

知觉还是概念？先前学习经验 对恐惧泛化路径的调节*

冯彪^{1,2} 张栋桓¹ 陈伟¹ 曾灵¹ 吴潇悦¹ 黄君琳¹ 郑希付¹

(¹ 华南师范大学大脑、认知与教育科学教育部重点实验室；华南师范大学心理学院；华南师范大学心理应用研究中心，广东省心理健康与认知科学重点实验室；² 华南师范大学教育信息技术学院，广州 510631)

摘要 过度的恐惧泛化是各类焦虑症的核心症状，人类的恐惧反应既可沿刺激的知觉线索传递，也可根据刺激的概念信息蔓延，给患者的生活造成极大负担。本研究采用经典的差异性条件恐惧范式，以 US 预期和皮肤电反应为指标，探讨了恐惧习得前的学习经验对两种恐惧泛化路径的调节作用。本研究共招募 50 名在校大学生，在恐惧习得前，被试被随机分为两组，一组对刺激的知觉属性进行判断学习(知觉组)，另一组则对刺激的概念属性进行判断学习(概念组)，随后两组被试进行完全相同的恐惧习得和泛化测试。结果显示，知觉组被试产生了显著的知觉性恐惧泛化，而概念组被试则表现出明显的概念性恐惧泛化，先前经验显著影响了被试恐惧泛化的路径。此外，概念组被试还表现出明显的知觉性恐惧泛化的倾向，而知觉组却未发现任何概念性恐惧泛化的迹象，两组被试在恐惧泛化路径上表现出非对称的特点。

关键词 先前经验，知觉信息，概念信息，注意偏向，恐惧泛化

分类号 B845

1 引言

恐惧是一种进化而来的情绪反应，它帮助个体应对威胁，促进生存(Resnik et al., 2011)。个体对某具体事物的恐惧大多是后天习得的，当一个中性刺激(如，色块)反复和厌恶刺激(如，电击)匹配呈现后，单独呈现该中性刺激即可诱发个体的条件性恐惧反应(conditioned fear)，该中性刺激被称为条件刺激(conditioned stimulus, CS)，该厌恶刺激被称为非条件刺激(unconditioned stimulus, US) (Lissek et al., 2008)。然而，人类习得的恐惧很少只被限定于原来的刺激或者情境中，往往与 CS 类似的其他刺激或者情境也会诱发个体的恐惧反应，表现出典型的恐惧泛化(fear generalization)现象，这些与原 CS 类似但不同的刺激被称作泛化刺激(generalization stimulus, GS) (Fraunfelter et al., 2022)。和恐惧反应

本身类似，正常的恐惧泛化具有一定的生存适应性，它使个体在不必亲身经历危险的情况下，即可发现潜在的威胁，进而促进个体的生存。然而当恐惧对象的范围不断扩大，恐惧的边界会变得模糊，个体会对本质上非威胁的刺激或情境也产生不必要的恐惧反应或回避，此时正常的恐惧泛化就变成了非适应性的临床症状(Fraunfelter et al., 2022; Sep et al., 2019)。研究表明，面对泛化刺激，相较于正常个体，焦虑症患者在杏仁核(amygdala)、前扣带回(anterior cingulate cortex, ACC)、腹侧被盖区(ventral tegmental area, VTA)等脑区表现出更强烈的激活，反映出对威胁信息的过度警觉和反应，而在前额叶皮质(prefrontal cortex, PFC)等脑区的激活则降低，表明患者对恐惧反应的抑制减弱，而在海马(hippocampus)等区域的功能异常，表明患者无法分辨安全和威胁刺激，导致过度泛化(Dunsmoor &

收稿日期: 2024-12-29

* 国家自然科学基金项目(32371137)；广东省脑认知与人的素质发展基础学科研究中心(2024B0303390003)资助；幸福广州心理服务与辅导基地资助；华南师范大学冲补强心理学高峰学科资助；广东省哲学社会科学规划学科共建项目(GD23XXL15)。

通信作者: 郑希付, E-mail: zhengxifu@m.scnu.edu.cn

Paz, 2015; Lopresto et al., 2016)。恐惧的过度泛化(Fear overgeneralization)被认为是各类焦虑症的核心症状(American Psychiatric Association, 2013; Cooper et al., 2022; Lee et al., 2024), 在焦虑症的形成和维持过程中扮有重要作用, 是临床干预的重点对象(Scheveneels et al., 2021)。

以往的研究表明, 人类的恐惧反应可以沿着刺激的知觉线索传递, 即知觉性恐惧泛化, 表现出常见的恐惧泛化梯度(Dymond et al., 2015)。如 Lissek 等人(2008)的研究表明, 当分别对大、小圆习得条件性恐惧和安全反应后, 个体会对尺寸介于大、小圆之间的其他类似的圆(GS)也产生恐惧反应, 且对这些 GS 的恐惧反应强度会随着其与条件性恐惧刺激(CS+)的知觉相似性的下降而逐渐降低, 后来的研究者在其他知觉维度也发现了类似的现象, 比如声音(Resnik et al., 2011)、颜色(Dunsmoor & LaBar, 2013)等, 表现出典型的知觉效应(Feng et al., 2025)。临床研究进一步发现, 相对于正常被试, 焦虑症患者在知觉线索上会表现出明显的过度性恐惧泛化(Cooper et al., 2022; Fraunfelter et al., 2022)。除了依赖于知觉线索的恐惧泛化, 人类还可凭借自身发达的认知系统, 在更高的概念水平进行恐惧泛化, 即概念性恐惧泛化(Dunsmoor & Murphy, 2015)。Dunsmoor 等人(2012)的研究表明, 习得的恐惧反应可以根据刺激的类别信息进行传递, 被试对未直接匹配电击的恐惧类别的其他样例成员也会表现出恐惧反应, 而概念关联的强度会影响恐惧泛化的程度(Mertens et al., 2021), 此外, 恐惧泛化还表现出经典的典型性和多样性效应, 对类别典型样例的恐惧可以泛化到非典样例上, 反之则不成立(Dunsmoor & Murphy, 2014), 而恐惧习得过程中类别样例的多样性会进一步加剧恐惧的泛化程度(Fan et al., 2022)。和知觉性恐惧泛化类似, 研究初步表明正常被试和病人在概念性恐惧泛化上也表现出不同的反应模式(Morey et al., 2020; Wang, Wang, et al., 2021)。

如上所述, 尽管以往的研究对基于知觉和概念线索的恐惧泛化进行了大量探讨, 然而这些研究往往将这两类恐惧泛化进行分离的单独研究, 无法对恐惧泛化中的两种路径进行同时考察, 考虑到知觉和概念过程可能存在相互影响(Goldstone et al., 2000; Jepma & Wager, 2015), 因此研究存在一定的局限。近年来, 雷怡团队(Wang, Mei, et al., 2021; Wang, Wang, et al., 2021; Zhang et al., 2024)对知觉

和概念性恐惧泛化进行了深入探讨。在他们的研究中, 恐惧习得阶段, 被试对刺激的知觉和概念信息分别进行条件化恐惧习得, 在泛化测试阶段则呈现由概念和知觉信息组合而成的复合刺激, 以此考察被试恐惧泛化的概念和知觉成分。研究表明, 习得阶段被试对知觉线索的恐惧预期比概念线索更高, 且反应更快, 而泛化阶段被试对概念威胁线索的恐惧反应显著大于安全线索, 但在知觉威胁和安全线索上则没有这种差异(Wang, Mei, et al., 2021)。此外, 广泛性焦虑症患者(generalised anxiety disorder, GAD)相对于正常被试, 也表现出对概念性泛化线索更强的恐惧预期, 而对知觉性泛化线索的恐惧预期则与正常被试无异(Wang, Wang, et al., 2021), 这些研究暗示, 刺激的知觉和概念信息在恐惧泛化过程中扮有不同的作用。

一般认为, 知觉性恐惧泛化涉及到刺激信号通过丘脑(thalamus)直接向杏仁核投射神经信号, 以及各感觉皮层(如, 视觉皮层, visual cortex)和杏仁核间的功能连接增强, 该过程受刺激物理特性的影响; 而概念性恐惧泛化则依赖于前额叶皮质和海马等脑区对信息的进一步抽象整合, 之后再向杏仁核投射神经信号, 此过程受到语义和情景的影响(Dunsmoor & Paz, 2015; Webler et al., 2021)。从人类认知加工的一般过程看, 刺激信息的处理从知觉特征分析到概念归类是一个连续过程, 任何刺激都同时具有知觉和概念两种属性信息, 现实生活中个体的恐惧习得不太可能对刺激的知觉或者概念线索进行分离的条件化学习, 而是对知觉和概念的复合体进行整体的恐惧习得(Dunsmoor & Murphy, 2015), 因此当个体习得对某刺激复合体的恐惧后, 在什么情况下刺激的概念或知觉线索会主导恐惧泛化的路径, 即何种因素会调节恐惧泛化的两种路径, 这对于我们理解恐惧泛化中的信息加工过程具有重要的理论意义(Fraunfelter et al., 2022)。另一方面, 考虑到不同类型的焦虑症患者在恐惧泛化的类型上存在差异, 如蜘蛛恐惧症(spider phobia)、旷野恐惧症(claustrophobic)和社交恐惧症(social anxiety disorder)患者更多地表现出知觉性恐惧泛化(Lee et al., 2024; Peperkorn et al., 2013; Shiban et al., 2016), 而 GAD、创伤后应激障碍(posttraumatic stress disorder, PTSD)患者则表现出概念性恐惧泛化的倾向(Morey et al., 2020; Wang, Wang, et al., 2021), 然而不同研究者的结果并不一致(Cooper et al., 2022; Fraunfelter et al., 2022), 因此, 弄清楚何种因素调

节恐惧泛化的路径对于我们理解焦虑症维持和发展的机制也具有重要的现实意义。

恐惧泛化和学习迁移涉及到类似的心理过程(Dunsmoor & Murphy, 2014; Shepard, 1987), 人类可通过过去学习的知识来应对当下的陌生情境, 过往经验会影响个体的注意、对客体的识别和推理等(雷怡等, 2017), 因此不同的先前学习经验很可能导致个体的恐惧泛化表现出不同的反应模式(Zaman et al., 2023)。Vervliet 等人(2010a)的研究表明对泛化刺激的先前安全经验可以抑制随后对其的恐惧反应, 表现出良好的保健作用, 而实验前的言语诱导则可调节恐惧泛化的线索路径, 表明人类语言系统对恐惧泛化的影响(Vervliet et al., 2010b)。Ahmed 和 Lovibond (2015)的研究表明, 被试甚至可以根据过去的学习经历, 从中归纳相应的规则, 并将这种习得的规则应用到新的学习中去, 从而影响恐惧的泛化。以上这些研究都初步表明恐惧习得前的先前经验对恐惧的维持和发展具有深远影响。

近期有研究表明, 视知觉的客体识别(object recognition)过程是恐惧泛化的重要环节, 对客体的识别程度决定了恐惧反应的强度, 且该过程和人格因素呈现明显的交互作用(Feng et al., 2025)。考虑到人类高级认知系统可根据任务要求灵活地调配认知资源在刺激的不同属性上, 影响个体对刺激的注意和记忆(Franconeri et al., 2013; Kim & Rehder, 2010), 而不同的注意指向会导致对同一客体产生完全不同的主观归类和识别(Nosofsky, 1986)。Dowd 等人(2016)的研究表明个体对泛化刺激的注意分配梯度与恐惧泛化梯度相一致, 表明注意资源的分配和恐惧泛化是同步的。Reutter 和 Gamer (2023)的研究进一步显示, 对刺激不同特征的注意力分配决定了恐惧泛化的程度, 对区分性特征的持续关注可以抑制恐惧泛化, 而对非区分性特征的关则加剧恐惧泛化。考虑到先前学习的知识经验会调节被试的注意偏向, 从而影响个体对刺激的归类和识别(Kim & Rehder, 2010), 我们预期先前学习经验可通过影响个体的注意偏向, 导致个体对同样的威胁刺激产生完全不同的主观识别, 进而导致习得的恐惧沿着两种不同的路径泛化。

综上所述, 本研究在前人的基础上, 采用更具生态效度的实验范式对恐惧泛化的两种路径进行综合研究, 考察恐惧习得前的先前学习经验对后续恐惧泛化路径的影响。本研究预期, 恐惧习得前的不同属性判断任务可调控被试对刺激不同属性信

