

论文 · SDE 发生学

何处正要发生

裂缝的定位、三大亏缺与九步发生机制

把"先造世界再行动"从一句口号，变成有合格判据的工序

王德生 著

德麦国际 · Demai International Pte. Ltd. (新加坡)

2026 年 7 月

本文据《SDE 智能体导论——从语言模型到发生世界模型的智能体革命》
(德麦国际出版, 2026, ISBN 978-1-970820-62-1) 第八章与第九章改写成文

摘要

关于智能系统的讨论几乎都从"如何更好地回答"开始，跳过了一个更靠前的问题：在生成任何东西之前，凭什么知道该往哪里使劲。本文论证，新结构的发生有其着力点——只有在旧结构正撑不住、新结构正要破土的位置上，发生才既可能又划算；本文称这个位置为裂缝，并把它从隐喻变为可操作的对象。文章分三层展开。定位层：给出裂缝的四个现象信号（说不清、对不上、绕不过、收不拢），并明确它们只是扫描仪而非诊断书；继而给出两项判断——深度的四个层级（术语、方法、理论、本体）与时机的五个阶段（微裂、扩展、显裂、竞争填补、新壳），据此判断一道缝值不值得下手、现在下手划不划算。诊断层：给出三大亏缺作为裂缝的机制级定义——关系建立亏缺、变化动力亏缺、路径展开亏缺，分别对应纠缠、差异、结构三维；并说明这三类的穷尽性并非经验归纳，而是从三维坐标构造性地继承而来，同时区分"维度正交"与"成因可叠加"这两件常被混淆的事。工序层：给出承载定位与生成的九步机制，逐步标出合格判据与常见失败模式，并用一个完整走查演示它如何把一场收不拢的争论培育成一个可检验的新框架。最后论证该机制不能被压缩为流水线：使它区别于执行链的是回流（下游发现问题可退回上游重做）与回写（产出改变下一轮的初始条件）。文章末尾诚实标注了演示案例与既有研究的近邻关系，以及三大亏缺穷尽性所依赖的上游承诺。

ABSTRACT

Discussions of intelligent systems almost always begin with how to answer better, skipping a prior question: before generating anything, how does one know where to apply effort? This paper argues that the emergence of new structure has a locus — it is both possible and economical only where an existing structure is failing and a new one is pressing to appear. That locus is here called a crack, and the paper turns it from metaphor into an operable object across three layers. Location: four phenomenal signals mark a crack (the inarticulable, the irreconcilable, the unavoidable, the unresolvable), which serve as a scanner rather than a diagnosis; two further judgments follow — four depths (terminological, methodological, theoretical, ontological) and five stages (micro-fracture, expansion, open fracture, competitive filling, new shell) — determining whether a crack is worth entering and whether now is the moment. Diagnosis: three deficits give the crack a mechanism-level definition — of relation, of driving difference, and of unfolding path — corresponding to the entanglement, difference, and structure dimensions; their exhaustiveness is shown to be inherited constructively from the three-dimensional coordinate rather than induced from cases, while distinguishing orthogonality of dimensions from superposability of causes. Procedure: a nine-step mechanism is specified, each step accompanied by its pass criterion and characteristic failure, and demonstrated end to end on an unresolved disciplinary controversy. Finally the paper argues the mechanism cannot be compressed into a pipeline: what distinguishes it from an execution chain is backflow and write-back. Limits, and the proximity of the worked example to existing findings, are stated openly.

关键词 裂缝 · 问题发现 · 三大亏缺 · 发生机制 · 合格判据 · 病态结构问题 · 智能体设计 · 知识创造

Keywords crack · problem finding · three deficits · genesis mechanism · pass criteria · ill-structured problems · agent design · knowledge creation

○ · 引言：一个被跳过的前置问题

关于智能系统的讨论，几乎都从同一处开始：如何回答得更准、更全、更快。这些讨论跳过了一个更靠前的问题——**在生成任何东西之前，凭什么知道该往哪里使劲？**

一般的任务型系统不问这个问题。它接到任务就执行，接到提问就回答，着力点由指令直接给定：用户说做什么，它就做什么。但一个以“发生新结构”为目标的系统不能这样。发生不是对着整片虚空均匀地使劲，发生有它的着力点：只有在旧结构正撑不住、新结构正要破土的那个精确位置上，发生才既可能又划算。对着一片早已稳固、毫无松动的旧结构硬造新东西，造出来的多半是没人需要的赘物，或者干脆是漂亮的空话。

