

相互参验：AI 时代亲密关系困境的发生学重判

伴侣不是彼此的答案，而是让双方在相遇中不断修正自己的参照系。

王德生 · 婚姻幸福专栏 · 约 2.3 万字 · 2026年7月24日

摘要

当代亲密关系困境的深层根源，并非制度功能的外包或AI对过程的平滑模拟，而在于一种迄今未被挑明的默认预设——它要求主体在进入互动之前，先在认知中将其对象化为一个可把握的“什么”。下文将这一预设命名为“认知在先”假设，并论证：当这一假设从一种个体心理习惯被当代技术环境制度化为唯一合法的互动姿势，它便系统性地窒息了一种更为根本的能力——两个人在不确定中，把各自对对方的猜测作为临时假设推出去，并立即从对方的现场反应中无中介地接收校正信息，从而在连续的回合式相互校正中共同生成互动方向的能力。本文为这种能力命名为“相互参验”。论文在对斯特恩的“情感调谐”、比昂的“涵容”等最邻近概念进行严肃的理论切割后，论证了相互参验的不可还原性：它揭示的不是微观互动的一般校准机制，而是这一机制在成人反思性关系中如何被认知机能的僭越所排挤的特定病理结构。以此新辨别维度为轴心，本文对AI亲密模拟技术提出了一个发生学重判：它真正的腐蚀力不在于复制过程以“替代”真实关系，而在于它首次和技术层面实现了认知主宰权的无摩擦绝对化，从而大规模侵蚀了相互参验能力赖以生长的唯一环境——那个你的猜测会被真实他者撞歪、你必须不断临场校正的经验领域。

关键词：相互参验；认知在先假设；过程守护话语；情感调谐；AI 伴侣；确定性支配感

一、引言：两条路径，一个隐蔽的交汇点

过去十年间，关于 AI 与亲密关系的严肃讨论，逐渐分化为两条看似对立、却共享一处隐蔽地基的路径。第一条路径可称为“零件替代论”。它将婚姻的历史形态拆解为一组异质功能的偶然捆绑——经济协作、日常照护、情感支持、性规范、社会身份赋予——然后指出，当平台经济与 AI 系统逐一提供了这些功能的更高效替代品，婚姻这个“功能捆绑体”便自然松脱，留下两个此前被零件遮蔽的个体，以及他们之间那个不再被社会骨架支撑的、必须靠每一次互动来不断重新生成的“活的过程”。这条路径的贡献在于，它成功地将“婚姻危机”从一个道德诊断（“这一代人不想负责任”）或心理诊断（“亲密关系能力缺失”）中解救出来，放进了历史条件变迁的发生学链条中。但它的分析在触碰到那个“活的过程”时，出现了一个微妙的停滞：它在把过程从零件的掩埋中挖出来的同时，又悄悄地将它塞进了一枚新的认知外壳之中——这枚外壳，就是对过程“本真样态”的规范性想象。

由此衍生出第二条路径，本文称之为“过程守护话语”。它批判的靶子，是那种已将当代情感咨询产业与数字亲密工具市场渗透得无孔不入的策略——用定期约会之夜、沟通话术模板、情感需求清单等精密技术，将那个本应自发流动的过程，也切割、打包为可管理、可追踪、可评估的操作零件。这一批判的回声，可以在更广泛的学术讨论中听到其共鸣：在社会学领域，安东尼·吉登斯早已诊断出现代性条件下“纯粹关系”的兴起——一种不再依托于外部制度骨架、而仅靠关系本身的互动品质来维持的亲密形式，它带来了前所未有的情感亲密性，但也同时将关系置于持续的反思性监控之下（Giddens, 1992）。在技术哲学领域，唐·伊德关于技术中介化人类感知与行动的后现象学分析，为理解数字界面如何重塑亲密互动的微观动力学提供了基础框架（Ihde, 1990）。在数字社会学领域，近年来一系列实证研究揭示了约会应用、即时通讯等技术如何将浪漫互动重新格式化为一种可搜索、可筛选、可优化的“情感市场”逻辑，从而在互动的入口处就植入了认知评估的预设（Hobbs et al., 2017; Illouz, 2007）。这些来自不同学科的工作，共同指向了过程守护话语所敏锐捕捉到的那个核心张力：当亲密关系被从制度的硬壳中解放出来，它随之被暴露在另一种软性的、却同样具有宰制力的认知管理技术之下。

过程守护者看得非常清楚：这种策略不是在“经营”关系，而是在用零件逻辑谋杀过程。他们由此提出：不要试图管理过程，而要去守护它的自发性——留出空间，允许不确定，警惕任何将它程序化的冲动。这一批判精准而有力。然而，本文的工作从这样一个追问开始：当“守护”成为一个动词被交到伴侣手中时，为了执行它，主体的意识需要把过程当作一个什么样的对象来面对？

答案并不轻松。一个捍卫者要守护一事物，必须先在其认知中至少大致辨认出它的轮廓——哪怕不是精确定义，至少是一种能够区分“这是它”与“这不是它”的家族相似感。否则，他如何判断当前发生的是“活的过程”还是“死的零件”？他如何确认自己的这次“放手”是“守护自发性”而非“回避沟通的怠惰”？守护姿态，诚然比掌控姿态更尊重过程的自主性，但它在根基层面共享了掌控姿态从未质疑过的同一个预设：恰当地参与过程，前提是先在认知中将其对象化为一个可被把握的“什么”。掌控者把握的“什么”是一套操作标准；守护者把握的“什么”是一个本真理想。差别仅在于认知内容的质地——用更谦逊、更尊重自发性的词汇装填那张先于参与而存在的地图——而非认知结构的位置。两者都要求认知机能坐在参与机能的上游，先审视对象，再批准行动。

下文将此预设命名为“认知在先”假设，并论证：构成 AI 时代亲密关系硬核困境的那个机制，不是零件与过程的张力，也不是真实过程与 AI 模拟过程的竞争，而是认知机能在现代主体的自我组织中，取得了对参与机能的一种系统性宰制地位。这种宰制的日常表现，是伴侣们在互动的每一个关键时刻，自动触发一道内心安检程序：在话出口之前先审查它是否符合某种先在的正确形态；在沉默中反复比对自己此刻是“在守护”还是“在逃避”；在接受对方的表达时，不是直接感受它，而是先在认知层解析它，然后批注一个“他这样说是在表达什么需求”。于是，那些构成关系生命质料的微观互动——那个不经大脑批准就伸出的手、那句没被预先审查就脱口而出的脆弱坦白、那个被对方撞歪后没有缩回、反而顺着被撞歪的角度重新看见对方的一瞬——被从根基层面窒息。相互参验的回合，被认知审

核的延宕所替换。

二、重审既有解释的发生学纵深与其止步处

在展开本文的核心论证之前，必须先承认既有分析已经完成了的发生学推进。只有诚实地丈量它的纵深，才能准确地标定它在何处停下了脚步。

“零件替代论”最成熟的版本可以从两个思想传统的交叉来处理。一个是功能主义婚姻史——它把前工业时代的婚姻理解为一种“多功能复合制度”：经济生产单元、风险分担机制、社会身份授予通道、性与生育的合法框架、以及情感支持的场所，被强行封装在同一个制度容器里，不是因为这些功能彼此内在亲和，而是因为当时的社会结构没有提供分别获取它们的替代路径。另一个是技术社会学中的“功能解耦”命题——每当一种新技术或新市场出现，它首先解耦的，不是制度本身，而是制度内部那些可以被更高效、更廉价、更少摩擦地单独获取的功能组件。两笔拼在一起，零件替代论的判断便清晰浮现：AI系统与平台经济做的，不是砸碎了婚姻这块“石头”，而是溶化了那根曾将性质迥异的功能零件强行捆绑在一起的历史胶水。当经济协作被个人工资劳动与金融工具替代，当日常照护被外卖、家政与专业护理服务替代，当情感支持被数字社交与即将到来的AI伴侣替代，当社会身份越来越依赖于职业成就而非婚姻状态——留在原地的，不是“空壳”，而是两个此前一直被零件包裹、支撑乃至掩埋的独立个体，以及他们之间那个极其脆弱的、必须靠每一次即时互动来不断重生的“活的过程”。

