

面向智能驾驶控制感的交互设计影响要素研究

计越 宫晓东*

(北京理工大学设计与艺术学院, 北京 100081)

摘要:智能驾驶技术为构建更加智慧、安全、高效的交通系统提供了可能,但技术能否惠及更广泛群体还受到用户接受度的影响。通过对技术接受模型及相关理论的分析,指出控制感是影响用户技术接受度的关键因素之一,而智能驾驶场景下,驾驶者角色由“操控者”转变为“监督者”,控制权的让渡势必会对个体控制感产生负面影响。因此,通过探究智能驾驶情境下影响用户控制感的因素以提升智能驾驶技术的接受度。研究通过整合相关经典理论框架,将控制感解构为可预测与可归因两个中介变量,分别引入可解释性、透明度、一致性和充分性四个外源变量,构建了针对智能驾驶情境的控制感结构方程模型(SEM)。收集了358份有效问卷,采用路径分析对模型进行验证。结果表明:可预测与可归因确为影响控制感的关键中介变量;其中,可预测显著受可解释性与透明度影响,而可归因则显著受一致性与充分性影响,研究假设基本得到验证。该研究揭示了智能驾驶情境下影响个体控制感的核心因素,为通过交互设计干预智能驾驶控制感、推动智能驾驶技术的广泛应用提供了理论依据。

关键词:智能驾驶;控制感;可预测;可归因

中图分类号:B848

文献标志码:A

文章编号:1003-5184(2026)01-0013-10

1 引言

控制感(Perceived Control)是指个体在面对环境和自身行为时,相信自己能够影响事件进程以实现期望结果的主观信念(Haggard & Chambon, 2012; 孟可强等, 2023)。作为一种与客观控制相对应的心理体验,它反映了个体对生活情境、行为结果及情绪状态的掌控认知与信念(Berberian, 2019)。研究指出,维持较高水平的控制感是人类基本的生存需求(Shapiro et al., 1998)。此外,控制感也是决定用户是否采纳和使用新技术的关键变量。这一观点在技术接受理论体系中的计划行为理论(Theory of Planned Behavior, TPB)、解构计划行为理论(Deconstructing the Theory of Planned Behavior, DTPB)以及技术接受模型3(Technology Acceptance Model 3, TAM3)等经典理论模型中均得到了充分强调(Ajzen, 1991; Taylor & Todd, 1995; Venkatesh & Bala, 2008)。因此,维持适度的控制感对于保障用户的心理舒适度及促进技术接受至关重要。

在智能驾驶情境下,驾驶员的角色发生了根本性转变,即从传统的“主动操控者”转变为系统的“被动监督者”。这种控制权的让渡导致个体行为与系统结果之间的因果关联被解耦,从而显著削弱

了驾驶员的控制感(喻丰等, 2024)。Le Goff等人(2018)指出,控制感的削弱常伴随着焦虑、怀疑和困惑等负面情绪,这些心理障碍构成了用户接受智能驾驶技术的关键壁垒。然而,现有研究多局限于从系统的有用性和易用性视角探讨技术接受度,缺乏对控制感缺失及其缓解机制的系统性考察。因此,深入探究智能驾驶情境下控制感的影响因素,对于提升用户接受度及推动技术落地具有迫切的理论价值与实践意义。

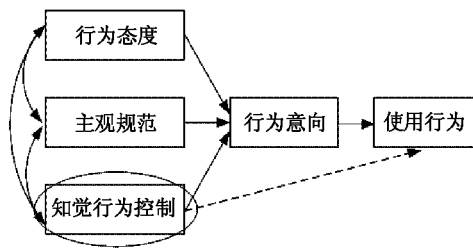
2 相关理论及假设的提出

2.1 控制感是技术接受模型理论体系中的共性要素

技术接受度(Technology Acceptance)是指个体或群体对新技术的接受、采纳和使用的态度与行为,反映了用户对新技术的接纳意愿(Davis, 1989)。为了更深入地剖析影响这种意愿与行为的心理机制, Ajzen(1991)提出的计划行为理论(TPB)模型具有极高的解释力。该模型首次引入了“知觉行为控制”(Perceived Behavioral Control, PBC)这一概念,并将其确立为预测个体行为意向与实际使用行为的关键变量,如图1所示。

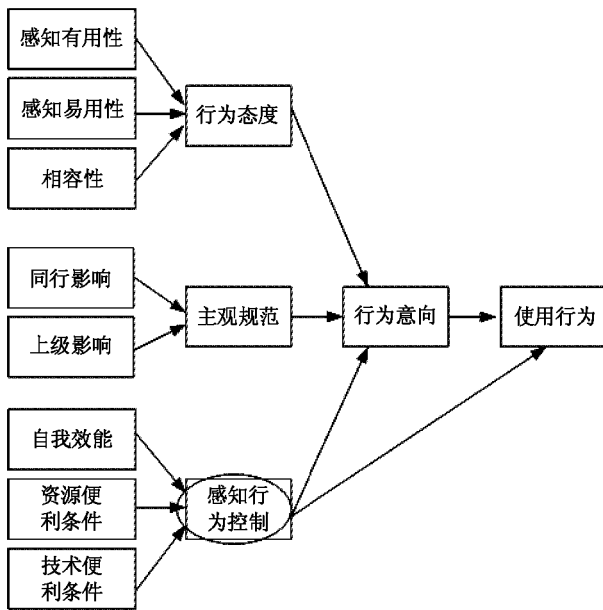
1995年, Taylor和Todd在TPB模型的基础上提出了DTPB理论模型。该模型不仅继续将知觉行

* 通信作者:宫晓东, E-mail: hildagxd@bit.edu.cn.