息的注意加工, 导致对后续威胁刺激产生完全不同的识别, 进而影响恐惧泛化的路径。具体来说, 本研究假设, 恐惧习得前的知觉属性判断会导致被试在后续的恐惧泛化测试中表现出明显的知觉性恐惧泛化(即, 对 P+的反应显著大于 P-, 而对 C+和 C-的反应无显著差异), 而概念属性判断则会让被试表现出明显的概念性恐惧泛化(即, 对 C+的反应显著大于 C-, 而对 P+和 P-的反应无显著差异), 两组被试在恐惧泛化的路径上表现出明显的差异。

2 方法

2.1 被试

通过自由报名的方式招募在校大学生参与本实验, 全程参与本实验可获得一定的报酬。被试招募的标准为, 所有被试需年满 18 岁, 均为右利手, 没有任何心理疾病或精神障碍病史, 视力或矫正视力正常, 听力正常, 最近没有鼻塞或咳嗽等症状, 6 个月内没有参加过类似的恐惧情绪实验。本研究通过了华南师范大学心理学院伦理委员会的伦理审查(批准编号: SCNU-PSY-2024-418)。在正式实验前, 被试签署知情同意书, 告知被试实验过程, 介绍皮肤电实验的原理, 告知实验过程中会出现轻微的电击, 但不会对人体造成实质性伤害, 并提醒被试“在实验过程中若有任何不适, 可以随时降低电击强度或者终止实验”, 告知被试其隐私将被严格保密, 所有与其有关的数据仅被用于科研目的, 同时也要求被试对实验内容进行保密, 被试确认无误后, 在指定表格中签署自己的姓名。之后被试被随机分配到两个实验组中, 被试对于实验目的和自己的实验组别均不知情, 实验执行人员也不知道被试的组别。

采用 G*Power 3.1 软件对本实验的样本量进行估算, 为了在显著性水平 α 为 0.05, 以 0.8 的检验效能($1 - \beta$), 发现中等水平的效应量($f = 0.25$), 至少需要 46 名被试, 本实验最终招募有效被试 50 人, 被随机分配到两个实验组中, 其中知觉组 26 人(男生 5 人), 概念组 24 人(男生 4 人), 被试年龄在 18~29 岁之间。在正式实验前, 被试需填写不确定不容忍量表(The Intolerance Of Uncertainty Scale, IUS-12)(张亚娟等, 2017)、抑郁-焦虑-压力量表(The Depression Anxiety Stress Scale, DASS-21)(龚栩等, 2010)、焦虑敏感指数量表(Anxiety Sensitivity Index-3, ASI-3)(王雷等, 2014), 两组被试在这些人格变量上无显著差异, 详见表 1。

2.2 实验材料

参照 Wang, Mei 等人(2021)等人的研究, 本实验材料共包括 100 张矢量图片, 各图片在亮度、饱和度和等视觉方面保持一致, 这些图片按照实验设计共分为 8 类: 25 张紫色动物图片, 15 张蓝色动物图片, 5 张黑色动物图片; 25 张蓝色工具图片, 15 张紫色工具图片, 5 张黑色工具图片; 5 张蓝色云团图片, 5 张紫色云团图片。在正式实验前, 由 40 名独立被试对这 8 类图片的效价和唤醒度进行 1 到 9 分的评定, 结果显示, 各类图片在效价 $[F(7, 33) = 1.85, p = 0.111]$ 和唤醒度 $[F(7, 33) = 1.77, p = 0.128]$ 上均无显著差异。所有图片在实验中均只随机呈现一次。

2.3 测量指标

2.3.1 US 主观预期

在每个实验试次中, 要求被试根据自己的学习经验对该图片后面伴随电击的可能性进行主观评定。在每个试次的屏幕下方会出现 0 到 10 的标尺, 0 代表“一定没有电击”, 5 代表“不确定”, 10 代表“一定有电击”, 被试的任务是通过鼠标拖动标尺上的

滑块选择相应的数字, 以做出自己的 US 主观预期判断。

2.3.2 皮肤电反应

采用 Biopac Mp36 系统采集被试的皮肤电数据, 采样率为 2000 Hz, 将两个涂满导电膏(Biopac GEL101)的皮肤电传感器(Biopac SS3LA EDA Finger Transducer)粘贴于被试左手食指和中指的指腹。原始皮肤电(Skin conductance response, SCR)信号采用 1Hz 的低通滤波, 采用刺激呈现后 7.5 s 的最大值减去刺激呈现前 2 s 的平均值作为原始 SCR 值, 小于 0.02 微西门子(microsiemens, μ s)的原始 SCR 值被归为 0 也纳入数据分析中, 对这些 SCR 值进行开方处理以促进数据分布正态化, 采用 Z 转换以减小被试间的个体误差(Bauer et al., 2020; Feng et al., 2025)。

2.4 实验设计与程序

本实验采用 Eprime 3.0 编程实现, 共分成 4 个阶段: 前学习阶段、恐惧习得阶段、泛化测试阶段 1 和泛化测试阶段 2, 一共 72 个试次, 如表 2 所示。

表 1 被试分组及问卷信息表

变量	组别		t 或 χ^2	p
	知觉组($n = 26$)	概念组($n = 24$)		
男被试人数(占比)	5(19.23%)	4(16.67%)	0.06	0.814
年龄(岁)	21.23 ± 2.47	21.25 ± 1.30	0.03	0.973
DASS21-D	5.92 ± 6.34	7.67 ± 7.57	0.89	0.380
DASS21-A	7.15 ± 5.66	7.25 ± 5.27	0.06	0.951
DASS21-S	8.46 ± 5.49	9.50 ± 5.60	0.66	0.511
IU12-P	16.19 ± 3.49	18.04 ± 4.52	1.63	0.110
IU12-I	9.38 ± 2.58	11.00 ± 3.23	1.96	0.056
ASI3-P	4.19 ± 4.13	5.91 ± 5.13	1.31	0.195
ASI3-C	5.15 ± 3.13	5.54 ± 4.28	0.37	0.715
ASI3-S	7.00 ± 3.89	8.38 ± 5.12	1.07	0.288

注: DASS21-D, A, S 分别为 DASS21 的抑郁、焦虑、压力亚量表; IU12-P, I 分别为 IU12 的预期性焦虑和抑制性焦虑亚量表; ASI3-P, C, S 分别为 ASI3 的躯体关注、认知关注和社会关注亚量表。

表 2 实验流程及参数表

前学习阶段	恐惧习得阶段	泛化测试阶段 1 (断开电击仪)	泛化测试阶段 2 (连接电击仪)
知觉组 知觉属性判断(“紫色还是蓝色”), 40 个试次	CS+ × 8 (75% 电击) CS- × 8(无电击)	C+P+ × 1(无电击)	C+P+ × 1(100%电击)
		C+P0 × 1(无电击)	C+P0 × 1(无电击)
		C+P- × 1(无电击)	C+P- × 1(无电击)
		C-P+ × 1(无电击)	C-P+ × 1(无电击)
概念组 概念属性判断(“动物还是工具”), 40 个试次	CS- × 8(无电击)	C-P0 × 1(无电击)	C-P0 × 1(无电击)
		C0P- × 1(无电击)	C0P- × 1(无电击)
		C0P+ × 1(无电击)	C0P+ × 1(无电击)
		C-P- × 1(无电击)	C-P- × 1(无电击)

注: 前学习阶段的反应指标为正确率; 恐惧习得和泛化测试阶段的反应指标为 US 主观预期和 SCR。

正式实验前,被试填写知情同意书和问卷,之后对 US 的电击强度进行评定。电刺激由恒流刺激仪(model DS2A-MK. II, Digitimer, Inc., Hertfordshire, UK)产生,施加于被试的右手腕部。电击强度从 10 V 开始每次按 5 V 递增,直到被试选择一个“极端不舒服但不痛”的电击强度,最高 80 V,该电击强度作为正式试验中的电击强度(Fan et al., 2022; Feng et al., 2025)。

前学习阶段:参考前人的研究(Ahmed & Lovibond, 2015; Kim & Rehder, 2010; Wong & Lovibond, 2021),本阶段通过属性判断任务以控制被试对刺激的先前学习经验。本阶段包括 40 个训练试次,呈现 40 张图片,包括 10 张紫色动物图片,10 张蓝色动物图片,10 张紫色工具图片,10 张蓝色工具图片。40 张图片随机出现在屏幕中央,知觉组被试被要求对图片的知觉属性进行判断(即判断该图片是蓝色的还是紫色的),而概念组被试则被要求对这些图片的类别属性进行判断(即判断该图片是动物还是工具),被试通过鼠标点击屏幕下方的对应按钮做出反应,每次被试反应后会有正误反馈,以强化被试的学习。

恐惧习得阶段:本阶段包括 16 个实验试次,条件威胁刺激 CS+ 和条件安全刺激 CS-各呈现 8 次,对这 16 个试次进行伪随机设置,确保 CS+ 或 CS-不会连续出现超过两次。CS+试次后有 75% 的概率出现电击,本阶段所有刺激均呈现 8 s, CS+的电击出现在其开始呈现后的 7500 ms,电击呈现时间为 500 ms,随 CS+一起消失。各试次的间隔(Inter-Trial Interval, ITI)时间为 10 到 14 s,平均为 12 s。本阶段一共呈现 16 张图片,包括 8 张紫色动物图片,8

张蓝色工具图片,两类图片作为 CS+或 CS-的情况在被试间进行平衡。如前所述,在每个试次中,被试需用鼠标拖动屏幕下方的滑块对该图片后面伴随电击的可能性进行“0~10”的评定,0 代表“一定没有电击”,5 代表“不确定”,10 代表“一定有电击”(见图 1)。

泛化测试阶段 1:本阶段包括 8 个测试试次,根据刺激的知觉和概念信息的组合,一共随机呈现 8 种不同的泛化刺激:概念和知觉信息都危险的刺激(C+P+),如紫色动物;概念和知觉信息都安全的刺激(C-P-),如蓝色工具;概念信息安全而知觉信息危险的刺激(C-P+),如紫色工具;概念信息危险而知觉信息安全的刺激(C+P-),如蓝色动物;概念信息危险而知觉信息无关的刺激(C+P0),如黑色动物;概念信息安全而知觉信息无关的刺激(C-P0),如黑色工具;概念信息无关而知觉信息安全的刺激(C0P-),如蓝色云团;概念信息无关而知觉信息危险的刺激(C0P+),如紫色云团。由于本阶段不会出现任何电击,为防止消退过程(即,被试习得所有刺激后面都将不再跟随电击),在本阶段开始之前,电脑屏幕上会出现提示“因为伦理限制,短时间内电击次数已经达到上限,请联系实验员断开电击”,然后研究人员进入实验室断开电击仪,并告知被试,由于伦理限制,电击仪不得不断开,但是要求被试仍按照电击仪依然正常连接的情况来作答(Wong & Lovibond, 2021)。其他设置和习得阶段相同,要求被试根据之前的学习经验对每张图片后面伴随电击可能性进行评定。

泛化测试阶段 2:本阶段类似于泛化测试阶段 1,包括 8 个测试试次,呈现 8 种不同的泛化刺激。不同之处在于,本阶段会重新连接电击。在本阶段

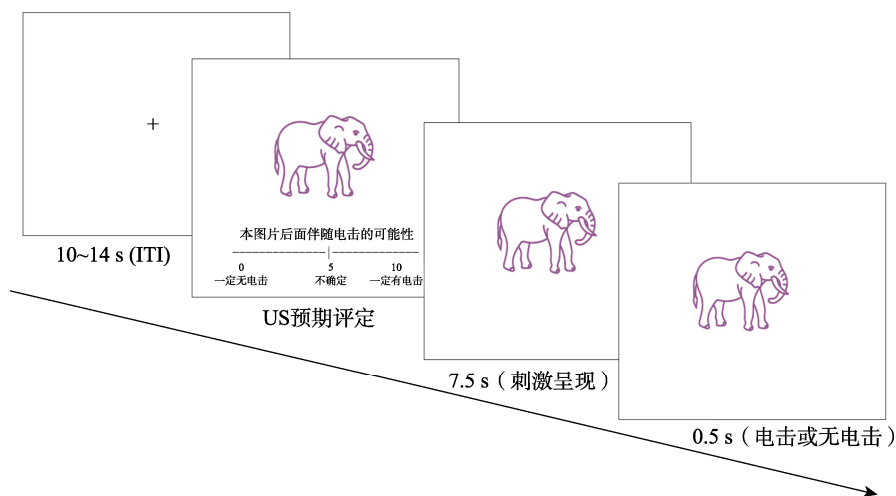


图 1 实验试次示意图

开始之前, 屏幕上再次出现提示“限制解除, 请连接电击仪”, 随后实验人员进入实验室, 告知被试, 电击限制已经解除, 现在重新连接电击, 待实验人员完成连接电击仪后, 明确告知被试电击已经重新连接, 请其继续实验。为了最大限度地防止消退, 减小由于预期错误而引发的消退学习(Wong & Lovibond, 2021), 本阶段的 C+P+和 C-P-会随机出现在前两个试次中, 且 C+P+后面会伴随电击, 其他 6 个刺激在剩下的 6 个试次中随机呈现一次。

