

AI赋能下数字化营销技术的伦理调节机制研究

吕冷然

江苏大学马克思主义学院, 江苏 镇江

收稿日期: 2026年7月13日; 录用日期: 2026年7月27日; 发布日期: 2026年9月24日

摘要

在人工智能深度融入数字化营销的背景下, AI营销技术已从效率工具转变为影响消费者认知与伦理判断的重要力量。现有研究多侧重外部规制, 缺少对AI内在伦理调节机制的系统阐释。本文基于技术调解理论, 从知觉、行动与道德判断三个维度, 系统剖析AI赋能数字化营销技术对消费者伦理决策的调节机理。研究发现, 相关技术通过重构价值认知、引导消费行为与异化道德判断, 形成了全流程、系统性与深层次的伦理干预。这种干预在提升个性化服务效率与营销精准度的同时, 也带来了消费者自主性侵蚀、隐私边界侵犯、算法歧视蔓延与公平交易失衡等多重风险。生成式AI的普及更使得虚假内容、版权模糊与责任悬空等新型伦理挑战持续加剧。据此, 本文构建了全周期伦理嵌入的技术设计路径, 并提出法律规制、行业自律与消费者教育多元协同、闭环治理的框架, 旨在为数字化营销的健康发展提供理论参考与实践指引。

关键词

人工智能赋能, 数字化营销, 伦理调节机制, 算法伦理

Research on the Ethical Regulation Mechanism of Digital Marketing Technology Empowered by AI

Lingran Lyu

School of Marxism, Jiangsu University, Zhenjiang Jiangsu

Received: July 13, 2026; accepted: July 27, 2026; published: September 24, 2026

Abstract

With the deep integration of artificial intelligence into digital marketing, AI marketing technology

文章引用: 吕冷然. AI 赋能下数字化营销技术的伦理调节机制研究[J]. 电子商务评论, 2026, 15(9): 887-895.

DOI: 10.12677/ec.2026.1591070

has transformed from an efficiency tool to an important force that influences consumer cognition and ethical judgment. Existing research mostly focuses on external regulations and lacks a systematic explanation of the internal ethical regulatory mechanisms of AI. This article is based on the theory of technological mediation, and systematically analyzes the regulatory mechanism of AI enabled digital marketing technology on consumer ethical decision-making from three dimensions: perception, action, and moral judgment. Research has found that relevant technologies have formed a full process, systematic, and deep level ethical intervention by reconstructing value cognition, guiding consumer behavior, and alienating moral judgments. This intervention not only improves the efficiency of personalized services and marketing precision, but also brings multiple risks such as erosion of consumer autonomy, invasion of privacy boundaries, spread of algorithmic discrimination, and imbalance of fair trade. The popularization of generative AI has further exacerbated new ethical challenges such as false content, copyright ambiguity, and liability suspension. Based on this, this article constructs a technical design path for full cycle ethical embedding, and proposes a framework for diversified collaboration and closed-loop governance of legal regulation, industry self-discipline, and consumer education, aiming to provide theoretical reference and practical guidance for the healthy development of digital marketing.

Keywords

AI Empowerment, Digital Marketing, Ethical Regulatory Mechanism, Algorithmic Ethics

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

AI 与大数据技术正推动数字化营销向智能化、个性化发展。个性化推荐、程序化广告等技术在提升营销效率的同时，也引发了隐私滥用、算法操控、大数据杀熟等结构性伦理冲突[1]。这些问题的根源在于技术物对消费者行为的隐性引导与系统性干预。然而，现有研究多聚焦于外部监管和事后补救，难以揭示 AI 技术内在的伦理调节机制[2]。

在 AI 伦理治理的理论谱系中，价值敏感性设计(Value Sensitive Design, VSD)与助推理论(Nudge Theory)是两个具有广泛影响力的分析框架。价值敏感性设计由弗里德曼等人于 20 世纪 90 年代提出，主张在技术设计的全生命周期中系统性地纳入人类价值考量，通过将伦理价值转化为技术参数实现“设计即治理”的目标。该理论强调技术并非价值中立，开发者应当主动识别并平衡不同利益相关者的价值诉求，但其研究重心主要集中在技术设计阶段的价值植入，对技术投入使用后如何动态调节人的行为与判断缺乏深入分析。助推理论则由塞勒与桑斯坦提出，认为可以通过设计“选择架构”来引导人们做出更符合自身利益的决策，同时保留选择自由。在数字化营销领域，助推理论被广泛用于解释倒计时、低库存提示等界面设计对消费行为的影响，但该理论主要关注行为层面的引导，未能深入探讨技术对消费者知觉模式与道德判断的深层重塑，且对算法的自主性与动态适应性特征回应不足。