(根据 Ajzen(1991)的研究整理绘制)

图1 计划行为理论模型



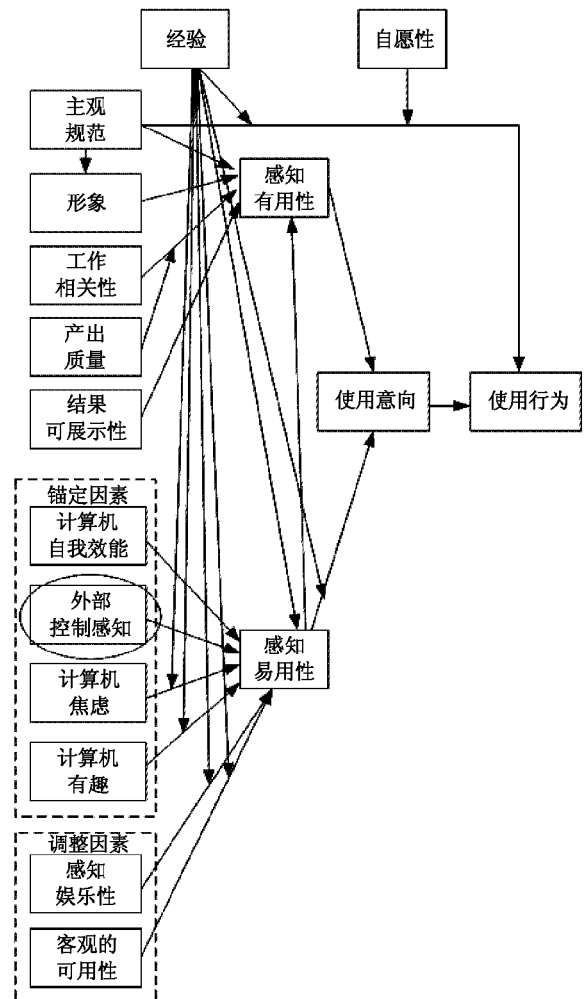
(根据 Taylor & Todd(1995)的研究整理绘制)

图2 解构计划行为理论模型

为控制视为影响个体行为意向和使用行为的关键变量,更进一步对其进行了多维度的解构,将其细分为自我效能和资源便利条件等维度,以提高模型的解释力,如图2所示。

随后,Venkatesh 和 Bala(2008)在技术接受模型的基础上构建了 TAM3 理论模型。该模型将“外部控制感知”确立为影响感知易用性的关键前因变量。这意味着,控制感通过作用于感知易用性,间接影响着用户的使用意愿与实际行为,如图3所示。

综上所述,“控制感”作为技术接受度的影响要素在技术接受理论模型体系中形成共识,成为其中的理论共性要素。从 TPB 的“知觉行为控制”到 TAM3 的“外部控制感知”,这种词汇的转变代表着学界对“控制感”这一概念的认识不断深化,其衡量方式也从最初对自身行为能力的感知,转向对客观条件的一种直觉判断和主观评估。这表明控制感的来源不仅包括个体内部的能力,也延伸至其所处的客观环境和技术支持,这为研究智能驾驶个体接受



(根据 Venkatesh & Bala(2008)的研究整理绘制)

图3 技术接受模型3

度提供了更全面的视角。

智能驾驶技术的核心功能依赖于人工智能算法、传感器数据处理、车联网通信等信息技术,符合技术接受理论体系最初针对的“信息技术”范畴。其核心观念在智能驾驶场景中仍具解释效力。因此,提出如下假设:

H1:用户对智能驾驶系统的感知控制对其使用意向具有显著的正向影响。

2.2 控制感相关理论模型

深入解析控制感的生成机制是构建有效干预策略的关键。尽管学界已提出了包括比较器模型(Comparator Model)、表观心理因果关系理论(Apparent Mental Causation Theory)、两阶段模型(Two-Stage Model)、线索整合理论(Cue Integration Theory)及优化线索整合模型(Optimal Cue Integration)在内的五种主流认知模型(Frith et al., 2000; Moore et al., 2009; Synofzik et al., 2008; Synofzik et al., 2013;

Wegner & Wheatley, 1999), 但纵观这些理论, 其核心逻辑均殊途同归地指向了两个基础维度: 可预测 (Predictability) 与可归因 (Attributability)。其中, 可预测性衡量了系统或现象能被精确预判的程度 (Lorenz, 1984; Krishnamurthy, 2019); 而可归因性则解释了个体对事件结果的因果溯源过程 (Malle, 2022)。

在众多理论中, 比较器模型与两阶段模型最为清晰地揭示了这两个维度在控制感形成中的协同机制。

比较器模型的核心机制在于“预测-反馈”匹配。大脑在执行动作前会生成对动作结果的预期。当实际感官反馈与预期的结果状态高度一致时, 个体即体验到“如我所料”, 这是产生控制感的关键信号 (Frith et al., 2000)。此外, “预测状态”与“估计实际状态”的匹配, 最终服务于确认“动作由我发起”, 这本身也是一种因果归因。比较器模型通过其内在的计算逻辑将结果的可预测性验证与动作的自我归因确认紧密耦合, 共同服务于控制感的产生。

两阶段模型则在比较器模型的基础上, 明确地将控制感的形成拆解为预测阶段与归因阶段 (Synofzik et al., 2008): 预测阶段是指事件发生前生成行为-结果的因果预期, 个体基于先验知识、意图和目标, 对即将发生的行为及其可能产生的结果进行预测, 依赖于对行为-结果链的可预测性评估。归因阶段是指事件发生后对结果进行归因判断, 确认因果关系是否由自身引发。

