

交互式问题解决测验中学习效应的分析： 过程数据测量模型的拓展与应用*

陆翔宇 陈 平

(北京师范大学中国基础教育质量监测协同创新中心, 北京 100875)

摘 要 在交互式问题解决测验中, 问题情境不是一次性呈现完整, 被试需要进行探索逐渐积累信息。这使得被试当前状态的行为选择, 不仅受到其问题解决能力的影响, 还受到其对问题情境的了解程度的影响(即学习效应)。针对现有模型方法的缺陷, 在单参数行动序列模型(1P-ASM)的基础上引入当前状态在作答序列中的位置这一变量, 对被试在问题解决过程中的学习效应进行建模, 提出考虑学习效应的 1P-ASM 拓展模型(1P-ASM-R*), 并通过实证研究和模拟研究评估新模型 1P-ASM-R* 的表现。结果显示: (1) 相比于 1P-ASM, 1P-ASM-R* 能更好地拟合实证数据; (2) 在模型中引入学习效应不影响其捕捉问题解决任务的特征。总之, 在问题解决能力过程数据测量模型中引入学习效应能够获得更加准确的问题解决能力估计值, 为过程数据的分析提供新的、有价值的方法。

关键词 过程数据, 问题解决, 项目反应理论

分类号 B841

1 引言

随着心理与教育研究的不断深入, 复杂问题解决等传统纸笔测验难以测量的高阶认知能力逐渐得到更多关注, 这对统计测量方法也提出更高要求。相比于纸笔测验, 计算机化的交互式问题解决测验通过计算机仿真模拟的方式开展测试, 在测验过程中能够实时记录个体行为表现并存储在日志文件中。记录的数据包括题目作答时间和答案修改次数等题目层面的数据以及行动序列与操作间隔时间等操作层面的数据, 这类数据被统称为过程数据(process data; Hao et al., 2015; 袁建林 等, 2023)。深入分析过程数据有助于研究者了解被试的思维过程、解题策略和作答风格等。

目前对过程数据分析方法的研究仍处于探索阶段, 尚未得到广泛认可的明确方法分类。根据统

计分析的逻辑, 学者将现有的过程数据分析方法分为两类: 数据挖掘法和统计建模法(刘耀辉 等, 2022; 韩雨婷 等, 2022)。数据挖掘法是数据驱动的自下而上的探索性方法, 它采用统计学中算法模型(algorithm modeling culture)的思路从过程数据当中发掘出有意义的信息, 具体包括 N-grams、编辑距离等自然语言处理方法(Hao et al., 2015; He et al., 2021; Qiao & Jiao, 2018; Tang et al., 2020; Ulitzsch et al., 2022)和聚类分析、神经网络等统计学习方法(Chen et al., 2022; He et al., 2023; Wang et al., 2023; Xu et al., 2024)。统计建模法则是理论或模型驱动的自上而下的验证性方法, 它采用统计学中数据模型¹ (data modeling culture)的思路, 依据被试作答过程相关的理论假设来构建统计模型, 并借助模型估计被试的能力参数, 从而实现测量的目的。这种方法主要包括多维 IRT 模型、诊断分类模型等心理

收稿日期: 2024-11-06

* 国家自然科学基金项目(32071092)资助。

通信作者: 陈平, E-mail: pchen@bnu.edu.cn

¹ 数据模型核心理念是从数据中建模, 重视数据与变量之间的关系及其解释能力, 重在理解数据生成的机制, 揭示因果关系或特定变量的影响。算法模型核心理念是基于算法的预测和优化能力, 最大化模型的预测准确性或性能, 而不是专注于解释变量之间的关系。对于二者的比较详见 Breiman (2001)。

测量模型(Zhan & Qiao, 2022; 李美娟 等, 2020)和隐马尔可夫模型、动态贝叶斯网络等随机过程模型(Levy, 2019; Xiao et al., 2021)以及结合随机过程思想的测量模型(Chen, 2020; Han et al., 2022; LaMar, 2018; Shu et al., 2017; Wang & Liu, 2024; Xiao & Liu, 2024; 付颜斌 等, 2023)。相较于数据挖掘法, 统计建模法具有与传统心理与教育测量学相同的研究逻辑, 即根据潜在特质和外显行为之间的因果关系构建潜变量模型, 通过数据拟合对理论模型进行验证, 并通过估计模型参数获得对个体潜在特质水平的测量(Borsboom et al., 2003)。统计建模法着力于对数据和现象进行解释, 能够更好地揭示行为背后的心理学机制²。

在统计建模法中, 最具代表性的是结合随机过程思想的测量模型。这类模型综合了随机过程模型和传统心理测量模型的优势, 利用相对完整的过程数据信息获得能力估计值(韩雨婷 等, 2022)。针对具有有限问题状态的交互式测验情境, 这类模型将测验中的每个状态视为一道题目, 将每个状态下可执行的行为视为选项, 并由专家事先判断每个状态下每一种行为的正确性, 依此能够将人机交互过程转化为完成多道具有正确答案的“题目”的过程。另外, 在给定潜在能力的条件下, 这类模型将每名被试的反应过程都视为一个具有条件一阶马尔可夫性的离散时间的随机过程(建立起局部独立性), 并将被试在测验过程中的问题状态序列构建为“作答序列”, 进而使用心理测量模型进行分析。在 Shu 等人(2017)首次在研究中引入马尔可夫性假设后, 后续研究者均以此为基础来建构模型。但在交互式问题解决测验过程中, 问题情境并未一次性完整呈现, 被试需要进行探索从而逐渐积累信息。这使得被试当前状态的行为选择, 不仅受到其问题解决能力的影响, 还受到其对问题情境的了解程度的影响, 本研究将其称为学习效应(learning effect)。过去的研究均忽略了这一点。

以国际成人能力评估项目(Programme for the International Assessment of Adult Competencies, PIAAC)中技术增强环境下的问题解决能力测评

(Problem Solving in Technology-Rich Environments, PSTRE)的 Job Search 测验为例进行说明。如图 1 所示, 交互式测验界面包含丰富的信息, 点击各个链接/按钮后会打开不同的页面, 对应不同的问题状态。注意: 被试在测验开始时对问题情境的了解是不完整的, 他/她做出正确行为的概率受到其对问题情境的了解程度的影响, 因此未将这一点纳入模型考虑可能会使得对被试问题解决能力的估计出现偏差。

针对这一问题, 目前已有学者从不同角度对其进行研究。比如, Wang 等人(2023)综合自然语言处理、神经网络等统计学习方法提出子任务分析(subtask analysis), 将行为序列分割为对应不同子任务的子序列, 其中分离出被试对于问题情境的探索过程。Tang (2024)在隐马尔可夫模型的基础上引入潜变量, 通过隐状态也将行为序列切分为子序列, 并在其中识别出被试对于问题情境的探索阶段。但是, 现有研究均未从结合随机过程思想的测量模型这类方法的角度对该问题进行分析。鉴于此, 本文拟对现有的测量模型进行拓展, 引入当前状态在作答序列中的位置这一变量, 对被试在交互式测验过程中的学习效应进行建模, 以期获得更加准确的问题解决能力估计值。

接下来的内容按照以下结构进行组织: 首先, 对模型对应的问题情境进行界定并详细介绍拓展模型的建构思路, 说明其参数估计方法; 其次, 通过实证研究比较原模型与拓展模型的拟合表现, 并验证拓展模型在分析真实数据时的有效性; 再次, 开展模拟研究探索拓展模型的参数估计返真性及其影响因素; 最后, 对研究结果进行总结并在此基础上讨论当前研究局限以及未来可能的研究方向。

2 引入学习效应的过程数据测量模型

2.1 问题情境说明

对于具有有限情境的交互式测验, 被试的完整作答过程可以用状态序列来表征。假设有 n 名被试共完成 K 个交互式测验任务, 在任务 k 中共有 R 种可能的问题状态, 记为 $\mathbf{S}_k = \{s_1, s_2, \dots, s_R\}$ 。将被试 i 在任务 k 上的状态序列记为 $\mathbf{Y}_{ik} = \{Y_{ik1} = s_1, Y_{ik2} = s_2, \dots, Y_{ikj} = s_1, \dots, Y_{ikl_{ik}} = s_R\}$, 代表被试 i 在完成任务 k 的过程中在时刻 1 处于状态 s_1 , 在时刻 2 处于状态 s_2 , 在时刻 j 处于 s_1 , 在时刻 l_{ik} 处于状态 s_R , 其中 l_{ik} 表示被试 i 在任务 k 上状态序列的长度。此外, 记 $t_{iks} = j$ 为当前状态 $Y_{ikj} = s$ 在被试 i 完

² 从学科历史上看, 心理学研究的重点在于解释行为背后的因果机制, 因此在研究中通常采用数据模型的统计思路, 本文在这部分的论述中也遵循这种传统。当前, 以机器学习为代表的算法模型在各研究领域的应用快速增长, 对两种思路在心理学研究领域的比较已有研究探讨, 详见 Levy (2020)、Yarkoni 和 Westfall (2017)以及 Zheng 等人(2024)。

