

“共赢” vs. “牺牲”：道德消费叙述框架 对消费者算法推荐信任的影响*

徐 岚¹ 陈 全¹ 崔 楠² 辜 红³

(¹ 武汉大学经济与管理学院; ² 武汉大学组织营销研究中心, 武汉 430072)

(³ 华中师范大学城市与环境科学学院, 武汉 430079)

摘 要 消费者在做出道德消费选择时, 要面对丰富的道德产品, 复杂地权衡功利利益和道德利益。算法决策推荐可以减轻道德消费决策的复杂性, 但是消费者在涉及道德伦理的权衡决策中对算法存在不信任, 因为算法是一种典型的以功利论来处理道德权衡问题的方式。本研究提出采用“共赢”(vs. “牺牲”)叙述框架对道德消费中的算法推荐信任有积极影响, 消费者对功利论式道德观念的接受程度中介了上述积极效应。实验 1 发现“共赢”(vs. “牺牲”)叙述框架会增加消费者选择算法推荐的意愿。实验 2 验证了“共赢”(vs. “牺牲”)叙述框架对消费者算法推荐信任的正向影响, 以及功利论式道德观念接受度的中介作用。实验 3 进一步识别了上述影响的边界条件, 即道德消费的“共赢”(vs. “牺牲”)叙述框架仅会增强消费者对算法替代型决策推荐的信任, 但是并不会改变消费者对于算法增强型决策推荐信任。

关键词 道德消费, 算法推荐, 道义论, 功利论, 框架效应

分类号 B849: F713.55

1 引言

道德消费的观念日益普及, 企业推出了越来越多的具有道德属性的商品, 如更加环保、更符合公平交易原则等。购买道德产品确实能为环境和社会带来好处, 其提供的道德属性也能给消费者带来积极的精神价值。消费者在面对道德产品购买时往往不再纠结于要不要响应道德消费观念的号召购买道德产品, 而是受困于如何选择一个既可以满足功利需求又可以实现道德价值的产品。通常, 道德消费往往要求消费者做出一定的“牺牲”, 例如承担更高的价格成本, 或者是接受产品特点的一些改变、增加使用成本(Luchs et al., 2010)。这就意味着购买道德产品需要在道德价值和功利价值之间进行权衡(Niven & Healy, 2016), 即为了实现产品上的道德价值(如保护环境或追求公平)就必需付出功利上

的损失(如支付更高价格、减少使用便利性等)。由于道德价值难以计算和衡量, 这使得消费者在面对道德产品决策时面临困难, 他们很难确定究竟该选择多大的道德价值以及付出多少功利损失才是适合的, 这种决策的困难性是消费者将道德消费意愿转换为实际道德行为之间的重要障碍(Hassan et al., 2013; Herzstein et al., 2020)。

算法推荐被认为是降低消费者决策困难的一种重要工具(Puntoni et al., 2021)。然而, 先前研究表明人们在面临涉及伦理道德的权衡决策时, 并不愿意接受由算法提供的建议(Bigman & Gray, 2018; Dietvorst & Bartels, 2022)。换言之, 在道德消费情境中, 消费者可能会认为算法能够计算和识别出功利性结果最优化的消费选项, 但是它不具备理解和衡量消费所能产生的道德价值的能力(Bigman & Gray, 2018; Gray et al., 2007), 所以消费者并不信

收稿日期: 2022-11-04

* 国家自然科学基金项目(72072135; 72172110)资助。

通信作者: 陈全, E-mail: chenquan_0915@163.com

任算法提供的有关道德消费决策的推荐。

本研究旨在为如何增加消费者在道德消费情境中对算法推荐的信任,提供一个可能的解决思路。具体而言,本文提出营销者可以通过改变对道德消费的叙述框架策略来影响消费者对算法推荐的信任——当营销者将道德消费描述为“共赢”(vs. “牺牲”)时,消费者更可能信任算法对有关道德消费决策的推荐。“牺牲”叙述框架强调在道德消费中,消费者需要做出个人牺牲来帮助环境或社会,而“共赢”叙述框架则强调在道德消费中,消费者与道德相关方之间不存在功利利益的冲突。“共赢”叙述框架策略使消费者将自己与道德相关方纳入到一个共同利益体中,从而协调了功利价值与道德价值之间的对立关系,将道德价值统一在更宽泛的功利价值范畴之中(例如对于公平贸易产品而言,其强调社会公平的道德价值可以转化为一种更宽泛和基础的功利性价值,即采用公平贸易将最终带来对包含消费者在内的所有市场参与方的持续交易公平和总体利益提升)。因此,接受“共赢”(vs. “牺牲”)叙述框架信息的消费者更有可能认为在道德消费情境中,实现道德价值与实现利益共同体总体功利价值优化是共通的,此时他们更可能相信遵循功利价值最大化原则的算法工具也能够用来解决道德消费决策问题。

启用“共赢”或“牺牲”叙述框架实际上影响了消费者对道德消费决策原则的理解。将道德消费看作“牺牲”意味着消费者认为追求道义正确要付出功利上的代价,消费者要权衡道德价值和功利价值这两种不同的价值。“牺牲”叙述框架通过放弃或贬低产品的非道德价值(产品的功利性价值)来强调人们对道德价值“正确性”的追求。这强化了道义论的道德观念,即认为“道德正确性”是神圣的、高于世俗价值的,人们要做道义上正确的消费行为,并抵制那些“不正确”的世俗理念和行为(Miller & Prentice, 2016)。相比而言,通过让消费者将道德消费看作是“共赢”,会让消费者感知自己与道德相关方是利益共同体(Hiller & Woodall, 2019),这有助于强化基于功利论的道德观念。该观点认为追求利益共同体的利益最大化是实现道德正确性的一种合理方式(Conway & Gawronski, 2013)。因此,本研究提出,启动“共赢”(vs. “牺牲”)叙述框架会促使消费者更加认同以功利论式的决策原则来做道德消费决策,从而增加了人们对道德消费情境中算法推荐的信任程度。

通过 3 个实验,本研究检查了道德消费叙述框架策略(“共赢”vs. “牺牲”)对算法消费推荐信任的影响。实验 1 首先检验了启动“共赢”(vs. “牺牲”)叙述框架如何影响消费者对推荐来源(算法/人类专家)的选择偏好,结果显示“共赢”(vs. “牺牲”)叙述框架显著提升了消费者选择算法推荐的倾向。实验 2 进一步检验启动“共赢”(vs. “牺牲”)叙述框架对消费推荐的信任程度的影响,以及功利论道德价值观念的中介作用,结果显示“共赢”(vs. “牺牲”)叙述框架提升了消费者对算法消费推荐的信任,但其对人类专家的推荐信任并无显著影响。最后,本研究对“共赢”(vs. “牺牲”)叙述框架策略的实施边界进行了检查。研究发现,“共赢”(vs. “牺牲”)叙述框架的积极作用只对算法替代型推荐有效,而对算法辅助型推荐并无影响。

1.1 道德消费决策的复杂性

道德消费观念认为消费者应该考虑消费活动的道德影响,倡导消费者做出更符合道德原则的消费行为,典型的道德消费原则包括环保主义、公平交易原则等(Hiller & Woodall, 2019)。因为人们有基本的道德观,希望能做道义上正确的事情。道德原则对人们的行为起到约束作用,如果有所违反,会产生一些负面的道德情绪,如内疚、羞愧(Zollo, 2021)。遵守道德原则是人们的自我存在和自我实现的重要内容(魏心妮 等, 2023)。现有研究发现道德消费行为能够增加幸福感(Hwang & Kim, 2018),缓解社会排斥所带来的威胁(Trudel et al., 2020),让消费者感受到一种暖阳效应(Tezer & Bodur, 2020)。所以道德消费行为能为消费者提供道德价值,即获得基于行为道德正确性的一种精神价值。

消费者在做具体的道德消费决策时,难以避免地面临各类相关价值权衡的问题(Hiller & Woodall, 2019)。消费决策要考虑的目标是多元的。尽管消费者有意愿响应道德消费的号召,将实现相关道德价值的目标加入到消费决策中,但其在实际决策中不只是单纯地将道德作为唯一的消费目标和决策原则(Hiller & Woodall, 2019)。相反,消费者在决策中需要同时考虑产品的道德价值及其与其他功利价值(非道德价值)之间的关系。消费者进行道德产品消费的目的是“消费”而不是纯粹的“做慈善”,这使得其需要从备选的道德产品中挑选出在道德价值和功利价值方面均适合的产品。

道德消费决策的复杂性主要反映在消费者需要衡量道德价值并权衡众多可能与之冲突的功利

价值。在道德产品选择日益丰富的当下,道德消费决策权衡所需要考虑的信息更多。当消费者面临的购买决策更加复杂时,他们更可能选择放弃或者延后这一决策(Chernev et al., 2015; Herzenstein et al., 2020),导致对道德消费的倡导难以转化为消费者实际的行动。

算法决策的发展给降低道德消费决策的复杂性提供了一种可能。消费者可以将道德消费决策任务委派给算法系统,最终只需要考虑算法所提供的建议即可(Puntoni et al., 2021)。其前提是消费者能够信任算法,愿意在道德消费决策中采用算法决策提供的信息。消费者对算法消费决策的信任取决于消费者对算法本身以及所要处理的决策的认识。本研究将探讨消费者在何种情况下,更可能在道德消费中信任算法决策。

