value of LLMs may lie not in replacing discriminative models for dense prediction tasks, but in supplying reliable weak supervision at a fraction of the cost of manual annotation—precisely where discriminative models' efficiency, determinism, and reproducibility can be put to full use. Building on this, the study proposes a methodological framework in which annotation strategy is determined by gold-label availability, offering guidance for corpus linguistics practice that integrates generative and discriminative models.

中美新能源车企 ESG 话语的搭配词对比研究 ——以比亚迪与特斯拉为例

林博泽、林敏奋 福州大学至诚学院

ESG(环境、社会与治理)话语已成为企业阐述可持续发展承诺的关键场域,但现有研究多集中于金融与管理学视角,基于语料库方法的ESG话语研究仍较为匮乏,中美新能源车企ESG话语的跨文化建构对比更是鲜有涉及。本研究以比亚迪与特斯拉2024年度英文可持续发展报告为初始语料(后续拟扩展至多家中美车企),聚焦三个问题:(1)中美车企ESG话语的核心主题词有何差异?(2)核心节点词的搭配模式如何反映双方在环境、治理与社会责任维度的话语侧重点与语义倾向?(3)GenAI辅

助标注与人工标注在一致性与效率上有何差异？研究运用 AntConc 与 Sketch Engine 提取搭配模式，结合 MI 值与 T-score 筛选显著搭配，并引入 GenAI 工具辅助语义标注，系统对比人机标注的一致性与效率。初步分析表明，中美车企在 ESG 话语的主题聚焦上呈系统性分化：美方突出排放与供应链话语，中方侧重治理与发展话语，搭配证据进一步揭示双方在 ESG 三个维度的语义倾向存在显著差异。本研究填补新能源车企 ESG 话语跨文化对比的语料库研究空白，为 GenAI 辅助语料标注的实证评估提供参考，并为理解中美企业 ESG 话语中的文化差异与修辞策略提供语言学证据。

次短语的计算建模：基于非典型宾语结构的自然语言生成探究

林晨希 南京大学

汉语非典型宾语结构存在句法语义错配特征，现有工程视角研究仅停留在结构分类识别层面，基于“语义特征识别 + 数据库匹配”的生成逻辑存在底层预设偏差，无法支撑完整自然语言生成任务。研究基于“次短语”理论，明确该结构兼具短语可组合性与词汇凝固性的中间性质，否定纯词库化与常规动宾变体两种处理路径。在词条表征层面，将可进入该结构的单音节动词统一设为独立黏合式动词词条，标注 [+NONCANON-OBJ] 特征，绑定其可搭配的非典型宾语语义类型；宾语保留单一词条，将关系化、被动化等句法语义限制转移至黏合式动词，引入基于生活经验的匹配度机制处理表达可接受度的中间状态，同时明确不同宾语类型的形式约束，仅工具、时间、处所类宾语可带指示代词。在结构生成层面，于依存语法框架下建构新型动宾依存关系，收束双音节复合动词边界，抑制错误分词，在结构整体节点编码修饰限制，排除方式副词与“了”“着”体貌成分，保留频率副词与“过”的修饰合法性。最终构建完整规则系统，通过测试与现有模型对比，实现非典型宾语结构与常规动宾短语的精准区分，排除不合法表达。

Narrating the Corporate Self Across Institutional Divides: A Corpus-Based Comparison of Huawei's and Apple's 2025 Annual Reports

林敏奋 福州大学 谢 钦 闽江大学

Corporate annual reports are not merely financial disclosures but strategically crafted instruments for constructing corporate identity, managing stakeholder impressions, and legitimating organizational existence. How leading high-technology multinationals narrate themselves across divergent institutional systems, however,

remains under-described. This study addresses that gap through a corpus-assisted comparative analysis of the 2025 annual reports of Huawei and Apple, two leading technology corporations operating within state-coordinated and market-driven institutional environments respectively. Drawing on two self-built corpora of corporate annual reports(HW2025 vs AP2025), the analysis adopts a three-stage corpus analytical framework that integrates keyword extraction (log-likelihood keyness), collocation and semantic-network mapping (mutual information and t-score association, read through KWIC concordances), and Biber's multidimensional register analysis operationalized with the Multi-Dimensional Analysis Tagger(MAT). Huawei's discourse emerges as collective and first-person in reference and future-oriented in outlook, organized around technological innovation, ecosystem-building, and partnership, while remaining informational and technical-rational in register. Apple's discourse, by contrast, is markedly impersonal and risk-foregrounding, consistent with the stringent regulatory regime of the U. S. financial reporting. These profiles index contrasting legitimation strategies: Huawei advances a capability- and vision-driven narrative, whereas Apple enacts a responsibility- and risk-management one. The study contributes to institutionally embedded, cross-cultural business-discourse research and offers practical implications for Chinese multinational corporations seeking to strengthen their international corporate communication and global reach.

非洲媒体报道中的中国 AI 形象研究——基于语料库的及物性分析

林树义 北京航空航天大学

随着人工智能成为全球科技的新前沿，国内人工智能的海外传播与形象建构的研究价值日益凸显。然而，当前学界相关研究多聚焦西方媒体的 AI 叙事，对全球南方尤其是非洲媒体关于中国 AI 的报道关注不足。基于此，本研究旨在以非洲媒体为切入口，依托 LexisNexis 数据库采集非洲媒体有关中国人工智能的相关英文报道，并自建语料库。以韩礼德系统功能语言学及物性系统为核心理论框架，结合大模型辅助标注的研究方法，系统探究非洲媒体对中国人工智能的话语建构。研究发现，非洲媒体的报道核心聚焦中美 AI 领域竞争，中国 AI 发展的国际影响，以及中国 AI 产业、企业与技术三大维度。宏观层面，非洲媒体塑造了中国 AI 作为行业引领者、美国的竞争者、技术壁垒“破壁者”等多重形象；微观层面，国产大模型的代表 DeepSeek 被赋予技术企业主体与中国 AI 发展标杆的双重形象。在全球南方崛起、中非 AI 合作持续深化的背景下，本研究系统呈现了非洲媒体对中国 AI 的话语建构与形象塑造，能够为中国 AI 的发展、中非合作、中国科技形象传播等提供重要理论参考。

Exemplification in GenAI-Generated and Expert-Written Research Articles: A Local Grammar-Based Investigation

刘彩晴、马 蓉 江苏师范大学

Since late 2022, the rapid uptake of generative artificial intelligence (GenAI) in academia has made its linguistic features a nascent area of inquiry. Comparative research on AI-generated and human-authored academic writing has increasingly shifted from formal linguistic features toward discourse-level features and rhetorical functions, revealing systematic differences between the two. Yet the pragmatic realization of specific discourse acts remains underexplored. Accordingly, this study focuses on exemplification, a frequent discourse act in academic writing, and combines local grammar analysis with a discourse-functional perspective to develop a three-dimensional framework for systematically examining exemplification in ChatGPT-generated and expert-written Linguistics research articles. Specifically, the study identifies exemplification markers and local grammar patterns, traces their distribution across article sections (introduction, literature review, methodology, results, and discussion), and interprets their rhetorical functions through Hyland's (2000) knowledge construction-knowledge promotion distinction, with knowledge construction realized through concept elaboration, research illustration, and argument support, and knowledge promotion through reader guidance and authorial engagement. Results show that (1) both corpora are dominated by a few high-frequency canonical patterns, with the expert corpus showing a richer set of low-frequency long-tail patterns; (2) exemplification in the expert corpus clusters in the literature review and methodology sections, reflecting section-specific rhetorical needs, whereas the ChatGPT corpus displays a more even distribution and uses examples mainly as a general explanatory device; and (3) both corpora use exemplification mainly for knowledge construction, especially concept elaboration and research illustration, while the expert corpus more frequently extends to argument support and knowledge promotion. These findings help integrate local grammar analysis with academic writing research by incorporating rhetorical distribution and functional analysis, while offering empirical implications for EAP pedagogy and AI-text detection in the GenAI era.

商务语境中指责回应话语策略及其社会语用理据

刘风光、柴宜林 大连外国语大学

本研究基于社会语用学视角，探讨在线商务纠纷第三方评审语境下，商家指责回应话语策略的言语特征、互动结构特征及其社会语用理据。研究发现，在言语特征方面，商家频繁使用第一人称代词自指，而以机构称谓语指称消费者，建构交际双方语用距离，并选取原因说明、缓和等词汇表达实施辩解、否认等策略。在互动结构特征方面，商家回应指责后的互动发展呈现三种主要动态模式。以上两个维度的特征隐含说话人的社会语用动因，包括映射、体现及建构以维护一般性道德规范，此外，上述话语选择还具有维护规则性道德规范的考量。通过考察数字化时代商务话语实践的社会语用特征及理据，本研究为揭示社会行为者在多方互动中的权益维护与声誉管理提供有益参考。

生成式人工智能赋能语用现象识别： 海外中文餐厅评论中编码负面情感的检测、分类与解码

刘苗苗 华中师范大学

生成式人工智能的发展为语料库语言学带来新的理论与方法议题，其中之一是大语言模型对超越字面语义的语用现象的理解能力。本研究以海外中文使用者在 Google Maps 平台撰写的中文餐厅评论为语料，考察一类编码负面情感现象：评论者通过谐音替代、字形替代、反讽、藏头等策略表达真实不满，以规避平台对负面评论的自动审核机制。该现象对人类读者意义清晰，但对依赖表层语义的传统情感分析模型构成挑战，可作为检验大语言模型语用推理能力的样本。语料以 Wan 等（2026）构建的 CodedLang 数据集为基础，并结合中文网络语境补充本土化编码类型标注。研究方法分为三阶段：第一，构建“编码语言检测—编码策略分类—标准中文转写”三层标注体系，形成结构化平行语料；第二，基于该语料设计对应的大语言模型识别任务，分别评估模型在检测、分类与意图解码三项任务上的表现，并与 BERT 基线及词典方法对比；第三，结合评论文本中的语言学线索，探索性分析编码策略的使用是否与评论者作为海外中文社群成员的身份处境相关。研究表明，大语言模型在表层模式识别任务上可达到较高准确率，但在深层意图重建任务上存在局限，揭示了当前模型在语用推理层面的能力边界，并为编码语言现象的群体性语用动因提供初步证据。

Evaluative Collocational Patterns of Menstruation-Related Node Words in Chinese Social Media Weibo and Xiaohongshu: A Corpus-Based Study

刘诗睿 中国人民大学

Menstruation has long been socially constructed as a stigmatised and concealable bodily experience, and the language in which it is discussed reflects this stigma. This study employs corpus-based collocation analysis and semantic prosody theory to investigate the evaluative co-occurrence patterns of four menstruation-related node words: 月经 (menstruation), 姨妈 (aunt), 例假 (period leave), and 生理期 (physiological cycle) in Chinese social media Weibo and Xiaohongshu. The study draws on a specialised corpus of 2,428 user-generated texts (144,544 words) between January 2023 and February 2026 and analyzes the top 50 high-frequency collocates of each node word across eight sub-corpora and examines their evaluative orientations using a hybrid dictionary-driven (NTUSD sentiment lexicon) and manual annotation approach. The findings reveal that negative collocates dominate across all eight sub-corpora on both platforms, consistently constituting half or more of all high-frequency co-occurring terms, while positive collocates remain marginal. Negative collocates groups around two main areas: physical symptoms and expressions of aversion and shame. This suggests that menstrual stigma in Chinese digital discourse works in two ways, where bodily discomfort and moral judgment strengthen each other. Stigma is embedded in the broader collocational environment that surrounds menstrual vocabulary. De-stigmatising menstrual discourse requires sustained accumulation of new co-occurrence patterns and structural and multisectoral change.

