

基于自建语料库的轨道车辆实操口译模块 特异性与问题诊断研究

——以2025轨道车辆技术国际训练营三大模块为例

王朝英

吉林外国语大学高级翻译学院, 吉林 长春

收稿日期: 2026年6月17日; 录用日期: 2026年8月17日; 发布日期: 2026年8月31日

摘要

笔者基于2025年轨道车辆技术国际训练营英语陪同口译实践, 聚焦“模块A: 受电弓检修与控制”“模块B: 客室车门安装与调试”“模块C: 车辆整车故障排查与处理”三大核心实操模块, 自建模块化的专门语料库——“2025轨道车辆技术国际训练营实操口译语料库(RVTC-PI Corpus, 2025)”。该语料库总库容10,800形符, 其中模块A 3600形符、模块B 3800形符、模块C 3400形符, 属于单语语料库, 内容严格限定于实操口译转写文本, 针对当前轨道车辆口译研究中“泛化场景多、模块细分少”的现状, 本研究采用“模块标注 + AntConc检索 + 人工统计”的分析方法, 系统诊断各模块口译问题。研究发现三大模块问题特征分化显著: 模块A因术语高度依赖实物语境, 术语错误率最高(12.9%), 典型错误集中于“碳滑板”等部件名称; 模块B因操作属性词的语境特殊性, 指令误译率最为突出(9.8%); 模块C受代码与原因双重认知负荷影响, 非流利频次远高于其他模块(9.2次/百词)。本研究验证了细分模块语料库在口译问题诊断中的有效性与精准性, 据此提出“实物 - 术语”对照卡片、属性词表、代码 - 原因关联表等针对性训练对策, 构建了“诊断 - 训练 - 验证”一体化的实操口译提升路径, 拓展了语料库翻译学在技术口译领域的应用维度。本研究深度依托吉尔口译精力分配模型、语料库口译量化分析范式展开实证分析, 通过标准化人工标注流程规避自动转写偏差, 实现不同实操模块口译认知负荷、错误类型的差异化量化归因, 弥补现有技术口译研究认知理论落地不足、操作流程不透明的短板。

关键词

轨道车辆实操口译, 模块特异性, 自建口译语料库, 口译问题诊断

A Study on the Specificity and Problem Diagnosis of Practical Interpretation Modules for Rail Vehicles Based on a Self-Built Corpus

—A Case Study of the Three Core Modules of the 2025 International Training Camp on Rail Vehicle Technology

Chaoying Wang

Graduate School of Translation and Interpretation of Jilin International Studies University, Changchun Jilin

Received: June 17, 2026; accepted: August 17, 2026; published: August 31, 2026

Abstract

Based on the English liaison interpreting practice at the 2025 International Rail Vehicle Technology Training Camp, this study centers on three core practical modules: “Module A: Maintenance and Control of Pantograph”, “Module B: Installation and Commissioning of Passenger Compartment Door”, and “Module C: Fault Finding and Repair of Vehicle”. A self-built, modular, specialized interpreting corpus, the “2025 Rail Vehicle Technology Camp Practical Interpreting Corpus (RVTC-PI Corpus, 2025)”, was constructed, with a total size of 10,800 tokens (Module A: 3600; Module B: 3800; Module C: 3400). As a monolingual interpreting corpus, the study comprises exclusively English target-language texts from practical interpreting scenarios. In response to the prevailing issue in rail vehicle interpreting research, which often features “generalized scenarios and lacks module-specific analysis”, this study adopted a “module-specific annotation + AntConc retrieval + manual analysis” approach to systematically diagnose interpretation issues across the modules. The findings reveal distinct problem profiles for each module: In Module A, due to the high contextual dependency of terminology on physical objects, the terminology error rate was the highest (12.9%), with typical errors concentrating on component names like “carbon contact”. In Module B, attributed to the context-specific nature of operational attribute words, the instruction misinterpretation rate was most prominent (9.8%). Module C, influenced by the dual cognitive load of associating fault codes with their causes, exhibited a significantly higher frequency of disfluencies (9.2 instances per hundred words) compared to the other modules. This research validates the effectiveness and precision of a segmented module-specific corpus in diagnosing and interpreting problems. Consequently, targeted training countermeasures are proposed, including “physical object-term” reference cards, attribute word lists, and code-cause association tables. These contribute to building an integrated “diagnosis-training-verification” pathway for enhancing practical interpreting skills and expanding the application scope of corpus-based translation studies in the domain of technical interpreting. Drawing on Gile’s Effort Model and corpus-driven quantitative interpreting research paradigm, this study conducts empirical analysis, avoids automatic transcription bias through a standardized manual annotation process, realizes differentiated quantitative attribution of interpreting cognitive load and error types across practical modules, and addresses the shortcomings of insufficient application of cognitive theory and opaque operational procedures in existing technical interpreting research.