所以第一步不是生成，而是定位。这一判断并非本文独有。Getzels 与 Csikszentmihalyi 在 1970 年代对艺术学生的研究中，把“问题发现”与“问题解决”分开，并发现前者的行为特征比后者更能预测长期的创作成就；Dewey 更早就把探究的起点安放在“不确定的情境”上——探究不是从一个已经成形的问題开始，而是从一片尚未成形的困惑开始，把它转化为一个可作业的问题，本身就是探究最要紧的一步。认知科学的实证也指向同一处：Chi、Glaser 与 Rees 关于专家与新手的经典比较表明，专家的优势很大程度上不在解题步骤，而在**他们对问题的表征方式不同**——他们先把问题重新组织了一遍。

本文要做的，是把这一被广泛承认、却始终停在洞见层面的判断，落成可操作的工序：把“裂缝”从隐喻变为可扫描、可分层、可诊断的对象，再给出承载它的完整机制。文章分三层——第一至三节是定位（信号、深度、时机），第四节是诊断（三大亏缺），第五至八节是工序（九步、走查、为什么不能压成流水线、与坐标系的对应）。第九节收束并交代边界。

一 · 裂缝不是瑕疵，而是产道

1.1 一个必须先扭转的直觉

在“发现”范式里，旧理论中说不清、对不上、绕不过的地方，被当作瑕疵——是这个理论还不够完善的证据，是要被修补、被抹平、被绕开的麻烦。一个好的研究者、一个好的系统，应当尽量让这些地方消失。

本文的判断相反：**这些地方不是瑕疵，而是产道**——是新结构即将诞生的位置。旧结构在某处撑不住，恰恰意味着那里有一个旧结构容纳不下的新结构正在挤压。达尔文面对的不是物种学说的瑕疵，而是化石序列对不上、地理分布绕不过、家养变异收不拢这几道产道，他顺着它们接生出了演化；弗洛伊德面对的不是理性心理学的瑕疵，而是口误说不清、梦说不清、症状收不拢这几道产道，他顺着它们接生出了无意识。他们的过人之处不在于抹平了裂缝，而在于没有抹平它、反而顺着它接生。

Kuhn 在科学史一侧给出了同构的观察：范式转换起于反常——那些在既有范式内反复出现却安置不下的现象。本文与之同向，但要往前推一步：Kuhn 描述的是范式为何转换，是一个事后的历史判

断；本文要给的是在事情发生之前、当场可用的定位工序——如何在此刻的问题域里，把那道正在裂开的缝找出来、量出深浅、诊断出它缺什么。

1.2 抹平机与接生机

这个扭转对系统设计是决定性的。如果把裂缝当瑕疵，系统会被设计成一台抹平机——把问题里所有不顺、所有矛盾、所有说不清都圆滑处理掉，给出一个工整、全面、毫无棱角的回答。这恰恰是生成模型默认最擅长的姿态：落在已有判断的加权平均上。

举一个后文将贯穿使用的例子。面对"为什么某种训练方法在某些技能上效果显著、在另一些上几乎无效，而现有理论始终收不拢"这个问题，抹平机会写："训练效果取决于多种因素，包括目标、反馈、强度等"——四平八稳，把那道最要命的缝（为什么随技能而异）抹得干干净净。

如果把裂缝当产道，系统会被设计成一台接生机——主动去寻找那道最说不清、最对不上的缝，把全部力气压上去。抹平机生产平庸的正确，接生机生产新的结构。区别不在能力强弱，而在着力点的选择；一个不会找裂缝的系统，无论生成能力多强，都是在错误的位置上浪费它的强大。

二 · 四个现象信号：裂缝如何暴露自己

裂缝既然要被定位，就得有可辨认的标记。在现象层面，它通过四个信号暴露自己。

信号	现象	它通常意味着
说不清	人人都在用、却谁也定义不清的概念；一追问"它到底是什么"就各说各话	缺一个尚未显露的结构——现有语言还罩不住它
对不上	两套各自成立的说法、两组各自可靠的证据，摆在一起却相互矛盾	两个被分别处理、实则纠缠的结构没有被打通
绕不过	一条所有人都默认要走、却谁走都走不通的路；反复被绕开却绕不干净的难点	缺一条以前不存在的路径
收不拢	旷日持久、各方都有道理却始终收不拢的争论；怎么也归不到一起的现象	缺一个能把分歧统摄起来的更高结构

把四把尺子对准那个贯穿案例：最刺眼的信号是**收不拢**——主张训练量近乎决定一切的一派与强调天赋、情境、运气的一派各有可靠证据，争论多年不收敛。次刺眼的是**对不上**——同一套方法在棋类、乐器上效果显著，在即兴创作、临床判断上几乎无效，两组可靠结果摆在一起对不上。两个信号叠加，把候选裂缝从泛泛的"这个方法好不好"，精确到了"有效性为何随技能而异"这个点上。定位的第一份收益就在这里：它把问题换了一个提法。