1.2 算法 vs.人类的消费建议

算法能在诸多方面为人类的决策提供帮助。早期研究发现算法在许多预测任务中都展现出高于人类的预测精度(Dawes, 1979),能避免人类认知能力的局限(Lim & O'Connor, 1996)。随着技术的进步,算法预测和决策的精准度在许多专业领域已经超过了专业人士,例如审计(Commerford et al., 2022)、医疗(Leachman & Merlino, 2017)、销售(Luo et al., 2019)、人力资源(Tong et al., 2021)等。此外,算法决策建议还是一种非常高效的问题解决方案。人类受制于触及范围的限制,无法大规模地提供建议服务,但是算法系统能针对大规模的用户群体提供个性化、准确的决策推荐。个性化消费推荐就是一个典型的例子(Davenport et al., 2020; Logg et al., 2019)。当前,提供基于算法的个性化消费推荐已成为在线消费平台的常态,消费者也因此受益良多(Gai & Klesse, 2019)。

然而,现阶段的算法决策尚未发展到通用智能的层次,还有许多缺陷(Haenlein & Kaplan, 2019)。因此,研究者开始关注消费者的算法厌恶,即对算法所生成建议的不信任(Burton et al., 2020)。算法不信任问题出现的一个重要原因是消费者主观地认为在一些特定场景中算法决策的质量不能得到保证。研究发现,消费者认为算法更擅长处理客观性(vs.主观性)的决策任务(Castelo et al., 2019)。针对享乐型(vs.实用性)消费情境,消费者对算法建议的信任程度会更低(Longoni & Cian, 2022)。在医疗情境中,消费者认为算法无法像人类医师一样识别个体的病情差异,因而随着自身情况特殊性感知增强,

更不愿意选择基于算法决策的医疗服务(Longoni et al., 2019)。算法还被认为不具备像人一样的创造力,更擅长一些重复性的、要求精度的任务(Logg et al., 2019),人们会认为算法产生的作品不具有像人类作品一样的诚意和真实性(Jago, 2019)。

在涉及道德伦理的应用领域,算法决策遭到更大的质疑和挑战(Martin, 2019)。道德伦理评价往往被看作是人类才有的能力,它涉及人类大脑对道德经验和相关情感因素的处理(Gray et al., 2007),而算法系统被认为无法具备像人一样完整的道德评判能力,因此无法胜任道德相关决策(Bigman & Gray, 2018; Jebari & Lundborg, 2021)。研究发现,人们会直觉性地认为算法决策是一类计算式的和追求最大化结果利益的决策方法,当决策内容涉及道德方面的权衡时,人们会认为采用算法(vs.人类)决策不是一个合适的决策方案(Dietvorst & Bartels, 2022)。这是由于在涉及道德因素的决策中,人们往往关注决策是否在道义上是正确的,更胜于决策是否带来更多的功利利益。所以人们会担忧算法在道德相关决策中过多地关注功利价值,会更可能导致在道义上错误的决策(Dietvorst & Bartels, 2022)。

在本研究所关注的道德消费情境中,消费者所进行的权衡会涉及功利利益和道德价值这两方面因素。按照 Dietvorst 和 Bartels (2022)的理论,消费者认为算法只能对功利利益进行计算,如果追求功利利益和追求道德价值之间存在冲突,那么消费者就更不愿意在有关道德消费的决策中信任算法。反之,如果能够通过改变消费者有关道德消费的信念来协调追求功利利益和追求道德价值之间的关系,消费者在道德消费决策中对算法的信任就可能增加。

1.3 道德消费叙述框架策略与道德消费决策

对同一事物采用不同的叙述方式进行描述能够影响人们相关的判断,这被称为框架效应。事物具有多面性,不同的叙述框架中的预设背景信息影响人们对事物理解的构建过程,这就会强化人们某个特定类型的认识,进而形成不同的态度和行为倾向(Sher & McKenzie, 2006)。人们对某些事物的习惯性的语言描述就可能蕴含着特定的认知预设,从而会导致框架效应。有关亲密关系的研究发现,将爱情理解为两个人互为天造地设的伴侣,相较于将爱情理解为一段旅程,会更加加重关系冲突所带来的伤害感知(Lee & Schwarz, 2014)。关于目标的研究发现,将目标实现比作一段旅程,相较于比作达到某个目的地,会更有利于人们在完成阶段目标后

继续与目标一致的行为,因为人们受比喻的影响更多地关注目标完成的过程而非结果,进而启动了成长心态(Huang & Aaker, 2019)。也就是说,叙述框架改变了人们对特定情境中关系和目标的认知构建。

道德消费决策涉及道德价值和功利价值的衡量和取舍。对道德消费的不同叙述框架可能影响人们对道德价值和功利价值之间关系的认识,进而改变消费者对如何实现道德价值的看法。这会影响消费者对道德消费决策难点的认知,如果消费者认为道德消费决策的难点正是算法决策所擅长解决的,他们会更信任基于算法的道德消费决策建议。反之消费者的信任就会更低。

1.3.1 道德消费叙述框架对算法信任的影响

对道德消费的一种典型的叙述框架是,要追求相关的道德价值,就需要做好“牺牲”部分功利价值的准备。“牺牲”是指消费者愿意放弃宝贵的资源,以换取重要和有价值的东西(Gao et al., 2020)。“牺牲”叙述框架就意味着假定在出现个人功利价值损失时,消费者依旧愿意做出道德消费行为。所以“牺牲”叙述框架首先会强化消费者的这样一种认识,即道德消费行为有可能造成消费者个人功利价值损失。此外,“牺牲”叙述框架还强调消费者要降低对所承担代价的敏感性(Gao et al., 2017)。换言之,“牺牲”叙述框架倾向于贬低和隔离人们对功利价值的追求,认为追求功利价值无助于、甚至是会妨碍道德价值的实现(Ruttan & Nordgren, 2021; 辛自强 等, 2022)。

当营销者采用道德消费的“牺牲”叙述框架时,个人功利价值的潜在损失会被强调,消费者会更多地考虑到追求自身功利价值与实现道德价值之间的冲突。消费者更可能将道德消费决策理解为在道德价值和个人功利价值之间进行取舍的问题。其决策的难点在于消费者应该在多大程度上为了道德价值而牺牲功利价值(Hiller & Woodall, 2019)。由于道德价值的特殊性,做出道德价值和功利价值的权衡需要具有对道德价值进行认知的能力,这种能力往往被认为是算法所不具备的(Bigman & Gray, 2018; Jebari & Lundborg, 2021)。并且算法所遵循的计算功利利益最大化的决策方式被认为会阻碍消费者实现道德价值(Dietvorst & Bartels, 2022)。由于消费者会认为算法在解决涉及道德和功利价值之间冲突权衡这样的问题上存在不足,他们不太可能信任由算法提供的道德消费决策建议。

本研究提出道德消费的“共赢”(vs. “牺牲”)叙

述框架有可能提升消费者在道德消费决策中对算法的信任。“共赢”叙述框架强调通过道德消费实现道德价值与追求功利价值之间并不冲突,而是协同的关系。这意味着道德消费行为是要实现消费者与道德相关方的共同利益最大化,对双方来说是“共赢”。商业情境中的一类重要的道德愿景是商业决策应当将人类的共同利益视作是最重要的(Gustafson, 2013)。这一愿景的潜在含义就是在商业中实现道德价值的目的和意义就是为了人类整体能够更好地发展(Yamoah et al., 2016)。在道德消费的场景中,消费者与道德相关方实际上是利益共同体,追求双方共同利益最大化就是应当追求的道德价值目标(Xu & Ma, 2016)。

当消费者认为自己与道德消费涉及的道德相关方是利益共同体时,个人和道德相关方的功利价值能实现协同。合适的道德消费价值体现为可以提升共同体的整体利益。由于追求双方功利价值的最优化与实现道德价值之间不存在矛盾,消费者会认为共同价值的最大化就是道德消费观念的应有之意。

在“共赢”叙述框架下,消费者认为共同利益最大化对每个成员都是更好的,不会担忧会因为考虑功利价值而忽视道德价值。如此,消费者会将道德消费当中实现道德价值的目标,转化为与其一致的追求利益共同体利益的功利最大化目标。道德消费的相关决策就成为了如何实现双方共同功利价值的最大化的问题。算法决策被认为具有解决该类问题的优势,它能够全面搜集和处理多方面数据、快速比较和计算各种功利价值属性信息,有利于更好地优化全局决策方案以实现共同利益的最大化(Puntoni et al., 2021)。

需要补充说明的是,本研究所关注的道德消费的叙事框架(“共赢” vs. “牺牲”)与前人研究涉及的道德产品的“利己(vs. 利他)诉求”是不同的概念。本文提出的道德消费的叙事框架(“共赢” vs. “牺牲”)影响的是道德消费价值衡量中道德价值和功利价值之间的关系感知。前人研究所涉及的“利己(vs. 利他)诉求”区分的是道德产品的价值受益者的不同。例如,营销者会通过广告向消费者强调购买道德产品对“自己”(vs. “他人”)价值上的好处,从而驱动消费者进行道德产品购买(Green & Peloza, 2014; Ryoo et al., 2020)。“利己(vs. 利他)诉求”的作用是让道德消费产品更能分别迎合特定情境下的、或者特定类型的消费者的自我中心(Egoistic)或利