Keywords

Rail Vehicle Practical Interpreting, Modular Specificity, Self-Built Interpreting Corpus,

Interpretation Problem Diagnosis

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

1.1. 研究背景

2025 轨道车辆技术国际训练营是交通运输行业面向海外技术人员开展的专项实操培训项目，旨在通过“理论 + 实操”相结合的方式，提升学员的轨道车辆检修能力。其中，“模块 A：受电弓检修与控制”“模块 B：客室车门安装与调试”“模块 C：车辆整车故障排查与处理”是训练营的三大核心实操模块。模块 A 重点包括受电弓碳滑板更换、升降压力调试等操作；模块 B 涵盖车门平行度校准与紧急解锁装置安装等内容；模块 C 则涉及故障代码读取与整车电路故障定位。

2025 年 5 月 22 日至 27 日，笔者作为该训练营英语口译员，全程参与共计 68 小时的实操模块陪同口译，发现轨道车辆实操口译与传统会议口译存在显著差异：第一，术语高度实操化，如模块 A 中的“pantograph carbon contact wear (受电弓碳滑板磨损)”和模块 B 中的“door leaf parallelism deviation (车门扇平行度偏差)”等表达，需结合具体部件与操作动作进行理解；第二，指令传递即时性强，例如技师指令“adjust the lifting pressure to 0.6 MPa (将升降压力调至 0.6 兆帕)”“check fault code 0x2F (检查故障代码 0x2F)”等需在 1~2 秒内准确传译，否则将影响实操进度；第三，模块间语言特征差异显著，具体表现为受电弓模块的“机械部件术语”、车门模块的“安装调试指令”以及整车故障模块的“代码 - 故障关联表达”。

实践后的初步复盘发现，传统的“记忆复盘”方法仅能模糊识别“受电弓术语不熟”“车门指令翻译慢”等问题，难以量化各模块间的具体差异(如模块 C 的非流利频次是否高于模块 A)，更无法定位具体错误类型(如受电弓术语错误集中于碳滑板还是升降机构)。鉴于此，本研究参考张威(2019)提出的“实操场景个人语料库应聚焦模块特异性”观点[1]，自建以三大实操模块为核心的个人语料库，借助简易工具实现模块问题的精准诊断，以弥补当前实操口译研究中“泛化场景多、细分模块少”的不足。

1.2. 研究目的与意义

本研究旨在构建以 2025 轨道车辆技术国际训练营为唯一语料来源的实操口译语料库，系统分析三大模块的口译语言特征与核心问题；通过量化比较模块 A (受电弓)、模块 B (车门)与模块 C (整车故障)在口译问题上的差异，识别各模块的特异性短板，并据此提出针对性优化建议，为同类实操口译训练建立一套可复制的“问题诊断 - 训练提升”路径。在方法论层面，本文完整披露语料清洗、ASR 转写校对、多层级人工标注全流程细则，配套标准化标注手册，解决当前自建小型口译语料库研究流程透明度不足、判定标准模糊的共性缺陷；在理论层面，本研究填补了轨道车辆实操模块细分口译研究的空白，验证了在单一训练营场景下构建细分模块语料库的诊断价值，拓展了语料库翻译学在实操技术口译领域的应用范围；在实践层面，研究结果为口译学习者提供了“模块针对性”训练依据，例如模块 A 应强化机械部件术语，模块 B 需重点优化操作指令传译，从而有效提升训练效率，同时也为国际训练营等实操培训项目的口译服务提供了实用参考。

2. 文献综述

2.1. 语料库翻译学与口译研究的发展

语料库翻译学为口译研究的科学化发展提供了关键的方法论支持。胡开宝(2012)指出,该学科以真实语料为基础,融合定量与定性分析方法,系统揭示翻译活动的内在规律,兼具方法论与研究范式的双重属性[2]。在学科发展中,语料库建设是核心前提。Sinclair (1991)曾强调:“任何语料库研究的起点是语料库本身建设。关于语料选材与组织方式的决策,将深刻影响后续研究的开展。”[3]这一观点在口译语料库研究中尤为重要,由于口译语料具有即时性与多模态特征,其构建规范直接决定后续分析的有效性。

从学科发展历程来看,国外语料库翻译学兴起于20世纪90年代,以Baker团队建设的“翻译英语语料库”(TEC)为代表,主要关注翻译共性与译者风格等问题[4]。国内研究起步稍晚,但在本土化语料库建设方面成果显著,以上海交通大学建设的“汉英会议口译语料库”(CECIC)为例,逐渐形成了会议口译、外交口译等特色研究方向。

语料库口译研究的发展大致可分为三个阶段:1998~2012年为理论引介阶段,陈菁与符荣波(2014)指出,该阶段主要探讨口译语料的转写与基础标注规范,尚未形成系统研究体系[5];2015年后进入快速发展阶段,张威与刘宇波(2021)通过CiteSpace分析发现,语料库类型从单一会议口译扩展至多模态与学习者语料库,研究方法更注重结合计量工具[6];2010年以来则呈现“微观化、多维度与实用化”趋势,开始关注副语言特征及教学应用。

在核心研究领域中,语料库建设与标注技术是基础支撑。刘剑与胡开宝(2015)提出,多模态口译语料库应整合语音、视频与文本信息,构建多层次标注体系,以弥补传统文本语料库在非语言信息上的缺失[7];张威则制定了学习者语料库的副语言标注标准,涵盖停顿、重复等参数,为口译教学评估提供了有效工具[1]。在口译产品分析方面,王斌华(2015)指出,语料库方法推动口译研究从“主观描述”转向“客观分析”,能够精准识别显化、简化等语言特征及译员风格差异[8];戴朝晖(2011)基于PACCEL-S语料库研究发现,学习者的非流利现象主要集中于有声停顿(如“er”)和单词重复,且低分组学生的停顿频率显著高于高分组[9]。此外,语料库也推动了口译评估模式的创新,王斌华(2012)提出“规范导向型”评估框架,将口译质量分解为术语准确性与流利度等可量化维度,通过对比标准语料库实现客观评估,从而突破传统主观评分的局限[10]。