中美主流媒体生成式 AI 伦理风险话语图景建构和中国应对研究

刘文欣、边琳杰、杜易霖、骆屹霄、鲁正豪、金碧希 北京航空航天大学

2022 年底, OpenAI 推出的 ChatGPT 是人工智能技术领域的一项开创性成就。然而, 这项技术的伦理影响日益凸显, 引发了社会公众的广泛关注。中美媒体对这一现象展开了激烈讨论, 形成了一个关键的跨国媒体话语交互场域。关于人工智能的讨论已成为全球话语领域中最突出的议题之一。通过基于语料库的分析, 研究中美媒体如何定义生成式人工智能的伦理风险, 并分析这些描述背后的意识形态动机, 对于解释该话语的内在特征具有实际价值。现有的人工智能话语研究主要基于“全球北方”的媒体环境, 这种环境与对技术伦理风险的批判性审视脱节, 并且未能结合中国媒体的认知特征和话语实践来解

释人工智能技术的动态。本研究以大量新闻语料为基础,结合风险话语的社会建构性和科学传播理论,并充分运用框架分析、话语主体分析和情感分析相结合的方法,分析中美媒体在报道生成式人工智能伦理问题时的话语建构异同。实证结果揭示了新兴技术在话语建构动态中的跨文化差异,优化 AI 伦理风险话语模式的分析框架,并为中国媒体在处理技术风险传播问题上提出了本地化的理论和实践建议。

中华学术外译中语言学著作英译本的短语化模式及其规范趋同与差异 ——一项大语言模型辅助的语料库研究

柳 玥 华中师范大学

“中华学术外译项目”是中国学术成果对外译介的重要制度化实践。语言学著作英译本中的短语化模式不仅体现译语文本的意义组织方式,也关涉中国学术话语进入英语学术语境时的规范化表达。本文以学术话语规范与短语化研究为理论参照,并将研究对象置于国家翻译实践语境中加以考察。研究以 2015 年—2018 年该项目中语言学学术著作英译本为对象,构建英译本语料、汉语原文本语料和英语原创语言学学术文本参照语料,考察制度化学术外译实践中译语学术文本的短语化模式及其与英语学术话语规范之间的趋同与差异。研究借助语料库工具提取高频词簇和特征性词簇,并引入大语言模型辅助词簇话语功能分类,重点识别其在指称表达、文本组织、立场表达和学科术语建构等方面的功能分布;在此基础上,通过人工复核和代表性案例分析,进一步阐释短语化模式如何参与译语学术文本的意义组织和规范化表达。本文旨在探索“大语言模型辅助分类—研究者复核阐释”相结合的人机协同语料库研究路径,并从短语化模式入手,揭示中华学术外译项目中语言学著作英译本在英语学术语境中的文本表现及其话语规范意义。

生成式人工智能辅助科技语体“的”字功能标注研究 ——基于小型历时语料的初步考察

刘育含 北京航空航天大学

“的”字是现代汉语中高频而复杂的结构成分,其功能分化与汉语语法化、名词化结构及书面语体发展密切相关。对于科技语体而言,“的”不仅承担定中连接功能,也参与术语结构、名词链建构和抽象概念表达,是观察汉语科学话语形式化的重要入口。既有历时语法研究和语料库语言学研究为“的”的功能描写提供了基础,但在跨时期、跨语体的细粒度标注中,人工分类往往面临工作量大、边界模糊和一致性不足等问题。本文拟选取晚清、民国和当代三个时期的科技文本各约 1 万字,构建小型科

技汉语历时语料库，并将“的”字用例初步分为定语标记、名物化标记、术语结构成分和固定表达成分四类。在此基础上，比较生成式 AI 直接标注、给定规则后标注和人工复核标注的差异，分析 AI 在古今文本识别、术语边界判断和名物化结构处理中的优势与不足。研究旨在探索生成式 AI 参与汉语历时语料功能标注的可行路径，提出适用于小型专题语料库的流程，即先进行规则提示，再由 AI 完成预标注，然后开展人工校订，并在此基础上进行频率统计，从而为科技语体历时研究和语料库标注方法更新提供参考。

Negative Semantic Prosody in Weibo Discourse: A Corpus-Based Study of a Viral Public Controversy

刘宇欣 广东外语外贸大学

This study investigates the semantic prosody of three high-frequency lexical items, *daoqian* (apology), *kaihe* (doxxing), and *xinxi* (information) in a Weibo comment corpus surrounding the “Baidu Vice President’s Daughter Doxxing” incident, in which a Baidu executive’s daughter was accused of disclosing another Internet user’s personal information. Drawing on a corpus of over 17,000 Weibo comments, segmented and analyzed using AntConc 4.3.1, the three node words were identified as the most frequent content words and examined through collocate statistics within a 5-word span, followed by manual concordance (KWIC) analysis. Preliminary findings indicate that all three node words exhibit a predominantly negative semantic prosody, realized through distinct discursive mechanisms: *daoqian* frequently co-occurs with rhetorical questions challenging its sincerity (e.g., “有用吗,” “不配”); *kaihe* collocates with victim-identifying terms (e.g., “孕妇,” “素人”) and explicit evaluative descriptors (e.g., “可怕,” “风险”), reflecting its status as an already pejorative neologism; *xinxi* is consistently framed as a violated entity through association with “泄露” (leaked) and “隐私” (privacy).

社交媒体文本的自动化叙事标注 ——一项使用大语言模型的探索性研究

陆娇娇 北京外国语大学

社交媒体文本碎片化、口语化与隐含叙事特征显著。传统人工叙事标注耗时长、主观性强；自然语言处理则难以精准捕捉隐性叙事逻辑。上述技术局限大大制约了数字时代社交媒体叙事传播研究的

规模化推进。针对该问题，本研究基于大语言模型（LLM）开展社交媒体文本自动化叙事标注的探索性实验。本研究结合叙事要素特征，通过结构化提示工程与少样本学习优化模型标注逻辑，实现叙事角色、时序、因果关联和价值取向等多维度要素的自动识别与标注。本研究以推特社交媒体真实文本为实验数据，将 LLM 标注结果与人工黄金标注标准对比检验。结果显示，大语言模型能够有效地开展叙事分析和标注，但在不同维度上的表现存在差异：在较为简单的标注任务方面，比如行为体识别和时序分析，大语言模型的表现与人类分析结果较为接近；但在处理更为复杂的标注任务时，比如判断角色类型和价值取向等，大语言模型则表现相对逊色。社交媒体短文本中的种种模糊之处与复杂性对大型语言模型来说是一大挑战，这也在一定程度上限制了其在此类情况下准确解读文本的能力。本研究验证了大语言模型用于文本叙事自动化标注的可行性与有效性，为大规模社交媒体叙事挖掘、舆情演化分析提供可复制的技术方案。

"He's Indeed Not as Handsome as Me": Local Grammar Patterns and Semantic Prosody of Negative Comparison in Putonghua Chinese

卢明远 四川外国语大学

This study investigates negative comparisons in Chinese online interactive discussions through local grammar and semantic prosody/preference analysis. While negative comparison has been extensively studied from grammatical and lexical-semantic perspectives, its pragmatic realizations and evaluative orientations through corpus-based analysis remain understudied. The study identifies seven negative comparison markers (*buru*, *mei*, *buji*, *bibushang*, *buxiang*, *meiyou*, *bubi*) and uses them to retrieve instances of negative comparison from a specialized sub-corpus compiled from the NetworkCONE corpus (Xu & Sun, 2026, forthcoming). The study proposes a set of analytical terms for local grammar description of negative comparison, including Comparand-inf, Comparand-sup, Parameter, and Value. Based on these terms, 618 validated instances are annotated and analyzed, yielding 27 distinct local grammar patterns. Using Wei and Li's (2014) prosodic strength formula, a coefficient of 0.906 is obtained, demonstrating that although negative comparison can express favorable evaluations in context, the markers predominantly convey strong negative semantic prosody in actual language use. Furthermore, the study examines the semantic preference of the "Parameter" and "Value" slots through a word embedding-based clustering approach, revealing that the lexical items filling these slots cluster into semantically coherent categories that reflect the evaluative domains most frequently implicated in negative comparison. The study advances local grammar research by establishing an integrated methodological

framework combining local grammar, semantic prosody, and semantic preference analysis, enabling fine-grained examination of language meaning and function with evaluative orientations. The findings also have implications for corpus-informed language teaching and pragmatic description.

LLM 辅助的 RT 巴以冲突报道标题关键形态研究

卢婷婷 北京外国语大学

本研究结合关键形态分析与认知语言学理论，以“今日俄罗斯”（RT）俄语版巴以冲突报道标题为语料，围绕格选择如何配置参与者角色，揭示 RT 叙事建构的微观机制。研究采用“定量—语境—解释”三层整合框架：在定量层，引入大语言模型（LLM）对关键词语法格进行标注（准确率 98.58%），并在双重参照框架下识别具有统计显著性的关键形态。具体而言，语料库内部使用标准化残差（Standardized Residuals）评估差异性，外部参照库引入对数似然检验（Log-likelihood tests）衡量相对突显性，从而实现关键形态筛选的可验证与稳健化；在语境层，借助 AntConc 考察关键形态的词汇与语义环境，描述其共现与搭配所指向的语义框架；在解释层，以 L. Janda 提出的俄语格认知语义体系为基础，阐释格标记所承载的角色配置机制及其背后的认知图式。结果表明，格标记使用呈现“中心—边缘”层级结构，RT 在标题层面构建了三组核心对立图式：军事行动者—人道主义受难空间、活跃主体—受限主体、域外主导力量—域内被动场景。本文提出的“LLM 辅助标注—双参照统计识别”路径为形态丰富语言的语法特征自动识别与规模化分析提供了可复用的技术支持，并深化了对语言形态认知功能的理解。

Human-Machine Collaborative Annotation and Validation of Multimodal Corpora: Generative-Model-Driven Image-Text Logico-Semantic Relations in Digital News

吕兴克 国防科技大学

The large-scale co-production of images and text in digital news strains the small-data, manual-annotation paradigm of traditional multimodal discourse analysis. Generative AI offers a new path for building and annotating multimodal corpora, yet it raises a key methodological problem: when a model both annotates and classifies the data, how do we avoid the circularity of validating machine output with machine output? Drawing on 1,184 title-body-image triads from United Nations News, this study develops a generative transcription plus

human-machine collaborative validation pipeline. First, BLIP-2 transcribes each image into an English caption, textualizing the visual modality so that it can enter lexical co-occurrence analysis alongside the verbal text. Second, CLIP contrastive similarity provides a computable index of image-text alignment. Third, coders blind to the similarity scores independently annotate a sample using Martinec & Salway's image-text relation system; we report inter-coder reliability (Kappa) and then test how well CLIP similarity predicts the independently coded relation types (confusion matrix, ANOVA, classification accuracy). Preliminary results reveal three stable relations, differentiated by transitivity and visual representation; CLIP similarity broadly aligns with human coding but blurs at category boundaries, demonstrating both the promise and the limits of metric-based annotation. The study contributes a reproducible framework for human-machine collaborative multimodal corpus annotation and reflects on the division of labor between generative AI and linguistics.

人机协同语料库模式下多元评论话语特征探究 ——以“京师心理大学堂”公众号为例

马彩霞 西北民族大学

近些年,互联网的快速发展加速了新的社交媒体的涌现,依附于微信的心理主题公众号成为大众获取信息和分享情绪的新途径。微信用户的评论话语能够真实体现读者的心理认知状态、情感表达倾向与言语行为特征。本文以北京师范大学心理学部官方运营的“京师心理大学堂”微信公众号为研究对象,选取该公众号2021年—2026年发布的最低点赞量不小于100的235篇优质文章评论话语为语料来源,通过人工筛选整理得到有效语料3747条,依据文章评论话语属性将所有语料划分为个人成长、情感关系、社会思考、家庭关系四大主题。全程采用人机协同语料库探究模式,通过手动分类与清洗文本,依托ROSTCM6.0软件开展实证分析,从词频特征、语义网络可视化分析和情感极性特征分析三个维度,逐层剖析四大主题下读者评论话语的言语行为模式与心理表达状态、情感表达特征,实证结果既可以丰富新媒体话语研究与网络心理语言学的实证研究成果,也能够为心理类新媒体内容的创作和大众心理健康的传播提供真实可参考的实证依据。

基于语义向量空间模型的语义空间演变历时研究 ——以“V the crap out of NP”构式为例

马金辉 河南师范大学

研究基于分布语义学采用语义向量空间模型,以“V the crap out of NP”构式为例,借助COCA和

COHA 语料库、使用 PyCharm 和 R 软件，从语义角度分析构式的分布情况，尝试回答构式对动词的选择偏好以及动词嵌入类型，并通过构建语义密度指标与时间交互项的混合效应模型，探究 1950 年至今动词语义聚集的演变规律及其对词汇创新的驱动机制。研究结果表明：语义空间演化呈现出累积、上升、动态平衡的三阶段特征，K 近邻语义密度对动词涌现具有显著正向效应。该研究不仅揭示了动词从局部聚集到全局结构化的演化路径，也验证了语义密度作为句法能产性量化指标的有效性。

Revisiting the Extended Unit of Meaning Based on Dependency-Relation: A Distributional Semantic Approach to Semantic Prosody

明志龙、董 敏 北京航空航天大学

Due to the lack of operable criteria for identification of the sequence of co-occurring items as the extended unit of meaning, traditional semantic prosody analysis is characterised by labour-intensive work and accordingly substantial inconsistency. This study proposes an operationalized framework involving three key procedures: identifying the unit of meaning via dependency relations, employing an automated extraction algorithm, and utilising distributional semantics and clustering algorithms to categorise semantic prosodies. Focused on the use of the word global in a Belt and Road Initiative (BRI) news corpus, this article successfully identifies ten distinct semantic prosodic patterns grouped into negative, positive, dynamic, and neutral orientations. The classification of “dynamic prosody” particularly reveals that the putative reader’s stance may play an important role. By formalising sequences performing attitudinal discourse functions, the new approach enhances the consistency, reproducibility and granularity of semantic prosody analysis in corpus linguistics.

