

AI诊断情境下技术焦虑量表的修订及信效度检验

刘 祥¹, 牛海月¹, 刘红梅², 颜 颜¹, 吴淼菲¹, 史晓普¹, 王 瑞¹, 唐 翠³, 杨 磊^{1*}

¹湖州学院生命健康学院(体育部), 浙江 湖州

²湖州学院附属南太湖医院护理部, 浙江 湖州

³浙江鑫达医院, 浙江 湖州

收稿日期: 2026年7月12日; 录用日期: 2026年8月11日; 发布日期: 2026年8月19日

摘 要

目的: 基于压力认知评价理论(Cognitive Appraisal Theory), 修订适用于医疗人工智能(Artificial Intelligence)诊断场景的技术焦虑量表, 并检验其信效度, 为护理人员评估青年群体AI诊断焦虑水平提供可靠工具。方法: 收集有效样本1416份, 随机分半进行探索性与验证性因子分析; 3周后, 选取其中247名进行重测。采用技术焦虑量表、一般自我效能感量表(GSES)作为效标工具。结果: 量表最终包含效能失调、风险感知、伦理隐私3个因子, 共13个条目, 累计解释方差77.20%。验证性因子分析拟合良好($\chi^2/df = 2.66$, RMSEA = 0.08, CFI = 0.96, TLI = 0.95)。量表及维度得分与技术焦虑量表、GSES均显著相关($|r| = 0.33 \sim 0.63$)。总量表Cronbach's α 系数为0.93, 重测信度ICC为0.85。结论: 修订后的量表信效度良好, 可作为评估青年群体AI诊断焦虑的有效工具。

关键词

技术焦虑, 人工智能诊断, 量表修订, 信效度检验, 护理评估

Revision and Reliability and Validity Testing of the Technology Anxiety Scale in the Context of AI Diagnosis

Xiang Liu¹, Haiyue Niu¹, Hongmei Liu², Yan Yan¹, Miaofei Wu¹, Xiaopu Shi¹, Rui Wang¹, Cui Tang³, Lei Yang^{1*}

¹School of Life and Health Sciences, Huzhou College, Huzhou Zhejiang

²Nursing Department, Nantaihu Hospital Affiliated to Huzhou University, Huzhou Zhejiang

³Zhejiang Xinda Hospital, Huzhou Zhejiang

*通讯作者。

文章引用: 刘祥, 牛海月, 刘红梅, 颜颜, 吴淼菲, 史晓普, 王瑞, 唐翠, 杨磊. AI 诊断情境下技术焦虑量表的修订及信效度检验[J]. 护理学, 2026, 15(8): 309-318. DOI: 10.12677/ns.2026.158275

Abstract

Objective: Based on the Cognitive Appraisal Theory of stress, we revised a technology anxiety scale tailored for medical AI diagnostic scenarios and tested its reliability and validity, providing a trustworthy tool for healthcare workers to assess AI diagnosis anxiety among young people. **Method:** We collected 1,416 valid samples, randomly split them in half for exploratory and confirmatory factor analysis; three weeks later, 247 participants were retested. The technology anxiety scale and the General Self-Efficacy Scale (GSES) were used as criterion measures. **Results:** The final scale included three factors—efficacy dysfunction, risk perception, and ethical privacy—covering 13 items and explaining 77.20% of the total variance. Confirmatory factor analysis showed a good fit ($\chi^2/df = 2.66$, RMSEA = 0.08, CFI = 0.96, TLI = 0.95). The total scale and dimension scores were significantly correlated with both the technology anxiety scale and GSES ($|r| = 0.33\sim 0.63$). The total scale Cronbach's α was 0.93, and retest reliability ICC was 0.85. **Conclusion:** The revised scale has good reliability and validity and can serve as an effective tool for assessing AI diagnostic anxiety in young people.

