

内容权

论决定表现的不是控制多少、也不是注意朝哪，而是「下一个动作是什么」这项权限落在谁手里

运动学习 × 音乐即兴神经科学 × 临床心理学

作者 王德生 + Claude

栏目 学科通融 · 站外碰撞

篇幅

发表 2026年8月1日

摘 要

运动学习说，把意识放到自己身上会干扰动作，要做好就别向内看；爵士即兴的脑成像说，创造恰恰需要内部驱动的自我表达增强，同时逐音符的在线审查减弱；正念研究又说，向内觉察可以在叙事性自我下降的同时增强。三条各自有实验支撑，放在一起却互相否定——而且两种被公认优秀的状态在前额叶上的模式近乎相反。本文认为僵局来自一个三方共有的前提：**意识对行为的介入被当成了一个量，三家只在争这个量该往哪个方向调。** 本文撤销它，提出决定表现的**不是控制的多少、也不是注意的朝向，而是一项权限落在谁手里：「下一个动作、下一个音、下一个念头是什么」由谁定——本文称之为内容权**，与之相对的框架权是：什么算完成、哪条线不许越、什么时候评。把内容权与「局部回路能不能自己发现并修好误差」交叉，得到四格：生成性掌握、接管损耗、放任性漂移、校准性支架。本文与最接近的六种既有说法的分界，全部压在同一个可做的对照上：**提示文本逐字相同，只改变「下一步能否在不经一次确认的情况下发生」**——若这个对照做不出差异，本文应当被并入已有的信息负荷解释。三处让步各削掉本文一块，其中母语习得这个反例最贵：当环境能让误差自己变得可见，回路可以在完全没有内容接管的条件下长起来。

这一篇由哪三个领域的争论撞成

三边对「一个人做得最好的时候，意识应该在哪里」给出互相排斥的答案：运动学习说把注意放到自己身上会干扰动作组织；爵士即兴的脑成像说创造需要内部驱动增强而在线自我审查减弱；正念研究说向内觉察可以在叙事性自我下降的同时增强。三个「向内」其实是三条不同的轴（来源、对象、层级、所有权），本文把它们拆开之后剩下的那个变量，是权限的归属。

运动学习 · 效果聚焦优势 ↔ 2024 年的稳健再分析

[Wulf, Chiviacowsky, Schiller & Ávila, Frequent External-Focus Feedback Enhances Motor Learning \(Frontiers in Psychology, 2010\)](#)

四十八名十至十二岁儿童练习足球掷界外球，比较两种提示对象与两种反馈频率：持续获得效果导向反馈的一组在无反馈的迁移测验中动作形式最好（5.95，其余三组 5.05–5.30），准确性无组间差异。这类结果累积成「有意识控制身体会限制自动调节」的判断；而 2024 年一项稳健贝叶斯再分析（Psychological Bulletin 150(11)）认为该文献存在中度至强的发表偏倚，校正后多个平均效应接近零。

<https://doi.org/10.3389/fpsyg.2010.00190>

音乐即兴神经科学 · 自我表达增强而在线审查减弱

[Limb & Braun, Neural Substrates of Spontaneous Musical Performance: An fMRI Study of Jazz Improvisation \(PLoS ONE, 2008\)](#)

六名职业爵士钢琴家在扫描仪中即兴演奏，与过度学习的固定演奏对照：背外侧前额叶与外侧眶额区广泛下降，额极内侧前额叶局部升高，感觉运动区参与增加。研究者的解释是内部驱动的自我表达仍活跃，而持续监视、评价、有意操纵每一个音符的过程减弱——与运动学习的结论正相反。

<https://doi.org/10.1371/journal.pone.0001679>

临床心理学 · 向内不等于自我纠缠

[Farb, Segal, Mayberg, Bean, McKeon, Fatima & Anderson, Attending to the Present \(Social Cognitive and Affective Neuroscience, 2007\)](#)

比较「把当下刺激接到个人特质与故事上」与「只留意此刻的想法、情绪与身体感觉」两种指令，并对照受过八周训练者与新手：受训者在后一种条件下出现更广泛的内侧前额叶下降，同时右侧外侧前额叶、脑岛、次级躯体感觉区与顶下小叶参与增强。内部感觉可以被清楚观察，却不必被编进关于「我是谁」的叙述。