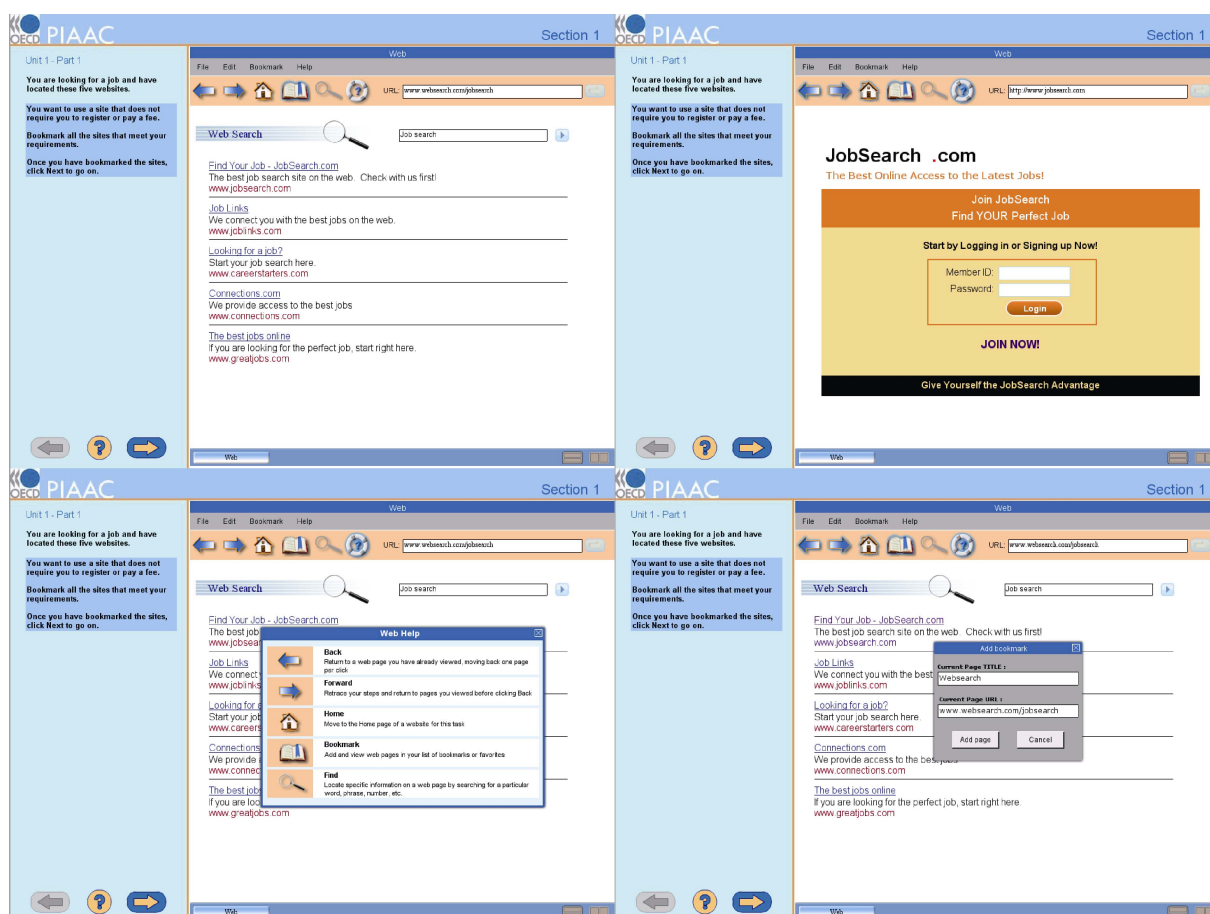


图1 PIAAC 中 PSTRE 的 Job Search 测验的部分界面

注: 图片来源为 <https://piaac-logdata.tba-hosting.de/public/problemsolving/JobSearchPart1/pages/jsp1-home.html>

成任务 k 的整个状态序列中的位置。比如, 在状态序列 ABCDEFG 中, 状态 C 对应 $t = 3$ 。

若被试 i 在时刻 j 处于状态 s , 在时刻 $j+1$ 处于状态 s' (即 $Y_{ikj} = s$ 且 $Y_{ik(j+1)} = s'$), 则将被试 i 在时刻 j 的状态转移记为 $Y_{ikj}^* = ss'$, 这时可进一步对应状态转移序列 $\mathbf{Y}_{ik}^* = \{Y_{ik1}^*, Y_{ik2}^*, \dots, Y_{ikj}^*, \dots, Y_{ik(l_k-1)}^*\}$ 。对每个状态转移定义一个正确性指标, 即将正确状态转移编码为 1, 错误状态转移编码为 0, 于是可以将 $\mathbf{Y}_{ik}^*$ 编码为二分的状态转移正确性序列 $\mathbf{A}_{ik} = \{A_{ik1}, A_{ik2}, \dots, A_{ikj}, \dots, A_{ik(l_k-1)}\}$ 。

2.2 模型构建

首先采用解释性项目反应理论³(Explanatory Item Response Theory Model, EIRT; De Boeck & Wilson, 2004; 陈冠宇, 陈平, 2019)的视角对现有的过程数据测量模型结构进行简要阐述, 以便更清晰地说明本研究提出模型的构建思路。现有各模型

对应的随机成分为多变量伯努利分布(multivariate Bernoulli distribution), 连接函数为基线类别 logit (baseline-categories logit), 而系统成分则建构了状态参数和能力参数(即状态和被试水平的截距)用于解释被试在各状态下做出正确状态转移的概率。

而在交互式测验的过程中, 被试开始作答时获得的问题情境信息是不完全的, 他/她并不掌握问题解决所需要的完全信息, 需要在与测验系统交互的过程中逐渐获取与问题解决相关的信息。被试对这些信息的掌握程度会影响其在测验过程中做出正确状态转移的概率。由于无法直接测量得到被试对问题情境信息的掌握程度, 因此本研究提出在模型中纳入当前状态在作答序列中的位置这一变量作为对问题情境信息掌握程度的间接表征。当前状态的位置从宏观上反映了被试对问题解决任务情境探索的程度: 随着交互的进行, 被试在过程中会逐渐积累问题情境背景的相关信息。

基于上述分析, 我们对现有的过程数据测量模型进行拓展。这里我们选择付颜斌等人(2023)提出的单参数行动序列模型(One-Parameter Action

³ EIRT 偏向于从“解释”的角度来看待模型, 将模型分解为随机成分(random component)、连接函数(link function)和系统成分(systematic component), 其中系统成分通常包含题目特征和被试特征来解释作答反应如何受到这些变量的影响。

Sequence Model, 1P-ASM)进行拓展,拓展思路也可直接迁移至其他模型。1P-ASM 是在 Han 等人 (2022)提出的序列反应模型(Sequential Response Model, SRM)的基础上约简得到,它在减少模型复杂度的同时,还能提供与原模型接近的能力估计精度。公式(1)给出的是 1P-ASM 的表达式(付颜斌 等, 2023),公式(2)呈现的是 1P-ASM 在 EIRTM 框架下的等价表达式:

$$p(A_{ikj} = 1|Y_{ikj} = s, \theta_i, \beta_{ks}) = \frac{\exp(\beta_{ks} + \theta_i)}{1 + \exp(\beta_{ks} + \theta_i)} \quad (1)$$

$$A_{ikj} \sim \text{Bernoulli}(\pi_{ikj}),$$

$$\pi_{ikj} = p(A_{ikj} = 1|Y_{ikj} = s, \theta_i, \beta_{ks}), \quad (2)$$

$$\text{logit}(\pi_{ikj}) = \log\left(\frac{\pi_{ikj}}{1 - \pi_{ikj}}\right) = \beta_{ks} + \theta_i,$$

其中, β_{ks} 为任务 k 中状态 s 的容易度参数,表示在任务 k 中状态 s 下做出正确状态转移的容易程度; θ_i 为被试 i 的能力参数。

在 1P-ASM 的基础上,加入当前状态在作答序列中的位置 t_{iks} 这一变量,将其斜率 γ_k 定义为任务 k 的学习效应参数,得到拓展模型:

$$p(A_{ijk} = 1|Y_{ijk} = s, \theta_i, \gamma_k, \beta_{ks}) = \frac{\exp(\beta_{ks} + \theta_i + \gamma_k \cdot t_{iks})}{1 + \exp(\beta_{ks} + \theta_i + \gamma_k \cdot t_{iks})} \quad (3)$$