Keywords

Technology Anxiety, Artificial Intelligence Diagnosis, Scale Revision, Reliability and Validity Testing, Nursing Assessment

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

当前, 人工智能(Artificial Intelligence, AI)技术正快速融入医学影像[1]、辅助诊断[2]、健康管理[3]等临床环节, 其在提升诊疗效率的同时[4], 也深刻改变了患者就医的决策环境[5]。其中, 作为智慧医疗核心环节的 AI 诊断, 因其直接关乎生命健康, 具有高专业性、高风险性及强隐私性等特征[6] [7], 使得青年群体在面对这一新兴技术时, 普遍存在紧张、担忧及行为回避的“技术焦虑”现象[8]。特别是在医疗诊断情境下, 此种焦虑不仅源于对技术操作的不熟悉, 更与误诊责任、生命健康风险、数据安全及算法伦理等特殊因素密切相关[9]。

在临床护理实践中, 护士是健康宣教、心理评估与人文关怀的主要执行者。面对服务对象日益增长的 AI 焦虑, 护士往往难以仅凭经验进行有效识别[10], 若能提前识别潜在使用者的 AI 诊断焦虑水平与核心诉求, 便可在技术应用的早期开展健康教育与心理干预, 帮助服务对象适应智慧医疗模式, 提升就医体验。目前, 国际应用较为广泛的是技术焦虑量表(Technophobia Scale), 该量表由 Khasawneh [11]编制, 国内经孙尔鸿[12]等针对老年人群完成汉化与信效度检验, 可用于评估个体面对现代技术时产生的焦虑情绪。但该量表维度设置聚焦传统信息技术, 缺乏对人工智能场景的针对性测量, 并不适用于医疗 AI 诊断特定情境下的焦虑评估。此外, 王点等引入的人工智能焦虑量表(AIAS) [13] [14]虽聚焦 AI 相关焦虑问题, 但其适用对象以医学生为主, 测评内容侧重学业、职业替代等教育场景, 与本研究面向 AI 诊断潜在服务使用者的测评目标及护理评估场景适配性不足。

为了精准捕捉青年群体在面对 AI 诊断时的心理应激机制, 本研究引入压力认知评价理论[15], 将其

作为量表修订的核心框架。对技术焦虑量表进行 AI 诊断情境化修订并完成信效度检验，以期为护理人员开展心理评估、健康宣教及人文干预提供可靠测评工具，助力智慧医疗护理服务开展。

2. 对象与方法

2.1. 研究对象

采用探索性因子分析与验证性因子分析的方法确定量表的结构，建议样本量至少是条目数量的 5~10 倍[16]，验证性因子分析需要更大的样本量。本研究共 13 个条目，故需要 130 份有效问卷。纳入标准：鉴于当前医疗人工智能诊断技术尚未全面普及，绝大多数人群尚未实际使用过该项技术，因此本研究以 AI 诊断潜在服务使用者为研究对象，即未来可能接受 AI 辅助诊断的青年高校学生。考虑到中青年群体是互联网医疗及 AI 诊断的主要潜在受众，且其认知功能完善、能独立理解并作答，所以收集对象限定被试年龄为 18~60 岁，且排除有精神病史者。

采用方便取样，通过在线问卷调查平台“问卷星”收集数据。线上发放问卷 1495 份，剔除无效(未通过测谎题、规律作答、作答时间过短)问卷 79 份，有效问卷 1416 份。并在 3 周后，方便选取其中 250 名被试进行重测，获有效问卷 247 份。具体见表 1。

Table 1. Demographic information
表 1. 人口学信息

人口学变量	类别	总样本(n = 1416)	重测子样本(n = 247)
年龄(岁)	范围	18~30	18~24
	$\bar{x} \pm s$	22.34 $\pm$ 4.15	21.45 $\pm$ 2.23
性别	男	660	116
	女	756	131
文化程度	大专/本科	1344	230
	研究生	72	17

2.2. AI 诊断技术焦虑量表的改编

2.2.1. 理论基础

本研究以压力认知评价理论[15]为框架开展量表修订工作。该理论是研究个体压力与情绪反应的经典理论，核心用于解释人们面对外界压力事件时的认知变化与情绪产生过程，适用于分析青年群体在接触新型、高风险医疗技术时的心理状态，能够为本研究探索 AI 诊断焦虑结构、修订量表提供扎实的理论支撑。如图 1 所示，该理论主张个体面对压力事件时需经历双重认知评价：

