

《心理学报》审稿意见与作者回应

题目：连贯关系调节语义冲突加工的脑网络机制——来自因果与让步关系的 fMRI 证据
作者：贾程，徐晓东

第一轮

感谢《学报》编辑老师与两位审稿专家的细致评审与建设性意见。我们已对全部意见逐条认真研读，并据此对稿件进行了系统性修订（见修改稿正文蓝色标注）。在逐条回复之前，谨将本次修订的核心要点扼要说明如下，以便快速了解修改全貌。

核心修订要点

（一）术语规范与概念澄清。将核心术语统一规范为“连贯关系”，覆盖中英文标题、摘要及正文共 30 处，使全文表述与心理语言学传统一致；并在引言中新增概念说明段落，对关键理论术语作精准界定。

（二）行为结果补充报告。新增 3.1 行为结果一节，报告可理解性评分的 2×2 重复测量方差分析，完善实验设计的行为层面证据。

（三）阴性结果的稳健性论证。针对跨句式泛化的阴性结果，新增三项工作：基于因果句内部定位的 super mask 跨句式泛化分析；新增 2.4.4 贝叶斯统计分析独立子节，以贝叶斯因子 ($BF_{01} = 8.50$) 量化支持零假设的证据强度；在 4.3.2 节增加针对替代解释的排除分析，从特征信噪比与认知负荷两个角度进一步排除假阴性可能。

（四）方法学描述完善。补充实验流程与扫描参数描述，增强论文的方法学自足性，便于结果复现与同行评议。

（五）结论审慎性提升。系统弱化“证明”“挑战经典模型”等强表述，改为“支持”“提示”“限定条件”等中性学术表达；明确区分显著、边缘显著与趋势性结果；进一步弱化中英文摘要中关于跨句式泛化的理论推论强度。

（六）理论贡献凸显。在 5 结论部分明确列出本研究四项独立理论贡献，并扩充“局限性”讨论，综合行为评分、独立选择实验数据与同领域文献先例进行全面论述，凸显与已发表研究 (Xu et al., 2022, Brain and Language) 在研究问题、分析方法、结果变量与理论结论四个层面的实质性区别。

我们衷心感谢编辑老师与两位审稿专家在评审过程中给予的细致指正与建设性建议，使本稿在数据分析、理论表述、术语规范及结论审慎性等方面均获显著改进。以下按审稿人和意见顺序逐条回复，每条回复均说明：(a) 我们对意见的理解与认同；(b) 具体回应；(c) 相应的修改位置与新增/修改内容。

审稿人 1 意见：

本研究采用 fMRI 结合图论网络分析与多变量模式分析，考察因果与让步句式处理语义冲突的神经关联。脑网络分析在多种指标上均发现了两类句式下的差异，多变量模式分析发现了句式内语义冲突的神经表征，但未能实现跨句式泛化。作者据此推论，语义冲突的脑网络加工机制深度依赖于句法 - 语义框架，因果推理维持模块化的高效网络架构，而让步推理

需要打破模块边界实现跨系统整合。

总的来说，这篇文章尝试结合脑网络分析以及多体素识别等高阶脑数据分析方法，来揭示网络级别上句法框架对语义冲突表征的影响，这在方法学上提供角度是比较丰富的，是本文的优点。然而评审人在阅读后，对于实验材料的操控，分析之间的关联，跨句式解码结果的假阴性可能性存有一些疑问。如果这些节问题能够得到妥善解决，评审人认为是适宜发表的。

感谢审稿人的整体肯定与具体关切。下面针对您提出的三项主要意见逐条回复。

意见 1:

实验材料中，让步关系下的无冲突示例句“外婆从海南迁到了哈尔滨，尽管她喜欢那里冬天暖和。”这句话存在一定的指代不明，即“那里”到底是指代前后哪个城市？评审人在读这个句子时感觉到明显的疑惑，无法自动的理解这句话，需要经过反思才能知道所指。在初次阅读中，我很容易将那里的指代对象理解为哈尔滨，因为距离“那里”更近，且在第一句中是宾语目的地，重要性突出。或许这句话更自然的说法是“外婆从海南迁到了哈尔滨，尽管她并不喜欢那里冬天寒冷”，或“外婆从海南迁到了哈尔滨，尽管她喜欢海南冬天暖和。”作者的原始材料很可能存在歧义和解歧的加工，需要后置的“冬天暖和”来为“那里”究竟是哪里解歧。

