

《心理科学进展》审稿意见与作者回应

题目：从利己到伤害：负创造力向恶意创造力的演化路径与干预

作者：周淑金 朱静怡 涂洪悦 李燕

第一轮

审稿人 1 意见：

本文讨论了从负创造力向恶意创造力的转化路径，并据此，提出了早期干预框架。作者总结认为，创造力由新颖性与适用性的平衡构成，但在儿童早期创造力可能呈现偏离适用性的“负创造力”表现，甚至进一步演变为“恶意创造力”，基于发展的视角作者提出应将干预前移至负创造力阶段，提出了“认知—动机—发展”干预模型。这对于恶意创造力的早期干预与和谐社会的建设具有非常重要的意义。全文引用文献新颖，反映了当前的研究进展，表述流畅，可读性强，在梳理文献的基础上提出了自己的观点或干预模型，观点鲜明，写作规范。

同时，建议在以下几个方面做些微调：

意见 1：引言部分可以简化，概括说明相关的社会背景和学术背景，引出本文所论述的主题即可，这样也可以缩短文章篇幅。

回复：我们已按您的建议对引言部分进行了大幅精简。修改后的引言聚焦于：（1）创造力被误用的现实危害；（2）负创造力与恶意创造力的概念混淆问题；（3）儿童早期作为干预窗口期的特殊意义。删减了与国家战略相关的背景铺陈，使主题切入更加直接。具体修改见稿件引言部分（红色字体标注）。

意见 2：个别用语希望再核对一下，例如将独创性与原创性并提，是否符合原作者的意思。

回复：感谢您的这一重要提醒。基于您的建议，在创造力定义部分，我们参考 Mayer（1999）的经典定义，并将“独创性”与“原创性”统一表述为“新颖性”，以避免术语混用。相关修改见稿件“2.1 创造力：新颖性与适用性的平衡”部分（红色字体标注）。

意见 3：第二部分可以简化分析创造力、负创造力与恶意创造力的关系，包括区别与联系。

回复：感谢您对本研究提出的宝贵意见和建议，我们已对第二部分进行了全面精简。修改后的 2.1 节聚焦于创造力的核心定义，2.2 节则在此基础上集中阐述负创造力与恶意创造力的概念内涵与动机差异，删减了与主题关联较弱的文献综述，使论述更加紧凑清晰。具体修改见稿件第二部分（红色字体标注）。

意见 4：另见文中批注。

回复：感谢您对本研究提出的宝贵意见和建议，我们已逐一处理文中的所有批注，并对相应内容进行了修改。所有修改均以红色字体标注，便于您复核。

审稿人 2 意见：

作者关注恶意创造性行为对未成年个体的潜在危害，并构建了恶意创造性行为的演化模

式,同时提出了针对未成年个体的干预模式构想。这些探索为该领域的发展提供了新的思路,具有一定的实际应用价值。然而,研究中也存在一些问题,例如主要概念的定义和区分不够明确,部分关键结论缺乏足够的实证支持。

意见 1: 当前研究的两个核心变量: 负创造力与恶意创造力, 在概念上存在一定重叠。

恶意创造力的完整概念是指, 通过多种新颖方式蓄意伤害他人(或自己、组织和社会)的创造力, 包括伤害身心健康、剥夺财产、破坏某个过程或象征物等。基于这一完整概念, 某些表面上未直接对他人或社会造成明显伤害的行为, 实质上仍可能属于恶意创造力的范畴:

(1) 个体为满足个人目的而逃避惩罚或利用法律漏洞的行为, 可视为对规则系统或程序公正的“过程”的破坏;

(2) 在自我防卫情境下产生的伤害性创新方案, 其构思过程本身已符合恶意创造力中“以新颖方式伤害外界对象”的特征。此时, 自我防卫反映的是行为动机, 而非对行为是否属于恶意创造力的否定;

(3) 若负创造力中所指的“幽默”特指对他人造成心理或社会性损害的戏弄或嘲讽, 则其与恶意创造力中“捉弄他人”这一类别在内涵上高度接近。

因此, 建议作者在研究中进一步明确负创造力与恶意创造力在概念边界上的区别, 这是作者所提出的干预模型中最基础的部分。

回复: 感谢您提出的这一关键问题。确实, 厘清负创造力与恶意创造力的概念边界是本研究的逻辑起点, 也是后续干预模型建构的基础。根据您的建议, 我们在修改稿中对二者的概念内涵与关系进行了系统梳理与明确界定, 具体修改如下:

第一, 明确了二者的概念内涵与分野标准。我们在 2.1 节首先界定了创造力的核心构成——新颖性与适用性(有效性), 强调创造力作为一种中性认知能力, 本身不涉及道德判断(Mayer, 1999; Runco & Jaeger, 2012)。在此基础上, 2.2 节将负创造力与恶意创造力定位为创造力在道德价值维度上的两种变体: 二者共享“新颖性+适用性”的认知基础, 其根本分野在于伤害在行为动机中的位置(Kapoor & Khan, 2016; Cropley et al., 2008)。负创造力以利己为目标, 伤害仅为行为的副产品; 恶意创造力以伤害为核心目标, 利己退居为伴随动机(Dou et al., 2022; Kapoor, 2025)。

第二, 逐一回应了您提出的三个边界案例, 明确其在本文概念框架下的归属(详见 2.2 节):

关于逃避惩罚或利用法律漏洞的行为: 若个体意图仅为自保或获利, 未主动希望破坏规则系统, 则属负创造力; 若意图包含藐视规则、挑战权威, 则属恶意创造力。区分关键在于意图本质, 而非行为表面形式。

关于自我防卫情境下产生的伤害性创新方案: 自我防卫的核心动机是保护自身安全, 伤害是手段而非目的, 因此即使方案具有伤害性, 仍应归入负创造力范畴——这正体现了“伤害为副产品”的核心特征。

关于幽默表达中对他人造成损害的戏弄或嘲讽: 若意图逗乐但判断失误造成伤害, 属负创造力; 若意图贬低、伤害对方, 则属恶意创造力。Kapoor (2025) 以“新颖的八卦”与“原创性谋杀”的对比生动阐释了这一区分。

第三, 引入了神经科学证据作为区分的生物学基础。Perchtold-Stefan 等人(2021)的研究表明, 恶意创造力的产生伴随前额叶与杏仁核的协同激活, 形成“认知-情感”双系统耦合模式; 而 Kapoor 和 Khan (2018) 的 EEG 研究则发现负创造力可能与左额叶 alpha 去同步化(与欺骗相关)存在关联。这两项研究分别揭示了恶意与负向创造力可能涉及的神经基础, 为二者的动机差异提供了初步的生物学线索(详见 2.2 节)。

第四, 在表 1 中直观呈现了三者的区别, 将“适用性”与“道德价值”分离, 清晰展示创造

力、负创造力与恶意创造力在新颖性、适用性、道德价值与动机本质四个维度上的异同（详见 2.2 节表 1）。

以上修改已在稿件中第 2 部分以红色字体标注。再次感谢您的宝贵意见，您的建议使本文的概念界定更加清晰、严谨，为后续演化路径与干预模型的建构奠定了坚实基础。

意见 2：作者在引入“适用性”概念时，似乎对其内涵的解释有一定的偏差。在常规的创造力研究中，适用性通常指创造性观念在现实中成功实施的程度。然而，在当前研究中，适用性更像是衡量创造性观念在道德上的可接受性，而非一般意义上的适用性。建议作者对“适用性”进行更为详细的定义和解释，并相应地调整全文的相关论述，以确保概念的一致性和准确性。

回复：感谢您的精准指正。我们完全认同：在创造力研究的学术传统中，“适用性”特指想法或产品在特定情境下的有效性、可行性、可实施性，即能否成功实现预期目标（Mayer, 1999; Runco & Jaeger, 2012）。它关乎“手段-目的”的有效性，而非道德层面的善恶判断。将“适用性”与“道德可接受性”混用，确实会造成概念混淆，进而影响理论建构的严谨性。

基于您的意见，我们对全文进行了系统性修正，具体包括：

第一，重新界定“适用性”的内涵。在修改后的 2.1 节中，我们明确将适用性定义为“想法在特定情境下的有效性、可行性，即能否成功实施并达成预期目标”。同时强调，“此处的适用性聚焦于‘手段-目的’的有效性，本身不涉及道德层面的善恶判断——创造力作为一种认知能力，既可服务于建设性目标，也可能被用于自利或伤害性目的。”（详见 2.1 节，红色字体标注）

第二，引入“道德价值”作为独立维度。既然创造力本身是中性能力，其消极面的区分就不能在“适用性”内部寻找，而需要在另一个维度上实现。因此，我们引入道德价值作为独立的第三维度，用以评估创造性行为在意图与结果上的伦理属性。这一调整为后续区分负创造力与恶意创造力奠定了清晰的概念基础（详见 2.2 节）。

第三，基于新框架重新界定负创造力与恶意创造力（详见 2.2 节）：

- 创造力：新颖性 + 适用性 + 道德价值正向（建设性）
- 负创造力：新颖性 + 适用性 + 道德价值负向（结果）+ 意图非伤害
- 恶意创造力：新颖性 + 适用性 + 道德价值负向（意图+结果）+ 意图伤害

第四，相应修正表 1。（详见 2.2 节表 1）。

第五，全文一致性调整。我们对全文进行了系统性检索与修正，确保“适用性”一词在所有语境中均指代“有效性/可行性”，涉及道德判断时统一使用“道德价值”“社会价值”“意图”等术语。具体修改已用红色字体标注于全文各处。