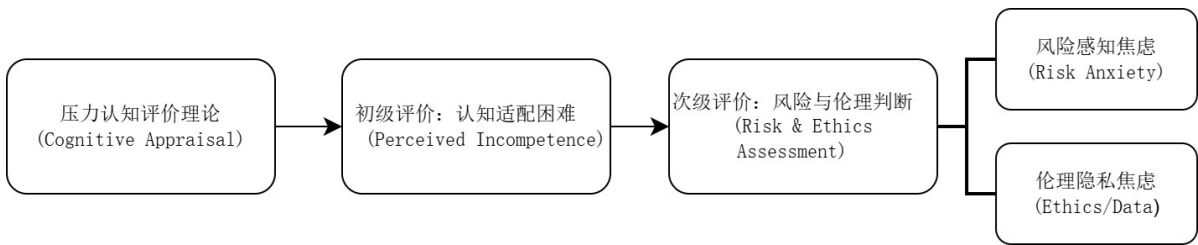


Figure 1. AI diagnosis anxiety structural model based on the pressure-cognition evaluation theory
图 1. 基于压力认知评价理论的 AI 诊断焦虑结构模型

首先是初级评价(Primary Appraisal)阶段。该阶段聚焦于外部环境属性的判定,即评估事件是否具有挑战性、威胁性或造成损伤。对于无医学及技术背景的潜在使用者而言,AI 诊断作为新兴智慧医疗技术,存在较高的认知门槛、应用不确定性与安全风险,属于典型的压力应激场景。当青年群体面对 AI 诊断的“算法黑箱”与复杂流程时,若判定自身无法理解或适配该技术,即产生认知层面的失控感与适配压力,本研究将其界定为“效能失调焦虑”。

其次是次级评价(Secondary Appraisal)阶段。该阶段指向内部资源评估,即判断自身是否具备应对该事件的资源与能力。鉴于医疗行为直接关系到生命健康,个体对 AI 诊断的评估不再停留于操作层面,而是向纵深发展:一方面,担忧 AI 误诊隐患、技术故障等会对健康造成实质性损害,形成“风险感知焦虑”;另一方面,担忧数据隐私泄露、算法伦理缺失等社会性风险,形成“伦理隐私焦虑”。个体最终的焦虑情绪与技术接纳态度,正是这两个阶段认知评价共同作用的结果。

基于上述理论逻辑,本研究选取孙尔鸿等汉化的技术焦虑量表[12]为修订蓝本,结合医疗 AI 诊断的场景特征对原量表条目进行情境化优化与筛选,最终确立效能失调焦虑、风险感知焦虑、伦理隐私焦虑三维结构。

2.2.2. 量表修订

1) 条目具体修订

本研究并未直接套用原量表,而是基于压力认知评价理论,结合 AI 医疗的技术特征,对原有条目进行“情境化”重构,以确保测量的准确性:情境化剔除:删除了“使用搜索引擎”等与医疗场景无关的通用技术条目,并将剩余条目中的测量对象统一锚定为“AI 诊断工具”,以消除无关变量干扰。专业化修正:将宽泛的表述(如“技术故障”)修正为具体的医疗后果(如“AI 误诊造成的身体伤害”),提升量表在特定情境下的构念效度。理论化增补:基于责任归属理论[7]与数据推断风险理论[17],自主新增了“误诊责任认定困难”与“健康数据被深挖”两个条目,以弥补传统量表在算法伦理维度的缺失,具体见表 2。

2) 专家评定

为保证改编质量,修订稿形成后,本研究共邀请 8 名专家参与德尔菲咨询[18][19],其中医学教育专家 3 名、心理学专家 3 名、信息技术专家 2 名。专家均具有副教授及以上职称。经过两轮咨询后,条目重要性评分趋于稳定,专家权威系数为 0.87。

3) 初始问卷

随后,在小范围目标人群(n=30)中进行预测试与认知访谈,以进一步优化条目的可读性与可理解性。共 13 个条目,采用 Likert 5 点计分法,从“完全不符合”(1 分)到“完全符合”(5 分),得分越高表示技术焦虑水平越高。

Table 2. Comparison of revised scale items
表 2. 量表修订条目对比表

原量表条目(技术焦虑量表)	修订后条目(AI 诊断技术焦虑量表)	修订理由/依据
1.1 当我必须使用新的通讯设备(新手机、新掌上电脑)时,我会感到不安	Q1.1 当必须使用 AI 诊断工具时,我内心会感到不安。	情境化修订
1.2 我尽可能地避免更换通讯设备(如换手机),因为这使我感到紧张	Q1.2 因为使用 AI 诊断工具让我感到紧张,我会尽量避免使用它。	情境化修订
1.3 我尽可能地避免使用新的技术,如智能手机、智能手环等	Q1.3 我会尽可能避免依赖 AI 技术进行健康评估或疾病诊断。	情境化修订