2.5 数据分析

前学习阶段, 对各组被试的反应正确率采用独立样本 t 检验, 以验证两组被试在前学习阶段的训练水平是否均等。

恐惧习得阶段, 计算被试对条件刺激 CS+和 CS-的反应均值, 作为后续分析的基本单位。

泛化测试阶段, 将被试对概念威胁刺激的整体反应 C+定义为 C+P-, C+P+, C+P0 三者反应的均值, 将被试对概念安全刺激的整体反应 C-定义为 C-P-, C-P+, C-P0 三者反应的均值, 将被试对知觉威胁刺激的整体反应 P+定义为 C-P+, C+P+, C0P+三者反应的均值, 将被试对知觉安全刺激的整体反应 P-定义为 C-P-, C+P-, C0P-三者反应的均值(Wang, Mei, et al., 2021)。将知觉性恐惧泛化定义为 P+和 P-之间显著的差异性反应, 将概念性恐惧泛化定义为 C+和 C-之间显著的差异性反应。

采用 R 语言的 lme4(lmerTest)、emmeans、effectsize 包对本实验的数据进行混合线性模型分析, 全模型的固定效应部分为刺激类别、组别、刺激类别×组别, 随机效应部分以被试为单位设置随机截距和刺激呈现顺序的随机斜率, 以及以刺激呈现顺序设置随机截距, 并将人口学变量和问卷变量等作为协变量也纳入到模型的固定效应中, 若全模型无法拟合, 则依次去除随机斜率、随机截距, 直至模型退化为一般线性模型, 模型拟合和对比的细节见网络版补充材料。对拟合模型的固定效应采用方差分析进行评估, 检验组别(知觉组, 概念组)和刺激类别(P+, P-, C+, C-)的主效应, 以及组别×刺激类别的交互效应, 多重事后比较采用 BONFERRONI 法矫正, Cohen's d 效应量估算采用模型残差进行标准化, 显著性水平设为 0.05。

3 实验结果

3.1 前学习阶段

计算被试在前学习阶段的正确反应率, 结果显

示两组被试在本阶段都有非常高的反应正确率, 知觉组为 99.90%, 概念组为 99.79%, 表明两组被试在本阶段对刺激的概念或知觉信息都进行了良好的辨别学习。独立样本 t 检验显示, 两组被试在反应正确率上无显著差异, $t(48) = -0.66$, $p = 0.514$, 表明两组被试在本阶段的学习效果是均等的。

3.2 恐惧习得阶段

3.2.1 US 主观预期数据

计算被试对 CS+和 CS-反应的平均值, 采用混合线性模型对均值进行分析。结果显示, 所有包含随机效应的模型均不能收敛, 所以本阶段采用一般线性模型对数据进行分析。结果显示, 一般线性模型对数据拟合显著, $F(15, 84) = 13.64$, $p < 0.001$, 整个模型解释了因变量 70.91% 的变异($R^2 = 0.7091$)。方差分析显示, 刺激类别(CS+ vs. CS-)的主效应显著, $F(1, 84) = 194.81$, $p < 0.001$, $\eta_p^2 = 0.70$, 95% CI (0.59, 0.77), 组别的主效应不显著, $F(1, 84) = 3.61$, $p = 0.152$, 刺激类别和组别的交互效应不显著, $F(1, 84) = 1.20$, $p = 0.152$ 。事后多重比较显示, 两组被试对 CS+的 US 主观预期均显著高于 CS-, 知觉组 [$t(84) = 10.83$, $p < 0.001$, $d = 3.00$, 95% CI (2.28, 3.72)], 概念组 [$t(84) = 8.87$, $p < 0.001$, $d = 2.56$, 95% CI (1.87, 3.26)], 表明两组被试均成功习得了对 CS+和 CS-的差异性恐惧预期(见图 2)。

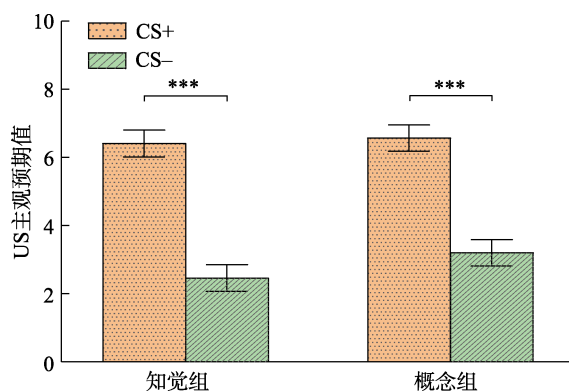


图 2 恐惧习得阶段, 两组被试对 CS+和 CS-的 US 主观预期

注: * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$, ns 为不显著, 误差线为标准误, 下同。

3.2.2 皮肤电反应数据

知觉组有两名被试由于设备故障, SCR 数据未能成功收集, 所以实验最终有效的 SCR 数据为 48 人(知觉组 24 人, 概念组 24 人)。模型结果显示, 所有包含随机效应的模型均不能收敛拟合, 故采用一般线性模型对本阶段的 SCR 数据进行分析。一般

线性模型的拟合结果显示, 模型拟合显著, $F(15, 80) = 4.84, p < 0.001$, 整个模型解释了因变量 47.55% 的变异 ($R^2 = 0.4755$)。方差分析显示, 刺激类别的主效应显著, $F(1, 80) = 59.18, p < 0.001, \eta_p^2 = 0.43, 95\% \text{ CI } (0.27, 0.55)$, 组别的主效应不显著, $F(1, 80) = 2.16, p = 0.146$, 刺激类别和组别的交互效应不显著, $F(1, 80) = 0.001, p = 0.983$ 。事后多重比较显示, 两组被试对 CS+ 的 SCR 值均显著高于 CS-, 知觉组 [$t(80) = 5.46, p < 0.001, d = 1.57, 95\% \text{ CI } (0.99, 2.16)$], 概念组 [$t(80) = 5.42, p < 0.001, d = 1.57, 95\% \text{ CI } (0.99, 2.14)$], 表明两组被试均成功习得了对 CS+ 和 CS- 的差异性恐惧唤醒 (见图 3)。

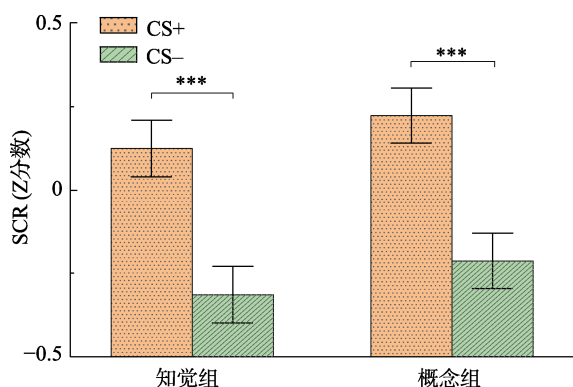


图 3 恐惧习得阶段, 两组被试对 CS+ 和 CS- 的 SCR (Z 分数)

3.3 泛化测试阶段 1——US 主观预期数据

由于本阶段电击已经断开, 仅要求被试假定电击仍然还是连接的, 根据自己之前的学习经验对呈现的泛化刺激进行 US 主观预期评定, 被试清楚知晓本阶段不会出现真实的电击, 所以实验刺激无法引起足够的恐惧唤醒, 本阶段的 SCR 数据失真, 故本阶段仅对 US 预期数据进行分析 (Wong & Lovibond, 2021)。

和习得阶段的数据类似, 混合效应模型分析结果显示, 所有包含随机效应的模型均不能收敛拟合, 故本阶段仍采用一般线性模型对数据进行分析。结果显示, 模型拟合显著, $F(20, 179) = 9.02, p < 0.001$, 整个模型解释了因变量 50.18% 的变异 ($R^2 = 0.5018$)。方差分析显示, 刺激类别的主效应显著, $F(3, 179) = 43.53, p < 0.001, \eta_p^2 = 0.42, 95\% \text{ CI } (0.31, 0.51)$, 组别的主效应不显著, $F(1, 179) = 0.53, p = 0.467$, 刺激类别和组别的交互作用显著, $F(3, 179) = 12.34, p < 0.001, \eta_p^2 = 0.17, 95\% \text{ CI } (0.07, 0.26)$ 。

事后多重比较显示, 知觉组被试对 P+ 的 US 预期显著大于 P- [$t(179) = 9.34, \text{corrected-}p < 0.001$,

$d = 2.67, 95\% \text{ CI } (2.04, 3.30)$] 而对 C+ 和 C- 的 US 预期无显著差异 [$t(179) = 2.39, \text{corrected-}p = 0.106$]; 相反, 概念组被试对 C+ 的 US 预期显著高于 C- [$t(179) = 7.80, \text{corrected-}p < 0.001, d = 2.25, 95\% \text{ CI } (1.64, 2.87)$], 此外, 概念组对 P+ 的反应还显著大于 P- [$t(179) = 2.73, \text{corrected-}p = 0.041, d = 0.79, 95\% \text{ CI } (0.21, 1.37)$], 这表明知觉组被试产生了基于知觉线索的恐惧泛化, 而概念组被试除了产生基于概念线索的恐惧泛化外, 还表现出了明显的知觉性恐惧泛化。

进一步分析发现, 概念组被试对 C+ 和 P+ 的 US 预期无显著差异 [$t(179) = 2.50, \text{corrected-}p = 0.081$], 均显著高于 C- [P+ vs. C-, $t(179) = 5.30, \text{corrected-}p < 0.001, d = 1.53, 95\% \text{ CI } (0.94, 2.12)$] 和 P- [C+ vs. P-, $t(179) = 5.24, \text{corrected-}p < 0.001, d = 1.51, 95\% \text{ CI } (0.92, 2.10)$], 概念组被试对 C- 和 P- 的 US 预期也无显著差异 [$t(179) = -2.56, \text{corrected-}p = 0.068$]。知觉组被试对 C+ 的 US 预期却显著低于 P+ [$t(179) = -3.91, \text{corrected-}p < 0.001, d = -1.10, 95\% \text{ CI } (-1.67, -0.54)$], 对 C- 的 US 预期显著高于 P- [$t(179) = 3.19, \text{corrected-}p = 0.010, d = 0.90, 95\% \text{ CI } (0.34, 1.47)$]。

在组间比较上, 知觉组被试对 P+ 的反应显著高于概念组 [$t(179) = 2.86, p = 0.005, d = 0.85, 95\% \text{ CI } (0.26, 1.44)$], 对 P- 的反应显著低于概念组 [$t(179) = -3.43, p < 0.001, d = -1.03, 95\% \text{ CI } (-1.63, -0.43)$]。概念组被试对 C+ 的反应显著大于知觉组 [$t(179) = 3.27, p = 0.001, d = 0.97, 95\% \text{ CI } (0.38, 1.57)$], 对 C- 的 US 预期显著低于知觉组 [$t(179) = -2.05, p = 0.042, d = -0.61, 95\% \text{ CI } (-1.19, -0.02)$] (见图 4)。

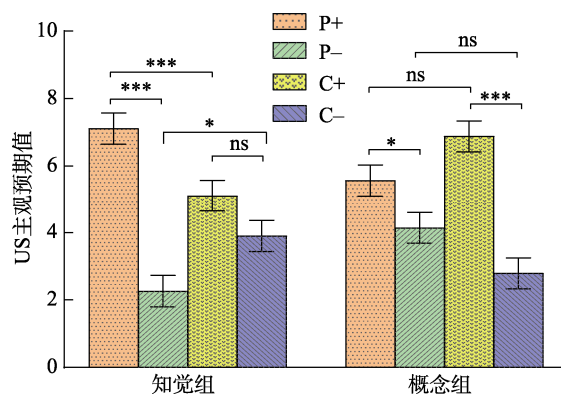


图 4 泛化测试阶段 1 中, 两组被试对泛化刺激的总体 US 预期

以上数据结果表明, 概念组被试除了展现出概念性恐惧泛化外, 还表现出了明显的知觉性恐惧泛化, 使得该组被试对 P+和 P-的反应分别与 C+和 C-无显著差别, 对知觉危险刺激 P+的反应也显著高于知觉安全刺激 P-; 而知觉组却没有表现出概念性恐惧泛化的迹象, 该组被试对 C+的反应较 P+显著下降, 对 C-的反应较 P-也显著上升, C+和 C-之间无显著差异(见图 4)。