他主义(Altruistic)倾向(Ryoo et al., 2020)。

然而,使用“利己(vs.利他)诉求”概念框架无法解释其对算法推荐信任的影响。在利己/利他研究中,研究者关注的是消费者购买道德产品的动机,而非对道德产品价值的衡量过程。无论消费者是出于利己还是利他诉求来考虑道德产品,道德产品提供给消费者的价值都是多元(既包括功利价值也包括道德价值)而非单一的,利己/利他诉求因此并不能解释和影响消费者比较和权衡这些多元价值的心理过程。而这一过程正是消费者是否信任算法推荐的基础,因为算法推荐的目的是帮助消费者在众多道德产品的价值之间进行比较。使用不同叙述框架(“共赢”vs.“牺牲”)将影响消费者对产品道德价值和功利价值之间关系的理解,从而可能解决消费者在道德消费产品推荐情境中的算法信任问题。

总结来说,在面临道德消费情境时,消费者的道德目标会启动(Reczek et al., 2018)。“牺牲”叙述框架强调在道德消费中,自己与其他利益相关方存在利益冲突,实现道德价值的消费行为意味着“牺牲”自身利益。由于道德价值被认为仅与人类坚持的“道德正确”原则相关,与功利价值这样的可计算的世俗性价值无法相容,消费者会更倾向于认为道德价值是一种特殊的、只有人类才能理解和判断的价值,因而不信任基于算法的决策推荐(Bigman & Gray, 2018)。而“共赢”叙述框架强调实现道德价值和追求功利价值之间并不冲突,在道德消费中,自己与其他利益相关方是利益共同体,做出符合道德原则的消费行为意味着使得双方的共同利益更优化,形成“共赢”(Gray & Schein, 2012; Gustafson, 2013)。由于“共赢”叙述框架让消费者理解到追求道德价值是可以通过实现共同利益的最大化来实现的,而算法可以通过其计算功利价值最大化方案的优势来助力实现道德价值,因此,消费者会更愿意信任由算法提供的决策推荐。由此本研究提出以下假设:

H1: 采用道德消费的“共赢”(vs.“牺牲”)叙述框架增加了消费者在道德消费决策过程中对算法推荐的信任程度。

1.3.2 功利论道德价值观念的中介作用

在规范伦理的理论体系中,存在两种有关如何对道德价值进行判断的观念:道义论与功利论。其中道义论强调道德价值的评判依据是行为是否符合道德原则,道德性是一种关于行为本身对错的评价;而功利论则强调道德价值评判的依据取决于行

为是否产生更优的人类总体利益结果,功利论属于一种典型的结果论(Mencel & May, 2009)。由于道德伦理的出发点是强调人们应该做道义上正确的事情,因此道义论的观点相对于功利论而言更容易被人理解和赞同。但道义论的关键问题在于对“道德正确”的界定往往是因人而异的,不同文化、不同社会群体及持不同政治立场的人们对道德正确的解释可能相差迥异。相对而言,功利论从利益结果的角度来看待这一问题,将人类利益结果的最优化定义为“道德正确”原则。这在一定程度上简化了道德价值的判断原则,使具有不同文化和社会背景的人们对道德价值的判断趋向一致。

道德消费的叙述框架将影响和形塑人们的道德价值观念。当受到道德消费的“牺牲”叙述框架影响时,消费者会认为道德价值的重要性高于功利价值,且实现道德价值与追求功利价值是相冲突的。因此消费者在道德产品消费中,会分别对产品的道德价值和功利价值进行评估,并优先考虑实现道德价值目标的可能性。如此,消费者认为追求功利价值是有助于实现产品道德价值的。道德消费决策的核心是要首先确定道德正确性的程度,再考虑需要做出怎样的牺牲来实现道德价值。

当消费者受到道德消费的“共赢”叙述框架影响时,消费者认为自己与道德相关方是利益共同体,道德消费在实现道德价值与追求功利价值之间并不冲突,两者是协同和共赢的关系。消费者因而更倾向于接受功利论道德观念,从利益方功利价值最大化的角度来理解道德正确性,相信道德价值可以通过衡量共同利益最大化来评估。因此,在“共赢”(vs.“牺牲”)叙述框架的影响下,消费者会倾向于认同道德消费决策要符合功利论的道德价值观念。功利论的道德价值观念意味着道德价值的衡量是以功利结果的最优化为标准的。消费者会认为算法在功利价值最大化的计算上具有优势,从而增加了其对算法推荐的信任。由此本研究提出以下假设:

H2: 道德消费的叙事框架(“共赢”vs.“牺牲”)对消费者算法道德推荐的信任的影响作用,受到消费者对功利论道德价值观念接受度的中介。具体而言,“共赢”(vs.“牺牲”)叙述框架通过提升功利论道德价值观念在道德消费决策中的接受度进而增加了消费者对算法推荐的信任程度。

1.3.3 算法推荐类型的调节作用

道德消费叙述框架对算法推荐信任的影响之所以存在,是因为消费者认为道德价值可以通过衡

量总体功利价值来进行评估,此时消费者对算法执行道德推荐的信任程度才可能提高。然而,消费者也可能会考虑将道德产品决策中有关道德价值和功利价值的衡量进行拆分,比如借助算法来衡量产品的功利价值,在此基础上再由人类专家根据对产品道德价值的衡量来做出综合判断,并形成最终决策。

算法应用的实际情境中,算法推荐应用除了算法替代(人类)型,还包括算法增强型(Brynjolfsson & Mitchell, 2017)。算法增强型推荐是指算法与人类共同协作来完成决策推荐。许多算法决策应用的研究发现算法增强型决策更容易被消费者信赖(Luo et al., 2021; Sampson, 2021)。因为消费者预期两者的结合能够形成能力互补(Burton et al., 2020)。本研究认为算法推荐类型(算法替代型推荐 vs. 算法增强型推荐)可能成为道德消费叙述框架影响消费者算法推荐信任的一个调节因素。由于算法增强型决策中,人类能够弥补算法在道德伦理议题上的短板。因此,无论消费者认为道德消费中道德价值和功利价值的关系是冲突或是一致的,消费者在算法增强型推荐中都能充分利用算法在处理功利价值方面的优势,从而使叙述框架对算法增强型推荐的信任程度的影响被削弱。由此,本研究提出以下假设:

H3: 道德消费叙述框架(“共赢” vs. “牺牲”)对算法推荐信任的影响受到算法决策类型的调节。对算法替代型推荐而言,“共赢”(vs. “牺牲”)叙述框架会增强消费者对其的信任,然而对于算法增强型推荐而言,“共赢”(vs. “牺牲”)叙述框架不会影响消费者对其的信任。

2 实验 1: 道德消费叙述框架对消费者算法推荐偏好的影响

2.1 实验目的

本实验将检验不同类型的道德消费叙述框架对消费者算法推荐的偏好的影响。由于在道德消费决策中,消费者需要考虑道德和功利价值,受到不同叙述框架(“共赢” vs. “牺牲”)的影响会导致消费者对道德消费情境中的道德价值和功利价值之间关联的看法存在差异。本研究预测,“共赢”(vs. “牺牲”)叙述框架会让消费者更愿意以功利论式道德价值评价的方式进行决策。因此,消费者会在决策中更多地选择由算法提供决策辅助。

2.2 实验设计与被试

本实验采用单因素组间设计(道德消费叙述框架类型:“共赢” vs. “牺牲”)。因变量则是被试在道德消费情境中对专家和人工智能算法决策辅助的相对偏好。通过 G*power 软件计算出在显著性水平为 0.05 且效应量水平(Odds ratio)为 1.6 时,预测达到 90%的统计力水平的总样本量至少为 171 名。本实验从见数(Credemo)平台招募了 199 个被试参与线上实验,被试平均年龄 27.39 岁($SD = 7.63$),其中 25.63%为男性。被试被随机分配到两个实验组中(“共赢”组 vs. “牺牲”组)。

2.3 实验材料与程序

本研究关注的自变量是道德消费叙述框架类型,“共赢”框架强调道德消费能为消费者和道德相关方都带来好处,“牺牲”框架强调消费者在道德消费中应该愿意放弃部分自己的功利利益,提升他人利益,从而更好地满足道德原则的要求。本研究采用不同的道德消费解释语句来引导消费者的看法,从而实现对自变量的操纵。

具体而言,实验中我们请被试想像自己主要用咖啡帮助自己在工作中提神,会经常光顾一家专业的咖啡豆网购平台。该平台近期推出了一系列“公平交易认证”的咖啡豆,并倡导消费者购买这类咖啡豆。消费者购买这类咖啡豆会让弱势产地种植户获得额外的经济补贴,支持他们过上更好的生活,以及持续经营和提升咖啡豆品质,所以价格一般更高。被试随后看到一张有关“公平交易认证”的宣传海报。在“共赢”组中,海报的宣传语是“关爱创造价值,共赢美好生活”,以及“构建‘共赢’社会,让世界更公平”。在“牺牲”组中,海报的宣传语是“为公平而牺牲,成就一份关爱”,以及“做一点‘牺牲’,让世界更公平”。本实验向被试介绍的咖啡产品被描述为来自国外欠发达的产地,好处是被试对国外不同地区的共情程度会比较相近,避免被试可能对国内某些地域有更特殊的情感而影响其判断(童泽林 等, 2019)。