续表

1.4 当我必须学习新的操作系统时(如更新手机系统、更新电脑 windows7 到 8)时, 我会感到不安	Q1.4 当需要学习如何操作 AI 诊断平台时, 我会感到心神不宁。	情境化修订
1.5 我不敢使用百度、360 搜索、搜狗等网络搜索引擎	—(删除)	条目净化
2.1 我害怕新技术未来会取代我的工作	Q2.1 我担心 AI 诊断技术最终会完全取代医生进行疾病诊断。	情境化修订
2.2 我害怕迟早有一天新技术会让我们(人类)被淘汰	Q2.2 我担心随着 AI 诊断的普及, 医生在医疗中的地位会逐渐削弱。	情境化修订
2.3 我担心新技术, 是因为技术会在情感上、生理上和心理上影响我的生活	Q2.3 我担心 AI 诊断的普及, 会损害医患之间的情感沟通与信任。	构念聚焦
2.4 我非常害怕技术会改变我们的生活、沟通、爱、甚至评判他人的方式		
2.5 我担心新技术, 是因为一旦技术出现故障(如某些原因导致技术失灵), 我会回到石器时代而寸步难行	Q2.4 我惧怕 AI 诊断系统出现的失误, 可能会对我的身体健康造成无法挽回的伤害。	情境化修订
—(新增)	Q2.5 如果发生 AI 误诊, 我很难找到谁应该为此负责	基于责任归属理论
3.1 我害怕有人用技术时刻监听、监视我的所作所为	Q3.1 我担心 AI 诊断系统会在未经我同意的情况下, 收集或使用我的健康数据。	情境化修订
3.2 我非常害怕连接互联网后, 在网络上留下痕迹可能会被人追踪	Q3.2 我害怕在使用 AI 诊断后, 我的健康信息会被泄露或被他人不当利用。	情境化修订
3.3 我害怕百度、360 搜索、搜狗等搜索网站, 是担心它们更容易让别人跟踪监控我	Q3.3 我担心 AI 诊断平台的数据安全措施, 不足以保护我的个人隐私。	情境化修订
—(新增)	Q3.4 我害怕 AI 通过分析我的健康数据, 能推断出我其他的敏感信息。	基于数据推断风险理论

2.3. 校标工具

本研究共使用两个校标工具量表:

1) 技术焦虑量表: 评估个体使用数字健康技术过程中的技术焦虑情况。共 13 个条目, 采用 Likert 5 级评分“完全不符合”(1 分)到“完全符合”(5 分), 得分越高表示技术焦虑水平越高。本研究中, 该量表的 Cronbach's α 系数为 0.88。

2) 一般自我效能感量表(GSES) [20]: 共 10 个条目, 采用 Likert 4 级评分“完全不正确”(1 分)到“完全正确”(4 分), 得分越高表示一般自我效能感越高。本研究中, 该量表的 Cronbach's α 系数为 0.85。

2.4. 统计方法

采用 SPSS27.0 软件进行条目分析、探索性因子分析、效标关联效度分析和信度分析; 采用 AMOS26.0 软件进行验证性因子分析。

采用 SPSS27.0 中的随机抽样将总样本随机分为两组。先将一组($n = 708$)被试用于条目分析和探索性因子分析, 以初步确定量表的因子结构。然后将另一组($n = 708$)被试用于验证性因子分析, 以验证因子结构的适应性和泛化能力。总样本用于效标关联效度和内部一致性信度分析。

3. 结果

3.1. 条目分析

首先, 对各条目得分与 AI 诊断技术焦虑量表总分进行 Pearson 相关分析, 结果表明, 所有条目得分与总分的相关系数(r)在 0.67~0.79 (均 $P < 0.01$)。其次, 采用临界值法, 分别取总分 27%与后 27%作为低分组与高分组, 进行独立样本 t 检验, 结果发现高分组与低分组的 13 个条目得分差异均有统计学意义(均 $P < 0.01$)。

3.2. 效度分析

3.2.1. 结构效度

1) 探索性因子分析: KMO 值为 0.91, Bartlett 球形检验 χ^2 值为 2421.76 ($P < 0.001$), 表明数据适合进行因子分析。采用主轴因式分解因子提取与直接斜交因子旋转法, 提取特征根大于 1 的因子。结果显示, 共提取特征根大于 1 的因子 3 个, 其特征根分别为 3.76、3.31、2.97, 累计方差解释率为 77.20%。各条目的因子负荷见表 3。

