

《经济学人》社论的机器翻译及译后编辑问题探析

袁媛

四川大学外国语学院, 四川 成都

收稿日期: 2026年6月28日; 录用日期: 2026年7月28日; 发布日期: 2026年8月11日

摘要

本文旨在研究基于LLM的机翻系统在《经济学人》社论英译汉中所出现的典型问题及译后编辑重点。本文首先简要回顾了机器翻译与译后编辑的相关研究, 以及新闻社论的文本特点与翻译标准, 继而采用个案研究法与对比研究法, 以《经济学人》2026年5月刊社论“From cyber-security to biosecurity”的英译汉实践为研究对象, 对比DeepSeek-V4-Pro、豆包Doubao-Seed-2.1-Turbo模型的输出译文与人工译后编辑版本, 从专业术语、搭配与表达、句式结构、隐喻修辞等维度系统考察了AI翻译在新闻社论中的典型问题。研究发现, 两个系统在术语的概念化处理、语序调整、被动语态转换、定语从句逻辑显化、对隐喻修辞的把握等方面存在不足, 而译后编辑的核心任务可概括为“去欧化、显逻辑、调语气”。本文指出, 在“数智翻译时代”, 人类译者应与机器翻译系统保持“竞合”状态, 一方面充分利用技术带来的效率提升, 另一方面持续强化自身在英汉对比意识、语体敏感度和语境推理能力方面的优势。

关键词

《经济学人》, 机器翻译, 译后编辑, 新闻社论, 英汉对比

An Analysis of Machine Translation and Post-Editing Issues in *The Economist's* Editorials

Yuan Yuan

College of Foreign Languages and Cultures, Sichuan University, Chengdu Sichuan

Received: June 28, 2026; accepted: July 28, 2026; published: August 11, 2026

Abstract

This paper aims to investigate the typical problems and post-editing priorities of LLM-based

machine translation systems in the English-to-Chinese translation of *The Economist* editorials. It begins with a brief review of relevant studies on machine translation and post-editing, as well as the textual features and translation standards of news editorials. Adopting a case study approach and a comparative research method, the study takes the English-to-Chinese translation practice of the editorial “From cyber-security to biosecurity” published in the May 2026 issue of *The Economist* as its research object. It compares the output translations of the DeepSeek-V4-Pro and Doubao-Seed-2.1-Turbo models with a human post-edited version, and systematically examines the typical problems of AI translation in news editorials from four dimensions: terminology, collocation and expression, syntactic structure, and metaphorical rhetoric. The findings reveal that both systems exhibit deficiencies in the conceptualization of terminology, word-order adjustment, passive-voice conversion, logical explication of attributive clauses, and the handling of metaphorical rhetoric. The core tasks of post-editing can be summarized as “de-Europeanization, logical explication, and tonal adjustment”. This paper argues that, in the age of digital intelligence translation, human translators should maintain a state of “coopetition” with machine translation systems: on the one hand, they should make full use of the efficiency gains brought by technology; on the other hand, they should continuously strengthen their advantages in English-Chinese contrastive awareness, stylistic sensitivity, and contextual reasoning ability.

Keywords

The Economist, Machine Translation, Post-Editing, News Editorials, English-Chinese Contrastive Analysis

Copyright © 2026 by author(s) and Hans Publishers Inc.

This work is licensed under the Creative Commons Attribution International License (CC BY 4.0).

<http://creativecommons.org/licenses/by/4.0/>



Open Access

1. 引言

近年来,大语言模型(LLM)技术的飞速发展正在深刻重塑机器翻译的范式。以 DeepSeek、豆包等为代表的 LLM 驱动型机翻系统,凭借其在上下文理解、多任务泛化等方面的显著优势,逐渐被应用于各类文本的翻译实践之中。然而,现有研究多聚焦于 LLM 在通用领域或文学翻译中的表现,对于其在新闻社论这一特殊文本类型中的适用性与局限性,尚缺乏系统性的考察。新闻社论兼具信息传递与观点表达的双重功能,其语言高度凝练、逻辑严密、修辞丰富,对翻译的准确性、流畅性与风格再现均提出了较高要求,因而成为检验机翻系统综合能力的理想试金石。

《经济学人》作为全球最具影响力的时政财经周刊之一,其社论代表了英语新闻评论写作的高水准,亦是中国读者了解国际视野的重要窗口。本研究以《经济学人》2026 年 5 月刊社论 “From cyber-security to biosecurity” 的英译汉实践为研究对象,采用个案研究法与对比研究法,系统对比 DeepSeek-V4-Pro、豆包 Doubao-Seed-2.1-Turbo 模型(测试时间为 2026 年 6 月 27 日)的输出译文与人工译后编辑版本,从专业术语、搭配与表达、句式结构、隐喻修辞等维度考察机翻在新闻社论中的典型问题。研究进一步提炼了译后编辑的核心任务,即“去欧化、显逻辑、调语气”,并在此基础上探讨“数智翻译时代”人机关系的重构路径。