意见 3：负创造力到恶意创造力的动态演变

作者需澄清一个潜在的理论矛盾：当前主流观点认为，幼儿及未成年个体的道德功能发展尚不成熟，可能与其前额叶等负责执行控制功能的脑区发育未完成有关。同时，创造性观念生成作为一种复杂的认知过程，高度依赖于执行控制功能的参与。依此逻辑推论，处于该发展阶段的个体应较难表现出结构相对复杂的恶意创造性行为；其可能表现的伤害行为，更倾向于执行控制功能不足时所引发的、缺乏新颖性特征的冲动性及一般性伤害行为。

回复：感谢您提出这一深刻且富有挑战性的问题。我们完全认同审稿人的基本逻辑：执行功能确实是创造性思维的重要认知基础，幼儿阶段执行功能尚未成熟，因此复杂、结构化的恶意创造力在早期确实罕见，而常见的伤害行为多为缺乏新颖性的冲动性攻击。然而，这一逻辑并不否定负创造力在少数个体中存在的可能性，更不否定早期识别与干预的必要性。我们在修改稿中从以下四个方面对这一理论矛盾进行了澄清：

第一，幼儿具备生成“新颖+有效”想法的认知基础，但呈个体差异分布。创造力研究早已证实，学前及小学阶段儿童已具备可被稳定测量的创造性思维能力，其创意产出能够同时体现新颖性与情境适切性(Runco & Acar, 2012; 叶平枝, 马倩茹, 2012)。正如林崇德(1992)所言，人人乃至儿童都有创造力——尽管受限于知识经验与认知成熟度，儿童的创造力通常表现为对个体自身有意义的 mini-C 或 little-C 水平，而非具有社会价值的产品(Kaufman & Beghetto, 2009)。正是这一基础性能力的存在，使得儿童在特定条件下也可能表现出创造力的负向变体。只是高创造力与负创造力并非该年龄段的普遍现象，而是少数个体在特定条件下的特殊表现(Barbot et al., 2020) (详见引言第二段，红色字体标注)。

第二，负创造力与恶意创造力具有不同的认知复杂性层级。儿童早期的负创造力主要表现为认知监控不足下的“机会主义”行为(Barbot et al., 2020)，其认知加工以低监控、偶然性为主。而恶意创造力则需要更高阶的认知参与：在愤怒等情绪唤起状态下，个体需要同时激活攻击动机(杏仁核)和目标计划(前额叶)(Perchtold-Stefan et al., 2021)，并能够在“报复性思维”与“建设性重评”之间做出选择(Cheng et al., 2021)。因此，从负创造力到恶意创造力的演变，不仅是动机的“恶质化”，也是认知加工从“低监控-偶然性”向“高唤醒-工具性”的升级过程。这也解释了为何结构复杂的恶意创造力更多出现在青少年后期及成年期，以及我们要前置干预负创造力的初衷。

第三，“不具备新颖性的伤害行为”不应被归入负创造力范畴。负创造力的定义明确包含“新颖性”要件——只有同时满足“新颖+有效(对自身)+伤害(结果)”的行为才构成负创造力(Kapoor & Khan, 2016)。幼儿阶段常见的冲动性攻击、模仿性伤害等行为，虽可能具有伤害后果，但因缺乏新颖性，属于一般性行为问题，而非负创造力表现。这一区分在修改稿 3.1 节中予以明确(详见负创造力定义部分，红色字体标注)。

第四，正因为负创造力是少数个体的特征，早期识别与干预才具有特殊价值。Barbot 等人(2020)对青少年犯的研究发现，负性意念具有高度的个体差异与情境特异性——并非所有人都表现出负创造力，且即使在表现出负创造力的个体中，其表现也随任务类型而变化。这意味着，负创造力在发展中始终是少数个体的特征，而恶意创造力更为罕见。然而，正是这种罕见性，凸显了早期识别与干预的特殊价值——若能在少数表现出负创造力的个体中实现有效引导，或可阻断其向恶意创造力演化的潜在路径，避免未来更严重的行为问题。

综上所述，我们承认审稿人逻辑的合理性；但我们也通过引入“个体差异”“认知层级区分”“新颖性要件”等概念，澄清了负创造力在少数个体中存在的可能性及其与一般性伤害行为的本质区别。这些澄清不仅没有削弱干预的必要性，反而恰恰论证了针对少数高风险群体开展早期识别与引导的独特价值。再次感谢您的深刻质疑，您的意见促使我们对发展视角下的理论假设进行了更严谨的梳理与表达。

意见 4：负创造力到恶意创造力的动态演变

根据作者的论述，负创造力在特定条件下可能转化为伤害意图更直接的恶意创造力。然而，该转化过程依赖于两个尚未充分阐明的前提：（1）负创造力与恶意创造力之间并非并存关系，而是一种此消彼长的替代关系；（2）恶意创造力应被理解为负创造力在特定情境下发展出的更“高级”形式。