Table 3. Factor loadings of each item on the AI diagnostic technology anxiety scale
表 3. AI 诊断技术焦虑量表各条目因子载荷($n = 708$)

条目	效能失调 焦虑	风险感知 焦虑	伦理隐私 焦虑
Q1.1 当必须使用 AI 诊断工具时, 我内心会感到不安。	0.88		
Q1.2 因为使用 AI 诊断工具让我感到紧张, 我会尽量避免使用它。	0.87		
Q1.3 我会尽可能避免依赖 AI 技术进行健康评估或疾病诊断。	0.86		
Q1.4 当需要学习如何操作 AI 诊断平台时, 我会感到心神不宁。	0.83		
Q2.1 我担心 AI 诊断技术最终会完全取代医生进行疾病诊断。		0.68	
Q2.2 我担心随着 AI 诊断的普及, 医生在医疗中的地位会逐渐削弱。		0.66	
Q2.3 我担心 AI 诊断的普及, 会损害医患之间的情感沟通与信任。		0.76	
Q2.4 我惧怕 AI 诊断系统出现的失误, 可能会对我的身体健康造成无法挽回的伤害。		0.69	
Q2.5 如果发生 AI 误诊, 我很难找到谁应该为此负责		0.69	
Q3.1 我担心 AI 诊断系统会在未经我同意的情况下, 收集或使用我的健康数据。			0.79
Q3.2 我害怕在使用 AI 诊断后, 我的健康信息会被泄露或被他人不当利用。			0.83
Q3.3 我担心 AI 诊断平台的数据安全措施, 不足以保护我的个人隐私。			0.88
Q3.4 我害怕 AI 通过分析我的健康数据, 能推断出我其他的敏感信息。			0.78

2) 验证性因子分析: 采用最大似然法检验模型拟合, 结果显示, 量表单因子模型的各项拟合指数基本符合统计要求($\chi^2/df = 2.66$, RMSEA = 0.08, CFI = 0.96, TLI = 0.95, SRMR = 0.09)。各条目的标准化因子负荷在 0.68~0.90 之间, 见图 2。