具体来说,拓展模型在 β_{ks} 和 θ_i 的基础上,在模型的系统成分中加入协变量 t_{iks} 。 t_{iks} 与原有参数(即 β_{ks} 和 θ_i)共同解释被试在各状态下做出正确状态转移的概率。在测验过程中,随着与计算机的不断交互,被试会逐渐积累对于问题情境的有关知识。因此,当前状态在作答序列中的位置这一变量能够用来衡量被试在测验过程中所积累的信息量,其斜率 γ_k 则代表在任务 k 中这种信息积累的速度。需要说明的是,学习效应在理论上可能受到题目情境设置的影响(表现为不同任务中学习情境知识的难度不同),也可能受到被试学习能力的影响。两者分别对应将参数定义在任务水平与被试水平上。本研究将新增参数 γ_k 定义在任务水平上,主要是考虑到模型能够在捕捉学习效应的同时还具有最大程度的简洁性,也考虑到实证研究中使用的任务情境较易(被试间的情境学习能力水平不存在较大差异)。实际上, γ_k 也可以被解释为被试全体在任务 k 上的学习能力。为叙述方便,本文将公式(3)对应的拓展模型记为 1P-ASM-R。

此外,值得注意的是:理论上,被试在交互式

测验中对问题情境相关知识的积累速度应当不是始终不变的,而是“在任务一开始时能够获得更多信息,当有了一定积累后,再能够从情境中习得的知识增量也会逐渐下降”,即对问题情境知识的学习和积累存在边际收益递减(diminishing marginal returns)。这在统计中体现为被试积累的信息量与互动次数关系的函数为凹函数(即二阶导数小于 0)。因此,在 1P-ASM-R 的基础上,本研究进一步建构 1P-ASM-R*模型来考虑学习效应的边际收益递减:

$$p(A_{ijk} = 1|Y_{ijk} = s, \theta_i, \gamma_k, \beta_{ks}) = \frac{\exp(\beta_{ks} + \theta_i + \gamma_k \cdot t_{iks}^*)}{1 + \exp(\beta_{ks} + \theta_i + \gamma_k \cdot t_{iks}^*)} \quad (4)$$

其中 $t_{iks}^* = \sqrt{t_{iks}}$, 即对 t_{iks} 进行平方根转换后再纳入模型。这种做法基于以下两点考虑:(1)平方根函数是一个典型的凹函数,符合边际收益递减的理论假设,是经济学中常用的收益函数之一;(2)平方根函数的结构简单,它在能够有效捕捉收益变化的非线性特征的同时,还便于后续的分析 and 计算。

2.3 参数估计

采用贝叶斯方法对模型参数进行估计,具体采用哈密尔顿蒙特卡洛(Hamiltonian Monte Carlo, HMC)方法。相比于传统的基于马尔可夫链的采样策略,HMC 利用哈密顿力学中的概念,通过模拟物理系统中的轨迹来探索参数空间,使用梯度信息指导采样过程,从而实现高维空间中的高效采样。模型参数的后验分布为:

$$p(\theta, \gamma, \beta | Y, \mathcal{R}) \propto p(Y | \theta, \beta, \mathcal{R}) p(\theta, \gamma, \beta) = \prod_{i=1}^n \prod_{k=1}^K \prod_{j=1}^{l_{ik}-1} p(A_{ijk} | Y_{ijk} = s, \theta_i, \gamma_k, \beta_{ks}) p(\theta_i) p(\gamma_k) p(\beta) \quad (5)$$

其中 $p(\theta_i)$ 、 $p(\beta)$ 和 $p(\gamma_k)$ 分别为 θ_i 、 β 和 γ_k 的先验分布。参考前人研究(Fu et al., 2024; Xiao & Liu, 2024)对参数先验分布的设置,这里将 $p(\theta_i)$ 设为标准正态分布 $\mathcal{N}(0, 1)$; 将 $p(\beta)$ 设为标准多元正态分布 $\mathcal{MVN}(0, \mathbf{E})$, 其中 $\mathbf{E}$ 为单位矩阵; 将 $p(\gamma_k)$ 设为标准正态分布 $\mathcal{N}(0, 1)$ 。

3 实证数据分析

3.1 数据描述

研究数据来源于 PISA 2012 的计算机化问题解决测验 TICKETS 模块(CP038),在该任务模块中被试需要按照要求操作一台虚拟的铁路售票机购买车票(OECD, 2013)。虚拟测试系统采用有限状态自

动机(finite state automata)的框架搭建, 其初始界面如图2所示。

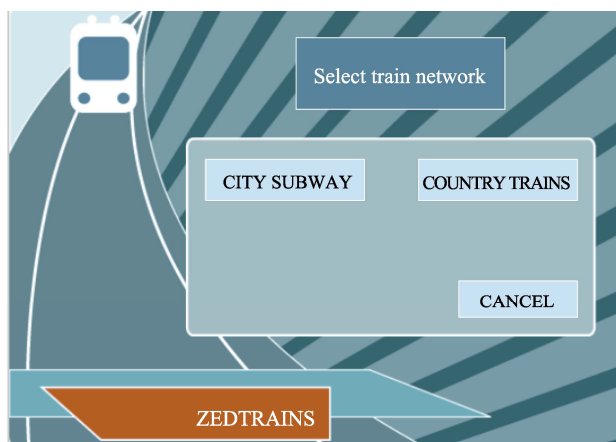


图2 PISA 2012 TICKETS 任务界面

在这台虚拟售票机上, 被试首先需要选择购买的铁路网络类型: 城市地铁(CITY SUBWAY)或乡村铁路(COUNTRY TRAINS); 其次选择购票的优惠类型: 全价票(FULL FARE)或优惠票(CONCESSION); 而后再选择需要购买的车票类型: 日票(DAILY)或次票(INDIVIDUAL); 最后在购买界面上会显示票价, 如果是次票则需要被试进一步选择购买的车票张数。被试在每个界面都可以点击取消(CANCEL)

回到系统的初始界面, 在购买界面点击购买(BUY)则完成测验任务。

这个问题解决任务模块共有3个子任务, 本研究对其中的一个子任务(CP038Q02)进行分析, 研究方法可以拓展至多个任务联合分析的情境中。在该任务中, 被试需要购买两张城市地铁的优惠次票。该问题情境具有有限数量的可能状态, 本研究采用 Han 等人(2022)的设置, 将任务情境分解为11种可能的问题状态, 依据问题情境共有27种可能的状态转移。如图3所示, 在11个问题状态中, 初始状态记为A, 结束状态记为K, 其余均为问题解决过程的中间状态。状态间的箭头表示状态转移可能的方向, 其中实线表示正确状态转移, 虚线表示错误状态转移, 而线段上的文字则标注状态转移所对应的操作。比如, 从初始状态A转移到正确的铁路网络状态B, 这一转移是正确的状态转移, 被试在初始状态A选择“COUNTRY TRAINS”会实现该状态转移。此外, 从图中可以看出, 该问题解决任务的最优操作序列为: (初始状态A)→ COUNTRY TRAINS →(正确的铁路网络状态B)→ FULL FARE →(正确的购票优惠状态C)→ INDIVIDUAL →(正确的车票类型状态D)→ 2 trips →(正确的乘车次数状态E)→ BUY →(结束状态K)。