2.2. 轨道车辆实操模块口译研究现状

在专业口译领域,轨道交通口译的术语处理是一大难点。李菲菲(2013)基于自建语料库研究发现,该领域术语普遍存在“一词多义”现象,如“bogie”可对应“转向架”或“台车”,需结合具体技术语境确定译法[11];韩雅雅与韩帅(2018)强调,语料库能够提供大量真实用例,有助于保障术语翻译的准确性与一致性,为行业口译实践提供重要支持[12]。然而,现有关于轨道车辆口译的研究多集中于“宏观术语译法”与“会议口译场景”,对实操模块的细分研究较为缺乏。在受电弓研究领域,韩雅雅与韩帅(2018)系统梳理了“pantograph(受电弓)”“catenary(接触网)”等基础术语的译法,但未涉及“carbon contact wear(碳滑板磨损)”等实操部件术语的动态使用特征[12];在车门调试方面,李菲菲(2013)提及“door leaf(车门扇)”的译法,但未深入探讨“adjust parallelism(调整平行度)”等操作指令的口译特点[11]。

在实操指令口译研究方面,Gile(2009)提出的“精力分配模型”指出,实操场景中“听辨(指令)-记忆(参数)-产出(翻译)”三重任务叠加容易导致错误,但该模型未进一步区分不同模块间的差异[13];戴朝晖(2011)基于机械实操口译语料发现,“动作+参数”类指令(如“tighten to 30N·m”)错误率最高,但其研究未覆盖轨道车辆三大实操模块的特异性问题[9]。

2.3. 口译错误与非流利现象研究

近年来,“模块特异性语料库”研究逐渐兴起。例如,王璐、丁幕菲、周鹤等(2025)构建了“医学手术模块口译语料库”,发现“开腹模块”与“缝合模块”的口译问题存在显著差异,但该研究聚焦于医学领域,且依赖 TreeTagger 等复杂词性标注工具,不符合实操口译对“简单易行”的需求[14]。在轨道车辆领域,张威(2019)提出“实操模块语料库应聚焦‘术语-指令-动作’关联”,但尚未开展相关实证研究,亦未涉及 2025 轨道车辆技术国际训练营中的三大核心模块[1]。

2.4. 研究空白与本研究定位

现有研究存在以下三方面空白:一是轨道车辆口译研究未充分关注三大实操模块的特异性,未能系统区分受电弓、车门与整车故障模块的语言特征差异;二是现有实操语料库研究多针对泛化场景,缺乏以单一训练营为背景的细分模块语料支持;三是分析方法普遍较为复杂,尚未形成“简单工具+人工统计”的实操路径,同时现有同类研究普遍缺失标准化标注判定细则、语料预处理完整流程,自动语音转写工具带来的数据偏差未被充分讨论,且极少结合认知口译理论对模块差异化错误进行分层归因。为应对上述挑战,本研究以 2025 轨道车辆技术国际训练营为唯一研究场景,围绕三大实操模块构建了一个模块化的专门语料库,其内容严格限定于轨道车辆技术实操领域,采用极简分析方法,配套完整标注手册、细化语料清洗规范、讨论 ASR 工具局限,并深度依托吉尔精力分配模型完成多模块认知动因对比分析[1],旨在通过结构化语料精准诊断各模块的口译问题以探索针对性的口译能力提升路径。

3. 研究设计

3.1. 语料库构建

本研究所建语料库属于专门语料库,内容严格限定于轨道车辆技术实操领域;同时,作为一个模块化的单语译语语料库,其语料按实操模块划分,且仅包含英语译语文本,旨在通过结构化语料精准诊断各模块的口译问题。

3.1.1. 语料库基本信息

本研究构建了“2025 轨道车辆技术国际训练营实操口译语料库(RVTC-PI Corpus, 2025)”,其中“RVTC”代表“Rail Vehicle Technology Camp”,“PI”指“Practical Interpreting (实操口译)”。语料来源于 2025 年 5 月 22 日至 27 日举办的轨道车辆技术国际训练营,内容为中方技师面向来自马来西亚、哈萨克斯坦、乌兹别克斯坦和赞比亚的技术选手与专家开展的三大模块实操指导口译,未掺杂其他场景语料。模块语料具体分布如下:模块 A(受电弓检修与控制)共 3600 形符,涵盖碳滑板更换、升降压力调试及受电弓姿态校准;模块 B(客室车门安装与调试)共 3800 形符,包括车门扇安装、平行度调整和紧急解锁装置调试;模块 C(车辆整车故障排查与处理)共 3400 形符,覆盖故障代码读取、电路故障定位与制动系统故障排除。模块语料分布见表 1。

在语料清洗阶段,研究借助豆包辅助工具初步识别并剔除中文及无关信息,随后进行人工校对,具体流程包括:首先删除“休息时间”“工具领取”等非实操对话内容,共计 240 形符;其次去除重复指令的相同翻译,如“检查碳滑板”重复译述 2 次仅保留 1 次,共清理 160 形符;再次清除飞书妙计转写后保留的中文指令(如“把扳手递过来”)约 1800 形符。为真实反映实操口译状态,完整保留了术语错误、指令误译及非流利现象(如“the (DISFL: 停顿) carbon contact”),未作人为修正。最终获得有效英文语料 10,800 形符,按“模块-实操环节”划分为 9 个 TXT 文本(如“ModuleA-CarbonContact-522.txt”),每行存储单句,便于 AntConc 进行检索,全程无需复杂格式处理。