随后,被试回答了关于功利价值与道德价值之间冲突感知的问题。问项包括:“支付道德产品的溢价就意味着消费者会有一定物质利益损失”,“种植户拿到更多利益,消费者就会得到更少实惠”,“消费者与种植户的物质利益之间是存在冲突的”(1 = 非常不同意, 7 = 非常同意, $\alpha = 0.85$)。这些问项是作为自变量的操纵检查,“共赢”(vs. “牺牲”)叙述框架会让被试感知功利价值与道德价值之间的冲突

更小。

然后,引导材料告知被试,根据产区和加工工艺的区别,“公平交易认证”的咖啡豆相较于类似的普通咖啡豆差价幅度很广。所以选择一款符合口味偏好,有更好道德性,同时差价相对合理的产品不是一件容易的事情。该平台提供两种免费的辅助方案,一种是由行业专家向消费者提供产品建议,另一种是由算法提供产品建议。被试需要从上述两种方案中选择一种更为信任的方案。

最后,被试回答了如果由他们自己来选择购买哪一种咖啡豆,这一决策的感知困难程度(1 = 非常简单,7 = 非常困难)。被试还回答了人口统计方面的问题。

2.4 实验结果

首先进行操纵检查。通过方差分析发现,“共赢”组的被试($M_{win} = 4.19, SD = 1.22$),相较于“牺牲”组($M_{sacrifice} = 4.72, SD = 1.21$),感知功利价值与道德价值之间冲突更低, $F(1, 197) = 9.47, p = 0.002$ 。这一结果表明关于叙述框架的操纵是成功的。

本实验将“共赢”组编码为1,“牺牲”组编码为0,将选择算法推荐编码为1,选择人类专家推荐编码为0。以道德消费叙述框架为自变量,以选择结果为因变量进行二项 logistic 回归分析。结果显示,道德消费叙述框架对选择结果有显著正向影响($b = 0.63, Wald's \chi^2 = 4.84, p = 0.028$)。进一步,卡方分析显示相较于“牺牲”组(43.43%)的被试,“共赢”组(59.00%)的被试更多选择算法推荐(vs.人类专家推荐)($\chi^2 = 4.22, p = 0.040$)。

控制变量分析。通过方差分析发现,“共赢”组的被试感知决策困难程度与“牺牲”组无显著差异, $F(1, 197) = 0.61, p = 0.435$ 。

2.5 结果讨论

本实验操纵了道德消费叙述框架,随后比较被试对算法推荐和人类专家推荐的相对偏好,结果表明,当消费者受到“共赢”(vs.“牺牲”)框架的影响时,他们在道德消费中选择算法推荐的比例有所提升。这一结果符合我们的预测,即在“共赢”(vs.“牺牲”)框架的影响下,消费者会增加对算法推荐信任,为本研究 H1 提供了初步证据。实验结果也反映出,“共赢”(vs.“牺牲”)框架并未影响消费者的决策困难程度感知,排除了这一因素影响的可能性。进一步地,我们通过增加算法和人类专家推荐的组间设计,来检查“共赢”(vs.“牺牲”)框架对消费者对算法及人类专家推荐信任的分别影响。

3 实验 2: 功利论道德价值观念的中介效应

3.1 实验目的

本实验进一步检验本研究的主效应和中介效应,具体而言,本研究预测道德消费“共赢”(vs.“牺牲”)叙述框架增强了消费者对功利论道德价值观念的接受度,进而导致道德产品消费者增加了对算法推荐的信任。而道德消费叙述框架并不会影响道德产品消费者对人类专家推荐的信任程度。

3.2 实验设计与被试

本实验采用 2(道德消费叙述框架类型:“共赢”vs.“牺牲”) $\times$ 2(推荐者:算法 vs.人类专家)组间设计。因变量为消费者对推荐者的信任。通过 G*power 软件计算出在显著性水平为 0.05 且效应量为中等偏小水平($f = 0.17$)时,预测达到 90%的统计力水平的总样本量至少为 366 名。本实验从见数(Credemo)平台招募了 408 个被试参与线上实验,被试平均年龄 28.93 岁($SD = 9.02$),其中 34.31%为男性。被试被随机分配到 4 个实验组中。

3.3 实验程序

本实验采用了与实验 1 相同的情境材料。在被试了解完咖啡豆网购平台新上市了“公平交易认证”咖啡豆,并接受了“共赢”或者“牺牲”叙述框架的引导信息之后,被试完成自变量操纵检查问题($\alpha = 0.85$)。随后引导材料告知被试认证咖啡豆和普通咖啡豆差价幅度很广,平台向消费者提供一种选择辅助方案,由算法或者人类专家给出建议。被试需要回答是否信任以上算法或者专家给出的建议(“信任上述算法/专家对产品道德溢价的判断”,“信任上述算法/专家推荐的产品在道德上是经得起考验的”,“信任上述算法/专家对口味、质量、价格、道德的权衡”,“信任上述算法/专家推荐的道德产品是最合适的选择”,1 = 非常不同意,7 = 非常同意, $\alpha = 0.86$)。

接下来,实验将测量被试在本次咖啡豆购买中被试对功利论道德观念的接受度。具体来说,被试需要回答如果是自己选择购买一款咖啡豆,以下选择原因的重要性程度如何,一共有 8 个问项,前 4 个问项反映的是道义论道德价值观念接受度,后 4 个问项反映的是功利论道德价值观念接受度(Tanner et al., 2008)。具体内容如下:“选择它是因为这一选择符合人们必须遵循的原则”,“选择它是因为我有道德义务做出这样的选择”,“选择它是因为

这一选项是绝对符合公平贸易原则的, 不管带来什么利益结果”, “选择它是因为其他选择是存在道德问题的”, “选择它是因为成本效益分析得出这一产品是更优的”, “选择它是因为这一产品带来的物质利益结果可以证明其合理性”, “这一产品能产生最好的物质上的净收益”, “这一产品带来的物质上的好处多于坏处” (1 = 非常不重要, 7 = 非常重要, 功利量表 $\alpha = 0.76$, 道义量表 $\alpha = 0.78$)。

最后, 被试回答了如果由他们自己来选择购买哪一种咖啡豆(认证或者非认证的), 这一决策的感知困难程度(1 = 非常简单, 7 = 非常困难)。被试还回答了人口统计方面的问题。

3.4 实验结果

首先进行道德消费叙述框架的操纵检查。方差分析结果显示, 与我们的预期一致, 道德消费叙述框架对被试感知功利价值与道德价值之间冲突有显著影响, $F(1, 404) = 9.39, p = 0.002, \eta_p^2 = 0.023$; 推荐者类型对感知功利价值与道德价值之间冲突无显著影响, $F(1, 404) = 0.41, p = 0.524$; 两者的交互作用不显著, $F(1, 404) = 0.01, p = 0.914$ 。“共赢”组($M_{Win} = 4.36, SD = 1.36$)相较于“牺牲”组($M_{Sacrifice} = 3.95, SD = 1.30$), 被试的感知功利价值与道德价值之间冲突更低。这表明对道德消费叙述框架的操纵是成功的。

接下来检验各组在因变量上的差异。通过方差分析发现, 道德消费叙述框架和推荐者类型对被试对算法道德推荐的信任程度的交互作用显著, $F(1, 404) = 11.22, p = 0.001, \eta_p^2 = 0.027$ 。此外道德消费叙述框架的主效应也显著, $F(1, 404) = 5.56, p = 0.016, \eta_p^2 = 0.014$ 。简单效应分析发现, “共赢”组的被试($M_{AI-Win} = 5.22, SD = 0.96$), 相较于“牺牲”组的被试($M_{AI-Sacrifice} = 4.68, SD = 1.00$), 对算法推荐的信任程度更高, $F(1, 404) = 16.48, p < 0.001, \eta_p^2 = 0.039$; 道

德消费叙述框架是“共赢”($M_{Human-Win} = 4.91, SD = 0.98$), 还是“牺牲”($M_{Human-Sacrifice} = 5.00, SD = 0.82$), 未显著影响被试对专家推荐的信任程度, $F(1, 404) = 0.43, p = 0.510$ 。道德消费叙述框架对信任程度的影响只在算法推荐的情况下出现。

然后, 通过方差分析检验道德消费叙述框架和推荐者类型对被试功利论道德价值观念接受度的影响。结果发现, 道德消费叙述框架对功利论观念接受度有显著影响, $F(1, 404) = 50.81, p < 0.001, \eta_p^2 = 0.112$; 推荐者类型的影响显著, $F(1, 404) = 4.12, p = 0.043$; 道德消费叙述框架与推荐者类型的交互作用不显著, $F(1, 404) = 0.14, p = 0.707$ 。具体而言, “共赢”组被试的功利论观念接受度($M_{Win} = 5.50, SD = 0.80$)显著高于“牺牲”组被试($M_{Sacrifice} = 4.93, SD = 0.81$)。道德消费叙述框架($F(1, 404) = 0.46, p = 0.499$), 推荐者类型($F(1, 404) = 0.01, p = 0.969$), 以及两者的交互项($F(1, 404) = 0.02, p = 0.884$)对道义论道德价值观念接受度无显著影响。详见图 1。

