

生成式人工智能重构电商生态的伦理挑战与多维度治理

吕冷然

江苏大学马克思主义学院, 江苏 镇江

收稿日期: 2025年10月11日; 录用日期: 2025年10月28日; 发布日期: 2025年11月26日

摘要

生成式人工智能发展迅速, 正深层改变电子商务的整体运作模式。在显著提升电商运营效率与个性化服务精准度的同时, 其“自主内容创造”特性引发一系列伦理问题: 内容真实性缺失、用户个性化选择被操纵、责任归属模糊及环境代价凸显等。本研究以“技术-社会双向形塑”理论为核心分析框架, 融合价值敏感设计与负责任创新理论, 结合案例分析与消费者行为的理论分析, 系统解析伦理问题的表现形式与作用机制。研究发现, 若人工智能生成的内容没有标注清楚来源, 消费者会觉得自己受到欺骗的比例可能显著上升; 而算法对用户选择的操控, 会削弱用户自主做决策的能力。基于此, 本研究构建了由“技术防控、组织自律、社会监管”三部分组成的协同治理策略。同时提出可通过提升算法的可理解性、对内容真实性进行验证、加强行业自身的规范管理, 以及完善相关法律制度等方式, 让生成式人工智能在电子商务领域的应用既符合责任要求, 又能实现长期健康的更新迭代。

关键词

生成式人工智能, 电商生态, 伦理挑战, 多维度治理, 消费者感知

Ethical Challenges and Multi-Dimensional Governance of Generative Artificial Intelligence in Reshaping the E-Commerce Ecosystem

Lengran Lyu

School of Marxism, Jiangsu University, Zhenjiang Jiangsu

Received: October 11, 2025; accepted: October 28, 2025; published: November 26, 2025

文章引用: 吕冷然, 生成式人工智能重构电商生态的伦理挑战与多维度治理[J]. 电子商务评论, 2025, 14(11): 2506-2512. DOI: 10.12677/ec.2025.14113715

Abstract

The rapid advancement of generative artificial intelligence is profoundly reshaping the overall operational model of e-commerce. While significantly enhancing operational efficiency and the precision of personalized services, its capacity for “autonomous content creation” has triggered a series of ethical issues: lack of content authenticity, manipulation of user choices, ambiguity in accountability, and growing environmental costs. Anchored in the theoretical framework of “bidirectional shaping of technology and society”, this study integrates Value-Sensitive Design and Responsible Innovation theories, combining case analysis with theoretical examination of consumer behavior to systematically investigate the manifestations and mechanisms of these ethical challenges. The findings indicate that when AI-generated content is not clearly labeled, the proportion of consumers feeling deceived may rise significantly; meanwhile, algorithmic manipulation of user choices can undermine their autonomy in decision-making. Based on these insights, the study constructs a collaborative governance strategy comprising three pillars: “technical safeguards, organizational self-regulation, and societal oversight”. It further proposes measures such as improving algorithmic interpretability, verifying content authenticity, strengthening industry self-discipline, and refining legal frameworks to ensure that the application of generative AI in e-commerce aligns with responsible norms and enables sustainable, healthy evolution.

Keywords

Generative AI (GenAI), E-Commerce Ecosystem, Ethical Challenges, Multi-Dimensional Governance, Consumer Perception

Copyright © 2025 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

生成式人工智能不断进步，正促使电子商务生态从仅在局部提供辅助，转变为实现全链条的革新。这种技术能够生成极为逼真的文本、图像和视频，大幅提升了电商在运营方面的效率，也显著提高了个性化服务的水准，还催生了虚拟模特、拟人客服以及元宇宙商场等全新的交互形式[1]。

然而技术创新背后潜藏多重伦理风险：内容真实性存疑、用户决策受隐性操纵、责任链条断裂及环境成本激增等。这些问题不仅侵犯消费者知情权与自主权[2]，更侵蚀电商行业信任根基，甚至与全球碳中和目标形成冲突[3]。

在此背景下，本研究采用理论分析与案例结合的方法，以技术－社会双向形塑理论为核心，结合价值敏感设计理论，系统解构四大伦理挑战的生成逻辑，构建多维度治理框架，对推动电商人工智能的可持续发展具有重要理论与现实意义。