Table 1. Modular distribution and core practical contents of the RVTC-PI corpus
表 1. RVTC-PI corpus 模块分布与核心实操内容

模块类型	实操环节	形符	核心术语/指令示例	海外学员反馈问题
模块 A (受电弓)	1. 碳滑板更换 2. 升降压力调试 3. 姿态校准	3600	1. carbon contact wear 2. lifting pressure 0.6 MPa 3. pantograph alignment	碳滑板术语理解困难
模块 B (车门)	1. 车门扇安装 2. 平行度调整 3. 紧急解锁调试	3800	1. door leaf installation 2. parallelism adjustment ± 0.5 mm 3. emergency unlock	平行度指令传递慢
模块 C (整车故障)	1. 故障代码读取 2. 电路定位 3. 制动故障排除	3400	1. fault code 0x2F 2. circuit fault location 3. brake system troubleshooting	故障代码与原因关联难
总计		10,800		

1) 语料清洗与预处理详细规范

预处理全程遵循“分层剔除、保真留存、人工复核”三大核心原则，完整操作细则如下：

无关场景文本剔除规则：仅保留技师实操指令对应口译译文，完全删除休息闲聊、工具申领、场地协调、后勤通知等无技术实操内容，原始转写文本中该类无效文本共计 240 形符，全部剔除；非技术类日常对话(如学员问候、茶水沟通)统一清除，不纳入统计；

重复文本精简规则：同一操作指令反复复述产生的重复译文，仅保留单次标准译句，其余重复内容删除；仅当重复伴随术语修正、句式调整时，两段译文全部留存用于错误对比。本次清洗去除无差异重复译文共 160 形符；

源语中文干扰项清理规则：飞书妙计 ASR 转写会同步记录技师中文原指令，所有中文句子、中文词汇全部删除，仅留存英文口译译文；原始转写混杂中文文本约 1800 字符，完成全量清除；

口译产出保真原则：术语误译、参数记错、停顿、填充词、重复等非流利、错误表达不做人工修正，完整保留原始口译状态，仅通过标注标签标记问题，还原真实实操口译认知产出特征；

文本拆分规范：清洗完成后按“模块 - 实操子环节”拆分 9 份独立 TXT 文档，单句分行存储，统一 UTF-8 编码，适配 AntConc 批量检索需求。

2) ASR 转写工具局限性说明

本研究采用飞书妙计完成口译音频自动转写，工具基础识别准确率 92%，存在三类固有偏差，可能对研究数据产生轻微干扰，处理方案如下：

专业术语识别偏差：轨道车辆低频专业词汇(carbon contact, pantograph alignment)易被识别为通用词汇 carbon plate、lift frame，该偏差通过人工对照实训实物、行业标准术语表逐句校对修正，消除术语识别错误；

短时停顿、填充词漏识别：机器转写常省略“um、er”等填充词与 2 秒内短停顿，笔者对照原始录音逐段人工补标非流利标签，弥补 ASR 对副语言特征捕捉不足的缺陷；

数字、故障代码识别错乱：0x2F、0.6 MPa 等参数易出现数字错写，人工核对实训记录单完成参数校正。

综上，ASR 仅作为初转写工具，全部关键技术信息、副语言特征均通过人工二次校验，最大限度削弱工具偏差对统计结果的影响。

3.1.2. 语料标注方案

结合三大模块的实操特点，设计包含 3 类核心标签的极简标注体系，所有标签均直接在 TXT 文件中

手动插入，具体标注规则如下(见表 2)：

Table 2. Annotation system of the RVTC-PI corpus

表 2. RVTC-PI Corpus 标注体系

标注类型	标签格式	模块适配示例	说明
术语错误	〈TERM-ERR: 模块 - 类型〉	模块 A: 〈TERM-ERR: A-误译〉 carbon plate (应为 carbon contact) 模块 B: 〈TERM-ERR: B-漏译〉 door lock (漏译“调试”)	类型包括误译/漏译; 标注模块便于后续对比
指令误译	〈INST-ERR: 模块〉	模块 B: 〈INST-ERR:B〉 adjust level (应为 adjust parallelism) 模块 C: 〈INST-ERR:C〉 check code 0x2F (漏译“故障”)	仅标注“动作 + 参数/对象”类指令错误
非流利现象	〈DISFL: 模块 - 类型〉	模块 C: 〈DISFL: C-长停顿〉 fault code 0x2F 模块 A: 〈DISFL: A-填充词〉 um, carbon contact	类型包括短停顿(<2 秒)/ 长停顿(≥2 秒)/填充词

标注原则：1) 仅标注实操相关问题，非技术类犹豫不标注；2) 每个标签必关联模块，确保模块间对比分析。

3.2. 研究工具与方法

本研究采用操作简便且贴合实操需求的两类工具：飞书妙计作为转写工具，完成英语实操指令的自动转写，初始准确率为 92%，后通过人工结合实物的方式校对了模块特异性术语(如将初录为“carbon plate”的表述修正为“carbon contact”)；语料分析则使用 AntConc 3.5.9，启用其三个核心功能，包括通过 Word List 统计各模块错误频次(见表 3)、借助 Concordance 查看含标签索引行以分析问题语境，以及利用 Cluster 提取模块指令的核心搭配以辅助误译分析。

Table 3. Comparison of High-frequency words in each module

表 3. 模块高频词对比表

排名	模块 A (受电弓)	词频	模块 B (车门)	词频	模块 C (整车故障)	词频
1	pantograph	84	door	112	fault	78
2	carbon	76	parallelism	65	code	71
3	contact	73	leaf	58	system	69
4	pressure	68	adjust	55	circuit	52
5	lift	55	install	48	brake	47