冰雪旅游专题多模态语料库的设计与建设

潘晓琪、李 琳 浙江外国语学院

本研究以冰雪旅游这一特色文旅业态为具体研究对象，系统描述了一个冰雪旅游专题多模态语料库的设计和建设过程。该语料库聚焦东北三省核心城市，涵盖近五年政府公文、官方媒体报道和自媒体内容三大类语料，以文本、图像、音频、视频等多模态形式呈现，共计 600 万字左右。本研究使用 ELAN、UAM Image Tool 等工具对图片、音频、视频等语料进行了标注和分析，并基于大语言模型与 Python 构建了关键词检索索引，实现了跨模态数据的关联与检索。依托该语料库，研究者可系统开展

政府公文、官方媒体与自媒体的三维对比分析，深度剖析官方政策导向、主流舆论塑造与自媒体真实评价之间的互动机制与话语协同。本研究所建设的专题语料库可为旅游话语分析、文旅领域智能服务开发等相关研究与实践提供高质量的语言数据支撑，亦是数智时代专用领域语料库赋能文旅产业、推动语言资源向产业价值转化的一次有益探索。

算法深描：基于 Vibe Coding 的欧盟对华产业政策批评性话语分析

彭伊寒 武汉理工大学

生成式人工智能为语料库语言学的研究范式革新提供了全新技术路径，传统人工主导的建库与话语分析模式存在效率偏低、流程割裂等局限，难以适配海量政策话语的深度挖掘需求。为推动批评性话语分析与量化语料研究的智能化转型，本研究以欧盟对华产业政策英文话语为研究对象，创新性地引入 Vibe Coding 智能研发范式，依托 Claude Code 搭建全流程自动化研究技术框架。研究整合智能爬虫、自动化清洗标注、语料库构建、BERT 主题建模与情感分析技术，打通从原始话语采集、标准化语料加工到多维度话语量化分析的完整研究链路。本研究最终构建规模约 50 万字的欧盟对华产业政策专用英文语料库，配套开发可视化检索的网页交互平台，并结合量化数据与批评性话语分析理论，系统阐释欧盟对华产业政策话语的主题架构、情感倾向与隐性意识形态特征。Vibe Coding 技术可有效解决传统语料研究流程碎片化、人工干预度高、数据分析浅层化等痛点，构建了语料库建设与批评性话语分析的智能化全新研究范式。本研究不仅为国际贸易政策话语研究提供精准的数据支撑与多维分析视角，更为语料库语言学领域的技术迭代与方法创新提供可复用、可拓展的实践路径。

Constructing the "Other": Representations of Traditional Chinese Medicine in the British Press (1800 – 1980)

钱毓芳 浙江工商大学

The historical reception of Traditional Chinese Medicine (TCM) in Britain is characterized by a prolonged period of obscurity punctuated by moments of intense, often sensationalist, curiosity. While the post-1980s era witnessed the rapid professionalization and mainstreaming of TCM, the preceding century laid the discursive groundwork for its acceptance or rejection. Based on the historical news corpora accessible via the Lancaster University CQPweb platform, this study argues that the British press, ranging from metropolitan dailies to regional weeklies, functioned as the primary site where cultural myths about Chinese medicine were forged and

circulated. By examining these early representations, we can trace the deep structures of Orientalist thought that framed TCM as fundamentally “other” to Western biomedical norms.

A Local Grammar of Relevance Statements Across Engineering and Medical Sciences Research Articles

邱辛新 中国科学院大学

This study examines the metadiscoursal act of making relevance statements through a local grammar framework. Prior studies have examined the lexico-grammatical features of relevance in isolation within established academic genres; yet, little is known about the recurrent semantic role patterns of relevance statements in underexplored genres or how such patterns vary across disciplines. To address this gap, two discipline-specific corpora were compiled, each comprising 300 significance statements drawn from the fields of engineering and medical sciences respectively, sourced from the Proceedings of the National Academy of Sciences (PNAS). A set of context-specific functional terminologies was proposed to capture the semantic configurations of relevance statements, enabling a cross-disciplinary comparison of how the two disciplines structure relevance claims and modulate epistemic certainty through the deployment of hedging devices. The study then compared the commonalities and differences in these recurrent patterns to illuminate how members of the two disciplinary communities evaluate the importance of their research and calibrate epistemic certainty in distinct ways. The findings revealed that, despite sharing common patterns for evaluating the relevance of their research, the two disciplines diverge systematically in how they calibrate epistemic certainty. Beyond their shared use of modal hedges (e.g., could, may), engineering favors value indicators that directly link research outcomes to the target domain (Direct Relevance), whereas the medical sciences favor those that merely point to potential value for the target domain (Projected Relevance), with this divergence further shaped by the epistemic risk inherent in the target domain rather than disciplinary affiliation alone. The findings offer theoretical and pedagogical implications for EAP/ESP genre research, writing instruction, and the teaching of disciplinary academic literacy.

Local Grammar of Cause and Effect in Economic Discourse: A Corpus Analysis of the Chinese Report on the Work of the Government

邱音竹、张芸洁（共同一作） 北京航空航天大学

This study employs a local grammar framework to investigate causal constructions in the English translations of the Chinese Report on the Work of the Government (2021—2025), a prototypical macroeconomic discourse genre. Drawing on Allen's (2005) local grammar of cause and effect and the New Development Concepts, this study examines a corpus of 80,669 tokens, from which 292 instances are randomly sampled for manual analysis. The manually annotated instances served as training examples and evaluation data for GPT-5, which was subsequently prompted to annotate the remaining corpus. To assess annotation reliability, the two authors independently re-annotated a randomly selected 20% of the corpus, and agreement between the human annotations and the overall LLM-generated annotations reached Cohen's $k=0.87$. Through manual coding of functional categories (Cause, Effect, Hinge, *etc.*) and economic semantic parameters (Economic Agents, Economic Instruments, Economic Objectives, External Factors), twelve configurations of local grammar of cause and effect in economic discourse, based on form-meaning mapping, are identified across twenty grammar patterns. Multiple Correspondence Analysis reveals systematic associations between grammar patterns, local grammar configurations, and five semantic orientations, namely innovative, coordinated, green, open, and shared development. Findings suggest that macroeconomic discourse in these reports is centered on a government-driven causal paradigm, with each grammar pattern articulating nuanced connotations of economic development. Furthermore, the present study reveals how lexico-grammatical choices systematically encode causal relationships in Chinese economic development as constructed in the five reports. By extending local grammar theory to the analysis of economic discourse, this study provides empirical insights into the lexico-grammatical mechanisms through which institutional policy texts construct causal relations.

局部语法框架下韩国媒体“朝鲜族”多文化报道的图文共选机制与多模态话语表征分析（2009–2023）：基于BERTopic主题建模与Vertex AI图片标注

齐兆贤 延世大学（韩国）

本研究以韩国国立国语院新闻语料库（2009—2023年）中包含“朝鲜族”、“中国同胞”、“中国侨

胞”等关键词的 7257 篇新闻（含报道图片）为研究对象，结合 BERTopic 主题建模、Vertex AI 图片标注、局部语法与视觉语法框架，分析韩国媒体“朝鲜族”多文化报道的图文共选机制与多模态话语表征。研究发现：第一，BERTopic 主题建模共识别 44 个主题，其中 Topic 1（多文化教育，271 篇）、Topic 5（移民人口与社会治理，141 篇）和 Topic 33（多文化家庭与社会融合，40 篇）共同构成“朝鲜族”多文化报道的三个子主题，并据此构建新闻文本—图片多模态语料库。第二，利用 Vertex AI 标注新闻图片，并结合局部语法与视觉语法开展跨模态分析，三个子主题分别建构了“朝鲜族”作为多文化教育参与者、移民人口与社会治理对象以及多文化家庭成员的不同表征，文本局部语法模式与视觉资源呈现稳定的跨模态对应关系。第三，文本局部语法模式与新闻图片视觉资源之间形成相对稳定的图文共选机制，共同实现了不同多文化形象的建构。本研究借鉴图文共选思想，构建融合 BERTopic 主题建模、Vertex AI 图片标注、局部语法与视觉语法的人机协同分析框架，丰富了图文共选思想在新闻多模态语料分析中的应用，为语料库语言学与多模态话语分析的融合提供实证路径。

语境中的“中国”：美国国情咨文涉华话语的语义聚类与框架分析（1972 年—2026 年）

任俊强 北京外国语大学

美国国情咨文是观察美国国家叙事演变的核心文本。本文基于 1972 年—2026 年国情咨文语料，运用语义聚类与框架分析方法，系统考察“中国”相关概念在美国总统国情咨文中的语义角色与修辞功能。研究发现，涉华话语并非围绕“中国”本身形成单一叙事，而是分散嵌入多种框架之中。当美国总统提及中国时，他们真正谈论的并非中美关系本身，而是美国自身的治理议程。“中国”在美国国家叙事中承担的是“功能性修辞角色”——在不同历史时期被征用为外交成就的证明、经济竞争的理由、技术安全的威胁或国内困境的解释。涉华话语演变的本质，并非美国对华态度从友好走向强硬的线性过程，而是中国在美国国家叙事中的“语法位置”发生了根本性变化：从被描述的外部对象，转变为定义美国自身的“构成性参照系”。

A UFA—Approach to Diachronic Investigation into the Identity Construction of Veganism in *China Daily*

任纹漩 四川外国语大学

Recently, veganism has become a prominent trend in the West, prompting growing scholarly interest in

media representations of vegetarians and vegans (Buttny *et al.*, 2020; Brookes & Chalupnik, 2022; Wilczek-Watson and Brickley, 2025). However, such research in the Chinese context remains limited. This study investigates the identity construction of vegetarians and vegans and their diachronic changes in a corpus of 981 China Daily news reports published between May 2001 and May 2026 through Usage Fluctuation Analysis (UFA). UFA examines the changes in the usage of the target word using automatic collocation comparison (Brezina, 2018:255) with four steps (McEnery *et al.*, 2019; Liu *et al.*, 2026). First, the corpus is divided into 290 subcorpora using overlapping sliding-windows. Second, collocates of vegetarians and vegans were identified using Mutual Information (MI). Third, Python was used to calculate Gwet's AC1 coefficient (Gwet, 2002, 2008; Brezina, 2018) to measure diachronic collocational similarity. Finally, Generalized Additive Models (GAMs) (Hastie and Tibshirani, 1987) were implemented in R to generate UFA graphs based on AC1 values. The results are four types of collocates and corresponding UFA graphs. Retrieved collocates were classified into six semantic categories: Food and Diet, Health and Body, Business and Market, Religion and Belief, Culture and Ideology, and Ethics and Environment. Results indicate that vegetarian is primarily associated with health and religion, whereas vegan is more closely linked to health debates and global food trends. Meanwhile, the UFA graphs show that vegetarian usage is relatively stable, while vegan discourse remains more dynamic and less consolidated. Overall, representations of vegetarians and vegans in China Daily are shaped by health concerns, cultural traditions, market forces, and, to a lesser extent, environmental awareness. Methodologically, the study demonstrates the value of UFA as an innovative approach to exploring the diachronic changes of identity construction in discourse.

From Visual–Textual Relations to Linguistic Choices: A Probe into Chinese Print Advertising

任卓璇 北京外国语大学

Meaning in print advertisements (hereafter, ads) is built through the interaction between visual and verbal modes. Previous studies of text–image relations in print ads have mainly followed a bottom-up approach, inferring relation types from fine-grained modal resources. Less attention has been paid to how relation types shape lexicogrammatical choices. Using a relation-first approach, this study investigates ideational visual–textual relation patterns, their lexicogrammatical realisations, and the co-selection between them. The analysis draws on a corpus of 300 print ads collected from the 2000–2025 volumes of the IAI China Advertising Yearbook. Visual–textual pairs were categorised using a framework adapted from Martinec and Salway (2005),

and annotation followed a human–AI–human collaborative procedure. Extension and elaboration emerge as the two most frequent relation types and form the most common pairing. The four relation types co-occur with distinct lexicogrammatical features. Enhancement draws mainly on temporal nouns and non-Chinese or numerical expressions; Projection tends to select causal and explanatory clause complexes; Elaboration favours nominal structures such as modifier–head and classifier constructions; and Extension co-occurs with verb–object and coordinate constructions. These findings reveal systematic co-selection between visual–textual relations and linguistic features in print ads. They also offer insights for copywriting pedagogy and the design and evaluation of advertising content, including content generated by large language models.

媒体高频词语频序轨迹的生命周期模式及话题变化机制 ——基于 2010 年—2024 年词表的软聚类研究

宋婧婧 厦门理工学院外国语学院

词汇历时变化检测是应用语言学的前沿方向，但对汉语当代媒体词汇的轨迹刻画不足。由此引出两个问题：汉语媒体高频词语的频序动态能否被还原为少数几类可辨认的生命周期模式，从而呈现媒体词汇兴起、退场与更替的整体规律；这种刻画能否为社会及媒体的发展变化提供语言学旁证，并区分真正的语言系统变化与话题热度波动。研究依据《中国语言生活状况报告·年度媒体高频词语表》(2010 年—2024 年)，取 15 个年度词表交集得 16,935 个跨年稳定共现词，以频序为使用强度的秩统计量，对各词轨迹作 z-score 标准化后进行 Mfuzz 模糊 C 均值聚类，并以 K-means、Ward、DTW-KMeans、K-Shape 交叉验证：多种方法在一阶结构上高度一致，确认升 / 降二元分化稳健；进而采用能保留边界词模糊归属的 Mfuzz 软聚类，兼顾统计质量与解释粒度。研究识别出五类生命周期模式：持续上升型（党政治理话语，24.5%）、后期兴起型（经济产业话语，16.3%）、中期峰值型（法治反腐议程，11.6%）、U 型反弹型（文艺生活话语，14.8%）与持续下降型（日常口语退出，32.7%），整体呈现“官方与专业话语体系化兴起、日常口语化表达系统性退出”的结构性替代。研究进一步检测到 2020 年前后一次完整的话题腾挪—恢复循环：479 个公共卫生词批量跃升、304 个社会新闻词同步退场，进退非对称表明危机期媒体注意力存在短期弹性扩张，且退场话语多在 2023 年后回升，呈现可逆征用而非永久替代。本研究将 Mfuzz 软聚类迁移至汉语媒体词汇生命周期刻画，并尝试引入“词汇更替”框架，为媒体话题变化的实证检验、舆情监测与新词语生命周期追踪提供方法与数据支持。

Exemplification in University Student Assignments: A Local Grammar Analysis Across Levels of Study

宋畏林、林 铃 上海交通大学

While local grammar analysis of exemplification in academic discourse has recently attracted increasing scholarly attention (Su & Zhang, 2020), it has been mainly applied to examining either published research articles or students' degree dissertations (e.g., Su & Lu, 2022). To the best of our knowledge, no study has traced the developmental patterns of exemplification in university student assignment genres. To address this gap, the present study adopts the local grammar approach to explore the commonalities and differences in students' patterned use of exemplification across four levels of study (from first-year undergraduate to master's level) based on the British Academic Written English (BAWE) corpus. The analysis has identified 42 local grammar patterns in 13,900 instances of exemplification retrieved from the corpus. The findings show that student writers consistently prefer PA1 (Exemplified + Indicator + Exemplification-SC) and its variants as their primary strategies for realising exemplification in coursework assignments. Cross-level comparisons further reveal a distinctive V-shaped developmental trend in the overall frequency of exemplification patterns: their use gradually decreases across the undergraduate phase but increases markedly at the postgraduate level. The study demonstrates the value of local grammar for documenting developmental changes in the patterned realisation of exemplification in university student assignment writing. It also provides an inventory of exemplification patterns that can inform pedagogical support for students' more appropriate and strategic use of exemplification.