这四个信号是"测体温"的工具——它们告诉系统"这里可能有缝"，是最值得下手的候选位置。一个善于扫描四信号的系统，能在使用者自己都还没意识到的地方，先嗅出正要发生的位置；而这恰恰是被动等待指令的系统永远做不到的事。

但必须立刻划清它们的边界：**四信号只是裂缝的现象，不是裂缝的本质**。发烧告诉你"这里有病"，却不告诉你"是什么病、该怎么治"。同样，四信号告诉系统"这里可能有缝"，却不告诉它"这道缝在机制上是什么、缺的是哪一样、该补什么"。从"这里好像有缝"精确到"这是哪一类缝"，需要一套机制层面的定义——那是第四节的工作。四信号是扫描仪，三大亏缺是诊断书。

三 · 两个判断：这道缝有多深，现在动手划不划算

3.1 深度：裂缝的四个层级

同样标记着"说不清、对不上"，缝与缝的深浅天差地别。扫到一道缝之后，第一项判断是量它的深度。

层级	缝在哪里	顺着它能接生出什么
术语缝	某个词没定义清楚	换个说法、补个定义就能补上；产出很小
方法缝	现有方法在某类问题上失效	一套新工序
理论缝	现有框架在某个边界上系统性失效	一套新理论
本体裂缝	一个被整个领域当作不证自明的房角石 其实站不住	最难补，但产出最大——分量最重的创新无一例外抓住了这一层

这给出一条实战纪律：**抓到一道缝时，别满足于它是不是术语缝就动手，要往下追问"这道缝下面，是不是还连着一道更深的缝"**。

贯穿案例正好演示这一层层下探。表层看是一道方法缝——某个训练方法在某些技能上失效。但往下追问会触到一道更深的本体裂缝：整个争论默认了"技能"是一个统一的、可用同一套训练逻辑对待的东西，而这个房角石其实站不住。不同技能的内在结构根本不同，把它们当成同一类来争论"这个方法有没有用"，问题本身就提错了。能追到这一层，就站在了产出最大的位置上。

3.2 时机：裂缝的五个阶段

第二项判断是时机。裂缝不是静止的，它沿一条曲线生长：**微裂期**（极少数人隐隐觉得"这里不对"，说不清，旧框架表面依然牢固）→ **扩展期**（越来越多人在同一处撞上同样的别扭，私下议论增多）→ **显裂期**（这道缝被公开承认，旧框架权威松动）→ **竞争填补期**（各种新方案蜂拥而至争相填补）→ **新壳形成期**（某个方案胜出、凝固成新框架，并在别处埋下新的微裂，循环重启）。

认清阶段，就有了下手时机的判断力。**最有价值、也最考验眼力的，是在微裂期下手**——那时几乎无人察觉，抢先抓住就握住了一片几乎无人竞争的产道。到了竞争填补期，缝早已人尽皆知，涌进来抢着填补的人挤成一团，此时下手事倍功半。

这也解释了一个更大的判断：强大的新工具几乎总会制造新裂缝——它让大量旧框架同时进入微裂期。而当前的大语言模型正是这个时代最大的一台造缝机：它让无数领域的旧框架几乎同时齐齐裂开崭新的、还无人填补的微缝——"什么算原创""什么算理解""什么算专业能力""教育在评什么"，这些问题在两三年前都还不是问题。站在这台造缝机刚启动的当口，帮人在微裂期就精准抓住这些缝，是本文所述方法最现实的用武之地。

四 · 三大亏缺：从"这里有缝"到"这里缺什么"

4.1 裂缝的机制级定义

四信号让人在现象层"闻"到缝，深度与时机判断值不值得、该不该现在下手。但要真正补对，还差一步：**精确说出这道缝在机制上缺的究竟是什么。**

亏缺类型	缝的根源	补法的方向
关系建立亏缺 [E]	某个本该被建立的纠缠没有被建立——某要素被孤立起来，当成与其他要素无关的外生常量	把被孤立的要素重新接回它真实纠缠的网络
变化动力亏缺 [D]	只有静态的结构描述，缺一个解释"为什么会这样变、被什么驱动"的动力	找出驱动这个结构发生与维持的差异序列
路径展开亏缺 [S]	土壤够厚、动力也在，但没有一条工序把隐而未现的结构一步步接生出来，它停在"隐约觉得有这么回事"	铺设那条展开路径，让弱结构显露、稳定、可表达