尽管 4 个条件下都存在指代不明的问题，但是读者理解指代的关系的难度似乎不同。因果关系下的无冲突例句“外婆从哈尔滨迁到了海南，因为她喜欢那里冬天暖和。”，对于评审人来说很容易理清指代关系。这可能是因为“海南”与“那里”距离更近，更符合就近指代的原则；并且第一句话搬家的目的地是海南，它在句法上是宾语，在语义上的重要性也更加凸显，所以容易将其默认为作为第二句所指带的地点。

上述问题所可能导致的结果是，指代不明所诱发的歧义很可能混淆作者想要操纵的自变量以及交互作用。甚至，评审人个人觉得 AC 这个条件读起来并不连贯，是会引起一定的语义冲突的。这有可能意味着作者对语义连贯性/冲突的操控也并不完全符合主观的预期。

评审人建议作者开展一个对指代目标歧义程度的主观评分实验，或者其他有效的评估方法，来确认是否存在这样的问题。如果问题确实存在，作为弥补，可考虑将这个难度作为控制变量，从脑数据分析中以回归因子等方式控制掉。

此外，作者提到曾经对刺激做了句子整体的“可理解性评分”，这个评分虽然并不直接针对评审人所提出的问题，但也具有一定的参考价值，可惜作者并未呈现这部分行为结果，建议完整报告出这部分评分的结果以及在条件间的差异。

回应:

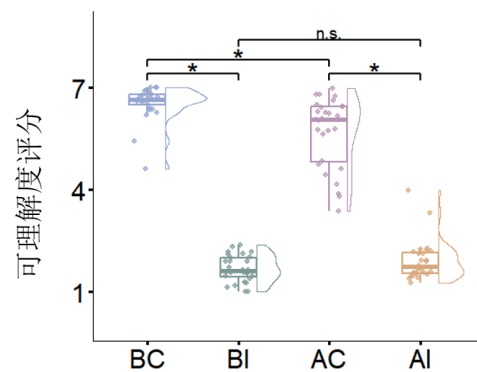
非常感谢审稿人敏锐地指出了实验材料中可能存在的指代歧义问题。我们认真考虑了您的意见，并承认：在孤立阅读示例句时，“那里”的指代确实可能产生短暂歧义，且该歧义在让步条件（AC）中可能比因果条件（BC）更明显。以下从四个方面进行系统回应。

（一）补充报告可理解性评分结果

按照审稿人的建议，我们对 fMRI 扫描中采集的 7 点可理解度评分进行了 2（连接词类型：因果 vs. 让步） $\times$ 2（语义一致性：一致 vs. 冲突）重复测量方差分析（见下图）。

结果显示，连接词类型与语义一致性存在显著交互效应 ($F(1, 27) = 15.62, \eta^2G = 0.15, p < 0.001$)。关键的是，让步一致条件（AC）的均值 $M = 5.70, SD = 1.04$ ，远高于量表中心点（4），且远高于让步冲突条件（ $M = 1.92, SD = 0.59$ ）。如果指代歧义严重到使 AC 句本身产生实质性的语义冲突，其评分应接近 AI 水平，但事实并非如此。此外，因果冲突（BI）与让步冲突（AI）之间的可理解度无显著差异（ $\beta = 0.24, p = 0.260$ ），排除了因指代解歧难度不

同导致语义冲突操纵强度被混淆的可能。



(补充为修改稿中的图 2)

(二) 已有的指代选择实验结果

Xu 等（2015）的工作中已通过独立的选择实验验证了指代的可控性。该实验采用与本研究完全相同的实验材料，由 24 名母语者选择代词“那里”的指代对象。结果显示：在因果句中，96.9% 的被试将“那里”判断为目的地（B 地/第二个地点）；在让步句中，93.1% 的被试将其判断为出发地（A 地/第一个地点）。两组指代偏向无显著差异（ $p > 0.5$ ）。该结果表明，尽管“那里”在句法上更靠近第二个地点，被试仍能依据连接词类型稳定且无歧义地确定指代对象，指代解歧具有较高稳定性。

Table 2. Percentages of referents for the demonstrative pronoun *nali* to refer to the closer place (place B) in the causal construction and the distant place (place A) in the concessive construction.

