

趋近化视域下生成式人工智能教育话语的风险与价值建构：高校学术规范文本与AI写作工具广告的对比研究

班亚凯

新疆大学外国语学院，新疆 乌鲁木齐

收稿日期：2026年7月21日；录用日期：2026年8月28日；发布日期：2026年9月14日

摘要

高校学术规范文本与AI写作工具广告围绕生成式人工智能形成了不同的风险与价值叙事。本研究以趋近化理论为框架，选取六所高校的生成式人工智能使用指南和六个AI写作工具官网广告作为语料，采用定性文本细读方法，从空间、时间和价值三个维度比较两类话语对生成式人工智能风险与价值的建构。研究重点考察高校文本对学术失范、能力减损与责任边界的呈现，以及企业广告对写作困难、效率提升和未来发展的表述，并分析不同机构立场对学生身份与技术使用正当性的塑造。

关键词

趋近化理论，生成式人工智能，学术规范，批评性话语分析

Risk and Value Construction in Generative AI Education Discourse from the Perspective of Proximization: A Comparative Study of University Academic Norms and AI Writing Tool Advertisements

Yakai Ban

School of Foreign Languages, Xinjiang University, Urumqi Xinjiang

Received: July 21, 2026; accepted: August 28, 2026; published: September 14, 2026

文章引用：班亚凯. 趋近化视域下生成式人工智能教育话语的风险与价值建构：高校学术规范文本与 AI 写作工具广告的对比研究[J]. 现代语言学, 2026, 14(9): 287-293. DOI: 10.12677/ml.2026.149925

Abstract

University academic norms and advertisements for AI writing tools frame generative artificial intelligence through different accounts of risk and value. Drawing on Proximity Theory, this study examines generative AI guidelines issued by six universities and advertisements published on the official websites of six AI writing tools. Through qualitative close reading, it compares how the two genres construct the risks and values of generative AI across spatial, temporal, and axiological dimensions. The study focuses on representations of academic misconduct, diminished capability, and responsibility in university texts, as well as writing difficulties, efficiency, and future development in commercial advertisements. It also considers how institutional positions shape student identities and legitimize particular forms of technology use.

Keywords

Proximity Theory, Generative Artificial Intelligence, Academic Norms, Critical Discourse 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

生成式人工智能进入高校写作场景后，高校治理文本与商业推广文本对技术形成了不同定位。高校通过课程规定、学术诚信条款和使用指南界定可接受的辅助范围，强调作者责任、学习过程及违规后果。AI 写作工具企业则把产品嵌入选题、起草、修改和投稿等环节，以节省时间、改善表达和促进成功作为主要诉求。

这种分歧使同一技术获得两种相异的意义：在高校文本中，它可能侵入学术规范；在企业广告中，它被塑造成应对任务压力和能力竞争的解决方案。两类文本在说明“能否使用 AI”的同时，还借助时空关系和价值判断，引导学生识别正在逼近的危险，并确认相应行动的正当性。分析这种意义竞争，可以说明机构立场如何通过具体语言选择介入技术使用争议。

现有生成式 AI 教育研究主要围绕教育机会与使用风险展开。相关研究讨论了大语言模型在个性化支持、反馈和学习辅助方面的潜力，同时指出虚假信息、偏见和依赖等问题[1] [2]。据此，生成式 AI 的教育意义需要结合具体任务、使用方式和治理条件加以判断，单一的威胁或工具定位难以涵盖其复杂影响。

学术诚信研究较多关注代写边界、作者身份和作弊识别，并主张从课程设计、透明披露及制度规范等方面回应新型失范[3] [4]。高校 AI 政策教育框架还把技术、伦理和教学治理纳入统一的分析体系[5]。这类研究阐释了生成式 AI 带来的教育问题及政策需求，对于不同机构如何借助语言把风险或收益转化为学生当下可感知的事实，仍缺少充分讨论。治理文本与商业推广围绕行动正当性展开的竞争也有待比较。

趋近化理论可用于解释上述话语过程。Cap 将趋近化界定为一种认知语用策略：话语主体通过压缩空间、时间和价值距离，使外部实体、未来情境或异质价值进入中心指示空间，进而为防范或介入行动提供合法化依据[6]-[8]。国内研究系统梳理了趋近化的理论基础、分析范畴、研究方法及其在批评性话语分析中的适用范围[9]。批评认知语言学有关理论基础、认知机制和研究路径的讨论，则扩展了话语、认知与社会关系的解释框架[10] [11]。这些成果为结合认知知识解和社会功能分析教育 AI 话语提供了方法论