本研究的意义体现在理论与实践两个层面。在理论层面,研究丰富了 LLM 机翻在新闻评论类文本中的实证研究,为翻译技术研究提供了新的观察样本与分析框架;在实践层面,研究揭示的典型问题与译后编辑策略,可为新闻社论类文本的机翻选型、译后编辑规范制定以及译者培训提供参考,亦有助于推动人工智能时代人机协作翻译模式的优化与完善。

2. 文献综述

机器翻译(Machine Translation, MT)是指利用计算机系统将内容从一种自然语言自动转换为另一种自然语言的过程[1], 已在诸多领域辅助或替代人工翻译[2]。然而, 机译输出往往难以直接达到专业出版标准, 由此催生了译后编辑(Post-Editing, PE)环节——即对机器翻译输出进行校订和修正。围绕机译错误的识别与分类、译后编辑应遵循的原则及译者所需的能力, 学界已积累较为丰富的研究成果。

2.1. 机器翻译错误类型研究

早期国内研究多关注统计机器翻译与规则机器翻译阶段的技术缺陷, 针对不同文本类型的机译错误展开了持续考察。罗季美与李梅(2012)通过对科技类文本的分析指出, 机器翻译的错误表现具有多层次性, 主要体现在词汇选择不当、句法结构错乱、标点符号使用不规范以及语篇衔接失效四个层面[3]。该研究首次在国内系统构建了“词汇-句法-符号-语篇”四层次的错误分析框架, 强调机器翻译问题不仅停留在微观语言单位层面, 更会蔓延至宏观文本组织。

在这一分析框架的基础上, 具体错误类型的归纳成为后续研究的重点。崔启亮与李闻(2015)基于科技文本的英汉机器翻译, 提炼出 11 种错误类型, 涵盖术语翻译错误、过译、欠译、漏译、多译、冗余、形式错误、格式错误、词性判断错误、从句翻译错误以及受英语句子结构的束缚[4]。同期, 包凯(2017)则从谷歌翻译英译汉的实践出发, 从句法完整性角度归纳出 10 种常见错误类型, 包括缺少主语、缺少谓语、句子不全、动词重复、连词重复、短语重复、主谓不一致、短语错误、词性误用、百科知识不一致[5]。比较而言, 崔启亮与包凯的分类维度略有不同, 但两者在“术语准确性”、“词性判断”、“句法完整”三大核心维度上高度重合, 这一共识在后续研究中得到延续与验证。李国兵和张敬(2022)针对物流论文摘要的机器英译问题再次确认了术语翻译错误和逻辑连贯性问题是机器翻译的主要短板[6]。

2022 年以来, ChatGPT 等大型语言模型(LLMs)的出现使机器翻译技术实现质的跃升。与传统的神经机器翻译(NMT)相比, LLMs 具备更强的上下文理解能力和文本生成能力, 这使得机器翻译错误类型发生了结构性变化。传统研究中被频繁报告的明显语法错误显著减少, 取而代之的是更为隐蔽的“语义幻觉”与“事实性偏差”等问题。

在此背景下, 学界开始系统开展 ChatGPT 翻译质量的实证评估研究。其中, 耿芳和胡健(2023)以 ChatGPT 生成的英汉与汉英翻译实例为研究对象, 从错误修正、句法结构、篇章连贯及语域调整四个维度考察了 ChatGPT 的译后编辑表现, 发现其能够根据指令修正翻译错误、增强语篇连贯性、调整语域风格、提高句法复杂度, 但也指出 ChatGPT 在处理文化负载词和复杂句式时仍存在不稳定性[7]。整体而言, 在语言流畅度和语篇连贯性上显著优于传统神经机器翻译, 但对复杂论证结构、文化负载词及特定文体的处理仍存在稳定误差。更为关键的是, 该研究揭示了大语言模型翻译错误的双重性: 既存在拼写、用语、句法等表达层面的“形式性误译”, 也存在因语境理解不足或文化知识缺失导致的“内容性误译”。后者表现为表面通顺但实质偏离原文意义的“隐性错误”, 其识别难度更高、隐蔽性更强。这一发现表明, LLM 的错误类型与以往机译既有重叠亦有新变, 为界定 LLM 翻译的译后编辑层级、重新定义译后编辑者的工作重点提供了重要实证依据。