而其余相关理论模型虽然在侧重点和描述上各有不同, 但均指出控制感的产生植根于可预测与可归因的协同作用。

基于这一理论共识, 研究将这两个维度确立为探究智能驾驶情境下用户控制感的核心解构要素。

2.3 智能驾驶场景下可预测与可归因对控制感的影响

在社会互动与环境适应过程中, 个体倾向于对周遭事件的结果进行有意识或无意识的预测与解释, 以探寻背后的因果机制, 从而实现对环境的有效掌控 (刘世奎, 1991)。Weiner (1985) 指出, 当用户能将系统行为归因于明确因素时, 其对系统的理解与预期将显著增强, 进而提升控制感。此外, 控制感补偿理论进一步表明, 当个体遭遇控制感缺失时, 会本能地通过重建因果关系来恢复心理秩序 (白洁等, 2017)。

这一机制在驾驶情境的变迁中尤为关键。在传统驾驶中, 驾驶员通过直接操控车辆实时响应路况, 这种控制感源于对自身行为及其后果的明确预判与归因。然而, 在智能驾驶情境下, 驾驶员的角色由“物理操控”转向“认知监控”, 这种转变导致个体行为与系统反馈之间的因果链条被解耦。Pagliari 等人 (2022) 认为, 这种解耦削弱了个体对系统行为的预测能力与归因清晰度, 是导致控制感降低的根本原因。

基于上述理论推演, 若个体在智能驾驶中能够对系统行为进行有效预测, 并对结果进行合理的解释与归因, 将能在认知层面补偿因物理操控权让渡而断裂的行为-结果关联, 从而显著修复并提升其对系统的控制感。

据此, 提出如下研究假设:

H2: 智能驾驶系统行为的可预测对用户的控制感具有显著的正向影响。

H3: 智能驾驶系统行为结果的可归因对用户的控制感具有显著的正向影响。

2.4 系统行为可预测的影响因素

系统行为的可预测通常受到以下几个因素的影响:

一是系统行为的规律性 (Regularity), 即系统行为在时间或空间上呈现可重复、有序的模式或统计特性, 规律性为预测提供了可能 (Scriven, 1962)。具体到智能体 (Agent) 研究领域, 规律性特指智能体在与环境交互或执行任务过程中, 表现出的可被建模的稳定行为模式 (Chakraborti et al., 2019)。这种行为的规律性通常源于系统底层的设计约束, 或由其学习算法的内在属性所决定。

二是系统的稳定性 (Stability), 即系统在扰动或初始条件变化下保持平衡状态或收敛到预期行为的能力 (方庆霞, 2008)。稳定性确保系统状态或输出在扰动下不会无限发散, 是可预测性的重要基础 (Bazzi et al., 2018), 系统的稳定性越高其可预测性也就越强 (Sternad, 2017), 而系统的稳定性通常需要一些技术手段如控制算法、结构设计、容错机制等来保证。

三是系统的透明度 (Transparency), 即指系统内部结构、数据流转及决策逻辑对用户的可见性与可访问程度 (Kim & Hinds, 2006)。在智能驾驶情境中, 透明度通过揭示系统的内部状态与运行逻辑, 能够有效消解智能算法的“黑箱”效应。这种机制帮

助用户理解系统状态的演变过程,从而显著增强其对系统行为的可预测性感知(Chen et al., 2014; Mercado et al., 2016);且研究表明,这种正向影响会随着信息透明度的提升而增强。Naujoks 等人(2019)在 L3 级自动驾驶研究中进一步证实,当系统提供决策依据信息时,驾驶员的情境意识(Situation Awareness)得到显著改善,不仅能更精准地预测系统行为,还强化了个体对系统的掌控感。

四是系统行为的可解释(interpret-ability),即系统模型具有的使受众明了其运行机制的能力,强调系统或模型的内部机制、决策逻辑和行为模式能够清晰地传达并被个体理解(Hong et al., 2020)。当用户理解系统的决策规则时,他们能更准确地预测系统行为。以智能驾驶为代表的数据挖掘和机器学习场景下,可解释性被定义为以人类能理解的方式向人类解释或提供意义的能力(Doshi-Velez & Kim, 2017)。对智能驾驶系统决策逻辑和行为原因的解释,可以传达对系统未来行为的预期,提升智能体的可预测性,直接影响了个体对其可控性的感知(Koo et al., 2015)。

综合上述分析,系统的规律性与稳定性主要依赖于底层技术内核的演进;相比之下,人工智能系统的可解释性和透明度则属于人机交互的核心议题,旨在增强用户对系统的理解,从而构建可信、可控的人机协同认知生态(Walmsley, 2021)。鉴于此,立足于人机交互视角,通过设计手段提升系统的可预测性具有重要的理论与实践意义。因此,将研究焦点置于透明度和可解释性这两个关键维度,据此提出如下假设:

H2a:智能驾驶系统行为的可解释正向显著影响用户对系统行为的感知可预测。

H2a₁:智能驾驶系统行为的可解释正向显著影响用户对系统的控制感。

H2b:智能驾驶系统的透明度正向显著影响用户对系统行为的感知可预测。

H2b₁:智能驾驶系统的透明度正向显著影响用户对系统的控制感。

2.5 系统行为结果的可归因影响因素

系统行为结果的可归因可以帮助个体判断系统行为应归因于系统自身特性还是外部环境因素,本质上是个体对系统行为原因的归责判断。系统的可归因性可以通过构建与个体已有认知模型相吻合的认知链路和提供信息线索两种方式实现。