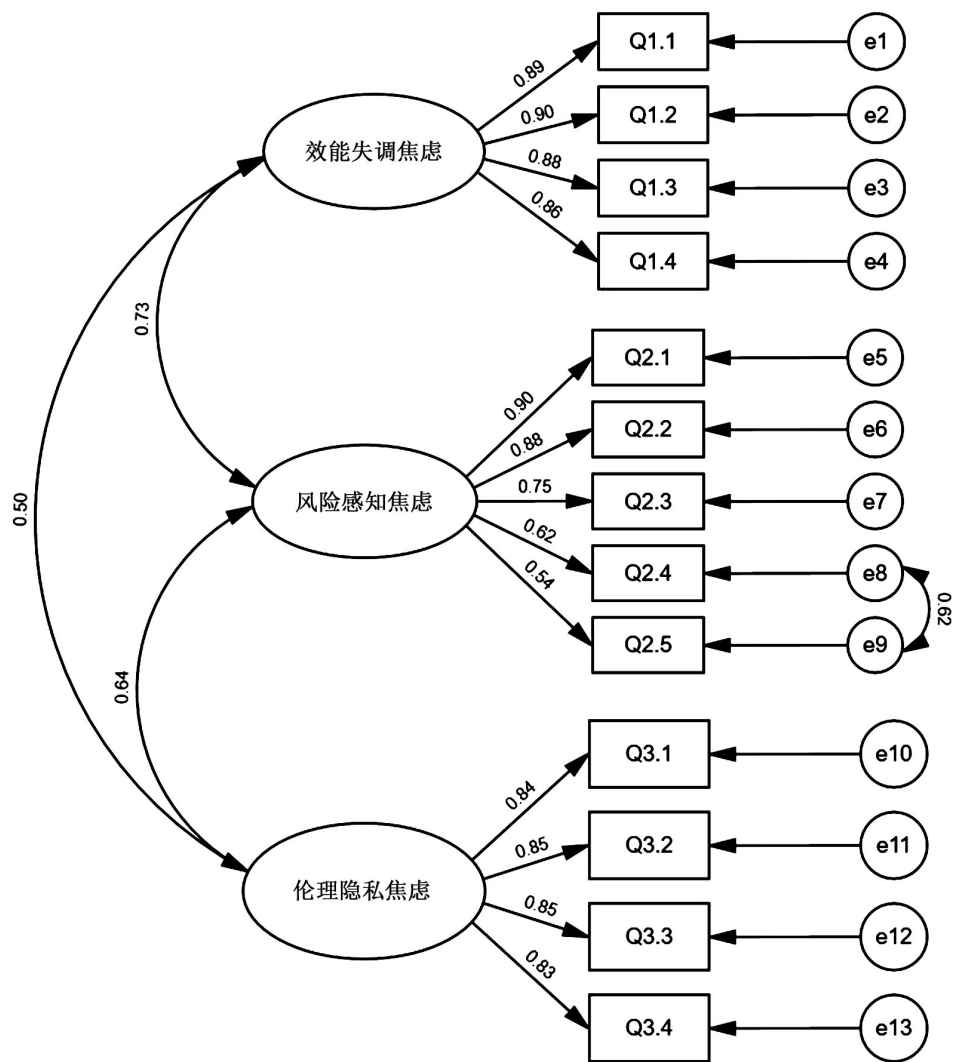


Figure 2. Confirmatory factor analysis model of the AI diagnosis technology anxiety scale after revision
图 2. AI 诊断技术焦虑量表修正后验证性因子分析模型

3.2.2. 校标效度

各变量数据满足正态分布，经 Pearson 相关分析表明，AI 诊断技术焦虑量表总分及各因子得分与技术焦虑量表得分成正相关，与一般自我效能感得分成负相关，见表 4。

Table 4. Correlation between total scores and each dimension of AI diagnostic technology anxiety scale and criterion scores (r, n = 708)

表 4. AI 诊断技术焦虑量表总分及各维度与效标得分的相关(r, n = 708)

维度	$\bar{x} \pm s$	技术焦虑量表得分	一般自我效能感得分
效能失调焦虑	13.35 ± 4.50	0.37**	-0.63**
风险感知焦虑	17.21 ± 4.93	0.51**	-0.38**
伦理隐私焦虑	14.11 ± 4.01	0.43**	-0.33**
总分	44.67 ± 11.52	0.53**	-0.48**

注：** $P < 0.01$ 。

3.3. 信度分析

1) 内部一致性信度: AI 诊断技术焦虑量表的 Cronbach's α 系数为 0.93, 效能失调焦虑、风险感知焦虑、伦理隐私焦虑 3 个因子的 Cronbach's α 系数分别为 0.93、0.88、0.91。

2) 重测信度: AI 诊断技术焦虑量表的组内相关系数(ICC)为 0.85, 效能失调焦虑、风险感知焦虑、伦理隐私焦虑 3 个因子的重测信度(ICC)分别为 0.76、0.72、0.82。

4. 讨论

4.1. 量表的三个维度的理论基础与验证

本研究以压力认知评价理论为框架,以青年高校学生为主要研究对象,对技术焦虑量表进行医疗场景的情境化修订与验证。本研究依据“初级评价-次级评价”的理论基础,将青年群体面对 AI 诊断时的复杂焦虑,系统性地解析为三个逐层递进的维度:在初级评价中,因感知自身能力与技术要求不匹配而产生的效能失调焦虑;在次级评价中,对技术可能造成的健康损害与系统风险的风险感知焦虑;以及对数据安全、医患关系等社会性后果的伦理隐私焦虑。探索性因子分析结果支持该三维结构,累计解释方差 77.20%;验证性因子分析表明,除 SRMR 值(0.09)受大样本特性影响略高于理想阈值外[21], χ^2/df (2.66)、CFI (0.96)、TLI (0.95)及 RMSEA (0.08)等核心指标均达到优异标准[22],有力证实三因子模型在医疗 AI 诊断场景下的适配性与稳健性。

4.2. 量表的信度与效度检验的结果与解读

在测量学属性方面,量表展现出良好的信效度。条目分析显示各条目与总分相关系数在 0.67~0.79 之间且高低分组差异显著,表明量表条目区分度良好,无冗余信息[23]。信度检验中,总量表 Cronbach's α 系数高达 0.93,重测信度 ICC 为 0.85,均显著优于心理测量学标准,满足临床护理评估对工具精度的要求。尤为重要的是效标关联效度的发现:量表总分与一般自我效能感呈显著负相关,且效能失调焦虑维度的负相关性最强($r = -0.63$)。这一实证结果有力地支撑了理论预设,即个体对 AI 技术的自我效能感越低,越倾向于在初级评价中判定自己无法适配技术,从而产生焦虑情绪。这与庞华等[24]的研究发现相呼应,即 AI 焦虑本质上是个体在认知评价过程中产生的负面情感交互结果。本研究进一步证实,这种评价机制在医疗诊断这一高风险场景下同样适用,从而验证了压力认知评价理论在青年群体医疗诊断场景中的适用性。