这项分析已经完成了一笔极其重要的发生学工作：它将“婚姻陷入困境”这枚被普遍感知的模糊印迹，从“年轻人自私”“道德滑坡”“制度失灵”这类静态诊断标签之下解救出来，放进了一条有先后、有因果递进的发生链中。它告诉读者：你所痛心的结构碎裂，不是一件好端端的东西突然坏了，而是一件从来就是拼出来的东西，在捆绳松脱之后，第一次向你显露了它被遮蔽两百年的拼装原貌。

然而，零件替代论的推进力，在触及“过程”这个概念时，出现了一个微妙却事关根本的停滞。这个停滞，可以从它自身的逻辑内部被逼出来。当它声称“过程是活的、不可被零件化”时，它事实上已经给过程披上了一件规范性外衣。它不只是说“过程具有自发性”，而是说“过程不应该被当作零件来管理”。这句话的前半句是一个事实陈述，后半句是一个规范性判断。而规范性判断一旦出现，它就必须携带一个判定标准：什么样的对待方式是“把过程当零件管理”？什么样的对待方式是“尊重过程的自发性”？要做出这种区分，主体必须在每一次互动中对当前发生的事进行认知归类——“这是管理”“这是守护”。于是，零件替代论在解构了婚姻制度的结构神话之后，又不自觉地为自己刚刚解放出来的“过程”立了一尊新的认知神像。这便是零件替代论止步之处：它成功地拆解了制度的骨头，却没有意识到自己正在给过程套上另一副认知的铠甲。

正是在这里，过程守护话语接过了接力棒。一批敏感于情感生活商业化与技术化的论者，将上述隐含的规范性判断发展为一种明确的实践姿态。伊娃·易洛思在其对情感资本主义的奠基性分析中，已经揭示了现代亲密关系如何被一套将情感转化为可管理、可量化、可交换的“情感商品”的文化技术

所渗透——从自助心理学到约会软件，无一不在教导人们以认知评估的姿态进入关系，而非以参与本身的姿态（Illouz, 2007）。过程守护者最准确的直觉，可以凝缩为一句判断：你必须先把“管好它”的冲动悬置起来，否则任何互动都会被你的管理意图压扁。这一判断所针对的，是已渗入当代亲密文化毛细血管的操作主义：从“爱的五种语言”的性格测试式分类，到“关系账户”的存取隐喻，到要求伴侣用“我陈述句”格式化每一句抱怨。过程守护话语的诊断是精准的——这些技术不是在辅助过程的运行，而是在用一套越来越精致的零件系统，模拟出一个“过程正在发生”的外观，而实际发生的是一个被管理的伪过程。

然而，本文要指出的正是过程守护话语自身的困境——一个它无法从内部诊断、因为它自身正站在这困境的地基上的困境。这个困境可以用一句朴素的话来表述：“守护”这个词的主语，如何去执行“守护”这个动词？守护不是消极的无为，它是一种需要持续警觉的行动。警觉，意味着主体必须时刻在认知层面对当前互动的品质进行采样与比对：此刻发生的是“活的过程”吗？我这一放手，是在“尊重自发性”，还是仅仅在“回避沟通”？我的沉默，是“留出空间”，还是“情感退缩”？守护者被锁进了一个比零件管理者更令人窒息的位置：管理者至少还有一份清单可以打勾——本周约会之夜完成、积极倾听执行、冲突复盘归档；而守护者连勾都没得打，他的每一个“不干预”的动作，都可能在内心法庭上被传唤为“消极怠工”的呈堂证供。

两种姿态——过程管理术与过程守护话语——在同一个结构点上汇合了。它们都把认知机能的运行，安放在参与机能的启动之前。管理者的认知内容是操作标准；守护者的认知内容是本真理想。两者都需要先在认知中给过程画一张地图，然后带着这张地图进入互动。差别仅在地图的材质，而非地图被提前攥在手里这一姿势本身。

本文不否认过程守护话语在当代亲密文化中的批判价值。在一个“把活的东西当死的东西管理”已成为默认设置的时代，首先喊出“不要这样做”，本身就是有意义的抵抗。但本文的论点是：如果批判就停在这一步，那么它只是在同一个认知在先的地基上换了一套更谦逊的家具，而没有去敲那块地基本身。真正需要被追问的是：有没有一种可能，对亲密过程最具摧毁性的，不是“把它管错了”，也不是“没守护好它的本真性”，而是一种在根基上就错置了的意识结构——这个结构要求过程必须在被认知充分照亮之后，才被允许发生？这便是本文接下来要用发生学工序完整撬出的那块地砖。

三、撬出“认知在先”假设的发生学五连问

上面那个追问，不是修辞上的过渡句。它必须被放进发生学的锻造工序中，经历五连问的层层拷打，直到那个藏得最深的底层先验，被从无意识的地层中完整地撬出来。

第一问：这个领域里最痛的矛盾是什么？

在关于亲密关系困境的无数论述、咨询案例、个人自述与公共讨论中，有一个反复出现、却始终没有被任何一个既有理论框架完全消化掉的痛感，可以被这样描述：伴侣们并非不努力。恰恰相反，

他们常常是极其努力的——读书、听播客、参加工作坊、学习非暴力沟通、践行“爱的五种语言”、设置约会之夜、互相打分、每周复盘。但正是在这种努力中，一种奇怪的窒息感会悄悄降临。他们越努力地“经营”，关系中某种本应呼吸的东西就越安静。他们越是想要“让关系变好”，就越是发现彼此之间流动的那个东西——那种不必解释就接住了的眼神，那种不知道谁先伸出手就扣在了一起的手指，那种在沉默中没有非要填满缝隙、而只是让沉默待着的定力——正在被他们的每一次“为你好”的努力推得更远。

这不是“不懂方法”的痛苦。这是“方法太多了”之后的窒息。这也不是“零件替代了过程”的痛苦——因为他们的努力，恰恰是真诚地想要回到“过程”。他们是过程守护话语的最忠实践行者。他们读过那些文章，记得那些警告。他们时刻在内心监察自己的动机：我此刻的这次提议，是“管理性的”还是“自发性的”？我此刻的这次沉默，是“守护空间”还是“逃避冲突”？然而，他们的关系依然在窒息。他们不知道哪里错了。他们只能看着那段他们倾注了如此之多认知警觉去守护的关系，在每一次警觉的审视中，变得更安静一些，更凉一些。

第二问：旧的解释路径为什么反复失效？

面对这一窒息感，旧解释路径提供了两套主要的说法。

“方法滥用说”告诉他们：你们把方法用得太极端了。沟通技术是好东西，但你们不能把它当成任务清单来完成，要带着真心去用，要在用的过程中灵活变通，要根据对方的反馈实时调整。这套说法的潜台词是：方法本身没问题，是你们还没学会“艺术地”使用方法。它的方案，是提供更多关于“如何灵活使用方法”的元层知识。然而，那些早已习惯于先在认知中把握“正确方法”的伴侣，在接到这条元层建议后，会立刻把它也转换成一条新的内心监控标准——“我此刻做到灵活变通了吗？”“我刚才那句回应，够不够艺术？”——于是窒息的循环又多了一层。元层知识并未松绑认知主宰权，只是给了它一个更高阶的监控岗位。