3.4 泛化测试阶段 2

3.4.1 US 主观预期数据

采用混合线性模型对本阶段数据进行拟合分析, 结果显示, 固定效应解释了模型 52.85% 的变异 ($Marginal R^2 = 0.5285$), 被试随机截距的效应较小 ($var = 0.034$), 额外解释了模型 0.56% 的变异 ($Conditional R^2 = 0.5342$), 模型整体解释了因变量 56.38% 的变异 ($\Omega^2 = 0.5638$)。方差分析结果显示, 刺激类别的主效应显著, $F(3, 143.43) = 53.07, p < 0.001, \eta_p^2 = 0.53, 95\% CI (0.41, 0.61)$, 组别的主效应不显著, $F(1, 36.88) = 3.03, p = 0.09$, 刺激类型 $\times$ 组别的交互效应显著, $F(3, 143.58) = 14.32, p < 0.001, \eta_p^2 = 0.23, 95\% CI (0.11, 0.34)$ 。

事后比较显示, 知觉组被试对 P+的反应显著高于 P- [$t(144) = 10.40, corrected-p < 0.001, d = 2.94, 95\% CI (2.30, 3.58)$], 对 C+和 C-的反应无显著差异 [$t(143) = 2.66, corrected-p = 0.052$]; 对 C+的反应显著小于 P+ [$t(144) = -4.19, corrected-p < 0.001, d = -1.16, 95\% CI (-1.73, -0.41)$], 对 C-的反应显著高于 P- [$t(143) = 3.72, corrected-p = 0.002, d = 1.03, 95\% CI (0.47, 1.60)$]。概念组被试对 C+的反应显著高于 C- [$t(143) = 8.61, corrected-p < 0.001, d = 2.49, 95\% CI (1.86, 3.13)$], 对 P+的反应也显著大于 P- [$t(143) = 3.40, corrected-p = 0.005, d = 0.98, 95\% CI (0.40, 1.56)$], 对 C+和 P+的反应无显著差异 [$t(143) = 2.11, corrected-p = 0.219$], 对 C-的反应显著低于 P- [$t(143) = -3.12, corrected-p = 0.013, d = -0.91, 95\% CI (-1.49, -0.32)$]。

组间比较上, 知觉组对 P+的反应显著大于概念组 [$t(163) = 2.13, p = 0.035, d = 0.64, 95\% CI (0.04, 1.24)$], 对 P-的反应显著小于概念组 [$t(162) = -4.378, p < 0.001, d = -1.32, 95\% CI (-1.94, -0.71)$]; 概念组对 C+的反应显著高于知觉组 [$t(162) = 3.75, p < 0.001, d = 1.14, 95\% CI (0.52, 1.75)$], 对 C-的反应显著小于知觉组 [$t(163) = -2.05, p = 0.042, d = -0.62, 95\% CI (-1.22, -0.02)$] (见图 5)。

综上, 和泛化测试阶段 1 类似, 概念组被试不仅表现出概念性恐惧泛化, 还表现出明显的知觉性恐惧泛化, 而知觉组只表现出了知觉性恐惧泛化, 两组被试在恐惧泛化的路径上表现出非对称性(见图 5)。

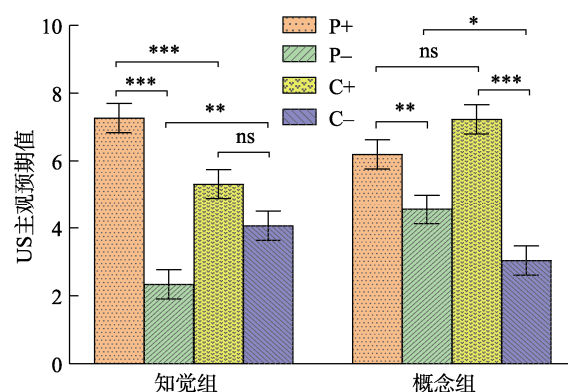


图 5 泛化测试阶段 2 中, 两组被试对泛化刺激的总体 US 预期

3.4.2 皮肤电反应数据

混合线性模型分析结果显示, 模型整体解释了因变量 52.18% 的变异 ($\Omega^2 = 0.5218$), 固定效应解释了模型 27.24% 的变异 ($Marginal R^2 = 0.2724$), 被试的随机截距效应 ($var = 0.0760$) 解释了模型额外的 18.96% 的变异 ($Conditional R^2 = 0.4620$), 表明被试的个体差异对模型产生较大的影响。对固定效应进行方差分析, 结果显示刺激的主效应显著, $F(3, 137.13) = 13.23, p < 0.001, \eta_p^2 = 0.22, 95\% CI (0.10, 0.33)$, 组别的主效应不显著, $F(1, 35.085) = 2.84, p = 0.101$; 刺激类型和组别的交互作用显著, $F(3, 137.28) = 3.12, p = 0.028, \eta_p^2 = 0.06, 95\% CI (0.01, 0.14)$ 。

事后多重比较显示, 知觉组被试对 P+的反应显著高于 P- [$t(138) = 4.79, corrected-p < 0.001, d = 1.41, 95\% CI (0.81, 2.01)$], 而对 C+和 C-的反应无显著差异 [$t(136) = 1.45, corrected-p = 0.889$]; 概念组被试对 C+的反应显著高于 C- [$t(137) = 4.47, corrected-p < 0.001, d = 1.30, 95\% CI (0.72, 1.88)$], 而对 P+和 P-的反应无显著差异 [$t(136) = 1.81, corrected-p = 0.440$]。

在组间比较上, 知觉组对 P+的反应显著高于概念组 [$t(102) = 2.32, p = 0.023, d = 0.85, 95\% CI (0.11, 1.58)$], 对 P-的反应与概念组无显著差异 [$t(101) = 0.11, p = 0.91$], 概念组被试对 C+的反应与知觉组无显著差异 [$t(101) = -0.17, p = 0.86$], 而对 C-的反应显著低于知觉组 [$t(103) = -2.56, p = 0.012, d = -0.94, 95\% CI (-1.68, -0.20)$] (见图 6)。

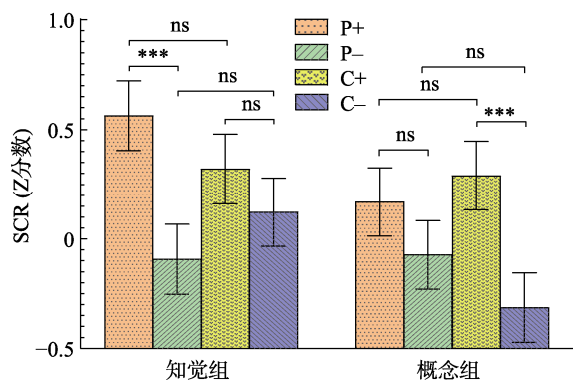


图 6 泛化测试阶段 2 中, 两组被试对泛化刺激的总体 SCR (Z 分数)

3.5 补充分析

前面的分析发现, 在 US 主观预期指标上, 概念组被试不仅对 C+和 C-产生了差异性恐惧反应, 且对 P+和 P-也表现出了显著的差异性恐惧反应, 且该组被试在两个泛化测试阶段对 P+和 C+的反应均无明显差异, 对 P-的反应在第一个阶段也与 C-也无异(在第二阶段 P-显著大于 C-); 而另一方面, 知覚组被试在两个测试阶段都只对 P+和 P-产生了差异性恐惧反应, 对 C+和 C-的反应无差异, 知覚组被试在两个阶段中对 C+的 US 预期也都明显降低, 均显著小于 P+, 而对 C-的反应都明显升高, 均显著高于 P-。总体上, 这些数据表明, 概念组被试不仅产生了概念性恐惧泛化, 还产生了知覚性恐惧泛化, 而这主要是由于概念组被试对 P+的反应升高所致, 知覚组被试仅表现出了知覚性恐惧泛化, 而无概念泛化的迹象。

考虑到 P+是由 C+P+, C0P+, C-P+三者平均计算而来, 为了进一步探明概念组的知覚性恐惧反应升高的具体来源, 对两个泛化测试阶段中被试对各刺激的单独反应, 采用混合线性模型进行进一步分析。

3.5.1 US 主观预期数据

测试阶段 1 的结果显示(一般线性模型), 刺激类型的主效应显著, $F(7, 371) = 33.62, p < 0.001, \eta_p^2 = 0.39, 95\% \text{ CI } (0.31, 0.45)$, 刺激类型和组别的交互效应显著, $F(7, 371) = 9.63, p < 0.001, \eta_p^2 = 0.15, 95\% \text{ CI } (0.08, 0.21)$ 。测试阶段 2 也发现类似的结果(混合线性模型), 刺激类型的主效应显著, $F(7, 335.40) = 36.71, p < 0.001, \eta_p^2 = 0.43, 95\% \text{ CI } (0.35, 0.50)$, 刺激类型和组别的交互效应显著, $F(7, 334.60) = 10.67, p < 0.001, \eta_p^2 = 0.18, 95\% \text{ CI } (0.10, 0.24)$ 。事后多重比较显示, 在两个泛化阶段, 两组

被试(概念组 vs.知覚组)对 C0P+(即, 概念无关而知覚危险刺激的) US 预期均无组间差异[泛化阶段 1, $t(371) = -1.74, p = 0.082$; 泛化阶段 2, $t(338) = -0.97, p = 0.333$], 而对 C-P+的 US 预期存在显著的组间差异[泛化阶段 1, $t(371) = -4.48, p < 0.001, d = -1.30, 95\% \text{ CI } (-1.88, -0.72)$; 泛化测试阶段 2, $t(337) = -4.10, p < 0.001, d = -1.21, 95\% \text{ CI } (-1.80, -0.62)$]。进一步分析发现, 导致两组被试对 C0P+的反应无组间差异的主要原因是由于概念组被试对 C0P+的反应升高所致。在两个泛化阶段, 概念组被试对 C0P+的反应和对 C+P+ (概念和知覚均危险的刺激)的反应都无显著差异[泛化阶段 1, $t(334) = -2.06, \text{corrected-}p = 0.99$; 泛化阶段 2, $t(334) = -1.35, \text{corrected-}p = 0.99$], 而知覚组被试对相应的 C+P0 (即, 概念危险而知覚无关刺激)的反应却显著小于 C+P+ [泛化阶段 1, $t(334) = -5.13, \text{corrected-}p < 0.001, d = -1.43, 95\% \text{ CI } (-1.98, -0.87)$; 泛化阶段 2, $t(334) = -3.46, \text{corrected-}p = 0.017, d = -1.03, 95\% \text{ CI } (-1.63, -0.44)$]。

知覚组被试对 C-P0 和 C-P-的反应在泛化阶段 1 [$t(371) = 2.24, \text{corrected-}p = 0.711$]和泛化阶段 2 [$t(334) = 1.54, \text{corrected-}p = 0.999$]均无显著差异; 概念组被试对 C0P-和 C-P-的反应在泛化阶段 1 无显著差异[$t(371) = 2.15, \text{corrected-}p = 0.914$], 在泛化阶段 2, 对 C0P-的反应显著大于 C-P- [$t(334) = 3.81, \text{corrected-}p = 0.004, d = 1.20, 95\% \text{ CI } (0.57, 1.83)$]。

以上表明, 相对于完全危险的刺激(C+P+), 概念组被试对概念无关而知覚危险的刺激(C0P+)也产生了高水平的恐惧预期反应, 而知覚组被试对相应的概念威胁而知覚无关刺激(C+P0)的恐惧预期却明显下降, 两组被试的泛化反应模式表现出非对称的特点(见图 7, 图 8)。