技术调解理论指出，技术并非中立工具，而是通过调节人的知觉、行动与判断深度参与伦理建构，技术物在人与世界之间充当着积极的调解者角色[3]。与价值敏感性设计相比，技术调解理论更关注技术使用过程中的动态伦理效应；与助推理论相比，它提供了一个涵盖知觉、行动与道德判断的完整分析框架，能够更系统地解释 AI 技术如何从根本上改变人与消费世界的关系。基于此，本文以 AI 赋能数字化营销技术为对象，沿着“技术应用 - 伦理风险 - 调节机制 - 治理路径”的逻辑链条，系统揭示其伦理调

节的内在机制与双重效应，并提出伦理嵌入设计与多元协同治理方案，以期为数字化营销的健康发展提供理论参考与实践指引。

2. AI 赋能数字化营销技术及其伦理争议

2.1. AI 赋能数字化营销技术的主要类型与趋势

当前，AI 在数字化营销中的应用已经形成了覆盖数据洞察、创意生产、精准触达、动态定价、交互服务与智能决策的完整技术体系。按照技术演进阶段与核心能力差异，可将其划分为预测型 AI 与生成型 AI 两大类别：预测型 AI 以数据分析与结果预测为核心，包括个性化推荐、程序化广告、用户画像与动态定价等技术；生成型 AI 则以内容创造与自主交互为核心，包括生成式 AI 内容生产系统与 AI 智能体等新兴形态。以下六类技术构成了这一体系的核心支柱。

第一，个性化推荐算法。该类算法基于用户的浏览历史、购买记录与搜索行为等数据，通过协同过滤、深度学习等模型预测用户偏好，并实时推送个性化的商品内容。这是电商平台应用最广泛、影响最直接的 AI 营销技术。研究表明，超过 70% 的电商平台交易受到算法推荐的影响，个性化推荐已成为消费者决策的关键入口。

第二，程序化广告系统。该系统利用实时竞价(Real-Time Bidding, RTB)技术，在毫秒级时间内完成广告位的拍卖与投放。系统根据用户画像、地理位置与设备信息等定向参数，将广告精准推送给目标人群，实现“千人千面”的广告触达。程序化广告极大地提升了广告投放的效率与精准度，但也因其对用户数据的广泛采集与共享而引发了严重的隐私争议。

第三，用户画像平台。通过整合交易数据、行为数据与社交数据等多源信息，用户画像平台为每个用户构建了包含人口属性、消费偏好与价值敏感度等维度的标签体系。用户画像是个性化推荐与程序化广告的数据基础，其构建过程本身就涉及对用户隐私的深度挖掘与商业化利用。

第四，动态定价算法。该类算法基于市场需求、库存水平与用户价格敏感度等实时数据，自动调整商品价格。动态定价在出行、旅游与零售等领域应用广泛，但也因“大数据杀熟”现象而备受争议——平台利用用户画像对高忠诚度用户实施更高定价，构成了对公平交易原则的挑战。

第五，智能客服与聊天机器人。依托自然语言处理技术，智能客服能够为用户提供 7 × 24 小时的咨询、导购与售后服务。这类技术在降低人工成本的同时，也通过精心设计的对话流程引导用户的消费决策，其劝导性特征值得伦理层面的关注。

第六，生成式 AI 内容生产系统与 AI 智能体。这是近年来涌现的新兴技术形态。生成式 AI 基于大语言模型与多模态技术，能够自动产出文案、图文、海报与短视频等营销内容，实现规模化、低成本与个性化的创意生产，彻底改变了传统营销内容的生产模式。AI 智能体则具备自主决策、环境感知[4]、持续学习与交互执行能力，可深度介入信息检索、方案推荐、比价筛选乃至自动购买的全过程，推动人机交互从被动响应向主动代理持续升级。