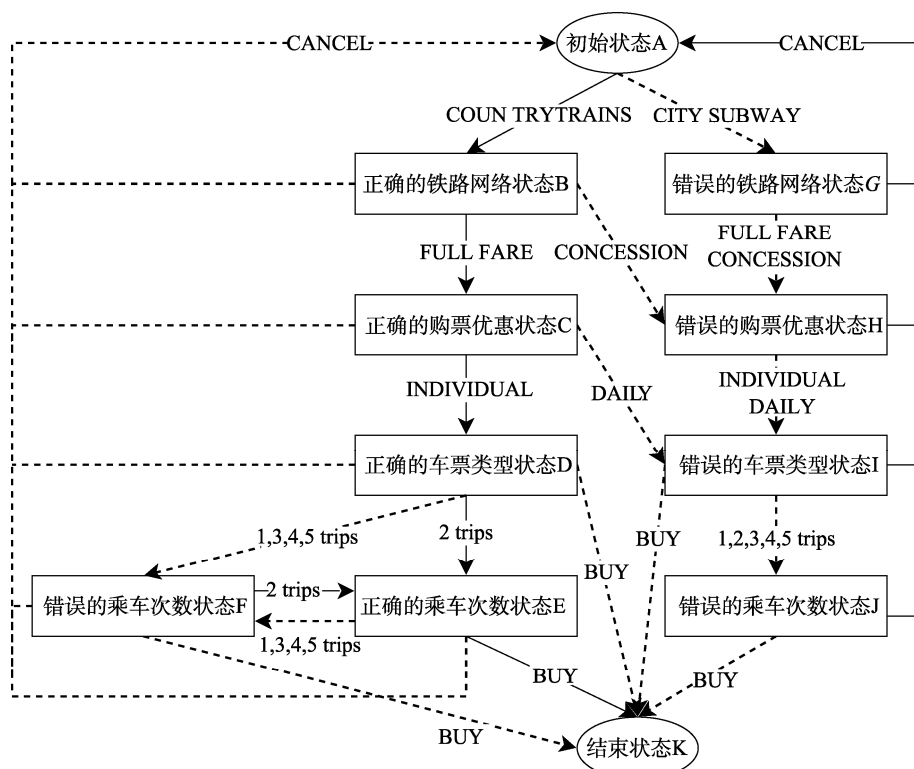


图3 问题解决任务 CP038Q02 的问题状态图

数据从 PISA 官方网站⁴上获得。根据对 CP038Q02 的问题状态的定义,对原始数据进行重新编码,并参考付颜斌等人(2023)的标准对数据进行清理:(1)删除提前终止作答的序列,如未点击购票就结束任务的序列;(2)删除包含不可能存在的状态转移模式的序列。最后,得到 27,711 名被试的行动序列,其中最小长度为 5,最大长度为 110,平均长度为 6.97;包含 1330 种行动序列,其中 542 种行动序列完成问题解决任务(对应 14925 名被试完成任务,其中 10348 名被试按照最优行动序列完成任务)。考虑到算力资源与研究效率的问题,本研究采用简单随机抽样的方法从 27,711 名被试中抽取 4,000 名被试的过程数据用于分析,样本中行动序列的最小长度为 5,最大长度为 80,平均长度为 6.98;包含 339 种行动序列,其中 144 种行动序列完成问题解决任务(对应 2115 名被试,其中 1472 名按照最优行动序列完成任务)。

3.2 分析方法

分别采用 1P-ASM、1P-ASM-R 和 1P-ASM-R* 拟合数据,对三者的表现进行比较分析,并评估在模型中加入学习效应以及在拓展模型中考虑学习效应边际递减的合理性。

模型的参数估计和拟合检验通过 Stan(Carpenter et al., 2017)及对应的 R 语言程序包 rstan(Stan Development Team, 2024)与 loo(Vehtari et al., 2017)自编代码实现⁵。在进行参数估计时,设置 4 条采样链,每条链长 3000 次,其中预热(warm-up)2000 次,最终每个参数得到 4000 个采样点。所有程序均在配置为 Intel® Xeon® Ice Lake @ 3.2GHz / 3.5GHz 和 32G 内存的服务器上运行。

链的收敛性通过潜在比例缩减因子 R (Potential Scale Reduction Factor, PSRF; Gelman & Rubin, 1992)进行评价。依据拇指原则,若 $R < 1.1$,则认为模型参数估计收敛。在模型参数估计收敛的基础上,再使用后验预测检验(Posterior Predictive Check, PPC; Gelman et al., 1996)评估模型的绝对拟合情况。PPC 从后验预测分布中抽取样本,将这些样本与观测到的数据进行对比,然后计算后验预测 p 值(Posterior Predictive p-values, PPP),以此来评估模型适用或失效之处。如果模型对数据的拟合良好,那么 PPP 值将接近 0.5 (Gelman et al., 2014)。

本研究对观测数据与抽样数据中得到的状态转移正误进行比较,计算得到模型的 PPP 值(付颜斌等, 2023)。对于模型的相对拟合情况,采用留一法交叉验证(Leave-One-Out cross-validation, LOO; Gelfand et al., 1992)和 Watanabe-Akaike 信息准则(Watanabe-Akaike Information Criterion, WAIC; Watanabe, 2010)两个指标进行分析。LOO 和 WAIC 越小,表明模型对数据的拟合程度越好。

3.3 结果

三个模型中所有参数的 PSRF 值均小于 1.1,说明三个模型的贝叶斯估计均收敛,可以依据贝叶斯参数估计的结果进行进一步的分析。

表 1 呈现了三个模型对数据的拟合情况。结果显示:三个模型的 PPP 值均接近 0.5,说明三个模型均拟合该数据。LOO 和 WAIC 的结果表明:与 1P-ASM 相比,1P-ASM-R 和 1P-ASM-R* 均能更好地拟合数据,意味着在模型中纳入当前状态在作答序列中的位置变量并考虑被试的学习效应能更好地反映数据的情况,即被试在作答过程中存在对问题情境的学习。进一步对 1P-ASM-R 和 1P-ASM-R* 的结果进行比较,可以看出 1P-ASM-R* 的拟合表现优于 1P-ASM-R。表 2 呈现了 1P-ASM-R 和 1P-ASM-R* 模型中学习效应参数 γ 的估计结果,可以发现:1P-ASM-R 中的学习效应参数接近于 0,未体现出被试对问题情境的学习。上述结果说明,在模型中考虑对问题情境学习的边际效应的影响更符合实际情况。因此,接下来仅对比 1P-ASM 和 1P-ASM-R* 的结果。

表 1 三种模型对数据的拟合情况

模型	LOO	WAIC	PPP
1P-ASM	21647.9	21626.5	0.520
1P-ASM-R	21637.7	21616.4	0.496
1P-ASM-R*	21595.2	21573.5	0.507

注:1P-ASM 为单参数行动序列模型,1P-ASM-R 为引入学习效应的单参数行动序列拓展模型,1P-ASM-R* 为考虑学习边际效应的单参数行动序列拓展模型;LOO = 留一法交叉验证,WAIC = Watanabe-Akaike 信息准则,PPP = 后验预测概率。

表 2 1P-ASM-R 和 1P-ASM-R* 模型中学习效应参数的后验估计结果

模型	均值	SD	95%HPDL	95%HPDU
1P-ASM-R	0.010	0.004	0.002	0.018
1P-ASM-R*	0.158	0.029	0.101	0.214

注:1P-ASM-R 为引入学习效应的单参数行动序列拓展模型,1P-ASM-R* 为考虑学习边际效应的单参数行动序列拓展模型;SD 为标准差,95%HPDL 为 95%最高后验密度(highest posterior density)区间的下界,95%HPDU 为 95%最高后验密度区间的上界。

⁴ <https://www.oecd.org/pisa/pisaproducts/database-cbapisa2012.html>

⁵ 研究程序代码及数据已上传至 https://osf.io/nb9ug/?view_only=0b444c6148ed4942ae984386a7c21b26

表 3 1P-ASM 和 1P-ASM-R* 的容易度参数的后验估计结果

状态	1P-ASM				1P-ASM-R*			
	均值	SD	95%HPDL	95%HPDU	均值	SD	95%HPDL	95%HPDU
A	0.872	0.033	0.809	0.935	0.667	0.049	0.568	0.764
B	1.606	0.048	1.511	1.699	1.331	0.068	1.198	1.459
C	1.415	0.052	1.316	1.515	1.093	0.076	0.947	1.244
D	1.452	0.059	1.339	1.567	1.089	0.088	0.916	1.266
E	1.940	0.074	1.800	2.088	1.536	0.101	1.336	1.733
F	0.329	0.118	0.099	0.553	-0.086	0.138	-0.353	0.180
G	-1.721	0.078	-1.875	-1.571	-1.952	0.091	-2.130	-1.773
H	-2.086	0.083	-2.252	-1.927	-2.375	0.100	-2.571	-2.182
I	-0.860	0.053	-0.964	-0.757	-1.201	0.082	-1.364	-1.047
J	-0.398	0.111	-0.617	-0.178	-0.798	0.130	-1.056	-0.544

注: 1P-ASM 为单参数行动序列模型, 1P-ASM-R* 为考虑学习边际效应的单参数行动序列拓展模型; SD 为标准差, 95%HPDL 为 95% 最高后验密度区间的下界, 95%HPDU 为 95%最高后验密度区间的上界。

表 3 中呈现的是 1P-ASM 和 1P-ASM-R* 的容易度参数 β_{ks} 的后验估计结果。结果显示, 相较于 1P-ASM 的容易度参数, 1P-ASM-R* 的容易度参数均有一定程度的降低, 这是由于在模型中纳入当前状态在作答序列中的位置变量后, 新增的变量解释了一部分原本全部由容易度参数来解释的变异。值得注意的是, 1P-ASM-R* 的容易度参数仍然保留原 1P-ASM 模型的参数关系, 比如在正确问题解决路径上(A、B、C、D 和 E)的问题容易度参数都大于 0。这表明当被试处于正确路径上的问题状态时, 其更容易继续呈现正确状态转移(Han et al., 2022; 付颜斌 等, 2023), 而纳入位置变量后, 模型依然能够捕捉到这一问题解决任务情境特点。

图 4 呈现了 1P-ASM 和 1P-ASM-R* 的问题解决能力参数(θ_i)估计后验均值的对比散点图及直方图。散点图的结果表明: 两种模型得到的问题解决能力估计值之间具有很高的相关性, 相关系数高达 0.998。从直方图和概率密度曲线也可以看出, 两个能力参数的概率密度分布也近似。值得注意的是: 在低问题解决能力区间中, 1P-ASM-R* 模型估计得到的能力参数值则略低于 1P-ASM, 说明: (1)在考虑学习效应后, 相对于高问题解决能力被试组, 低问题解决能力被试组的问题解决能力会得到更多的校正; (2)相对于具有高问题解决能力的被试, 低问题解决能力被试的能力参数估计受学习效应的影响更大。