2.2. 译后编辑的原则与能力模型

国际标准化组织(ISO)对译后编辑(Post-Editing, PE)给出明确定义: “译后编辑是对机器翻译输出进行校订和修正的过程”[8]。这一定义强调译后编辑并非独立的翻译行为, 而是建立在机器翻译产出基础之上的二次加工活动, 其本质是以经济高效的方式达到预设质量目标的人机协同过程。

在 ISO 定义基础上, 国内学界结合本土实践逐步完善了译后编辑的原则体系。早期研究普遍认同的

核心原则包括：忠实原文、语义正确、尽量保留机译产出、不随意修改风格等[9]。一些学者尝试将纽马克“翻译四层责任”、“信达切”等传统翻译理论引入译后编辑的指导框架[10][11]，译后编辑的原则既可从行业标准中提取，也可借鉴成熟翻译理论的洞见。然而，也有学者提醒注意原则适用的情境约束。

在译后编辑原则逐渐清晰的同时，对其从业者能力构成的系统研究也成为学术热点。冯全功和张慧玉(2015)较早从语言服务行业角度切入，探讨了译后编辑能力的构成要素及高校人才培养方案，强调译后编辑者需兼具语言敏感度与技术适配力[12]。在此基础上，冯全功和刘明(2018)正式提出译后编辑能力的“三维模型”，即语言能力、工具能力、策略能力的整合框架[13]。随着大语言模型技术的快速渗透，原有的三维能力模型面临新的挑战。王律和王湘玲(2023)在《ChatGPT时代机器翻译译后编辑能力培养模式研究》一文中，从理论上构建了“ChatGPT + MTPE”交互式译后编辑能力培养模式，并以MTI笔译硕士为研究对象进行了实证探索[14]。杨艳霞和魏向清(2023)则从认知范畴观视角，对机器翻译译后编辑能力进行了系统的解构与培养研究，为能力模型的深化提供了认知层面的理论支撑[15]。然而，在LLM技术深度嵌入翻译实践的背景下，译后编辑能力的内涵如何进一步拓展、各能力维度如何适应人机协作的新要求，仍有待深入探究。

2.3. 本研究的定位

综观上述成果，学界在LLM翻译错误类型研究方面已取得显著进展，但仍存在若干值得关注的局限。其一，研究对象集中于科技文本、通用文本等，对新闻社论这类修辞密集、论证复杂的评论性文体关注不足。其二，现有错误类型分类多以技术文本为语料基础，其分类体系对修辞性文体中常见的隐喻处理、逻辑显化等问题覆盖不全。其三，多数研究止步于错误类型的归纳，未将错误分析与文本功能理论进行有效对接，缺少“何种错误损害何种功能”的系统分析。

针对上述局限，本研究的学术定位如下：在研究对象上，本研究将LLM翻译错误分析的文体视野从科技文本、通用文本拓展至新闻社论这一修辞密集、论证复杂、评论性“高功能复合型文本”。在分析框架上，本研究引入赖斯的文本功能理论，将错误分析从“语言形式的正误”提升至“文本功能的损害”层面，系统考察每个翻译问题对信息功能、表情功能、操作功能的不同损害机制，弥补现有研究“有分类无功能”的不足。在研究发现上，验证崔启亮与牛硕(2023)关于LLM“隐性错误”的核心判断在新闻社论文本中的普遍性的同时，补充该理论在修辞处理、逻辑显化等维度上的具体表现形态，将LLM翻译错误类型研究从通用描述推向文体细化。

3. 《经济学人》社论的文本特征

德国功能派翻译理论家赖斯(Katharina Reiß)在《翻译批评：潜力与制约》一书中提出，等值的概念不应只停留在字、词、句的微观层面上，而应该涉及文本(语篇)层面[16]。赖斯在布勒(Karl Bühler)语言功能三分法基础上，将文本划分为信息型(informative)、表情型(expressive)和操作型(operative)三种主要类型，并指出文本类型从根本上决定了翻译策略的取向。赖斯进一步强调，多数文本兼具多种功能，但三者往往并不均衡，识别文本的主导功能是翻译的关键前提，因为传递原文的主要功能是衡量译文的决定性因素。识别主导功能并非否定其他功能的存在，而是为了在翻译中确立策略重心，从而更有效地处理不同层级的信息[17]。