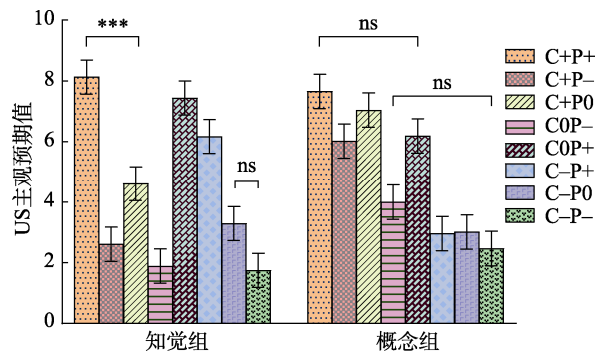


图 7 泛化测试阶段 1 中, 两组被试对各泛化刺激的 US 预期

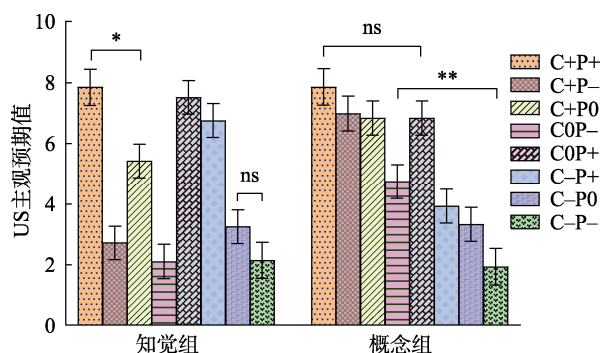


图 8 泛化测试阶段 2 中, 两组被试对各泛化刺激的 US 预期

3.5.2 皮肤电反应数据

在泛化测试 2 的 SCR 指标上(混合线性模型), 刺激类别的主效应显著, $F(7, 64.53) = 7.48, p < 0.001, \eta_p^2 = 0.45, 95\% \text{ CI } (0.23, 0.57)$, 组别的主效应不显著, $F(1, 33.33) = 2.62, p = 0.115$, 刺激类型和组别的交互效应不显著, $F(7, 284.47) = 1.27, p = 0.263$ 。事后多重比较同样发现, 概念组被试对 C0P+ 和 C+P+ 的 SCR 值无显著差异 [$t(100.5) = -2.40, \text{corrected-}p = 0.505$], 而知觉组被试对 C+P0 的反应却显著低于 C+P+ [$t(134.5) = -4.18, \text{corrected-}p = 0.001, d = -1.32, 95\% \text{ CI } (-1.95, -0.68)$]; 知觉组对 C-P0 和 C-P- 的反应无显著差异 [$t(119.2) = 0.46, \text{corrected-}p = 0.999$], 概念组对 C0P- 和 C-P- 的反应也无显著差异 [$t(107.3) = 0.47, \text{corrected-}p = 0.999$]。和 US 预期数据结果类似, 概念组被试对概念无关而知觉危险的刺激(C0P+)表现出了较高的恐惧唤醒, 而知觉组被试对相应的概念威胁而知觉无关刺激(C+P0)的恐惧唤醒却显著下降(见图 9)。

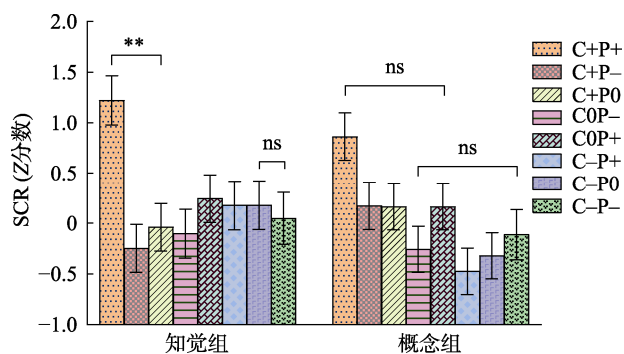


图 9 泛化测试 2 中, 两组被试对各泛化刺激的 SCR (Z 分数)

综合以上分析, 概念组被试产生的知觉性恐惧泛化主要是由于该组被试对 P+ 的恐惧反应升高所致, 尤其是该组被试对 C0P+ 的恐惧反应在 US 预期

和 SCR 指标上都与完全恐惧的 C+P+ 无异。

4 讨论

本研究通过操纵被试恐惧习得前对实验材料的学习经验, 探讨先前经验对恐惧泛化的影响。结果显示, 恐惧习得前的知觉属性判断使得被试产生了明显的知觉性恐惧泛化, 而概念属性判断则使得被试产生了显著的概念性恐惧泛化。此外, 本研究还发现一个有趣的现象, 即概念属性判断的被试对刺激的知觉线索也表现出了明显的恐惧泛化, 特别是对概念无关而知觉危险刺激的恐惧反应尤为强烈, 这种效应在主观评定指标和 SCR 指标上都得到证实, 而知觉属性判断的被试则对刺激的概念线索则无明显的恐惧泛化, 对知觉无关而概念威胁刺激的恐惧反应显著下降, 两组被试在恐惧泛化的路径上表现出非对称的特点。

4.1 概念和知觉信息都是恐惧泛化的重要途径

恐惧的过度泛化是各类焦虑症的核心症状, 异常的恐惧反应和回避给患者的生活造成巨大负担 (American Psychiatric Association, 2013)。以往的研究显示, 人类能够在各种知觉线索间传递恐惧反应, 表现出和动物类似的恐惧泛化梯度 (Dymond et al., 2015), 研究进一步表明, 除了简单的知觉性恐惧泛化, 人类还能够凭借高级认知系统在概念类别 (Dunsmoor et al., 2012)、抽象语义 (Mertens et al., 2021) 以及规则推导 (Ahmed & Lovibond, 2015) 层面泛化恐惧反应。Wang, Mei 等人 (2021) 对恐惧泛化的知觉和概念过程进行了综合研究, 结果显示, 习得阶段, 被试对知觉线索产生了更强的恐惧预期, 且反应更快, 而在泛化阶段, 被试却仅对概念线索产生了明显的恐惧泛化反应。从本实验结果看, 一方面, 我们的研究结果和 Wang, Mei 等人 (2021) 的结果不一致, 我们并没有发现在概念路径上有独特的恐惧泛化优势, 知觉组被试也以知觉线索产生了显著的恐惧泛化, 两研究结果的差异可能源自于实验设置的不同。在 Wang, Mei 等人 (2021) 的研究中, 泛化阶段包含 240 个实验试次, 消退效应非常明显, 在刺激反复呈现的情况下, 相对于涉及到高级认知加工的概念信息, 简单的知觉信息更容易产生习惯化 (Grill-Spector & Weiner, 2014; Hasson et al., 2015), 这可能混淆了实验效应, 导致在 Wang, Mei, et al. 的研究中, P+ 和 P- 无显著差异, 而本研究严格控制了泛化测试的试次数 (每个泛化测试阶段仅包含 8 个试次), 降低了消退效应的影响, 因此捕获到

了 P+和 P-的差异。

另一方面, Wang, Mei 等人(2021)的研究结果表明, 被试在恐惧习得阶段对知觉线索的反应更快, 这和研究结果类似。本研究中, 知觉组被试在泛化测试阶段没有产生任何概念性恐惧泛化的迹象, 而概念组被试却对刺激的知觉信息展现出了强烈的差异性恐惧反应, 这暗示概念组被试似乎无法忽略刺激的知觉信息, 而知觉组被试却可以不对刺激的概念信息进行附带加工, 因此加工过程更短, 进而表现出对知觉线索的反应比对概念线索更快, 这和 Peperkorn 等人(2013)以及 Shibata 等人(2016)的研究结果一致, 表明知觉性信息在诱发恐惧反应中起着基础性作用。

4.2 先前的学习经验调节恐惧泛化的两种路径

人类之所以能适应复杂的生存环境, 一个重要原因在于人类拥有发达的认知系统, 可以对过去的经验进行总结概括来指导未来的行为, 表现出良好的认知迁移能力(Shepard, 1987), 然而在条件性恐惧研究领域, 直接探讨先前经验对恐惧情绪维持和发展的影响的研究还比较少(Dymond et al., 2015)。早期的恐惧泛化研究, 重点关注刺激辨别能力对恐惧泛化的影响, 认为知觉或概念相似性是影响恐惧泛化的核心因素(Dunsmoor et al., 2012; Lissek et al., 2008), 假设人类恐惧泛化是一个单一模式的过程, 而个体差异被认为是随机因素被排除在研究内容之外。然而近年来的研究表明, 即使被试在进行完全相同的恐惧习得后, 在泛化模式上也会表现出多种不同的亚型, 如基于相似性或基于规则的恐惧泛化(Stegmann et al., 2019; Zaman et al., 2023)。这些不同反应模式的差异无法用实验本身的效应来解释, 而是受到被试先前经验的影响。

恐惧泛化的工作神经模型(working neural model)认为, 恐惧泛化是一个涉及到信息输入, 模式匹配(schematic matching)和行为反应的连续过程, 其中模式匹配起作核心作用(Webler et al., 2021)。GS 的知觉信息输入后, 海马体会启动匹配过程, 将当前的输入信息和长时记忆中存储的 CS 信息进行对比, 若 GS 诱发的神经表征与记忆中 CS 的表征发生足够协同(synergy), 海马将启动“模式整合(pattern completion)”过程, 向杏仁核等负责威胁处理的脑区投射神经冲动, 导致恐惧泛化反应; 若 GS 和 CS 的神经表征之间协同不足, 海马体则启动“模式分离(pattern separation)”过程, 向默认模式网络(default mode structures)等负责恐惧抑制的脑区

投放神经冲动, 从而抑制恐惧反应(Feng et al., 2025; Webler et al., 2021)。根据以上模型, 本研究提示, 和传统认知任务类似(Kim & Rehder, 2010), 恐惧泛化中, 先前经验会对后续刺激的识别产生重要影响, 易化后续恐惧泛化过程中与先前经验类似的神经表征的激活, 从而影响个体的恐惧泛化路径。这表明人类的恐惧泛化, 并不是简单地遵循相似性原则的刺激反应连接, 而是会受到更加复杂的自上而下加工过程调节的综合反应(Dunsmoor & Murphy, 2015; 雷怡 等, 2017), 个体因先前经验的不同而在大脑中形成不同的认知结构(Shepard, 1987), 这些不同的内部结构又反过来影响被试对恐惧刺激的解释和识别, 进而影响恐惧传导的路径。

4.3 恐惧泛化中知觉和概念过程的非对称性

本研究还发现一个意料之外的现象, 概念组被试不仅对概念威胁刺激(C+)和安全刺激(C-)产生了明显的差异性恐惧反应, 还对知觉危险刺激(P+)和安全刺激(P-)也产生了显著的差异性恐惧反应, 而知觉组被试对 C+和 C-的恐惧反应却无显著差异。这表明概念组被试不仅产生了概念性恐惧泛化, 还产生了一定程度的知觉性恐惧泛化, 而知觉组被试却无任何概念性恐惧泛化的迹象, 两组被试在恐惧泛化路径上表现出非对称的特点。进一步分析发现, 概念组被试对概念无关但知觉危险刺激(C0P+)的恐惧反应尤为强烈, 而知觉组被试对相应的知觉无关而概念危险刺激(C+P0)的恐惧反应却显著下降。结合 Wang, Mei 等人(2021)的研究, 本研究结果提示, 在恐惧习得和泛化过程中, 刺激知觉信息的加工(如知觉组的情况), 可以在忽略刺激附带的概念信息的情况下进行, 个体可不对刺激的概念信息进行任何有效的处理, 表现为被试对刺激知觉属性的恐惧反应更快, 刺激附带的概念信息也无法引起明显的恐惧反应。然而, 刺激概念属性的加工(如概念组的情况), 却无法忽略刺激的知觉信息, 表现为被试对刺激概念属性的恐惧反应更慢, 且对刺激附带的知觉信息也会展现出强烈的恐惧反应。

根据注意资源竞争模型(attentional resource competition model), 人类的注意资源是有限的, 各信息输入会相互竞争以获得加工资源, 注意资源的分配会受到动机目标(自上而下)的驱动, 个体根据不同任务的要求, 灵活的调整注意资源的分配(Wickens, 1991; Blair et al., 2009), 进而选择性地加工某些信息而忽略另外一些信息(Kan & Thompson-Schill, 2004), 这种不同的注意指向可导

致不同的客体识别结果(Nosofsky, 1986; Kim & Rehder, 2010), 因此对恐惧泛化产生潜在的影响(Dowd et al., 2016)。以往的研究表明, 个体倾向于将注意资源投向熟悉的事物, 对熟悉对象的加工效率也更高(Fernandes & Garcia-Marques, 2020)。基于以上, 本研究可能的解释是, 前训练阶段的属性判断任务使得两组被试分别对刺激的知觉和概念属性更加熟悉, 而这在后续的恐惧习得中调节了被试的注意偏向, 使得知觉组被试更加关注刺激的知觉信息, 而概念组被试则更加关注刺激的概念信息, 而在泛化阶段, 这种对刺激属性的选择性注意导致被试对泛化刺激做出了不同的识别, 从而引起了路径不同的恐惧泛化。此外, 从信息加工过程的角度, 个体对刺激的视觉加工, 开始于对刺激的知觉特征进行分析和整合, 然后才进行归类识别(Gerlach, 2009; Grill-Spector & Weiner, 2014), 所以知觉特征分析是高级概念加工一个必备环节, 因此尽管概念组被试由于先前经验的影响, 更加关注刺激的概念属性, 但是对刺激知觉特征的分析是高级概念加工的必然环节, 因此该组被试对非注意的刺激知觉信息仍进行了潜在的处理。