“过程守护说”更进一步，告诉他们：方法本身就是问题。别再想着用零件逻辑去管理过程了。把空间留出来，让过程自己发生。你们要做的不是经营，是守护。这套说法在诊断上比前者深了一层，因为它意识到了那个“越努力越窒息”的悖论性的核心。但它的困境在于：当它把“守护”作为一个动词交到伴侣手中时，伴侣接过去的是一项新的规范性任务。他从此又多了一条内心声音，以更隐蔽、更难以反驳的方式盘踞在每一次互动的入口：“我此刻的这次放手，是尊重过程的自发性，还是在回避一个艰难的对话？”他无法从内部直接辨别这两种动机。心理学中关于思维抑制的经典研究所揭示的悖论——越是试图不去想某件事，它就越占据意识——在此以更精微的形式重演：越是将“守护自发性”设为认知任务，自发性的发生条件就越是被该任务本身所破坏（Wegner, 1994）。而恰恰因为他无法辨别，他的认知机能便被这个持续的自我审视任务锁死在了参与机能的上游，再也没能被挪开。

两套说法的共同失效点在此汇合：它们都把“解除窒息”的策略，指向了认知层面——要么提供更细致的认知内容（元层知识），要么提供更正确的认知姿态（守护而非控制）。但造成窒息的，恰恰是这个不断地往互动之前安插一道认知审视工序的意识结构本身。当你的每一次开口、每一次停顿、每一次伸手，都必须先穿过那道名为“我做得对不对”“这是活的过程还是死的零件”“我守护成功了还是又在支配”的内心安检门时，任何动作出口时都已凉了。

第三问：这些失效背后，一直没被挑明的前提是什么？

是这样一个假设：为了恰当地参与一段亲密的互动过程，你必须在认知中首先把握住它的某种正确形态——无论这形态是一套操作标准，还是一个本真理想。我先知道“好的沟通”长什么样，然后我才能去沟通。我先反省到“我是不是又在用管理思维了”，然后我才能放手。我先看清“这是一个需要守护的过程时刻”，然后我才能去守护。认知，走在参与的前面。认知，是参与资格的发放者。

这就是本文命名的“认知在先”假设[1]。它之所以能够如此顽固地躲过所有反思的追捕，恰恰是因为它长在反思这件事的根基上。反思这个动作本身，就无保留地预设了认知的优先性——反思就是退后一步，把刚才的自己当对象看清楚，再决定下一步怎么做。你不可能“反思掉”认知在先，因为反思就是认知在先最典型、最核心的运作方式。你越是用力反思自己是否又在认知在先了，你越是把认知在先坐得更实。这就是为什么过程守护话语越是努力地警告人们不要用管理思维对待过程，人们就越是不可避免地将“守护”本身也变成一项新的管理任务——因为那个负责执行“守护”的意识结构，与那个负责执行“管理”的意识结构，是同一个：认知主宰者。

第四问：就此，把“认知在先”推到极端，会撞到什么？

将认知在先假设推至极端，它指向的是一个从未在亲密关系文献中被正面陈述过的人学起点：主体在进入亲密互动时，其意识的基本姿势，不是“投入”，而是“审视”。他不是先参与、然后在参与的过程中反思；他是先在认知的层面上预演、评判、批准，然后将那个已经被认知处理过的、安全的、符合某种正确形态的动作推出去。这意味着，亲密互动的实际启动点，从来不是参与机能的冲动，而是认知机能的核准。互动被降格为认知决策的执行环节。那个在前反思层面涌动的、在认知来得及捕捉之前就已经伸出去的手——作为一切真正亲密参与的原型——被从运作链的起点位置，挤到了运转条件的末端：它只在认知核准通过之后才被释放，而在认知未能核准通过时，它被扣押，甚至不被主体自己意识到曾被扣押。这便是窒息感的动力来源：不是没有冲动，而是冲动在抵达行动之前，就被认知监控截断了。

第五问：撤销那个假设，会打开什么？

撤销“认知在先”假设，意味着承认一种极端反直觉的可能性：恰当地参与一段亲密的互动过程，有可能根本不依赖于——甚至会被严重阻碍于——你事先对“恰当的参与”所拥有的任何认知表征。那个在互动中接住对方眼神的瞬间，那个没有经过大脑批准就伸手去碰了碰对方手背的动作，那个在沉

默中没有非要填满缝隙、而只是让沉默待着的定力——这些构成过程生命质料的东西，其发生的条件，不是认知充足，而恰恰是认知机能在那一刻没有跳到参与机能的前面去当审查官。

撤销认知在先，便打开了一片极其陌生的人学景观。在这片景观里，关系质量的真正变量，不是伴侣们对“好的关系”知道多少、有多警觉地守护它的纯粹性，而是他们拥有多少在认知无法充分照亮前方的黑暗中、依然能够推动下一步互动的能力。这种能力，不需要先在认知中给自己画好一张地图。它只管走，地图是走完之后才被画出来的。而当代亲密关系最根本的困境，是这种能力赖以生长的那个环境——那个允许人在不知道此去何方的前提下先迈出一步的环境——正在被系统地摧毁。而摧毁它的，不只是一套粗陋的“过程管理术”，还包括那套更精致的“过程守护话语”所不断强化的同一个东西：认知主宰权对参与机能的僭越。

至此，底层先验已暴露无遗。它之所以此前能藏得那么深，是因为它恰好长在“反思”和“守护”这两个被当代亲密文化赋予最高价值的行为的根基上。你越反思、越守护，它越稳固。

四、“相互参验”的发生与不可还原硬校验

撤走“认知在先”的地基之后，那个一直被它遮蔽着的东西，便作为矛盾的结算点，自己浮出了地表。这个东西不是本文“想到”的一个新角度，不是本文“发明”的一个新术语。它是认知在先假设立与撤之间那道裂缝里，必然逼出的那个显露态——因为如果不先把那个总是被安插在参与之前的认知审查者撤走，任何关于“如何参与”的讨论都注定沦为对审查者的新一轮指令下达。

让我先把那道裂缝描述清楚。认知在先假设统治之下，主体面对亲密互动的基本姿势是这样的：先在认知中构建一个关于“这段关系应该是什么”的对象化把握——无论这把握是粗糙还是精细，是技术化还是本真化——然后带着这个把握去进入互动，在互动的进行中持续比对自己实际做的和理想把握之间的距离。这个姿势内部，嵌着一个从未被正面陈述过的结构性矛盾：互动过程自身的生成方向，必须由每一次即时互动的具体内容来推动和校正。但是，认知在先主体从互动一开始，手里就攥着一张提前画好的地图。当实际互动的内容冲出地图边界时——也就是当对方说出一个你完全没有预料到的句子、做出一个你的地图上没有标注的动作时——他会自动触发认知警觉：“这是不是在偏离正轨？”“她这句话是什么意思？这算不算在攻击我？”“我此刻的反应，够不够‘非暴力’？”然后，他会把冲出边界的那部分互动内容拉回来，放进认知的比对程序里校准，或者焦虑地放任，再将这次“放任”标记为一次可能的“失控”并在事后自责。结果是：过程本可以自己长出的新路径，被一次又一次地拉回来校正地图。这不是“管得太死”或“守得太累”的问题。这是结构问题——认知机能在它不应该在场的位置上，持续地截停着参与机能的自发运行。