就现有论述而言，作者所提及的条件可能同时促进负创造力与恶意创造力，但尚未提供充分的理论或实证依据来说明二者之间存在前者向后者演变的关系。

回复：感谢您提出这一深刻的理论性质疑。您提出的这两个前提非常关键。我们在修改稿中对二者关系的理论定位进行了重新梳理与明确表述，核心观点如下：

第一，负创造力与恶意创造力的关系并非简单的“并存”或“替代”，而是“在不同认知-动机水平上的层级关系”。这一观点在修改稿第 3 节末尾和第 4 节开篇均有明确表述（详见第

3 部分的倒数第二段，红色字体标注）。具体而言，恶意创造力在负创造力基础上叠加了“伤害意图”这一动机成分。负创造力的核心是“利己”，伤害仅为副产品；而恶意创造力的核心是“伤害”，利己退居为伴随动机。Dou 等人（2022）进一步提出，“增值效应”是推动负创造力向恶意创造力转化的关键刺激——当个体意识到伤害行为能带来额外收益（如竞争优势、心理满足）时，利己动机可能扩张为伤害动机。

第二，您所提及的“条件可能同时促进负创造力与恶意创造力”这一观察非常敏锐，但这恰恰印证了演变的“多路径协同”特征，而非否定演变关系本身。我们在修改稿中明确指出，负创造力向恶意创造力的演变由认知控制失效、情境诱发、动机强化、道德推脱四类内部因素与社会监管不足这一外部因素协同驱动（详见第 3 部分）。这些因素既可独立发挥作用，也常常相互叠加——当多个因素同时出现时，演变的概率显著升高；而当所有因素协同作用时，从负向恶的转化便可能加速完成。因此，所谓“同时促进”并非问题，而是演变的累积效应机制的体现：同一个因素（如愤怒情境）既可能直接诱发恶意创意，也可能通过强化负创造力模式间接推动演变。

第三，Dou 等人（2022）提出的双向联结模型进一步支持了二者之间存在动态转化关系的可能性。该模型明确指出，负创造性思维可能转化为恶意创造性思维，也可能始终保持负性状态；而恶意创造性思维在生成后，也可能因风险规避、道德约束或后果评估等因素，在实际执行中“降级”为负创造力行为或终止实施（Dou et al., 2022）。同时，我们也体现在演变机制及相应的图上（详见第 3 部分内容，红色字体标注）。

综上所述，我们并不主张负创造力必然向恶意创造力演变，而是强调在特定条件下存在这种转化的可能性，且这一转化具有“认知前体—动机质变”的层级特征。再次感谢您的深刻质疑，您的意见促使我们对演变模型的理论基础进行了更严谨的梳理与表达。

意见 5：负创造力的干预前提：测量与识别

作者在论述三个观察指标时，其引用的实证研究主要集中于论证儿童会出现这些行为本身。然而，对于这些行为与负创造力之间的理论或机制关联——即它们为何及如何作为负创造力的有效指标——现有论证尚缺乏充分的文献支持。建议作者补充相关研究或理论，明确阐释各项行为指标与负创造力在认知、动机或行为表现层面的具体联系。

回复：感谢您的精准指正。我们完全认同：仅证明儿童存在某些行为，并不足以将这些行为界定为负创造力的有效指标——必须阐明这些行为与负创造力核心特征（新颖性+适用性+利己动机，伤害为副产品）之间的理论关联。基于您的建议，我们对 4.1 节进行了系统性重构，具体修改如下：

第一，明确了测量困境与建构路径。我们在 4.1 节开篇系统梳理了现有测量工具的局限。在此基础上，我们指出：由于该领域研究尚处起步阶段，目前并无成熟的儿童负创造力测量理论可资借鉴。一个可行的建构路径是借鉴 Hao 等人（2016）开发恶意创造力行为量表（MCBS）的方法论思路——从行为表征出发，通过系统梳理目标群体中的典型行为模式，结合理论特征进行指标归纳，再通过实证研究检验结构效度（详见 4.1 节，红色字体标注）。

第二，将“观察指标”重新定位为“可供未来研究探索的观察维度”。为避免过度承诺理论的成熟度，我们将原稿中的“三项观察指标”修改为“基于文献梳理，以下几个方向可作为儿童负创造力识别的潜在观察窗口”，并明确强调这些维度的探索性与待验证性（详见 4.1 节，红色字体标注）。

第三，为每个维度补充了与负创造力核心特征的理论关联（详见 4.1 节）。

第四，我们在 4.1 节末尾明确指出只有当行为同时满足新颖性与有效性标准，且伤害并非行为者的目标导向时，方可被界定为负创造力表现。这一限定既避免了将一般性行为问题误判为负创造力，也与前文“负创造力是少数个体特征”的论述形成呼应。