	Percentages of referents	
	Mean (%)	SD (%)
Causal construction – place B	96.9	4.4
Concessive construction – place A	93.1	7.7

（表格来源：Xu et al., 2015, LCN）

(三) 原实验材料设计的控制逻辑

在材料设计阶段，作者已考虑了潜在的指代问题，并采取了以下措施：（1）两个地点具有明确的特征对比（如海南暖和、哈尔滨寒冷），使得“冬天暖和”这一语义属性只能与海南唯一匹配；（2）采用正性态度动词（喜欢、偏好等），进一步约束指向。诚如审稿人所言，“那里”在句法上更靠近第二个地点，需要读者结合语义进行解歧；但这正是本研究关注的核心机制——连接词作为自上而下的加工框架，系统调节读者对后续代词的判断，因果框架引导将“那里”解读为目的地、让步框架则引导将其解读为出发地。指代解歧过程本身、特别是由连接词类型引导的解歧差异，正是本研究探讨的焦点。

(四) 关于将指代解歧难度作为协变量的考量

审稿人建议将指代解歧难度作为控制变量纳入回归分析。我们仔细评估了该方案的可行性，但综合判断后未予采纳，理由如下：

(1) 代词解歧评测结果已显示指代解歧具有较高稳定性（正确率 > 93%），个体间变异极小，统计上不适合作为有效协变量；

(2) 本研究的核心结论建立于交互效应模式（如 AC vs. AI 的悬殊差异、BI vs. AI 的无显著差异），细微的指代解歧变异不太可能解释这些核心模式；

(3) fMRI 实验时未逐试次采集指代解歧的反应时或难度评分，难以构建可靠的事后协变量；

(4) 同类范式的已发表研究（Xu et al., 2015, 2018, 2022）均未将指代解歧难度作为协变量纳入分析，沿用其分析框架有助于结果的可比性。

基于上述四点，我们未在本研究中实施事后回归控制，但充分认识到这是本研究的方法局限，并已在 论文讨论中的“局限性”部分增设第四条进行专门讨论。

具体添加内容如下：

最后，关于本实验材料的可能指代歧义问题，行为评分显示让步一致条件的可接受度均值（ $M = 5.70$, $SD = 1.04$ ）远高于量表中点和让步冲突条件，且让步冲突与因果冲突的可接受度无显著差异；同时，Xu 等（2015）采用相同材料的独立选择实验显示，因果句中 96.9% 的被试将“那里”判断为目的地、让步句中 93.1% 的被试将其判断为出发地，提示指代解歧具有较高稳定性。同类范式的已发表研究（Xu et al., 2015, 2018, 2022）亦显示不存在指代歧义问题。尽管如此，未来研究仍应采用名词重复等更严格的控制材料，以进一步分离指代加工与话语关系加工。

意见 2:

作者采用了多种分析方法，旨在从各个角度揭示条件间的差异。作者总结说，“四个层面的结果相互印证、逐层深入”。但是这些分析目前看来是平行且独立的，无法证明有存在互相印证的关系，至少在数据分析层面不存在将它们联系起来的分析和结果。建议作者考虑如何从数据分析角度将不同方法下的结果建立联系来支持上述说法。否则，建议宜用“多角度揭示”这类表述替代“相互印证、逐层深入”以及“一条连贯的证据链”这类措辞。

回应:

非常感谢审稿人指出我们在论证表述上的不足。我们完全同意：四类分析在统计上是独立进行的，不应使用“相互印证、逐层深入”等暗示统计或因果递进的措辞。我们全盘采纳您的修改建议。同时进一步说明：四类分析虽统计独立，但在理论逻辑上分别回答了四个不同层次的研究问题，其结论共同汇聚于同一核心发现——即不同连贯关系（因果 vs. 让步）从根本上调节了语义冲突加工的脑网络机制。

修改部分

【1.4 研究问题与假设·末段补充】

上述四个层面的分析从全局拓扑、局部节点、核心枢纽到表征模式，共同构成了对脑网络组织的多角度考察框架。该框架采用逐层递进的策略：首先通过全局网络指标判断整体组织策略是否存在差异；若存在，则进一步通过局部节点和富人俱乐部分析识别关键脑区及核心节点的变化；最后通过 MVPA 独立检验神经表征的泛化性。各层分析在统计上独立，但在理论逻辑上相互衔接。通过这一框架，本研究期望从网络视角揭示连接词类型对语义冲突加工的调节机制，为理解语篇理解中连贯关系与语义加工交互作用的神经基础提供新的证据。

【措辞修订】

【1.4 末段】“多层次考察框架”→“多角度考察框架”

【4.3.2 第二段】“形成了一条连贯的证据链”→“形成了多角度的证据链”

【4.3.2 第二段末】“四个层面的结果相互印证、逐层深入，共同表明...”→“四个层面的结果从多个角度表明...”