构建与个体经验认知一致的归因链路,当个体试图对系统行为或结果进行归因时,通常会对照自己过往的经验和已经形成的因果图式,去理解其背后的因果关系(Kelley, 1973; 宫晓东等, 2021)。因此,如果系统行为与用户既有的认知模式相符合,个体就会更容易理解和解释系统的行为及其产生的原因,增强了系统的可归因性。此处的一致性原则指的是智能驾驶系统以与用户已有的因果认知图式相吻合的方式传递其行为逻辑,就可以降低用户归因时的认知负荷,使其能够快速、自然地归因,提升系统的可归因性与用户控制感。

在智能驾驶系统中,可归因层面的一致性研究可以解释为用户主观认知模型与机器决策逻辑的认知对齐。如果二者偏差超出用户认知弹性阈值时,即归因失败,则会触发控制感的断裂与信任的指数级衰减(Lee & See, 2004)。

提供信息的充分性,系统或环境提供的信息是否充分在个体归因过程中发挥着至关重要的作用,直接影响个体是否能够准确、全面地识别并理解事件的原因(Nisbett & Ross, 1980)。当信息充分时,个体能够更全面地了解事件的背景、相关因素以及因果关系,从而做出准确归因。如果信息有所遗漏,将会导致归因偏差。

综上所述,提出以下假设:

H3a:智能驾驶系统行为结果与个体经验认知结果的一致性正向显著影响系统行为结果的可归因。

H3a₁:智能驾驶系统行为结果与个体经验认知结果的一致性正向显著影响用户对系统的控制感。

H3b:智能驾驶系统提供信息的充分性正向显著影响系统行为结果的可归因。

基于上述理论阐释与提出的假设,关于个体对智能驾驶系统控制感的研究框架得到确立,如图4所示。

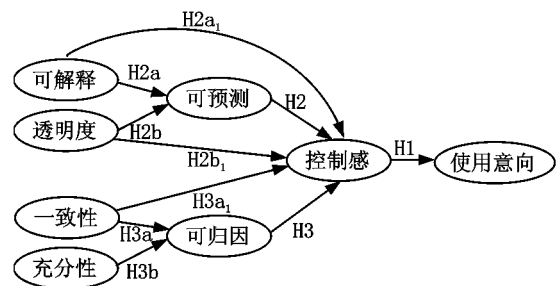


图4 面向智能驾驶控制感的问卷结构模型

3 研究方法

3.1 问卷设计

问卷包含 8 个潜变量和 24 个显变量。所有测量题项均来源于已有文献中的成熟量表或经过适当的情境化改编。每个潜变量包含 2-4 个显变量,形成完整的测量模型,如表 1 所示。”在量表设计方

面,采用 5 点李克特量表进行测量,量表选项从 1(完全不赞同)到 5(完全赞同)。

此外,为确保问卷作答的有效性和科学性,在问卷正式作答前,研究者向被试提供关于智能驾驶系统的详细视频材料,确保被试在作答过程中具备充分的认知基础和对系统的准确理解。

表 1 潜变量与对应显变量及问项来源

潜变量	显变量	来源
使用意向	1. 未来我会使用这样的智能驾驶汽车	Choi 等(2015)
	2. 我盼望未来使用这样的智能驾驶汽车	
	3. 我计划好了在未来使用这样的智能驾驶汽车	
	4. 在智能驾驶过程中我很少感到无助	
控制感	5. 在智能驾驶中我没有想要手动控制车辆的冲动	Lachman 等(1998) Tapal 等(2017)
	6. 我有办法解决智能驾驶中的任何问题	
	7. 在智能驾驶过程中,我感到对自己的安全有保障,不需要更多的车辆控制权	
	8. 智能驾驶系统具有一定规律性,使我能够预测车辆的下一步行动	
可预测	9. 智能驾驶系统的行为在我的意料之中	Rempel(1985) Endsley(2003)
	10. 我能够根据过去的经验预测车辆未来的表现	
可归因	11. 我能够知道智能驾驶系统的决策依据	Westaby 等(2005) Fishbein(1977) Liao 等(2023)
	12. 我能够理解智能驾驶系统为什么会做出某些行为	
	13. 智能驾驶系统的行为结果可以归结为特定的因素	
	14. 智能驾驶系统提供了足够的解释帮助我理解其决策	
可解释	15. 我能够理解智能驾驶系统的工作原理	Schank(2014) Liao 等(2023)
	16. 我可以清楚地知道智能驾驶系统是如何得出结论的	
	17. 智能驾驶系统的操作和决策过程是透明的	
透明度	18. 我可以轻松获取智能驾驶系统的相关信息	Schneider 等(2021) Saito 等(2021)
	19. 智能驾驶系统对外部是开放的,信息是易获取的	
一致性	20. 智能驾驶系统的行为通常符合我的驾驶习惯	Wilson(1989) Yampolskiy(2013)
	21. 智能驾驶系统的行为与我预期的一致	
	22. 我可以获取足够的信息来理解当前的情况	
充分性	23. 我可以获取足够的信息做出明智的决策	Kelley(1967) Nisbett(1980)
	24. 系统提供的信息与我当前任务或情境高度相关	

3.2 研究对象

整个量表编制过程包含初测、正式施测两个阶段。其中初测阶段采用随机取样,选取 15 名具有驾驶经验的被试(男性 9 人,女性 6 人;年龄 $M = 32.4$ 岁, $SD = 8.7$,学历均为本科及以上),根据反馈结果对问卷内容进行优化和修订,从而形成正式施测问卷。正式施测阶段采用问卷星网络平台与线下纸质问卷发放两种途径,以方便取样方式获取问卷数据,共回收问卷 374 份,根据作答情况剔除无效问卷 16 份,最终获得有效问卷 358 份,有效回收率为 95.72%。其中男性 214 人(59.18%),女性 144 人(40.22%),18~30 岁 204 人(56.98%),31~50 岁 104 人(29.05%),51~60 岁及以上 50 人(13.96%),学历均为高中及以上,且持有效机动车驾驶证(实际驾驶经