研究方法采用“模块对比分析法”，共分三个阶段实施。第一阶段为语料导入与预处理，包括批量导入按模块分组的 9 个 TXT 文件、在 AntConc 中设置为按空格分词(不进行词性标注)，并建立模块专属检索词表(如模块 A 包括 carbon contact、lifting pressure 等)。第二阶段开展模块问题检索与统计，分别对术语错误、指令误译和非流利现象进行检索与计算，包括区分误译与漏译类型，统计错误率、误译次数及每百词非流利频次。第三阶段进行模块问题对比与动因分析，通过横向比较三大模块在错误率与非流利频次等方面的差异识别模块特异性问题，结合吉尔精力分配模型[13]区分单一认知负荷、多重叠加认知负荷对译员产出的差异化影响，依托 Concordance 索引行结合实训录音复盘解析认知动因，如模块 C 的长停顿是否源于故障代码与原因之间的关联困难。

3.3. 核心变量定义

为确保三大模块口译问题的横向对比具有统一的量化标准，本研究对模块术语错误率、模块指令误译率和模块非流利频次三项核心变量进行明确的操作定义，并规定相应的统计方法，具体见下表。各变量的操作定义均基于语料库标注体系，通过 AntConc 检索与人工计数相结合的方式获取数据，以保证统计结果的可复现性与模块间可比性。

变量类型	操作定义	统计方法
模块术语错误率	模块内术语错误次数/模块内术语总出现次数 × 100%	AntConc 检索 “〈TERM-ERR: 模块〉” + 人工计数
模块指令误译率	模块内指令误译次数/模块内指令总条数 × 100%	AntConc 检索 “〈INST-ERR: 模块〉” + 人工计数
模块非流利频次	模块内非流利次数/模块词数 × 100%	AntConc 检索 “〈DISFL: 模块〉” + 人工计数

4. 研究结果

4.1. 模块 A (受电弓检修与控制)：术语错误主导

模块 A 语料共计 3600 形符，术语出现 420 次，错误 54 次，总错误率为 12.9%；指令共 180 条，误译 12 次，误译率为 6.7%；非流利现象出现 216 次，每百词 6.0 次(见表 4)。

Table 4. Statistics of core interpreting problems in Module A (maintenance and control of pantograph)
表 4. 模块 A (受电弓检修与控制)核心问题统计

问题类型	次数	占比	典型示例
术语错误	54	68.4%	1. 〈TERM-ERR: A-误译〉 carbon plate (应为 carbon contact) 2. 〈TERM-ERR: A-漏译〉 lifting mechanism (漏译“升降”)
指令误译	12	15.2%	〈INST-ERR: A〉 set pressure to 0.8 MPa (应为 0.6 MPa)
非流利现象	216	16.4%	〈DISFL: A-短停顿〉 carbon contact wear

关键发现：术语错误主要集中在“碳滑板(carbon contact)”“升降机构(lifting mechanism)”等实物部件术语上，原因在于这些术语需结合受电弓实物进行理解，仅依靠文字记忆易导致误译；指令误译多涉及压力参数(如 0.6 MPa)漏记，反映出参数即时记忆与动作指令叠加所带来的认知负荷。

从吉尔精力分配模型视角解读[13]，模块 A 译员主要承担“听辨 + 术语记忆 + 短时参数存储”三重并行任务，但任务均依附单一实物认知场景，属于单层叠加认知负荷，因此非流利现象控制在较低水平，认知冲突集中体现在实体术语语义匹配环节，最终表现为高术语错误率。

4.2. 模块 B (客室车门安装与调试)：指令误译突出

模块 B 语料共计 3800 形符，术语出现 380 次，错误 35 次，总错误率为 9.2%；指令共 220 条，误译 21 次，误译率为 9.8%；非流利现象出现 247 次，每百词 6.5 次(见表 5)。

关键发现：指令误译率为三大模块最高，错误多集中于“平行度(parallelism)”“对称(symmetrical)”等操作属性词。例如，“parallelism”被误译为“level”，尽管两者在日常语境中意义相近，但在车门调试中，“parallelism”特指“前后缝隙均匀”，需精准区分。非流利现象则多因“操作步骤叠加”(如“先调平行度再锁螺栓”)导致，连续传递多步指令易引发犹豫。

Table 5. Statistics of core interpreting problems in Module B (installation and commissioning of passenger compartment door)
表 5. 模块 B (客室车门安装与调试)核心问题统计

问题类型	次数	占比	典型示例	实操环节分布
术语错误	35	37.2%	〈TERM-ERR: B-误译〉 door panel (应为 door leaf)	车门扇安装(71.4%)
指令误译	21	22.3%	1. 〈INST-ERR:B〉 adjust level (应为 adjust parallelism) 2. 〈INST-ERR:B〉 tighten 2 bolt s(漏译“对称”)	平行度调整(66.7%)
非流利现象	247	40.5%	〈DISFL: B-填充词〉 er, parallelism adjustment	紧急解锁调试(52.2%)

依托精力分配模型分析[13]，模块 B 核心认知压力来自连续多阶操作指令的序列记忆，译员需要同步区分通用词汇与工程语境限定属性词，听辨、序列记忆、语义筛选三类精力持续挤占产出资源；属性词语义边界模糊进一步提升语义检索成本，使得指令语义匹配失效概率显著上升，最终形成全模块最高指令误译率。该结论补充戴朝晖关于机械指令口译的研究结论[9]，证明连续多步骤指令会放大属性词语境特异性带来的认知损耗。