意见 3:

作者在跨句式泛化解码中得到了阴性结果，即未能发现显著泛化的脑区。这个无法泛化的发现被作者作为“句法框架依赖性的关键证据（讨论部分 4.3.2 小标题）”，以及“本研究最具理论意义的发现之一”。

然而，这个结果是一个阴性结果。在经典统计推断中，考虑到主要控制假阳性率，而非假阴性率，因此不能基于此类发现做出统计推断，或者说不能作为支持 0 假设的证据。因此，从原则上讲，跨句式解码失败难以充当关键证据和重要发现，这个阴性结果有多大概率是假阴性结果难以得知。

此外，作者在跨句式分析中所使用的方法，如果理解正确，是全脑探照灯分析。这类全脑分析方法为了克服 FWE 问题，会采用多重比较校正，因而对假阳性率的控制非常严格，从而间接提升了假阴性率。

建议做 2 方面调整

（1）在跨句式分析中，建模部分同时进行建模和体素筛选（比如筛选出因果句子中能够有效区分句法冲突是否存在的全部体素），然后以这部分体素作为单一 super ROI（或根据空间关系区每个 cluster 为一个 ROI），以 ROI 为单位做跨句式推广分析。这样无需或仅需很少量的多重比较校正，可以提升检验对阳性结果的敏感性。

（2）如果结果经过上述分析仍为阴性，建议采用贝叶斯检验替代传统统计检验，从而对支持结果为真阴性结果的证据强度给出评估。注意，常用的多重比较校正方法都旨在控制假阳性结果，其代价是会提高假阴性概率，是不利于论证真阴性结果的存在的。

回应:

感谢审稿人提出这两项关键的方法学建议。审稿人的统计推断逻辑完全正确——基于 NHST 的全脑探照灯分析在严格多重比较校正下，假阴性率被间接放大，阴性结果不能直接等同于零假设的证据。我们完全采纳两项建议，并进行了系统性的方法学补充。

（一）基于因果句内部定位的 super mask 跨句式泛化分析

我们首先在因果句式内部（BC vs. BI）进行基于搜索灯（searchlight）的 MVPA，对每位被试生成体素级分类准确率图（准确率减去随机水平 0.5）；然后将个体准确率图纳入组水平单样本 t 检验（检验均值是否显著大于 0），取未校正 $p < 0.001$ 且 $t > 0$ 的体素，合并为一个二值超级掩膜（super mask）。该掩膜代表因果句框架下能稳定区分语义一致与冲突的神经活动空间。对每位被试，提取该 super mask 内所有体素的 beta 值作为特征向量，在因果句（BC vs. BI）上训练线性 SVM 分类器，然后在让步句（AC vs. AI）上测试泛化准确率。此分析无需全脑多重比较校正，显著降低了假阴性风险。

（二）贝叶斯统计推断

根据审稿人的建议，我们对跨句式泛化准确率进行了贝叶斯统计评估。考虑到研究假设是单向的（检验准确率是否高于机会水平 0.5），我们主要报告了单侧贝叶斯 t 检验结果（JZS 先验， $r = 0.707$ ；Morey & Rouder, 2011）。结果显示，支持零假设（ $\mu \leq 0.5$ ）的贝叶斯因子 $BF_{01} = 8.50$ ，达到中等强度证据（参照 Lee & Wagenmakers, 2013， $BF_{01} \in [3, 10]$ 为中等证据）。先验敏感性分析显示，当 Cauchy 先验尺度参数 r 分别取 0.5、0.707、1.0 时，

BF₀₁ 分别为 6.19、8.50、11.83，均 > 3，表明结论对先验设定不敏感（见下表）。综上，贝叶斯分析为“语义冲突的神经表征无法从因果句式泛化到让步句式”提供了量化证据支持。

单侧贝叶斯 t 检验结果（H1 : $\mu>0.5$ ，，JZS Cauchy 先验）

r	BF ₀₁	后验均值（准确率）	95% HDI	零假设证据强度
0.5	6.19	0.491	[0.472, 0.512]	中等证据
0.707	8.50	0.491	[0.471, 0.511]	中等证据
1.0	11.83	0.491	[0.471, 0.511]	强证据