<https://pmc.ncbi.nlm.nih.gov/articles/PMC2566754/>

目 录

一、导论：三个不该同时为真的答案.....	5
二、三条既有路径与它们各自停在哪里.....	6
三、冲突定位：被压在一根轴上的四件事.....	7
四、内容权.....	8
五、二维辨别表.....	11
六、逐领域兑现，以及三处让步.....	11
七、判据：一个提示逐字相同的实验.....	14
八、与最近的六种说法的分界.....	16
九、三家各自为何到不了这一条.....	19
十、推论、适用边界与反噬预言.....	19
十一、本判断对它自己适用吗.....	20

结 论.....21

这一次作废了多少条.....22

注 释.....22

一、导论：三个不该同时为真的答案

三场争论，分属运动学习、音乐即兴的脑成像研究与临床心理学。每一场都有厚实的实验支撑，而它们对同一个问题给出的答案不能同时为真。

那个问题是：一个人做得最好的时候，意识应该在哪里。

第一场在运动学习。 一个十岁的孩子练习足球掷界外球。教练可以说「手臂从头顶越过」，也可以说「让球沿着头顶上方的弧线飞出去」。两句话给出的技术信息几乎一样，落点却不同：前一句要求他盯住自己的身体部位，后一句要求他盯住球的去向。一项把四十八名十至十二岁儿童分成四组、比较两种说法与两种反馈频率的实验发现，持续获得「效果导向」反馈的那一组，在事后没有任何反馈的迁移测验中动作形式最好（5.95 分，其余三组在 5.05 至 5.30 之间），而准确性没有组间差异。这类结果累积成一个强势判断：**把意识钉在自己身上会干扰本可以更自动的动作组织；要做好，就别向内看。**

第二场在音乐即兴。 六名职业爵士钢琴家在脑成像设备里即兴演奏。即兴不是把背熟的乐句放一遍，而是在调式、和声、节拍与风格的边界内，即时生成此前不存在的旋律。与过度学习的固定演奏相比，即兴时出现一组分离的变化：背外侧前额叶与外侧眶额区广泛下降，额极内侧前额叶局部升高，感觉运动区域参与增加。研究者的解释是：**内部驱动的自我表达仍然活跃，而持续监视、评价、有意操纵每一个正在发出的音符的过程减弱了。**这与第一场正相反：内部来源不但不是障碍，还是创造的条件。

第三场在临床心理学。 一组参加过八周正念训练的人与新手作对照，分别在两种指令下接受扫描：一种要求把当下的刺激接到「我一贯如何、这说明我是什么人」的故事上，另一种要求只留意此刻的想法、情绪与身体感觉。受过训练者在后一种条件下出现更广泛的内侧前额叶下降，同时右侧外侧前额叶、脑岛、次级躯体感觉区与顶下小叶参与增强。**向内并没有把人拖进自我纠缠：内部感觉可以被清楚地观察，却不必被编进关于「我是谁」的叙述。**

三条各自成立。放到一起就断了。

第一场说：**向内有害。** 第二场说：**内部驱动是创造的来源。** 第三场说：**向内可以不带自我。** 更尖锐的是后两场：爵士即兴与当下觉察都被描述为流畅、当下、高质量的状态，前额叶的模式却近乎相反——一个是内侧升高而外侧广泛降低，另一个是内侧降低而部分外侧升高。

于是三条常见的解释轴同时失效：**注意朝内还是朝外**解释不了（三场里的「内」根本不是同一件事）；**控制多还是少**解释不了（即兴并没有丢掉和声与节拍，觉察也没有丢掉注意维持）；**某个脑区开得多还是少**解释不了（两种优秀状态在同一片区域上方向相反）。

本文认为，僵局来自一个三方共有、却从未被单独说出的前提：**意识对行为的介入是一个量**——多了会干扰（第一场）、必须有一部分留住（第二场）、可以只减掉其中一种（第三场）。三家争的是这个量该往哪个方向调，没有人问它是不是那个起作用的东西。

本文撤销这个前提，提出：

决定表现的**不是控制的多少，也不是注意的朝向，而是一项权限落在谁手里**：「下一个动作、下一个音、下一个念头是什么」这件事，由谁定。本文称这项权限为「**内容权**」。