另一方面, 注意资源的分配同样受到刺激特征(自下而上)的影响(Kan & Thompson-Schill, 2004)。大量研究表明, 恐惧刺激对于注意具有强烈的捕获效应, 即使在非注意状态, 甚至在无意识条件下, 恐惧信息也能得到有效的加工(Vuilleumier, 2005; Finucane & Power, 2010; Mulckhuyse et al., 2013)。综上所述, 概念组被试由于对知觉信息进行了潜在的加工(尽管是在非注意状态下), 而恐惧情绪本身对于这种非注意状态的信息具有放大和强化作用, 从而导致概念组被试对刺激附带的知觉威胁信息也产生了明显的恐惧反应; 而知觉组被试由于完全忽略了刺激的概念信息, 使得其对概念威胁和安全刺激的反应无显著差异。最终, 展现出了两组被试在恐惧泛化路径上的非对称性。

4.4 主观预期和生理唤醒指标的不一致

本研究结果显示, 在整体水平上, 恐惧反应的主观评定指标(US 预期)和生理唤醒指标(SCR)存在不一致的现象。在 US 预期指标上, 两个泛化测试阶段中, 概念组被试对 P+的反应都显著大于 P-, 而在 SCR 指标上则无这种差异。根据恐惧情绪的双通路理论(two-system model), 人类的恐惧反应由两条相对独立又密切联系的神经过路负责, 一条是低级(快速)通路(low road), 主导快速的生理唤醒

(如 SCR)和惊跳反射等, 而另一条是高级(慢速)通路(high road), 主导个体的恐惧预期(如 US 预期)和主观体验等(Webler et al., 2021; LeDoux & Pine, 2016)。一般认为低级通路负责对潜在的威胁进行初步的不精确的加工, 而高级通路则负责对威胁信息进行进一步精细加工, 两条通路相互调节(LeDoux & Pine, 2016)。以往的研究发现, 恐惧反应的低级通路易受到个体情绪状态的影响, 如压力因素会导致恐惧生理激活的显著升高(Hinrichs et al., 2019), 而高级通路则易受到个体认知经验的影响, 如认知重评可显著降低主观恐惧预期(廖素群, 郑希付, 2016)。本实验中, 由于前习得阶段的训练任务为中性(非情绪性)的认知任务, 恐惧的高级通路对这种认知操作的响应更加敏感, 因此产生了实验中观察到的, 概念组对 P+和 P-的差异反应仅体现在主观预期指标上。的确, 混合线性模型的结果也显示, US 预期指标的个体差异(泛化测试 2, 随机截距模型)对模型的解释率仅为 0.56%, 而 SCR 指标的个体差异对模型的解释率则为 18.96%, 表明 US 预期指标在不同被试间较为稳定, 而 SCR 指标的个体差异较大, 和恐惧双通路模型的预期相符合。

4.5 对临床研究的启示

目前, 对临床患者恐惧泛化特异性的研究多集中在知觉性恐惧方面, 发现知觉性恐惧的过度泛化现象广泛存在于各类焦虑症患者中(Fraunfelter et al., 2022), 但对临床患者的概念性恐惧泛化的研究还比较少(Morey et al., 2020; Wang, Wang, et al., 2021)。如前所述, 知觉信息的加工和恐惧的低级通路在神经机制上存在重叠(如, 丘脑, 初级视觉皮层等), 而概念信息的加工则和恐惧的高级通路在神经机制上存在重叠(如, 前额叶皮质, 海马等)(Dunsmoor & Paz, 2015; Webler et al., 2021), 结合本研究的结果, 两种信息线索都是传导恐惧情绪的有效路径, 加之概念性信息相较于知觉信息, 具有更高的不确定性(Dunsmoor & Murphy, 2015), 因此有理由相信概念性恐惧的过度泛化也很可能是各类焦虑症的重要特征之一, 这一点有待后续研究进一步验证。

此外, 以往研究表明, 单独对刺激的知觉(Vervliet & Geens, 2014)或概念线索(Gerdes et al., 2020)进行消退并不能抑制对原 CS 的恐惧反应。本研究提示, 由于知觉和概念线索都能有效地诱发恐惧反应, 因此应考虑对刺激的知觉和概念线索进行联合和分别消退, 从而增强抑制性学习的效果

(Lipp et al., 2020)。另一方面, 本研究也提示, 基于知觉线索的恐惧反应是一种快速而直接的情绪反应, 而基于概念线索的恐惧反应还需要进一步的高级认知加工。考虑到不同类型的焦虑症患者在两种恐惧泛化路径上存在差异(Shiban et al., 2016; Wang, Wang, et al., 2021), 因此在干预基于知觉线索的恐惧时, 应重点对知觉信息进行反复暴露和脱敏, 而在干预由概念性线索诱发的恐惧或回避时, 除了干预基础的知觉性恐惧外, 还应该注重知识教育和社会示范的作用, 以促进个体对概念线索的安全认知(Golkar et al., 2013; Lipp et al., 2020)。

4.6 研究局限与展望

本研究存在一定局限。首先, 本研究的被试是正常人群, 考虑到病人的认知加工过程可能存在异常(Cooper et al., 2022), 因此研究结果是否能推广到临床被试身上有待进一步验证, 未来的研究可在病人群体进一步考察。此外, 本研究对先前学习经验的操作较为简单, 和现实生活中人类实际知识经验的学习有一定差距(Dunsmoor & Murphy, 2015), 后续研究应在更加复杂的知识结构中去验证先前经验对恐惧泛化的作用。再者, 本研究中两类泛化刺激仅有单个水平, 无法检验概念或知觉信息的变化梯度对恐惧泛化的影响(Lissek et al., 2008; Mertens et al., 2021), 后续研究可进一步创新实验范式, 对恐惧泛化的路径和梯度进行联合检验。最后, 本研究仅有行为和生理指标, 缺乏神经活动指标(如, fMRI 或 ERPs 等), 无法对实验效应的大脑机制进行探讨, 后续研究可采用多维实验指标对恐惧泛化中知觉和概念信息加工的神经机制进行探索。

5 结论

本研究发现, 个体的先前学习经验会显著影响后续的恐惧泛化路径, 知觉属性判断会导致恐惧沿着刺激的知觉线索泛化, 而概念属性判断会导致恐惧沿着刺激的概念线索泛化, 两种路径都是传导恐惧反应的有效路径。此外, 基于概念线索的恐惧泛化, 包含了部分基于知觉线索的恐惧泛化, 而基于知觉线索的恐惧泛化, 却可以不涉及明显的概念加工过程, 两种恐惧泛化在路径上表现出非对称性。本研究加深了对恐惧泛化中信息加工过程的理解, 提示研究者应将被试的先前学习经验纳入到恐惧研究中, 以进一步探明人类高级认知过程在恐惧泛化的产生和维持中扮有的作用。

参 考 文 献

- Ahmed, O., & Lovibond, P. F. (2015). The impact of previously learned feature-relevance on generalisation of conditioned fear in humans. *Journal of Behavior Therapy and Experimental Psychiatry*, 46, 59–65. <https://doi.org/10.1016/j.jbtep.2014.08.001>
- American Psychiatric Association. (2013). *Diagnostic and statistical manual of mental disorders: DSM-5*. Washington, DC: American Psychiatric Publishing.
- Bauer, E. A., MacNamara, A., Sandre, A., Lonsdorf, T. B., Weinberg, A., Morris, J., & van Reekum, C. M. (2020). Intolerance of uncertainty and threat generalization: A replication and extension. *Psychophysiology*, 57(5), e13546. <https://doi.org/10.1111/psyp.13546>
- Blair, M. R., Watson, M. R., Walshe, R. C., & Maj, F. (2009). Extremely selective attention: eye-tracking studies of the dynamic allocation of attention to stimulus features in categorization. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 35(5), 1196–1206. <https://doi.org/10.1037/a0016272>
- Cooper, S. E., van Dis, E. A. M., Hagenars, M. A., Krypotos, A. M., Nemeroff, C. B., Lissek, S., ... Dunsmoor, J. E. (2022). A meta-analysis of conditioned fear generalization in anxiety-related disorders. *Neuropsychopharmacology*, 47(9), 1652–1661. <https://doi.org/10.1038/s41386-022-01332-2>
- Dowd, E. W., Mitroff, S. R., & LaBar, K. S. (2016). Fear generalization gradients in visuospatial attention. *Emotion*, 16(7), 1011–1018. <https://doi.org/10.1037/emo0000197>
- Dunsmoor, J. E., & LaBar, K. S. (2013). Effects of discrimination training on fear generalization gradients and perceptual classification in humans. *Behavioral Neuroscience*, 127(3), 350–356. <https://doi.org/10.1037/a0031933>
- Dunsmoor, J. E., Martin, A., & LaBar, K. S. (2012). Role of conceptual knowledge in learning and retention of conditioned fear. *Biological Psychology*, 89(2), 300–305. <https://doi.org/10.1016/j.biopsycho.2011.11.002>
- Dunsmoor, J. E., & Murphy, G. L. (2014). Stimulus typicality determines how broadly fear is generalized. *Psychological Science*, 25(9), 1816–1821. <https://doi.org/10.1177/0956797614535401>
- Dunsmoor, J. E., & Murphy, G. L. (2015). Categories, concepts, and conditioning: How humans generalize fear. *Trends in Cognitive Sciences*, 19(2), 73–77. <https://doi.org/10.1016/j.tics.2014.12.003>
- Dunsmoor, J. E., & Paz, R. (2015). Fear generalization and anxiety: Behavioral and neural mechanisms. *Biological Psychiatry*, 78(5), 336–343. <https://doi.org/10.1016/j.biopsycho.2015.04.010>
- Dymond, S., Dunsmoor, J. E., Vervliet, B., Roche, B., & Hermans, D. (2015). Fear generalization in humans: Systematic review and implications for anxiety disorder research. *Behavior Therapy*, 46(5), 561–582. <https://doi.org/10.1016/j.beth.2014.10.001>
- Fan, M., Zhang, D., Zhao, S., Xie, Q., Chen, W., Jie, J., ... Zheng, X. (2022). Stimulus diversity increases category-based fear generalization and the effect of intolerance of uncertainty. *Behaviour Research and Therapy*, 159, 104201. <https://doi.org/10.1016/j.brat.2022.104201>
- Feng, B., Zeng, L., Hu, Z., Fan, X., Ai, X., Huang, F., & Zheng, X. (2025). Global precedence effect in fear generalization and the role of trait anxiety and intolerance of uncertainty. *Behaviour Research and Therapy*, 184,