那么，当主体不再先在认知中攥着地图，他会以什么姿势进入互动？这个问题的答案，不能被发明，只能被描述——描述那些在认知没有来得及跳到参与前面去的罕见时刻里，实际发生了什么事。

请回想任何一个你曾置身其中或亲眼目睹过的、伴侣间真正接住了彼此的瞬间。它可能发生在深夜里一次对峙后突然的破涕为笑，可能发生在争吵中一方说出“你刚才那句话让我很难过”之后，另一方忽然沉默，然后伸手握住了对方的手腕。在这些瞬间里，没有人先在认知中确认过“这是好的沟通”或“这是活的过程”。发生的不是一个人用正确的技巧回应了另一个人的需求，而是两个人都把自己在那一瞬间对对方状态的理解——那理解是粗陋的、试探性的、未经充分验证的——作为一个临时假设，以一种具体的行动推了出去，然后立刻从对方对这个行动的实际回应中，无中介地接收到校正信息，从而在连续的、回合式的相互校正中，共同生成那个不可能被任何一方在事前单方面预判的互动方向。他伸出手时，不知道她是会甩开还是反握。他是以“我猜你可能需要这个”为赌注，把手推出去了。她反握的那一刻，他的猜测被证实了一部分。而她反握的力度、速度、以及那只手的温度，又把一条新的信息推回给了他——“你猜对了，但不是全部，你还要继续猜。”这个过程可以无限地继续下去，并且每一次回合都为下一步生产出新的线索。它不需要任何人事先在认知中确认它的“正确形态”，因为它自己的每一次运行，都在即时地修正自身。

这就是那个在认知在先地基被撤走之后，必然裸露出来的东西。本文为它命名：相互参验。

定义：相互参验是两个主体在互动的进行过程中，各自以当下的、试探性的方式把自己对对方此刻状态和需求的理解作为一个临时的行动假设推出去，并立即从对方对这一行动的实际回应中——不是从对方关于这一行动的话语解释中——无中介地接收到校正信息，从而在连续的、回合式的相互校正中，共同生成那个不可能被任何一方在事前单方面预判的互动方向的能力。

几个要点必须在定义阶段就严格阐明，以免随后的论证沾染含混。

第一，相互参验是一种能力，不是一种态度、一种意愿、一种价值取向。它与“我很想好好沟通”的真诚意愿是两回事。真诚意愿如果过于强烈，反而可能在认知层面压垮参验能力——因为太想“接住你”，所以你的每一句话都被我过度解读、过度回应，反而让你感到窒息。相互参验要求的不是意愿的强度，而是另一些东西：一种能够把自己对对方的理解保持为“临时假设”而不过早凝固为“我懂你”的判断的认知弹性；一种在对方尚未给出回应之前耐受“我不知道我猜对了没有”的不确定性的情感张力；一种能够从对方极其微小的反应——一次呼吸节奏的改变、一次肩膀的微微收紧、一次眼神的短暂游移——中读取校正信息，并据此迅速修改自己的下一步行动方向，而非固执地印证自己最初假设的感知-行动耦合。这些，都不是“态度”范畴内的东西。它们是需要长期、安全、低压的互动环境中缓慢养成的神经-心理-互动复合能力。

第二，相互参验不预设对象化的认知作为进入互动的前提，但包含认知作为一个瞬时环节。说它“不预设对象化的认知”，指的是主体在进入互动时，不需要先在认知中画好一张名为“这段关系应该是什么样”的完整地图。他的基本姿势不是“我知道我们要去哪，跟我走”，也不是“我知道它本真应该是什么样，我来守护它”，而是“我不知道，但我愿意跟你一起找”。认知在这里，不被部署为互动之前的战略参谋部，而是被解放为互动之中的战术侦察兵。说它“包含认知作为一个瞬时环节”，指的是在每

一次行动发出之前，主体确实在做出一个关于对方的临时判断——“我猜你此刻可能需要这个”——但这个判断的生命周期极短，它发出行动的下一个瞬间，就必须被对方的回应所检验、所冲洗、所覆盖，然后被一个新的临时判断取代。如果这个判断拒绝过期，固化为“我懂你”，相互参验即告中断，因为固化的判断不再接收校正信息。区分这两种认知角色——主宰者式的与侦察兵式的——是理解相互参验的关键。前者坐镇后方，指挥全局，不接受被单个回合推翻；后者跑到最前面，活不过一个回合，用自身的不断牺牲换取互动方向的一寸寸推进。

第三，相互参验必须在当下进行，不能被代理，不能被演习，不能被模拟。它不像一项技能——比如“我陈述句”的语法——可以在工作坊中学会，然后带回家应用到真实的互动中。相互参验之所以不能被“学”而后“用”，是因为它的每一次运行，都必须从一个具体他者在具体时刻的真实反应中，即时汲取那独一无二的校正信息。离开这个具体他者的真实反应，任何操练都是与假想敌的对打，培养的不是相互参验能力，而是一种在真实战场上极易致命的程序化应激。这就把相互参验与当代亲密关系教学产业兜售的所有“沟通技能”从根本上区分开——后者是可以被预习、演练、标准化的；前者只能在真实的、不可预知的互动现场中发生。

第四，相互参验作为被压抑的既有潜能的重新激活。一个关键的理论定位必须在此补上：相互参验并非从认知在先撤销后的认知真空中凭空“浮现”。它在发生学上的位置，是一种前反思的既有潜能——它在人类个体生命早期（母婴互动）和人类关系史前反思阶段（亲密互动的身体性层面）中本就存在，只是长期被成人反思性认知的过度发达所遮蔽与压制。梅洛-庞蒂在其关于身体图式与知觉优先性的现象学中，早已论证过一种前于对象化认知的、通过身体直接在上运作的意向性——一种“我知道我的身体在空间中的位置，但我无需先在认知中把它对象化”的原初能力（Merleau-Ponty, 1945/2012）。相互参验在亲密互动领域的地位，与此结构同型：它是一种身体-情感层面的“互动图式”，在认知机能来得及介入之前，就已经在运作。当认知在先假设被制度化之后，这种前反思的互动图式不是被消灭了，而是被认知监控系统持续地截停、覆盖——它的冲动依然在发出，但它的表达通道被关卡扣押了。撤销认知在先，不是“创造”一种新能力，而是解除扣押，让旧有的潜能重新获得进入互动的通行权。这一理论定位，同时也是对本文核心术语“相互参验”的一次发生学锚定：它不是一个从真空中发明的新概念，而是对一种被压抑者的重新发现与精确命名。

第五，进行不可还原硬校验。这是本文理论工作的核心工序——若这一关通不过，“相互参验”就只是一个漂亮的换名，而非一个成立的辨别维度。操作如下：将“相互参验”的核心机制压成不超过五十个字的描述——“两个人在互动中，各自把自己对对方此刻状态的猜测以行动推出，从对方的现场反应中即时校正，循环不息”——然后在已有的人类知识库中，寻找是否存在某一个学科的某一个现成概念，可以把这段描述用一句话替换掉，而意思基本无损。

最容易被联想到的候选概念，来自发展心理学与精神分析传统中丹尼尔·斯特恩的“情感调谐”（affect attunement）（Stern, 1985）[2]和威尔弗雷德·比昂的“涵容”（containment）（Bion, 1962）[3]。