与它相对的另一项权限，本文称为**框架权**：什么算完成、哪条线不许越过、多快、什么时候评、出了什么信号就中断、失败之后从哪里重来。

三场争论在这条线上是同一个形状。掷界外球时，「手臂从头顶越过」这句话把内容权收了上去——肩、肘、腕此刻各自怎么动，由意识逐项指定；「让球沿着弧线飞出去」这句话把内容权留在原处，只给了一个完成条件。即兴时，和声与节拍是框架权，逐音符的审批是内容权，而钢琴家**保留了前者、交出了后者**。当下觉察时，注意的维持、走神的识别、安全的边界是框架权，而「我现在应该有什么感觉、这种感觉说明我是什么人」是内容权——训练要减掉的正是后者。

三家看上去在争「向内还是向外」，其实各自摸到了同一件事的一角：**意识必须在场，但不能替局部把每一步定下来**。

第二节梳理三条既有路径并指认它们各自停在哪里；第三节把「内部注意」拆成四条互不相同的轴，说明僵局如何由压轴造成；第四节界定内容权，并把它与「指导多少」这个量分开——这是全文最要紧的一刀；第五节给出二维辨别表；第六节在若干领域兑现，其中三处反过来迫使本文收窄；第七节给出判据，包括一个提示文本逐字相同、只有权限归属不同的实验设计；第八节与最近的六种说法逐条分界；第九节说明三家各自为何到不了；第十节给出推论、边界与反噬预言；第十一节处理本判断对自身的适用。

二、三条既有路径与它们各自停在哪里

2.1 干扰论：意识介入会打断自动化

第一条路径把优秀理解为「别插手」。它的经典表述是：对身体动作的有意识控制会限制本来更快、更反射性的自动调节；把注意放到动作效果上，误差就交给视觉、前庭、本体感觉与肌肉协同一起消化，意识不必在每一毫秒里决定肩、肘、腕各自怎样移动。压力下的失常也由此得到解释：一个熟练者在关键时刻重新逐项检查自己的动作，于是变慢、变僵。

它的长处是给了一条可操作的指令：把话说到效果上去。

它停在哪里？停在两处。

其一，它把「**向外**」当成了原因。可外部提示同样可以逐项指定内容——教练通过耳机在运动员起跑的每一步报出脚落点，这是彻头彻尾的外部提示，也是彻头彻尾的逐项接管。而某些内部提示只给一个整体锚点（「让呼吸保持连续，不去追某一次吸气的长度」），并不规定任何一个内容单元。方向与权限是两件事，这条路径把它们绑成了一件。

其二，它的**证据本身正在被修正**。2024 年一项稳健贝叶斯再分析与系统综述认为，这一族文献存在中度至强的发表偏倚；校正之后多个平均效应接近零，同时保留了明显的异质性。本文不把这当成「早期实验

无效」，而当成一个提示：**真正起作用的变量藏在情境、任务与学习阶段里，而不在「内外」这条轴上。**——顺带说明，本文并不把「推翻外部聚焦的普遍优势」算作自己的贡献，那是那份再分析做的事。

2.2 来源论：创造需要从我自己长出来

第二条路径把优秀理解为「让它自己长出来」。即兴的脑成像结果被通俗地读成「创造时关掉理性」，但这个读法解释不了一件事：如果理性关掉了，音乐为什么仍然合乎和声、节奏与风格？钢琴家没有失去规则，也没有失去对整体走向的把握。改变的是规则进入内容的方式。

因此这条路径真正说对的是：**内部驱动与内部监视必须分开。**前者回答作品从哪里长出来，后者回答每一个片段是否立刻接受评价。一个人可以高度地表达自己，同时不在每个音符上追问「这够不够聪明、够不够像我、观众认不认」。

它停在哪里？停在**它只描述了一种状态，没有给出条件。**六名被试、固定效应分析、单一任务，都限制了推广；更要紧的是，后续的创造力研究显示发散思维并不是简单地降低执行控制——默认网络与执行网络可能按阶段发生耦合。于是「关掉审查」这条处方在新手身上会失败：一个尚未把技术语法沉进听觉与手指的人撤掉内容控制，多半只会重复熟悉的片段或随机乱按。**这条路径缺的是「什么时候可以这样做」。**

2.3 关系论：可以观察内部而不占有它

第三条路径把优秀理解为「换一种与内部经验相处的方式」。它最有价值的贡献是把「内部注意」内部的一条裂缝显示了出来：同样是向内，一种是叙事占有——感觉一出现就被接进「我一贯怎样、这说明我怎样、别人会怎样看」的时间链；另一种是当下显现——感觉被辨认、定位、允许变化，却不被立即转成身份判断。二者都发生在身体与心灵内部，对行为的后果却相反。