把三把尺子用在贯穿案例上：那道"为何随技能而异"的缝，机制上主要是**关系建立亏缺**——整个争论把"技能"当成了一个孤立的、统一的外生常量，没有把"训练的有效性"与"技能所处环境的结构稳定性、反馈的即时性、情境的可重复性"这些本该深度纠缠的要素建立起关系。争论双方都在问"这个方法对技能有没有用"，却没人把"技能"这个被孤立的整体，拆回它与环境结构的真实纠缠里。诊断到这一步，补法就清楚了：建立那个被切断的纠缠——而不是继续在"有用/没用"之间加权。

诊断决定补法，这是本节全部的分量所在。缺动力却去补关系，等于药不对症：你把要素之间的关系画得再密，也解释不了它为什么会变。

4.2 为什么恰好是三类

一个必须正面回答的理论问题：为什么恰好是三类，不是两类、不是五类？不答清楚，三大亏缺就只是一份经验清单。

回答是：三类不是数出来的，而是推出来的。一道裂缝，本质上是某次发生卡住了；而一次发生只由三样东西构成——结构、差异、纠缠。发生卡住，逻辑上就只可能卡在这三处之一。**三大亏缺的穷尽性，是从三维的穷尽性继承来的**，不是"目前只见过三类"的归纳。

这条论证有一个诚实的代价，必须写明：它把穷尽性问题上推给了一个更上游的承诺——发生是否恰好有三个维度。要挑战"恰好三类"，得去挑战上游的三维，而不是在亏缺这一层数第四类。而"有没有第四维"本身是一个开放问题，本文不假装已经关闭它。把一个看似可疑的"恰好三类"，锚定到一个明确的、可被追究的上游承诺上——这本身就是本文所主张的方法对自己的一次应用。

4.3 正交与叠加：两件常被混淆的事

第二个理论问题：三类互斥吗，还是会重叠？回答分两面，而这两面常被混为一谈。

作为维度，三者正交、不可互相还原。"关系没建立""动力缺失""路径缺失"是三个独立方向，补一维不自动补另一维。这种正交性正是诊断有用的前提：正因为三维不可互相还原，"诊断为哪一维"才携带真正的信息、才决定补法。

作为一道具体裂缝的成因，三者可以叠加共存。一道真实的深缝，常常同时缺关系、缺动力、缺路径。

两面并不矛盾：维度的正交，说的是三个方向互不重叠；成因的叠加，说的是一道缝可以在多个方向上同时亏空。好比一个人可以同时缺三种维生素，但"缺 A""缺 B""缺 C"仍是三个不可互相替代的诊断。把"维度正交"误读成"一道缝只能属一类"，诊断会变僵硬；把"成因可叠加"误读成"三类其实是一类"，诊断会失去分辨力。正确的用法是：用三把正交的尺子分别量一道缝在每一维上各亏空多少，再针对亏空最重的那一维或几维定向地补。

这里还有一个跨学科的镜像，可以让三分显得不再像私设。在系统科学里，一个动态系统的失效也常被归到三个不可互相还原的方向：结构性失效（组件间的连接出了问题）、动力学失效（驱动演化的机制出了问题）、可达性失效（系统到不了某个本应可达的状态）。这三者与本文的三类亏缺几乎一一对应。这不是巧合：**任何"系统为何卡住"的追问，最终都会落到连接、驱动、可达这三个不可再省的面上。**一个在别的学科里被分头发现、各自命名的三分，在这里得到了一个统一的根据。

五·九步发生机制：每一步的合格判据与常见失败

定位与诊断解决了"往哪里使劲、该补什么"，还没有回答"具体怎么走"。下面这套九步机制，就是承载它的工序。

先立一句总纲，再展开：先造一个世界，再在这个世界里行动，并让行动的产物回过头来改造这个世界。口号与机制的差别在于：口号告诉你方向，机制告诉你每一步做什么、做到什么程度算合格、做不到会出什么症状。所以下表的三列里，**后两列才是这套机制真正的分量**——没有合格判据的流程，只是一份好看的清单。