以Leaders栏目为代表的《经济学人》社论正是典型的多功能复合型文本——既传递专业信息，又展现独特风格，更以说服读者接受其立场为根本目的。

信息功能是社论的基础。Leaders栏目以深度报道探讨全球重大议题，涉及政治、经济、科技等领域，须清晰传达事实脉络与专业知识。以本文考察的社论为例，文中密集出现“reverse-engineering a cell type

from raw DNA data”、“synthesise viruses”、“mirror life”、“DNA synthesis”等专业术语，以及“90% of the novice participants”等具体数据。这些术语与数据是论证链上的关键节点，任何误译都可能导致逻辑断裂。此外，《经济学人》自诩语言以“clarity, style and precision”为特征，其中“precision”不仅要求术语与数据准确，还体现在认知模糊限制语的精心运用上，以调节命题确定程度，保持论述严谨性。

表情功能在社论中尤为突出。该刊以英式文风著称，语言凝练犀利、幽默睿智，善用隐喻、反讽、双关等修辞。隐喻不仅是修辞点缀，更是其独特的认知框架。本文中，标题“From cyber-security to biosecurity”暗示威胁焦点的转移；“book smart”与“at the laboratory bench”形成知识与实践的对照；“crack open the ‘black box’”以“黑箱”指代神经网络的可解释性问题；结语“Petrifying dishes”则以双关同时指向“培养皿”与“令人恐惧”。这些修辞手段承载着杂志的立场与态度，是其品牌辨识度的核心。句法上，社论呈现典型的英语名词化倾向，多用抽象名词、被动语态、复杂从句，同时频繁使用碎片句和祈使句以增强强调效果。这些共同构成了《经济学人》独特的风格。

操作功能构成社论的核心意图。Leaders 栏目的根本目的并非中立报道，而是集中表达立场并试图说服读者。本文中“Gambling the future of humanity on such defences would be a mistake”等表述，以强烈语气传递警示与呼吁。论证策略上，通过“worse than”、“far less malleable”等比较构建紧迫感，通过“One option is...Another measure is...A third idea is...”的递进结构梳理方案并逐一指出不足，引导读者认同“scientific breakthroughs will therefore be needed”的结论。这种劝说功能要求译文不仅传达“说了什么”，更要让目标语读者产生与源语读者相近的紧迫感与认同感。

综上，该刊社论的三重功能相互交织。精确信息为说服奠基，鲜明修辞为观点注入感染力，强烈劝说意图统摄信息选择与表达方式。这种复合性使其翻译难度远超纯信息型文本，译者须兼顾信息准确、风格再现与功能等效，机翻系统亦面临术语处理、修辞识别与逻辑传递的综合挑战。这正是本研究选取其作为考察对象的根本原因。

4. 《经济学人》社论的机器翻译的典型问题及译后编辑的优先级原则

4.1. 专业术语的处理

术语翻译是《经济学人》社论翻译中的基础性环节，也是 AI 翻译问题的高发区域。AI 系统在处理专业术语时，倾向于字面对应而非概念重组，这在新闻社论翻译中尤为突出。应当承认，字面对应在处理部分约定俗成的术语时具有一定合理性，如“mirror life”直译为“镜像生命”，能够准确传递概念。然而，当术语承载复杂背景或抽象含义时，字面对应往往无法实现概念层面的重组，导致译文表面通顺而实质偏离原文意义。

【例 1】

原文：**Responsible researchers** should be able to use AI to advance the frontiers of science-DeepMind’s Isomorphic Labs is developing novel cancer therapies, for example-but under security protocols.

审校版：对于**有正当研究需求的相关人员**，应允许其利用 AI 推动科学前沿的发展……但必须在安全协议下进行。

分析：DeepSeek 将“Responsible researchers”译为“负责任的科研人员”，豆包译为“负责任的研究者”，均取 responsible 前置定语的首选义“可信赖的、尽职的”。然而，这一译法与上下文逻辑存在冲突：前文强调政府必须限制系统访问权限以防止生物恐怖主义，后文以“but”引出转折——“但必须在安全协议下进行”。若译为“负责任的(人品好的)人员”，“负责任”与“受协议限制”之间并不构成天然的转折关系，人品好不代表可以豁免监管。该误译主要损害文本的信息功能。原文中“Responsible”在

此语境中的精确含义是“有正当研究权限的”，而非通用的道德评价。机译将其泛化为“负责任的”，使读者误解作者对研究人员群体性质的判断，论证链条中“权限管控”的核心信息被“人品评价”所遮蔽，导致信息焦点从“准入制度”偏移至“道德判断”，信息功能的精确性受到损害。人工译后编辑将其处理为“有正当研究需求的相关人员”，其优势体现在三个层面：在语法层面，准确把握了 *responsible* 在权限分配语境中的实际指向——即“被授权从事相关研究的人员”；在逻辑层面，“相关”与后文“*but*”形成严密的转折链：即使是确有研究需求的相关人员，也必须受协议约束；在语体层面，“有正当研究需求的相关人员”的译法显化了原文隐含的准入条件，符合学术文献中“*responsible authorities*”约定俗成译法。

【例 2】

原文：If one engineered virus may cause billions of deaths, humanity has no room to learn from mistakes.