它停在哪里？停在**它是一门关于内部经验的学问，没有理由把同一结构推到动作与作品上。**它也没有回答一个临床上极重要的问题：对某些人而言，向内注意会触发创伤记忆或惊恐循环，此时「接受一切」不是解脱而是把人独自留在过强的刺激里。这条路径知道这件事，却是把它当作禁忌症处理的，而不是当作同一个变量的另一端。

2.4 三条路径共享的东西

三条路径分别把「意识介入的量」「内部驱动的有无」「与经验的关系」当作自变量。它们共享的是**自变量的类型**：都是关于意识状态的描述——多少、朝向、模式。本文提出的自变量不是状态，是一项**权限的归属**：在一个具体的时刻，下一个内容单元由谁定。

三、冲突定位：被压在一根轴上的四件事

「内部注意」与「外部注意」这一对词，把至少四件互不相同的事压进了同一条轴。

其一，来源。一个行为是由自身意图生成的，还是由外部命令触发的。

其二，对象。注意落在身体感觉、动作效果、环境目标，还是他人的评价上。

其三，层级。意识是在设定整体规则，还是在逐项指定微观内容。

其四，所有权。一段经验被当作正在发生的事件，还是被立即追认为「关于我的证明」。

只要这四条不分开，同一个「向内」就会同时指身体觉察、动作微管、情绪反刍与自我评价；同一个「向外」也会同时指动作效果、环境约束、教练命令与观众目光。三场争论正是在这四条被压成一条之后才互相顶撞的：

运动学习说的「向内」主要是**层级**问题（逐关节指定）；即兴研究说的「内部」主要是**来源**问题（从自己长出来）；临床心理学说的「向内」主要是**所有权**问题（是否被迫认为身份）。三个「内」是**三条不同的轴，而它们被同一个字装着**。

拆开之后，三条冲突各自消解：

冲突一（向内有害 vs 内部驱动有益）。有害的是层级过深，有益的是来源在内。二者可以同时成立：一个人完全可以从自己长出来，同时不逐音符审批自己。

冲突二（向内有害 vs 向内可以不带自我）。有害的是所有权的立刻追认，无害的是对象落在身体感觉上。二者可以同时成立：感觉可以被看得很清楚，而不必被解释成「我又要失败了」。

冲突三（前额叶降低 vs 前额叶升高）。两种优秀状态减掉的是不同的东西：即兴减掉的是逐音符的在线审查，觉察减掉的是叙事占有；而两者各自增强的部分（内侧的自我表达、外侧的注意维持）恰恰都是框架侧的功能。**方向相反的是被减掉的那一项，不是「控制」这个总量。**

还需要一个时间标记。同一句提示，在行动之前、行动之中、行动之后，属于完全不同的东西。赛前想清楚手臂路径，未必破坏自动化；出手的几十毫秒内逐段检查，才造成延迟；结束后看录像复盘，属于离线的模型更新。艺术里的草图、即兴与修订，训练里的设定意图、体验过程与事后反思，具有相同的分层。**缺少时间标记的研究，会把「有意识地思考」当成同一种干预。**

四、内容权

4.1 定义

内容权：在一个具体的时刻，「下一个内容单元是什么」这件事由谁决定。

框架权：什么算完成、哪条线不许越过、以什么节奏推进、什么时候评价、什么信号触发中断、失败后从哪里重来。

「内容单元」按任务定：动作里是一个关节的一次调整，音乐里是一个音或一个乐句，念头里是一次感受的处置，写作里是一个句子，管理里是一封邮件的措辞。

两项权限不是抽象与具体的差别。一条极具体的提示可以停在框架层（「把球送过第二棵树」，只给完成条件），一条极笼统的话也可以是内容接管（「每一步都先想清楚再动」，要求每个单元都经过一次审批）。**判据只有一条：执行的那一侧，还有没有选择与修正的余地。**