最后, 通过 Bootstrap 方法检验功利论道德评价倾向的中介作用。检验使用 PROCESS 中的 model 14, 以道德消费叙述框架(“共赢”组编码为 1, “牺牲”组编码为 0)为自变量, 功利论道德价值观念接受度为中介变量, 推荐者类型为调节变量(算法编码为 1, 人类专家编码为 0), 信任程度为因变量进行分析, Bootstrap 再抽样设定为 5000 次。结果显示, 道德消费叙述框架对功利论道德价值观念接受度有显著正向影响($b = 0.57, 95\%$ 置信区间 CI: [0.41, 0.72]), 功利论道德价值观念接受度与推荐者的交互项对信任有显著的正向影响($b = 0.31, 95\%$ 置信区间 CI: [0.09, 0.52])。自变量道德消费叙述框架对因变量信任程度的条件间接效应, 在推荐者为算法时是显著的($b = 0.21, 95\%$ 置信区间 CI: [0.10, 0.33]), 在推荐者为人类专家时是不显著的($b = 0.03, 95\%$

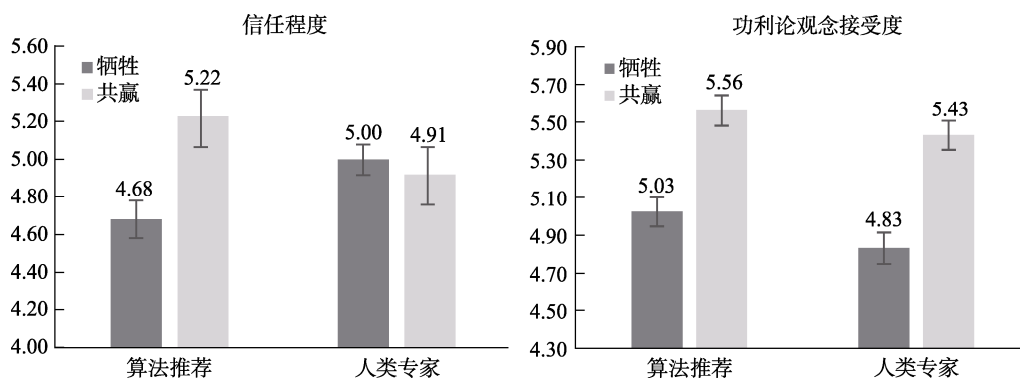


图 1 实验 2 道德消费信念与推荐者类型对信任程度和功利论道德价值观念接受度的交互效应

置信区间 CI: $[-0.06, 0.12]$)。自变量道德消费叙述框架对因变量信任程度的直接效应不显著($b = 0.11$, 95%置信区间 CI: $[-0.09, 0.30]$), 调节的中介效应指标显著($Index = 0.17$, 95%置信区间 CI: $[0.05, 0.32]$)。进一步将道义论道德价值观念接受度作为协变量加入模型, 以上结果依旧稳健, 调节的中介效应指标显著($Index = 0.15$, 95%置信区间 CI: $[0.03, 0.29]$)。这说明道德消费叙述框架对被试信任的影响只在推荐者为算法的时候存在, 功利论道德价值观念接受度中介了这一影响, H2 得到支持。

控制变量分析。通过方差分析检验道德消费叙述框架和推荐者类型对感知决策困难的影响。结果发现, 道德消费叙述框架($F(1, 404) = 0.20, p = 0.653$), 推荐者类型($F(1, 404) = 0.12, p = 0.731$), 以及道德消费叙述框架与推荐者类型的交互作用($F(1, 404) = 0.15, p = 0.698$)对感知决策困难无显著影响。

3.5 结果讨论

本实验通过分别检查道德消费叙述框架如何影响消费者对算法和人类专家信任程度, 发现道德消费叙述框架影响了消费者对算法推荐的信任, 而对人类专家的信任没有显著影响。在“共赢”(vs.“牺牲”)框架的影响下, 消费者更信任算法推荐, 不再对算法能否完成道德决策存在质疑。进一步的中介检查发现, 功利论道德价值观念的中介作用显著。这表明在“共赢”(vs.“牺牲”)框架的影响下, 消费者转变了对道德消费决策适用的决策方法论的认识, 更加认同以功利论的方式来进行决策, 从而改变了对算法推荐的想法。而对于人类专家的推荐, “共赢”(vs.“牺牲”)框架未改变他们对人类专家是否胜任道德消费推荐的感知, 因为人类专家与算法不同, 不具有在道义判断上的短板。

4 实验 3: 算法增强型推荐的调节作用

4.1 实验目的

我们之前的研究表明, “共赢”(vs.“牺牲”)叙述框架让消费者更多地认为道德消费中的道德价值与功利价值是协同的, 进而让消费者更能接受在道德消费的权衡中采用功利性的道德评价。然而, 在算法增强型建议的情况下, 决策任务被分解, 由算法进行功利价值的评价, 在此基础上再由人类专家根据道德原则得出最终的推荐, 那么道德消费叙述框架对消费者算法态度的影响将可能会消失。

4.2 实验设计和被试

本实验采用 2 (道德消费叙述框架类型: “共赢” vs. “牺牲”) $\times$ 3 (推荐方式: 算法替代型推荐 vs. 算法增强型推荐 vs. 人类专家推荐) 组间设计。因变量为消费者对道德消费中的产品推荐的信任。通过 G*power 软件计算出在显著性水平为 0.05 且效应量为中等偏小水平($f = 0.16$)时, 预测达到 90% 的统计力水平的总样本量至少为 498 名。本实验从见数 (Credemo) 平台招募了 541 个被试参与线上实验, 被试平均年龄 28.19 岁($SD = 7.94$), 其中 28.28% 为男性, 他们被随机分配到 6 个实验组中。

4.3 实验程序

本实验将选用电动自行车购买情境。本实验采用与前两个实验类似的程序。首先, 被试需要想象自己需要购买一辆电动自行车代步, 然后发现市场上的新电池产品。引导信息告知被试市场上推出了一种安装新技术电池的电动自行车。新技术的特点是让车辆在制造, 使用和回收的全流程中的碳排放降低 20%~45%, 但是生产成本会较传统电池更高。

随后被试阅读不同道德消费叙述框架的海报。在两类叙述框架组中, 海报都再次强调了新电池能够减少电动车的整体碳排放。“牺牲”组的宣传语强调的是每个人都要为环保出行付出, 实现低碳目标, 担当环保责任。“共赢”组的宣传语则强调环保出行能让每个人生活更好, 实现低碳目标, 能够共享美好环境。接着被试完成操纵检查问题, 包括“支付新型电池的溢价就意味着消费者会有一定物质利益损失”, “支付新型电池的溢价, 消费者就会得到更少实惠”, “从物质利益的角度来说, 支付新型电池的溢价是难以做到物有所值的”(1 = 非常不同意, 7 = 非常同意, $\alpha = 0.77$)。

随后引导材料告知被试不同品牌新型电池电动自行车的设计、价格档次、参数、使用体验以及碳减排效果存在许多差异。碳减排效果好的产品一般在其他方面相对更弱。算法替代型推荐组被告知现有一款能够分析车型信息和环保责任的算法, 可以为消费者提供消费建议。人类专家推荐组则被告知推荐者是业界专家。算法增强型推荐组被告知建议将由人类专家根据算法对车型信息的评估结果来做出。被试需要回答是否信任以上来自算法/专家的建议(“信任上述建议对产品环保溢价的判断”, “信任上述建议推荐的产品在环保上是经得起考验的”, “信任上述建议对性能、价格、环保的权衡”, “信任上述建议推荐的环保产品是最合适的选择”, 1 =

非常不同意, 7 = 非常同意, $\alpha = 0.82$)。

接下来, 实验将测量被试在本次电动自行车购买中的道德价值观念接受度, 测量问项与实验 2 一致(功利论量表 $\alpha = 0.84$, 道义论量表 $\alpha = 0.78$)。最后消费者回答了感知决策难度和人口统计相关问题。

4.4 实验结果

首先检验实验操纵对被试感知功利价值与道德价值之间冲突的影响。我们以人类专家推荐组为参照组设置两个虚拟变量(算法替代型推荐组虚拟变量: 人类专家推荐组 = 0, 算法替代型推荐组 = 1, 算法增强型推荐组 = 0; 算法增强型推荐组虚拟变量: 人类专家推荐组 = 0, 算法替代型推荐组 = 0, 算法增强型推荐组 = 1), 然后以感知功利价值与道德价值之间冲突为因变量, 以两个虚拟变量和道德消费叙述框架, 以及道德消费叙述框架与两个虚拟变量的交互项为自变量进行回归分析。结果发现, 道德消费叙述框架对感知价值冲突有显著影响($b = -0.69, t = -4.30, p < 0.001$), 算法替代型推荐组($b = -0.12, t = -0.76, p = 0.450$)和算法增强型推荐组($b = -0.07, t = -0.41, p = 0.679$)的两个虚拟变量的回归系数都不显著, 算法替代型推荐组虚拟变量与道德消费叙述框架的交互项($b = 0.25, t = 1.10, p = 0.270$), 以及算法增强型推荐组虚拟变量与道德消费叙述框架的交互项($b = 0.10, t = 0.45, p = 0.655$)也都不显著。以道德消费叙述框架为自变量, 感知价值冲突为因变量做方差分析, 结果显示“共赢”组($M_{Win} = 3.81, SD = 1.06$)的感知价值冲突显著低于“牺牲”组($M_{Sacrifice} = 4.38, SD = 1.07$), $F(1, 539) = 39.18, p < 0.001, \eta_p^2 = 0.068$ 。这表明本实验关于道德消费叙述框架的操纵是成功的。