Expressing Opinions in Spoken English Across L1 Backgrounds and Task Types: A Human–LLM Collaborative Exploration

宋瑛明 北京外国语大学

Research on second language (L2) pragmatics has increasingly examined how L2 English speakers perform speech acts, yet large-scale annotation of such pragmatic phenomena remains methodologically challenging. This study investigates how English learners from five first language (L1) backgrounds—Chinese, French, German, Greek, and Polish—express opinions in spoken interview tasks. Following Edmondson, House, and Kádár (2023), “opinions” are operationalized as instances of the speech act *Opine*. The study develops and

evaluates a human-LLM collaborative annotation procedure consisting of three stages: protocol development, prompt design and optimization, and human oversight of model-generated annotations. The results show that the GPT-5.4-based procedure identified Opine instances with consistently strong performance across task types and L1 groups, with F2 scores around 0.90. Subsequent analysis further reveals systematic variation in the distribution and realization of opinion-expressing speech acts across interactional contexts and learner backgrounds. Methodologically, the study demonstrates the feasibility of using LLMs, under principled human supervision, for speech act annotation in learner corpus research. It argues for a human-LLM collaborative paradigm that can support scalable, transparent, and theoretically informed analyses of speech acts.

How Do Student Interpreters of Different Proficiency Levels Handle Mode Change? A Corpus-Based Entropy Study on Consecutive and Simultaneous Interpreting

苏玥玥 北京外国语大学

Consecutive (CI) and simultaneous interpreting (SI) differ markedly in their temporal features and working-memory demands. Corpus-based studies of how the two modes shape interpreting output have generally compared the lexical and syntactic diversity across modes, whereas how output varies across interpreters of different proficiency remains relatively underexplored. Adopting entropy-based measures, this study draws on a learner corpus of 124 Chinese–English outputs by 62 master’s student interpreters to examine how their output differs in lexical variability and syntactic flexibility across the two modes, and whether such differences depend on proficiency. The results reveal that student interpreters’ outputs differ significantly between the two modes and across proficiency levels. In SI, outputs show greater lexical variability than in CI (paired $d_z = 0.70\text{--}1.02$), and lexical variability increases with proficiency in both modes. Grammatical-category use remains relatively stable. The sharpest mode- and proficiency-related differences appear in syntactic flexibility: SI outputs are more flexible than CI outputs (paired $d_z = 0.57\text{--}1.34$), and the gap between modes widens at higher proficiency. These findings suggest that students of higher interpreting proficiency may be more flexible in organizing sentences, which points to the importance of developing learners’ ability to organize sentences flexibly under time constraints, rather than merely expanding their lexical and grammatical inventory.

Exploring a Large Language Model–Assisted Approach to Annotation of Appraisal as Co–Selection

孙文青、董 敏 北京航空航天大学

Appraisal Theory (Martin and White 2005) draws a distinction between the appraisal of human feeling or behavior categorised into Affect and Judgement, and the appraisal of things defined as Appreciation. On this basis, Su and Hunston (2019) argue that appraisal is in essence instantiated by co-selection of both attitudinal targets (i.e. a human or a thing target) and attitudinal lexis (i.e. emotion or opinion lexis). Regarding the source of appraisal, Martin and White (2005: 72) point out that normally speakers and writers are interpreted as the source of evaluation, i. e. averral, unless attitude is projected as the speech or thoughts of an additional appraiser, i.e. attribution. Using the primary data of a corpus comprising Belt and Road Initiative news articles, which totals 10,016 texts and 6,381,224 word tokens, this study explores a new large language model-assisted approach to the annotation of appraisal as co-selection. Based on an annotation framework encompassing target, source, and attitudinal category, we designed a transferrable annotation workflow and validated a human-AI collaborative approach to annotation of appraisal. The approach divides the overall task into three stages, implemented over five rounds. First, after two rounds of error analysis and prompt designs for the annotation of targets, satisfactory annotation results were achieved in the third round. Building on this outcome, another round of error analysis and prompt refinement led to adequately acceptable annotation results for attitudinal categories in the fourth round. Finally, through explanation enhancement, substantial consistency in the annotation of sources was achieved in the fifth round. The findings demonstrate the robust performance of LLMs in the annotation tasks of appraisal as co-selection by virtue of an error analysis-based prompt refinement and a quantitatively-driven iterative annotation process.

Headedness, Overlap, and Salience: The Multi–dimensional Tension in English Blends

汤慧桃 复旦大学

A recent contribution to Digital Scholarship in the Humanities by Meng and Hilpert (2024) applied distributional semantics to English lexical blends but conflated three conceptually independent dimensions: semantic headedness, semantic overlap, and semantic salience, while relying on static word embeddings that

cannot adequately handle polysemy. To address these gaps, the present study distinguishes the three metrics and quantifies them using the dynamic word embedding model BERT on a dataset of 87 English blends. The research findings are threefold: 1) Semantic headedness is right-headed in 52.9% and left-headed in 47.1% of the 87 blends, supporting a general right-headedness tendency (Renner 2019: 38) while revealing that left-headedness is not rare; Semantic overlap is left-oriented in 43.7% and right-oriented in 56.3% of the 87 blends; 2) In terms of semantic salience, 56.3% of the 87 blends are left-oriented and 43.7% are right-oriented; 3) The three metrics exhibit systematic directional inconsistency, revealing that English blends function as multi-dimensional semantic objects in which categorical inheritance, distributional similarity, and contextual prominence need not converge, which calls for a more empirically grounded and multi-faceted account of blend semantics. Theoretically, our findings corroborate the gradient view of Meng and Hilpert (2024), moving blend research beyond taxonomic classification toward a usage-grounded, multi-dimensional model that requires a multivariate approach. Methodologically, the study demonstrates the necessity of cross-validating conclusions using multiple metrics in morphological research, illustrating how digital methods can address longstanding problems in word formation.

基于依存句法的大语言模型翻译质量优化研究

田文杰 武汉理工大学

随着大语言模型在翻译领域的广泛应用，其译文质量的量化评估与优化成为重要研究方向。当前，围绕大语言模型在实践中的应用以及与传统机器翻译的对比探讨正逐步成为国内外研究热点，然而，关注大语言模型生成译文的句法特征，尤其是以平均依存距离等句法指标探究其与翻译质量间关系的研究较为局限。本研究选取新闻、政治等不同体裁语料，以平均依存距离、依存关系类型分布为核心量化指标，从英译汉、汉译英双向翻译维度，计算大语言模型生成译文的依存句法特征，验证其是否符合“折中假设”，并分析核心语法指标与人工翻译质量评分的相关性，同时对比大语言模型译文与人工译文在平均依存距离、依存关系类型占比上的差异。通过量化依存句法特征，揭示大语言模型译文与权威译本存在差距的原因。据此提出构建基于依存句法结构的量化评估标准，并针对大语言模型优化译文质量的具体路径，助力大语言模型翻译的准确性与自然性，推动机器翻译与依存句法理论的交叉研究。

数智时代背景下汉语构式搭配分析的智能路径探索

王诚文 中央财经大学

在数智化浪潮的推动下，以大语言模型为代表的人工智能技术正深刻重塑语言学的研究范式，推动其向更纵深、更系统、更智能的方向转型。构式搭配分析作为构式研究中的主流实证方法，在构式使用规律及特征方面发挥着不可替代的作用。然而，当前汉语构式搭配分析在实践中面临显著困境：传统依赖语料库检索构式实例的方法智能化程度低，所获数据噪音干扰严重，致使研究者需投入大量时间与精力进行人工筛选与清洗，严重制约了研究效率与深度。大语言模型所展现出的卓越语言理解与生成能力，为突破上述瓶颈提供了全新的技术路径。其能够对海量文本进行深层语义解析与模式识别，从而为实现汉语构式搭配的自动、精准识别与提取提供强大支持。在人机协同的模式下，大语言模型能够有效助力研究者从复杂真实语料中高效析出有价值的构式实例。本文将首先系统梳理构式搭配分析的理论基础与操作方法，明确其方法论价值。进而，重点剖析当前汉语构式搭配分析在实践层面所面临的具体挑战与成因。在此基础上，将核心探讨如何利用大语言模型技术，构建汉语构式搭配分析的智能化研究新路径。最后，将以汉语法律文本中高频动词的搭配构式识别为具体案例，通过实证研究设计，详细阐释在人机协同模式下，实现构式搭配分析流程智能优化的可行方案与显著成效，以期为数智时代的汉语研究提供切实可行的新思路。

从语言资源库到互动知识库： 生成式人工智能时代汉语对话语料库建设的探索与实践

王德亮、孙钰涵 北京师范大学

目前国内外汉语对话语料库的可用性、开放性和时效性还不够理想，多数对话语料库主要服务于语言资源保存、语音技术开发或自然语言处理任务，对真实互动过程的关注相对不足。在生成式人工智能迅速发展的背景下，现有资源已难以满足人工智能时代互动语言研究与对话建模的需求。基于“对话转向”（Dialogic Turn）的理论视角，本研究团队正在建设一个以自然互动为核心的汉语对话语料库，目前语料库近 1500 万字，含有近 60 万个话轮。该语料库以真实互动为原则，涵盖日常对话、影视对白、网络互动、访谈、课堂交流、多模态语料及古代汉语对话等多个子库，构建了包括话轮结构、对话行为、共鸣现象、话语立场等在内的多层级标注体系。与传统语料库相比，该语料库突出互动性、动态性和过程性特征，强调语言形式与互动功能之间的实时关联。本文进一步探讨大语言模型在语料收集、自动转写、辅助标注和质量校验等环节中的应用潜力，并分析人工智能辅助语料库建设面临的

数据可靠性、标注一致性和伦理规范等问题。我们认为，未来语料库建设应从“语言资源库”走向“互动知识库”，推动语料库语言学由静态描写转向动态互动研究，为生成式人工智能时代的语言理论创新提供新的基础设施和研究范式。

大模型辅助下的美媒移民话语趋近化自动标注与人机协同分析

王 婕 中国传媒大学

本研究探讨生成式人工智能（GenAI）在批评话语分析（CDA）多层语义标注中的应用效能。研究以《纽约时报》中涉及“unauthorized/undocumented/illegal immigrant”的美国移民政策报道为语料构建目标语料库。传统话语分析在识别隐性政治修辞和意识形态话语时高度依赖人工编码，难以实现大规模量化分析。鉴于此，本研究基于彼得·卡普（Piotr Cap）的趋近化理论（Proximity Theory），将空间、时间和价值三个维度的威胁话语特征转化为精细化的多步思维链（CoT）提示词策略，引导大模型（如 DeepSeek-R1、GPT-4o）对语料进行自动分层标注。通过引入人工标注作为标准，计算 Krippendorff's Alpha 系数以验证模型标注的评判者间信度。研究发现，大模型在空间和时间维度的物理威胁识别上展现出极高的准确率，但在价值维度隐性修辞的推理上存在一定的偏见与对齐局限。基于此，本文提出了一种“人工设定标准—大模型批量预标注—人工纠错反思”的人机协同语料库分析模式，为 GenAI 时代的计算话语学研究提供了方法论创新。

Can Machine Translation Keep the Speaker's Stance? Appraisal and Interpersonal Meaning in English-Chinese Renderings of TED Talks

王 丽 上海师范大学

Translation does more than transfer propositional content: it re-voices a speaker's stance and relationship to an audience. As machine translation shifts from neural systems to generative large language models, it remains an open question whether these systems preserve this interpersonal layer or quietly rewrite it. Drawing on Appraisal Theory, we examine how three architecturally distinct engines—a generative model (ChatGPT), a commercial neural system (DeepL), and an open-source multilingual system (NLLB-200)—render evaluative meaning from English into Chinese across a large sentence-aligned corpus of TED talks. We track four resources through which speakers calibrate stance—negation, modality, emphasis, and concession—and find that the

systems fail in architecture-specific ways. Neural models show a certainty bias, sharpening modality and emphasis and, at times, reversing negation; the generative model, chasing fluency, falls into a simplification trap that flattens concessive and dialogic reasoning. In different registers, both reshape the speaker's persona in the target culture. We argue that translation fidelity must be judged at the level of interpersonal meaning rather than semantic content alone—an argument that carries particular force for high-stakes public discourse.