2. 生成式人工智能重构电商生态的技术图景

生成式人工智能的发展推动电子商务实现了从传统算法到创造性生成的范式转换。传统算法依赖规则驱动与统计模型，虽能提升推荐与定价等环节的效率，但其生成能力有限。而生成式人工智能基于生成对抗网络(GANs)与大型语言模型(LLMs)等技术，能够创造高真实度的虚拟模特、广告文案及拟人对话，重塑了营销、客服与用户体验三大核心环节，重构了电商生态中内容生产与消费的基本逻辑。

在技术特性上,生成式人工智能与传统推荐算法存在本质差异:后者基于协同过滤与内容匹配,结果可解释性强;而生成式人工智能具备主动创造能力,依赖复杂神经网络与多模态数据,呈现“黑箱”特性,在透明度与可控性方面略显不足[4],于是加剧了数据隐私与安全风险。这一转变标志着电商从“被动推荐”迈向“主动创造”,并呼唤建立与之适配的伦理治理框架,以平衡创新与责任、个性化与操纵风险。

3. 科技伦理学理论框架

生成式人工智能在电商生态中的伦理挑战,需依托科技伦理学框架审视。本研究以科技伦理学为元框架,并重点引入“技术-社会双向形塑”理论,作为分析问题的核心视角。科技伦理学为识别和归类伦理问题提供了基本范畴,而“技术-社会双向形塑”理论则指导我们探究这些问题何以在电商这一特定社会情境中产生、放大并系统化。

从科学、技术、计算机技术特殊伦理三维度看,其既涉及数据知情同意、隐私保护等科学伦理问题,拟人化应用还引发真实性偏差与用户自主性侵蚀,而创造性输出的不可预测性、责任归属模糊性,更凸显传统伦理法律框架应对不足。其中,人机关系重构是核心议题:虚拟客服、个性化推荐模糊人机边界,加剧信息不对称与消费决策操纵,责任难界定也挑战现有问责机制。

从技术-社会双向形塑视角看,本研究的整体分析框架正是基于此视角构建。生成式人工智能与电商模式、用户认知、社会规范互动共生。具体而言,电商平台固有的政治经济学逻辑——即通过数据垄断、注意力变现和最大化用户粘性来实现利润增长的商业模式——系统性地催生并加剧了前述伦理问题。平台对流量和转化的追求,内在地偏好能够高效操纵注意力的生成式内容,而非可能损害短期商业利益的透明与公正。这种资本逻辑不仅塑造了技术的应用方向,也构成了伦理治理的根本性阻力。技术重塑电商权力结构与消费信任,如虚拟模特替代真人、算法影响购买决策。同时,社会对技术的反向塑造同样显著:资本逻辑主导技术应用方向——平台对流量转化的追求促使人工智能优先生成“注意力导向型内容”[5];监管政策推动技术规范——欧盟《人工智能法案》倒逼电商平台完善人工智能内容标识机制[4];用户需求驱动技术迭代——消费者对“算法知情权”的诉求促使可解释人工智能落地[6]。这种双向作用表明,技术治理需突破单一路径,构建动态协同的多维框架,将伦理转化为系统内生属性,实现技术创新与社会价值的良性融合。

4. 生成式人工智能的伦理挑战

4.1. 真实性失真挑战

生成式人工智能的高保真内容创造能力引发显著的真实性失真风险,核心表现为内容虚假性与消费者认知偏差的双重叠加。某大型电商平台引入生成式人工智能创建虚拟模特用于服装展示,这些模特形象逼真,能适配不同体型、肤色,显著降低了摄影与模特成本。然而,平台初期未对人工智能生成内容进行标识,导致部分消费者误以为是真人模特,对服装上身效果产生不切实际的期望,收到实物后落差巨大,引发大量投诉与退货。据《2024 中国电商 AI 应用伦理白皮书》披露,此类未标识 AI 内容的消费者误判率达 68%,退货率较真人展示组上升 210% [7]。

由此可见,生成式人工智能在电子商务生态里的应用,靠着自身很强的内容创造能力,带来了较为显著的“真实性失真”的伦理问题。这个问题主要表现在两个方面:一是人工智能生成内容存在虚假性,二是这种虚假内容会让消费者产生认知偏差。