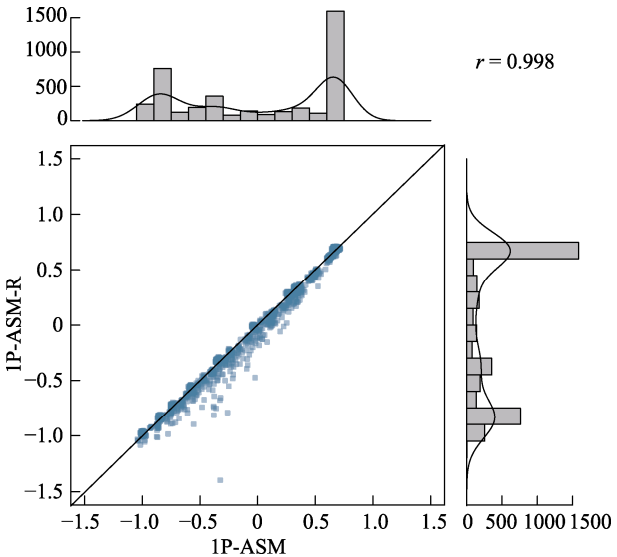


图 4 1P-ASM 和 1P-ASM-R* 的能力参数估计后验均值的对比散点图及直方图

注: 1P-ASM 为单参数行动序列模型, 1P-ASM-R* 为引入学习效应的单参数行动序列拓展模型; r 为 Pearson 相关系数。

4 模拟研究

4.1 研究设计与数据生成

通过模拟研究进一步比较 1P-ASM 和 1P-ASM-R* 两种模型在各种条件下的心理测量学表现, 并考查样本量(200 和 1000)、序列长度(短和长)和学习效应大小(0、0.1 和 0.3)三个操纵变量对参数估计返真性的影响。采用的模拟情境设计与实证研究相同。其中, 参考付颜斌等人(2023)的模拟研究设计, 通过控制模拟任务中的被试返回初始点

(即点击“CANCEL”)的概率来操纵序列长度。最终生成的短行为序列和长行为序列的平均长度分别为 7.65 和 14.80。

上述三个操纵变量共产生 2 (样本量: 200/1000) × 2 (序列长度: 短/长) × 3 (学习效应大小: 0/0.1/0.3) = 12 种模拟条件。为减少随机误差的影响, 每种模拟条件下均重复生成 50 组数据。在每次重复中, 状态转移参数的真值不变, 被试能力的真值从标准正态分布中随机生成。根据 1P-ASM-R*通过 R 语言代码模拟生成反应序列, 其中当学习效应为 0 时, 相当于采用 1P-ASM 生成反应序列。值得注意的是, 通过 1P-ASM 和 1P-ASM-R*无法直接生成被试解决任务所呈现的行为序列, 只能生成正确与错误操作对应的 0-1 向量。因此, 在研究中生成及分析的数据均为状态转移正确性序列。

4.2 分析方法

分别采用 1P-ASM 和 1P-ASM-R*拟合数据, 进行参数估计, 估计过程与实证研究保持一致。通过 PSRF 对估计程序的收敛性进行评价, 在程序收敛的基础上进一步从两方面来评估模型的表现: (1)采用 LOO 和 WAIC 两个指标对模型的整体拟合情况进行分析比较; (2)使用偏差(bias)和均方根误差(Root Mean Squared Error, RMSE)量化问题解决能力参数的估计返真性:

$$bias = \frac{1}{R} \frac{1}{n} \sum_{r=1}^R \sum_{i=1}^n (\hat{\theta}_{ri} - \theta_{ri}) \quad (6)$$

$$RMSE = \sqrt{\frac{1}{R} \frac{1}{n} \sum_{r=1}^R \sum_{i=1}^n (\hat{\theta}_{ri} - \theta_{ri})^2} \quad (7)$$

其中 R 为重复模拟的总次数, $\hat{\theta}_{ri}$ 和 θ_{ri} 分别是被试 i 在第 r 次重复时的能力估计值和真实值。

4.3 结果

在所有条件下, 两个模型中所有参数的 PSRF 值均小于 1.1, 说明所有模型参数的估计均收敛, 可以依据当前参数估计结果进行后续分析。

表 4 呈现了不同模拟条件下两个模型的拟合结果, 提供的是每种条件下 50 次重复中 1P-ASM-R*拟合优于 1P-ASM 的百分比。图 5 中的直方图呈现的是重复模拟中两个模型的拟合指标差值 (1P-ASM-R*减去 1P-ASM) 的分布情况。从表/图的结果来看: 首先, 在不存在学习效应时, 1P-ASM 的拟合结果更有可能优于 1P-ASM-R*, 不过两者在实际指标上的差异很小。因此可以认为, 在不存在

学习效应时, 原模型与拓展模型对模拟数据的拟合情况相近; 而在存在学习效应时, 1P-ASM-R*的拟合结果更有可能好于 1P-ASM, 而且随着学习效应增大, 这种优势会增加。其次, 样本量越大或行动序列越长, 学习效应带来的 1P-ASM-R*拟合优势也会增加。

表 4 1P-ASM 和 1P-ASM-R*的模型拟合结果

样本量	序列长度	学习效应大小	LOO	WAIC
200	短	0	38%	38%
		0.1	70%	72%
		0.3	98%	98%
	长	0	26%	26%
		0.1	100%	100%
		0.3	100%	100%
1000	短	0	28%	28%
		0.1	94%	94%
		0.3	100%	100%
	长	0	26%	26%
		0.1	100%	100%
		0.3	100%	100%

注: 1P-ASM 为单参数行动序列模型, 1P-ASM-R*为引入边际学习效应的单参数行动序列拓展模型; LOO = 留一法交叉验证, WAIC = Watanabe-Akaike 信息准则。

表 5 呈现的是能力参数估计的返真性情况。从表中可以发现: (1)被试样本量对能力参数估计不存在影响。样本量的大小不影响两个模型的参数估计返真性; (2)当存在学习效应时, 1P-ASM-R*的参数估计返真性优于 1P-ASM, 这一优势主要体现在“长序列、大的学习效应”条件下。此外, 网络版附录给出了容易度参数估计的返真性情况, 从总体趋势上看, 当存在学习效应时, 1P-ASM-R*的容易度参数估计返真性也要优于 1P-ASM。

表 6 提供了学习效应参数估计的返真性结果。总的来看, 1P-ASM-R*能够得到准确和可靠的学习效应参数估计值。对比三个操纵变量不同水平下的结果, 可以发现: (1)被试样本量对学习效应参数估计存在一定的影响。样本量越大, 参数估计返真性更高; (2)从总体趋势上看, 序列长度越长, 学习效应参数估计的返真性越高, 这一影响主要体现在 RMSE 指标上; (3)学习效应参数估计的返真性受学习效应参数真值的影响不大。