4.3. 模块 C (车辆整车故障排查与处理)：非流利频次最高

模块 C 语料共计 3400 形符，术语出现 320 次，错误 38 次，总错误率为 11.9%；指令共 160 条，误译 15 次，误译率为 9.4%；非流利现象出现 313 次，每百词 9.2 次(见表 6)，为三大模块中最高。

Table 6. Statistics of core interpreting problems in Module C (fault finding and repair of vehicle)
表 6. 模块 C (车辆整车故障排查与处理)核心问题统计

问题类型	次数	占比	典型示例	实操环节分布
术语错误	38	29.7%	〈TERM-ERR: C-误译〉 circuit line (应为 circuit fault)	电路故障定位(65.8%)
指令误译	15	11.7%	〈INST-ERR: C〉 read code 0x2F (漏译“故障”)	故障代码读取(80.0%)
非流利现象	313	58.6%	1. 〈DISFL: C-长停顿〉 fault code 0x2F (关联“制动压力低”) 2. 〈DISFL: C-短停顿〉 brake system troubleshooting	制动故障排除(72.2%)

关键发现：非流利频次显著高于其他模块，其中 72.2%的长停顿集中于“故障代码 - 原因关联”场景。例如，传译“fault code 0x2F”时需同时关联“制动系统压力低”，双重任务(代码翻译 + 原因回忆)导致停顿时间延长(≥2 秒)。术语错误则多表现为“故障(fault)”漏译，如在故障排查中“circuit fault (电路故障)”常被简化为“circuit”，造成核心语义缺失。

结合 Gile 双重认知负荷假说展开深度阐释[13]：模块 C 存在并行双轨认知任务，译员一方面需要完成故障代码符号转译，另一方面同步调取代码对应的故障机理、系统故障表现，两类独立认知加工同时占用短期记忆与语义检索资源，直接突破译员精力分配阈值；听辨、符号解码、故障知识库调取、译文产出四项任务相互争夺认知资源，产出环节频繁中断，最终表现为高频长停顿、填充词等非流利特征。该实证结果为精力分配模型在工程故障口译细分场景提供全新量化证据。

4.4. 三大模块问题对比

横向对比三大模块核心指标(见表 7)，模块特异性问题显著：
术语错误率：模块 A (12.9%) > 模块 C (11.9%) > 模块 B (9.2%)，模块 A 因受电弓部件术语较多且依赖实物理解，错误率最高；

指令误译率：模块 B (9.8%) > 模块 C (9.4%) > 模块 A (6.7%)，模块 B 因车门操作指令中“属性词(如平行度)”的精度要求高，误译率突出；

非流利频次：模块 C (9.2 次/百词) > 模块 B (6.5 次/百词) > 模块 A (6.0 次/百词)，模块 C 因故障代码与原因关联的认知负荷最大，非流利现象最为频繁。

Table 7. Comparison of core interpreting problems across the three modules
表 7. 三大模块核心问题对比

模块类型	术语错误率	指令误译率	非流利频次(次/百词)	核心问题
模块 A (受电弓)	12.9%	6.7%	6.0	实物部件术语误译
模块 B (车门)	9.2%	9.8%	6.5	操作属性指令误译
模块 C (整车故障)	11.9%	9.4%	9.2	故障代码 - 原因关联非流利

5. 讨论

5.1. 模块特异性问题的动因解读

5.1.1. 模块 A (受电弓)：实物依赖型术语的认知挑战

模块 A 术语错误率最高，主要动因为“受电弓部件的强实物依赖性”。例如，“carbon contact (碳滑板)”“lifting mechanism (升降机构)”等术语需结合部件形态(如碳滑板的弧形结构)与功能(如升降机构的液压驱动)进行理解，而译前准备仅依赖文字记忆，未实地观察实物，导致“carbon contact”误译为“carbon plate”(plate 意为“平板”，无法体现碳滑板的弧形特征)。此外，受电弓检修多在车顶进行，环境噪音(如风噪)干扰“lifting pressure 0.6 MPa”等参数指令的听辨，叠加参数即时记忆负荷，进一步引发指令误译(如将 0.6 MPa 误记为 0.8 MPa)。

基于认知口译理论进一步延伸：实物绑定术语属于具身认知范畴，单纯文本背诵无法建立术语与实体形态、功能的神经关联[13]，译员在无实物参照的口译现场，术语语义激活速度大幅下降，极易选用外形近似通用词汇替代专业术语；风噪外部输入损耗进一步压缩听辨精力池，有限认知资源被迫向语音识别倾斜，术语匹配资源被挤占，双重作用推高术语错误率。

5.1.2. 模块 B (车门)：操作属性词的语境特异性

模块 B 指令误译率最高，归因于“平行度(parallelism)”“对称(symmetrical)”等操作属性词的“语境特异性”。在日常英语中，“parallelism”与“level”均可表示“水平”，但在车门调试语境中，“parallelism”特指“车门扇与门框前后缝隙差 $\leq 0.5\text{ mm}$ ”，而“level”仅表示“上下高度一致”，二者不可互换。海外学员多为技术人员，对属性词的精度要求高，一旦误译(如将“parallelism”译为“level”)，易导致车门安装后缝隙不均，需重新调试。此外，车门调试涉及“多步指令连续传递”(如“先调平行度→再锁螺栓→最后检查锁舌”)，连续指令的记忆负荷加剧了非流利现象(如频繁使用填充词“emm”)。从精力分配模型序列记忆维度分析[13]，串行多步骤指令会持续占用短期记忆存储空间，当叠加属性词语境语义辨析任务时，译员产出精力被持续稀释；通用词汇与工程专属属性词存在语义重叠，语义检索阶段易出现语义混淆，造成指令核心技术参数、限定条件丢失，形成高频指令误译，验证了机械实操口译中“动作 + 精度属性”指令的认知脆弱性。