4.2 与「指导多少」的分离——本文最要紧的一刀

一个自然的怀疑是：内容权不过是「指导多不多」的另一种说法。指导得细，就是内容权在上；指导得粗，就是内容权在下。若真如此，本文只是给一个已知的量换了名字。

它们可以被干净地分开，因为**内容权可以在信息量完全不变的情况下改变归属。**

设想同一段提示，一字不改：「起跑的前三步，重心压在前脚掌，步幅比平时略小。」

条件甲：说完这句话，让他跑。

条件乙：说完这句话，要求他每跑完一步在心里确认一次「这一步做到了没有」，做到了再进入下一步。

两个条件的文字逐字相同，信息量、句长、术语密度、工作记忆的载入内容都一样。改变的只有一件事：**下一步能不能在没有一次确认的情况下发生。**甲把内容权留在了跑的那一侧，乙把它收了上去。

这个操纵是本文与「指导量」一族的分岔点，也是第七节实验设计的核心。若表现只随信息量变化，甲乙应当没有差别；若表现随权限归属变化，熟练者在乙条件下应当变慢、变僵，而初学者在乙条件下可能反而受益。

4.3 第二个变量：局部回路充分性

只有内容权还不够，因为同一条指令对新手和熟练者的后果相反。第二个变量是：**在没有逐项指令的情况下，执行的那一侧能不能自己发现与任务有关的误差、在允许的时间内修好它、并把这次修正留到下一次。**本文称之为**局部回路充分性**。

它不是一般能力，也不是人格特质，而是「某个人、在某个任务、某个环境、某个时刻」的关系变量。它至少包含四个成分：反馈是否清楚可得；误差信号能否对上合适的修正；修正是否快过任务的要求；这套修正能否在距离、速度、情绪或器材改变时仍然有效。

一个钢琴家可能在熟悉的和声上充分、在陌生的拍号上不足；一名运动员在训练里充分、在受伤或高压下暂时下降；一个练习者在轻微焦虑时能观察、在创伤被触发时需要更具体的外部支撑。**因此这不是一次性的毕业，而是一条会来回移动的线。**

4.4 一处必须自己说破的循环风险

这里有一个危险，本文必须先说破，否则整条论证会变成一句同义反复。

局部回路充分性被定义为「在没有逐项指令时能自己发现并修好误差」，而本文的预测是「充分性高时收走内容权有害」。若不加约束，**这两句话之间只隔着一层纸：**一个已经能自己修好的系统，逐项指令当然是多余的——多余是定义带来的，不需要任何实验。

本文的主张比这一步更强，而正是这一步强出来的部分承担了全部风险：**多余不等于有害**。一条被证明多余的指令，完全可以只是被忽略掉、什么也不改变。本文断言它不会被忽略，而会主动地把表现拉下来——变慢、变僵、变刻板、把变异压掉。这一断言不能由定义担保，只能由证据担保。

因此本文接受两条自缚：

其一，**充分性必须由与结果不同的独立任务测量**：无提示条件下的保持、轻度扰动后的恢复斜率、双任务代价、自发纠错率。绝不允许用同一场表现同时定义自变量和因变量。

其二，**接管损耗必须以正向证据出现**：不是「收走内容权之后没有变好」，而是「明显变差」，且差的方向可以预先说出（时间变长、变异结构收缩、扰动恢复变慢）。若只观察到「没有变好」，本文的主张就只剩下定义，应当撤回。

4.5 两种损耗

由此得到一对方向相反的失败：

接管损耗：局部回路已经能自组织与自纠错时，内容权仍被上层持续握着，由此造成的变慢、变僵、变异收缩、情绪反刍与创造阻塞。

放任损耗：局部回路还不能可靠运行时，内容权就被过早交下去，于是随机错误被反复练成习惯、危险信号无人识别、形式尚未建立就被称作自由。

二者不是「控制多」与「控制少」的两端，而是**同一处错配的两个方向**。这也解释了为什么「放松」「别想」「进入状态」这类话常常无效：它们表面上要求减少控制，实际却制造了一个新的内容任务——执行者开始持续检查自己是否足够放松、是否已经无我。**控制没有减少，只是升格为对「不控制」的控制**。本文称之为**元接管**。

4.6 框架权也会退化

最后一条，是本文对自己早先那套写法的一处更正。

框架权容易被想成「多多益善」：目标写得越清楚、边界列得越全、指标定得越细，越好。这是错的，而且错法可以精确说出：**当框架的条目多到必须在行动过程中逐项核对时，它在功能上已经变成内容接管**。组织里的原则清单、训练里的检查表、治疗里的操作手册，常常发生这种伪装：名字叫原则，实际是几十项实时审批。