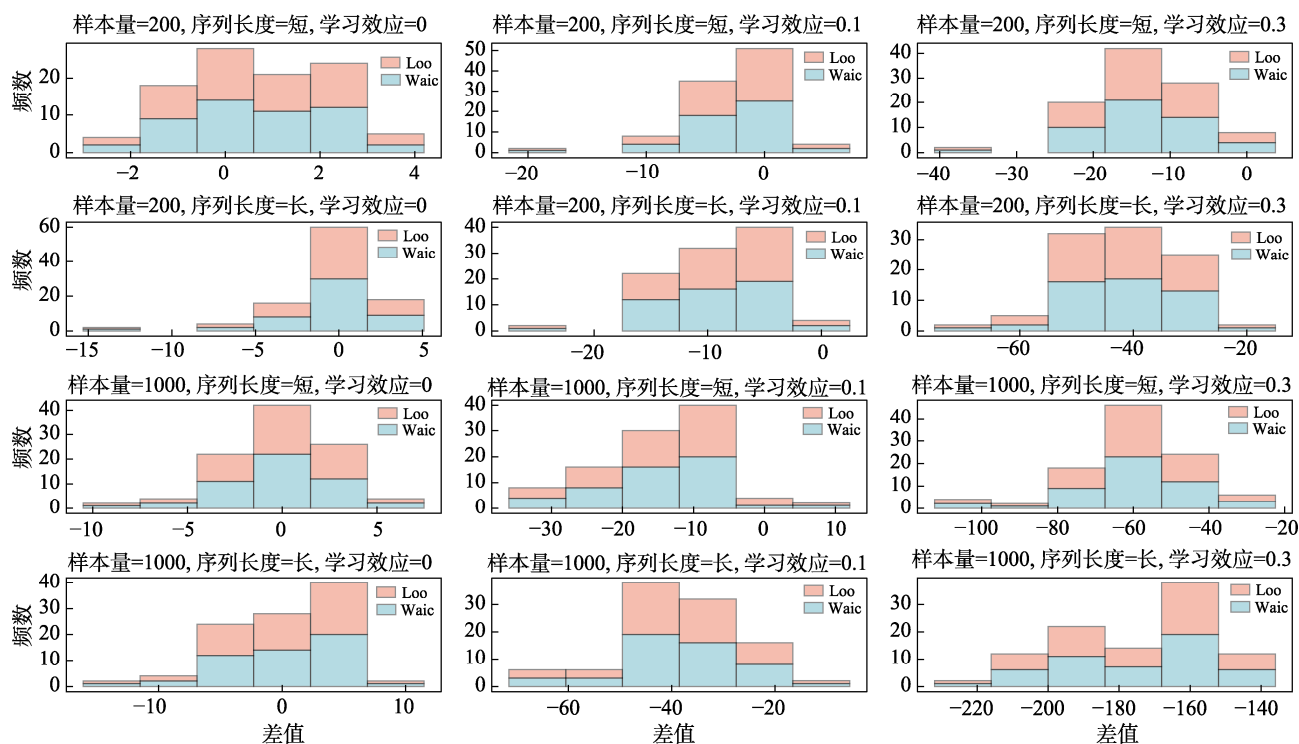


图5 重复模拟中 1P-ASM-R*和 1P-ASM 的拟合指标差值分布

表5 能力参数估计的返真性情况

样本量	序列长度	学习效应大小	Bias		RMSE	
			1P-ASM	1P-ASM-R*	1P-ASM	1P-ASM-R*
200	短	0	0.018	0.017	0.580	0.581
		0.1	0.025	0.015	0.584	0.583
		0.3	0.040	0.017	0.601	0.596
	长	0	0.027	0.027	0.449	0.449
		0.1	0.040	0.026	0.472	0.464
		0.3	0.058	0.025	0.524	0.501
1000	短	0	-0.023	-0.023	0.586	0.586
		0.1	-0.021	-0.022	0.590	0.589
		0.3	-0.023	-0.017	0.609	0.605
	长	0	-0.020	-0.020	0.447	0.447
		0.1	0.012	-0.009	0.471	0.466
		0.3	-0.012	0.005	0.525	0.502

注: 1P-ASM 为单参数行动序列模型, 1P-ASM-R*为引入边际学习效应的单参数行动序列拓展模型。

表6 学习效应参数估计的返真性情况

样本量	序列长度	学习效应大小	$\bar{\gamma}$	Bias	RMSE
200	短	0	0.013	0.013	0.006
		0.1	0.128	0.028	0.009
		0.3	0.334	0.034	0.009
	长	0	0.003	0.003	0.001
		0.1	0.120	0.020	0.001
		0.3	0.335	0.035	0.002
1000	短	0	0.002	0.002	0.001
		0.1	0.104	0.004	0.002
		0.3	0.312	0.012	0.002
	长	0	-0.005	-0.005	0.000
		0.1	0.103	0.003	0.000
		0.3	0.303	0.003	0.000

注: $\bar{\gamma}$ 为模拟研究中 50 次重复得到的 1P-ASM-R*的学习效应参数估计值的均值。

5 总结与讨论

相较于传统的作答数据, 过程数据蕴含着更丰富的信息, 为研究者推断在计算机动态测验中测量的被试复杂能力真实水平提供了更加细致的测量证据, 有助于实现对被试能力的全面评估。在已有研究中, 研究者结合随机过程的思想开发一系列过

程数据测量模型, 但它们均基于一阶马尔可夫性假设, 未考虑到“被试在进行问题解决时需要未完整呈现的问题情境进行探索和学习”的情况。鉴于被试在进行交互式问题解决任务时会逐步探索问题情境, 本文在 1P-ASM 模型的基础上, 引入当前状态在作答序列中的位置这一变量 t_{iks} 以及学习效应参数 γ_k , 得到能描述不同任务情况下被试群体

水平学习效应的拓展模型 1P-ASM-R*。

实证研究结果发现：(1)相较于 1P-ASM, 1P-ASM-R*在实证数据上有更佳的拟合表现, 说明被试在作答过程中确实存在对问题情境的学习；(2)学习效应具有边际递减的特点, 即随着被试逐渐深入探索待解决问题的情境, 对问题情境的学习对正确行动选择的影响增量会逐渐减少；(3)引入学习效应并不影响模型捕捉问题解决任务的特征。当被试已经在问题解决路径上时, 他们更倾向于保持在这一正确的轨道上；而一旦进入错误的解决路径, 他们也更容易在错误的方向上继续前进；(4)在考虑学习效应后, 相对于高问题解决能力被试组, 低问题解决能力被试组的问题解决能力会得到更多的校正。另外, 模拟研究结果表明：(1)当存在学习效应时, 1P-ASM-R*能够提供与 1P-ASM 近似的拟合效果, 而且能够正确地将学习效应参数估计为 0。而当存在学习效应时, 1P-ASM-R*能够更好地拟合数据, 而且这种优势在学习效应越强时越明显；(2)序列长度是影响 1P-ASM-R*参数返真性的原因之一。序列长度越长, 数据中所含信息就越多, 因此对参数的估计也就更准。

尽管本文提出一种可以有效分析问题解决测验中学习效应、提高问题解决能力估计精度的模型, 但仍存在一些不足值得今后进一步探索。

首先, 在本研究中, 学习效应被定义在任务水平上, 反映了“任务”的特征, 因此学习效应大小对所有被试而言都是一致的。这样设计的目的是为了模型能够在捕捉到学习效应的同时具有最大程度的简洁性。但也有研究指出, 问题解决本身就包括建立问题表征(problem representation)和生成解决方案(generating problem solution)两个认知过程(Goode & Beckmann, 2010; Novick & Bassok, 2005)。若从这一理论视角着手, 本研究拓展模型中的问题解决能力参数所代表的实际上是狭义的“生成解决方案能力”, 而任务水平的学习效应参数则是被试全体在该任务上“建立问题表征”的能力。因此, 未来的研究可考虑进一步在被试水平上对学习效应进行建模, 从而获得被试水平的学习能力。

其次, 本研究选择当前状态在作答序列中的位置这一较为笼统的变量来刻画被试的学习过程, 而过程数据中不仅包含被试解决问题的行动顺序, 还记录进行各种行为的时间戳(time stamp)。后续研究可以尝试利用行为时间戳信息记录的被试在特定状态上的停留时间, 从而对被试的问题情境学习进

行更为深入的分析。

再次, 本研究根据边际效益递减将学习轨迹界定为凹函数, 从简便性角度出发具体探索了平方根函数在模型应用中的效果。未来研究也可进一步根据问题解决任务的特征有针对性地设计学习轨迹。另外, 在当前凹函数形式的学习轨迹基础上, 还可以进一步探索其他凹函数(如对数函数)在拓展模型中的应用。

最后, 虽然本研究所采用的 PISA 2012 的 TICKET 问题解决任务是当前过程数据研究领域的经典素材, 但其任务结构相对简单, 未来可以尝试采用更加复杂的问题解决任务数据来对这一问题进行更加深入的研究。

参 考 文 献

- Borsboom, D., Mellenbergh, G. J., & Van Heerden, J. (2003). The theoretical status of latent variables. *Psychological Review*, 110(2), 203–219.
- Breiman, L. (2001). Statistical modeling: The two cultures. *Statistical Science*, 16(3), 199–215.
- Carpenter, B., Gelman, A., Hoffman, M. D., Lee, D., Goodrich, B., Betancourt, M., ... Riddell, A. (2017). Stan: A probabilistic programming language. *Journal of Statistical Software*, 76(1), 1–32.
- Chen, G., & Chen, P. (2019). Explanatory item response theory models: Theory and application. *Advances in Psychological Science*, 27(5), 937–950.
- [陈冠宇, 陈平. (2019). 解释性项目反应理论模型: 理论与应用. *心理科学进展*, 27(5), 937–950.]
- Chen, Y. (2020). A continuous-time dynamic choice measurement model for problem-solving process data. *Psychometrika*, 85(4), 1052–1075.
- Chen, Y., Zhang, J., Yang, Y., & Lee, Y. (2022). Latent space model for process data. *Journal of Educational Measurement*, 59(4), 517–535.
- De Boeck, P., & Wilson, M. (2004). *Explanatory item response models: A generalized linear and nonlinear approach*. New York, NY: Springer.
- Fu Y., Chen, Q., & Zhan, P. (2023). Binary modeling of action sequences in problem-solving tasks: One- and two-parameter action sequence model. *Acta Psychologica Sinica*, 55(8), 1383–1396.
- [付颜斌, 陈琦鹏, 詹沛达. (2023). 问题解决任务中行动序列的二分类建模: 单/两参数行动序列模型. *心理学报*, 55(8), 1383–1396.]
- Fu, Y., Zhan, P., Chen, Q., & Jiao, H. (2024). Joint modeling of action sequences and action time in computer-based interactive tasks. *Behavior Research Methods*, 56(5), 4293–4310.
- Gelfand, A. E., Dey, D. K., & Chang, H. (1992). *Model Determination using predictive distributions with implementation via sampling-based methods* (technical report, No. SOL ONR 462). Palo Alto, CA: Department of Statistics, Stanford University.
- Gelman, A., Carlin, J. B., Stern, H. S., Dunson, D. B., Vehtari, A., & Rubin, D. B. (2013). *Bayesian data analysis* (3rd ed.). Boca Raton, FL: Chapman & Hall/CRC Press.
- Gelman, A., Meng, X.-L., & Stern, H. (1996). Posterior