5.1.3. 模块 C (整车故障)：代码 - 原因关联的双重认知

模块 C 非流利频次最高，根源在于“故障代码 - 故障原因”的双重认知任务。例如，中方技师下达“check fault code 0x2F”指令后，海外学员需理解“0x2F 对应制动系统压力低”，而译员需同步完成“代码转译(0x2F→故障代码 0x2F)”与“原因关联(0x2F→制动压力低)”。此类双重任务超出了 Gile 所提出的“听辨 - 记忆 - 产出”精力分配上限[13]，导致长停顿。此外，整车故障涉及“电路、制动、受电弓”等多系统，术语跨模块使用(如“brake system”源自模块 A，“circuit fault”为模块 C 专属)，术语系统的复杂性进一步加剧了非流利现象。

认知层面可划分为符号加工、知识库检索两条独立认知通路，两类通路并行加工会造成严重认知资源竞争[13]：符号解码需要识别十六进制故障代码，故障机理检索需要调取跨模块术语与故障逻辑，两项任务无法串行处理，只能挤占译文产出的认知资源，口译产出频繁中断，长停顿、填充词等非流利行为成为译员缓解认知过载的补偿策略，为本模型多重认知负荷下的口译产出特征提供量化实证支撑。

5.2. 简单语料库方法的实操价值验证

本研究基于单一训练营场景，仅采用“3 类模块标签 + AntConc 基础检索 + 人工统计”方法，完成三大模块问题诊断，证明：1) 场景聚焦性提升诊断精度：单一训练营语料有效避免了多场景干扰，能够精准捕捉模块间特异性(如模块 A 的实物术语与模块 C 的代码关联)，较泛化场景语料库更贴合实操需求；2) 模块标注简化分析流程：标签与模块直接关联(如〈TERM-ERR:A〉)，无需额外分类，即可实现模块对比，较无模块标注语料库节省约 50%的分析时间；3) 人工统计确保方法可行性：无需依赖 Excel 建模或 SPSS 分析，仅通过 AntConc 检索与纸质记录即可完成数据统计，适用于无统计背景的口译学习者。同时，本研究完整公开语料清洗细则、标准化标注手册、ASR 工具偏差处理方案，解决了现有小型自建口译语料库研究流程不透明、判定标准主观化的缺陷，整套低成本、可复刻的量化诊断流程，可为工程口译学习者自主开展口译复盘研究提供标准化范式。

5.3. 与现有研究的呼应与补充

本研究结果与 Gile (2009)“精力分配模型”及戴朝晖(2011)关于“动作 + 参数指令错误率高”的结论相呼应[9] [13]，同时进一步补充了“模块特异性”维度，证实即使在同一领域内，不同实操模块的口译问题动因也存在显著差异(实物依赖 vs.属性词语境 vs.代码关联)。此外，本研究构建的“模块 - 问题 - 动因”对应关系(见表 8)，弥补了现有研究“泛化结论多、模块细分少”的不足，为轨道车辆实操口译提供了更具体的理论支持。

现有认知口译理论多停留在宏观场景对比，极少细分同一行业下不同实操模块的认知负荷分层差异，本文通过量化数据区分单层叠加认知负荷、序列记忆认知负荷、双轨并行认知负荷三类认知损耗模式，拓展了吉尔精力分配模型在细分技术口译场景的解释边界[13]，实现理论与行业实操数据双向对话。

Table 8. Correspondence of “problems-causes-targeted solutions” of the three modules

表 8. 三大模块“问题 - 动因 - 解决方案”对应关系

模块类型	核心问题	深层动因	针对性解决方案
模块 A (受电弓)	实物部件术语误译	术语学习脱离实物， 依赖文字记忆	结合受电弓实物， 制作“部件 - 术语 - 功能”对照卡片
模块 B (车门)	操作属性指令误译	属性词语境特异性理解不足	整理“车门操作属性词表” (如 parallelism = 平行度)
模块 C (整车故障)	故障代码 - 原因关联 非流利	代码 - 原因双重认知负荷	记忆“故障代码 - 原因”对照表， 译前强化关联训练