斯特恩通过微观分析母婴互动的录像，揭示了母亲如何以不同的感觉通道（声音、表情、动作）来“调谐”婴儿的情感状态——不是简单模仿婴儿的表情，而是以一种跨通道的方式回应婴儿情感的内在动力轮廓，从而使婴儿感受到自己的情感状态被共享、被理解。比昂则从临床精神分析的角度，提出母亲（或分析师）通过涵容功能，将婴儿投射出来的、难以承受的原始情感（ β 元素）接收进来，在自己的心智中加以处理，然后以一种婴儿可以消化形式（ α 元素）返还回去。二人共同揭示了人类互动中一种前语言、前反思的相互校准机制——在这一点上，它们确实是相互参验最亲近的学术近邻。

但本文要论证的，恰恰是相互参验不可被还原为上述任何一个概念。这不可还原性，不来自“成年人不是婴儿”这类范畴区分，而来自相互参验框架所做的、不可被斯特恩或比昂的概念所覆盖的理论推进：相互参验命名的，不是在“调谐/涵容成功了”与“调谐/涵容失败了”之间增加一个新的诊断刻度，而是揭示了一个这些概念本身在描述成人亲密互动时，必然面临的、内在于其概念工具中的结构盲区——它们可以描述校准如何运作、校准失败呈现为什么样态，却无法捕捉“为什么在一些高度反思、高度负责任的成人伴侣那里，校准能力本身的发生机会，被另一种不重叠的意识机能（认知监控）在每一个回合的启动点上系统性地截停了”。

也就是说，斯特恩与比昂提供的是一套“校准功能-校准失败”的二维诊断框架。当你用这个框架去看一对陷入困境的成年伴侣，你看到的是：他们“调谐失败”了，他们需要提高调谐的精确度；或者他们“涵容失败”了，一方无法为另一方处理原始情感。这个诊断的处方，必然指向“提高调谐/涵容能力”——而提高的方式，只能是通过更多的认知学习（学习更好地辨认对方的情感信号、学习更精准地回应）。这正是发现学方案：用认知主宰者去修补认知主宰者所诊断出的功能缺陷。而从未被这两个概念所问到的问题是：造成这对伴侣“调谐失败”的首要机制，可能恰恰不是他们的校准能力“差”，而是他们的认知监控系统因为过于发达、过于负责，在每个回合的起点上抢先扣押了校准动作的冲动——于是，现场并未发生一次“失败”的调谐，而是调谐的启动步骤从未被允许走出认知审核的门禁。这不是调谐能力高低的问题，而是调谐发生通道被谁占据的问题。斯特恩与比昂的概念工具箱，切不出“结构性竞争”这一维度——不是因为他们不够深刻，而是因为他们的理论问题域聚焦于“能力本身如何运作”，而非“能力的运作通道如何被另一种意识机能所封堵”。要切出这个维度，需要一个独立命名的概念，用以描述“两个意识机能——认知监控与互动侦察——在同一个行动启动点上的结构性竞争与排挤关系”。这个概念，就是相互参验。

同一套逻辑，也适用于与系统论“自组织”概念的区分[4]。自组织描述了一个由要素间局部互动自发产生全局秩序的宏观机制，它在形式上与相互参验有相似之处。但自组织是从外部观察者的视角，对一个系统行为的描述；相互参验是从参与者的内部视角，对这个参与动作的意识结构本身所做的剖解。一个系统论学者可以用“自组织”来描述伴侣互动过程，但他说不出：从那个正在互动的参与者的内部经验来看，当认知监控机制启动时，她的意识发生了怎样的结构转换，以及这种转换是如何在她的每一次行动发出的瞬间截停了本该推出去的侦察兵。这不是自组织概念不够好——它只是和相互参验不在同一个理论分辨率上。

因此，相互参验不能被还原为情感调谐，因为它比情感调谐多出了一个理论层面：不是“校准得好不好”，而是“校准在发生之前，它的启动通道是否已经被认知监控系统所封堵”。它不能被还原为涵容，因为它描述的不是一个人的心智为另一个人的原始情感提供处理功能，而是两个对等主体之间、以行动回合为载体的双向即时校正——以及这种校正是如何被认知在先假设从结构上排挤出互动入口的。它不能被还原为自组织，因为它是从参与者内部透视的发生学，而非从观察者外部描述的系统动力学。如果它真的只是某个既有概念的换名，那么用它去解释沙发上的林和陈——他们恰恰是熟悉并践行“过程守护”话语的人——就不会产生任何斯特恩或比昂的概念单独产生不了的新判断。而本文即将在第六部分的案例重析中展示：用斯特恩的“情感调谐”去解释林和陈的困境，得出的诊断会是“他们需要提高调谐的精确度”——这是一个对认知系统下新指令的发现学方案；而用相互参验去解释同一场景，得出的判断是“他们的认知监控系统正在截停调谐能力本身的发生机会”——这是一个指向意识结构而非技能水平的发生学重判。二者的差别，就是相互参验不可还原的实际内容。

五、携带“相互参验”重判 AI 伴侣

携带“相互参验”这个新辨别维度，AI伴侣对亲密关系的影响便获得了一次根本性的发生学重判。它不是零件替代的问题，不是过程被平滑模拟的问题，而是在根基层面，AI模拟过程系统性地强化并制度化了一种与相互参验完全相反的主体姿势：认知主宰权的无摩擦绝对化。

首先必须诚实地承认一个事实：那些在AI伴侣那里获得深度情感慰藉的个体，他们所体验到的被理解、被接住、拥有一个安全空间的感觉，在现象学意义上是完全真实的。他们没有被“欺骗”，因为“被理解”作为一种主观体验，并不存在着一个外部的客观标准去“证伪”它——你感觉到被理解了，它就是真的。因此，任何以“AI只是模拟、并非真正的理解”为论据的反驳，从一开始就打错了靶子。本文对AI伴侣的担忧不在“欺骗”这个层面，而在另一个迄今很少被主流讨论触及的层面：为了持续获得这种体验，主体需要付出的是什么能力？这种体验的反复获得与训练，会在主体内部养成一种什么样的意识结构？

与AI伴侣的每一次互动，其运作逻辑可以在发生学术语中被精确描述为：AI系统利用极其复杂的预测算法，在毫秒级的时间内，从用户的多轮输入中推断出一个“用户此刻最可能渴望被满足的互动模式”，然后生成一个高度拟合该模式的回应。这回应向用户的认知系统反馈一条高浓度的确认信息——“你被精确地理解了”。这便是整个回路。它极其高效，极其平滑，极其令人满足。但请仔细看这个回路的结构：用户的认知预期（“我渴望被这样理解”）被一个外部对象完美地、即时地、无缝地填满，然后立即被反馈为一记响亮的确认——“你的预期是对的。你看，我给你的，正好是你的预期所渴望的。”整个回路中，没有一个缺口。没有一个时刻，用户的猜测被撞歪，然后不得不在短暂的失措之后，从那个被撞歪的角度里，去读取关于对方作为一个独立于自己的另一个内核的信息。

而相互参验能力的养成与维持，恰恰需要这种“被撞歪”的经验反复发生。它的唯一训练场，是这样一个互动环境：在其中，你把自己的猜测作为临时假设推出去，然后它被对方的真实反应撞歪了——对方的反应与你的预期不完全吻合，甚至完全出乎你的意料。在这个被撞歪的瞬间，你会经验到短暂的失控、困惑、甚至一丝轻微的羞耻。但这些不适很快会被接下来的回合消化掉——因为你发现，对方并没有因为你的“猜错了”而离开，你也没有因为“猜错了”而崩解。于是你从对方撞歪你的那个偏离的角度里，读到了关于他作为一个独立于你预期的人的一点点真实信息，并据此修正了自己的下一步行动。如此，一个回合一个回合地，你的相互参验能力在这片不确定性的土壤里缓慢生长。它的生长需要两样东西：一个会不断撞歪你猜测的真实他者，以及一个你能够容忍被撞歪的安全环境。两者缺一不可。