从虚假性来看,生成式技术能高效制作出非常逼真的图像、视频和文本,像虚拟模特展示商品、个性化广告这些都是典型例子,这会让消费者怀疑内容的真假。理论分析与案例观察表明,如果不把人工

智能生成内容的来源明确标注出来，消费者会明显更觉得自己受到了欺骗。这种情况不仅会损害消费者的合法权益，还会动摇电商平台的信用根基。

再看认知偏差方面，生成式内容因为仿真度高，容易让消费者把虚假信息当成真实信息，进而干扰他们的购买决策。心理学领域的研究表明，人类本身就容易相信看到的视觉信息和听到的语言信息，而人工智能技术还在不断放大这种认知上的弱点。比如电商平台的虚拟试穿功能，可能会让消费者的商品实际使用效果产生不切实际的期待，这种情况长期积累下去，就会引发消费者对平台的信任危机。

4.2. 个性化操纵挑战

算法黑箱与精准数据采集的结合催生个性化操纵问题，其本质是平台通过人工智能技术实现对用户决策的隐性控制。某主流电商平台的生成式推荐系统通过分析用户浏览轨迹、停留时长及情绪关键词，动态生成“紧迫感文案”(如“您关注的商品仅剩3件”)，使非计划购买率显著提升。阿里研究院2023年调研显示，此类操纵策略可使非计划购买率提升42% [8]。

通过此案例，可见生成式人工智能在电商的个性化推荐和广告制作领域的应用，带来了明显的个性化操纵问题。其核心在于算法的“黑箱”特性——决策过程不透明，削弱了消费者的自主决策能力。生成式人工智能会分析用户过去的浏览记录和偏好数据，进而制作出针对性极强的个性化内容。虽然这类推荐能让用户在使用电商平台时体验更好，但算法背后的运行逻辑通常是不公开的，消费者很难搞清楚平台为什么会推荐这些内容，只能被动接受这种信息不对称的情况。现有理论研究和案例分析指出，当消费者面对这类个性化推荐时，他们自主决策的能力会受到显著削弱，这也体现出算法对消费者行为产生的实际影响。

除此之外，生成式人工智能还能精准捕捉用户的情绪状态和潜在需求，并且会动态调整推荐策略，通过这种方式潜移默化地影响消费者的购物心理。这种操纵行为不只是体现在商品推荐上，还有可能延伸到差异化定价和促销策略中，使得电商平台和消费者之间的权力不平衡问题更加突出。而这种做法的背后，其实是对消费者知情权和自主选择权的系统性忽视，这也引发了深层次的伦理方面的争议。

4.3. 责任归属模糊挑战

生成式人工智能的自主决策特性导致责任链条断裂，形成“开发者-平台-用户”三方权责真空。某电商平台的人工智能客服在售后纠纷中提供错误退货政策，导致消费者损失2300元，平台以“内容为人工智能自主生成”推诿责任，技术供应商则归咎于“平台训练数据不足”，最终消费者维权耗时超90天。该案例已纳入中国消费者协会2024年电商AI消费纠纷案例库[9]。

此类问题的核心在于，技术开发者和平台运营者之间的责任范围很难划分清楚，再加上技术更新速度快、而伦理规范跟不上，两者之间的差距进一步加剧了这种复杂性。

从技术开发者的角度来看，生成式人工智能模型在训练过程中，到最终输出内容结果，都存在一定的不可预测性。当这些生成的内容引发争议时，要确定开发者该承担什么责任，会同时面临法律和伦理两方面的难题。虽然现在的法律大多把主要责任归到电商平台身上，但模型本身在设计上的偏差、以及训练数据选择时出现的问题，也让开发者没办法完全摆脱责任。

对平台运营者来说，他们作为使用这些技术的一方，本来就有审核内容、保护用户的义务。但生成式内容数量多且动态变化，靠人工根本没办法实现全面监管。比如在个性化推荐商品、虚拟客服和用户互动这些场景里，平台常常拿“内容是系统自己生成的”当借口，推脱自身责任；而技术开发者又会把问题归咎于平台对技术的配置不当，这样就形成了一个双方互相推卸责任的灰色地带。