审校版：一种人为设计的病毒就可能夺去数十亿人的生命，人类没有任何试错的空间。

分析：DeepSeek 的译文“工程病毒”，语义含混，且不符合汉语表达习惯；豆包的译文“人工改造的病毒”，虽语义准确，但“改造”一词偏中性，未能传递原文隐含的“人类主动制造危险”这一语义焦点。这一信息功能的缺损使读者对病毒来源属性产生模糊认知。“改造”的平缓语气同时削弱了原文的警示力度，读者难以产生与源语读者同等的紧迫感，而正是这种紧迫感才使“没有试错空间”的结论具有说服力，操作功能因此受损。人工译后编辑将其处理为“人为设计的病毒”，其优势体现在两个方面：在语义层面，“人为”点明责任主体，与后文“人类没有任何试错的空间”形成紧密的因果逻辑，“设计”则比“改造”更具意图性，强化了“主动为之”的警示意味；在语体层面，七个字的篇幅短促有力，契合社论警句的文体要求。

4.2. 搭配与表达的自然度

中英两种语言在搭配习惯上存在显著差异，而 AI 系统倾向于保留英语的搭配结构进行直译，导致机译文中大量出现不符合中文惯用搭配的“翻译腔”，极易产生“语义可解但表达生硬”的译文。搭配层面的翻译问题虽不似术语误译那般直接导致信息错误，却深刻影响译文的自然度与可读性。

【例 3】

原文：Artificial intelligence (AI) will soon add biology to its list of superhuman abilities.

审校版：人工智能(AI)很快或将把生物学能力也纳入其“超人”本领之中。

分析：DeepSeek 的处理是“将生物学也纳入其超人能力的清单”，而豆包的处理是“将生物学领域纳入自身的超能力范畴”。两个译文均未能处理 *biology* 的转喻含义——英语习惯用学科名称(如 *biology*、*physics*、*chess*)直接指代该领域的能力或技能，例如“*He added math to his strong subjects*”中的“*math*”实指数学能力。这是一种英语中非常自然的转喻思维，读者会自动补全“能力”的语义。但汉语译文如果没有“能力”二字，直译为“把生物学加入其超人本领清单”，会产生逻辑跳跃。因为“生物学”是一门学科，不是一种“本领”。汉语读者需额外推理才能理解此处指“生物学能力”，这种理解障碍降低了信息传递效率，损害了信息功能。在搭配方面，DeepSeek 的“纳入……清单”和豆包的“纳入……范畴”均保留了英语的搭配结构，较为生硬，与《经济学人》社论追求的风格相悖，损害了表情功能。人工译后编辑将其处理为“把生物学能力也纳入其‘超人’本领之中”，此修改包含两处调整：一是增补“能力”一词，将抽象的学科名称具体化为可被“拥有”的能力范畴；二是以“纳入……之中”替代“纳入……清单”，因为“纳入……之中”在汉语中本身已暗示集合或范畴的存在，无需“清单”一词赘述。

【例 4】

原文: Software can be fixed quickly, but human biology is far less malleable.

审校版: 软件尚可快速修补, 但人体生物系统远没有软件那样容易修补。

分析: DeepSeek 译文为“软件可以快速修复, 但人类生物学远没有那么强的可塑性”, 而豆包的译文是“软件漏洞可以快速修复, 但人类的生理机制却远没有这般可塑性”。两者均用两个不同的词“快速修复”、“可塑性”处理了原文中语义相近的两个词“fixed quickly”、“malleable”, 这本身在语义上没有问题, 但割裂了原文的对比逻辑。英语喜用不同词汇避免重复, 而汉语则喜用重复词汇增强节奏和对比[18]。原文正是通过“fixed quickly”与“malleable”的语义接近性构建类比反差, “软件可快速修补”与“人体系统不易修补”构成鲜明对比, 这是说服读者接受“AI 安全风险远高于软件安全”的核心修辞手段。机译用两个不同动词削弱了这一反差, 句法平行强化论证的操作功能未能有效传递。人工译后编辑用同一个动词“修补”进行搭配, 既传达了原意, 又使前后对比更加紧凑有力。同时, 将 human biology 具体化为“人体生物系统”, 使“修补”的搭配在汉语中得以成立, 因为“生物学”本身无法被“修补”, 而“生物系统/生理机制”则可以。

4.3. 句法结构的逻辑显化

英语重形合, 多用右分支结构、物称与被动; 汉语重意合, 多用左分支结构、人称与主动。AI 系统在处理《经济学人》社论的长句和复杂句式时, 往往过度保留英语的句法特征, 导致译文逻辑关系隐晦。而这个问题的主要体现在于英语名词化倾向的机械直译。英语名词化倾向是其静态特征的集中体现, 在《经济学人》社论中尤为突出——抽象名词承载着比较、因果、判断等逻辑关系, 使句式凝练, 但对汉语读者而言却显得抽象而隔阂。

【例 5】名词性比较结构的动词化转换

原文: The threat from new pathogens is an even graver danger than AI-backed hackers.