步	动作	合格判据	常见失败
1	问题输入：把问题当发生入口	能识别问题背后的纠缠，而不只是字面诉求	急着应答 ——跳过对问题性质的判断，后面全部建在错位的地基上
2	三界画像：建立临时世界	理念、现实、自我三面齐备且不偏废	偏建 ——只建现实（盲目）、只建理念（漂浮）、只建自我（主观）
3	世界猜想：允许它先不确定	猜想有足够张力与可能性，未过早收敛	过早确定 （第一个稳妥结构当定论）／ 放任发散 （一堆互不收敛的可能）
4	参数确定化：把猜想拉向现实	关键未定参数被文献、数据、工具、实验、审稿收束	跳过确定化直接输出 ——语言流畅却经不起追问
5	纠缠识别：找到真正锁死的地方	识别出的纠缠能解释"为什么旧结构在这里失效"	表面归因 ——归到最显眼但最浅的因素上
6	差异生成：制造有方向的差异	差异能撬动上一步识别的纠缠点	伪差异 ——只换措辞、包装、类比，没触动纠缠
7	候选显露：生成多个新结构	候选之间有实质差异，各自指向不同方向	单点迷恋 ——过早爱上第一个候选，失去比较空间
8	三重审查：逻辑、现实、价值	同时通过三审，或据审查意见修正再审	三审走形式 ——只看语言是否流畅
9	输出与回写：闭合循环	产出既落地又回写，使系统进入下一轮	无回环 ——一次产出之后再无后续，系统不进化

三步值得单独说明，因为它们承担了这套工序最大的分辨力。

第三步的"允许它先不确定"是反直觉的纪律。常见的做法是让系统尽快给出一个可靠答案，而这里要求它先容纳一个带着未定参数、尚不能落地的猜想。理由在于：过早收敛会在差异还没被制造出来之前就关闭差异空间——那个稳妥的第一答案，往往正是已有判断的加权平均，也就是第一节说的抹平。

第四步是全程分量最重的一步，也是这套工序区别于一般生成的关键。它要做的不是"补充资料"，而是把猜想里每一个空着的位置逐一钉到现实上：概念是否准确、脉络是否可靠、数据是否真实、工具能否算、边界在哪里。少了这一步，前面的张力会直接坍缩为漂亮的空话。

第七步的"要多"同样反直觉。它要求在有了一个不错的候选之后，继续生成其他候选。设计研究中的"设计固着"实验提供了一个可对读的旁证：向被试展示一个示例方案，会显著收窄他们后续方案的空间，且被试往往并未察觉自己被限住了。单点迷恋不是懒惰，它是一种不自知的收窄；对治它的唯一办法是把"多候选"写进工序，而不是指望当事人自觉。

六 · 完整走查：一个收不拢的争论如何被走成新框架

把九步放回贯穿全文的案例，从头到尾走一遍，机制就从抽象变得可操作。

第一步，问题输入。研究者问“为什么这套训练方法在某些技能上几乎无效”。系统不把它当检索键，而识别出它背后是一场收不拢的理论争议——这是一个发生入口，不是一道待答的题。

第二步，三界画像。理念一面：这属于技能习得理论，已有强调训练量的一派与强调天赋、情境的一派；现实一面：有跨领域的训练研究、可测的技能维度、可查的文献；自我一面：研究者想立一个能收拢争论的新框架，而不是再写一篇综述。

第三步，世界猜想。生成一个带张力的初始猜想——“也许差别在于技能所处环境的某种结构”。此时“某种结构”还是一个空壳，按纪律允许它先不确定。

第四步，参数确定化。查文献，确认已有研究是否零散提过类似区分、各派的真实边界在哪；把“技能环境”操作化为可测维度（环境稳定度、反馈即时性、情境可重复性）；先让审稿攻一轮：“临床判断里也有可重复的部分，你的区分扛得住吗？”空壳被逐步填实。

第五步，纠缠识别。在确定化后的画像里识别出真正的纠缠点：训练的有效性与“技能所处环境的结构稳定度”深度纠缠——环境越稳定、规则越固定（棋类、乐器），训练越有效；环境越多变、情境越不可重复（即兴、临床判断），越无效。这个纠缠点能解释“为什么旧结构在这里失效”——旧争论恰恰把技能当统一整体，切断了这个纠缠。

第六步，差异生成。不再问“这个方法有没有用”，而是换一个视角：把技能按环境结构稳定度分层，问“有效性如何随这个分层而变”。

第七步，候选显露。显露多个候选：一个把技能分为高/低结构稳定两层的简版框架；一个引入连续“结构稳定度”维度并给出有效性预测曲线的进阶框架；一个再叠加“反馈即时性”的双维框架。

第八步，三重审查。逻辑审查它们是否自洽；现实审查哪个能被现有数据检验、能给出可证伪的新预测；价值审查哪个真能收拢争论、对领域有用。假设双维框架通过。

第九步，输出与回写。把它改写成技能习得领域的本地语言，输出为一个可投稿的框架；这个新结构回写进画像——它改变了“技能”这个概念本身，激发出新的追问（这个分层能否解释专家直觉、能否指导训练设计），为下一轮准备了起点。

