

Doi: doi.org/10.70693/rwsk.v1i2.539

语料库驱动的网络隐性仇恨言论的语用分析

——以“专家建议”话题评论为例

徐俊扬¹

(暨南大学, 广东 广州 510632)

摘要: 本文基于 KH coder 语料库工具的词汇共现网络功能和言语行为理论, 对网络媒体中关于“专家建议”话题的评论进行分析。研究发现网络隐性仇恨言论存在以谓词为驱动的表现方式, 其语用功能的表现与谓词原来功能不一致, 从而隐性地达成了评论者的目的。

关键词: KH coder; 仇恨言论; 言语行为理论

一、研究背景

1.1 现实背景

2023年1月27日, 联合国秘书长古特雷斯在纽约联合国总部举行的国际大屠杀纪念日之际发表讲话, 回顾了这段可怕的历史, 并对当今世界现状表示担忧。他着重提到了“错误信息、偏执的阴谋论和不受控制的仇恨言论 (Hate Speech) (Brown, 2017)”的猖獗, 反犹太主义的兴起以及白人至上主义和新纳粹主义意识形态的扩散, 世界必须对“蛊惑人心的仇恨声音”保持警惕。同一天, 在白宫举行的农历新年招待会上, 针对最近在亚洲社区发生的两起大规模枪击事件, 美国总统拜登承认反亚裔仇恨犯罪的增加, 并指出亚裔社区期间所经历的“深刻的仇恨、痛苦、暴力和损失”, 此类犯罪部分是由与新冠病毒流行有关的煽动性言论引起的。

可见, 仇恨言论这一现象在国际社会深受重视且具有极大的潜在危害:一方面, 仇恨言论的发生促使了暴力的煽动, 多样性和社会凝聚力的破坏;另一方面, 仇恨言论的发酵威胁到国际社会急需凝聚的核心价值观和原则。

近期, 我国主流社交媒体——新浪微博关于“建议专家不要建议”、“年轻人为什么不爱听专家建议”、“年轻人反感专家问题究竟出在哪”等话题走热, 话题底下频频出现各种看似温和却令人寒心的评论, 而这些“温和”的评论的点赞数和回复量都相对较高, 这意味着我国网络仇恨言论现象也处于发酵状态。

1.2 网络仇恨言论研究回顾

近几年, 关于网络仇恨言论的研究基本集中在国外英文期刊的发表, 主要分为宗教、人权、性别和政治四大类, 不同类型的网络仇恨言论都带有其特有的意识形态特征, 包括伊斯兰恐惧症 (Islamophobia)、另类右派 (Alt-Right) 和白人民族主义 (White Nationalism) 等 (Castaño-Pulgarín et al., 2021)。在仇恨言论的表达上, Răileanu (2016)着手于宗教话题探讨仇恨言论发布者的言语特征以佐证当事人普遍与其他人存在对事件的不同看法, 并且这种特立现象的表现会在社交网络上被放大。Sundén 和 Paasonen (2018)在设身处地地观察和分析关于女性的网络仇恨言论后发现相关话语倾向于羞辱的表达方式。Trajkova 和 Neshkovska (2018)在研究出于政治目的的仇恨言论中发现事主偏好于扮演分析人员或是仲裁者来表达负面评论并对他人施以权力和支配。He et al. (2021)收集不同种类的言论并分析得出仇恨言论和反言论 (Counterspeech) 的相互关系, 印

作者简介:

徐俊扬 (1998—), 男, 广东东莞人, 暨南大学外国语学院硕士研究生, 研究方向为语料库语言学。

证了反言论对减少仇恨言论的积极效应。Parvaresh (2023) 通过大量语料的观察发现网络仇恨言论的形式和意义都表现得比较隐晦, 并根据语用功能来尝试对相关言论作具体分类以为后续研究其隐晦特征为参考。本研究基于前人的研究基础, 以言语行为理论为指导, 尝试对我国媒体相关语料做语料库数据整理和语用分析, 并就相关现象提出菲薄见解。