验 ≥ 2 年)。

3.3 统计分析

采用 SPSS 27.0 进行信度分析和探索性因子分析,采用 Amos 26.0 完成验证性因素分析以及结构方程模型分析。

4 研究结果

4.1 问卷分析

4.1.1 信度分析

采用 Cronbach's α 系数对问卷信度进行检验。结果显示,总量表的 Cronbach's 系数为 0.854,各维度的 Cronbach's α 系数均在可接受范围内。具体而言:可解释性(0.766)、透明度(0.785)、一致性(0.691)、充分性(0.761)、可预测性(0.830)、可归因性(0.831)、控制感(0.914)及使用意向(0.801),

数据表明测量工具整体具有良好的内部一致性。

4.1.2 探索性因素分析

对样本进行探索性因素分析(EFA), KMO = 0.827, Bartlett 球形检验结果表明, $\chi^2(276) = 3578.677, p < 0.001$, 表明数据可进行因子分析。采

用主成分分析法以及方差极大正交旋转法, 提取出 8 个公因子, 累计解释率为 74.071%。各条目的因子负荷范围为 0.726 ~ 0.860, 具体负荷情况如表 2 所示。

表 2 探索性因素分析各条目因子负荷结果

	RC1	RC2	RC8	RC6	RC4	RC3	RC5	RC7	Communality
EX1							0.726		0.738
EX2							0.818		0.735
EX3							0.811		0.712
TR1				0.804					0.681
TR2				0.786					0.705
TR3				0.810					0.739
FE1								0.860	0.771
FE2								0.840	0.751
CO1						0.799			0.674
CO2						0.820			0.724
CO3						0.801			0.665
PD1			0.828						0.748
PD2			0.848						0.778
PD3			0.827						0.736
AB1		0.824							0.732
AB2		0.828							0.749
AB3		0.836							0.769
PC1	0.796								0.778
PC2	0.815								0.810
PC3	0.841								0.826
PC4	0.805								0.794
BI1					0.828				0.723
BI2					0.826				0.711
BI3					0.831				0.728

4.1.3 验证性因子分析

对样本进行验证性因子分析(Confirmatory Factor Analysis, CFA), 结果显示 CFI = 0.981, TLI = 0.978, SRMR = 0.044, RMSEA = 0.028, 均符合统计学标准。

4.1.4 区分效度分析

采用平均方差提取量(Average Variance Extracted, AVE)的平方根检验问卷的区分效度。Fornell 和 Larcker(1981)指出, 变量的 AVE 的方根需要大于与其他潜在变量间的相关值。结果显示 AVE 平方根值均大于同一列中其他变量的相关系数, 表明量表具有良好的区分效度, 具体结果如表 3 所示。

表 3 区分效度分析表

维度	可解释	透明度	一致性	充分性	可预测	可归因	控制感	使用意向
可解释								
透明度	0.401							
一致性	0.121	-0.031						
充分性	0.332	0.217	-0.020					
可预测	0.349	0.372	0.020	0.139				
可归因	0.152	0.057	0.380	0.310	0.052			
控制感	0.448	0.479	0.296	0.259	0.439	0.454		
使用意向	0.182	0.195	0.120	0.105	0.179	0.185	0.406	
AVE 的平方根	0.723	0.744	0.730	0.730	0.789	0.788	0.854	0.757

4.1.5 聚敛效度分析

为评估问卷的聚敛效度,采用平均方差提取量 (Average Variance Extracted, AVE) 和组合信度 (Composite Reliability, CR) 验证问卷的聚敛效度。结果显示,各维度的 AVE 均大于 0.5,除“一致性”

维度其余维度的 CR 值均大于 0.7,处于理想水平,而“一致性”维度的 CR 值在 0.6~0.7 可接受水平区间,表明本问卷具有较好的聚敛效度,具体结果如表 4 所示。

表 4 聚敛效度分析表

项目	可解释	透明度	一致性	充分性	可预测	可归因	控制感	使用意向
CR 值	0.766	0.787	0.694	0.770	0.831	0.831	0.915	0.801
AVE 值	0.523	0.553	0.533	0.533	0.622	0.622	0.729	0.573

4.2 结构方程模型分析

采用 Amos 26.0 软件进行结构方程模型分析。选取比值拟合指数 (Comparative Fit Index, CFI)、均方根残差 (Root Mean Square Residual, RMSR)、均方根误差近似 (Root Mean Square Error of Approximation, RMSEA) 调整的拟合优度指数 (Adjusted Goodness of Fit Index, AGFI)、图克-刘易斯指数 (Tucker-Lewis Index, TLI) 等拟合指标进行分析。模型拟

合度检验结果显示:GFI(0.936)和 AGFI(0.919)均大于 0.9 的优良阈值, RMR (0.044) 和 RMSEA (0.028) 均显著低于 0.05 的严格标准;在增值拟合指数方面, NFI(0.918)、IFI(0.981)、CFI(0.981) 和 TLI(0.978) 均超过 0.9 的推荐值。综合来看,各项适配指标均在推荐范围内,模型构建合理,路径分析结果如表 5 所示。