除此之外，生成式人工智能的更新迭代速度，远远超过了伦理规范和法律条文的更新速度，这就导

致遇到一些新型纠纷时，没有明确的依据来判断责任归属，进一步增加了责任认定的不确定性。

4.4. 可持续性矛盾挑战

生成式人工智能在电商领域快速发展并广泛应用，这凸显出技术进步与伦理规范之间存在难以持续协调的矛盾。这种矛盾主要体现在三个方面：一是技术更新时伦理规范相对滞后，二是缺乏对长期社会影响的评估，三是忽视了技术本身的环境成本。

伦理规范滞后的问题在于，生成式人工智能的更新速度远远超过了伦理规则的制定速度。比如，虚拟模特展示、个性化广告推送等具有高度真实感的内容已经在电商平台上大量使用，但针对这些内容的真实性验证和披露机制还不完善。这使得平台在没有明确伦理指导的情况下就急于应用新技术，进而加剧了消费者对内容真实性的疑虑。而且，生成式人工智能具有创造性，其行为的边界很难提前界定，这就进一步放大了监管和治理的滞后风险。

另外，目前还没有建立起评估生成式人工智能长期社会影响的体系。现在，技术的应用大多只关注短期的商业利益，比如提高转化率、优化用户体验等，但对潜在的社会成本研究不够。像虚拟试穿技术可能会改变人们的消费习惯，冲击实体零售业；算法个性化推荐可能会让用户陷入信息茧房，影响公众认知的多样性。这些长远的影响还没有被系统地纳入技术开发和治理的框架中。

此外，当前讨论普遍忽视了生成式人工智能的环境可持续性問題。生成式人工智能在电商领域的规模化应用引发严峻的环境代价，核心表现为高能耗与高碳排放对可持续发展的冲击。北京大学大数据分析与应用技术国家工程实验室2025年研究显示，全球电商行业人工智能模型的年耗电量达8.3~12.7 TWh，占全球人工智能总能耗的19%，对应碳排放量约440~780万吨，相当于300万辆燃油车的年排放量[3]。此外，大型模型的训练和推理过程消耗巨大的算力与电力，产生显著的“碳足迹”。在电商这一追求效率与规模的领域，人工智能的广泛应用无疑会加剧能源消耗与碳排放，这与全球范围内的可持续发展目标存在潜在冲突。将环境成本纳入评估框架，是“可持续性矛盾”中一个不可或缺的维度。

造成这种矛盾的根本原因是技术开发者、电商平台和监管机构的目标不一致：技术开发者追求创新，电商平台看重商业利益，而监管机构则关注社会福祉，这三者之间缺乏有效的协同。因此，构建一个跨学科、多方利益相关者共同参与的动态伦理框架显得尤为重要。

5. 多维度治理框架构建

5.1. 技术防控维度

生成式人工智能在电商生态中的应用带来了独特的伦理挑战，技术防控作为治理框架的第一维度，旨在通过技术手段直接应对这些挑战。算法透明度技术是其中的核心，通过设计可解释的人工智能模型，使消费者能够理解生成内容的逻辑来源和决策依据。研究表明，当用户能够追溯人工智能生成内容的生成路径时，其信任度可以得到有效提升，同时感知欺骗性显著降低。同时，需要在技术设计环节加入可解释性机制，让算法的运行过程更透明。

真实性验证机制则是另一项关键技术，通过数字水印、区块链等技术手段对人工智能生成内容进行标记和溯源。这种机制能够有效区分人工创作与人工智能生成内容，从而减少消费者因信息不对称而产生的认知偏差。案例分析表明，引入真实性验证有助于增强消费者对商品描述的信任。

用户控制权设计是技术防控的第三个重要方向。通过赋予用户对人工智能生成内容的定制权和选择权，例如允许用户调整个性化推荐的强度或屏蔽特定类型的人工智能生成广告，可以显著缓解个性化操纵带来的伦理问题。理论推演与行业实践显示，提供控制选项的电商平台中，用户自主决策能力的满意度通常更高。这些技术手段共同构成了生成式人工智能伦理治理的第一道防线，为后续的组织自律和社