基础。高校学术规范与 AI 工具广告如何围绕同一技术争取行动正当性, 现有研究涉及较少。

本文选取六所高校的官方英语生成式 AI/学术诚信文本和六个 AI 写作工具官网广告, 采用目的性抽样与定性文本细读, 考察两组语料对指示中心内部实体和外部实体的配置方式, 以及空间、时间和价值趋近化如何建构学术失范风险、学习落后风险与效率价值。分析还关注这些策略如何支持限制性规范或产品采用, 并比较两类话语所塑造的学生身份及其责任归属。高校文本倾向于把学生定位为可问责的学术主体, 企业广告则召唤持续优化的效率主体。本文的分析对象是话语建构, 不涉及 AI 工具的实际功效; 小样本所得解释也不外推至所有高校或企业。研究旨在通过批评性对照说明机构立场、商业利益和学生责任如何在教育 AI 话语中获得自然化表达。

2. 理论框架与研究设计

2.1. 趋近化理论

趋近化理论考察话语如何重组认知距离, 使原本处于空间、时间或价值远端的对象进入受话者的当下经验。Cap 在 Chilton 话语空间研究的基础上, 将参与者区分为指示中心内部实体(inside-deictic-center entities, IDC)和指示中心外部实体(outside-deictic-center entities, ODC)。IDC 通常包括说话者、受话者及获得认可的制度和价值; ODC 则指被表征为异质、危险或构成阻碍的主体、行为与情势[6]-[8] [12]。

话语可以通过描写 ODC 向 IDC 移动、施加影响或形成侵入, 把尚未贴近的威胁识解为迫在眉睫, 进而为防范、限制和干预提供合法化依据。IDC 与 ODC 是机构依据交际目的配置的语篇位置, 并无现实世界中的固定阵营与之直接对应。同一 AI 工具可能被设定为越过学术边界的外部力量, 也可能以受控助手的身份进入中心。趋近化框架因此能够解释技术角色如何随机构立场发生转换。

趋近化包含空间、时间和价值三个相互关联的维度。空间趋近化借助指示语、实体命名、位移动词及影响过程, 呈现外部对象与中心成员之间的边界和靠近关系。时间趋近化把未来后果、期限或发展路径投射到现在, 使受话者形成即时行动的判断。价值趋近化以诚信/失范、进步/落后、成功/失败等对立配置 IDC 与 ODC, 并赋予价值冲突以现实后果[6]-[8]。

距离压缩可以围绕威胁展开, 也可以把效率、能力和未来成功塑造造成可接近的正面机会。空间、时间和价值策略共同参与行动合法化: 高校借此说明规范与问责的必要性, 企业据此建构采用服务的合理性。国内关于趋近化理论的系统引介, 以及批评认知语言学对识解机制与社会批判关系的梳理, 表明该框架能够分析公共话语中的距离调节、价值定位和行动合法化[9]-[11]。

2.2. 语料与方法

本研究采用目的性抽样, 选取截至 2026 年 8 月 14 日仍可公开访问的两组英语官网文本。高校组包括牛津大学、哈佛大学教育学院、斯坦福大学、多伦多大学密西沙加校区、密歇根大学和悉尼大学发布的生成式人工智能或学术诚信页面。广告组包括 Grammarly、Paperpal、Writefull、Jenni、Trinka 和 Wordtune 的首页或学生页面。纳入语料须直接面向学生、教师或学术写作者, 并明确涉及责任、风险、未来、效率、能力或成功。两组文本在语种、媒介和教育写作主题上具有可比性, 制度功能的差异则有助于观察规范主体与商业主体如何围绕同一技术争夺解释权。

分析以具有完整命题意义的句段为基本单位。IDC、ODC 及其角色关系构成识别起点; 名词短语、指示表达、情态形式、位移和影响过程等语言线索, 则用于判断空间、时间和价值趋近化的实现方式。在此基础上, 本研究追踪这些线索如何连接作弊、能力减损、学习落后或效率提升等风险与价值路径, 再通过组内归纳和跨组比较, 说明不同路径如何支持限制性规范或产品采用, 并塑造“可问责的学术主体”与“持续优化的效率主体”。

3. 分析与讨论

两组语料均通过调整人物、工具、任务和后果之间的距离，为特定行动赋予合理性。高校文本关注 AI 在何种条件下越过辅助边界，企业广告则将未经工具优化的写作状态建构为需要即时改变的问题。以下分析从空间、时间和价值三个维度展开，相关判断限于所选官网文本。