- predictive assessment of model fitness via realized discrepancies. *Statistica Sinica*, 6(4), 733–760.
- Gelman, A., & Rubin, D. B. (1992). Inference from iterative simulation using multiple sequences. *Statistical Science*, 7(4), 457–472.
- Goode, N., & Beckmann, J. F. (2010). You need to know: There is a causal relationship between structural knowledge and control performance in complex problem solving tasks. *Intelligence*, 38(3), 345–352.
- Han, Y., Liu, H., & Ji, F. (2022). A sequential response model for analyzing process data on technology-based problem-solving tasks. *Multivariate Behavioral Research*, 57(6), 960–977.
- Han, Y., Xiao, Y., & Liu, H. (2022). Feature extraction and ability estimation of process data in the problem-solving test. *Advances in Psychological Science*, 30(6), 1393–1409.
- [韩雨婷, 肖悦, 刘红云. (2022). 问题解决测验中过程数据的特征抽取与能力评估. *心理科学进展*, 30(6), 1393–1409.]
- Hao, J., Shu, Z., & von Davier, A. (2015). Analyzing process data from game/scenario-based tasks: An edit distance approach. *Journal of Educational Data Mining*, 7(1), 33–50.
- He, Q., Borgonovi, F., & Paccagnella, M. (2021). Leveraging process data to assess adults' problem-solving skills: Using sequence mining to identify behavioral patterns across digital tasks. *Computers & Education*, 166, Article 104170.
- He, Q., Borgonovi, F., & Suárez-Álvarez, J. (2023). Clustering sequential navigation patterns in multiple-source reading tasks with dynamic time warping method. *Journal of Computer Assisted Learning*, 39(3), 719–736.
- LaMar, M. M. (2018). Markov decision process measurement model. *Psychometrika*, 83(1), 67–88.
- Levy, R. (2019). Dynamic Bayesian network modeling of game-based diagnostic assessments. *Multivariate Behavioral Research*, 54(6), 771–794.
- Levy, R. (2020). Implications of considering response process data for greater and lesser psychometrics. *Educational Assessment*, 25(3), 218–235.
- Li, M., Liu, Y., & Liu, H. (2020). Analysis of the problem-solving strategies in computer-based dynamic assessment: The extension and application of multilevel mixture IRT model. *Acta Psychologica Sinica*, 52(4), 528–540.
- [李美娟, 刘玥, 刘红云. (2020). 计算机动态测验中问题解决过程策略的分析: 多水平混合 IRT 模型的拓展与应用. *心理学报*, 52(4), 528–540.]
- Liu, Y., Xu, H., Chen, Q., & Zhan, P. (2022). The measurement of problem-solving competence using process data. *Advances in Psychological Science*, 30(3), 522–535.
- [刘耀辉, 徐慧颖, 陈琦鹏, 詹沛达. (2022). 基于过程数据的问题解决能力测量及数据分析方法. *心理科学进展*, 30(3), 522–535.]
- Novick, L. R., & Bassok, M. (2005). Problem solving. In K. J. Holyoak & R. G. Morrison (Eds.), *The Cambridge handbook of thinking and reasoning* (pp. 321–349). Cambridge University Press.
- OECD. (2013). *PISA 2012 assessment and analytical framework: Mathematics, reading, science, problem solving and financial literacy*. PISA, OECD Publishing, Paris.
- Qiao, X., & Jiao, H. (2018). Data mining techniques in analyzing process data: A didactic. *Frontiers in Psychology*, 9, Article 2231.
- Shu, Z., Bergner, Y., Zhu, M., & Hao, J. (2017). An item response theory analysis of problem-solving processes in scenario-based tasks. *Psychological Test and Assessment Modeling*, 59(1), 109–131.
- Stan Development Team. (2024). *RStan: The R interface to Stan* (Version 2.32.6) [R]. <https://mc-stan.org/>
- Tang, X. (2024). A latent hidden Markov model for process data. *Psychometrika*, 89(1), 205–240.
- Tang, X., Wang, Z., He, Q., Liu, J., & Ying, Z. (2020). Latent feature extraction for process data via multidimensional scaling. *Psychometrika*, 85(2), 378–397.
- Ulitzsch, E., He, Q., & Pohl, S. (2022). Using sequence mining techniques for understanding incorrect behavioral patterns on interactive tasks. *Journal of Educational and Behavioral Statistics*, 47(1), 3–35.
- Vehtari, A., Gelman, A., & Gabry, J. (2017). Practical Bayesian model evaluation using leave-one-out cross-validation and WAIC. *Statistics and Computing*, 27(5), 1413–1432.
- Wang, P., & Liu, H. (2024). Polytomous effectiveness indicators in complex problem-solving tasks and their applications in developing measurement model. *Psychometrika*, 89(3), 877–902.
- Wang, Z., Tang, X., Liu, J., & Ying, Z. (2023). Subtask analysis of process data through a predictive model. *British Journal of Mathematical and Statistical Psychology*, 76(1), 211–235.
- Watanabe, S. (2010). Asymptotic equivalence of Bayes cross validation and widely applicable information criterion in singular learning theory. *Journal of Machine Learning Research*, 11, 3571–3594.
- Xiao, Y., He, Q., Veldkamp, B., & Liu, H. (2021). Exploring latent states of problem - solving competence using hidden Markov model on process data. *Journal of Computer Assisted Learning*, 37(5), 1232–1247.
- Xiao, Y., & Liu, H. (2024). A state response measurement model for problem-solving process data. *Behavior Research Methods*, 56(1), 258–277.
- Xu, X., Zhang, S., Guo, J., & Xin, T. (2024). Biclustering of log data: Insights from a computer-based complex problem solving assessment. *Journal of Intelligence*, 12(1), Article 10, 1–32.
- Yarkoni, T., & Westfall, J. (2017). Choosing prediction over explanation in psychology: Lessons from machine learning. *Perspectives on Psychological Science*, 12(6), 1100–1122.
- Yuan, J., Li, M., & Liu, H. (2023). The development and challenge of contextualized testing. *Journal of China Examinations*, 3, 17–26.
- [袁建林, 李美娟, 刘红云. (2023). 情境化测验的进展与挑战. *中国考试*, 3, 17–26.]
- Zhan, P., & Qiao, X. (2022). Diagnostic classification analysis of problem-solving competence using process data: An item expansion method. *Psychometrika*, 87(4), 1529–1547.
- Zheng, Y., Nydick, S., Huang, S., & Zhang, S. (2024). MxML (exploring the relationship between measurement and machine learning): Current state of the field. *Educational Measurement: Issues and Practice*, 43(1), 19–38.

Analysis of learning effect in interactive problem-solving test: Extension and application of process data measurement model

LU Xiangyu, CHEN Ping

(Collaborative Innovation Center of Assessment for Basic Education Quality, Beijing Normal University, Beijing 100875, China)

Abstract

In the past decade, computer-based interactive problem-solving tests have become increasingly popular in large-scale assessments. Such tests require examinees to interact with a computer, explore virtual scenarios, and solve practical problems, thus making it possible to record the sequences of actions performed by examinees (i.e., process data). Process data contain rich information about the problem-solving process and can help to gain a deeper understanding of examinees' problem-solving strategies. Methods for analyzing such process data are still under development. For example, Han et al. (2021) proposed a sequential response model (SRM) that combines comprehensive information from the response process to infer problem-solving ability. Fu et al. (2023) replaced the multinomial logistic modeling in SRM with binary logistic modeling and proposed 1P-ASM with relatively lower model complexity. However, existing studies have ignored the fact that students gradually gather information while completing problem-solving tasks (i.e., learning effect). The probability that an examinee performs the correct behavior is affected by their understanding of the problem situation. If the model does not take this into account, it may result in biased estimate of examinee's problem-solving ability. To address this issue, this paper puts forward a new model (denoted as 1P-ASM-R*), which extends 1P-ASM to incorporate this learning effect to obtain more accurate ability estimates.

An empirical study was performed to compare 1P-ASM and 1P-ASM-R* in a real-world interactive assessment item (i.e., "Tickets") in the PISA 2012. The results showed that: (1) the extended model introducing learning effect fitted the empirical data better than the original model; (2) as the examinees delve deeper into the problem, the impact of learning effect on the accuracy of behavioral choices in problem-solving tasks decreased, reflecting a trend of diminishing marginal effect; and (3) introducing learning effect into the model does not affect its ability to capture the characteristics of the problem-solving tasks.