4.3. 量表的适用性局限与研究展望

本研究聚焦于青年高校学生群体,修订了适用于该群体的 AI 诊断焦虑量表。该量表以其简明、高效和情境专用的特点,能够契合门诊健康宣教、体检中心初筛等快节奏场景的快速评估需求,其三维结构(效能失调、风险感知、伦理隐私)能够精准区分焦虑的心理来源,为后续实施针对性心理疏导、认知干预或信息支持提供直接依据。然而,本研究存在以下局限性:首先,样本集中于 18~30 岁高校学生,大专/本科占比 94.9%,且多为尚未实际接受 AI 诊断的潜在使用者,这一高度同质化的结构意味着结论无法直接外推至老年人及其他低数字素养群体,且“前置性焦虑”与“使用后焦虑”的结构是否一致尚无证据,未来应在老年慢性病患者、基层社区人群以及已实际接受 AI 辅助诊断的患者中开展多群组验证性因子分析,并根据目标群体的语言习惯对条目做适应性调整。其次,本研究为横断面设计,仅能揭示相关关系,无法推断焦虑与技术接纳行为之间的因果时序,未来可尝试采用混合研究方法,在量化追踪的基础上结合半结构化访谈。再者,尚未形成配套的干预方案,三维结构的临床价值尚未兑现,未来应基于三维度分别开发标准化护理干预模块:针对效能失调焦虑设计 AI 诊断流程模拟体验与操作可视化指引,针对风

险感知焦虑制作“AI-医生双审核”流程科普卡并明确误诊责任归属,针对伦理隐私焦虑在就诊前签署数据使用知情同意承诺。最后,本研究仅报告了样本均值,尚未建立轻度/中度/重度焦虑的划界分数,限制了量表在临床筛查中的分级应用,未来应扩大样本(覆盖多省市、多年龄段),建立常模与划界分数,使量表从研究工具真正转化为临床筛查工具。