6.1 关于这个案例的一句诚实交代

必须写明一件事，否则这次演示就会误导读者：**上述走查得出的结论方向，并不是本文的新发现。**

它与既有研究高度同构。Kahneman 与 Klein 在 2009 年那篇著名的“求同”论文中给出的结论是：可靠的专家直觉需要两个条件——一个具有足够效度、规律性可学的环境，以及通过长期练习获得这些规律的机会；在低效度环境里，长期经验并不带来可靠判断。这与走查中“环境结构稳定度决定训练有效性”几乎是同一个判断的两种说法。Macnamara 等人 2014 年的元分析则提供了量化的一侧：刻意

练习所能解释的表现差异在不同领域间差别极大——在规则固定、反馈明确的领域较高，在专业职业等情境多变的领域极低。

所以本文用它做演示，不主张其为新结构。演示的是工序，不是结论——正因为这个结论已被独立地确证过，它才是一个好的走查样本：读者可以对照文献，逐步检查这条工序有没有把人带到它该到的地方。**一套方法论最需要的样本，不是它自己产出的孤例，而是一个终点已知的路段。**若换一个终点未知的问题来演示，读者将无从判断走查的每一步是真有效，还是叙事上显得有效——那正是本文所反对的那种自证。

七 · 为什么它不能被压成流水线：回流与回写

九步很容易被误读成一条九段流水线：问题从第一步进，依次穿过九段，产出从第九步出。这个误读很自然，却会把发生机制阉割成一台执行机。使它根本区别于流水线的，是两样东西。

7.1 回流：下游发现问题，退回上游重做

流水线只能前进，每一段做完交给下一段，不回头。这套机制不是这样：**任何一步发现前面出了问题，都可以、且应当退回去重做。**参数确定化若发现猜想的某个关键参数根本锚不住，要退回第三步重新生成猜想，甚至退回第二步重建画像；三重审查若否决了全部候选，要退回第六步重新制造差异。

一条会因为下游的发现而退回上游重做的链条，本质上不是流水线，而是一个带反馈的搜索过程——它在前进与回退中收敛。一个只会从第一步走到第九步、从不回流的系统，恰恰是把发生降格成了执行：它带着早期的错误一路走到底，产出一个建立在错位地基上的成品。

这一点在设计与规划研究里早有对应的论述。Rittel 与 Webber 指出，社会规划面对的问题无法被一次性完整地定义——你对问题的理解，要到你尝试解决它时才充分显现；Schön 则把专业实践刻画为“问题设定”与“行动中的反思”的交替，从业者在行动中不断重新框定他所面对的问题。**两者说的都是同一件事：问题的定义与问题的解决无法被切成前后两段。**回流正是把这一事实写进了工序。

7.2 回写：产出改变下一轮的初始条件

这同样更要命。流水线的产出是终点，下线即终止，系统状态不因这次产出而改变。这里的第九步不是终点，而是回写的起点：新结构一旦显露，会反过来改变第二步建立的画像、第五步识别的纠缠——新结构重置了纠缠场，新的纠缠场又激发下一轮的差异，下一轮的差异继续显露新的结构。

这意味着九步不是一次走完就结束的链，而是一个会自我重启的循环：**每一次产出都改写了下一次发生的初始条件。**一个没有回写的系统，做完一次就停在那里，它产出的新结构没有改变任何东西，系统不进化；有回写的系统，每一次产出都让下一轮站在一个被重置过的、更高的起点上。

这里要与工业界熟悉的循环划清界限。PDCA 一类的改进循环同样是闭环、同样强调反馈，但它的目标是让既定过程运行得更稳、更好——循环收敛于对标准的逼近。本文所说的回写不是这个方向：它

要的不是把同一个过程做得更好，而是**让下一轮面对的问题本身发生改变**。两者的判据因此不同：改进循环的成功判据是偏差缩小，发生循环的成功判据是初始条件被改写。

7.3 "能不能压成更少的几步"是个问错方向的问题

把回流与回写合起来看，这套机制的真实形态就清楚了：它不是一条直线，而是一个带内部反馈、又整体循环回写的结构。九这个数，标的不是九道必须顺次通过的工位，而是一次完整发生应当覆盖的九个不可省的环节——它们有逻辑上的先后（没有画像无处发生，没有确定化的猜想不能审查），但在实际运行中彼此回流、整体回写。