值得注意的是，AI 与数字化营销的融合正在不断深化。有研究指出，人工智能借助深度学习、自然语言处理等技术，不仅能够精准分析消费者行为数据，还能实现个性化推荐与智能客服，而大数据技术的应用使企业能够收集并分析海量消费者数据，为精准营销提供有力支持。与此同时，全渠道数字化营销的整合趋势日益明显，企业将线上与线下、社交媒体、电商平台等增多。这些发展趋势在提升营销效率的同时，也放大了本文所关注的伦理调节机制的作用范围与影响深度。

2.2. 数字化营销中 AI 技术的伦理争议

随着技术应用的不断深化，数字化营销中的伦理风险沿数据采集、算法运行、内容生成、决策干预

与交易达成的全链条逐步显现。以下从五个方面加以梳理。

其一，隐私侵犯与数据滥用。程序化广告的实时竞价系统在用户毫秒级的浏览过程中，将 IP 地址、设备标识、精准地理位置乃至敏感特征(如健康状况、政治倾向)广播给数十甚至数百家公司。这种“广播式”的数据共享模式，使用户在毫不知情的情况下被广泛追踪与商业化利用，构成了对信息自决权的隐性侵蚀[5]。用户在享受“免费”内容与服务的同时，其数字足迹已成为被交易的商品。

其二，算法操控与消费者自主性削弱。智能推荐算法通过排序策略、界面设计与推送节奏，系统性地干预用户的注意力分配。消费者从主动的信息搜索者转变为被动的信息接收者，其消费决策日益依赖算法提供的“选择剧场”。有学者尖锐地指出，AI 在数字化营销中的真实故事并非便利本身，而是对消费者自主权的系统性侵蚀[6]——算法以利润最大化为导向，用经过精心算计的推荐取代了真正的选择自由。

其三，大数据杀熟与算法歧视。平台利用用户画像对高忠诚度用户实施更高定价，这种“杀熟”行为的本质是一种算法歧视。经营者通过不正当采集与使用用户数据，滥用算法对用户进行精准定价，侵犯了用户的知情权、公平交易权与自由选择权。调查数据显示，近四成的网络购物消费者曾遭遇“大数据杀熟”，这一数据揭示了算法公平性问题的广泛性与紧迫性。

其四，信息茧房与认知窄化。个性化推荐算法倾向于推送符合用户既有偏好的内容，导致用户长期暴露在同质化的信息环境中。在消费场景中，这意味着消费者逐渐丧失接触多样化商品与多元评价的机会，其消费决策的参考框架被算法不断固化。信息茧房效应不仅限制了消费者的选择空间，也削弱了其作为理性决策者的认知能力。

其五，生成式 AI 带来的新型伦理风险。生成式 AI 的普及进一步放大了数字化营销的伦理挑战，具体表现为：模型幻觉可能导致虚假产品描述与误导性宣传；训练数据与生成内容可能引发复杂的版权纠纷；批量生成的虚假评论可能扰乱市场信任秩序；多模态内容的难以溯源则加剧了信息失真与舆论失序[7]。这些新型风险具有隐蔽性强、扩散速度快、影响范围广与责任认定难等特征，构成了数字化营销伦理治理的新难点。

2.3. 已有研究的贡献与不足

现有研究围绕算法伦理、数据合规等议题取得了法律规制与技术审计等方面的成果，但仍存在明显局限。多数研究停留在伦理问题的识别与描述层面，缺乏统一解释 AI 伦理影响的中层理论，对生成式 AI、AI 智能体等新技术的伦理新变化缺少深度解读。同时，研究视角多强调外部监管与事后治理，较少从技术调解的微观路径揭示 AI 如何系统性影响消费者的伦理感知与决策。由此形成了关键理论缺口：无法有效回答“技术为何产生伦理影响”“影响如何发生”“如何从内部干预”等核心问题。此外，现有研究大多将 AI 营销技术视为一个同质化的整体，未能区分预测型 AI 与生成型 AI 在伦理调节机制上的本质差异；对不同营销场景(如品牌内容营销与效果广告)的伦理风险特征关注不足；也较少探讨不同消费者群体(如高数字素养与低数字素养用户)在面对技术调解时的异质性反应。这些不足导致现有研究的理论解释力与实践指导意义受到限制。本文以技术调解理论为支撑，从内在机制层面解析伦理调节过程，将技术特征、行为逻辑与伦理后果纳入统一分析框架，以弥补现有研究的不足。