意见 6：负创造力的教育干预模型提出

作者当前提出的干预模型在理论建构与方法设计两方面均存在文献支持不足的问题，这在一定程度上削弱了模型的理论价值与对应干预方法的可靠性。具体而言，文中对干预方法的描述多停留于对基础原则的转述，未能充分整合现有实证证据与理论框架，导致其可操作性与指导性有限。建议作者系统梳理国内外相关领域的经典及前沿研究成果，基于此对模型的构成要素、作用机制及实施流程进行细化与深化，从而提升模型的解释力与实践意义。

回复：感谢您的深刻洞察。您指出的问题确实存在——原稿中的干预方法描述较为原则化，缺乏系统的理论整合与实证支撑。根据您的建议，我们对 CMD 模型进行了全面重构，具体修改如下：

第一，引入 Yeh（2004）创造力生态系统理论作为模型的宏观框架。这一框架将影响创造力发展的因素划分为四个子系统：个体系统、经验系统、环境系统、文化系统。CMD 模型的认知层对应个体系统，动机层对应经验系统，发展层整合环境系统与文化系统，形成与演变路径精准呼应的三层干预结构（详见 4.2 节，红色字体标注）。

第二，整合社会信息加工理论、创造力成分模型中任务动机理论、以及道德推脱理论三大理论资源，为各层干预提供坚实的理论基础（详见 4.2 节）。

第三，将干预靶点精准定位于演变路径中的关键节点。基于前文梳理的多路径演变机制（认知控制失效、情境诱发、动机强化、道德推脱、社会监管不足），我们为每一层明确了对应的干预靶点（详见 4.2 节及相关的图）。如，认知层的干预目标是针对认知控制失效与情境信息误读，修复信息加工的完整性与准确。

第四，细化各层干预策略，提升可操作性与发展适宜性。我们为每一层设计了具体、可操作的干预方法，并特别考虑了早期儿童的发展特点（详见 4.2 节各小节）。如认知层中培养情境信息的完整编码能力，采用“情境重演”游戏，用玩偶回放冲突场景。

第五，明确了模型的整合机制与实践路径（详见 4.3 节）。CMD 模型的三层干预并非线性执行，而是动态协同：认知层的修复为动机层的价值重构提供认知基础，动机层的转化需要发展层的环境支持予以巩固，而发展层的生态构建又反过来强化个体的认知与动机调节能力。这种“个体—经验—环境—文化”四系统联动，正是 Yeh（2004）所强调的创造力发展的核心机制。

再次感谢您的宝贵意见，您的建议使 CMD 模型从原则性框架发展为具有扎实理论根基、明确干预靶点、可操作实践路径的系统性干预方案。

意见 7：作者在引言部分引入关键研究变量之前的内容应更为简洁。建议作者在简要讨论创造力被用于负面用途的危害后，直接阐明选择未成年群体的恶意创造力作为研究对象的理由。

回复：感谢您的精准建议。我们已对引言部分进行了系统性精简与重构，具体可见修改稿中引言部分红色字体标注的内容。

意见 8：作者在提出自身观点时，经常会使用了大量推论，但缺乏实证研究的支持（或者仅以相同且数量较少的一两篇研究作为支持）。建议作者重新审视全文，并在给出结论时增加相关实证研究成果的引用，以增强论证的说服力。需要增加引用条目的部分包括但不限于：

（1）引言，“此外，儿童在童年早期正处于思维活跃、想象力丰富的阶段……从而及时采取干预措施，避免儿童的负创造力进一步演变为恶意创造力。”

（2）负创造力到恶意创造力的动态演变，“恶意创造力常是在负创造力的基础上逐步演化而来（Dou et al., 2022）”

（3）负创造力到恶意创造力的动态演变，“尽管负创造力的出现带有偶然性，但其向恶意创造力的演变并非随机。若偶然生成的负性创意或行为带来直接收益（如快速解决问题、

获得同伴关注、宣泄负面情绪）与“增值性”收益（如伤害他人的同时提升自身竞争优势），该行为便会受到内在强化（Dou et al., 2022）。”

回复：感谢您的宝贵意见。您指出的“推论过多、实证支持不足”确实是原稿中需要加强的关键问题。我们已对全文进行了系统性审视，在核心论点处补充了相关实证研究成果，以增强论证的说服力。具体修改如下：

第一，在引言部分关于儿童早期干预必要性的论述中，我们补充了儿童创造力的相关实证研究（Zhou et al., 2024），论证儿童具备创造性思维基础且存在个体差异，同时结合道德发展理论（Kohlberg, 1984；Turiel, 2006）说明该阶段道德判断尚未成熟，从而为“儿童早期是负创造力萌芽期与干预窗口期”的核心论断提供实证与理论双重支撑。