所以"能不能把它压成更少的几步、做成一条更高效的流水线"这个问题本身就问错了方向。**压缩工位是流水线思维的优化；这里要的不是更短的流水线，而是更充分的发生**。少掉任何一个环节，省掉的不是一道工序的耗时，而是发生的一个不可省的面——省掉确定化，换来的是漂亮的空话；省掉回写，换来的是停滞。理解了这一点，才不会把一个发生机制，误建成一台跑得很快、却什么也没真正发生的执行机。

八 · 与坐标系的对应：机制不是凭空的清单

九步不是设计出来的清单，它与本文所依据的坐标严丝合缝——这一节把对应关系写明，是为了让读者能追究它，而不是让它显得体面。

步骤	在坐标系里做的事
第一至第四步	建造与确定化 纠缠 ——把发生所扎根的土壤铺厚、压实
第五步	定位纠缠中被锁死的那一处矛盾
第六步	制造 差异 ，撬动那处矛盾
第七步	显露 结构 候选
第八步	审查结构的成立度
第九步	回写 ——新结构反过来重置纠缠场，循环重启

整条机制因此就是一次完整的"在纠缠中，经差异，成结构，再回写纠缠"的循环，外加参数确定化这一现实化装置。这个对应关系也解释了为什么恰好是这几个环节：它们不是从实践里数出来的九件事，而是三个维度加上一个现实化装置、再加上闭环，展开之后的必然分节。同样地，这条论证把"什么是九步"的追究上推到了坐标本身——要挑战九步的划分，得去挑战三维坐标，这与第 4.2 节对三大亏缺穷尽性的处理是同一个动作。

最后把定位与工序接起来。第四节的三大亏缺，正落在第五步"纠缠识别"上：这一步不再是含混地"感受一下问题的纠缠"，而是一套可操作的工序——**先用四信号扫描问题域定位候选裂缝，再用四层深度**

与五阶段时机判断值不值得、该不该现在下手，最后用三大亏缺把选中的缝精确诊断为缺关系、缺动力、还是缺路径。定位是工序的第一段，不是它之外的准备活动。

九 · 结论与边界

把全文收成三句话。

其一，第一步不是生成，而是定位。发生有着力点：只有在旧结构正撑不住的位置上，新结构才既可能又划算。裂缝不是待抹平的瑕疵，而是产道；把系统设计成抹平机还是接生机，决定了它生产平庸的正确还是新的结构。四个现象信号让裂缝可被扫描，四层深度与五阶段时机让它可被估值——其中最有价值的是在微裂期抓住那道还无人察觉的缝，而当下这台造缝机正让无数领域同时进入微裂期。

其二，诊断决定补法。三大亏缺把"这里有缝"精确到"这里缺什么"：缺关系、缺动力、还是缺路径。它的穷尽性不是经验归纳，而是从三维坐标构造性地继承；代价是把可争议之处上推给了一个明确的上游承诺，这一点本文写在明处而非藏起来。诊断错了，补得再用力也是药不对症。

其三，工序的分量在判据，机制的分量在回环。九步的价值不在于罗列了九个动作，而在于每一步都有合格判据与常见失败——没有判据的流程只是好看的清单。而使它区别于执行链的是回流与回写：一条会退回上游重做、且其产出会改写下一轮初始条件的链，本质上不是流水线。因此"能不能压成更少的几步"问错了方向：要的不是更短的流水线，而是更充分的发生。

边界同样要写在明处。其一，三大亏缺的穷尽性依赖于"发生恰有三个维度"这一上游承诺；有没有第四维，是一个开放问题，本文不假装已经关闭它。其二，九步的合格判据中，有几条（如"猜想有足够张力""候选有实质差异"）目前仍需人的当场判断，尚未被操作化到可以机械核验的程度——本文能主张它们比"答得好"更硬，不能主张它们已经足够硬。其三，本文的走查演示选用了有一个终点已被独立确证的案例，这保证了可对照性，但也意味着它没有证明这套工序能在终点未知处同样奏效；那需要的不是更漂亮的演示，而是若干次公开的、可被外部复核的实际运行记录。

把这三条写出来，是这套方法对它自己的应用：**一个不敢标出自己会在哪里失手的工序，正是它自己第八步要打回的那种东西。**

本文据《SDE 智能体导论——从语言模型到发生世界模型的智能体革命》（王德生 著，德麦国际出版，2026，ISBN 978-1-970820-62-1）第八章《裂缝与三大亏缺：智能体如何定位"何处正要发生"》与第九章《九步发生机制：从问题入口到新 SDE 输出》改写成文，并补入文献定位、论证编排与边界声明。同组论文见《没有统一世界模型》与《伪发生》；原书导读与全文阅读见 专著页。