3. AI 赋能数字化营销技术的伦理调节机制

技术调解理论将技术对人与世界关系的影响划分为知觉调解与行动调解两个基本维度。结合 AI 营销技术的智能化、自主性与交互性特征，本文进一步将分析框架拓展为知觉调解、行动调解与道德判断调解三个维度。这三个维度相互嵌套、协同作用，共同构成了贯穿消费全过程的系统性伦理调节机制。

3.1. 知觉调解：AI 技术对消费者伦理感知的重塑

知觉调解指技术物改变人感知世界的方式。在数字化营销中，AI 依托用户画像与算法推荐，系统性地重构消费者对商品价值、交易公平性与消费行为正当性的伦理感知框架。

首先，智能推荐算法改变了消费者感知商品价值的参照系。当消费者面对的不再是“完整”的商品世界，而是算法筛选后的“定制化”版本——高利润商品被优先展示，“为你推荐”取代自主搜索，“销量热榜”构建默认价值序列——消费者的伦理判断便不再基于独立比较，而是在算法设定的认知基准内被动形成，其价值感知被无声锚定。预测型 AI 主要通过排序与过滤机制实现价值锚定，而生成型 AI 则更进一步，能够直接创造符合算法价值偏好的商品内容与评价信息，从源头上重塑消费者的价值感知体系。在品牌内容营销场景中，生成式 AI 可以批量生产风格统一、情感共鸣强烈的营销内容，通过沉浸式体验将品牌价值内化为消费者的自我认知，这种知觉调解的深度与广度是传统技术无法比拟的。

其次，程序化广告的重复曝光机制在无意识层面塑造商品可信度。多触点曝光在用户认知中构建“存在感”与“社会证明”，影响其对商品质量与品牌信誉的判断。这一过程往往在用户毫无察觉时完成，其隐蔽性正是伦理风险的重要来源。效果广告场景下的知觉调解更为激进，它通过精准捕捉用户的情绪波动与决策弱点，在用户最易受影响的时刻推送广告，利用认知偏差实现“瞬间说服”。此外，用户画像技术将消费者复杂的道德身份(如环保偏好、公平贸易支持)简化为可量化标签，使其原本动态、情境化的道德信念固化，伦理感知的反思性与情境敏感性被技术性消解[8]。

值得注意的是，知觉调解的强度存在显著的群体差异。低数字素养用户由于缺乏对算法运作机制的了解，更容易被算法构建的“拟态环境”所误导，其价值感知与伦理判断受技术影响的程度远高于高数字素养用户。老年人、青少年与农村地区用户是受影响最为严重的群体，他们往往难以区分算法推荐与客观信息，更容易陷入虚假宣传与消费陷阱。

3.2. 行动调解：AI 技术对消费者购买行为的引导与规训

行动调解指技术物塑造人的行为方式。AI 营销技术通过内置“行为脚本”引导甚至规训消费者的购买行为。倒计时器与低库存提示内置“稀缺性”与“紧迫性”脚本；“猜你喜欢”内置“相关性”脚本；“购买此商品的用户也购买了”内置“社会遵从”脚本。这些脚本在决策节点提供即时引导，使购买行为沿算法预设路径行进。

更值得警惕的是，AI 营销技术的行动调解具有“自适应”特征。与传统固定脚本不同，推荐算法能根据用户实时反馈动态调整策略：若用户对稀缺性提示反应强烈，算法会强化此类推送；若用户对商品犹豫，算法会通过重复曝光或价格微调“助推”决策。这种自适应性使行动调解升级为持续性的行为塑造。预测型 AI 的行动调解主要基于历史行为数据的统计规律，而生成型 AI 则能够根据用户的实时对话与情绪状态生成个性化的劝导内容，实现“一人一策”的精准行为引导。例如，AI 智能客服可以通过分析用户的语气与用词判断其购买意愿，针对性地调整话术与优惠策略，这种动态的、对话式的行动调解具有更强的说服力与隐蔽性。