第二，在负创造力向恶意创造力演变的核心论述中，我们补充了多项实证研究：Perchtold-Stefan 等人（2021）关于恶意创造力与认知重评能力受损的研究、Cheng 等人（2021）关于愤怒情绪通过内隐攻击路径促进恶意创造力的研究，为“恶意创造力常在负创造力基础上演化”提供了过程性证据；Harris 和 Reiter-Palmon（2015）关于不公正情境诱发恶意创意的研究、Qin 等人（2019）关于创意心态通过道德推脱引发伤害行为的经验取样研究，则为“内在强化与增值效应”机制提供了实证支持。

意见 9：负创造力与恶意创造力：新颖性与适用性的失衡，“事实上，负创造力与恶意创造力正是在个体品格特征或动机视角下对创造力的进一步界定。”

对于该部分的论述结构，建议作者首先明确说明负创造力和恶意创造力的概念内涵，然后基于各自的内涵进行差异化分析和讨论。

回复：感谢您的宝贵建议。您指出的论述结构问题确实非常重要——先明确概念内涵，再进行差异化分析，这是确保概念清晰、逻辑严谨的基本要求。根据您的建议，我们对 2.2 节的论述结构进行了重新组织，具体修改如下：

第一，调整论述顺序，先分别明确负创造力与恶意创造力的概念内涵。修改后的 2.2 节首先独立阐述负创造力的定义，明确其“利己为目标、伤害为副产品”的核心特征（Kapoor & Khan, 2016；Barbot et al., 2020）；随后独立阐述恶意创造力的定义，明确其“伤害为核心目标”的本质属性（Crompton et al., 2008；Hao et al., 2016）。两个概念的内涵界定清晰、独立，避免了混同论述。

第二，在概念界定清晰的基础上，系统进行差异化分析与讨论。我们随后从动机本质（利己 vs. 伤害）、伤害在行为中的位置（副产品 vs. 核心目标）、神经基础（前额叶主导 vs. 前额叶-杏仁核协同）等维度对二者进行了系统比较，并引入 Kapoor（2025）的“意图效价连续谱”观点，说明二者虽有明确分野但也存在动态关联的可能。

第三，整合形成概念关系框架。在上述分析基础上，我们提出负创造力与恶意创造力共享“新颖性+适用性”的认知基础，其分野在于动机本质，二者构成“认知前体—动机质变”的层级关系。这一框架为后续演变路径的论述奠定了清晰的概念基础。

.....

审稿人 3 意见：

稿件介绍了负创造力向恶意创造力的演化路径与干预，综述的视角具有创新性，但需要修改的地方较多。

意见 1：引言部分，段与段之间衔接存在问题；迟迟不能进入主题；引言部分最后一段提到主题，也没有讲明确。

回复：感谢您的宝贵意见。您指出的引言部分逻辑松散、主题切入不清的问题，确实影响了

文章的阅读流畅性与问题聚焦。根据您的建议，我们对引言部分进行了系统性重构，具体可见修改稿中引言部分红色字体标注的内容。

意见 2：对于 2.1 与 2.2 部分，创造力、负创造或恶意创造都包括新颖与适用两个维度，而适用性上均包括可行性与价值两个方面。创造力与负创造或恶意创造都是具有新颖性和可行性的，区别仅在于价值，那么把创造力、负创造力或恶意创造力表述为新颖性与适用性的平衡、新颖性与适用性的失衡，是否合适？另外“负创造力与恶意创造力虽然都具有较高的新颖性，但在适用性所包含的价值维度与行为意图方面存在本质差异”，这些提法需要进一步的文献支持。

回复：感谢您对 2.1 与 2.2 部分提出的深刻意见。您提出的问题非常精准，确实揭示了原稿在概念表述上的不严谨之处。根据您的建议，我们对 2.1 节和 2.2 节进行了系统性修正：

第一，重新界定“适用性”的内涵。在修改后的 2.1 节中，我们明确将适用性定义为“想法在特定情境下的有效性、可行性，即能否成功实施并达成预期目标”（Mayer, 1999; Runco & Jaeger, 2012; Diedrich et al., 2015）。同时强调，“此处的适用性聚焦于‘手段-目的’的有效性，本身不涉及道德层面的善恶判断”。这一修正使“适用性”回归创造力研究的学术传统，也为后续区分创造力、负创造力与恶意创造力奠定了清晰的概念基础。

第二，引入“道德价值”和“动机意图”维度。既然创造力本身是中性能力，其消极面的区分就不能在“适用性”内部寻找。因此，我们在 2.2 节引入道德价值与动机本质，用以评估创造性行为在意图与结果上的伦理属性。基于这一框架：