3.1. 高校话语中的学术失范风险

高校文本的空间趋近化主要用于划定 AI 进入学术活动的边界。Oxford 把“unauthorised use of AI”直接归入“cheating and plagiarism”(U01)，Stanford 将 AI 实质性完成作业类比为他人代做(U03)。技术是否构成 ODC，取决于文本赋予它的角色；当 AI 取代学生应完成的认知劳动时，它便被置于规范共同体之外。学生、原创作品和学习过程构成需要保护的 IDC，未经许可的生成内容则被表征为穿越作者身份边界的外部力量。禁止性条款因此具有维护学术主体真实性的防护功能，其意义超出一般程序规定。

“辅助”与“替代”的角色转换进一步强化了这条边界。Toronto 提出 AI 应当“support, but not supplant your learning”(U04)。“support”允许工具进入 IDC 并承担受控协助，“supplant”则将其推出边界，使之成为侵占学习位置的替代者。Harvard 所说的 AI 捷径会“rob you of the learning”(U02)，借助“剥夺”这一及物过程，把抽象的能力损失表述为学生当下可感知的受损事件。空间侵入与受益剥夺相互配合，这种配置保留了技术在一定条件下的可用性，同时把限制对象具体落在越界使用上。

时间趋近化把一次作业中的选择连接到尚未发生的学业与职业后果。Toronto 从当前课程延伸至后续课程和未来工作，Oxford 把未经授权的使用与考核失败乃至纪律处分置于同一前摄链条，Harvard 则以入学求学目的反映当下的“认知代劳”。这些表述缩短了当前使用与后续受损之间的时间距离，使能力不足、资格受疑或纪律处罚进入眼前判断。语料不能证明风险必然发生。文本所完成的是预期性识解，即要求学生在行动之前依据未来后果重新评价眼前的便利，从而为预防性治理提供理由。

价值趋近化进一步规定了学术共同体可以接纳何种 AI 使用。Michigan 以“enhance, not hinder, learning”(U05)构成正反谓词对照，将促进学习、事实可靠、引用准确和透明披露归入 IDC 价值，将阻碍学习、虚假内容与责任逃避推向 ODC。Sydney 以“shortcut”对照“genuine academic success”(U06)，把偏离正当学习路径的近路与经由努力获得的成功分置两端。Oxford、Toronto 等文本也反复把诚信、原创和诚实同学生责任连接起来。这类配置所塑造的“可问责学术主体”需要核查输出、说明使用方式并为最终作品负责。该主体的自主性受到课程许可的约束；价值语言承认学生使用工具的能动性，同时将违规风险和验证义务主要落实到个人。

3.2. 企业广告中的学习落后风险与效率价值

企业广告将困难、停滞和未经优化的写作状态设为主要 ODC，越界 AI 的位置随之淡化。Writefull 先断言“Academic writing is hard”，随即以“made easy”给出反向路径(A03)；Wordtune 把用户面对的障碍描述为“scattered thoughts”，并承诺将其转化为准确表达(A06)。写作困难从一般经验被拉近为此刻阻挡“你”的近身问题，产品随后以排除障碍的助手身份进入 IDC。这种空间安排把采用工具建构为恢复表达秩序的合理反应。文本未涉及写作困难还可能源于学科训练、语言资源或指导条件等产品外部因素。

广告的时间趋近化集中压缩任务起点与成功终点之间的过程。Grammarly 用“For every assignment and every ambition”(A01)把眼前作业与远期抱负并列，Paperpal 以“submission-ready paper”(A02)将资料、草稿、润色和投稿收束为连续终点，Trinka 的“From First Draft to Final Submission”(A05)把整个写作周期纳入服务范围。在这类表达中，原本需要反复学习、反馈和修改的过程被改写为可由工具加速的路径，截止压力、投稿门槛和未来竞争也随之靠近当前选择。产品价值由即时节时延伸至成绩、发表和

职业发展,采用工具因此被叙述为接近未来目标的前置条件。

这一时间路径依靠效率与可信的价值组合获得正当性。Jenni 以 “Your Intelligent Research Assistant” (A04) 中的所有格和助手身份缩短人机距离,把自动引文、来源追溯和不中断的编辑流程并置。Paperpal 将信心、学术标准和经过核验的语言支持相连,Trinka 把 “更好” “轻松” “快速” 纳入同一改进方向。广告诉求由速度扩展到准确、可控和成功等 IDC 价值。这里的 “可信” 与 “投稿就绪” 属于企业的营销性承诺,语料只能说明这些表述如何合法化采用或购买,不能据此认定工具已经产生相应的学习或发表效果。