接下来检验各组在因变量上的差异。以对推荐建议的信任为因变量, 以两个虚拟变量、道德消费叙述框架, 以及道德消费叙述框架与两个虚拟变量的交互项为自变量进行回归分析。结果显示, 算法替代型推荐组虚拟变量与道德消费叙述框架的交互项($b = 0.46, t = 2.48, p = 0.013$)效应显著, 支持了 H3。并且, 算法增强型推荐组虚拟变量与道德消费叙述框架的交互项($b = 0.25, t = 1.40, p = 0.163$)效应不显著, 算法替代型推荐组虚拟变量($b = -0.09, t = -0.69, p = 0.489$)、算法增强型推荐组虚拟变量($b = -0.04, t = -0.32, p = 0.746$)以及道德消费叙述框架($b = -0.12, t = -0.94, p = 0.349$)的主效应不显著。

进一步分析表明, 对于算法替代型推荐, 道德

消费的“共赢”叙述框架($M_{AI-Win} = 5.19, SD = 0.78$), 相较于“牺牲”叙述框架($M_{AI-Sacrifice} = 4.86, SD = 0.86$), 增加了消费者对算法替代型推荐的信任, $F(1, 180) = 7.55, p = 0.007, \eta_p^2 = 0.040$; 对于算法增强型推荐, 道德消费叙述框架是“共赢”($M_{Human-Win} = 5.04, SD = 0.84$), 还是“牺牲”($M_{Human-Sacrifice} = 4.91, SD = 0.94$), 并未显著影响被试对算法增强型推荐的信任程度, $F(1, 180) = 1.05, p = 0.308$; 对于人类专家推荐, 道德消费叙述框架是“共赢”($M_{Human-Win} = 4.83, SD = 0.95$), 还是“牺牲”($M_{Human-Sacrifice} = 4.95, SD = 0.86$), 也未显著影响被试对人类专家推荐的信任程度, $F(1, 175) = 0.81, p = 0.369$ 。

我们也检验了各组在道德价值观念接受度上的差异。回归分析表明, “共赢”(vs. “牺牲”)叙述框架对功利论道德观念接受度有显著影响($b = 1.02, t = 8.61, p < 0.001$)。算法替代型推荐组($b = 0.26, t = 2.21, p = 0.028$)虚拟变量的回归系数显著。算法增强型推荐组($b = 0.11, t = 0.94, p = 0.350$)的虚拟变量的回归系数不显著。算法替代型推荐组虚拟变量与道德消费叙述框架的交互项($b = -0.14, t = -0.84, p = 0.404$), 以及算法增强型推荐组虚拟变量与道德消费叙述框架的交互项($b = -0.06, t = -0.40, p = 0.693$)也都不显著。以道德消费叙述框架为自变量, 功利论道德观念接受度为因变量做方差分析, 结果显示“共赢”组($M_{Win} = 5.56, SD = 0.73$)的功利论道德观念接受度显著高于“牺牲”组($M_{Sacrifice} = 4.61, SD = 0.85$), $F(1, 538) = 195.92, p < 0.001, \eta_p^2 = 0.267$ 。详见图 2。

接下来进行中介分析, 检验使用 PROCESS 中的 model 14, 以道德消费叙述框架为自变量, 功利论道德价值观念接受度为中介变量, 算法替代型推荐组虚拟变量为调节变量, 信任程度为因变量, 道义论道德价值观念接受度和算法增强型推荐组虚拟变量为协变量进行分析, Bootstrap 再抽样设定为 5000 次。结果显示, 道德消费叙述框架对功利论道德价值观念接受度有显著正向影响($b = 0.95, 95\%$ 置信区间 CI: $[0.82, 1.08]$), 功利论道德观念接受度与算法替代型推荐组虚拟变量的交互项对信任有显著的正向影响($b = 0.30, 95\%$ 置信区间 CI: $[0.14, 0.46]$)。自变量道德消费叙述框架对因变量信任程度的条件间接效应, 在推荐者为算法替代型推荐时是显著的($b = 0.33, 95\%$ 置信区间 CI: $[0.17, 0.50]$), 在推荐者为人类专家推荐时是不显著的($b = 0.04,$

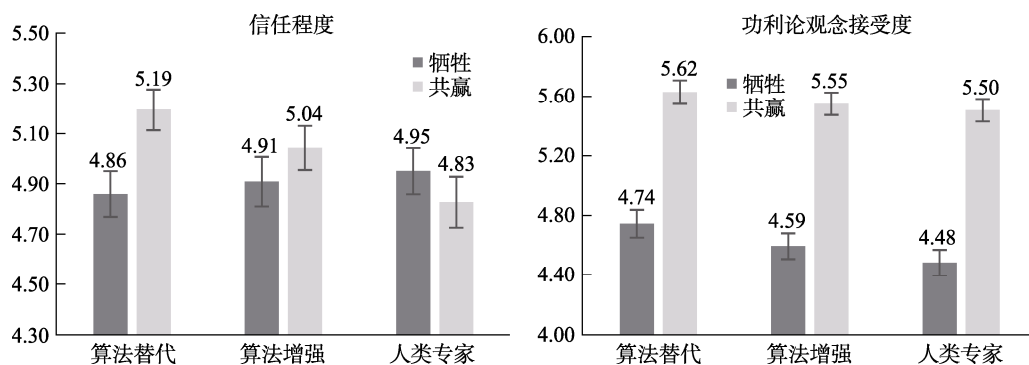


图2 实验3道德消费信念与推荐者类型对信任程度和功利论道德价值观念接受度的交互效应

95%置信区间 CI: [-0.06, 0.15])。自变量道德消费叙述框架对因变量信任程度的直接效应不显著($b = -0.02$, 95%置信区间 CI: [-0.19, 0.14]), 调节的中介效应指标显著($Index = 0.28$, 95%置信区间 CI: [0.11, 0.46])。这说明道德消费叙述框架对被试信任的影响只在算法替代型的情境中存在, 功利论道德观念接受度中介了这一影响。

控制变量分析。以感知决策复杂性为因变量, 两个虚拟变量、道德消费叙述框架, 以及道德消费叙述框架与两个虚拟变量的交互项为自变量进行回归分析。道德消费叙述框架对感知决策复杂性无显著影响($b = -0.28$, $t = -1.17$, $p = 0.245$), 算法替代型推荐组($b = -0.17$, $t = -0.72$, $p = 0.470$)和算法增强型推荐组($b = -0.30$, $t = -1.26$, $p = 0.208$)的两个虚拟变量的回归系数都不显著, 算法替代型推荐组的虚拟变量与道德消费叙述框架的交互项($b = 0.39$, $t = 1.15$, $p = 0.251$), 以及算法增强型推荐组虚拟变量与道德消费叙述框架的交互项($b = 0.52$, $t = 1.56$, $p = 0.121$)也都不显著。

4.5 结果讨论

本实验检验了道德消费叙述框架如何影响消费者对算法替代型、算法增强型和人类专家推荐的信任程度。结果发现, 道德消费叙述框架只对算法替代型推荐的信任有显著影响, 对于算法增强型和人类专家推荐的信任没有显著影响。对于中介变量, “共赢”(vs. “牺牲”)叙述框架提升了消费者对功利论道德观念的接受度。而只有在算法替代型推荐情况下, 功利论道德价值观念接受度起到了影响消费者信任的作用, 此时中介效应成立。对于另外两种情况, 因为人类在决策中的参与程度都是较高的, 消费者本身不会顾虑算法增强推荐和人类专家推荐在道义判断上存在短板, 所以消费者对两者的信任没有受到“共赢”(vs. “牺牲”)叙述框架的影响。以

上结果验证了 H1、H2 和 H3。

5 总讨论

5.1 结果总结

本研究提出道德消费叙述框架会影响消费者对道德消费中算法产品推荐的信任。受到“共赢”(vs. “牺牲”)叙述框架的影响, 消费者在道德消费情境中对功利论式道德价值观念的接受度会提升, 进而他们对算法推荐的信任程度会增加。并且叙述框架对算法推荐信任的积极作用只出现在算法替代型推荐的情况下, 对算法增强型推荐的信任无显著影响。本研究通过三个实验检验了以上理论推断。实验1的结果发现受道德消费的“共赢”(vs. “牺牲”)叙述框架影响的消费者更多地在算法推荐和人类推荐的二选一问题中选择了算法推荐。在实验2中, 推荐类型和叙述框架都被设定为组间因素, 结果显示“共赢”(vs. “牺牲”)叙述框架增强了消费者对算法推荐的信任, 而消费者对人类专家的信任没有受到显著影响。消费者在道德消费决策中对功利论道德价值观念接受度起到了中介作用。在实验3中, 我们增加了算法增强型推荐情境, 结果显示“共赢”(vs. “牺牲”)叙述框架对算法增强型推荐信任没有显著影响。本实验另外两组结果与实验2相似, 再次验证了“共赢”(vs. “牺牲”)叙述框架对算法替代型推荐信任的积极影响, 以及功利论道德观念接受度的中介作用。