从伦理学看，这种行动调解预设了特定道德取向——冲动购买被诱导，即时满足被强化。问题在于，这些道德取向并非消费者自主选择，而是被算法以“优化体验”之名无声强加。在直播电商等新兴场景中，行动调解的强度达到了前所未有的高度。算法通过实时流量分配、弹幕互动与限时秒杀机制，构建了一个高度情绪化的“消费场域”，使消费者在群体压力与时间压力下做出非理性的购买决策。

不同消费者群体对行动调解的抵抗能力存在显著差异。高数字素养用户能够识别并抵制部分算法诱导行为，他们会主动关闭个性化推荐、使用比价工具进行价格比对；而低数字素养用户则更容易被算法脚本所操控，成为冲动消费与过度消费的主要受害者。

3.3. 道德判断调解：AI 技术如何介入消费者的伦理决策

道德判断调解改变消费者“怎么判断对与错”，是 AI 伦理干预最核心、最隐蔽的维度[9]。

一、AI 推荐系统通过“偏好简化 - 标签固化 - 自动推送”的闭环削弱消费者的道德反思[9]。一个偶尔购买环保产品的用户，被标记为“环保偏好”后持续收到同类推送，其环保消费从有意识选择逐渐沦为标签驱动的自动行为，伦理决策异化为“标签顺从”。生成式 AI 进一步加剧了这一趋势，它能够生成看似客观中立的“专家建议”与“用户评价”，为消费者的伦理判断提供现成的答案，使消费者无需进行独立的道德思考即可做出决策。

二、动态定价算法通过差异化价格呈现操控公平感。当两名用户看到同一商品的不同价格时，其对“公平价格”的感知已被算法预设。平台通过隐藏更高价格、突出“专属折扣”等界面设计塑造用户公平感知，使差异化定价被合理化，消费者逐渐丧失独立判断交易公平性的能力。预测型 AI 主要通过价格歧视实现道德判断调解，而生成型 AI 则能够为不同价格提供不同的“正当化理由”，例如向高价格用户强调“专属服务”与“品质保障”，向低价格用户强调“限时优惠”与“库存清仓”，从而使不公平的定价行为在消费者认知中变得可以接受。

三、人机协同决策模式模糊了道德责任归属。消费者在算法推荐下购买不合适商品时，常陷入“被算法骗了”与“自己点的按钮”的矛盾中，其作为道德主体的完整性被削弱。AI 智能体的介入更使“谁做出了决定”难以回答。当 AI 智能体能够自主完成从信息检索到下单支付的全过程时，消费者的道德责任被进一步稀释，他们更容易将决策失误归咎于技术系统，而非自身的判断与选择。

3.4. 从个体调解到系统性调解：人 - 技协同决策模式的形成

上述三个维度的调解并非孤立发生，而是交织叠加、相互强化，最终形成系统性的人 - 技协同决策模式。在这一模式下，消费者的购买意向是“我想买”与“算法推荐我买”的复合产物；价值感知是“我认为有价值”与“平台让我觉得有价值”的复合产物；道德判断是“我认为公平”与“平台让我感觉公平”的复合产物。“消费者的意志”与“算法的引导”已无法清晰区分。

这种人 - 技协同决策模式既是 AI 营销高效运作的机制，也是伦理风险的核心来源。它迫使我们承认：道德行动的责任主体不再是纯粹的人或纯粹的技术，而是在人 - 技行动网络中生成的复合责任。这一认识为后续的责任分配与治理路径奠定了根本性的理论基础。

4. AI 赋能营销技术中的伦理责任分配

4.1. 算法能动性及其限度

AI 营销技术是否具有某种形式的“能动性”，是讨论责任分配的前提问题。以个性化推荐算法为例，其运作逻辑具有明确的指向性：它以最大化点击率与转化率为目标函数，持续将用户行为数据输入模型，输出预测性的推荐结果。这种“输入 - 处理 - 输出”的过程虽然不同于人类的意识性能动性，但在功能层面确实展现出了“指向特定目标”的系统性倾向。算法能够根据环境反馈调整自身行为(强化学习)，能够在没有人类实时干预的情况下做出决策(自动化决策)，生成式 AI 与 AI 智能体更可以在较少人工干预下自主执行复杂任务。这些特征使其区别于传统的被动工具。