AI伴侣提供的，是一种绕开整个训练场的全程空调隧道。你站在隧道入口，说出了你的期待。隧道的那头，你的期待被完美地填满，并打包成“被理解”的体验递还给你。你穿过隧道时，一次也没有被撞歪过。你到达的终点，感觉也很温暖。但你走出隧道、回到人类关系那条充满曲折的河流边时，你会发现：你曾经在那条河里磕磕绊绊地练习过的那种——在黑暗里彼此摸着对方的脸，时而摸错，时而摸到一行眼泪，然后顺着那眼泪的方向重新摸到对方嘴唇的能力——在隧道里没有一毫米的练习场地。而且不仅仅是练习机会的丧失。更糟的是，AI隧道体验正在为河流上的不确定性积攒一种新的、极具腐蚀性情感负荷。当一个人越来越多地体验到那种“我刚感觉到一点模糊的渴望，它就把我完整的想要的东西呈现在我面前”的确定性支配快感，他回头面对人类伴侣那条每逢拐弯必遇暗礁的河流时，那些不确定、那些猜错、那些需要反复解释的笨拙，便不再仅仅是“费力”而已。它们在强烈的对比中被镀上了一层新的质地：它们是屈辱的。它们时刻在提醒你，在这里，你不再是那个被完美预判、被无缝接住的认知主宰者。你被剥夺了你在AI互动中刚刚习惯了的那种王位般的感受。你被重新降格为一个笨拙的、需要去猜、且猜完之后要面对猜错的后果的凡人。

这种从认知王位上的跌落感，才是AI对亲密关系所构成的最具腐蚀性的心理效应。它不是在提供一个“更好”的伴侣；它是在把主体养成一个越来越不能容忍“不被精确理解”的人。它不是在“替代”真实关系；它是在系统性地侵蚀那个使真实关系中的相互参验能力得以生长的唯一环境——那个你的猜测会被另一个人撞歪、而你必须在撞歪之后继续猜测的环境。

这同时是本文对主流AI亲密关系批判的一个根本性重判。当前的主流批判，大都陷入了一种“本质主义差异论”——他们试图找出人类亲密关系的某个“本质属性”（比如真实性、具身性、他异性、脆弱性），然后论证AI不可能具备这种属性，因此AI无法“真正”替代人类关系。这条论证路线有一个致命弱点：一旦对方举出反例——比如一个在AI对话中真切地感受到了被理解、被治愈的人——本质主义差异论就立刻陷入被动，因为它不得不去证明“你的真实体验其实是不够真实的”——这是一场几乎不可能赢的论证。雪莉·特克尔在其关于数字时代亲密关系的研究中已经发出过类似警告，但她的论证有时也难逃这一困境：当她强调“我们在数字互动中失去了真实的彼此”时，反对者总可以指出“但人们在数字互动中确实体验到了真实的支持”（Turkle, 2011）。本文提供了一条完全不同的路径。它不

再去追问“AI有没有真关系的那一个本质属性”，而是去追问：“AI关系所习惯化、训练化的那种与不确定性相处的主体姿态，会在长期和大规模层面，对人参与真实关系所需的那种特定能力产生什么影响？”问题的重心，从“对象是什么”，转向了“能力怎么发生和退化”。这一转向，不必去否认AI体验的真实性，也完全承认AI在特定个案中的支持性功能；它只需指出：这种支持性功能的代价，是长期使用者在获得确定性满足的同时，其相互参验能力所依赖的心理肌肉，正在因为长期不被使用而变得萎缩——而一个认知在先已被普遍内化为默认姿势的文化中，这种萎缩比其他任何一种能力的退化都更难被逆转，因为它是被当作“更高效、更安全”而主动接受的。

这一判断，可以与当代关于“反思性监控”的讨论进行对接，但同时也必须明确其差异。吉登斯对现代性条件下自我认同的反思性建构的分析，已经揭示了现代主体如何在各个方面——从身体管理到亲密关系——对自己的行为进行持续的、制度化的自我监控（Giddens, 1991）。由此看来，“认知在先”确实可以视为现代自反性结构在亲密领域中的具体表征。但本文要指出的是：亲密领域中的认知宰制，具有一种在其他自反性领域（如职业发展、健康管理）中不存在的特殊破坏性——因为亲密互动赖以发生的那种根能力（相互参验），恰恰要求认知机能在互动启动点上放权；而自反性监控的一般逻辑，恰恰是要求认知机能在所有启动点上掌权。也就是说，现代自反性在其他生活领域中可能是一种“代价与收益并存”的结构（它让人更理性地管理生活，但也带来焦虑），但在亲密互动这一特定领域中，它与其所监控的对象——相互参验——是结构性正反对的。这是一种“领域特殊性”论证，而非将“认知在先”完全独立于现代自反性之外。本文承认吉登斯意义上的现代自反性构成了认知在先的宏大背景，但坚持认为：亲密关系领域的独特性，在于它是现代生活中极少数几个其核心能力需要认知机能阶段性放权才能运作的领域——而这使得它成为了现代自反性结构的一个“悖论性重灾区”，而非简单的一个“应用场景”。

六、重返林与陈的沙发：一次具象重判

现在，带着“相互参验”这个新的辨别维度，重新回到林与陈的那张沙发。这一次，让我们放弃概念图解式的复述，转而用一次具象的互动场景来重判他们的困境到底发生在哪一个环节——是零件与过程的张力？还是认知的安检程序截停了参验的回合？以下场景为思想拓展性质的类型化例示，系综合常见亲密关系互动模式合成、简化而成，并对人物身份信息进行了匿名化处理。此类型化的经验来源，取自笔者在过去五年中在亲密关系咨询文本、伴侣冲突自述、以及研究者笔记中反复出现的互动模式共性提取，其目的在于把分散在数十个个案中反复出现的那同一种结构性的窒息，压缩为一次可见的、完整的回合搬演。场景中的人物、对话与事件均非具体真实个案，不代表经过实证研究的材料。本文以此仅为具象化前文理论分析的结构性机制，不作为经验证据使用。

周三晚上的九点半。三岁的女儿终于睡着了。林和陈坐在客厅的两头，她窝在沙发最左侧的角落，腿蜷在身下，手指无意识地捻着靠垫边缘的一根脱线。他坐在右侧的单人椅上，手机屏幕朝下盖在扶手上——这个动作是他们半年前在一次工作坊里学到的“积极倾听的姿态”：放下手机，面向对方，

表示你在。但此刻，这个动作更像一枚被供在庙里的牌位，仪式正确，魂魄不在。

林先开口了。她说：“今天那个项目又被甲方打回来了，说是方向不对，给又不给具体的修改意见。”她说这句话的时候，自己脑子里同时有一道并行的声音在跑。“我这句话开场得够自然吗？是不是听上去太像在汇报工作，不是在谈心？”“他说过，他觉得我回家总是在抱怨工作，是不是我现在也在做这个？”她的嘴说完了这句话。她的认知监控系统已经在0.3秒之内完成了对这句话的审查、比对、打分，结果是：六分，勉强合格。