二、研究设计

2.1 研究问题

本文考察我国主流社交媒体中相关话题的评论并作语用分析, 主要研讨以下三个问题: 1) 网络隐性仇恨言论的相关表现形式及语义特征; 2) 相关话题中的网络仇恨言论者的具体攻击对象 3) 网络隐性仇恨言论行为的意义。

2.2 语料来源

哔哩哔哩弹幕网 (Bilibili), 简称 B 站, 是中国的一个弹幕视频分享网站, 其建立初衷为用户提供和谐而干净的弹幕视频分享平台。B 站用户年龄层主要集中在出生于 1990-2009 年区间的 Z 世代群体 (Generation Z), 青年用户倾向于依凭着匿名性、互动性极强的弹幕及评论功能达成情感认同 (糜舒涵, 2020)。为此, 笔者选择收集该网站中两则视频的一级评论作为语料制作语料库, 视频分为 1) 央视网的 2011 年春晚相声《专家指导》建议专家们多看、2) 央视网快看的央视网评为何年轻人越来越反感“专家”, 经语料库工具识别共计 2154 条中文语料。

2.3 研究工具和理论

KH coder(樋口, 2016, 2017)是一款用于定量内容分析和文本挖掘的免费软件, 可对中文 (简体中文)、法语、英语、日语等多种语言进行分析处理, 并预置词性赋码和可视化数据分析功能, 从形式和意义上帮助我们做进一步的文本分析。其中共现网络图 (Co-Occurrence Network) 能根据 Jaccard 系数对词与词之间 (根据词性划分不同词类的词) 关系及频率信息绘制图标, 并根据不同颜色划分相应的“关系社群” (subgraph)。

Searle 在 Austin 所提出的言语行为三分说 (顾曰国, 1989) 的基础上进一步提出直接言语行为 (Direct speech act) 和间接言语行为 (Indirect speech act) 的区别。其中, 直接言语行为即部分言语行为可直接实现交际目的, 而间接言语行为的交际目的需要通过另外一个言语行为表达时 (何兆熊, 1984)。

在多语篇的具体分析过程中, 一方面语料库语言性能提供自下而上的量化分析数据的视角, 从而补足在文本纵深分析中的人工观察语料中存在的短板; 另一方面, 语用分析为量化研究所遗留的问题提供理论兜底和细化探究, 进而高效地完成语篇分析目的。

三、研究过程

3.1 词汇共现网络分析

用 KH coder 语料库工具加载观察语料库, 对语料库进行词性标记并生成词汇共现网络, 如图 1 所示:

根据图 1, 围绕“专家”的关系群最多紧密, 主要包括不、建议 (Verb)、媒体、说、都, 可以看出相关评论主要涉及专家的一些行为措施, 并且“不”与“专家”的关系比较密切, 似乎表现出对“专家”的相关行为的持有否定态度较多, 具体观察下如下关系排名前三的索引行 (concordance):

(1) “不”与“专家”的关联语境:

1) 但很多媒体的专家没加名字, 也不知道是哪个专家, 或者就是几句话, 专家是流量的代名词。

2) 专不专家的我分得清, 相比起来我更关心 春晚节目能不能按这个水平 (对“专家建议”的讽刺性话题内容) 办。

3) 专家建议: 以后自媒体的标题不要使用“专家建议”。

4) (专家的相关事务) 我根本不管, 相声界各种什么争来争去的, 跟我没一毛钱关系, 我只想看到好节目。

5) 这是一个我认为很复杂的, 需要双向努力 的问题。有时候不是专家的问题
从中可以观察到, “不”在整体的语境中表现出 消极的语义韵 (Sinclair, 2004: 34), 对“专家”的身份、建议、责任等相关事项存在怀疑和不确定, 而这种所谓“怀疑和不确定”的模糊性概述似乎又藏有更深层次的语用表达。令人欣慰的是, 其中部分评论也用“不”来表达出视频内容的理性存疑,