审校版: 比起被 AI 强化的黑客, 更可怕的是新型病原体。

分析: DeepSeek 译为“新型病原体带来的威胁比 AI 支持的黑客更为严重”, 豆包译为“新型病原体的威胁比 AI 辅助的黑客更加严峻”, 两者均保留了英语的“A 比 B 更……”的语序, 主语“威胁”承载了过多修饰成分, 信息焦点分散, 读者需费力提取“病原体 > 黑客”这一比较关系的核心信息, 信息功能受损; 比较力度也因此被削弱, 原文“生物安全比网络安全更凶险”这一核心论点的说服力下降, 操作功能未能有效传递。而汉语更喜欢把已知对比项前置, 用“比起 B, A 更……”的结构, 这更符合汉语“先已知后焦点”的信息排列习惯。人工译后编辑将原文主语移到了后半句, 比较关系更清晰, 语气更强。

【例 6】名词性表语结构的动词转换

原文: One promising approach is the equivalent of brain surgery on models after they are trained.

审校版: 其中一种可行的方法, 相当于在模型训练完成后对其进行“脑外科手术”。

分析: DeepSeek 保留了“is the equivalent of”的名词性表语结构, 译为“一种有前景的方法是在模型训练后进行脑部手术的等同物”, 句式生硬且不符合汉语表达。这一译法使读者难以快速把握“方法 A 相当于手术 B”的比较关系, 信息功能受损; 而“等同物”的欧化表达与《经济学人》简明犀利的文风格格不入, 原文凝练有力的判断沦为生硬拗口的术语堆砌, 表情功能同样未能实现。豆包和人工译后编辑的处理类似, 以动词“相当于”取代整个名词性结构, 将静态名词比较转为动态动词判断。这一处理

符合英语惯用“is the equivalent of”、“is a reflection of”、“is an indication of”等名词结构表达比较、汉语则直接用“相当于”“反映”“表明”取而代之的转换规律，句式更简洁，逻辑更直接。

【例 7】定语从句的逻辑显化

原文：This matters especially for open-source models, which cannot be recalled once they have been disseminated, and whose use cannot be monitored.

审校版：这一点对于开源模型尤为重要，因为一旦开源模型被传播，便无法召回，其使用过程也难以监控。

分析：DeepSeek 译为“这一点对开源模型尤其重要，开源模型一旦被传播就无法召回，并且其使用也无法被监控”，以逗号连接“重要”与两个并列定语从句，虽语法正确，但因果关系仅靠逗号暗示，汉语读者不易察觉“为什么重要”与“不能召回、无法监控”之间的因果链条，因果链的隐化使原文“因为开源模型不可控，所以必须限制”的严密论证被模糊，说服力下降，操作功能未能有效传递。豆包处理和人工译后编辑相似，都以“因为”明确标示因果，并将“一旦……便……”的条件关系嵌入因果框架中，使“之所以重要，是因为风险不可控”的完整逻辑链条一目了然。

【例 8】被动语态的转换与施动者的补全

原文：Scientific breakthroughs will therefore be needed, to create new kinds of safeguards.

审校版：因此，人类仍需依赖科学突破，以建立新型安全防护机制。

分析：DeepSeek 和豆包均采用汉语无主句处理，回避了“被需要”的直译形式。此类译法在汉语中虽可接受，但缺失了“被谁需要”这一施动者信息。“被需要”的主体缺失使信息传递的精确性受损，信息功能未能充分实现；无主句的弱表达同时消解了原文“人类作为责任主体”的紧迫感，原文中“科学突破被需要”的深层语义是“人类必须主动寻求突破”，这一责任主体的强调正是唤起行动的核心机制，操作功能因此失效。人工译后编辑增补“人类”为主语，将被动转为主动，不仅补全了施动者信息，使语义完整，且“依赖”一词亦比“需要”更具力度，贴合原文隐含的紧迫感。

4.4. 修辞效果的再现

《经济学人》社论善用双关、拟人等修辞手法来增强文章的表现力和说服力。然而，AI 系统在处理这些修辞现象时，往往只能进行字面转换，缺乏对修辞效果的识别与再现能力。本节围绕双关和拟人两类修辞现象，分析 AI 翻译的典型问题及译后编辑的处理策略。