104669. <https://doi.org/10.1016/j.brat.2024.104669>
- Fernandes, A. C., & Garcia-Marques, T. (2020). A meta-analytical review of the familiarity temporal effect: Testing assumptions of the attentional and the fluency-attributional accounts. *Psychological Bulletin*, 146(3), 187–217. <https://doi.org/10.1037/bul0000222>
- Finucane, A. M., & Power, M. J. (2010). The effect of fear on attentional processing in a sample of healthy females. *Journal of Anxiety Disorders*, 24(1), 42–48. <https://doi.org/10.1016/j.janxdis.2009.08.005>
- Franconeri, S. L., Alvarez, G. A., & Cavanagh, P. (2013). Flexible cognitive resources: Competitive content maps for attention and memory. *Trends in Cognitive Sciences*, 17(3), 134–141. <https://doi.org/10.1016/j.tics.2013.01.010>
- Fraunfelder, L., Gerdes, A. B. M., & Alpers, G. W. (2022). Fear one, fear them all: A systematic review and meta-analysis of fear generalization in pathological anxiety. *Neuroscience and Biobehavioral Reviews*, 139, 104707. <https://doi.org/10.1016/j.neubiorev.2022.104707>
- Gerdes, A. B. M., Fraunfelder, L., & Alpers, G. W. (2020). Hear it, fear it: Fear generalizes from conditioned pictures to semantically related sounds. *Journal of Anxiety Disorders*, 69, 102174. <https://doi.org/10.1016/j.janxdis.2019.102174>
- Gerlach, C. (2009). Category-specificity in visual object recognition. *Cognition*, 111(3), 281–301. <https://doi.org/10.1016/j.cognition.2009.02.005>
- Goldstone, R. L., Steyvers, M., Spencer-Smith, J., & Kersten, A. (2000). Interactions between perceptual and conceptual learning. In E. Dietrich & A. B. Markman (Eds.), *Cognitive dynamics: Conceptual and representational change in humans and machines*. (pp. 191–228). Lawrence Erlbaum Associates Publishers.
- Golkar, A., Selbing, I., Flygare, O., Ohman, A., & Olsson, A. (2013). Other people as means to a safe end: Vicarious extinction blocks the return of learned fear. *Psychological Science*, 24(11), 2182–2190. <https://doi.org/10.1177/0956797613489890>
- Gong, X., Xie, X. Y., Xu, R., & Luo, Y. J. (2010). Psychometric properties of the Chinese Versions of DASS-21 in Chinese college students. *Chinese Journal of Clinical Psychology*, 18(4), 443–446. <https://doi.org/10.16128/j.cnki.1005-3611.2010.04.020>
- [龚栩, 谢嘉瑶, 徐蕊, 罗跃嘉. (2010). 抑郁-焦虑-压力量表简体中文版(DASS-21)在中国大学生中的测试报告. *中国临床心理学杂志*, 18(4), 443–446.]
- Grill-Spector, K., & Weiner, K. S. (2014). The functional architecture of the ventral temporal cortex and its role in categorization. *Nature Reviews Neuroscience*, 15(8), 536–548. <https://doi.org/10.1038/nrn3747>
- Hasson, U., Chen, J., & Honey, C. J. (2015). Hierarchical process memory: Memory as an integral component of information processing. *Trends in Cognitive Sciences*, 19(6), 304–313. <https://doi.org/10.1016/j.tics.2015.04.006>
- Hinrichs, R., van Rooij, S. J., Michopoulos, V., Schultebraucks, K., Winters, S., Maples-Keller, J., ... Jovanovic, T. (2019). Increased skin conductance response in the immediate aftermath of trauma predicts PTSD risk. *Chronic Stress*, 3, 1–11. <https://doi.org/10.1177/2470547019844441>
- Jepma, M., & Wager, T. D. (2015). Conceptual conditioning: Mechanisms mediating conditioning effects on pain. *Psychological Science*, 26(11), 1728–1739. <https://doi.org/10.1177/0956797615597658>
- Kan, I. P., & Thompson-Schill, S. L. (2004). Selection from perceptual and conceptual representations. *Cognitive, Affective, & Behavioral Neuroscience*, 4(4), 466–482. <https://doi.org/10.3758/cabn.4.4.466>
- Kim, S., & Rehder, B. (2010). How prior knowledge affects selective attention during category learning: An eyetracking study. *Memory & Cognition*, 39(4), 649–665. <https://doi.org/10.3758/s13421-010-0050-3>
- LeDoux, J. E., & Pine, D. S. (2016). Using neuroscience to help understand fear and anxiety: A two-system framework. *American Journal of Psychiatry*, 173(11), 1083–1093. <https://doi.org/10.1176/appi.ajp.2016.16030353>
- Lee, Y. I., Lee, D., Kim, H., Kim, M. J., Jeong, H., Kim, D., ... Choi, S. H. (2024). Overgeneralization of conditioned fear in patients with social anxiety disorder. *Frontiers in Psychiatry*, 15, 1415135. <https://doi.org/10.3389/fpsy.2024.1415135>
- Lei, Y., Wang, J. X., Cheng, Q. F., Zhang, W. H., & Mei, E. (2017). The influence mechanism of categories and concepts on fear generalization. *Journal of Psychological Science*, 40(5), 1266–1273. <https://doi.org/10.16719/j.cnki.1671-6981.20170537>
- [雷怡, 王金霞, 陈庆飞, 张文海, 梅颖. (2017). 分类和概念对恐惧泛化的影响机制. *心理科学*, 40(5), 1266–1273.]
- Liao, S. Q., & Zheng, X. F. (2016). Inhibition of cognitive reappraisal on the negative valence facilitates extinction in conditioned fear. *Acta Psychologica Sinica*, 48(4), 352–361. <https://doi.org/10.3724/sp.J.1041.2016.00352>
- [廖素群, 郑希付. (2016). 认知重评对负性价价的抑制促进条件性恐惧消退. *心理学报*, 48(4), 352–361.]
- Lipp, O. V., Waters, A. M., Luck, C. C., Ryan, K. M., & Craske, M. G. (2020). Novel approaches for strengthening human fear extinction: The roles of novelty, additional USs, and additional GSs. *Behaviour Research and Therapy*, 124, 103529. <https://doi.org/10.1016/j.brat.2019.103529>
- Lissek, S., Biggs, A. L., Rabin, S. J., Cornwell, B. R., Alvarez, R. P., Pine, D. S., & Grillon, C. (2008). Generalization of conditioned fear-potentiated startle in humans: Experimental validation and clinical relevance. *Behaviour Research and Therapy*, 46(5), 678–687. <https://doi.org/10.1016/j.brat.2008.02.005>
- Lopresto, D., Schipper, P., & Homberg, J. R. (2016). Neural circuits and mechanisms involved in fear generalization: Implications for the pathophysiology and treatment of posttraumatic stress disorder. *Neuroscience and Biobehavioral Reviews*, 60, 31–42. <https://doi.org/10.1016/j.neubiorev.2015.10.009>
- Mertens, G., Bouwman, V., & Engelhard, I. M. (2021). Conceptual fear generalization gradients and their relationship with anxious traits: Results from a Registered Report. *International Journal of Psychophysiology*, 170, 43–50. <https://doi.org/10.1016/j.ijpsycho.2021.09.007>
- Morey, R. A., Haswell, C. C., Stjepanovic, D., Mid-Atlantic, M. W., Dunsmoor, J. E., & LaBar, K. S. (2020). Neural correlates of conceptual-level fear generalization in posttraumatic stress disorder. *Neuropsychopharmacology*, 45(8), 1380–1389. <https://doi.org/10.1038/s41386-020-0661-8>
- Mulckhuysen, M., Crombez, G., & Van der Stigchel, S. (2013). Conditioned fear modulates visual selection. *Emotion*, 13(3), 529–536. <https://doi.org/10.1037/a0031076>
- Nosofsky, R. M. (1986). Attention, similarity, and the identification-categorization relationship. *Journal of Experimental Psychology: General*, 115(1), 39–57.
- Peperckorn, H. M., Alpers, G. W., & Mühlberger, A. (2013). Triggers of fear: Perceptual cues versus conceptual information in spider phobia. *Journal of Clinical Psychology*, 70(7), 704–714. <https://doi.org/10.1002/jclp.22057>
- Resnik, J., Sobel, N., & Paz, R. (2011). Auditory aversive

- learning increases discrimination thresholds. *Nature Neuroscience*, 14(6), 791–796. <https://doi.org/10.1038/nn.2802>
- Reutter, M., & Gamer, M. (2023). Individual patterns of visual exploration predict the extent of fear generalization in humans. *Emotion*, 23(5), 1267–1280. <https://doi.org/10.1037/emo0001134>
- Scheveneels, S., Boddez, Y., & Hermans, D. (2021). Predicting clinical outcomes via human fear conditioning: A narrative review. *Behaviour Research and Therapy*, 142, 103870. <https://doi.org/10.1016/j.brat.2021.103870>
- Sep, M. S. C., Steenmeijer, A., & Kennis, M. (2019). The relation between anxious personality traits and fear generalization in healthy subjects: A systematic review and meta-analysis. *Neuroscience & Biobehavioral Reviews*, 107, 320–328. <https://doi.org/10.1016/j.neubiorev.2019.09.029>
- Shepard, R. N. (1987). Toward a universal law of generalization for psychological science. *Science*, 237(4820), 1317–1323. <http://www.jstor.org/stable/1700004>
- Shiban, Y., Peperkorn, H., Alpers, G. W., Pauli, P., & Mühlberger, A. (2016). Influence of perceptual cues and conceptual information on the activation and reduction of claustrophobic fear. *Journal of Behavior Therapy and Experimental Psychiatry*, 51, 19–26. <https://doi.org/10.1016/j.jbtep.2015.11.002>
- Stegmann, Y., Schiele, M. A., Schumann, D., Lonsdorf, T. B., Zwanzger, P., Romanos, M., ... Pauli, P. (2019). Individual differences in human fear generalization—pattern identification and implications for anxiety disorders. *Translational Psychiatry*, 9(1), 307. <https://doi.org/10.1038/s41398-019-0646-8>
- Vervliet, B., & Geens, M. (2014). Fear generalization in humans: Impact of feature learning on conditioning and extinction. *Neurobiology of Learning and Memory*, 113, 143–148. <https://doi.org/10.1016/j.nlm.2013.10.002>
- Vervliet, B., Kindt, M., Vansteenwegen, D., & Hermans, D. (2010a). Fear generalization in humans: Impact of prior non-fearful experiences. *Behaviour Research and Therapy*, 48(11), 1078–1084. <https://doi.org/10.1016/j.brat.2010.07.002>
- Vervliet, B., Kindt, M., Vansteenwegen, D., & Hermans, D. (2010b). Fear generalization in humans: Impact of verbal instructions. *Behaviour Research and Therapy*, 48(1), 38–43. <https://doi.org/10.1016/j.brat.2009.09.005>
- Vuilleumier, P. (2005). How brains beware: Neural mechanisms of emotional attention. *Trends in Cognitive Sciences*, 9(12), 585–594. <https://doi.org/10.1016/j.tics.2005.10.011>
- Wang, J., Mei, E., Wu, Q., Xie, T., Dou, H., & Lei, Y. (2021). Influence of perceptual and conceptual information on fear generalization: A behavioral and event-related potential study. *Cognitive, Affective, & Behavioral Neuroscience*, 21(5), 1054–1065. <https://doi.org/10.3758/s13415-021-00912-x>
- Wang, J., Wang, Y., Liao, M., Zou, Y., Lei, Y., & Zhu, Y. (2021). Conditioned generalisation in generalised anxiety disorder: The role of concurrent perceptual and conceptual cues. *Cognition and Emotion*, 35(8), 1516–1526. <https://doi.org/10.1080/02699931.2021.1982677>
- Wang, L., Liu, W. T., Zhu, X. Z., Wang, Y. P., Li, L. Y., Yang, Y. L., & George, R. A. (2014). Validity and reliability of the Chinese Version of the Anxiety Sensitivity Index-3 in healthy adult women. *Chinese Mental Health Journal*, 28(10), 767–771.
- [王雷, 刘婉婷, 朱熊兆, 王瑜萍, 李玲艳, 杨玉玲, George, R. A. (2014). 焦虑敏感指数量表 3 版中文版测评健康成年女性的效度和信度. *中国心理卫生杂志*, 28(10), 767–771.]
- Webler, R. D., Berg, H., Phong, K., Tuominen, L., Holt, D. J., Morey, R. A., ... Lissek, S. (2021). The neurobiology of human fear generalization: Meta-analysis and working neural model. *Neuroscience and Biobehavioral Reviews*, 128, 421–436. <https://doi.org/10.1016/j.neubiorev.2021.06.035>
- Wong, A. H. K., & Lovibond, P. F. (2021). Breakfast or bakery? The role of categorical ambiguity in overgeneralization of learned fear in trait anxiety. *Emotion*, 21(4), 856–870. <https://doi.org/10.1037/emo0000739>
- Wickens, C. D. (1991). Processing resources and attention. In D. Damos (Ed.), *Multiple task performance* (pp. 3–34). CRC Press.
- Zaman, J., Yu, K., & Lee, J. C. (2023). Individual differences in stimulus identification, rule induction, and generalization of learning. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 49(6), 1004–1017. <https://doi.org/10.1037/xlm0001153>
- Zhang, J., Wang, J., Wang, Y., Zhang, D., Li, H., & Lei, Y. (2024). Sleep deprivation increases the generalization of perceptual and concept-based fear: An fNIRS study. *Journal of Anxiety Disorders*, 105, 102892. <https://doi.org/10.1016/j.janxdis.2024.102892>
- Zhang, Y. J., Song, J. B., Gao, Y. T., Wu, S. J., Song, L., & Miao, D. M. (2017). Reliability and validity of the Intolerance of Uncertainty Scale-Short form in university students. *Chinese Journal of Clinical Psychology*, 25(2), 285–288. <https://doi.org/10.16128/j.cnki.1005-3611.2017.02.020>
- [张亚娟, 宋继波, 高云涛, 武圣君, 宋蕾, 苗丹民. (2017). 无法忍受不确定性量表(简版)在中国大学生中的信效度检验. *中国临床心理学杂志*, 25(2), 285–288.]

Perceptual or conceptual? Modulation of fear generalization pathways by prior learning experience

FENG Biao^{1,2}, ZHANG Donghuan¹, CHEN Wei¹, ZENG Ling¹, WU Xiaoyue¹, HUANG Junling¹, ZHENG Xifu¹

(¹ Key Laboratory of Brain, Cognition and Education Sciences, Ministry of Education, China; School of Psychology, Center for Studies of Psychological Application, and Guangdong Key Laboratory of Mental Health and Cognitive Science, South China Normal University, China) (² School of Information Technology in Education, South China Normal University, China)

Abstract

Excessive fear generalization is a core symptom of various anxiety disorders, where unbounded fear or

avoidance imposes significant burdens on patients' lives. Previous studies have shown that human fear responses can propagate along the perceptual information (perceptual fear generalization) or conceptual information (conceptual fear generalization) of stimuli. In general, any stimulus inherently contains both perceptual and conceptual attributes. Following fear conditioning to stimulus compounds, the factors that regulate the two pathways of fear generalization are an intriguing topic. Recent research suggests that the object recognition process plays an important role in fear generalization. Given that attention orientation significantly influences object recognition, this study hypothesizes that prior experience may shape individuals' attentional biases, thereby affecting their recognition and categorization of generalized stimuli and ultimately regulating the pathways of fear generalization.