表 5 路径系数表

路径	标准化系数 β	非标准化系数	标准误差	z	P	置信区间下限	置信区间上限
H2a	0.239	0.264	(0.081)	3.258	0.001***	0.105	0.423
H2b	0.276	0.306	(0.081)	3.765	0.001***	0.147	0.465
H3a	0.386	0.445	(0.084)	5.271	0.001***	0.279	0.610
H3b	0.318	0.366	(0.077)	4.749	0.001***	0.215	0.517
H2a ₁	0.170	0.248	(0.087)	2.849	0.004**	0.077	0.418
H2b ₁	0.305	0.444	(0.091)	4.892	0.001***	0.266	0.622
H3a ₁	0.150	0.219	(0.087)	2.523	0.012*	0.049	0.389
H2	0.246	0.324	(0.075)	4.320	0.001***	0.177	0.471
H3	0.341	0.432	(0.075)	5.779	0.001***	0.285	0.579
H1	0.406	0.305	(0.050)	6.152	0.001***	0.208	0.402

注: * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

基于表 5 路径系数表的统计分析结果,本研究构建了智能驾驶系统中个体控制感的结构方程模型,如图 5 所示。模型验证了各路径的显著性,其中标准化系数均通过显著性检验,表明路径关系具有统计学意义。

该模型采用结构方程模型 (SEM) 分析了多个外源变量 (可解释、透明度、一致性、充分性) 对中介变量 (可预测、可归因) 对控制感的影响,以及控制感对使用意向的作用。路径系数表明了各变量之间的显著性关系及其相对影响强度。

5 讨论

研究表明在智能驾驶场景中,用户的控制感并非单一维度的心理概念,而是可以通过可预测性与可归因性这两个截然不同却又互补的路径来构建。

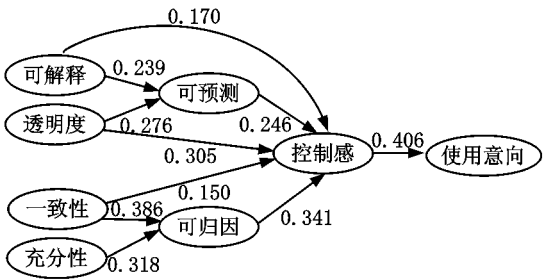


图 5 面向智能驾驶控制感的结构模型

其中,可归因对控制感的解释力显著高于可预测。这一结果具有深刻的理论意义。即在传统驾驶中,驾驶员作为直接“操控者”,其控制感高度依赖对车辆动态的实时预测和控制反馈。然而,在智能驾驶情境下,用户角色转变为“监督者”,直接的物理操

控被剥离,行为与结果间的即时、显性因果关系链断裂。此时,用户对“为什么系统会如此行为?”的理解变得比“接下来系统会做什么?”的预测更为关键。该结果有力地支持了这一角色转变带来的控制感内涵变化,凸显了可归因性在被动监控情境下的核心地位。

需要明确的是,研究结果并未否定可预测的重要性,而是揭示了在控制权让渡这一特定场景下,可归因的相对权重上升。这与线索整合理论和优化线索整合模型的理论结果相契合,不同线索的可靠性权重会因情境而动态调整。在智能驾驶的“监督者”模式下,归因线索的可靠性权重显著增加。

基于这一双路径机制,研究进一步揭示了不同交互设计要素的具体作用路径。

首先,针对基础的“可预测”路径,透明度与可解释性发挥了关键的前置作用。透明度通过实时揭示系统状态和意图,直接降低了环境的不确定性;而可解释性则更进一步,通过阐明决策的底层逻辑与规则,不仅提升了用户预判系统行为的能力,也为其理解行为原因奠定了基础。二者通过降低系统决策的“黑箱”特性,协同提升了行为预测的准确性。

其次,针对核心的“可归因”路径,一致性与信息充分性被证实是决定归因能力的关键变量。一致性要求系统行为符合用户基于经验形成的心理预期或社会规范,显著的不一致会阻碍用户将行为归因于合理的逻辑,从而导致失控感;而信息充分性则提供了因果推断所需的情境证据。充足且及时的信息支持用户快速构建因果链条,完成“为什么是这样?”的归因闭环。这与归因理论中关于信息充分性决定归因准确性的核心观点高度一致。

这些发现为面向智能驾驶系统的控制感设计提供了启示:其一,系统可以通过可视化决策逻辑(如实时显示变道原因或风险预判依据)提升透明度,以增强用户的可预测性感知;其二,设计应注重可解释性,例如以自然语言解释紧急制动或路径规划的底层规则,使用户能够理解并预测系统行为;其三,系统行为需符合用户既有驾驶习惯,并通过环境传感器数据等提供充分的情境信息,支持用户快速完成归因判断。这些策略可提升用户控制感,进而推动智能驾驶技术的广泛接受。

6 结论

(1)控制感显著正向影响个体使用意向。用户对智能驾驶系统的控制感越强,其使用该系统的意愿越高。

(2)系统行为的可预测与系统行为结果的可归

因是影响个体对智能驾驶系统控制感的核心因素。

(3)系统行为的可预测性可以通过提升系统透明度、以及系统决策规划的可解释信息,降低个体的认知不确定性,可以提高对系统的可预测性。

(4)系统行为与用户认知模式的一致性能够简化归因过程;提供充分的信息能够支持用户准确归因系统行为及结果。

研究明确了智能驾驶场景下用户控制感形成的“可预测”、“可归因”双路径机制,并从设计视角系统梳理了如何调控这些核心影响要素;然而,针对这些影响因素对应的具体设计方案形态、信息呈现方式及其实际交互效果,仍需通过用户测试与迭代设计开展深入验证与优化,未来研究可围绕上述具体设计落地与效果评估方向进一步拓展,为理论发现向实践转化提供更直接的支撑。