注：BF₀₁ 表示支持零假设 $\mu \leq 0.5$ 与支持备择假设 $\mu > 0.5$ 的贝叶斯因子比值。后验 HDI 为准确率均值的 95% 最高密度区间。

（三）针对替代解释的进一步排除

针对您关于“阴性结果可能仍是假阴性”的核心担忧，我们在第二轮修订中进一步在 4.3.2 节增加一段专门讨论，通过两步分析逐一排除若干替代解释：

(1) 是否因训练特征信噪比不足？super mask 分析已将分类特征限定于在因果句内部能稳定区分语义一致与冲突的体素集合，从源头降低了因训练特征不可靠而导致泛化失败的可能。

(2) 是否因让步条件认知负荷整体偏高，掩盖了泛化信号？我们的让步句式条件内 MVPA 成功区分了让步一致与冲突试次，表明让步句式本身存在可被解码的冲突相关神经活动，从而削弱“让步条件信号被整体噪声淹没”这一替代解释。

综合而言，泛化失败更可能反映两类语义冲突在表征层面采用了不同的编码维度，而非由信号强度或测量噪声所致。

修改部分

【新增 2.4.4 贝叶斯统计分析】

将原本散见于 2.4.3 节末段的贝叶斯方法论描述系统化整合为独立子节，明确报告：JZS Cauchy 先验（默认尺度参数 $r = 0.707$ ）、备择假设设定（H1: $\mu > 0.5$ ）、贝叶斯因子 BF₀₁ 的解释标准（参照 Lee & Wagenmakers, 2013）、以及先验敏感性分析方案（ $r \in \{0.5, 0.707, 1.0\}$ ）。该补充使方法学描述更加体系化，便于读者理解与复制。

【3.5.2 节·跨句式泛化结果】

新增基于 super mask 的跨句式泛化分析、贝叶斯检验、先验敏感性分析的结果及表 3。

【4.3.2 节·开篇语句弱化】

原“本研究最具理论意义的发现之一是跨句式泛化解码的失败”→ 修改为“跨句式泛化解码未能取得显著结果”。

【4.3.2 节·末段新增谨慎讨论】

此外，需要审慎指出，本研究的阴性结果是在特定的实验材料（汉语连接词、可理解性评分任务）、特定的分析框架（全脑搜索灯及后续基于 super mask 的泛化测试）以及有限的样本量下获得的。因此，我们无法完全排除在其他语言、其他范式或更大样本中观察到跨句式泛化的可能性。本研究的主要贡献在于，在当前条件

下未能检测到可靠的泛化证据，并结合贝叶斯检验量化了支持零假设的证据强度；未来研究可采用不同的实验设计和更广泛的被试群体，进一步检验连贯关系对语义冲突表征的调节作用。

【4.3.2 节·替代解释排除段落】

另外，本研究通过两步分析在一定程度上排除了若干替代解释。首先，基于因果句内部定位的 **super mask** 分析将分类特征限定于“能稳定区分语义一致与冲突”的体素集合，降低了因训练特征信噪比不足导致泛化失败的可能；其次，条件内 **MVPA** 成功区分了让步句式中的一致与冲突试次（数据未在表 2 中单独列出），表明让步句式本身存在可解码的冲突相关神经活动，这削弱了“让步条件认知负荷整体偏高、掩盖泛化信号”这一替代解释。综合而言，泛化失败更可能反映两类语义冲突在表征层面采用了不同的编码维度，而非由信号强度或测量噪声所致。

审稿人 2 意见：

本文基于因果关系与让步关系的语义一致/冲突条件，结合 **fMRI**、图论网络分析、富人俱乐部分分析和 **MVPA**，考察不同连贯关系下语义冲突加工的脑网络机制。选题具有一定新意，尝试从脑区激活差异推进到脑网络组织差异，整体研究问题较明确，分析层次也较丰富。稿件使用 2（因果关系和让步关系） $\times$ 2（语义一致和语义冲突）被试内设计，并采用与 Xu 等（2022）相同原始数据进行新的网络层面分析。作者报告了全局效率、模块化、局部节点指标、富人俱乐部连接及 **MVPA** 结果，并将跨句式泛化失败解释为“句法框架依赖性”的证据。总体而言，本文有一定发表潜力，但当前版本在理论界定、方法稳健性和结果解释方面仍存在较明显问题。