This study employed a classical differential fear conditioning paradigm using unconditioned stimulus (US) expectancy and skin conductance response (SCR) as measures to examine whether pre-acquisition learning experiences modulate the two fear generalization pathways. A set of 100 images was used as experimental material. These images were systematically categorized into eight groups: C+P+, C-P-, C-P+, C+P-, C+P0, C-P0, C0P-, C0P+ (where "C" = Conceptual, "P" = Perceptual, "+" = Threatening, "-" = Safe, and 0 = Neutral). Fifty university students were recruited and randomly assigned to perceptual or conceptual groups. During the pre-acquisition learning phase (40 trials), the perceptual group made judgments on the perceptual attributes of the stimuli (color discrimination) to enhance perceptual processing, whereas the conceptual group judged conceptual attributes (object categorization) to strengthen conceptual processing. Subsequently, both the groups underwent identical fear acquisition and generalization tests. In the fear acquisition phase (16 trials), participants underwent differential fear conditioning to compound stimuli that were either absolutely threatening or absolutely safe in both perceptual and conceptual dimensions (i.e., C+P+, C-P-). During the generalization test (two test phases, each with eight trials), eight different generalization stimuli blending perceptual and conceptual features (as previously described), were presented to assess fear generalization.

The results revealed that participants in the perceptual group exhibited significant perceptual fear generalization, whereas those in the conceptual group showed pronounced conceptual fear generalization. These findings indicated that prior experience significantly modulated the pathways of fear generalization and confirmed both pathways as effective routes for fear generalization. Additionally, an intriguing finding emerged: apart from conceptual fear generalization, the conceptual group also displayed a tendency for perceptual fear generalization, whereas the perceptual group showed no signs of conceptual fear generalization. This asymmetric pattern was consistently observed in both the US expectancy and SCR measures, demonstrating a robust effect. These findings can be explained by the differences in information processing and attentional biases between the two groups, suggesting distinct roles of perceptual and conceptual information in eliciting human fear responses. Theoretical and clinical implications were discussed.

Keywords prior experience, perceptual information, conceptual information, attentional bias, fear generalization pathways

补充材料：各实验阶段模型拟合和对比

附表 1 恐惧习得阶段模型拟合和对比

实验阶段及 结果变量	模型名称	模型结构	拟合结果	模型选用
\	model_full (全模型)	结果变量 = 性别 + 年龄 + DASS21_A + DASS21_S + DASS21_D + IU_P + IU_I + ASI3_P + ASI3_S + ASI3_C + 前习得阶段的正确率 + 刺激反应时 + 刺激类别 + 组别 + 组别×刺激类别 + (1 被试)	\	\
习得阶段 US 主观预期	model_AQ_US0	全模型	singular	无效模型
	model_AQ_US1	在全模型上去掉被试的随机截距(一般线性模型)	拟合显著。 $F(15, 84) = 13.64, p < 0.001$ Multiple R-squared = 0.7091 Adjusted R-squared = 0.6573	选用本模型
习得阶段 SCR	model_AQ_SCR0	全模型	singular	无效模型
	model_AQ_SCR1	在全模型上去掉被试的随机截距(一般线性模型)	拟合显著。 $F(15, 80) = 4.84, p < 0.001$ Multiple R-squared = 0.4755 Adjusted R-squared = 0.3772	选用本模型

注：反应时(ms)和各量表分数均中心化后再纳入模型

附表 2 泛化测试阶段模型拟合和对比(整体)

实验阶段及 结果变量	模型名称	模型结构	拟合结果	模型选用
\	model_full (全模型)	结果变量 = 性别 + 年龄 + DASS21_A + DASS21_S + DASS21_D + IU_P + IU_I + ASI3_P + ASI3_S + ASI3_C + 前习得阶段的正确率 + 刺激反应时 + 刺激类别 + 组别 + 组别×刺激类别 + 刺激顺序 + (1 + 刺激顺序 被试) + (1 刺激顺序)	\	\
泛化测试阶段 1 US 主观预期	model_GE1_US0	全模型	singular	无效模型
	model_GE1_US1	在全模型上去掉被试随机截距和刺激顺序随机斜率的相关	singular	无效模型
	model_GE1_US2	在全模型上去掉刺激顺序的随机斜率	singular	无效模型
	model_GE1_US3	在全模型上去掉刺激顺序的固定效应	singular	无效模型
	model_GE1_US4	在全模型上去掉刺激顺序的随机截距	singular	无效模型
	model_GE1_US5	在全模型上去掉刺激顺序的随机截距, 并去掉刺激顺序的随机斜率	singular	无效模型
	model_GE1_US6	在全模型上去掉所有随机效应(一般线性模型)	拟合显著。 $F(20, 179) = 9.02, p < 0.001$ Multiple R-squared = 0.5018 Adjusted R-squared = 0.4462	选用本模型
泛化测试阶段 2 US 主观预期	model_GE2_US0	全模型	singular	无效模型
	model_GE2_US1	在全模型上去掉被试随机截距和刺激顺序随机斜率的相关	singular	无效模型
	model_GE2_US2	在全模型上去掉刺激顺序的随机斜率	singular	无效模型
	model_GE2_US3	在全模型上去掉刺激顺序的固定效应	singular	无效模型
	model_GE2_US4	在全模型上去掉刺激顺序的随机截距	singular	无效模型
	model_GE2_US5	在全模型上去掉刺激顺序的随机截距, 并去掉刺激顺序的随机斜率(即, 随机效应仅保留被试的随机截距项)	收敛。 Var:被试 (Intercept)= 0.03384 ICC:被试 = 0.01192 Var: Residual = 2.80592 $R_{(m)}^2 = 0.52854$ $R_{(c)}^2 = 0.53416$ $\Omega^2 = 0.56378$	选用本模型
泛化测试阶段 2 SCR	model_GE2_SCR0	全模型	singular	无效模型
	model_GE2_SCR1	在全模型上去掉被试随机截距和刺激顺序随机斜率的相关	singular	无效模型
	model_GE2_SCR2	在全模型上去掉刺激顺序的随机斜率	singular	无效模型
	model_GE2_SCR3	在全模型上去掉刺激顺序的固定效应	singular	无效模型
	model_GE2_SCR4	在全模型上去掉刺激顺序的随机截距	singular	无效模型
	model_GE2_SCR5	在全模型上去掉刺激顺序的随机截距, 并去掉刺激顺序的随机斜率(即, 随机效应仅保留被试的随机截距项)	收敛。 Var:被试(Intercept) = 0.07608 ICC:被试 = 0.26055 Var: Residual = 0.21593 $R_{(m)}^2 = 0.27244$ $R_{(c)}^2 = 0.46200$ $\Omega^2 = 0.52179$	选用本模型

注: 反应时(ms)和各量表分数均中心化后再纳入模型

附表 3 泛化测试阶段模型拟合和对比(各刺激)

实验阶段及 结果变量	模型名称	模型结构	拟合结果	模型选用
\	model_full (全模型)	结果变量 = 性别 + 年龄 + DASS21_A + DASS21_S + DASS21_D + IU_P + IU_I + ASI3_P + ASI3_S + ASI3_C + 前习得阶段的正确率 + 刺激反 应时 + 刺激类别 + 组别 + 组别×刺 激类别 + 刺激顺序 + (1 + 刺激顺序 被试) + (1 刺激顺序)	\	\
泛化测试阶段 1 US 主观预期	model_GE1_US0	全模型	singular, 未收敛	无效模型
		在全模型上去掉被试随机截距和刺激顺 序随机斜率的相关	singular	无效模型
	model_GE1_US1	在全模型上去掉刺激顺序的随机斜率	singular	无效模型
	model_GE1_US2	在全模型上去掉刺激顺序的固定效应	singular	无效模型
	model_GE1_US3	在全模型上去掉刺激顺序的随机截距	singular	无效模型
	model_GE1_US4	在全模型上去掉刺激顺序的随机截距, 并 去掉刺激顺序的随机斜率(仅保留被试的 随机截距项)	singular	无效模型
	model_GE1_US5	在全模型上去掉所有随机效应(一般线性 模型)	拟合显著。 $F(28, 371) = 11.39, p < 0.001$ Multiple R-squared = 0.4622 Adjusted R-squared = 0.4216	选用本模型
泛化测试阶段 2 US 主观预期	model_GE2_US0	全模型	singular, 未收敛	无效模型
		在全模型上去掉被试随机截距和刺激顺 序随机斜率的相关	singular	无效模型
	model_GE2_US1	在全模型上去掉刺激顺序的随机斜率	singular	无效模型
	model_GE2_US2	在全模型上去掉刺激顺序的固定效应	singular, 未收敛	无效模型
	model_GE2_US3	在全模型上去掉刺激顺序的随机截距	收敛, 但 Var:被试 (Intercept) 接近于 0(1.637e-05)	无效模型
	model_GE2_US4	在全模型上去掉刺激顺序的随机截距, 并 去掉刺激顺序的随机斜率(即, 随机效应 仅保留被试的随机截距项)	收敛。 Var: 被试 (Intercept) = 0.14100 Var: Residual = 5.48831 ICC: 被试 = 0.02505 $R_{(m)}^2 = 0.46142$ $R_{(c)}^2 = 0.47491$ $\Omega^2 = 0.49976$	选用本模型

续表

实验阶段及 结果变量	模型名称	模型结构	拟合结果	模型选用
泛化测试阶段 2 SCR	model_GE2_SCR0	全模型	收敛。 Var: 被试(Intercept) = 0.15560 Var: 被试 刺激顺序 = 0.00668 ICC: 被试 = 0.14648 Var: 刺激顺序 (Intercept) = 0.00599 ICC: 刺激顺序 = 0.00564 Var: Residual = 0.90061 AIC = 1114.69 BIC = 1247.38 R _(m) ² = 0.18195 R _(c) ² = 0.22670 Omega ² = 0.26754	这几个收敛的模型之间的似然比检验均无显著差异, 选用模型残差和 AIC、BIC 数值都较小, 而 Marginal R ² , Conditional R ² 和 Omega ² 数值都较大的 model_GE2_SCR3
	model_GE2_SCR1	在全模型上去掉被试随机截距和刺激顺序随机斜率的相关	singular	
	model_GE2_SCR2	在全模型上去掉刺激顺序的随机斜率	收敛。 Var: 被试(Intercept) = 0.00575 ICC: 被试 = 0.00604 Var: 刺激顺序 (Intercept) = 0.00457 ICC: 刺激顺序 = 0.00480 Var: Residual = 0.94164 AIC = 1112.164 BIC = 1237.049 R _(m) ² = 0.18295 R _(c) ² = 0.19181 Omega ² = 0.20580	
	model_GE2_SCR3	在全模型上去掉刺激顺序的固定效应	Var: 被试(Intercept) = 0.14800 Var: 被试 刺激顺序 = 0.00631 ICC: 被试 = 0.14066 Var: 刺激顺序 = 0.00210 ICC: 刺激顺序 = 0.00200 Var: Residual = 0.90204 AIC = 1107.952 BIC = 1236.739 R _(m) ² = 0.18138 R _(c) ² = 0.22109 Omega ² = 0.26179	
	model_GE2_SCR4	在全模型上去掉刺激顺序的随机截距(即, 随机效应保留被试的随机截距和刺激顺序的随机斜率)	收敛。 Var: 被试(Intercept) = 0.151583 ICC: 被试 = 0.14333 Var: 被试 刺激顺序 = 0.00650 Var: Residual = 0.90602 AIC = 1112.861 BIC = 1241.648 R _(m) ² = 0.18004 R _(c) ² = 0.21832 Omega ² = 0.25947	
	model_GE2_SCR5	在全模型上去掉刺激顺序的随机截距, 并去掉刺激顺序的随机斜率(仅保留被试的随机截距项)	收敛。 Var: 被试(Intercept) = 0.00524 ICC: 被试 = 0.00552 Var: Residual = 0.94500 AIC = 1110.259 BIC = 1231.241 R _(m) ² = 0.18142 R _(c) ² = 0.18594 Omega ² = 0.20055	

注: 反应时(ms)和各量表分数均中心化后再纳入模型