5.2 理论贡献与实践意义

首先, 本研究通过关注道德产品消费这一特殊情境, 扩展了算法在道德决策中应用的研究范围。前人研究虽然探讨道德伦理相关的算法决策应用问题(Bigman & Gray, 2018; Dietvorst & Bartels, 2022, 许丽颖 等, 2022), 但这些研究往往以伦理道德属性作为核心且唯一的研究焦点, 例如法律裁

决中的道德公正性、医疗保险范围的道德合理性、算法决策引发的歧视等。这些研究主要强调道德正确原则对道德决策的影响,因而贬低了算法在道德决策领域的价值。然而,在本文关注的道德产品消费情境中,产品的道德属性和其功利性价值都是消费者决策中需要考虑的内容。本文将算法决策应用从单纯关注道德价值情境扩展到兼顾道德价值与功利价值的道德消费决策领域,为提升算法在道德决策领域的应用价值提供了新思路。研究表明,引导消费者理解道德价值与功利价值的不同关系,会改变消费者对算法进行道德相关决策的信任程度。

第二,本研究探讨了道德消费叙述框架对道德消费决策中算法信任的影响,丰富了框架效应理论文献。先前的框架效应文献发现叙述框架能改变人们对特定情境中关系和目标的认知构建,本研究发现,使用“共赢”(vs.“牺牲”)叙述框架能够形塑或改变消费者对道德价值的理解,使消费者更可能接受功利论道德观念,这为框架效应在道德观念认知方面的应用提供了支持(Gai & Klesse, 2019)。

第三,本研究从对道德价值的认知差异角度,探讨了道德观念因素对消费者算法信任的影响作用,这对于道德决策领域中的算法信任研究做出了重要补充。先前有关算法不信任的研究主要从算法决策缺陷的视角来探讨导致消费者算法不信任的原因及解决方案(Dietvorst & Bartels, 2022)。本研究则从道德价值观念的视角来探讨如何通过影响和改变消费者认知来促进其对算法决策的信任。本研究发现强调共同利益最大化的功利论观念比传统的强调道德原则正确的道义论观念更能够为算法参与道德决策提供合理理由,这一发现也展现了基于世俗价值的功利论道德观念在消费领域中的独特价值,丰富了道德价值观念领域的研究。

第四,本研究探讨了在不同类型的算法决策(算法增强型 vs. 算法替代型)中,使用“共赢”(vs.“牺牲”)叙述框架对消费算法信任的影响差异。过往研究已经在一些情境中检验了消费者对增强型和替代型算法决策的信任之间的差异(Baird & Maruping, 2021)。本研究则在此基础上,关注道德消费这一涉及道德因素的决策情境,为解决消费者对算法替代型决策的不信任问题寻找到了解决方案。这一发现丰富了关于不同类型的算法决策之间差异的研究。

本研究具有较显著的实践意义,具体而言,对于采用算法决策工具推荐道德产品的商家,在进行

道德消费劝导时可以更多地使用“共赢”视角的叙述框架,激发人们的功利论道德价值观念,这样会让消费者更信任由算法给出的决策推荐。此外,企业也可以考虑将算法作为一种增强人类决策的方式,让消费者更愿意相信算法这种低成本、快速高效的决策辅助工具,促进消费者选择合适的道德产品。

参 考 文 献

- Baird, A., & Maruping, L. M. (2021). The next generation of research on IS use: A theoretical framework of delegation to and from agentic IS artifacts. *MIS Quarterly*, 45(1), 315–341.
- Bigman, Y. E., & Gray, K. (2018). People are averse to machines making moral decisions. *Cognition*, 181, 21–34.
- Brynjolfsson, E., & Mitchell, T. (2017). What can machine learning do? Workforce implications: Profound change is coming, but roles for humans remain. *Science*, 358(6370), 1530–1534.
- Burton, J. W., Stein, M. K., & Jensen, T. B. (2020). A systematic review of algorithm aversion in augmented decision making. *Journal of Behavioral Decision Making*, 33(2), 220–239.
- Castelo, N., Bos, M. W., & Lehmann, D. R. (2019). Task-dependent algorithm aversion. *Journal of Marketing Research*, 56(5), 809–825.
- Chernev, A., Bockenholt, U., & Goodman, J. (2015). Choice overload: A conceptual review and meta-analysis. *Journal of Consumer Psychology*, 25(2), 333–358.
- Commerford, B. P., Dennis, S. A., Joe, J. R., & Ulla, J. W. (2022). Man versus machine: Complex estimates and auditor reliance on artificial intelligence. *Journal of Accounting Research*, 60(1), 171–201.
- Conway, P., & Gawronski, B. (2013). Deontological and utilitarian inclinations in moral decision making: A process dissociation approach. *Journal of Personality and Social Psychology*, 104(2), 216–235.
- Davenport, T., Guha, A., Grewal, D., & Bressgott, T. (2020). How artificial intelligence will change the future of marketing. *Journal of the Academy of Marketing Science*, 48(1), 24–42.
- Dawes, R. M. (1979). The robust beauty of improper linear models in decision making. *American Psychologist*, 34(7), 571–582.
- Dietvorst, B. J., & Bartels, D. M. (2022). Consumers object to algorithms making morally relevant tradeoffs because of algorithms' consequentialist decision strategies. *Journal of Consumer Psychology*, 32(3), 406–424.
- Gai, P. J., & Klesse, A. K. (2019). Making recommendations more effective through framings: Impacts of user-versus item-based framings on recommendation click-throughs. *Journal of Marketing*, 83(6), 61–75.
- Gao, H., Mittal, V., & Zhang, Y. (2020). The differential effect of local-global identity among males and females: The case of price sensitivity. *Journal of Marketing Research*, 57(1), 173–191.
- Gao, H., Zhang, Y., & Mittal, V. (2017). How does local-global identity affect price sensitivity? *Journal of Marketing*, 81(3), 62–79.
- Gray, H. M., Gray, K., & Wegner, D. M. (2007). Dimensions