衷心感谢审稿专家对本文的整体肯定与具体关切。下面针对您提出的三项主要意见和三项次要意见逐条回复。

主要意见 1：

题目和全文多处将核心操纵概括为“句法框架”，但从材料与理论分析看，作者实际操纵的不仅是句法形式，更涉及连贯关系类型、语义约束和语用推理方式。建议作者进一步澄清“句法框架”与“连贯关系/话语关系框架”的关系，避免理论概念界定过宽或不够准确。

回应：

非常感谢审稿人提出的这一关键性概念澄清意见。我们完全同意：原稿“句法框架”这一表述确实存在界定过宽、不够准确的问题。诚如审稿专家所指出，本研究的核心自变量是连接词类型（“因为” vs. “尽管”），连接词在语言学中确实不能被化约为纯句法形式，它同时激活三个层面的约束：（1）句法层面（引导从句结构类型）；（2）语义层面（建立不同的语篇连贯关系）；（3）语用层面（触发不同的推理策略）。

针对该意见，我们进一步反思了术语选择，并采用更精确、更契合心理语言学传统的术语“连贯关系（coherence relations）”，理由如下：

（1）与全文术语一致：稿件其他部分（摘要、引言、4.3.2 节等）反复使用“语篇连贯关系”这一表述，使用“连贯关系”作为核心理论术语可消除“话语关系”与“连贯关系”并行术语的张力。

(2) 更契合心理语言学传统：Sanders、Spooren、Noordman 等基于因果默认假说（causality-by-default hypothesis）的经典工作，以及 Xu 等（2015, 2018, 2022）研究汉语因果/让步连词的代表性文献，均采用“连贯关系”这一术语。

(3) 概念表达更紧凑：原“框架”一词意在强调连接词所构造的句法骨架；改为“连贯关系”后，由连接词建立的语篇结构内涵已自然蕴含其中，行文更简洁。

修改部分

【题目修改】

原“句法框架调节语义冲突加工的脑网络机制”→修改为“连贯关系调节语义冲突加工的脑网络机制——来自因果与让步关系的 fMRI 证据”

【全文术语规范化】

全文中“句法框架”及其派生表述（“句法框架依赖性”“句法框架特异性”等）经过两轮修订，统一规范为“连贯关系”及其派生表述（共 30 处）。修订范围涵盖中文标题、摘要、正文，以及英文标题与英文摘要。

【1 引言最后部分，新增概念说明】

值得注意的是，连接词作为语言标记并非单纯的句法成分，而是同时调用句法结构（从句类型）、语义关系（比如：因果 vs. 让步）与语用推理策略（证据评估 vs. 预期抑制与重构）三个层面的认知约束。本文沿用心理语言学中的通行做法（Sanders, 2005; Xu et al., 2022），以“连贯关系”统称由连接词所建立的这一多层次语篇结构，并将其作为本研究的核心理论建构。

意见 2:

作者说明本研究与已发表论文使用相同原始 fMRI 数据，但研究问题和分析方法不同。这个说明是必要的，但目前仍不够充分。建议作者更明确地说明：与 Xu 等（2022）相比，本文到底新增了哪些无法由既有结果推出的理论结论，以凸显本研究的独立贡献。

回应:

非常感谢审稿人提出的这一关键要求。我们完全同意，仅笼统说明“研究问题和分析方法不同”不足以凸显本研究的独立贡献。为此，我们做了以下三项补充。

修改部分

【2.1 被试部分·补充信息】

具体而言，Xu 等(2022)主要关注因果与让步关系在局部脑区激活强度及两两有效连接(DCM)上的差异；而本研究则从脑功能网络的整体拓扑组织(全局效率、模块化、富人俱乐部)和分布式神经表征(MVPA)的角度，考察连贯关系如何调节语义冲突的加工。换言之，前者回答“哪些脑区或通路参与”，后者回答“脑网络如何整体组织”。

【4 讨论部分·第一段补充】

在讨论开篇明确列出本文发现的、无法由 Xu 等（2022）结果推出的具体理论结论，包括全局拓扑差异、核心脑区中心性变化、富人俱乐部连接的扰动差异、跨句式泛化失败等。