- 创造力：新颖性 + 适用性 + 道德价值正向（建设性）
- 负创造力：新颖性 + 适用性 + 道德价值负向（结果）+ 意图非伤害
- 恶意创造力：新颖性 + 适用性 + 道德价值负向（意图+结果）+ 意图伤害

第三，修正“平衡/失衡”的表述。基于上述三维框架，原稿中“新颖性与适用性的平衡/失衡”表述确有不妥。修改后的 2.2 节将负创造力与恶意创造力定位为“创造力在道德价值维度上的特殊变体”，强调二者与普通创造力共享认知基础，差异体现在道德价值与行为动机层面，而非“适用性缺失”或“失衡”。具体表述为：“负创造力与恶意创造力共享‘新颖性+适用性’的认知基础，其根本分野在于伤害在行为动机中的位置——负创造力以利己为目标，伤害仅为行为的副产品；恶意创造力以伤害为核心目标，利己退居为伴随动机。”

第四，为“价值维度与行为意图”的提法补充文献支持。负创造力与恶意创造力在动机本质上的区分已有多项研究支撑：

- Kapoor 和 Khan（2016）明确指出，负创造力中的伤害仅为“副产品”，行为者并无故意伤害意图；而 Cropley 等人（2008）对恶意创造力的定义则强调“蓄意伤害”这一核心特征。
- Dou 等人（2022）进一步提出，“增值效应”（value-added）是推动负创造力向恶意创造力转化的关键，当个体意识到伤害行为能带来额外收益时，利己动机可能扩张为伤害动机。
- Kapoor（2025）在综述中指出，“意图的效价可以是一个从高尚、中性、模糊、自利，到罪恶、邪恶的价值区间”，这一框架为理解二者的分类关系提供了理论基础。

以上修改已在稿件以红色字体标注。再次感谢您的宝贵意见，您的建议使本文的概念框架更加严谨、表述更加精确。

意见 3：负创造力到恶意创造力的动态演变部分应该是稿件重点，建议重点丰富该部分的介绍，并增加从动机、情境、认知监控三个方面建构负创造力到恶意创造力演变路径的理论依据。而且建议文中演变路径的表述与摘要部分保持一致。

回复：感谢您的宝贵建议。我们完全认同“演变路径”部分是全文的理论核心，其论述深度直接影响后续干预模型的建构基础。根据您的意见，我们增加了从动机、情境、认知监控三个方面系统建构演变路径的理论依据；明确四类内部因素与外部因素的协同驱动机制（认知控制失效、情境诱发、动机强化、道德推脱四类内部因素与社会监管不足这一外部因素协同驱动的动态过程。这些因素既可独立发挥作用，也常常相互叠加）；并确保演变路径的表述与摘要完全一致。

意见 4：负创造力的提前教育干预部分，测量与识别应该略写，并且建议明确写出“提出早期儿童负创造力行为三项观察指标”的理论依据。建议干预模型与前面的动态演化路径相呼应，并且干预模型中“认知层的核心目标是帮助儿童理解创造力中“新颖性—适用性”的双重标准，避免片面追求新颖性而偏离创造力的本质。”“动机层强调以社会价值为核心，对儿童的自我中心动机进行价值重构”，这些核心目标需增加文献支持。

回复：感谢您对干预模型部分提出的建设性意见。根据您的意见，我们对第 4 部分进行了如下系统性修改：

第一，精简“测量与识别”部分的篇幅，将其定位为干预模型的前提性铺垫而非独立论述重点，使论述重心转向后续的干预模型。修改后的 4.1 节压缩了原有对测量工具的分类综述，聚焦于两个核心问题：一是现有测量工具在儿童群体中的适用性局限（Kapoor & Khan, 2016; Barbot et al., 2020）；二是借鉴 Hao 等人（2016）开发 MCBS 的建构逻辑，说明从行为表征出发、结合理论特征进行指标归纳是可行的探索路径。

第二，明确“三项观察指标”的理论依据，将其从“指标”重新定位为“可供未来研究探索的观察维度”。修改后的 4.1 节为每个维度补充了与负创造力核心特征（新颖性+适用性+利己动机）的理论关联：

- **规则突破维度：**以 Bridgers 等人（2025）发现的“漏洞行为”为例，阐明其兼具新颖性与有效性，且动机为自利而非伤害，符合负创造力的核心定义；其认知基础是儿童对规则的理解能力与道德评估的缺位，正是 Barbot 等人（2020）所指出的“认知监控不足”的表现。
- **危险探索维度：**Hansen Sandseter（2007）对儿童危险游戏的分类表明，这类活动往往需要突破常规思维，体现新颖性特征；Boles 等人（2005）的研究则揭示，儿童对危险的认知不足可能导致其产生对自身有效但后果不确定的行为，伤害仅为认知不足下的副产品——这正是负创造力的典型特征。
- **逃避责任维度：**Harvey 等人（2018）与 Talwar 和 Lavoie（2022）的研究表明，利己性谎言以利己为目标，可能对他人造成伤害，但其动机并非蓄意中伤，而是自我服务，与负创造力“伤害为副产品”的核心特征高度吻合。