参考文献

- 白洁,郭永玉,杨沈龙.(2017).人在丧失控制感后会如何?来自补偿性控制理论的揭示.《中国临床心理学杂志》,25(5),982-985,981.
- 方庆霞.(2008).两类系统的相关稳定性分析(硕士学位论文).湘潭大学.
- 宫晓东,张佳乐,陈立翰.(2021).老年人心智模型研究及在交互设计领域的应用.《包装工程》,42(24),84-92.
- 刘世奎.(1991).归因理论的历史、现状及发展趋势.《心理学探新》,(4),1-3.
- 孟可强,刘文雅,吴博文,李旺,李凤兰.(2023).心理健康问题是“富贵病”吗?社会经济地位与农村居民心理健康问题:一个链式中介模型.《心理学探新》,43(5),433-440.
- 喻丰,赵一骏,许丽颖.(2024).人工智能社会心理学的瓶与酒:范式与问题.《心理学探新》,44(6),483-492.
- Ajzen, I. (1991). The theory of planned behavior. *Organizational Behavior and Human Decision Processes*, 50(2), 179-211.
- Bazzi, S., Ebert, J., Hogan, N., & Sternad, D. (2018). Stability and predictability in human control of complex objects. *Chaos: An Interdisciplinary Journal of Nonlinear Science*, 28(10).
- Berberian, B. (2019). Man - Machine teaming: A problem of Agency. *IFAC - PapersOnLine*, 51(34), 118-123.
- Chakraborti, T., Kulkarni, A., Sreedharan, S., Smith, D. E., & Kambhampati, S. (2019, July). Explicability? legibility? predictability? transparency? privacy? security? the emerging landscape of interpretable agent behavior. In *Proceedings of the international conference on automated planning and scheduling* (Vol. 29, pp. 86-96). IEEE.
- Chen, J. Y., Procci, K., Boyce, M., Wright, J., Garcia, A., & Barnes, M. (2014). Situation awareness - based agent transparency. *US Army Research Laboratory*, (April), 1-29.
- Choi, J. K., & Ji, Y. G. (2015). Investigating the importance of

- trust on adopting an autonomous vehicle. *International Journal of Human – Computer Interaction*, 31(10), 692 – 702.
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 319 – 340.
- Doshi – Velez, F., & Kim, B. (2011). Towards a rigorous science of interpretable machine learning. *arXiv preprint arXiv:1702.08608*.
- Endsley, M. R., Bolt, B., & Joes, D. G. (2003). *Designing for situation awareness: An approach to user – centered design*. CRC Press.
- Fishbein, M., & Ajzen, I. (1977). *Belief, attitude, intention, and behavior: An introduction to theory and research*. IEEE.
- Fornell, C., & Larcker, D. F. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of Marketing Research*, 18(1), 39 – 50.
- Frith, C. D., Blakemore, S. J., & Wolpert, D. M. (2000). Abnormalities in the awareness and control of action. *Philosophical Transactions of the Royal Society of London. Series B: Biological Sciences*, 355(1404), 1771 – 1788.
- Haggard, P., & Chambon, V. (2012). Sense of agency. *Current biology*, 22(10), R390 – R392.
- Hong, S. R., Hullman, J., & Bertini, E. (2020). Human factors in model interpretability: Industry practices, challenges, and needs. *Proceedings of the ACM on Human – Computer Interaction*, 4(CSCW1), 1 – 26.
- Kelley, H. H. (1967). Attribution theory in social psychology. In *Nebraska symposium on motivation*. University of Nebraska Press.
- Kelley, H. H. (1973). The processes of causal attribution. *American Psychologist*, 28(2), 107.
- Kim, T., & Hinds, P. (2006, September). Who should I blame? Effects of autonomy and transparency on attributions in human – robot interaction. In *ROMAN2006 – The 15th IEEE international symposium on robot and human interactive communication* (pp. 80 – 85). IEEE.
- Koo, J., Kwac, J., Ju, W., Steinert, M., Leifer, L., & Nass, C. (2015). Why did my car just do that? Explaining semi – autonomous driving actions to improve driver understanding, trust, and performance. *International Journal on Interactive Design and Manufacturing (IJIDeM)*, 9, 269 – 275.
- Krishnamurthy, V. (2019). Predictability of weather and climate. *Earth and Space Science*, 6(7), 1043 – 1056.
- Lachman, M. E., & Weaver, S. L. (1998). The sense of control as a moderator of social class differences in health and well-being. *Journal of Personality and Social Psychology*, 74(3), 763.
- Le Goff, K., Rey, A., Haggard, P., Oullier, O., & Berberian, B. (2018). Agency modulates interactions with automation technologies. *Ergonomics*, 61(9), 1282 – 1297.
- Lee, C., & Cha, K. (2024). Toward the dynamic relationship between AI transparency and trust in AI: A case study on ChatGPT. *International Journal of Human – Computer Interaction*, 1 – 18.
- Liao, X., Li, X., Cheng, Z., & Yang, Y. (2023). Perceived control in human – agent interaction: Scale development and validation. *革新的コンピューティング・情報・制御に関する国際誌*, 19(2), 597.
- Lorenz, E. N. (1984). Irregularity: A fundamental property of the atmosphere. *Tellus A*, 36(2), 98 – 110.
- Moore, J. W., Wegner, D. M., & Haggard, P. (2009). Modulating the sense of agency with external cues. *Consciousness and Cognition*, 18(4), 1056 – 1064.
- Malle, B. F. (2022). Attribution theories: How people make sense of behavior. *Theories in Social Psychology, Second Edition*, 93 – 120.
- Mercado, J. E., Rupp, M. A., Chen, J. Y., Barnes, M. J., Barber, D., & Procci, K. (2016). Intelligent agent transparency in human – agent teaming for Multi – UxV management. *Human Factors*, 58(3), 401 – 415.
- Naujoks, F., Wiedemann, K., Schömig, N., Hergeth, S., & Keinath, A. (2019). Towards guidelines and verification methods for automated vehicle HMI. *Transportation Research Part F: Traffic – Psychology and Behaviour*, 60, 121 – 136.
- Nisbett, R. E., & Ross, L. (1980). *Human inference: Strategies and shortcomings of social judgment*. IEEE.
- Pagliari, M., Chambon, V., & Berberian, B. (2022). What is new with artificial intelligence? Human – agent interactions through the lens of social agency. *Frontiers in Psychology*, 13, 954444.
- Rempel, J. K., Holmes, J. G., & Zanna, M. P. (1985). Trust in close relationships. *Journal of Personality and Social Psychology*, 49(1), 95.
- Saito, H., Horie, A., Maekawa, A., Matsubara, S., Wakisaka, S., Kashino, Z., . . . Inami, M. (2021). Transparency in human – machine mutual action. *Journal of Robotics and Mechatronics*, 33(5), 987 – 1003.
- Schank, R. C. (2014). Explanation: A first pass. In *Experience, memory, and reasoning* (pp. 139 – 165). Psychology Press.
- Schneider, T., Hois, J., Rosenstein, A., Ghellal, S., Theofanous – Füllbier, D., & Gerlicher, A. R. (2021, May). Explain yourself! transparency for positive ux in autonomous driving. In *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems* (pp. 1 – 12). IEEE.
- Scriven, M. (1962). *Explanations, predictions, and laws*. Minnesota Studies in the Philosophy of Science.
- Shapiro, D. H., Astin, J., Shapiro, S. L., Soucar, E. A., & Santerre, C. (1998). Control therapy. *The Corsini Encyclopedia of Psychology*, 11, 404 – 406.
- Sternad, D. (2017). Human control of interactions with objects