广告中的第二人称还把潜在用户塑造成需要持续优化的 “效率主体”。Grammarly 覆盖 “每项作业与每个抱负”, Wordtune 则预先消除未来 “卡住” 的状态。理想使用者被设定为能够随时识别低效、借助工具修正状态并持续接近目标的人。产品经由人格化成为可随行的助手,企业服务也被纳入用户的能力配置,外部介入的属性随之淡化。相较于高校的问责逻辑,这种主体化较少追问任务是否适合由 AI 完成,关注点落在个体对工具的利用程度上。拒绝采用产品因而可能被隐含地等同于浪费时间、停留在困难中或落后于竞争节奏。

3.3. 机构立场、学生身份与合法化功能

两类话语均借助距离缩短形成紧迫感,并将受话者导向特定行动,但各自配置了不同的威胁来源。高校把未经授权的替代性使用推向 ODC,使处罚、能力减损和身份失真逼近当下,以此支持披露、核查和限制。企业把写作困难、流程迟滞和未来落后置于 ODC,同时将 AI 产品吸纳进 IDC,为持续采用提供理由。风险与价值在两组语料中同时存在:高校用学习促进等正面价值为受控使用留出空间,广告也借助投稿门槛与卡顿焦虑建构风险。相同的风险与价值耦合承担着不同制度功能,分别服务于规范治理和商业转化。

趋近化策略还使学生承担两套可以叠加的主体责任。高校要求学生成为可问责的作者,说明工具用途并承担事实、引文和原创责任。企业则要求用户成为持续优化的写作者,通过购买或调用助手主动缩短与成功之间的距离。两种身份均强调个人选择,课程设计、教师支持、语言资源不均和产品商业利益因而可能退居话语背景。上述批评性判断以承认个人责任和工具可能具有实用价值为前提。当风险与成功都被压缩为一次即时选择的后果时,影响 AI 使用的制度条件较难进入可见的因果链。从所选语料看,空间、时间和价值趋近化构成相互衔接的话语过程。AI 先被安排为越界替代者或贴身助手,未来处罚、能力损失或成功目标继而向当下移动,诚信或效率等价值再提供评价标准。这一过程解释了高校如何把限制表述为保护学习,也说明企业如何把产品采用表述为个人进步。

4. 结语

两类教育 AI 话语通过重置 IDC 与 ODC 的关系界定技术角色。高校的学术规范话语将未经许可且替代学生认知劳动的 AI 工具的使用行为界定为 ODC,学生、原创作品、学习过程和诚信规范构成 IDC;AI 整体仍保留在学术共同体的可控使用范围内。企业广告把困难、停滞和未经优化的写作推向 ODC,并以 “助手” 等称谓将产品纳入用户 IDC。同一技术因机构立场而获得受控工具、越界替代者或贴身伙伴等不同角色。

空间、时间和价值策略共同完成上述角色配置。高校借助边界、侵入与剥夺表述拉近 AI 代劳的影响,以纪律处分、后续课程和职业能力组成前摄后果链,并用诚信、原创、透明与责任划定可接受使用。广告把 “难” 与 “卡住” 呈现为当前障碍,通过 “每项作业与每个抱负” “从初稿到投稿” 等表述压缩改进路径,将轻松、快速、可信和成功组合为采用产品的理由。

参与者定位、未来投射和价值评价连续运作，分别为规范治理和服务采用提供合法化依据，并塑造出“可问责的学术主体”与“持续优化的效率主体”。两种主体位置都把技术选择集中到个体决定上，课程设计、教师支持、资源差异和商业利益等制度条件的可见性可能因此降低。

这一比较把趋近化分析由威胁表征扩展到风险与正面价值的协同建构，也表明规范文本和商业广告可以在目标不同的情况下共享距离压缩的认知结构。理论的适用边界仍需更多语料检验。本研究采用两组各六篇的目的性小样本，只考察英语官网中的机构自我呈现，研究范围未涵盖网页视觉、文本迭代、学生实际使用和受众接受，也不检验广告承诺的产品效果。

后续研究可以结合历时网页档案与多模态页面材料，考察政策更新、界面设计和产品功能如何改变趋近化路径。学习者访谈、写作过程记录及评论互动则有助于说明不同学生群体如何接受、协商或拒绝“可问责作者”与“效率主体”等位置，并检验官网话语与实际教育实践之间的关系。