声 明

本研究得到湖州学院附属南太湖医院医学科研与临床试验伦理委员会的批准(伦理批号:0401)。

基金项目

湖州学院 2025 年度校级教育教学改革研究项目,编号:JY66056;2025 年浙江省级大学生创新创业训练计划项目,编号:2025CXCY065。

参考文献

- [1] 周翔,王培军.中国医学影像人工智能发展现状及展望[J].同济大学学报(医学版),2025,46(1):1-7.
- [2] Cheng, Y., Azouzi, N., Laurent-Bellue, A., Guo, Z., Chung, T., Zeng, Q., *et al.* (2026) A Confidence-Based, Artificial Intelligence Pathology Model for Diagnosis of Intrahepatic Cholangiocarcinoma. *Annals of Oncology*, **37**, 974-985. <https://doi.org/10.1016/j.annonc.2026.02.018>
- [3] Zheng, X., Liu, Z., Liu, J., Hu, C., Du, Y., Li, J., *et al.* (2025) Advancing Sports Cardiology: Integrating Artificial Intelligence with Wearable Devices for Cardiovascular Health Management. *ACS Applied Materials & Interfaces*, **17**, 17895-17920. <https://doi.org/10.1021/acsami.4c22895>
- [4] 杨弋,黄燕.人工智能辅助诊断系统研究进展[J].实用医院临床杂志,2020,17(4):275-277.
- [5] 田林,任绪泽,涂峥程.人工智能、机器学习和深度学习在医学诊断中的应用进展[J].现代医学,2024,52(9):1480-1484.
- [6] Topol, E.J. (2019) High-Performance Medicine: The Convergence of Human and Artificial Intelligence. *Nature Medicine*, **25**, 44-56. <https://doi.org/10.1038/s41591-018-0300-7>
- [7] Char, D.S., Shah, N.H. and Magnus, D. (2018) Implementing Machine Learning in Health Care—Addressing Ethical Challenges. *New England Journal of Medicine*, **378**, 981-983. <https://doi.org/10.1056/nejmp1714229>
- [8] 李建春,刘雨丹,梁昌萍,等.数智赋能视域下老年人群就医技术焦虑研究进展[J].老年医学研究,2025,6(5):19-23.
- [9] Li, Y., Yi, X., Fu, J., Yang, Y., Duan, C. and Wang, J. (2025) Reducing Misdiagnosis in AI-Driven Medical Diagnostics: A Multidimensional Framework for Technical, Ethical, and Policy Solutions. *Frontiers in Medicine*, **12**, Article ID: 1594450. <https://doi.org/10.3389/fmed.2025.1594450>
- [10] Nirgiz, C., Sarı, M.K. and Çaylı, N. (2026) The Relationship between Nurses' Anxiety and Attitudes Towards Artificial Intelligence and Examination of Influencing Factors. *BMC Nursing*, **25**, Article No. 122. <https://doi.org/10.1186/s12912-026-04293-9>
- [11] Khasawneh, O.Y. (2018) Technophobia: Examining Its Hidden Factors and Defining It. *Technology in Society*, **54**, 93-100. <https://doi.org/10.1016/j.techsoc.2018.03.008>
- [12] 孙尔鸿,高宇,叶旭春.技术焦虑量表的汉化及其在老年群体中的信效度检验[J].中华护理杂志,2022,57(3):380-384.
- [13] 王点,商丽,孙莉,等.人工智能焦虑量表的汉化及信效度检验[J].护理研究,2026(10):1730-1735.
- [14] Wang, Y.Y. and Wang, Y.S. (2022) Development and Validation of an Artificial Intelligence Anxiety Scale: An Initial Application in Predicting Motivated Learning Behavior. *Interactive Learning Environments*, **30**, 619-634. <https://doi.org/10.1080/10494820.2019.1674887>
- [15] Folkman, S., Lazarus, R.S., Dunkel-Schetter, C., DeLongis, A. and Gruen, R.J. (1986) Dynamics of a Stressful Encounter: Cognitive Appraisal, Coping, and Encounter Outcomes. *Journal of Personality and Social Psychology*, **50**, 992-1003. <https://doi.org/10.1037/0022-3514.50.5.992>
- [16] Boateng, G.O., Neilands, T.B., Frongillo, E.A., Melgar-Quinonez, H.R. and Young, S.L. (2018) Best Practices for Developing and Validating Scales for Health, Social, and Behavioral Research: A Primer. *Frontiers in Public Health*, **6**, Article ID: 149. <https://doi.org/10.3389/fpubh.2018.00149>

-
- [17] Narayanan, A. and Shmatikov, V. (2008) Robust De-Anonymization of Large Sparse Datasets. 2008 *IEEE Symposium on Security and Privacy* (SP 2008), Oakland, 18-22 May 2008, 111-125. <https://doi.org/10.1109/sp.2008.33>
- [18] Dalkey, N. and Helmer, O. (1963) An Experimental Application of the DELPHI Method to the Use of Experts. *Management Science*, **9**, 458-467. <https://doi.org/10.1287/mnsc.9.3.458>
- [19] 屈京楼, 朱亚鑫, 曲波等. 德尔菲法在医学教育研究中的应用[J]. 中华医学教育杂志, 2019, 39(3): 227-230.
- [20] 王才康, 胡中锋, 刘勇. 一般自我效能感量表的信度和效度研究[J]. 应用心理学, 2001(1): 37-40.
- [21] Shi, D. and Maydeu-Olivares, A. (2020) The Effect of Estimation Methods on SEM Fit Indices. *Educational and Psychological Measurement*, **80**, 421-445. <https://doi.org/10.1177/0013164419885164>
- [22] 温忠麟, 侯杰泰, 马什赫伯特. 结构方程模型检验: 拟合指数与卡方准则[J]. 心理学报, 2004(2): 186-194.
- [23] Streiner, D.L., Norman, G.R. and Cairney, J. (2024) *Health Measurement Scales: A Practical Guide to Their Development and Use*. Oxford University Press.
- [24] 庞华, 朱玉廷, 李俊凯. 从技术认知到情感交互: 压力认识评价理论下人工智能焦虑对用户间歇性中辍行为的影响机制研究[J/OL]. 新媒体与社会: 1-22. <https://link.cnki.net/urlid/cn.20260128.0941.002>, 2026-06-12.