- variability, stability and predictability. In *Geometric and numerical foundations of movements* (pp. 301 – 335). Cham: Springer International Publishing.
- Synofzik, M., Vosgerau, G., & Newen, A. (2008). Beyond the comparator model: A multifactorial two – step account of agency. *Consciousness and Cognition*, 17(1), 219 – 239.
- Synofzik, M., Vosgerau, G., & Voss, M. (2013). The experience of agency: An interplay between prediction and postdiction. *Frontiers in Psychology*, 4, 127.
- Taylor, S., & Todd, P. (1995). Decomposition and crossover effects in the theory of planned behavior: A study of consumer adoption intentions. *International Journal of Research in Marketing*, 12(2), 137 – 155.
- Tapal, A., Oren, E., Dar, R., & Eitam, B. (2017). The sense of agency scale: A measure of consciously perceived control over one’s mind, body, and the immediate environment. *Frontiers in Psychology*, 8, 1552.
- Venkatesh, V., & Bala, H. (2008). Technology acceptance model 3 and a research agenda on interventions. *Decision Sciences*, 39(2), 273 – 315.
- Walmsley, J. (2021). Artificial intelligence and the value of transparency. *AI & Society*, 36(2), 585 – 595.
- Wegner, D. M., & Wheatley, T. (1999). Apparent mental causation: Sources of the experience of will. *American Psychologist*, 54(7), 480.
- Weiner, B. (1985). An attributional theory of achievement motivation and emotion. *Psychological Review*, 92(4), 548.
- Westaby, J. D. (2005). Behavioral reasoning theory: Identifying new linkages underlying intentions and behavior. *Organizational Behavior and Human Decision Processes*, 98(2), 97 – 120.
- Wilson, J. R., & Rutherford, A. (1989). Mental models: Theory and application in human factors. *Human Factors*, 31(6), 617 – 634.
- Yampolskiy, R. V., & Fox, J. (2013). Artificial general intelligence and the human mental model. In *Singularity hypotheses: A scientific and philosophical assessment* (pp. 129 – 145). Berlin, Heidelberg: Springer Berlin Heidelberg.

Research on the Influencing Factors of Interaction Design for Perceived Control in Intelligent Driving

Ji Yue Gong Xiaodong

(School of Design and Arts, Beijing Institute of Technology, Beijing 100081)

Abstract: Intelligent driving technology provides the possibility for building a smarter, safer, and more efficient transportation system, but whether the technology can benefit a wider group is still influenced by user acceptance. Through an analysis of the Technology Acceptance Model and related theories, it is pointed out that Perceived Control is one of the key factors affecting user technology acceptance, while in the intelligent driving scenario, the driver’s role transforms from “operator” to “supervisor,” and the transfer of control authority will inevitably have a negative impact on individual Perceived Control. Therefore, this study aims to enhance the acceptance of intelligent driving technology by exploring the factors affecting user Perceived Control in the intelligent driving context. The study integrates relevant classic theoretical frameworks, deconstructs Perceived Control into two mediating variables—Predictability and Attributability—and introduces four exogenous variables—Explainability, Transparency, Consistency, and Sufficiency—to construct a Structural Equation Model (SEM) of Perceived Control tailored to intelligent driving scenarios. A total of 358 valid questionnaires were collected, and path analysis was adopted to validate the model. The results indicate that: Predictability and Attributability are indeed key mediating variables affecting Perceived Control; specifically, Predictability is significantly influenced by Explainability and Transparency, while Attributability is significantly influenced by Consistency and Sufficiency, and the research hypotheses are basically verified. This study reveals the core factors affecting individual Perceived Control in the intelligent driving context, providing a theoretical basis for intervening in intelligent driving Perceived Control through interaction design and promoting the widespread application of intelligent driving technology.

Key words: intelligent driving; perceived control; predictability; attributability