参考文献 · References

下列文献为文中所涉思想的原典与代表性研究，供读者对读与查证；本文观点不代表所列作者立场。

1. Jacob W. Getzels & Mihaly Csikszentmihalyi, *The Creative Vision: A Longitudinal Study of Problem Finding in Art*, New York: Wiley, 1976.
2. John Dewey, *Logic: The Theory of Inquiry*, New York: Henry Holt, 1938（不确定情境与探究的起点）。
3. Donald A. Schön, *The Reflective Practitioner: How Professionals Think in Action*, New York: Basic Books, 1983；中译《反映的实践者》，夏林清译，教育科学出版社，2007。
4. Horst W. J. Rittel & Melvin M. Webber, "Dilemmas in a General Theory of Planning", *Policy Sciences*, Vol. 4, No. 2 (1973)（棘手问题）。
5. Herbert A. Simon, "The Structure of Ill-Structured Problems", *Artificial Intelligence*, Vol. 4, No. 3-4 (1973).
6. Allen Newell & Herbert A. Simon, *Human Problem Solving*, Englewood Cliffs: Prentice-Hall, 1972.
7. George Pólya, *How to Solve It*, Princeton: Princeton University Press, 1945；中译《怎样解题》，涂泓、冯承天译，上海科技教育出版社，2007。
8. Micheline T. H. Chi, Robert Glaser & Ernest Rees, "Expertise in Problem Solving", in R. J. Sternberg (ed.), *Advances in the Psychology of Human Intelligence*, Vol. 1, Hillsdale: Erlbaum, 1982.
9. Thomas S. Kuhn, *The Structure of Scientific Revolutions*, Chicago: University of Chicago Press, 1962；中译《科学革命的结构》，金吾伦、胡新和译，北京大学出版社，2003（反常与范式转换）。
10. Karl R. Popper, *Conjectures and Refutations*, London: Routledge, 1963；中译《猜想与反驳》，傅季重等译，上海译文出版社，2005。
11. Imre Lakatos, "Falsification and the Methodology of Scientific Research Programmes", in *Criticism and the Growth of Knowledge*, Cambridge University Press, 1970.
12. Daniel Kahneman & Gary Klein, "Conditions for Intuitive Expertise: A Failure to Disagree", *American Psychologist*, Vol. 64, No. 6 (2009).
13. Brooke N. Macnamara, David Z. Hambrick & Frederick L. Oswald, "Deliberate Practice and Performance in Music, Games, Sports, Education, and Professions: A Meta-Analysis", *Psychological Science*, Vol. 25, No. 8 (2014).
14. K. Anders Ericsson, Ralf Th. Krampe & Clemens Tesch-Römer, "The Role of Deliberate Practice in the Acquisition of Expert Performance", *Psychological Review*, Vol. 100, No. 3 (1993).
15. Gary Klein, *Sources of Power: How People Make Decisions*, Cambridge, MA: MIT Press, 1998（识别启动的决策模型）。
16. David G. Jansson & Steven M. Smith, "Design Fixation", *Design Studies*, Vol. 12, No. 1 (1991).
17. W. Edwards Deming, *Out of the Crisis*, Cambridge, MA: MIT CAES, 1986（PDSA 改进循环）。
18. Michael Polanyi, *The Tacit Dimension*, London: Routledge & Kegan Paul, 1966；中译《默会的维度》，许泽民译，上海人民出版社，2000。
19. Ludwik Fleck, *Genesis and Development of a Scientific Fact*, Chicago: University of Chicago Press, 1979（German original 1935）（思维风格与事实的生成）。

20. Robert K. Merton & Elinor Barber, *The Travels and Adventures of Serendipity*, Princeton: Princeton University Press, 2004 (意外发现的社会学)。
21. Ikujiro Nonaka & Hirotaka Takeuchi, *The Knowledge-Creating Company*, New York: Oxford University Press, 1995.
22. Jason Wei et al., "Chain-of-Thought Prompting Elicits Reasoning in Large Language Models", NeurIPS 2022, arXiv:2201.11903.
23. Shunyu Yao et al., "ReAct: Synergizing Reasoning and Acting in Language Models", ICLR 2023, arXiv: 2210.03629.
24. Emily M. Bender & Alexander Koller, "Climbing towards NLU: On Meaning, Form, and Understanding in the Age of Data", ACL 2020.
25. 王德生: 《SDE 智能体导论——从语言模型到发生世界模型的智能体革命》, 德麦国际出版, 2026, ISBN 978-1-970820-62-1 (本文第八章与第九章之改写底本)。
26. 王德生: 《SDE本体论入门: 世界是如何发生的》, 德麦国际出版, 2026 (三大方程、123 原理与三界结构的完整展开)。