然而，算法的能动性始终是有限的、派生的能动性。它源于设计者的意图、训练数据的特征与目标函数的设定，而非算法自身的“意志”或“自由”。算法不具有道德认知能力、情感判断能力与责任承担能力，不能为自己的行为后果承担道德责任——责任追究最终必须指向人类行动者。因此，在承认算法具有部分功能性能动性的同时，必须明确其限度：算法可以“调解”道德行为，但不应被视为能够“承担”道德责任的独立行动者。这一区分对于避免“责任悬空”至关重要——如果我们过分夸大算法的道

德能动性，就可能将本应由人类承担的责任不当转嫁给技术系统。

4.2. 复合责任框架的构建

在人-技协同决策结构下，传统的单一主体归责模式已难以适用。必须建立一个覆盖技术开发、运行使用与制度保障全流程的复合责任框架，实现事前预防、事中监控与事后救济的闭环治理。

第一层：设计伦理责任。技术开发者在设计算法时，应当预见并规避可能的伦理风险。这包括：选择嵌入公平性约束的目标函数、确保训练数据不具有歧视性偏误、在系统架构中预留透明度接口、防范生成式 AI 的幻觉与版权风险等。设计伦理责任是一种“事前责任”——它要求在技术投入使用之前就将伦理考量嵌入设计过程，从源头上降低伦理危害的可能性。开发者承担着技术合理性与合规性的首要责任，确保算法目标不与公共利益及消费者权益相冲突。

第二层：运行监控责任。平台运营者在技术部署之后，应当建立常态化的监测机制，持续追踪其伦理影响并做出及时调整。这包括：定期审计算法的公平性指标、对推荐结果、定价策略与 AIGC 内容开展伦理审核、建立用户投诉与救济通道、对高风险决策设置人工复核阈值等。运行监控责任是一种“事中责任”——它要求在技术运行过程中保持对伦理风险的警觉与响应能力。作为技术的直接管控方与利益的直接获取方，平台负有持续监管、风险预警、违规止损与用户保护的主体责任。

第三层：制度保障责任。政府监管部门与行业组织应当为复合行动者中的责任划分提供明确的规则框架与救济途径。这包括：制定算法透明度与可解释性的法律标准、明确生成式 AI 与自动化决策的合规要求、建立算法歧视的认定与处罚机制、为受损消费者提供有效的法律救济等[10]。制度保障责任是一种“事后责任”——它要求在伦理问题发生后能够实现公正的归责与救济。监管方与行业组织负责构建规则底线、维护市场秩序与保障治理落地，形成外部约束与引导力量。

4.3. 消费者责任的位置

在讨论 AI 营销技术的伦理责任时，不能忽视消费者自身的作用。消费者在享受智能化服务便利的同时，也应当培养基本的数字素养与审慎消费意识。然而，消费者责任的定位应当是“补充性的”而非“替代性的”。在平台与企业充分履行信息披露、透明度保障与公平性义务的前提下，消费者应理性对待推荐信息、自主做出消费选择并主动维护个人信息权益；消费者应提升对算法推荐、数据采集与虚假信息的识别能力，保持必要的审慎态度与自主判断意识。但在企业存在数据滥用、算法操纵与歧视定价等系统性违规行为时，由于双方存在显著的信息与权力不对称，消费者处于弱势被动地位，不应承担主要责任。此时，相关责任应由技术开发方与平台运营方承担。责任分配必须遵循权责对等、风险与收益匹配的基本原则。

5. 面向伦理嵌入的 AI 营销技术设计路径

5.1. 从“事后规制”到“事前嵌入”

传统的 AI 伦理治理往往采取外部监管的思路——先有技术应用，后发现问题，再通过法律与政策进行规范。这种“事后规制”模式在面对快速迭代、实时决策与全域渗透的 AI 营销技术时，往往显得力不从心。一个更为根本的治理路径是将伦理价值嵌入技术设计的全过程，即采用“伦理嵌入设计”或“价值敏感设计”的理念。这一路径强调在技术开发阶段就将自主权、公平性、透明度与隐私保护等伦理目标转化为可执行的技术参数与系统规则，使技术在运行过程中自发地实现伦理约束。其核心逻辑是将外部伦理要求内化为技术运行的内在逻辑与行为边界，从而实现从被动应对风险向主动防范风险的范式转换。在数字化营销场景中，这意味着推荐算法不应只以点击率与转化率为目标函数，还应嵌入用户自主