“语料库 + 大模型”双轮驱动的垂直领域语言智能教学范式构建 ——以轮机英语为例

王冉然 天津海运职业学院

通用大语言模型在专门用途英语（ESP）等垂直领域应用中普遍面临“通而不专”的困境，表现为专业术语准确性差、行业场景适配不足等问题。本研究基于语料库语言学与 AI 教育应用双重视角，以轮机英语为学科载体，提出“语料库 + 大模型”双轮驱动的垂直领域语言智能教学新范式。该范式以百万级轮机英语专业语料库为知识底座，通过领域持续预训练与指令微调，将通用大模型锻造为精通轮机领域知识、熟稔国际海事公约、理解机舱工作场景的“领域专家模型”。研究构建了“资源—模型—应用”三位一体的理论框架，并通过智能词典、题库讲解、阅读理解与口语对话四大功能模块的系统集成加以实证验证。这一范式不仅有效解决了通用模型在垂直领域“水土不服”的共性难题，也为航空英语、法律英语、医学英语等其他专门用途英语的智能化教学研究提供了可复制、可推广的方法论样板，推动了 ESP 教学研究从“静态语料资源”向“智能动态交互”的范式跃迁。

大语言模型在放置动词教学中的应用探索

王 珊、王少茗 澳门大学

放置动词作为描述物体在空间中的移动和定位的重要词汇，对语言学习者而言具有较高的学习价值。然而，现有研究在类型划分与受事分析方面仍不系统，影响了教学效果。本文基于受事与处所之间的接触方式，将放置动词划分为四类，并以“盖”“摆”“装”为例，归纳其受事语义特征与空间表达模式。结果发现，三类动词的受事均为具体物，但属性存在差异；结合 Pustejovsky 等人的空间关系分类，进一步总结出典型的空构型类型。在此基础上，本文探讨了大语言模型在汉语教学中的应用优势，主要体现在可视化展示、语言结构归纳、语境生成和即时反馈等方面，显示出其在提升学习效率与教学支持方面的潜力。

生成式 AI 辅助的小说—译本—电影多模态隐喻标注研究 ——以《了不起的盖茨比》为例

汪雄飞 广东外语外贸大学 朱宇吟 香港理工大学

生成式人工智能的发展为文学语料的跨语言、跨媒介标注提供了新的方法可能。本文以《了不起的盖茨比》英文原文、中译本及其 2013 年电影改编为研究对象，构建小说—译本—电影三层对齐的小型多模态语料库，考察小说中的核心隐喻如何在翻译和电影改编中被保留、显化、弱化、替换或重构。研究重点关注“梦想是光”“道德腐败是灰烬”“欲望是追逐”“阶层是空间距离”“凝视是权力/审判”等概念隐喻系统，分析其在英文文本、汉译文本以及电影色彩、光影、空间、音乐、服饰和镜头调度中的不同实现方式。方法上，本研究使用 Codex Agent 和 DeepSeek V4 Pro API 对候选隐喻单元、概念映射和跨模态对应关系进行初步标注，再由研究者进行人工复核与修订，以提高标注的可解释性和可靠性。研究发现，生成式 AI 能有效辅助隐喻候选项提取和场景对齐，但在识别弱显性隐喻、复杂概念映射和电影模态意义时仍需人工判断。本研究旨在为生成式人工智能参与文学翻译研究、多模态隐喻分析和人机协同语料库标注提供一个可操作的案例。

大模型辅助情感对话语料隐喻标注的人机对比研究 ——以中文线上心理咨询文本为例

王旭佳 北京外国语大学

生成式人工智能为大规模语料库自动化标注提供了高效路径，但在情感类专业对话文本的隐性隐喻识别中，仍存在精度不足、语境适配性差等现实问题，制约了专用认知语言学相关语料库的智能化构建。为探究大模型辅助隐喻语料标注的可行性与局限，本文以中文线上心理咨询对话为研究载体，筛选构建 5 万字规范标注语料库，完成精细化人工隐喻标注。同时，研究选取主流通用大模型（GPT-3.5、DeepSeek）开展全自动隐喻提取实验，贴合心理咨询语境，系统对比模型自动标注结果与人工标注数据，量化分析人机标注的一致性差异。实证结果表明，通用大模型普遍存在隐喻漏标、泛化误判、情境语义脱节三类典型误差，难以精准捕捉依附个体情绪、语境依赖性极强的心理咨询个体化隐喻，纯自动化标注无法满足认知语言学研究的严谨性要求。基于此，本文构建“大模型粗标+人工分层校准”的混合式标注工作流程，有效平衡语料标注效率与标注质量。本研究可为 GenAI 时代情感领域专用隐喻语料库的智能化建设、人机协同标注范式优化提供实证支撑与实操参考。

基于句向量的英语二语学习者作文语篇衔接评估

王雅慧 中国海洋大学

语篇衔接是二语写作质量的重要维度，但其测量长期面临挑战：人工评分信度有限，连贯性作为心理表征难以直接观测，现有 Coh-Metrix/TAACO 等方法仍偏重词汇频率与局部特征，深度学习方法又缺乏可解释性。本研究基于句向量语义相似度与图结构思想，从段内、段间与整篇三个层面构建 12 个结构化语篇语义指标，刻画句子与段落之间的语义推进与全篇组织关系。以 ELLIPSE 语料库 360 篇英语二语学习者作文为对象（低 / 中 / 高各 120 篇），结合 Kruskal-Wallis 检验、相关分析与多元层级回归进行验证。结果表明，结构化语义指标能够显著区分不同衔接水平作文，其中整篇层面指标区分力最强。在控制词数、句子数与段落数后，模型仍可解释衔接评分额外 6.2% 的变异量，具有稳定的增量效度。研究表明，将句向量语义表示与篇章结构建模结合，可构建更具解释力的结构化衔接评估框架，为二语写作自动评价与教学诊断提供新路径。

How Academic Writers Summarize: Towards a Local Grammar of Summarization

王一凡 四川外国语大学

The study adopts a local grammar approach to investigate the discourse act of summarising in academic writing. A corpus of research articles from five social science disciplines was compiled for the investigation, and summarising linguistic instances were extracted through lexical markers (e.g., summarise, in sum). The subsequent local analyses suggest 7 functional terminologies that are needed for a local grammar of summarisation and identify 22 local grammar patterns. The pedagogical implications and applications for EAP writing of these patterns are further discussed, concerning the development of teaching materials and practical teaching exercises. Overall, the study indicates that local grammar can be an effective approach to investigate and describe discourse acts by revealing functional elements behind syntactic structures and offering valuable insights into EAP writing pedagogy.

Affordances and Constraints on Qualitative Corpus Analysis in the GenAI Age: Several Epistemological and Methodological Complexities from Simulated Intersubjectivity and Participatory Sense-Making in Second Language Education

王 媛 郑州西亚斯学院

Qualitative corpus analysis delves into naturally occurring interactions between participants in actual, authentic and attested settings. GenAI produces a paradigmatic shift in qualitative research on interactions, particularly it transforms language education by providing both empirical and theoretical considerations about what sustained interactions look like and how language structures, strategies and patterns should be highlighted for pedagogical significance. While current quantitative metrics produce a veneer of measurable success, they push qualitative analysis to the forefront of evidence-based practice where the fact of what is actually happening in GenAI dialogues should be faithfully captured and accurately diagnosed. However, qualitative analysis has suffered from this acute shortage for a long time. The research explores the epistemological and methodological complexities of human and GenAI-generated dialogues for L2 education by positing that, as a unique form of simulated intersubjectivity, such interactions are grounded in a semblance of simulating language patterns using statistical probability to create a feeling of mutual understanding which is in fact ontologically empty. This affords a stark exposure of L2 learners' adaptive strategies for close scrutiny of dialogues and preconditions a review of participatory sense-making, through which meanings are supposed to be coordinated to invoke meaningful multimodal reactions to realize co-varying intentional behaviours, but for GenAI are fundamentally constrained. Methodologically, on the one hand, GenAI assumes the role of anxiety-relieving L2 learners, and on the other, its simulated character forestalls pragmatic authenticity of structural coupling in unidirectional incorporation. Qualitative corpus analysis is not a solitary task of corpora annotation, retrieval and interpretation but means entering a reflexive analysis of how GenAI dialogues reshape L2 learning outputs and make the perception-action loops inextricably intertwined. A methodological qualitative framework that treats GenAI corpora as a heuristic device and constitutive dialectic is proposed to investigate synthetic sociality for interpretive rigor in new communicative agency.

数字时代的算法规训与语言生活治理 ——基于大语言模型春节交际互动的话语分析

王悦 北京外国语大学

随着大语言模型的爆发式发展，人类社会逐步迈入自然人、机器人和数字人共生交互的新阶段，数字平台及其底层算法逐渐成为国家宏观语言治理体系中不可忽视的技术实施主体。本研究以国内五款主流大语言模型（元宝、豆包、千问、Kimi、文心一言）在春节交际语境下的互动语料为对象，运用批判性话语分析系统考察生成式人工智能如何通过算法机制执行隐性语言治理。研究发现：在语言地位规划层面，数字平台系统性强化标准书面语的优先地位，方言则被降格为可拼凑的修辞装饰，构成隐性语言等级；在语用实践层面，算法的模板化生成逻辑催生出形式多元而实质趋同的语言产品，消解个体真实的语言个性与创造力；在民俗禁忌与冲突情境中，算法通过安全对齐对民俗禁忌与冲突话语实施认知治理与道德规训，客观上压缩用户真实情绪表达的空间与话语主体性。本研究揭示了算法逻辑对语言多样性与个体话语权的隐蔽影响，提出应警惕算法驱动下的语言生态失衡，并为构建包容多元的数字语言生活治理体系提供理论依据。

Register Differences in the Syntactic Hierarchical Structure of L2 Chinese Written Production: A Synergetic Linguistic Perspective

王悦 中国海洋大学

Grounded in synergetic linguistics, this study employs the Menzerath–Altmann Law (MAL) as the core analytical tool to investigate self-organization in the sentence-clause hierarchy across two typical written registers (narrative and argumentative essays), based on a comparative analysis of written production by Chinese as a second language (CSL) learners and native Chinese speakers. Drawing on the HSK Dynamic Composition Corpus and the Lancaster Corpus of Modern Chinese (LCMC), the study performed sentence length distribution modeling, MAL fitting, parameter testing and hierarchical cluster analysis on four datasets. The results confirmed the validity of MAL across all four groups and both registers. Further comparative analyses revealed systematic group and register differences: (1) Significant differences existed between written production of CSL and native speakers in the core fitting parameter b (the rate at which mean clause length decreases as sentence length increases); (2) Native speakers exhibited clear register differentiation: their argumentative writing

exhibits a stronger Menzerathian shortening effect with tighter syntactic hierarchical constraints; narrative writing displays a weaker such effect alongside greater syntactic flexibility. No such register differentiation was observed in CSL learners. These findings align with synergetic linguistics' principle of self-organization and cognitive optimization, confirming that CSL learners' written production follows the inherent self-regulation logic of linguistic systems. However, CSL learners have not yet developed native-like register-specific syntactic adaptation. This study extends quantitative linguistic laws to L2 Chinese syntactic development and provides quantifiable indicators for CSL writing assessment and register competence characterization.

真实生活融入导向的网络语言筛选——以流行语榜单（2013 年—2025 年）及《人民日报》用例为依据

王宇婷 四川大学

本研究聚焦中国主流媒体《人民日报》对网络语言的使用与演变，以 2013 年至 2025 年间三大流行语榜单（《咬文嚼字》《语言文字周报》、国家语言资源监测与研究中心）中的网络热词为对象，以其在《人民日报》中的使用为语料来源，尝试用生成式人工智能（GenAI）辅助构建语料库并进行分析。研究运用 DeepSeek-R1 对《人民日报》历时语料进行智能分词、词性标注及语义特征标注，建立网络热词的高维标注语料库，并从历时活跃度、语境分布、语义演变三个维度系统考察网络语言进入传统主流媒体后的调整与融入。研究发现，网络语言在《人民日报》中的融入呈现阶段性递进特征，从低频试探、高频扩张到常态化渗透，各阶段在融入频次、引号标记率、语义调适程度等方面存在显著差异。DeepSeek-R1 标注在语义特征识别上与人工分析的一致性较高，验证了大模型辅助筛选的有效性。本研究对理解主流媒体如何以“真实生活融入”为导向筛选网络语言、推动语言规范与日常生活有机互动有一定的启示意义。

“一译三体”：大语言模型辅助的华兹生 *Basic Writings of Mo Tzu, Hsün Tzu, and Han Fei Tzu* 译者风格研究

王昭鉴 北京外国语大学

本研究以华兹生（Burton Watson）*Basic Writings of Mo Tzu, Hsün Tzu, and Han Fei Tzu* 为研究对象，构建小型汉英平行语料库，运用语料库工具与大语言模型辅助分析，考察同一译者面对墨家（论辩体）、荀子（说理体）、法家（论述体）三种不同思想文体时的翻译策略差异与风格一致性。研究发现，

华兹生在三篇选译中保持了高度一致的“少注”翻译策略和通俗化取向，但在处理三个核心概念时呈现出系统性偏移：墨家逻辑术语倾向归化、荀子伦理概念倾向保留、法家政治术语采取变通阐释。研究验证了大语言模型在典籍合辑翻译风格分析中的可行性，并为“合辑翻译”这一特殊翻译现象提供了初步的理论描述。

隐喻场景理论视域下的中美能源企业身份建构研究

王子玥 广东外语外贸大学

本文基于语料库方法，对比分析中美能源企业 CEO 致辞中语类特定的概念隐喻、隐喻场景及立场表达，探究两者如何通过隐喻建构差异化的企业身份。研究具体考察三个问题：1）中美能源企业 CEO 致辞中的语类特定概念隐喻有哪些？2）这些概念隐喻激活了怎样的隐喻场景？3）这些隐喻场景通过何种立场表达建构了何种企业身份？研究语料来自 2025 年《财富》世界 500 强中各 13 家中美能源企业 2010 年—2025 年可持续发展报告中的 CEO 致辞，分别构建中国能源企业 CEO 致辞语料库（CECL）和美国能源企业 CEO 致辞语料库（AECL）。研究借助 Wmatrix 7 语义标注工具提取关键语义域以确定目标域和源域，采用 MIPVU 隐喻识别程序识别语言隐喻，并通过卡方拟合度检验和对数似然检验进行隐喻的跨文化比较。研究共识别出 7 类主导源域：建筑、旅途、机器、游戏与运动、实体、力和人。研究发现，尽管中美能源企业共享“有能力的建设者”“行业领跑者”等身份，但隐喻场景的差异导致了差异化的身份建构。共性身份可由中美企业共同面临的行业挑战加以解释，差异性身份则根植于两国在政策导向及利益相关者沟通策略上的文化差异。