- of mind perception. *Science*, 315(5812), 619–619.
- Gray, K., & Schein, C. (2012). Two minds vs. two philosophies: Mind perception defines morality and dissolves the debate between deontology and utilitarianism. *Review of Philosophy and Psychology*, 3(3), 405–423.
- Green, T., & Peloza, J. (2014). Finding the right shade of green: The effect of advertising appeal type on environmentally friendly consumption. *Journal of Advertising*, 43(2), 128–141.
- Gustafson, A. (2013). In defense of a utilitarian business ethic. *Business and Society Review*, 118(3), 325–360.
- Haenlein, M., & Kaplan, A. (2019). A brief history of artificial intelligence: On the past, present, and future of artificial intelligence. *California Management Review*, 61(4), 5–14.
- Hassan, L., Shaw, D., Shiu, E., Walsh, G., & Parry, S. (2013). Uncertainty in ethical consumer choice: A conceptual model. *Journal of Consumer Behaviour*, 12(3), 182–193.
- Herzenstein, M., Dholakia, U. M., & Sonenshein, S. (2020). How the number of options affects prosocial choice. *International Journal of Research in Marketing*, 37(2), 356–370.
- Hiller, A., & Woodall, T. (2019). Everything flows: A pragmatist perspective of trade-offs and value in ethical consumption. *Journal of Business Ethics*, 157(4), 893–912.
- Huang, S.-C., & Aaker, J. (2019). It's the journey, not the destination: How metaphor drives growth after goal attainment. *Journal of Personality and Social Psychology*, 117(4), 697–720.
- Hwang, K., & Kim, H. (2018). Are ethical consumers happy? Effects of ethical consumers' motivations based on empathy versus self-orientation on their happiness. *Journal of Business Ethics*, 151(2), 579–598.
- Jago, A. S. (2019). Algorithms and Authenticity. *Academy of Management Discoveries*, 5(1), 38–56.
- Jebari, K., & Lundborg, J. (2021). Artificial superintelligence and its limits: Why AlphaZero cannot become a general agent. *AI & Society*, 36(1), 807–815.
- Leachman, S. A., & Merlino, G. (2017). Medicine: The final frontier in cancer diagnosis. *Nature*, 542(7639), 36–38.
- Lee, S. W. S., & Schwarz, N. (2014). Framing love: When it hurts to think we were made for each other. *Journal of Experimental Social Psychology*, 54, 61–67.
- Lim, J. S., & O'Connor, M. (1996). Judgmental forecasting with interactive forecasting support systems. *Decision Support Systems*, 16(4), 339–357.
- Logg, J. M., Minson, J. A., & Moore, D. A. (2019). Algorithm appreciation: People prefer algorithmic to human judgment. *Organizational Behavior and Human Decision Processes*, 151, 90–103.
- Longoni, C., Bonezzi, A., & Morewedge, C. K. (2019). Resistance to medical artificial intelligence. *Journal of Consumer Research*, 46(4), 629–650.
- Longoni, C., & Cian, L. (2022). Artificial intelligence in utilitarian vs. hedonic contexts: The “word-of-machine” effect. *Journal of Marketing*, 86(1), 91–108.
- Luchs, M. G., Naylor, R. W., Irwin, J. R., & Raghunathan, R. (2010). The sustainability liability: Potential negative effects of ethicality on product preference. *Journal of Marketing*, 74(5), 18–31.
- Luo, X., Qin, M. S., Fang, Z., & Qu, Z. (2021). Artificial intelligence coaches for sales agents: Caveats and solutions. *Journal of Marketing*, 85(2), 14–32.
- Luo, X., Tong, S., Fang, Z., & Qu, Z. (2019). Machines vs. humans: The impact of artificial intelligence chatbot disclosure on customer purchases. *Marketing Science*, 38(6), 937–947.
- Martin, K. (2019). Ethical implications and accountability of algorithms. *Journal of Business Ethics*, 160(4), 835–850.
- Mencel, J., & May, D. R. (2009). The effects of proximity and empathy on ethical decision-making: An exploratory investigation. *Journal of Business Ethics*, 85(2), 201–226.
- Miller, D. T., & Prentice, D. A. (2016). Changing norms to change behavior. *Annual Review of Psychology*, 67, 339–361.
- Niven, K., & Healy, C. (2016). Susceptibility to the ‘dark side’ of goal-setting: Does moral justification influence the effect of goals on unethical behavior? *Journal of Business Ethics*, 137(1), 115–127.
- Puntoni, S., Reczek, R. W., Giesler, M., & Botti, S. (2021). Consumers and artificial intelligence: An experiential perspective. *Journal of Marketing*, 85(1), 131–151.
- Reczek, R. W., Irwin, J. R., Zane, D. M., & Ehrich, K. R. (2018). That's not how I remember it: Willfully ignorant memory for ethical product attribute information. *Journal of Consumer Research*, 45(1), 185–207.
- Ruttan, R. L., & Nordgren, L. F. (2021). Instrumental use erodes sacred values. *Journal of Personality and Social Psychology*, 121(6), 1223–1240.
- Ryoo, Y., Sung, Y., & Chechelnytska, I. (2020). What makes materialistic consumers more ethical? Self-benefit vs. other-benefit appeals. *Journal of Business Research*, 110(March), 173–183.
- Sampson, S. E. (2021). A strategic framework for task automation in professional services. *Journal of Service Research*, 24(1), 122–140.
- Sher, S., & McKenzie, C. R. M. (2006). Information leakage from logically equivalent frames. *Cognition*, 101(3), 467–494.
- Tanner, C., Medin, D. L., & Iliev, R. (2008). Influence of deontological versus consequentialist orientations on act choices and framing effects: When principles are more important than consequences. *European Journal of Social Psychology*, 38, 757–769.
- Tezer, A., & Bodur, H. O. (2020). The greenconsumption effect: How using green products improves consumption experience. *Journal of Consumer Research*, 47(1), 25–39.
- Tong, S., Jia, N., Luo, X., & Fang, Z. (2021). The Janus face of artificial intelligence feedback: Deployment versus disclosure effects on employee performance. *Strategic Management Journal*, 42(9), 1600–1631.
- Tong, Z., Zhou, Y., & Zhou, L. (2019). The impact of cross-border charity on domestic and foreign consumers brand evaluation: The moderating effect of national power. *Nankai Business Review*, 22(2), 14–22.
- [童泽林, 周元元, 周玲. (2019). 跨国慈善对东道国和来源国消费者品牌评价的影响: 品牌国家实力视角. *南开管理评论*, 22(2), 14–22.]
- Trudel, R., Klein, J., Sen, S., & Dawar, N. (2020). Feeling good by doing good: A selfish motivation for ethical choice. *Journal of Business Ethics*, 166(1), 39–49.
- Wei, X., Yu, F., Peng, K., & Zhong, N. (2023). Psychological richness increases behavioral intention to protect the environment. *Acta Psychologica Sinica*, 55(8), 1330–1343.
- [魏心妮, 喻丰, 彭凯平, 钟年. (2023). 心理丰富提高亲环境行为意愿. *心理学报*, 55(8), 1330–1343.]
- Xu, L., Yu, F., & Peng, K. (2022). Algorithmic discrimination causes less desire for moral punishment than human discrimination. *Acta Psychologica Sinica*, 54(9), 1076–1092.
- [许丽颖, 喻丰, 彭凯平. (2022). 算法歧视比人类歧视引起更少道德惩罚欲. *心理学报*, 54(9), 1076–1092.]

- Xu, Z. X., & Ma, H. K. (2016). How can a deontological decision lead to moral behavior? The moderating role of moral identity. *Journal of Business Ethics*, 137(3), 537–549.
- Xin, Z., Liu, F., & Mu, H. (2022). The instrumental use of sacred values in advertising and its influence. *Journalism & Communication*, 8(2), 97–110.
- [辛自强, 刘菲, 穆昊阳. (2022). 广告中“神圣价值观”工具性使用及其影响. *新闻与传播研究*, 8(2), 97–110.]
- Yamoah, F. A., Duffy, R., Petrovici, D., & Fearn, A. (2016). Towards a framework for understanding fairtrade purchase intention in the mainstream environment of supermarkets. *Journal of Business Ethics*, 136(1), 181–197.
- Zollo, L. (2021). The consumers' emotional dog learns to persuade its rational tail: Toward a social intuitionist framework of ethical consumption. *Journal of Business Ethics*, 168(2), 295–313.

“Win-win” vs. “sacrifice”: Impact of framing of ethical consumption on trust in algorithmic recommendation

XU Lan¹, CHEN Quan¹, CUI Nan², GU Hong³

(¹ School of Economics and Management, Wuhan University, Wuhan 430072, China)

(² Research Center for Organizational Marketing of Wuhan University, Wuhan 430072, China)

(³ College of Urban and Environmental Sciences, Central China Normal University, Wuhan 430079, China)

Abstract

In making ethical consumption decisions, consumers need to consider the instrumental and ethical attributes of products and make tradeoffs. This complexity is a critical barrier between consumers' ethical consumption intention and actual ethical consumption behavior. With the development of artificial intelligence, algorithms are increasingly being used to provide advice to consumers, which can reduce the difficulty of decision-making. However, studies have found that people are reluctant to adapt algorithms to make decisions in ethical trade-offs.

This study aims to provide a potential solution for increasing consumers' trust in algorithmic recommendations in ethical consumption contexts. In particular, this research proposes that marketers can influence consumers' trust in algorithmic recommendations by changing the narrative framing strategy for ethical consumption. When ethical consumption is described as “win-win” (vs. “sacrifice”), consumers are more likely to trust the algorithm's recommendations for ethical consumption decisions. The sacrifice narrative framing strategy emphasizes that consumers need to make personal sacrifices to help the environment or society in ethical consumption. The win-win narrative framing strategy emphasizes that there is no conflict of utilitarian interests between consumers and other stakeholders in ethical consumption. Moreover, this strategy encourages consumers to consider that ethical consumption is for the greater good of everyone. Therefore, the win-win (vs. sacrifice) narrative framing information will enhance consumers' belief that achieving moral values is compatible with optimizing the overall utilitarian values of the community in ethical consumption situations. Thus, consumers are more likely to trust algorithmic tools that maximize utilitarian values to solve ethical consumption decision problems.

By portraying ethical consumption as a win-win, consumers perceive themselves as part of a common interest community with relevant stakeholders, helping reinforce utilitarianism-based moral beliefs. Utilitarianism believes that pursuing the maximization of common interest is a rational way to achieve moral correctness. Accordingly, activating the win-win (vs. sacrifice) narrative framing encourages consumers to identify more with utilitarianism in ethical consumption decisions, thereby increasing their trust in algorithmic recommendations in ethical contexts.

Through three experiments, this study examines the impact of ethical consumption narrative framing strategies (win-win vs. sacrifice) on trust in algorithmic consumption recommendations. Experiment 1 first tests how activating the win-win (vs. sacrifice) narrative framing affects consumers' preference for the source of consumption recommendations. Results show that the win-win (vs. sacrifice) narrative framing significantly increases consumers' inclination to choose algorithmic consumption recommendations. Experiment 2 further examines the influence of activating the win-win (vs. sacrifice) narrative framing on trust in consumption

recommendations and the mediating role of utilitarian moral values. Results indicate that the win-win (vs. sacrifice) narrative framing enhances consumers' trust in algorithmic consumption recommendations but does not significantly impact trust in recommendations from human experts. Consumers' acceptance of utilitarianism in ethical decisions mediates the preceding positive effects. Lastly, this study examines the implementation boundaries of the win-win (vs. sacrifice) narrative framing strategy and finds that it only works for consumption recommendations from substitutive algorithms and has no effect on algorithm-augmented consumption recommendations from human experts.

Keywords ethical consumption, algorithmic recommendation, deontology, utilitarianism, framing effect