参考文献

- [1] Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., *et al.* (2023) ChatGPT for Good? On Opportunities and Challenges of Large Language Models for Education. *Learning and Individual Differences*, **103**, Article 102274. <https://doi.org/10.1016/j.lindif.2023.102274>
- [2] Tlili, A., Shehata, B., Adarkwah, M.A., Bozkurt, A., Hickey, D.T., Huang, R., *et al.* (2023) What If the Devil Is My Guardian Angel: ChatGPT as a Case Study of Using Chatbots in Education. *Smart Learning Environments*, **10**, Article No. 15. <https://doi.org/10.1186/s40561-023-00237-x>
- [3] Cotton, D.R.E., Cotton, P.A. and Shipway, J.R. (2024) Chatting and Cheating: Ensuring Academic Integrity in the Era of ChatGPT. *Innovations in Education and Teaching International*, **61**, 228-239. <https://doi.org/10.1080/14703297.2023.2190148>
- [4] Perkins, M. (2023) Academic Integrity Considerations of AI Large Language Models in the Post-Pandemic Era: ChatGPT and beyond. *Journal of University Teaching and Learning Practice*, **20**, Article 7. <https://doi.org/10.53761/1.20.02.07>
- [5] Chan, C.K.Y. (2023) A Comprehensive AI Policy Education Framework for University Teaching and Learning. *International Journal of Educational Technology in Higher Education*, **20**, Article No. 38. <https://doi.org/10.1186/s41239-023-00408-3>
- [6] Cap, P. (2008) Towards the Proximization Model of the Analysis of Legitimization in Political Discourse. *Journal of Pragmatics*, **40**, 17-41. <https://doi.org/10.1016/j.pragma.2007.10.002>
- [7] Cap, P. (2010) Axiological Aspects of Proximization. *Journal of Pragmatics*, **42**, 392-407. <https://doi.org/10.1016/j.pragma.2009.06.008>
- [8] Cap, P. (2013) Proximization: The Pragmatics of Symbolic Distance Crossing. John Benjamins Publishing Company. <https://doi.org/10.1075/pbns.232>
- [9] 武建国, 林金容, 栗艺. 批评性话语分析的新方法——趋近化理论[J]. 外国语, 2016, 39(5): 75-82.
- [10] 张辉, 杨艳琴. 批评认知语言学: 理论基础与研究现状[J]. 外语教学, 2019, 40(3): 1-11.
- [11] 张辉, 张艳敏. 批评认知语言学: 理论源流、认知基础与研究方法[J]. 现代外语, 2020, 43(5): 628-640.
- [12] Chilton, P. (2004) *Analysing Political Discourse: Theory and Practice*. Routledge.

语料来源

[U01] UNIVERSITY OF OXFORD. Guidance on safe and responsible use of GenAI [EB/OL]. [2026-08-14].

<https://www.ox.ac.uk/students/life/it/genai-tools/guidance-on-safe-and-responsible-use-of-genai>

[U02] HARVARD GRADUATE SCHOOL OF EDUCATION. HGSE AI policy [EB/OL]. [2026-08-14].

<https://registrar.gse.harvard.edu/learning/policies-forms/ai-policy>

[U03] STANFORD UNIVERSITY. Academic honesty and Stanford's Honor Code [EB/OL]. [2026-08-14].

<https://teachingcommons.stanford.edu/teaching-guides/foundations-course-design/feedback-and-assessment/academic-honesty-and-stanford>

[U04] UNIVERSITY OF TORONTO MISSISSAUGA. Generative AI and academic integrity [EB/OL].

[2026-08-14]. <https://www.utm.utoronto.ca/academic-integrity/students/generative-ai>

[U05] UNIVERSITY OF MICHIGAN. Course policies & syllabi statements [EB/OL]. [2026-08-14].

<https://genai.umich.edu/resources/faculty/course-policies>

[U06] UNIVERSITY OF SYDNEY. Academic integrity [EB/OL]. [2026-08-14].

<https://www.sydney.edu.au/students/academic-integrity.html>

[A01] GRAMMARLY. Grammarly for students [EB/OL]. [2026-08-14].

<https://www.grammarly.com/students>

[A02] PAPERPAL. AI academic writing tool [EB/OL]. [2026-08-14]. <https://paperpal.com/>

[A03] WRITEFULL. Academic writing is hard [EB/OL]. [2026-08-14]. <https://www.writefull.com/>

[A04] JENNI. AI academic writer & research tool [EB/OL]. [2026-08-14]. <https://jenni.ai/>

[A05] TRINKA. AI writing and grammar checker tool [EB/OL]. [2026-08-14]. <https://www.trinka.ai/>

[A06] WORDTUNE. Express yourself with confidence [EB/OL]. [2026-08-14]. <https://www.wordtune.com/>