多模态大语言模型视觉语法标注能力评估 ——以药品包装再现结构识别为例

吴 悠 北京外国语大学

本研究考察多模态大语言模型（Multimodal Large Language Models, MLLMs）在理论驱动的视觉语法标注任务中的应用潜力，以探索其在多模态语料库建设中的方法价值。既有研究虽已广泛关注大模型在图像识别与文本标注中的表现，但其能否处理由特定理论框架定义的分析类别，仍有待进一步检验。本研究以 Kress 与 van Leeuwen 的视觉语法为理论框架，构建包含 500 余张药品包装图像的多模态语料库，系统评估 GPT、Claude、Gemini、通义千问（Qwen）和智谱 GLM 五款多模态大语言模型对包装主视觉界面再现结构的识别表现。研究通过逐轮优化提示词，将模型标注结果与人工标注进行

比较, 分析不同模型的标注表现。研究进一步关注视觉语法类别在自动标注过程中的操作化问题, 并讨论多模态大语言模型在视觉语料标注与视觉语篇分析中的方法价值与局限。

基于语料库标注分析的汉语完句性问题研究

吴术燕、周 卫 广西民族大学

完句性指句子具备自足表达意义所必需的各种条件, 或句子作为独立交际单位所必须具备的完句能力。缺乏时体标记、情态成分或特定句式支持的汉语词组结构常呈现“不完句”特征, 这给汉语语法语义研究及自然语言处理带来挑战。传统的研究多基于内省语料对完句条件进行定性描述, 难以系统揭示完句成分的分布规律与互动机制。本文主张采用语料库标注与量化分析的方法对汉语完句性问题进行重新审视, 确立多层标注维度, 统计不同类型完句成分的出现频率、共现模式以及语境偏好, 对比语体分布差异, 以揭示完句性作为一个动态范畴的连续性与梯度特征, 从而超越“完句/不完句”二分法, 建立基于特征加权的话段完句度评价模型。研究成果服务于汉语本体句法理论的实证分析, 并为中文信息处理、文本生成自然度评测等应用任务提供语言学支持。

A Study on Human-Machine Collaborative Annotation of a Multimodal Discourse Corpus of Classroom Teaching Videos

相贺延 北京语言大学

With the advancement of multimodal discourse analysis, classroom teaching videos have become a core data source for research on instructional discourse. Traditional annotation of multimodal corpora relies on frame-by-frame manual labeling via ELAN, which suffers from heavy time costs, low efficiency and massive repetitive labor. To explore a new research paradigm for corpus construction in the era of large language models (LLMs), this study experiments with a human-machine collaborative annotation framework combining LLM pre-annotation and manual refinement on classroom teaching videos. The results show that LLMs can efficiently complete basic modal recognition and rough annotation to greatly reduce repetitive manual work, while manual proofreading effectively remedies the machine's deficiencies in detail deviation and contextual comprehension. This research finally develops a replicable workflow for human-machine collaborative annotation of classroom multimodal corpora. It provides practical references to build multimodal corpora and conduct LLM-assisted discourse analysis, and offers a case study to support the human-machine collaborative research model in corpus linguistics.

Combining CDA and Topic Modeling: Analyzing the Evolution of Green Economy Media Discourse in China Daily (2001 – 2024)

向珮源 广东外语外贸大学

This study analyzes the evolution of the green economy discourse in China Daily's English-language reports from 2001 to 2024. Drawing on Fairclough's three-dimensional Critical Discourse Analysis (CDA) framework, we combine topic modeling with quantitative tools, including Jaccard similarity coefficients and Sankey diagrams, to identify discourse themes in three distinct phases, map the shifts in focus over time, and unpack the underlying power dynamics and ideological orientations. We argue that the transition in discourse reflects a dual influence of domestic policy shifts and global climate governance trends, while mature discourse also exerts a reverse constructive effect on social practice. Results reveal three stages: the initial phase (2001–2008) centered on “local governance and ecological awareness”; the second (2009–2016) expanded to “global governance and green trade” and “financial policies with renewable energy”; the final (2017–2024) emphasized global ecological discourse, notably “ecosystem and water resource governance” and “climate change leadership from the EU.” Topic numbers were reduced from nine to five, indicating a more focused structure. This transformation reflects China's national green strategy and growing engagement in global climate discourse. Meanwhile, China Daily has undergone a profound role transformation in international green communication, from a disseminator of local ecological practices to a translator of national green policies and finally to an active constructor and agenda-setter of a global green narrative with Chinese characteristics. By combining LDA topic modeling with CDA, this study not only highlights the potential of natural language processing techniques in media discourse analysis and enriches the diachronic application of the CDA framework in international environmental communication research but also provides valuable insights into the evolving strategies of green communication in China and the construction of international discourse power in ecological civilization.

Readability of Human–Written and LLM–Generated Summaries of Maritime Accident Reports: A Corpus–Based Comparative Study

谢 敏、张 滢 上海海事大学

Maritime accident reports serve as an important mechanism for safety communication and organizational

learning. Previous research has shown that the readability of maritime accident reports significantly influences the accessibility and usability of safety information, particularly for maritime professionals with diverse linguistic backgrounds (Mallam, Wahl, & Aas, 2022; Goerlandt & Liu, 2024). However, while recent advances in large language models (LLMs) have enabled the automatic generation of report summaries, little empirical evidence exists regarding whether AI-generated summaries are as readable as those written by human experts. This study compares the readability of human-written and LLM-generated summaries using a parallel corpus of 232 maritime accident reports from four investigation agencies (ATSB, MAIB, NTSB, and TSIB). Each human-written summary was paired with an AI-generated version produced by Doubao-Seed-2.0-pro model under a zero-shot setting, resulting in 464 summaries and over 1.1 million tokens. Readability was assessed using established readability indices, while semantic similarity and factual consistency were evaluated using ROUGE, BLEU, BERTScore, TF-IDF similarity, FActScore, and Sentence-BERT measures. Results show that AI-generated summaries are consistently less readable than human-written summaries across all four agencies. Although AI summaries maintain high semantic similarity and factual consistency with the source reports, they tend to employ longer sentence structures and denser vocabulary, increasing reading difficulty. Institutional differences are also observed, indicating that reporting conventions influence summary style. The findings suggest that current zero-shot LLM summarization systems preserve factual content effectively but do not automatically optimize readability, highlighting the need for readability-aware human-AI collaborative summarization in maritime safety communication.

“Such Comparison Fails To...”: The Neglected Open-Ended Negatives in the Literature Review Section of Social Science Research Articles

许晨琛 南京大学 徐 昉 北京航空航天大学

Despite the increased focus on negation in academic writing, studies mainly center on closed-class negatives (e.g., not, no, never). Nevertheless, negation can also be manifested through an open-ended array of forms (e.g., fail, lack, insufficient), which are acknowledged in natural language research yet often overlooked in academic discourse analysis. This study examines the frequency and communicative functions of open-ended negatives in contrast to closed negatives within the literature review section of 180 English research articles across six social science disciplines. Given the practical complexity of identifying open-ended negatives, the study retrieved and analyzed negation instances through a combination of a large language model and

manual inspection. It was found that open-ended negatives display a higher occurrence than closed negatives. In comparison to closed negatives, open-ended negatives are utilized more frequently in performing both informational and affective communicative functions and provide more alternatives for adjusting the force of propositions. These results suggest that the hitherto neglected open-ended negatives are crucial linguistic resources in academic writing, affording writers great flexibility in selecting expressions of negation and establishing positions with modulated force.

基于语料库的人机翻译对比研究 ——以中华学术外译项目文化类文本为例

徐道冉、翁义明 中南民族大学

本文以中华文化外译研究场景为基础，选取中华学术外译项目中的《中华文化简明读本》《中国文化要义》《中国文化通论》这三部文化类文本作为研究对象，从中挑选出思想信仰文化核心章节的 55 组节选内容，自行创建人工译本和 ChatGPT-4 机器译本的平行语料库，利用 Coh-Metrix 3.0 提取文本特征指标来开展人机翻译对比。经由研究得知，人工译本在篇章衔接、逻辑连贯以及文本可读性方面比机器译本更好，可以较好地传递出中华文化概念的深层含义，而机器译本则在词汇多样性、实词信息密度上较人工译本有优势，但是跨句语义贯通能力较差。因此本文提出以中华文化外译为基础的人机协同翻译模式，希望可以推动中华文化外译的深入发展。

生成式人工智能驱动的人智协同双语语料智能分析系统构建研究 ——基于 RAG 与多智能体框架的实践探索

徐铤伟 浙江越秀外国语学院

面向“大思政”外语教学对高质量双语平行语料与可解释性语言分析的需求，本文提出一种生成式人工智能驱动的人智协同双语语料智能分析系统，并构建基于检索增强生成（Retrieval-Augmented Generation, RAG）与多智能体协同的语料知识处理框架，以实现语料库研究从“检索工具”向“知识发现系统”的范式转型。与传统以关键词检索和频次统计为核心的语料库方法不同，该系统融合关键词检索、语义向量匹配与大语言模型推理能力，构建了“语料检索—语义理解—生成分析—知识反馈”的闭环式人机协同机制。系统支持多类型双语语料的自动导入、结构化对齐与可视化呈现，为复杂语料分析提供统一数据基础。在分析机制上，系统引入多智能体协同架构，使语料分析过程由单一工具

操作转变为多角色智能体协同完成，包括语料检索代理、语义分析代理与解释生成代理。用户可通过自然语言与系统进行交互式分析，基于检索证据驱动大语言模型生成多维度解释性结果，实现由“结果检索”向“证据驱动的语义解释”的转变。进一步地，本文提出人智协同语料分析的三阶段演进路径：数据检索（Data Retrieval）、知识发现（Knowledge Discovery）与人智协同知识建构（Human-AI Collaborative Knowledge Construction）。其中，人类研究者负责语料规范制定、术语边界判定与解释性校验，而人工智能系统承担语义扩展、模式识别与初步推理任务，从而形成可追溯、可解释的协同研究机制。研究表明，该系统不仅显著提升了双语语料分析的语义深度与交互效率，也为“大思政”外语教学语料资源建设提供了可复用的智能化技术路径与方法论框架，同时为生成式人工智能时代语料库语言学的范式转型提供了系统性证据支持。

A Corpus-Based Comparative Study of Hedges in English and Chinese Research Articles

徐迎春、娄宝翠 河南师范大学

As an essential linguistic device, hedges are widely adopted in academic writing to express standpoints and boost the persuasiveness of academic arguments. Based on two self-built corpora, this study conducts a comparative analysis of the frequency distribution and linguistic realizations of hedges in English and Chinese research articles. The results indicate that the overall frequency of hedges in English research articles is significantly higher than that in Chinese research articles, and that different hedge subcategories exhibit distinct preferences in usage between the two corpora. In terms of adaptors, intensifying adverbs are predominantly utilized in English texts to adjust the certainty or strength of claims, whereas mitigating adverbs are favored in Chinese texts to soften the tone of arguments. In terms of rounders, non-numerical types are significantly more frequent in English texts than in their Chinese counterparts. In terms of plausible shields, certain adjectives are incorporated in English texts while being notably absent in Chinese texts, indicating a difference in how tentative judgments are constructed in the two languages. In terms of attribute shields, discourse verbs are relied on more heavily in English texts to achieve mitigation, whereas research verbs are employed more frequently in Chinese texts to achieve the same hedging effect. This study not only enriches existing research on hedges in academic discourse but also offers a direct reference to the teaching of academic writing.