【例 9】双关的语体适配

原文：Petrifying dishes

审校版：培养皿里的危机

分析：原文中这个小标题是一个双关语，其修辞效果建立在两层含义之上。一是字面义“Petri dishes”（培养皿），指向生物学实验场景，二是修辞义“petrifying”（令人石化的、极度恐惧的），传达警示意味。原文通过利用英语的音形巧合，将“Petri”改写为“Petrifying”，用一个词的变形同时承载了“实验场景”和“恐怖警示”两层信息，在汉语中完全无法复制。豆包将其译为“惊魂培养皿”，试图保留警示语气，但“惊魂”一词偏口语化，与《经济学人》社论的严肃文体不够匹配。DeepSeek 未对此小标题进行单独处理。机译未能再现原文的音形双关，原文表情功能中智性的文字游戏与情感的预警信号随之流失，这正是双关修辞的精妙所在，其不可译性对表情功能构成了结构性损害。人工译后编辑将其处理为“培养皿里的危机”，舍弃了双关的修辞形式，但通过“培养皿”保留科学意象、通过“危机”传递警示语气，以功能等效策略实现了表情功能中“风格识别”与“情感传递”的核心目标，契合了社

论结语的点题作用。

【例 10】拟人化修辞的消解

原文: Another technique teaches models to favour wrong answers in some areas.

审校版: 另一种技术则是让模型在某些特定领域倾向于给出错误答案。

分析: DeepSeek 译为“另一种技术教会模型偏爱某些领域的错误答案”, 豆包译为“另一种技术训练模型在某些领域偏好错误答案”, 两者均保留了原文的拟人化修辞。“教会”暗示“有意识地教”, “偏爱”、“偏好”带有情感色彩, 用于机器显得怪异。英语科技文本中, 机器常被赋予人类行为动词 (teach, favour, think, know), 这在英语中完全自然, 但直译至汉语时会产生过度的拟人色彩。人工译后编辑将其处理为“让模型倾向于给出错误答案”, 包含两层调整: 其一, 将拟人动词“teach”替换为中性使令词“让”, 消解了“有意识教授”的拟人色彩; 其二, 将情感化动词“favour”替换为客观描述“倾向于给出”, 消除了“偏好”的情感指向。这一调整同时实现了表情功能的跨语言适配——英语科技文体允许适度的拟人化修辞以增强表达活力, 而汉语科技文体更倾向于客观中性表达, 译后编辑将拟人化表达转为客观描述, 使译文的表情功能从英语科技文体的“活力型”调整为汉语科技文体的“平实型”, 实现了语体规范的跨语言转换。

4.5. 译后编辑的优先级原则

基于本案例的实践, 译后编辑的介入可按优先级分为三个层次。第一优先级是必须修改的层面, 涉及影响信息准确性的错误, 包括术语误译、逻辑关系错误、语义偏差, 这类错误使读者无法准确获取原文的事实信息, 动摇翻译的“信”之根本。第二优先级是建议修改的层面, 涉及影响可读性与说服力的表达问题, 包括语序生硬、被动语态过度保留、名词化结构的直译等欧化表达, 这类问题虽不直接导致信息错误, 但损害了译文的操作功能, 降低了读者接受论点的意愿。第三优先级是可选修改的层面, 涉及风格层面的润色, 包括修辞力度的强化、插入语的位置调整等, 这类修改主要服务于表情功能的完善, 在时间有限时可暂缓处理。这三个层次的划分有助于译者在有限的时间内做出最有效的译后编辑决策, 实现效率与质量的平衡。

5. 结论

本研究以赖斯的文本类型学为框架, 以《经济学人》社论“From cyber-security to biosecurity”的英译汉实践为对象, 采用个案研究与对比研究法, 考察 DeepSeek 与豆包两个 LLM 机翻系统在新闻社论翻译中的典型问题及译后编辑策略, 得出以下发现。第一, 专业术语的翻译问题集中于概念重组的缺失, 主要损害文本的信息功能。LLM 系统倾向于字面对应而非概念重组, 当术语承载复杂背景或抽象含义时, 译文表面通顺但实质偏离原意。人工译后编辑的核心策略是依据语境进行概念重组。第二, 句法结构的翻译问题集中于英语逻辑编码方式的机械保留, 同时损害信息功能与操作功能。LLM 系统过度保留英语的名词化结构、隐性因果和被动语态, 导致译文逻辑关系隐晦、比较力度削弱、施动者缺失, 导致译文逻辑关系隐晦。人工译后编辑的核心策略是依据汉语表达习惯进行逻辑重构。第三, 修辞与搭配的处理问题集中于直译策略对表情功能与操作功能的消解。LLM 系统倾向于字面转换而非功能再现。基于上述发现, 译后编辑的介入可按优先级划分为三个层次: 第一优先级为必须修改的层面, 涉及术语误译、逻辑关系错置、语义偏差等直接损害信息功能的错误; 第二优先级为建议修改的层面, 涉及语序生硬、被动语态过度保留、名词化结构直译等影响可读性的欧化表达; 第三优先级为可选修改的层面, 涉及修辞力度的强化等服务表情功能的风格润色。