【5 结论·四项独立理论贡献】

为进一步明确本研究做出的独立理论贡献，我们在 5 结论部分新增“四项独立的理论贡献”清单：

具体而言，本研究做出了四项独立的理论贡献：（1）从脑功能网络全局拓扑层

面证实，因果与让步关系在全局效率和模块化上呈现系统性分离，反映了“模块化策略”与“跨模块整合策略”的对立；（2）富人俱乐部分析揭示让步冲突对核心枢纽连接（hub-hub）和馈线连接（feeder）的扰动显著大于因果冲突，刻画了核心网络架构的动态重构；（3）条件内 MVPA 表明因果冲突与让步冲突在左侧额-颞区域具有可区分的分布式神经表征；（4）跨句式泛化解码未取得显著结果，结合贝叶斯检验（ $BF_{01} = 8.50$ ，中等强度证据）提示语义冲突的神经表征可能具有连贯关系依赖性。

这四项贡献分别对应全局拓扑、核心网络架构、分布式表征、表征泛化性四个层面，每一项都不可直接由 Xu 等（2022）的局部激活与 DCM 分析推导出来。

意见 3:

对“跨句式泛化失败”的解释明显偏强。当前结果最多说明，在本研究样本、特征空间和分析框架下，未观察到可靠的跨句式泛化证据；但尚不足以直接推出“不存在独立于句法框架的抽象冲突信号”，更不足以据此强力挑战经典冲突监测模型。建议作者降低结论强度，并补充更谨慎的讨论。

回应:

感谢审稿人宝贵的意见。诚如您所指出的，阴性结果不能作为“不存在某效应”的直接证据。我们完全认同，并在修订中对相关论述进行了系统性的弱化处理（另请参见，对审稿人 1 意见 3 的回复）。

【数据分析】补充贝叶斯因子（ $BF_{01} = 8.50$ ）作为量化证据

【讨论部分】“对经典冲突监测模型提出了挑战” → “为经典冲突监测模型提供了重要的限定条件”

【讨论部分】“挑战将语义加工视为模块化独立过程的传统观点” → “与将语义加工视为模块化独立过程的传统观点不同”

【4.3.2 节开篇】“本研究最具理论意义的发现之一是跨句式泛化解码的失败” → “跨句式泛化解码未能取得显著结果”

【中文摘要】“上述发现表明...深度依赖于话语关系框架” → “上述发现提示...在很大程度上依赖于连贯关系”

【英文摘要】“syntactically dependent” → “discourse-relation-dependent”; “rather than abstract, domain-general signals” → “rather than purely abstract, domain-general signals”

【4.3.2 节·末段补充谨慎讨论】

见审稿人 1 意见 3 部分所引段落。

【4.3.2 节·替代解释排除段落】

针对您关于阴性结果稳健性的深层关切，新增一段专门排除若干替代解释（详见审稿人 1 意见 3 的回复部分）。

次要意见 1:

方法部分不宜过多依赖“详见 Xu 等（2022）”，建议补充必要的实验流程信息，以增强论文自足性。

回应:

非常感谢审稿人的这一重要建议。我们完全同意，论文应尽量自足，减少读者需要查阅

其他文献才能理解关键实验细节的依赖。为此，我们对方法部分进行了系统性补充，将原本仅引用 Xu 等（2022）的内容替换为对实验流程、刺激呈现、任务设置等的直接、完整描述。

修改部分

【2.2 实验材料和设计·末段补充】

本实验采用快速序列视觉呈现 (RSVP) 方式逐词呈现句子，每个词呈现 300 ms，随后呈现 300 ms 空屏。句子以白色字体呈现在黑色背景上，每个词的水平视角为 1°–3°，垂直视角为 1°。被试需要在安静状态下默读每个句子，理解其含义，并在每个句子结束后使用 7 点量表(1=最不可接受，7=最可接受)对句子的可理解性进行评分。不同试次之间的固定点“+”以“抖动”方式呈现(随机持续 2–8s)，以优化检测效率。每个列表包含 168 个句子(32 个每种实验条件 × 4 种条件 + 40 个填充句)，分为 3 个扫描阶段(run)，每个阶段持续约 15 分钟。正式扫描前，被试完成 28 个练习试次以熟悉实验流程。

次要意见 2:

讨论中应更清楚地区分显著结果、边缘显著结果和趋势性结果。

回应:

感谢审稿人提出的这一规范性建议。我们完全同意。为此，我们对全文进行了系统梳理，明确区分了显著、边缘显著与趋势性结果。

修改部分

【4.1 节·新增独立段落讨论边缘显著的小世界属性】

此外，小世界属性表现出边缘显著的交互效应。尽管未达到传统显著性阈值，其变化模式与全局效率、模块化保持一致：因果一致条件(BC)的小世界属性最高，让步冲突条件(AI)最低，且事后比较显示 因果一致条件(BC) 显著高于让步一致条件(AC) 和 因果不一致条件(BI)，让步一致条件(AC)显著高于让步不一致条件(AI) 。这一趋势提示，语义冲突同样削弱了脑网络局部聚类与全局整合之间的平衡，且效应方向与全局效率及模块化结果共同支持“连贯关系调节网络拓扑组织”的结论。

【4.2.1 节·关键脑区表述】

原“左侧额下回 (IFG) 是唯一在节点度、介数中心性和节点效率三项指标上均表现出显著交互效应的脑区”→ 修改为“左侧额下回 (IFG) 是唯一在节点度和节点效率两项指标上均表现出显著交互效应的脑区，同时在介数中心性指标上表现出边缘显著的交互效应”。

次要意见 3:

摘要及讨论中的“直接证明”“挑战经典模型”等措辞建议适当弱化。

回应:

十分感谢审稿人这一中肯建议。我们认识到，原稿中“直接证明”“挑战经典模型”等表述过于绝对。我们已在两轮修订中对摘要、讨论与结论中的相关措辞进行了系统弱化处理。

【中文摘要】: “MVPA 揭示...跨句式泛化解码失败” → “MVPA 显示...跨句式泛化解码

失败”

【中文摘要末段】“上述发现表明” → “上述发现提示”；“深度依赖于” → “在很大程度上依赖于”

【讨论】：“有力支持假设 4” → “为假设 4 提供了支持性证据”；“不存在独立于句法框架的抽象冲突信号” → “提示可能不存在独立于...的抽象冲突信号”；“挑战经典冲突监测模型” → “为经典冲突监测模型提供了重要的限定条件”

【4.3.2 节开篇】“本研究最具理论意义的发现之一是跨句式泛化解码的失败” → “跨句式泛化解码未能取得显著结果”

【结论】：“MVPA 进一步证明...句法框架特异性” → “MVPA 进一步提示...连贯关系特异性”

【英文摘要】：“Critically, cross-sentence-type generalization decoding failed entirely” → “failed within the current sample and analytical framework”；“demonstrated” → “indicated”；“advance coherence relation research” → “contribute to coherence relation research”

【英文摘要】“syntactically dependent” → “discourse-relation-dependent”；“rather than abstract, domain-general signals” → “rather than purely abstract, domain-general signals”

除上述具体修改外，我们还对全文进行了通盘检查，将“证明”“挑战”等强因果动词统一替换为“支持”“提示”“与...一致”“为...提供证据”“对...提出限定”等更中性的学术表达。修改后的措辞更符合 fMRI 研究作为相关证据而非因果证据的定位，也避免了过度推论。

回复中所引主要参考文献

- Jeffreys, H. (1961). *Theory of probability* (3rd ed.). Oxford: Clarendon Press.
- Lee, M. D., & Wagenmakers, E.-J. (2013). *Bayesian cognitive modeling: A practical course*. Cambridge University Press.
- Morey, R. D., & Rouder, J. N. (2011). Bayes factor approaches for testing interval null hypotheses. *Psychological Methods*, 16(4), 406–419.
- Xu, X., Jiang, X., & Zhou, X. (2015). When a causal assumption is not satisfied by reality: Differential brain responses to concessive and causal relations during sentence comprehension. *Language, Cognition and Neuroscience*, 30(6), 704–715.
- Xu, X., Chen, Q., Panther, K.-U., & Wu, Y. (2018). Influence of concessive and causal conjunctions on pragmatic processing: Online measures from eye movements and self-paced reading. *Discourse Processes*, 55(4), 387–409.
- Xu, X., Liu, J., & Zhou, X. (2022). Understanding an implicated causality: The brain network for processing concessive relations. *Brain and Language*, 234, 105177.

第二轮

审稿人 1 意见：

感谢作者的细致回复和修改工作。作者已经针对我第一轮所提出的问题做出了最大程度

的完善，新补充的内容显著增强了文章的可信性和严谨性。推荐发表。

审稿人 2 意见：

感谢作者对上一轮审稿意见的认真回应和修改。修订后的稿件质量已有明显提升，我此前提出的主要问题也已得到妥善解决。

编委意见：

建议录用。

主编意见：

同意发表。