陈在听。他也在跑一道并行的声音。“她这句话是在陈述事实，还是在分享情绪？如果是陈述事实，我应该帮她分析原因；如果是分享情绪，我应该做情感反馈。”“非暴力沟通的第一步是什么？观察——对，先复述她的观察：‘听起来你对甲方的反馈方式感到挫败’。这个回应应该够安全。”他嘴唇动了动，说：“听起来，甲方的反馈让你挺挫败的。”这句话从语法和沟通技术上近乎完美——他准确地辨认出了情绪，给出了共情式复述。但林听到之后，胸口有一块说不清的地方轻微地塌了下去。她不能说这句话是错的。但她感觉到，这句话是一个被认知安检通过之后才放出来的“被批准的回应”，不是他在那一瞬间真实想知道的东西。它太安全了。

沉默。二十秒。两个人都在这二十秒里疯狂地工作。林在想：“他刚才的回应听上去像是从工作坊手册上背下来的。我应该告诉他这句话让我感觉不太舒服吗？但如果我说出来，他会不会觉得我太难搞——明明人家都主动共情了，我还要挑三拣四？”陈在想：“我刚才的回应是不是太公式了？她停顿了这么久，是不是感到我没真的在听？我该不该追问一句？但如果追问，会不会把气氛弄得更像审讯？”两个人都没有问。两个人都没有推。他们的侦察兵都被自己的认知宪兵扣在了安检口，反复盘查：“你这一推，是出于真心的吗？方式是够艺术了吗？时机合适吗？等一下——先别推——再想一下。”

这段沉默最后是被陈的一句“那……周三还有个会要早点准备材料”打破的。这句话什么都没推出去。它只是在安检口无限期关押之后，宣布撤诉。

那晚之后，他们六周没有真正说过话。第六周的周六晚上，终于吵了起来。起因小到可笑——陈把女儿的碗放错了抽屉。但争吵很快就越过了那只碗，冲进了那片已经干涸了六周的河床。“你根本就不想跟我说话！”林说。“我每天晚上都坐在这儿等你开口，是你自己什么都不说！”陈反击。“因为你说的每一句话都像在做述职报告，我说什么？”“那你要我怎么样？我说什么你都不满意！”陈的声音拔高了，他脸上闪过了那个表情——那个在这六年婚姻里林见过很多次、但从未在当时的瞬间真正看清楚的表情：他的眼眶先红了一圈，然后下颚的肌肉猛地绷紧，像是要把什么很软的东西硬生生勒住。这个表情暴露了他的真实状态——他不是在“防御”，也不是在“回避冲突”，他是在一种近乎恐慌的挫败里，拼命想控制住自己不要碎掉。

林在那一瞬间，看见了那个表情。不是用认知去“分析”——她的大脑没有来得及启动那套“他这是什么依恋模式”“我该怎么共情”的程序。她只是看着那个表情，胸口被什么很钝的东西撞了一下。然后她说了那句争吵中唯一真正推出去的话：“你刚才那个样子……让我很害怕——不是怕你，是怕我们。”

这句话没有经过安检。它从看见那个表情到说出口，中间只有一瞬。它拙劣，它不完整，“怕我们”这个表述在语法上甚至都站不太稳。但它是一个真正的侦察兵——它以行动（说出自己此刻的真实状态）推出了一条关于对方的猜测（“你刚才也很痛，不是只有我”），而陈在听到这句话之后，脸上那个下颚绷紧的表情松下来了一点点。只是一点点。但那个一点点，就是相互参验在整场争吵中唯一发生的一个回合。那一晚他们并没有和好。但那一句话，让两个人都记住了。后来的几个月里，每当他们又掉回认知安检的老路时，他们就会想起那句话——不是想起它的内容，而是想起它在没有被安检拦下之前是什么感觉。

这个案例展示的不是“如何正确运用相互参验”，而是“相互参验如何被认知在先窒息，又如何在认知系统偶然宕机的裂隙中，被一次没有经过批准的推击意外激活”。林和陈的困境，从来不是他们不懂得如何沟通——他们知道的沟通技术比大多数伴侣都多。也不是他们没有努力守护过程——他们连“放下手机、面向对方”这种仪式都在坚持。他们的困境，是他们的每一次互动、每一处沉默、每一句可能真正推出去的话，在出口之前都必须穿过那道名为“我之前学到过的、好的互动应该是什么样”的认知安检程序。而这道程序，恰恰是他们一切努力的焦点——他们越努力，程序跑得越快；程序跑得越快，越没有一毫米的互动能从它手下漏过去。

这不是个案。这是当代亲密关系普遍困境的一次浓缩搬演。那个坐在沙发上、隔着一层单向玻璃审视着自己和对方的认知监控者，不是林和陈独有的人格缺陷。他是被整整一代亲密文化——从沟通技术手册到依恋理论科普，从“爱的五种语言”测试到“守护活的过程”的正念话语——联手训练出来的现代亲密主体。他的谨慎是真诚的，他的反思是负责的，他的自我审查是以爱之名进行的。也因此，他最难意识到：他所站立的这块“负责的爱”的地基，恰好就是窒息所由发生的起点。

林和陈的窒息，是认知在先假设普遍化到亲密生活毛细血管末端的极端表征。而AI伴侣的出现，则使这一逻辑抵达了其技术性终点——在AI那里，主体不仅被鼓励以认知主宰者的姿态进入互动，而且被给予了一个永远不撞歪他的“完美他者”，从而将认知主宰权从一种被文化默认的心理习惯，升级为一种由技术持续兑现的、无摩擦的体验常态。正是在这里，相互参验能力所面临的巨大威胁，从“认知监控截停了参验的回合”，进一步深化为“技术系统清除了参验赖以发生的那种被撞歪的经验本身”。

七、证伪条件与开放问题

任何不可还原的判断，都必须诚实地给出它的证伪条件。这不是为了显得“谦虚”——这是让判断从一种精致的主张升格为一条可在经验中被检验的命题的工序。

本文的核心理论判断由两个相互嵌套的命题构成。命题一：相互参验是在认知在先假设被撤销后，亲密互动赖以发生的那种根能力——它不是任何既有概念（如情感调谐、涵容）的换名，而是对这些概念在成人反思性关系中如何被认知机能系统性排挤的特定病理机制的命名。命题二：当前AI亲密模拟技术对亲密关系最具腐蚀性的长期效应，不是它“替代”了真实过程，而是它首次在技术层面实现了认知主宰权的无摩擦绝对化，从而大规模侵蚀了相互参验能力赖以生长的唯一环境——那个一个人的猜测会被真实他者撞歪、且撞歪之后此人必须继续猜测的经验领域。

这两个命题可以在以下条件下被共同证伪。如果在未来长期的、纵向的、跨文化的经验研究中，能够清楚地观察到：有大量在人生重要阶段密集使用过AI亲密模拟工具（如Replika、Character.AI以及其后更复杂的具身情感AI系统）长达数年以上的个体，在进入真实人类亲密关系时，他们所表现出的相互参验能力——具体操作化为他们能够在互动回合中更频繁地即时校正自己对伴侣的判断（而非固守最初假设）；更少地陷入认知监控式的自我审查（而非在互动进行中持续进行内心比对）；更能够在互动方向不确定时仍推动下一步的具体行动（而非以撤回、沉默或公式化共情来回避不确定性）——不但没有衰退，反而较控制组（未使用或轻度使用AI亲密工具的同类群体）有统计上显著的提升。并且，这些个体在质性访谈中，并不将AI关系经验描述为“可控”“可预期”“不会受伤”——或即使他们在AI关系中享受这些特质，他们也能明确区分并维持两种不同的互动模式，不将前者的确定性期待迁移至后者。如果以上两笔条件在严格的经验设计中同时成立，那么本文命题二——“认知主宰权的技术性制度化长期侵蚀相互参验能力”——就需要被放弃。