大语言模型学术论文结论生成的语体特征对比及预测研究

闫鹏飞 北京理工大学

本研究综合运用语料库分析和条件推断树方法,对比探究大语言模型 GPT 和专家学者所写学术论文结论的语体特性并对之进行区分预测。结果显示,两类文本均突显正式性和信息型书面语体特性,尤其是词汇的多样性、名词短语和句式结构的压缩度和复杂度;均紧扣结论部分的核心交际意图,凝练报道自身研究成果、评述自身研究优劣并指出后续研究方向,整体传递积极肯定的语义倾向,并清晰标识论述逻辑和语篇衔接。两类文本的语体特征也差异显著,并可借助多种语言特征及其组合型式进行有效区分预测。GPT 文本更为突显命题论述的抽象度、概指性与层级性以及评价立场鲜明外显特点,同时暴露出内容泛化、论调笃定、表述固化等明显不足;专家学者文本则尤为注重命题提炼的具象化和精准度、语篇逻辑的衔接性和互文性以及认识情态的思辨性和审慎度。

大语言模型辅助化学期刊英文学术论文摘要语步结构识别的研究

杨传鸣、孔祥瑞 东北农业大学

郭云洁 荆门通用航空职业技术学院 刘语轩 东北农业大学

英文学术论文摘要的语步结构识别对学术英语写作教学具有重要意义,但传统人工识别过程繁琐且耗时。尽管大语言模型为语步结构识别的自动化带来了新的机遇,但其在化学学科中的适用性尚缺乏系统评估。本研究聚焦化学 SCI 期刊的学术论文摘要,系统比较了 DeepSeek-V3.2, GPT-5.1, Gemini 2.5 Pro 和 Grok 4.1 Fast 在语步结构识别中的表现。结果显示,DeepSeek-V3.2 在识别精准性上最优, Gemini 2.5 Pro 则在识别全面性上更佳,二者优势互补;相较而言, GPT-5.1 与 Grok 4.1 Fast 虽未见明显优势,但仍具备较强的识别能力。研究表明,大语言模型在化学期刊英文学术论文摘要的语步识别中具有较高的有效性与可靠性,能够显著提升识别效率。本研究不仅有助于实现语步结构的自动化识别,更为化学学科的学术英语写作教学提供了可操作的技术支持和人机协同的教学应用路径。

Can LLMs Parse Classical Chinese? A Prompting-Based Evaluation

杨 茗 浙江大学

Dependency parsing is an essential and well-established Natural Language Processing (NLP) task. While

a number of dependency parsers have been developed, these tools are not accurate on Classical Chinese, and current Classical Chinese dependency treebanks remain predominantly manually annotated. As generative AI advances rapidly, a question arises: how reliably can these models support, or even automate, tasks traditionally requiring expert manual annotation, particularly for historical and low-resource language varieties? Recent work has investigated LLMs as dependency parsers and found them effective on contemporary, well-resourced languages, but their capacity to process Classical Chinese, a language variety central to a vast body of historical and literary heritage, remains largely untested despite growing interest in LLM-based understanding of Classical Chinese texts. This study addresses that gap by systematically evaluating state-of-the-art LLMs on Classical Chinese dependency parsing under various prompting methods, using a manually annotated Universal Dependencies treebank as the benchmark and Unlabeled Attachment Score (UAS) and Labeled Attachment Score (LAS) as evaluation metrics. We found that LLMs showed considerable variation in parsing performance, with only some substantially outperforming traditional rule-based parsers. Claude Opus 4.6 achieved the best overall performance, and chain-of-thought and in-text translation prompts significantly improved performance. These findings provide an empirical baseline for LLM-assisted parsing of Classical Chinese and inform discussions on whether and how generative AI tools can responsibly support corpus annotation work for historical language varieties.

系统功能语言学与人工神经网络的生成机制

杨明托 宁夏大学

系统功能语言学与人工神经网络在理论缘起、研究方法上有所不同，却在自然语言生成的语义建构机制上存在内在同构性。系统功能语言学以经验主义为依托，以系统网络作为语义潜势的刻画，注重语言选择的社会文化环境及概率属性；人工神经网络结合理性主义与经验主义，以词嵌入、反向传播、多层非线性变换在高维语义空间中完成语义表征及条件概率建模，与系统功能语言学在系统网络与高维向量、选择关系与权重激活、概率思想与参数化分布、功能标签与词嵌入训练上形成功能对应。系统功能语言学为神经网络提供简洁的语义功能框架及可解释性支持，神经网络为系统功能语言学形式化建模及语料自动标注提供计算路径和训练方法，跨学科融合可促进语言研究的计算化发展及语言学理论的形式化。

帕金森病语言及言语障碍研究现状与发展趋势 ——基于 Bibliometrix 的可视化分析

杨鑫悦 北京外国语大学

帕金森病的语言学研究是神经语言学领域的重要研究内容，已取得大量的理论及应用研究成果。本研究基于 SSCI 数据库文献，借助 R 语言的文献计量软件工具包 Bibliometrix 对 1985 年—2026 年期间的 433 篇帕金森病的语言学研究文献进行可视化分析。通过分析语言学视角下的帕金森病研究的高产期刊、高产作者、国家合作、关键词共现、历史被引文献、高被引文献等，本文梳理了帕金森病的语言学研究现状、热点和趋势。研究发现：（1）帕金森病的语言学研究论文数量整体呈现上升趋势。*Journal of Speech, Language, and Hearing Research* 为高产期刊，Kris Tjaden、Jessica E. Huber、Bruce E. Murdoch 等学者发文量较多。美国在该领域国际合作中的影响力最大；（2）研究热点主要集中于帕金森病患者的语言及言语障碍表现。语言研究方面主要包括句法与词汇加工能力、语篇与隐喻理解以及双语患者的母语选择性损伤；言语研究方面主要包括构音障碍，韵律与停顿异常。（3）实证研究占主导地位，主要采用声学分析、电磁发音描记仪、眼动追踪、行为任务及听觉反馈操作等研究手段。

A Quantitative Analysis of Image–Text Dynamics and Strategic Choices in Multimodal Advertising

么 娟、王一漫 北京工业大学

In the era of digital globalization, the paradigm of transnational advertising has shifted from simple linguistic conversion to sophisticated multimodal orchestration. This study explores the provocative question: Is traditional translation theory still relevant in an age of visual dominance? Utilizing DJI’s global social media advertisements as a case study, the research employs a mixed-methods approach grounded in Eco-translatology. A four-stage empirical design was implemented, involving a parallel corpus (N=202) and a technical pipeline that integrates GPT-4V and Sentence-BERT to quantify “modal contribution” through an ablation study paradigm. The findings reveal a pronounced “visual-dominance, textual-marginalization” logic: while image semantic contribution significantly and positively predicts communicative popularity, textual contribution exerts a significant negative effect. Crucially, hierarchical regression analysis demonstrates that after controlling for macro-ecological constraints—such as cultural distance and market GDP—translation strategy (domestication

vs. foreignization) per se exerts no independent significant effect on communicative outcomes. These results challenge the traditional text-centric focus of translation studies, suggesting that the visual mode now bears the primary “semantic load.” Methodologically, this research offers a replicable computational paradigm for the digital humanities, providing data-driven insights for global marketing managers to shift strategic focus from textual micro-adjustments to upstream multimodal content design.

情感问答语域中大语言模型与人类语言特征对比研究

尤 美 华中师范大学

本研究运用多维度分析法，对中文情感问答语域中 ChatGPT、通义千问与人类语言的差异进行实证分析。基于包含 1500 篇文本的平衡语料库及 51 项语言特征，研究提取出“高互动性强情感表达与低互动性弱情感表达”“主观阐释与客观建议”“主观意愿表达与客观信息传递”“高词汇复杂度整合性信息传递与低词汇复杂度浅层次信息传递”四个功能维度。分析发现，两类大语言模型在“客观建议”“客观信息传递”及“高词汇复杂度整合性信息传递”维度上表现趋同，与人类语言形成系统性对立。研究表明，大语言模型生成文本呈现“高信息整合性”与“低主观情境性”并存的特征，其语言是算法规约下的功能变体，与人类语言基于主体间性的情感表达存在本质差异。

大语言模型辅助的语步标注优化：聚焦模型间的不一致标注

余丹妮 北京外国语大学 喻儒辰 剑桥大学

Marina Bondi 摩德纳·雷焦艾米利亚大学 Ken Hyland 东安格利亚大学

近期研究表明，大语言模型在体裁语步标注方面具有潜力，但同时也提醒人们注意其可能存在的错误和不一致性。为提高大语言模型辅助标注的准确性，本研究提出了一种人机协作方法，即通过人工核验，重点解决不同轮次大语言模型生成标注中暴露出的不一致问题。在本研究中，我们考察了 GPT-4 与 GPT-4o 在标注四种应用语言学期刊中 800 篇摘要的修辞语步时所存在的不一致性。该方法揭示了两个模型在错误模式上的相似性与差异性，表明它们在语言识别方面具有不同侧重点，尤其是在处理语义模糊的实例时更是如此。这一结果部分证实了我们的猜想：不同模型标注的一致之处的正确率高，而不一致之处则提示错误标注的出现。研究结果表明，聚焦模型间不一致的人机协作方法能够显著提高标注准确性。此外，对模型不一致标注的分析有助于完善标注方案并进一步改进相关语言理论。

Rhetorical Moves and Lexico-Grammatical Features in Applied Economics Abstracts: An LLM-Assisted Corpus-Based Analysis

于银燕、张伶俐 中国地质大学

This study examines the rhetorical moves and lexico-grammatical features of research article abstracts in Applied Economics. Drawing on a self-built corpus of 60 empirical abstracts from three top journals in Applied Economics, it explores how rhetorical moves are organized and how discipline-specific lexico-grammatical features realize these moves. Using Hyland's (2004) five-move model, the rhetorical moves are mainly analyzed quantitatively by calculating move frequency, textual space, and move patterns after LLM-assisted annotation and manual validation. Lexico-grammatical features are mainly examined qualitatively through AntConc concordance lines and close textual analysis, supplemented by raw and normalized frequency counts of recurrent move-specific features. The present study reveals that Purpose, Method, and Product are obligatory moves, whereas Introduction and Conclusion are optional. Among all the five moves, product occupies the largest textual space, indicating a clear emphasis on reporting findings; meanwhile, move patterns tend to be flexible and embedded. Lexico-grammatically, the abstracts are characterized by frequent use of the simple present and active voice, selective use of passive voice in Method and Product, recurrent self-mentions in Purpose, Method, and Product, and limited but move-specific use of modal verbs. These findings have pedagogical implications for EAP instruction and offer useful guidance for novice writers seeking to produce more effective abstracts in Applied Economics.

大语言模型的构式知识表征能力研究

张 懂、蔡拂晓 北京航空航天大学

本文以大语言模型 DeepSeek 为对象,通过受控探测任务实验考察这类大模型对于不同图式性层级的英语构式表征能力。本研究采用的实验素材包括由 450 题构成的题库,由 3 类构式、3 类任务(目标构式识别、构式边界辨析、构式义解释选择)和 3 种提示条件交叉组成,每题要求模型同时输出一个答案选项和一句解释。结果显示:1)模型选择准确率为 88.4%,解释正确率仅为 49.8%,差异主要存在于实体构式和图式构式之间;2)在三种提示条件控制下,构式名称提示提高了选择概率,但并未改善解释表现;3)临造词的总体效应有限,但与构式组别发生交互,作用方向因构式而异。

Alternation Among Chinese Passive Constructions: A Multivariate Diachronic Analysis

张 懂、欧阳霖 北京航空航天大学

Chinese passive constructions can be realized through several markers, most notably *bei*, *jiao*, *rang*, and *gei*. Previous studies have primarily focused on the semantic and syntactic properties of individual passive markers, while research on their alternations has typically examined only a subset of these forms. As a result, relatively little is known about how the four major passive markers interact as competing constructions within a shared functional domain. Adopting a construction alternation perspective, this study investigates the distribution of these four Chinese passive markers and the linguistic factors associated with their selection. This study is based on a diachronic corpus spanning from the late nineteenth century to the present, with corpus data retrieved from the CCL corpus (2024 version). Passive constructions marked by *bei*, *jiao*, *rang*, and *gei* were extracted from the corpus and annotated for a range of variables. Given the relatively high frequency of *bei*, a stratified random sample was drawn to ensure comparability and analytical feasibility. Methodologically, the study combines corpus-based analysis with a multivariate framework, planning to employ Multiple Correspondence Analysis, Random Forest modeling, and logistic regression. These complementary methods enable the exploration of both overall patterns of variation and the relative contribution of individual factors to construction choice. Temporal information is incorporated as one of several explanatory variables, allowing the study to examine how construction preferences vary across historical periods and linguistic contexts. At the current stage, initial annotations were conducted with the assistance of large language models, followed by ongoing manual verification to ensure reliability and consistency across variables.

大模型辅助语料库标注框架的构建与应用： 以车载语音英中葡三语质检语料为例

张可人 香港城市大学

中国汽车品牌加速进入巴西市场，车载语音本地化质量成为用户体验的关键。然而，大模型在多语言内容生成中常出现各类系统性与语用性偏差，亟须系统化的语料库标注方案加以识别与评估。本研究基于约 2 万条英 - 中 - 葡三语车载语音质检真实语料，通过分层抽样构建 250 条子语料库，提出 PSP 原则（Partial Semantic Protocol）及 CES 四要素操作清单（操作类型、操作对象、方向 / 位置、数

值参数), 作为功能性文本翻译质量评估的统一框架。在此框架下, 本研究将错误首先区分为两个层面: 语义性执行错误(译文导致至少一个 CES 要素发生偏差, 使系统执行结果偏离用户意图, 构成“功能性错误”)与语用性语言错误(译文虽未偏离 CES 要素, 但在术语一致性、文化适应性、语域得体性等方面不符合巴西葡语规范, 构成“接受性/流畅性错误”)。此外, 对可追溯至中-英-葡三语流程中跨环节传导所致的错误, 标注“D 类(三语语流损耗)”以标记其成因。研究发现, 在 250 条标注语料中, 语义性执行错误共 53 条(21.2%), 主要集中在方向参数反转、数值参数放大及跨模态转换三类子型; 语用性语言错误共 172 条(68.8%), 覆盖 38 个独立指令, 同一指令存在多达 8 种术语变体, 严重影响语音识别引擎的关键词匹配稳定性。21 条错误可追溯至三语语流损耗, 其中约三分之二源于单一环节的翻译或生成失误, 三分之一源于跨环节的语义传导。研究表明, 基于 CES 四要素的“语义性-语用性”两级标注框架, 为大模型时代多语言语料库的系统化标注与质量评估提供了可操作、可推广的方法论支撑, 对出海产品的多语言内容质检具有直接应用价值。

A Local Grammar Approach to Evaluative Language in the Register of Fake News

张 岚、董 敏 北京航空航天大学

This study investigates how evaluative language in news of varying authenticity levels achieves persuasiveness. Integrating the Local Grammar of Evaluation, Strategic Rituals of Objectivity and Emotionality, and the concept of persuasiveness, this study combines both quantitative and qualitative methods. Utilizing a Large Language Model-assisted annotation method, it analyzes news along a five-point authenticity gradient in the Snopes Gold corpus (fully True, Mostly True, Mixture, Mostly False, and fully False). The findings indicate that news authenticity is not a simple binary opposition but a discourse continuum, along which Grammar Patterns and Local Grammar Patterns play a progressively distinguishing role. At the macro level, subjective evaluation correlates negatively with authenticity. Specifically, different authenticity levels exhibit differences in the use of evaluative grammar patterns. Fully False news prefers n v-link ADJ, and fully True news tends to use n v-link ADJ of/to, while Mostly True, Mixture, and Mostly False news rely heavily on n/it v-link ADJ that. At the micro level, local grammar analysis reveals how these patterns adapt or distort journalistic strategic rituals. Fully True news achieves rational persuasion through transparent logic and explicit source attribution, integrating objectivity and emotionality. In contrast, fully False news employs monoglossic assertions and concealed attributions, exploiting emotionality rituals for deception. Mostly True, Mixture, and Mostly False news exhibit cognitive manipulation by pseudo-objectifying subjective stances. Furthermore, analysis of shared

patterns shows that fake news employs grammatical disguise to mimic objective journalism, injecting fabricated causalities into seemingly objective structures. This study contributes to the understanding of strategic rituals in news discourse and offers potential linguistic features that can assist in automatic fake news detection.