第三，使干预模型与前面的动态演化路径形成清晰呼应。修改后的 4.2 节开篇即明确：CMD 模型的认知层针对演变路径中的“认知控制失效”与“情境信息误读”，动机层针对“动机强化（增值效应）”与“道德推脱启动”，发展层针对“社会监管不足”与“道德约束脆弱”。每一层的干预目标均与前文梳理的演变驱动因素一一对应，形成精准的反向阻断机制。

第四，为认知层与动机层的核心目标补充文献支持。

- **认知层核心目标：**原表述“帮助儿童理解创造力中‘新颖性—适用性’的双重标准”已调整为“促进社会信息加工的完整性与准确性”。修改后的认知层以社会信息加工理论（Crick & Dodge, 1994）为核心依据，强调干预目标是修复儿童在情境编码、解释、反应生成等环节的加工偏差，而非简单灌输“新颖-适用”概念。同时补充 Cheng 等人（2021）关于情绪调节与 Perchtold-Stefan 等人（2021）关于认知重评的研究，为具体干预策略提供实证支撑。

- 动机层核心目标：原表述“以社会价值为核心，对儿童的自我中心动机进行价值重构”现融合两大理论资源——任务动机理论（Amabile, 1983）解释动机如何从外在利己向内在亲社会转化，道德推脱理论（Bandura, 1999）揭示道德约束如何被解除。同时补充 Dou 等人（2022）关于“增值效应”的研究、Qin 等人（2020）关于道德推脱在创造力领域的研究，以及裘指挥和张丽（2006）关于幼儿利己-利他博弈的论述，使动机层的理论基础更加扎实。

第二轮

审稿人 3 意见：作者已根据修改意见进行了很好的修改，建议发表。

编委 1 复审意见：稿件修改情况良好，文章质量已得到显著改善，本人对该稿件持肯定意见，建议编辑部接受发表。

编委 2 复审意见：请审稿人 2 再看一遍。

审稿人 2 意见：论文做了大幅修改，包括结构、逻辑和具体内容，较好回答了评审专家意见，论文质量已有很大提升。建议文末加一段，对本论文提出的“动态转化路径模型”和“CMD 三层干预模型”的理论创新性和实践意义，进行凝练总结。

回复：非常感谢审稿专家对本文修改工作的肯定，也感谢专家提出的进一步完善建议。我们完全认同该意见，本文在前期修改中围绕负创造力向恶意创造力的动态演化路径，以及儿童负创造力的前移性干预框架进行了较大幅度重构，因此确有必要在文末对本文提出的两个核心模型进行更加凝练的总结，以突出其理论创新性与实践意义。根据您的建议，我们已在正文第 5 部分“总结与展望”中新增并凝练了总结性论述。此外，根据本轮修改，我们已对全文进行语言润色、术语统一，并对文中符号、图表等表述进行校对与完善，同时优化了摘要的凝练性，以增强全文的整体逻辑性与学术规范性。相应修改已在正文中以黄色背景标注。再次感谢审稿专家的建设性意见。

第三轮

编委 2 复审意见：同意评审意见。接受发表。

主编意见：

建议在理论概念上做一下调整：将创造力在总体上分为“创造”（正面的，亲社会的）与“负性创造”（负面的，不利于社会的）两大类。而在“负性创造”之下，又可存在两个亚类，即“以利己为主要目的的负性创造（类型 1）”和“以害人为主要目的的负性创造（类型 2）”（实际上，可能还存在某些混合类型，比如，兼具利己和害人双重特性的负性创造）。这样，论文的主题就是如何防止从“类型 1”向“类型 2”转化。修后发表。

回复：非常感谢编委专家提出的建设性意见。我们完全认同专家关于理论概念层级需要进一步清晰化的建议。根据该意见，我们对文中相关概念进行了调整：首先，从结果效价出发，将创造力区分为正性创造与广义负性创造；其次，在广义负性创造之下，根据动机属性进一步区分以利己为主要目的的类型 1 负性创造和以伤害为主要目的的类型 2 负性创造，并说明

现实中还可能存在兼具利己与伤害动机的混合型负性创造。相应地，我们已在摘要、概念界定部分、表格及其他部分对相关表述进行了修改（具体修改的内容以蓝色字体标记），以更清楚地呈现创造力、积极创造、消极创造、负创造力与恶意创造力之间的层级关系。修改后，文章的核心逻辑更加明确，即：创造力首先依据结果效价区分为积极创造力与消极创造力，而本文重点关注消极创造力内部由利己型向伤害型转化的机制及其前移性干预。再次感谢编委专家的宝贵建议。