或事件提出自身合理的见解或意见,使其向着更积极的方向发展,其语用功能是积极的,而这里第一个“建议”是让受事者去接受一个对其自身影响不好的事物——讽刺性话题视频,这明显与“建议”的出发点相违背。第二个“建议”出现在第三个“建议”之前,“专家”是第二个“建议”的受事者,按照“建议”原来的语用功能解释——“不要再建议”是评论者对受事者“向积极方向发展”的驱动,而第三个“建议”所驱动的成分本应也是积极的,而“不要再建议”明显否定了第二个“建议”所带来“积极方向”的驱动的解析,因此第二个和第三个“建议”所呈现的语义与其原本的语用功能也是相斥的。可见,整段话语所呈现的是话语的间接引语功能,评论者通过这种“违背”话语规则的表达方式间接地表达出对专家建议的讽刺和专家身份的贬毁,属于仇恨言论的隐性表达。

(2) 专不专家的我分得清,相比起来我更关心春晚节目能不能按这个水平(对“专家建议”的讽刺性话题内容)办。

此处从谓词“分”(区分)、“关心”和“办”(举办)来探讨,“分”所引导的小句呈现的语义是对“专家身份的真伪”的领悟,功能是推动话语向进一步话题发展或避谈相似话题,体现评论者对“专家身份的真伪”话题的冷漠态度。“关心”的施事对象是春晚节目,并且评论者对春晚节目的愿望是按“这个水平”举办,而材料中的“水平”是“专家建议”的讽刺性内容的表现方式,体现了对这类话题的娱乐性的赞誉,语用功能是表达对相关话题内容的肯定的态度。该则评论中两个小句通过呈现的相反语用功能表达了对专家身份的漠不关心,而对讽刺专家话题的内容报以热忱,间接表明了评论者对专家的冷嘲热讽,同样属于隐性仇恨言论的表达方式。

四、总结

本研究以使用了语料库工具中的词汇共现网络以分析网络仇恨言论的基本表现方式和具体目的,并结合言语行为理论对话语的表层和深层语义语用价值做进一步的剖析,发现其中存在着仇恨言论的隐性表达,回答了本文所提出的三个问题:

(1) 网络隐性仇恨言论的表现形式存在以谓词为驱动推动话语表现具体语用功能的趋向,语义特征多数表现为具体语境内的谓词语义与谓词原有语义发生偏差。

(2) 相关话题中的网络仇恨言论者的具体攻击对象一般为话题对象,但也存在与其关联的其他对象,比如词汇共现网络中的“媒体”也被予以攻击对象的身份。

(3) 网络隐性仇恨言论行为的意义是根据其话语的语用功能所驱动的,大多评论者都是对目标对象进行“攻击”行为,这种行为的表现很大程度上释放了评论者的情绪,从而引起媒体及社会层面的关注,推动更进一步的发酵反应。

本文所搜集的语料是来自官方媒体账号的视频评论,相较于其他评论渠道,该渠道对于评论中自带的“攻击性”的评论包容度较低,因此其自身普遍带有攻击性言论的自动或人工的筛选过滤机制,除去了大多与仇恨言论相关的评论内容,但本研究通过语料库数据整理和语用分析发现其中仍有一些“温和”的仇恨言论。笔者认为,这种隐性仇恨言论现象的发生也表明了,在“讲文明,树新风”的背景下仇恨言论在发酵的同时也发生了演化,评论者在释放情绪的过程中以隐性的表达方式有效地过渡到大多官方正式媒体平台。对于这种转变,笔者建议应当加大相关规定限制。虽然不否认互联网确实需要一定程度的包容性,但长时间的隐性仇恨言论的发酵和演化可能会产生新的不可控因素,对于这样的不可控因素,相关机构应该提前做好引导和控制以建立良性、可持续发展的网络文明生态。

参考文献:

- [1] Brown, A.(2017).What is hate speech? Part 1: The Myth of Hate.Law and Philosophy, 36(4), 419–468.<https://doi.org/10.1007/s10982-017-9297-1>
- [2] Castaño-Pulgarín, S.A., Suárez-Betancur, N., Vega, L.M.T., & López, H.M.H.(2021).Internet, social media and online hate speech.Systematic review.Aggression and Violent Behavior,58,101608.<https://doi.org/10.1016/j.avb.2021.101608>
- [3] Răileanu, R.C.(2016).Religion-Based User Generated Content in Online Newspapers Covering the Collective Nightclub Fire.Romanian Journal of Communication and Public Relations, 18(2), 55.<https://doi.org/10.21018/rjcp.2016.2.2090>

- [4] Sundén, J., & Paasonen, S.(2018).Shameless hags and tolerance whores: Feminist resistance and the affective circuits of online hate.Feminist Media Studies, 18(4), 643–656.<https://doi.org/10.1080/14680777.2018.1447427>
- [5] Trajkova, Z., & Neshkovska, S.(2018).Online hate propaganda during election period: The case of Macedonia.Lodz Papers in Pragmatics, 14(2), 309–334.<https://doi.org/10.1515/lpp-2018-0015>
- [6] He, B., Ziemis, C., Soni, S., Ramakrishnan, N., Yang, D., & Kumar, S.(2021).Racism is a virus: Anti-asian hate and counterspeech in social media during the COVID-19 crisis.Proceedings of the 2021 IEEE/ACM International Conference on Advances in Social Networks Analysis and Mining, 90–94.<https://doi.org/10.1145/3487351.3488324>
- [7] Parvaresh, V.(2023).Covertly communicated hate speech: A corpus-assisted pragmatic study.Journal of Pragmatics, 205, 63–77.<https://doi.org/10.1016/j.pragma.2022.12.009>
- [8] 糜舒涵. (2020). 新媒体环境下的青年亚文化狂欢——以哔哩哔哩弹幕网为例. 传媒论坛(10), 157-158.
- [9] 樋口耕一. (2016). A Two-Step Approach to Quantitative Content Analysis: KH Coder Tutorial using Anne of Green Gables (Part I). 立命館産業社会論集 = Ritsumeikan Social Sciences Review, 52(3), 77–91.
- [10] 樋口耕一. (2017). A Two-Step Approach to Quantitative Content Analysis: KH Coder Tutorial Using Anne of Green Gables (Part II). 立命館産業社会論集 = Ritsumeikan Social Sciences Review, 53(1), 137–147.
- [11] 顾曰国. (1989). 奥斯汀的言语行为理论:诠释与批判. 外语教学与研究, (1), pp. 30-39, 80. <https://kns.cnki.net/kcms/detail/detail.aspx?dbname=cjfd1989&filename=WJYY198901014&dbcode=cjfd>.
- [12] 何兆熊. (1984). 英语语言的间接性. 外国语(上海外国语学院学报), (3), pp. 11-15. <https://kns.cnki.net/kcms/detail/detail.aspx?dbname=CJFD1984&filename=WYXY198403002&dbcode=CJFD>.
- [13] Sinclair, J.(2004).Trust the Text: Language, Corpus and Discourse (R.Carter,编).Routledge, 34.<https://doi.org/10.4324/9780203594070>
- [14] 甄凤超. (2022). 词项框架下话语对象的意义建构及其路径分析. 外国语(上海外国语大学学报), 45(2), pp.37-46.<https://kns.cnki.net/kcms/detail/detail.aspx?dbname=DKFX2022&filename=WYXY202202004&dbcode=DKFX>.

A Corpus-driven Pragmatic Analysis of Implicit Hate Speech on the Internet: A Case Study on the topic of Experts' Advice

Xu Junyang

College of Foreign Studies, Jinan University, Guangzhou, Guangdong, China

Abstract: The paper analyzes comments on the topic of "experts' advice" in online media based on the lexical relational network function of the KH coder corpus tool and speech act theory. It is found that online implicit hate speech exist a predicate-driven expression, and its pragmatic function shows inconsistent with the original function of the predicate, thus implicitly achieving the reviewer's illocutionary purpose.

Keywords: KH coder; Hate Speech; Speech Act Theory