“China Travel” 视频评论的话语建构： 基于 BERTopic 与批评话语分析的研究

赵明圆、张继光 江苏师范大学

本文以 YouTube 平台 “China Travel” 相关视频评论为研究对象，构建英文评论语料库，综合运用 BERTopic 主题建模、VADER 情感分析、词频与共现词统计等语料库方法，结合 Fairclough 批评话语分析三维框架，从文本、话语实践与社会文化实践三个层面对海外用户的话语建构机制进行分析。研究发现：（1）在文本维度上，评论语料整体呈现积极情感倾向，围绕城市空间、日常生活、普通民众与现代化景观形成区别于传统政治报道的“中国日常叙事”。（2）在话语实践维度上，评论通过视觉经验、旅行见证与互动认同形成“平台化真实性”机制。（3）在社会文化实践维度上，相关评论在一定程度上修正了西方媒体涉华叙事，但其论述逻辑仍以反驳西方媒体叙事为主。研究有助于理解数字平台环境下中国形象的海外接受机制。

人机协同的庭审问答语料库构建与语用功能标注研究

赵 奕 中山大学

庭审话语以问答互动为核心，其语用功能（如确认事实、质疑证词、引导陈述、反驳主张）的识别与标注，是法律语料库语言学中的难点。传统人工标注耗时耗力且难以保证跨标注者的一致性。本研究探索以大语言模型为辅助工具，构建人机协同的庭审问答语料库语用标注框架。研究以中国刑事庭审中的法官—公诉人—被告人/证人三方问答轮次为语料，设计了一套涵盖七类核心语用行为（宣告、探询、确认、质疑、反驳、诱导、总结）的标注方案。在具体操作中，LLM 负责对问答话轮进行初步的自动化转写分割、语句对齐及语用行为候选标注，随后由法律语言学专家对模糊案例进行复核与仲裁。研究重点评估 LLM 在识别隐性语用功能（如通过重复发问实现的质疑功能）和权力指示语（如人称代词“我”“你”“本庭”所折射的机构性身份）方面的表现。研究发现，LLM 对显性语用标记（如“是否”“对吗”）敏感度较高（ $F1 > 0.8$ ），但对依赖会话隐含和语境推理的语用功能识别能力有限，需专家深度介入。本研究的贡献在于：（1）提出一套适用于中文庭审语篇的语用标注体系；（2）检验人

机协同模式在法庭话语标注中的效度与边界;(3)为后续基于法律口语语料库的机构性话语研究提供规范化数据资源。

From Statistical Filtering to Semantic Retrieval: An Automatic Pipeline for Metadiscourse Phrase Extraction

赵泽栋 广东外语外贸大学

Metadiscourse plays a pivotal role in academic writing, yet its automatic identification has long been constrained by precompiled wordlists and manual checking, a labour-intensive approach that largely confines analysis to single words or a small set of fixed expressions, thus overlooking the pervasive phrasal realisations of metadiscourse in academic texts. Based on systematic observations of metadiscourse usage in disciplinary writing, we argue that metadiscourse most frequently occurs not as isolated lexical items but as complete phrasal units. Accordingly, we propose the term “metadiscourse phrase” and develop a three-stage automatic extraction pipeline, integrating statistical filtering, large language model (LLM) screening, and semantic retrieval on a 3-million-word corpus of marketing journal articles. Stage 1 adopts a contiguous phrase extraction algorithm informed by corpus phraseology, using an adjusted mutual information measure, boundary entropy, and a local maxima method (Wei & Li, 2013). From over 6.3 million raw 2- to 6-word N-grams, we obtain about 305,000 high-quality candidates, removing over 95% of noisy combinations. Stage 2 feeds Ädel’s (2006) identification criteria and taxonomy into an LLM, which autonomously learns the standards and classifies the candidates, yielding an initial seed set of over 3,200 metadiscourse phrases with functional labels. Stage 3 performs semantic retrieval by encoding these seeds as sentence embeddings and retrieving semantically similar variants from the embedded corpus. Retrieved expressions are added to the seed set, and the retrieval loop repeats iteratively until no new phrases emerge, approaching exhaustive coverage. Preliminary results demonstrate that statistical measures effectively reduce noise, LLM performs context-sensitive classification, and semantic retrieval uncovers semantically related variants beyond fixed wordlists. This pipeline offers a scalable framework for large-scale metadiscourse phrase mining and validates the theoretical significance of metadiscourse phrases as an independent unit of analysis, paving the way for systematic investigations into disciplinary and diachronic variations in metadiscourse use.

生成式人工智能辅助的上古汉语句法树标注生产系统

郑宇熹 北京大学

本文以《论语》句法树语料修订为例，探讨生成式人工智能辅助上古汉语句法树标注的生产系统设计。传统句法树标注高度依赖人工检查，面对复杂括号结构、核心标记、词类标签、短语结构归类和跨句一致性问题时，容易出现遗漏，且难以直接转化为可训练、可复查、可扩展的数据流程。为解决这一问题，本文提出一套由标注流程、规则库和多模型协作构成的句法树标注生产系统。系统首先将完整句法树拆解为自顶向下的层次划分、底层词类标注和自底向上的短语结构分析，以降低人工检查和模型学习的难度；其次建设词类库、短语结构库、争议处理案例库、金标准库、大模型易错材料分析库和闭源模型提示词库；再次通过调度开源微调模型、闭源模型、程序检查和人类专家，形成“模型初标—第三方复核—专家审核—系统适配复核—规则与模型迭代”的闭环。本文强调，生成式人工智能在语料库标注中的价值不只是自动生成结果，而是参与规则显性化、错误发现、候选比较和系统迭代。该实践可为古汉语语料库建设和人机协同标注方法提供参考。

人机协同视角下语料库语步标注研究： 基于本硕英语论文摘要的 IMRD 分析

仲子懿 外交学院

随着大语言模型的发展，利用 AI 辅助语料库标注成为可能。本研究以本科与硕士英语学位论文摘要为语料，探讨大模型在语步自动标注中的可靠性与应用潜力。研究基于 IMRD 模型构建包含 100 篇摘要的语料库，首先进行人工语步标注，随后引导 DeepSeek-v4 进行自动标注。通过对比人工标注与大模型标注的异同，评估大模型在不同语步识别上的准确率。研究发现，大模型在结构明确的“方法”和“结果”语步上标注一致性较高，但在语步边界模糊的“背景”与“讨论”部分表现不稳定；同时，大模型能够有效捕捉本科生与研究生摘要在语步使用频率上的宏观差异。本研究为大模型辅助语料库标注提供了实证依据，展示了人机协同的语料库研究新模式，对提升语料库标注效率及学术写作教学具有启示意义。

基于少样本学习的隐喻自动识别研究

周穗芳 广东外语外贸大学

隐喻识别存在主观性强、耗时高、自动化难度大的困境。大语言模型可通过自然语言提示执行复杂语言分析，为突破这一局限提供了契机。本研究据此以精细化提示工程引导 GPT-5.1 遵循 MIPVU 流程，对 VUA 隐喻语料库小说语篇进行隐喻自动识别与评估，并补充学术、新闻、会话语篇以检验跨语域迁移能力。结果表明：增加兼顾典型与边缘的异质性示例，可提升模型的隐喻识别能力；少样本提示通过任务约束、能力激活、分步推理与输出收束等可供性促使模型稳定地遵循 MIPVU 流程，但在高度常规化的间接隐喻与隐喻 - 转喻边界上仍存在系统性漏标；跨语域结果则显示该提示具有较高的迁移能力。本研究为基于大语言模型的隐喻识别提供了有益启示。

藏东南区域语言关系计算研究 ——基于核心词的编辑距离算法与历史比较法的互补验证

周潇然 西藏大学

藏东南地区是我国语言多样性最丰富的区域之一，其语言关系研究对理解汉藏语系演化具有重要意义。本文收集了 27 种藏东南语言 / 方言，构成斯瓦迪士 100 核心词语料库，采用 ASJP 模式的莱文斯坦编辑距离计算语言相似度，并通过邻接法构建谱系树，同时辅以历史比较法进行交叉验证，希望能为一些语系归属尚不明确的语言提供归属建议。结果显示：(1) 藏东南语言可分为藏语方言、门珞语言、僜人语言三大语言群；(2) 系属不明的如美话、松林语在聚类中与僜人语言关系密切，建议归入景颇语支；(3) 扎话处于门巴语与僜人语言之间的过渡位置。研究表明，编辑距离方法可为语言分类提供客观的量化依据，可以与历史比较法构成一个互补的、适用于藏东南区域语言的语言关系计算框架。

生成式人工智能辅助下的多模态符际翻译与叙事节奏重构研究 ——以电影《银翼杀手》跨媒介转译为例

朱宇吟 香港理工大学 汪雄飞 广东外语外贸大学

1982 年，根据菲利普·迪克的小说《仿生人会梦见电子羊吗？》改编的科幻电影《银翼杀手》上映，影片开赛博朋克电影先河，在之后的几十年间被评为经典。本研究以《仿生人会梦见电子羊吗？》

至《银翼杀手》的跨媒介转译为考察对象，引入生成式人工智能（GenAI）辅助多模态互动分析。以“罗伊之死”等核心场景为例，本研究利用大语言模型对字幕（文本）、空间建构（视觉）与环境音效（听觉）进行结构化的精准对齐与参数化标注。结合多模态社会符号学，研究发现，相较于原著单一文本模态下的内聚焦叙事，电影通过视听模态的“互补”“强化”与“冲突”机制，重塑了仿生人的主体性体验。GenAI 的介入有效量化了模态间的张力关系，揭示了符际翻译如何通过拉长叙事停顿、转换叙事视角等手段来改变原著的叙事节奏。本研究不仅拓宽了多模态翻译的数字化研究路径，也为理解科幻经典的跨媒介重构提供了微观视角的实证支撑。

ChatGPT 看图作文叙事特征研究

朱周晔 北京外国语大学

本研究以记叙类看图作文为语料，以 ChatGPT-4o 为对象，从宏观结构和微观结构两个维度对大语言模型生成的叙事文本进行分析。研究发现：与汉语二语者相比，ChatGPT 在宏观结构上使用更多的人物、结局和尾声要素，叙事框架更为完整；在叙事顺序上使用更少的因果连接，逻辑推进更为隐性；在叙事评价上使用更多的感叹词、重叠与副词，评价手段更为密集。本研究有助于深化对大语言模型叙事能力的认识，也为人工智能赋能二语写作教学提供了实证参考。



教育部人文社会科学重点研究基地

北京外国语大学中国外语与教育研究中心

National Research Centre for Foreign Language Education
Beijing Foreign Studies University

北京市海淀区西三环北路 19 号 邮编: 100089

19 Xisanhuan Beilu, Beijing 10089, China

Tel: 010-88816824 Fax: 010-88810376

Email: wjyzx@bfsu.edu.cn

Web: <https://sinotefl.bfsu.edu.cn>

生成式人工智能时代的语料库语言学研讨会

邀请函

尊敬的 _____ 女士/先生:

您好!

经专家评审,您提交的论文摘要已被“生成式人工智能时代的语料库语言学研讨会”录用。谨向您表示热烈祝贺,期待您在会议期间的精彩分享与交流。

本次研讨会由中国英汉语比较研究会语料库语言学专业委员会与北京外国语大学中国外语与教育研究中心联合主办,外语教学与研究出版社承办,《语料库语言学》杂志协办。会议将于 2026 年 8 月 9 日至 11 日于北京大兴北京外国语大学国际会议中心举行。会议旨在提升生成式人工智能时代我国语料库语言学研究水平。诚邀各方学者莅临交流。现将相关事项通知如下:

报到时间: 2026 年 8 月 9 日 14:00—20:00; 2026 年 8 月 10 日上午 8:30 前

报到地点: 北京大兴北京外国语大学国际会议中心大堂

住宿地点: 北京大兴北京外国语大学国际会议中心宾馆

会议时间: 2026 年 8 月 10 日全天及 8 月 11 日上午

会议费用: 普通参会人员 1000 元/人,学生 500 元/人(交通住宿费用自理)。

期待您拨冗莅临,与学界同仁分享研究成果,探讨语料库语言学的前沿进展,共襄学术盛举。如有疑问,请与我们联系。

组委会邮箱: bfsucrg@sina.com

中国英汉语比较研究会语料库语言学专业委员会

北京外国语大学中国外语与教育研究中心

2026 年 7 月 15 日





语料天涯
北外语料库团队官方公众号