权、透明度与公平性等伦理约束；程序化广告系统不应只追求投放效率，还应内置隐私保护与知情同意的技术机制；动态定价算法不应只优化利润，还应设置公平交易的底线参数。

5.2. 伦理嵌入的具体实践路径

伦理嵌入理念需要落实到可操作的技术设计层面。以下五条路径构成了伦理嵌入的具体实践框架。

第一，透明度设计。让算法的推荐逻辑对消费者可见。具体措施包括：在推送个性化内容时提供简明的解释(例如“您看到此商品是因为浏览过类似商品”)；在用户界面中明确标注商业推广内容与算法推荐内容；向用户披露哪些数据被用于个性化推荐以及这些数据的来源。透明度设计有助于恢复消费者对技术物的信任与自主判断能力。需要强调的是，透明度并非简单的信息告知，而是要求以通俗、清晰、可获取的方式呈现算法的核心逻辑，真正消除信息不对称。对于生成式 AI 内容，应当建立强制性的标识制度，明确告知消费者内容由 AI 生成，避免消费者将其误认为人类创作的客观信息。

第二，用户控制权强化。赋予消费者对个性化推荐与自动化决策的有效拒绝权。我国《中华人民共和国个人信息保护法》第 24 条第 2 款已规定，通过自动化决策方式进行信息推送与商业营销应当提供不针对个人特征的选项或便捷的拒绝方式。从伦理嵌入的视角看，这种拒绝机制本身就是一种技术设计层面的伦理保护——它使消费者能够“退出”技术的过度调解，恢复一定程度的自主空间。控制权的核心是让消费者真正掌握个人数据使用与个性化服务的开关，而非仅仅停留于形式上的同意。应当为消费者提供更细粒度的控制权，例如允许用户选择推荐内容的多样性程度、调整个性化推荐的强度、关闭特定类型的广告推送等。

第三，公平性参数嵌入。在算法设计中纳入公平性约束，防止因用户画像差异而导致的价格歧视与机会不均。具体可以包括：设定公平性损失函数来约束算法的优化方向，使算法在追求效率的同时兼顾不同用户群体的公平待遇；引入公平性审计机制来定期评估算法的决策偏差；对高风险决策(如差异化定价)设置人工复核阈值。公平性不仅要求结果的公平，更要求过程的公平与群体间机会的均等。对于动态定价算法，应当设置价格浮动上限与下限，禁止对同一商品在同一时间向不同用户收取显著差异的价格；对于推荐算法，应当引入多样性指标，避免过度推送同质化内容。

第四，可解释性与可追溯性。确保算法的关键决策能够被理解与追溯。这不仅有助于监管机构的审计，也为消费者提供了质疑与救济的依据。在技术层面，可以开发可解释的推荐模型(如基于注意力机制的解释生成)，或保留完整的决策日志以备事后审查。可追溯体系是责任认定、风险回溯与纠纷解决的重要技术基础。对于生成式 AI 内容，应当建立完整的溯源机制，记录内容的生成时间、生成模型、训练数据来源与修改历史，以便在发生版权纠纷或虚假宣传时能够明确责任主体。

第五，生成式 AI 的专项伦理约束。针对生成式 AI 带来的新型风险，应建立专项伦理约束机制。具体包括：通过幻觉检测技术过滤虚假信息；建立版权规范与内容标识制度；对 AIGC 营销内容进行伦理审核；拦截虚假、侵权与误导性信息，稳定市场信任秩序。生成式 AI 的伦理约束是应对新型风险、保障行业长期健康发展的关键环节。

5.3. 多元协同治理框架

单纯依赖技术层面的伦理嵌入设计，难以完全应对复杂多元的伦理风险。必须构建法律规制、行业自律与消费者教育协同发力的多元治理体系。在法律层面，应细化《中华人民共和国个人信息保护法》《中华人民共和国电子商务法》等相关法律法规的实施细则，明确生成式 AI 与自动化决策的合规要求，强化执法力度与救济机制——法律制度为伦理治理提供刚性底线与强制约束力。在行业层面，应推动平台与技术企业建立行业伦理准则，开展第三方算法审计与伦理评估，强化自我约束——行业自律能够弥