会监管奠定了基础。

5.2. 组织自律维度

组织自律是生成式人工智能电商治理框架的核心支柱，强调企业通过内部机制建设与行业协作主动履行伦理责任。鉴于平台经济的垄断特性和利润至上逻辑，单纯依靠企业道德自觉是不现实的，必须通过具有约束力的行业标准与市场声誉机制来推动。企业应率先制定明确的人工智能伦理准则，界定生成式内容的应用边界与透明度要求，例如强制标注人工智能生成信息，以保障消费者知情权。中国电子商务协会 2024 年发布的《电商 AI 伦理自律公约》已推动 127 家平台承诺“人工智能生成内容 100%标识”[10]。同时，需设立伦理审查委员会及内部审计机制，对人工智能内容进行合规性与社会影响评估，并鼓励员工参与监督，构建全员伦理文化。在行业层面，由协会牵头制定统一伦理指南，通过案例共享与标准推广，推动形成行业共识，增强整体伦理水平与消费者信任。组织自律通过企业自律与行业共治，为生成式人工智能的负责任应用提供系统性保障。

5.3. 社会监管维度

在生成式人工智能电商治理体系里，社会监管是必不可少的外部约束手段。它的核心作用，是搭建一个多方共同参与的监管框架，以此弥补技术防控和组织自律存在的不足。

首先，要做的是完善相关法律法规。可以参考欧盟《人工智能法案》的思路，明确要求电商平台必须对人工智能生成的内容进行标识，同时对那些可能操纵用户的算法，按照风险程度实施分级管理。

其次，要建立第三方独立监督机制。可以组建由不同领域专家构成的评估机构，让这些机构从技术层面审查算法是否公平、是否容易被理解，也就是开展技术审计。同时，还要推动媒体和公众参与进来，形成多种力量共同监督的网络。

最后，提升公众的伦理素养是提高监管效果的基础。需要通过教育让消费者更清楚如何识别人工智能带来的风险，并且建立方便消费者反馈问题的渠道，从而形成社会各方共同治理的良性循环。

总的来说，把法治建设、监督机制和公众参与这三方面有机结合起来，形成社会监管体系，有助于生成式人工智能在电商领域的规范发展，提供全面系统的保障。

6. 结论与展望

生成式人工智能在重塑电商生态过程中，在提升效能的同时也引发了真实性失真、个性化操纵等多维度伦理挑战。案例分析表明，未明确标识的人工智能生成内容会显著提高消费者的欺骗性感知，算法操纵情境下用户自主决策能力亦明显下降。本研究融合科学伦理学、技术伦理学与计算机伦理视角，构建生成式人工智能伦理分析框架，深化了对人-机关系及技术-社会双向形塑机制的理解，并提出“技术防控-组织自律-社会监管”三维治理路径，为推进负责任创新提供理论依据与实践指南。未来研究可拓展跨文化样本、引入纵向设计与行为经济学等跨学科视角，持续优化动态治理框架，以协同技术发展与社会价值。

参考文献

- [1] 李华强, 张敏. 生成式 AI 在电商营销中的应用风险与治理对策[J]. 商业研究, 2024, 67(2): 45-53.
- [2] 陈昌凤, 石泽. 算法的陷阱: 个性化推荐对消费者决策的隐性控制与规制路径[J]. 新闻与传播研究, 2022, 29(3): 5-23.
- [3] 北京大学大数据分析与应用技术国家工程实验室. 当 AI 学会“创造”, 地球却在“碳息”?——探讨全球生成式人工智能背后的环境代价[R]. 2025.

- [4] 张欣. 算法解释权与人工智能治理路径研究[J]. 法学研究, 2020, 42(3): 1-20.
- [5] 张明, 李娟. 平台资本逻辑下 AI 伦理失范的生成机理与治理[J]. 改革, 2023, 37(9): 124-133.
- [6] 王莉, 刘阳. 电商 AI 生成内容的消费者感知欺骗机制及影响因素研究[J]. 消费经济, 2024, 40(1): 78-86.
- [7] 中国电子商务协会. 2024 中国电商 AI 应用伦理白皮书[R]. 2024.
- [8] 阿里研究院. 2023 年中国电商生成式 AI 应用发展报告[R]. 2023.
- [9] 中国消费者协会. 2024 年电商 AI 消费纠纷案例分析报告[R]. 2024.
- [10] 中国电子商务协会. 电商 AI 伦理自律公约[Z]. 2024.