6. 结语

6.1. 研究结论

本研究通过对 2025 轨道车辆技术国际训练营三大实操模块(模块 A: 受电弓检修与控制、模块 B: 客室车门安装与调试、模块 C: 车辆整车故障排查与处理)口译语料的系统分析,得出三方面核心结论: 首先, 模块特异性口译问题显著, 三大模块呈现鲜明的特征——模块 A 以实物部件术语错误为主, 错误率达 12.9%, 集中于“碳滑板(carbon contact)”“升降机构(lifting mechanism)”等依赖实物理解的术语; 模块 B 以操作属性指令误译最为突出, 误译率为 9.8%, 主要涉及“平行度(parallelism)”“对称(symmetric)”等语境特异性强的操作属性词; 模块 C 则以故障代码 - 原因关联引发的非流利现象频发, 每百词达 9.2 次, 源于代码翻译与原因回忆的双重认知负荷, 这一差异充分证实了细分模块研究的必要性。上述差异可依托吉尔口译精力分配模型完成分层认知归因[13], 区分单层叠加、序列记忆、双轨并行三类差异化认知损耗机制, 实现认知口译理论在轨道交通实操口译领域的落地应用。其次, 简单语料库方法具备高有效性, 研究以单一训练营场景为语料来源, 采用“模块关联标签 + AntConc 基础检索 + 人工统计”的分析路径, 配套完整标准化标注手册、细化语料清洗流程、讨论自动转写工具固有局限及校正方案, 全流程操作透明可复刻, 无需复杂工具即可精准诊断各模块特异性问题, 不仅操作门槛低, 且可复制性强, 为实操口译领域的语料库研究提供了“低成本 - 高实效”的实践范式。最后, 各模块问题的动因与针对性对策明确: 模块 A 的术语错误源于术语学习脱离实物场景, 模块 B 的指令误译源于对操作属性词语境特异性理解不足, 模块 C 的非流利现象源于代码 - 原因关联的双重认知压力; 而“实物 - 术语 - 功能”卡片、操作属性词表、故障代码 - 原因对照表等针对性解决方案, 可直接应用于后续口译能力提升训练, 为实操技能优化提供明确方向。

6.2. 研究局限与未来方向

本研究存在以下局限: 语料仅来源于笔者个人口译, 缺乏与其他训练营口译员的对比数据; 标注维度未涵盖“海外学员反馈对译员表现的影响”(如学员提问引发的临时非流利); 同时 ASR 自动转写工具仍存在无法完全捕捉微弱副语言停顿的局限, 人工补标虽能缓解偏差, 但无法完全消除工具本身带来的数据瑕疵。未来研究可: 1) 构建“多译员对比语料库”, 纳入其他训练营口译语料, 分析个体与群体在模块问题上的差异; 2) 新增“学员反馈标签”(如<FEEDBACK: 提问>), 深入探讨学员互动对实操口译的影响机制; 3) 引入多模态录音标注工具, 同步整合音频、文本标注, 弥补纯文本语料库对副语言信息捕捉不足的短板, 进一步完善认知负荷量化分析体系。

6.3. 实践启示

本研究为轨道车辆实操口译学习者提供了三点核心实践启示: 首先, 应开展模块靶向训练, 根据三大模块特点实施差异化训练策略——模块 A 强调整合实物与术语认知, 模块 B 聚焦操作属性词的精准传译, 模块 C 注重故障代码与原因之间的关联训练; 其次, 提倡以个人实操实践为唯一语料来源, 采用“模块标签 + AntConc”构建极简语料库, 参照本文附录标准化标注手册完成人工标注, 严格执行分层语料清洗规范, 规避自动转写工具偏差干扰, 无需追求大规模数据, 约 1 万形符规模即可有效支撑问题诊断; 最后, 应建立动态验证机制, 通过将后续实操语料持续补充至语料库, 系统对比各模块问题的改善情况(如模块 A 术语错误率是否降至 10% 以下), 从而形成“诊断 - 训练 - 验证”的闭环提升体系。

基金项目

该论文系吉林省教育厅人文社科研究项目“基于多模态语料库的《同声传译》教学模式研究(编号

JJKH20241527SK)”系列成果。

参考文献

- [1] 张威. 中国口译学习者语料库的语言信息标注: 策略及分析[J]. 外国语, 2019, 42(1): 83-93.
- [2] 胡开宝. 语料库翻译学: 内涵与意义[J]. 外国语, 2012, 35(5): 59-70.
- [3] Sinclair, J. (1991) *Corpus, Concordance, Collocation*. Oxford University Press.
- [4] Baker, P. (2021) *Corpus Linguistics and Translation Studies: Implications and Applications*. John Benjamins.
- [5] 陈菁, 符荣波. 国内外语料库口译研究进展(1998-2012)——一项基于相关文献的计量分析[J]. 中国翻译, 2014, 35(1): 36-42.
- [6] 张威, 刘宇波. 国内外口译研究最新进展对比分析——基于 CiteSpace 的文献计量学研究(2015-2019) [J]. 外国语, 2021, 44(2): 86-98.
- [7] 刘剑, 胡开宝. 多模态口译语料库的建设与应用研究[J]. 中国外语, 2015, 12(5): 77-85.
- [8] 王斌华. 语料库口译研究——口译产品研究方法的突破[J]. 中国外语, 2012, 9(3): 94-100.
- [9] 戴朝晖. 中国大学生汉英口译非流利现象研究[J]. 上海翻译, 2011(1): 38-43.
- [10] 王斌华. 从口译标准到口译规范: 口译评估模式建构的探索[J]. 上海翻译, 2012(3): 49-54.
- [11] 李菲菲. 新型轨道交通车辆科技英语术语释译[J]. 无线互联科技, 2013, 10(7): 175-176.
- [12] 韩雅雅, 韩帅. 如何翻译好轨道车辆的词汇[J]. 海外英语, 2018(1): 136-137+149.
- [13] Gile, D. (2009) *Basic Concepts and Models for Interpreter and Translator Training*. Revised Edition, John Benjamins, 90-108. <https://doi.org/10.1075/btl.8>
- [14] 王璐, 丁慕菲, 周鹤, 何倩倩, 宋江典. 医学大语言模型的研发与应用系统综述[J]. 智能系统学报, 2025, 20(6): 1295-1303.