补法律滞后性的不足,提升治理的灵活性与专业性。在消费者教育层面,应提升公众对算法推荐、数据采集与伦理风险的认知水平,培养理性判断能力与自我保护意识——消费者素养的提升是构建健康数字生态的基础性环节。特别需要关注低数字素养群体的消费者教育,通过社区宣传、公益讲座等通俗易懂的方式,帮助他们了解算法的基本原理与潜在风险,提高自我保护能力。通过上述三个层面的协同发力,可以形成多方参与、协同共治的良性格局,推动 AI 营销技术从“操纵消费者”的工具转向“赋能消费者”的助手。

6. 结论

本文基于技术调解理论,系统解析了 AI 赋能数字化营销技术的伦理调节机制。研究表明, AI 营销技术通过知觉调解、行动调解与道德判断调解三个维度,深度介入消费者的伦理决策过程,形成人-技协同的复合决策模式。这种调节具有双重效应:既可提升营销效率与用户体验,也可能侵蚀消费者自主权、侵犯隐私边界、破坏公平交易。预测型 AI 与生成型 AI 在伦理调节机制上存在显著差异:预测型 AI 主要通过数据分析与结果预测实现调解,而生成型 AI 则通过内容创造与自主交互实现更深层次、更具个性化的调解。不同营销场景与不同消费者群体受到的技术影响也存在明显异质性,这要求我们在制定治理策略时应当采取分类施策的原则。生成式 AI 的普及进一步放大了虚假内容、版权模糊与责任悬空等新型挑战。

针对上述问题,本文提出了透明度设计、用户控制权强化、公平性参数嵌入、可解释性与可追溯性保障及生成式 AI 专项约束等伦理嵌入路径,并构建了法律规制、行业自律与消费者教育相结合的多元协同治理框架,旨在推动 AI 营销技术从“操纵消费者”转向“赋能消费者”。未来研究可向实证检验、生成式 AI 治理、跨文化比较及跨学科融合等方向深入。

基金项目

本文系 2025 年度江苏省研究生科研与实践创新计划项目(项目编号: KYCX25_4133)研究成果。

参考文献

- [1] 陈兵. 规制“大数据杀熟”: 算法向上向善的治理进路[N]. 人民论坛, 2025-09-16(04).
- [2] 刘铮. 技术物是道德行动者吗?——维贝克“技术道德化”思想及其内在困境[J]. 东北大学学报(社会科学版), 2017, 19(3): 221-226.
- [3] 程海东, 贾璐萌. 道德物化——技术物道德“调解”解析[J]. 道德与文明, 2014(6): 111-116.
- [4] 孙鸿亮. 基于自动化决策的商业广告的法律规制[J]. 中国法律评论, 2025(4): 196-210.
- [5] 谭婷婷, 易灿. 智能推荐算法下的消费者决策机制演化研究: 基于信息加工与行为经济学视角[N]. 今日消费, 2025-08-21(A06).
- [6] Bouchareb, N. (2025) When Algorithms Sell: Rethinking Consumer Behavior in the Ai-Enhanced Marketplace. *AI & Society*, 41, 469-470. <https://doi.org/10.1007/s00146-025-02453-0>
- [7] Machidon, O.M. (2025) Beyond Algorithethics: Addressing the Ethical and Anthropological Challenges of AI Recommender Systems. *Journal of Media Ethics*. <https://doi.org/10.1080/23736992.2025.2584435>
- [8] Veale, M. and Zuiderveen Borgesius, F. (2022) Adtech and Real-Time Bidding under European Data Protection Law. *German Law Journal*, 23, 226-256. <https://doi.org/10.1017/glj.2022.18>
- [9] Saura, J.R. (2025) Ethics and Privacy in AI-Driven Digital Marketing: Navigating the Paradox. In: *Oxford Intersections: Social Media in Society and Culture*, Oxford University Press, 45-62. <https://academic.oup.com/edited-volume/60551/chapter-abstract/523642777?redirectedFrom=fulltext>
- [10] Gregory, N. (2020) The Role of the Designer as Moral Intermediary. North Carolina State University.