本研究是对现有 LLM 翻译错误类型研究的重要补充。崔启亮与牛硕(2023)关于 LLM “隐性错误”的核心判断在新闻社论文本中得到充分验证。在此基础上,本研究的贡献在于进一步揭示了这些“隐性错误”对新闻社论三重文本功能的不同损害机制,以及“形式错误减少”与“认知负荷转移”之间的内在关联,将 LLM 翻译错误类型研究从“错误是什么”推进至“错误损害了什么功能”的深度。

本研究存在局限性。语料集中于《经济学人》单篇社论,样本量有限,且所选文本属科技政策类议题,未能涵盖经济、政治等其他常见主题类型。此外,译后编辑策略的有效性分析以研究者判断为依据,尚未通过读者接受实验或译员行为实验加以验证。后续研究可扩大语料规模,纳入不同主题的社论以检验外部效度,亦可采用眼动追踪、键盘记录等认知实证方法,量化译后编辑者的认知负荷,为优先级原则提供实证支撑。

参考文献

- [1] International Organization for Standardization (2020) ISO 20771:2020 Legal Translation—Requirements. <https://www.iso.org/obp/ui/#iso:std:iso:20771:ed-1:v1:en>
- [2] 王湘玲, 杨艳霞. 国内 60 年机器翻译研究探索: 基于外语类核心期刊的分析[J]. 湖南大学学报(社会科学版), 2019, 33(4): 90-96.
- [3] 罗季美, 李梅. 机器翻译译文错误分析[J]. 中国翻译, 2012, 33(5): 84-89.
- [4] 崔启亮, 李闻. 译后编辑错误类型研究: 基于科技文本英汉机器翻译[J]. 中国科技翻译, 2015, 28(4): 19-22.
- [5] 包凯. 谷歌翻译汉译英错误类型及纠错方法初探[J]. 中国科技翻译, 2017, 30(4): 20-23.
- [6] 李国兵, 张敬. 机器英译物流论文摘要问题探究[J]. 中国科技翻译, 2022, 35(2): 24-27, 39.
- [7] 耿芳, 胡健. 人工智能辅助译后编辑新方向——基于 ChatGPT 的翻译实例研究[J]. 中国外语, 2023, 20(3): 41-47.
- [8] International Organization for Standardization (2017) ISO 18587:2017 Translation Services—Post-Editing of Machine Translation Output—Requirements. <https://www.iso.org/obp/ui/en/#iso:std:iso:18587:ed-1:v1:en>
- [9] 冯全功, 崔启亮. 译后编辑研究: 焦点透析与发展趋势[J]. 上海翻译, 2016(6): 67-74, 89, 94.
- [10] 朱慧芬, 赵锦文, 诸逸飞. 在线机器翻译的译后编辑原则研究: 以“八八战略”为例[J]. 中国科技翻译, 2020, 33(2): 24-27.
- [11] 陈胜, 田传茂. 在线翻译平台汉英翻译的问题及译后编辑: 以石油地质文献为例[J]. 中国科技翻译, 2021, 34(1): 31-34.
- [12] 冯全功, 张慧玉. 全球语言服务行业背景下译后编辑者培养研究[J]. 外语界, 2015(1): 67-73.
- [13] 冯全功, 刘明. 译后编辑能力三维模型构建[J]. 外语界, 2018(3): 63-70.
- [14] 王律, 王湘玲. ChatGPT 时代机器翻译译后编辑能力培养模式研究[J]. 外语电化教学, 2023(4): 16-23+115.
- [15] 杨艳霞, 魏向清. 基于认知范畴观的机器翻译译后编辑能力解构与培养研究[J]. 外语教学, 2023, 44(1): 90-96.
- [16] Reiss, K. (2000) Translation Criticism: The Potentials and Limitations: Categories and Criteria for Translation Quality Assessment. St Jerome Publishing, 16.
- [17] 张美芳. 文本类型理论及其对翻译研究的启示[J]. 中国翻译, 2009, 30(5): 53-60+95.
- [18] 连淑能. 英汉对比研究[M]. 北京: 高等教育出版社, 2010: 242-255.