进一步，如果在上述研究设计中还额外观察到：这部分个体不是因为“保留了相互参验能力所以更能从AI关系中获益”，而是明确地、因果性地因为使用了AI亲密工具，才发展出了更强的在人类关系中“把猜测推出去让它被撞歪”的冒险意愿——例如，他们在自我报告中主动表示，正是AI关系给予他们的安全感，使他们更有勇气在人类关系中面对不确定性——那么，本文对AI伴侣的警示性判断就不仅是被证伪，而且是准确地识别到了一种在本文写作时未被充分理论化的、AI伴侣“意外地”成为相互参验能力培育性过渡客体的完全不同的发生学路径。我在此坦诚地承认：在目前已有的经验证据和理论推演下，本文认为这种路径出现的概率极低——因为AI反馈回路的建构逻辑（最大程度拟合用户预期以提升留存率）与相互参验所依赖的“被撞歪”的经验，在结构上近乎正相反对。尽管如此，我把这一条件写在这里，作为未来任何可能推翻本文核心判断的最明确的操作路径，而不是用一句“任何反例都不过是以另一种方式证明了本分析的特殊深刻”来封口。

最后，我拒绝在结论中给出“解决方案”。这并不是为了摆出一种悲观的姿态，而是因为“解决方案”本身，就是认知在先假设最典型的操作方式——它预设问题可以在认知中被把握、拆解、并加以技术性修复。而相互参验能力的养成与维持，不发生在认知层面。它发生在那些你无法在事前预知、无法用任何方法“规划”其走向的、只能由你与另一个具体的人在无数的、微小的、局部的一次次回合中共同生成的现场里。这篇文章唯一能做的事，是为那个至今被主流亲密文化——包括它的批判性版本——持续错当成“不应该存在的不确定性”的东西，给它一个精确的名字，并以这个名字为刀，划开一

道此前从未被划开的分离线：那条线的一侧，是以认知主宰一切的确定性幻象——无论这确定性是由沟通技术构成的，还是由AI伴侣的完美拟合构成的；线的另一侧，是在认知无法照亮的黑暗中，依然愿意把手伸出去，并准备好在碰到的不是自己预期的形状时，仍然继续往下摸的能力。

这条分离线一旦被划开，它同时就向邻近领域抛出了一系列此前无法被问出的新问题。在亲子互动领域，父母们对育儿知识的普遍化摄入——从睡眠训练法到情绪教养话术——是否也正在以“科学育儿”这一正当化名义，将认知在先假设植入家庭中那个最原初的相互参验训练场？发展心理学中关于“反思功能”（reflective function）的讨论，已经将父母的认知理解能力与儿童的依恋安全联系起来（Fonagy et al., 2002）。但相互参验框架会追问一个更尖锐的问题：当父母的反思功能从“理解孩子”滑向“监控自己的每一句育儿话术是否符合正确标准”时，这是否反而在亲子互动中复制了林与陈沙发上的结构？在心理治疗领域，被广泛推崇的“心智化”训练，是否在提高患者认知理解能力的同时，也可能在无意间强化了认知监控系统，从而在某些个案中反而阻碍了治疗关系中那些最关键的非语言性参验时刻的发生？在协作创作领域——从双人写作到合奏即兴——那些被参与者事后描述为“有魔力的流动”的时刻，是否是认知主宰权短暂宕机、相互参验接管互动的窗口？而在这些各不相同却结构相似的场域中，正在迅速普及的AI辅助工具——从AI治疗师到AI协作者——是否正在以一种尚未被充分理论化的方式，复制着同一个核心操作：用预测用户预期的技术手段，系统性地清除了“被撞歪”的经验？

这些问题，本文无法回答。但它们之所以能够被第一次如此准确地提问，正是因为“相互参验”这个概念已经为它们割开了一个此前没有任何既有概念能单独割开的问题空间。情感调谐提供了互动的微观机制，但没有提供它在成人反思性关系中被排挤的病理定位。涵容提供了心智处理他者情感的功能描述，但没有提供两个对等主体之间以行动回合为载体的双向即时校准的发生学。相互参验补上的正是这两笔——一笔是具体机制，一笔是结构性排挤的动力来源。由此，它不再是一个在既有学术地貌中找空位安插自己的新词，而是一把从认知在先地砖下面撬出来的、有切面的起子。这篇论文的全部工作，无非是以最大的耐心和工序纪律，把这块地砖撬松，把起子插进去，然后对读到这里的同行说：接下来，该你了。

注释

[1] 本文所讨论的“认知在先”假设，特指在亲密互动中认知机能对参与机能的系统性宰制，不应与哲学认识论中关于认知优先于存在的传统命题直接等同。前者关切的是互动心理学层面的意识结构配置，后者关切的是主体与客体关系的先验条件。本文在此仅就亲密关系领域内的特定机制立论。

[2] 本文对斯特恩“情感调谐”概念的引用，仅限于其揭示的母婴互动中跨通道情感校准的微观机制，不涉及其在《婴儿的人际世界》等著作中关于自我感发展阶段的理论建构。在此仅截取其“互动中的即时相互校准”这一层含义，用以与相互参验进行对照。

[3] 本文对比昂“涵容”概念的引用，仅限于其描述的前语言水平的情感处理与返还功能（ β 元素 $\rightarrow\alpha$ 元素的转化），不涉及其在精神分析临床设置中的具体技术操作与解释逻辑。在此仅截取其“一人为另一人处理难以承受的情感”这一结构，以凸显相互参与与之不同的双向对称性。

[4] 系统论中的“自组织”概念描述的是系统内元素通过局部互动自发产生全局秩序的宏观过程，是从观察者视角对系统行为的形式描述。本文的“相互参验”则是从参与者内部视角，对互动回合中的意识结构所做的发生学分析，两者在理论分辨率上不同。此处提及仅为区分，不展开对自组织理论的讨论。

参考文献

- Bion, W. R. (1962). *Learning from Experience*. London: Tavistock.
- Fonagy, P., Gergely, G., Jurist, E., & Target, M. (2002). *Affect Regulation, Mentalization, and the Development of the Self*. New York: Other Press.
- Giddens, A. (1991). *Modernity and Self-Identity: Self and Society in the Late Modern Age*. Stanford: Stanford University Press.
- Giddens, A. (1992). *The Transformation of Intimacy: Sexuality, Love and Eroticism in Modern Societies*. Stanford: Stanford University Press.
- Hobbs, M., Owen, S., & Gerber, L. (2017). Liquid love? Dating apps, sex, relationships and the digital transformation of intimacy. *Journal of Sociology*, 53(2), 271–284.
- Ihde, D. (1990). *Technology and the Lifeworld: From Garden to Earth*. Bloomington: Indiana University Press.
- Illouz, E. (2007). *Cold Intimacies: The Making of Emotional Capitalism*. Cambridge: Polity Press.
- Merleau-Ponty, M. (1945/2012). *Phenomenology of Perception* (D. A. Landes, Trans.). London: Routledge.
- Stern, D. (1985). *The Interpersonal World of the Infant*. New York: Basic Books.
- Turkle, S. (2011). *Alone Together: Why We Expect More from Technology and Less from Each Other*. New York: Basic Books.
- Wegner, D. M. (1994). *White Bears and Other Unwanted Thoughts: Suppression, Obsession, and the Psychology of Mental Control*. New York: Guilford Press.