A simulation study was further conducted to explore the psychometric performance of the proposed model in different test scenarios. Three factors were manipulated, they are sample size (200 and 1000), average problem state transition sequence length (short and long), and strength of learning effect (0, 0.1, and 0.3). The problem-solving task structure in the empirical study was used here and 1P-ASM-R* was used to generate the action sequences of the examinees. The results indicated that: (1) when there was no learning effect, 1P-ASM-R* could provide similar fitting performance to the original model and correctly estimate the learning effect parameter as 0. However, when there was a learning effect, 1P-ASM-R* fits the data better, and this advantage became more pronounced as the strength of the learning effect increased; (2) sequence length is one of the factors affecting the parameter recovery of 1P-ASM-R*. The longer the sequence length, the more information the data contains and the higher the parameter estimation accuracy.

In summary, our proposed 1P-ASM-R* model incorporates the learning effect and demonstrates a strong ability to accurately analyze examinees' problem-solving abilities. The combination of simulation and empirical findings highlights the effectiveness of the model in a variety of contexts. Notably, when the task environment lacks a learning effect, the 1P-ASM-R* model exhibits comparable performance to the original 1P-ASM model. This finding underscores the excellent stability and adaptability of the model, indicating that it can function reliably under different conditions.

Keywords process data, problem-solving, item response theory

附录：模拟研究中容易度参数的返真性

表 A1、A2 分别给出了模拟研究中 1P-ASM 和 1P-ASM-R*在各模拟条件下容易度参数 $\beta_1 \sim \beta_{10}$ 的 Bias, 表 A3、A4 分别给出了模拟研究中 1P-ASM 和 1P-ASM-R*在各模拟条件下容易度参数 $\beta_1 \sim \beta_{10}$ 的 RMSE。指标的具体计算方法如下：

$$bias = \frac{1}{R} \sum_{r=1}^R (\hat{\beta}_r - \beta) \tag{A1}$$

$$RMSE = \sqrt{\frac{1}{R} \sum_{r=1}^R (\hat{\beta}_r - \beta)^2} \tag{A2}$$

其中 R 为重复模拟的总次数, $\hat{\beta}_r$ 和 β 分别是被试 i 在第 r 次重复时的能力估计值和真实值。

表 A1 模拟研究中 1P-ASM 在各条件下的 Bias

样本量	序列长度	学习效应大小	β_1	β_2	β_3	β_4	β_5	β_6	β_7	β_8	β_9	β_{10}
200	短	0	0.018	-0.078	-0.077	-0.086	-0.174	-0.083	0.130	0.203	0.066	0.059
		0.1	0.060	0.013	0.054	0.079	0.071	0.078	0.214	0.295	0.184	0.225
		0.3	0.108	0.310	0.404	0.484	0.537	0.431	0.465	0.524	0.556	0.614
	长	0	-0.034	-0.045	-0.037	-0.092	-0.093	-0.129	-0.046	-0.036	0.000	-0.044
		0.1	0.158	0.156	0.235	0.264	0.225	0.298	0.147	0.249	0.252	0.197
		0.3	0.409	0.557	0.657	0.724	0.821	0.509	0.474	0.544	0.583	0.467
1000	短	0	0.019	0.013	0.022	0.020	-0.014	-0.040	0.023	0.045	0.014	0.016
		0.1	0.068	0.154	0.199	0.238	0.236	0.187	0.143	0.199	0.202	0.215
		0.3	0.192	0.423	0.529	0.607	0.631	0.631	0.398	0.489	0.588	0.613
	长	0	0.012	0.017	0.011	0.016	0.022	0.002	0.017	0.002	-0.001	0.021
		0.1	0.197	0.235	0.268	0.299	0.342	0.341	0.209	0.242	0.281	0.283
		0.3	0.456	0.635	0.749	0.851	0.939	0.787	0.502	0.560	0.650	0.689

表 A2 模拟研究中 1P-ASM-R*在各条件下的 Bias

样本量	序列长度	学习效应大小	β_1	β_2	β_3	β_4	β_5	β_6	β_7	β_8	β_9	β_{10}
200	短	0	0.012	-0.094	-0.099	-0.112	-0.201	-0.103	0.116	0.185	0.043	0.038
		0.1	-0.009	-0.149	-0.154	-0.165	-0.199	-0.129	0.079	0.120	-0.038	-0.006
		0.3	-0.055	-0.113	-0.141	-0.154	-0.160	-0.067	0.135	0.093	0.006	0.045
	长	0	-0.038	-0.051	-0.046	-0.101	-0.103	-0.135	-0.051	-0.041	-0.006	-0.049
		0.1	-0.052	-0.107	-0.067	-0.076	-0.152	0.023	-0.085	-0.011	-0.043	-0.101
		0.3	-0.055	-0.108	-0.132	-0.171	-0.159	-0.142	-0.020	-0.017	-0.076	-0.101
1000	短	0	0.019	0.010	0.019	0.016	-0.019	-0.045	0.020	0.042	0.011	0.011
		0.1	0.006	0.014	0.021	0.027	-0.003	-0.036	0.018	0.038	0.008	-0.001
		0.3	0.018	-0.003	-0.017	-0.036	-0.091	-0.042	0.048	0.029	0.024	-0.011
	长	0	0.023	0.029	0.024	0.030	0.039	0.016	0.029	0.015	0.014	0.037
		0.1	0.006	0.000	-0.001	-0.004	0.004	0.035	-0.003	0.004	0.007	-0.013
		0.3	0.013	0.010	0.012	0.013	0.010	-0.043	0.021	0.011	-0.004	-0.003

表 A3 模拟研究中 1P-ASM 在各条件下的 RMSE

样本量	序列长度	学习效应大小	β_1	β_2	β_3	β_4	β_5	β_6	β_7	β_8	β_9	β_{10}
200	短	0	0.132	0.223	0.205	0.220	0.275	0.478	0.296	0.360	0.238	0.359
		0.1	0.149	0.182	0.180	0.238	0.298	0.447	0.321	0.386	0.289	0.434
		0.3	0.173	0.361	0.449	0.526	0.596	0.661	0.524	0.590	0.592	0.699
	长	0	0.083	0.093	0.133	0.180	0.209	0.550	0.163	0.139	0.137	0.281
		0.1	0.174	0.184	0.266	0.293	0.305	0.532	0.219	0.292	0.323	0.431
		0.3	0.416	0.564	0.664	0.734	0.843	0.709	0.496	0.570	0.637	0.644
1000	短	0	0.069	0.092	0.096	0.110	0.134	0.261	0.139	0.159	0.112	0.150
		0.1	0.090	0.183	0.222	0.265	0.273	0.302	0.206	0.244	0.243	0.273
		0.3	0.200	0.432	0.536	0.617	0.642	0.706	0.413	0.513	0.601	0.644
	长	0	0.038	0.052	0.052	0.066	0.092	0.242	0.069	0.072	0.088	0.153
		0.1	0.201	0.239	0.274	0.305	0.350	0.416	0.218	0.254	0.295	0.320
		0.3	0.458	0.637	0.751	0.853	0.945	0.855	0.508	0.567	0.662	0.720

表 A4 模拟研究中 1P-ASM-R*在各条件下的 RMSE

样本量	序列长度	学习效应大小	β_1	β_2	β_3	β_4	β_5	β_6	β_7	β_8	β_9	β_{10}
200	短	0	0.131	0.255	0.246	0.237	0.314	0.461	0.279	0.361	0.262	0.352
		0.1	0.154	0.262	0.275	0.305	0.381	0.458	0.264	0.316	0.266	0.364
		0.3	0.162	0.238	0.280	0.302	0.359	0.527	0.296	0.317	0.268	0.320
	长	0	0.109	0.116	0.151	0.204	0.214	0.561	0.178	0.163	0.167	0.274
		0.1	0.109	0.161	0.177	0.150	0.252	0.441	0.205	0.169	0.219	0.418
		0.3	0.113	0.163	0.184	0.230	0.267	0.508	0.162	0.190	0.283	0.484
1000	短	0	0.065	0.119	0.113	0.135	0.155	0.292	0.150	0.171	0.129	0.170
		0.1	0.072	0.132	0.132	0.151	0.166	0.272	0.162	0.181	0.145	0.207
		0.3	0.064	0.106	0.120	0.152	0.183	0.355	0.129	0.176	0.158	0.235
	长	0	0.050	0.065	0.071	0.084	0.108	0.251	0.082	0.081	0.097	0.154
		0.1	0.054	0.049	0.071	0.078	0.075	0.240	0.070	0.090	0.093	0.158
		0.3	0.071	0.070	0.083	0.092	0.119	0.329	0.082	0.096	0.152	0.235