所以框架的清晰度不是一个可以独立累加的好处，它**必须与内容权的归属相乘**：稀疏到能被整体地保持在心里，才是框架；密到需要一条条对照，它就落回内容层，并且伪装成了自由。

五、二维辨别表

把两个变量交叉，得到四种状态。它们不能靠一次成绩或一次主观流畅感区分，只能靠中断、扰动、迁移与延迟测验分辨。

表一 局部回路充分性 × 内容权归属

	内容权在执行的那一侧	内容权被上层握着
局部回路充分	甲格 生成性掌握。 目标与边界在场，具体内容自行生成与修复：运动员围绕效果组织动作，钢琴家在和声内生成，练习者维持觉察而不占有感觉	乙格 接管损耗。 高手被逐项监视（被別人，或被自己）：压力下僵硬、艺术陈词化、反刍、「想得太多」
局部回路不足	丙格 放任性漂移。 回路尚未成形却只给一句抽象目标或「自由发挥」，错误无人识别、无人修复，练一百次也不长进	丁格 校准性支架。 内容提示暂时承担教学、康复与安全功能—— 唯一正当的深控制，但必须自带撤出条件

这张表给出一个关键的反转：**同一条内容级指令，在低充分性时可能救人，在高充分性时会害人；同一份框架级自由，在高充分性时产生创造，在低充分性时产生随机。**因此任何一刀切的「少控制」「盯住外面」「相信直觉」都缺一根轴。

还有三件事只有这张表能说清。

其一，成绩幻觉。乙格在低压力、熟悉环境里可能暂时得高分——外部指导或内部微管把内容托住了；一旦提示撤去、任务变形或时间压缩，缺口才显出来。丙格有时也会偶然产生新颖结果，却没有重复与恢复的能力。所以判断落在哪一格必须看过程与迁移，不能看一次作品、一次比赛或一次舒适感。

其二，丁格不该被污名化。它是能力形成的必要过渡。真正的问题不是「用了支架」，而是**支架没有撤出条件**：什么读数达到之后减少提示，什么错误重新触发具体指导。没有撤出条件的支架会固化成依赖；没有返回入口的自由则会把失败全部归到个人身上。好的系统同时设计退出与回归。

其三，迁移必须可逆。稳定环境里的熟练不等于任何条件下都该停在框架层。场地湿滑、器材更换、身体疼痛、情绪跃升、规则突变，都可能让局部回路暂时失配。**成熟的标志不是永不下沉，而是能用最少的必要控制重新校准，并在回路恢复之后再次退出。**永久的内容控制与永久的放手，都缺这种可逆性。

交接的时机可以用三个信号识别：在没有提示的条件下错误会不会自己下降；轻度扰动之后能不能自行恢复；解释负担增加（加一个简单的并行任务）时表现会不会明显恶化。若一个人只有在**教练持续说话时才做得对**，内容权还不能交；若加一个简单双任务就崩，自动化还不充分；若轻度扰动后能快速重建，内容权就该交下去了。

六、逐领域兑现，以及三处让步

本节先给六处正面兑现，再给三处反过来迫使本文收窄的让步。让步部分比兑现部分更重要。

6.1 六处兑现

（一）训练。 训练计划的核心不是「提示要多要少」，而是**内容权什么时候交出去**。第一阶段，教练帮助建立感知与误差的对应，内容提示少而关键，每一条都对应一个可观察的结果；第二阶段，把同样的技术要求改写成效果、节奏或环境约束，让运动员在较完整的动作里自行修复；第三阶段，评价移到动作的间隙，比赛中只保留少量框架词与事件触发信号；第四阶段，用扰动、疲劳与压力测试检查回路是不是真的充分，而不是因为环境熟悉才显得流畅。这也改写了「压力下退回基本动作」的含义：**有效的退回不是在比赛中重新逐项想起全部技术细节，而是激活少数能恢复整体协同的框架锚点。**

（二）创作。 艺术教育长期在「技术规范」与「自由表达」之间摆动。这条线把争论改写为：哪些规则必须成为框架，哪些选择必须留给局部生成。没有和声、材料、主题、时间或受众边界的「自由」多半是丙格；边界深入到每个内容单元，则是乙格。成熟创作还包含一件本文认为被低估的事——**生成权与否决权的分时**：若否决权在每个微时刻同步运行，许多尚未成形的材料会被过早删掉；若否决权永久缺席，作品会失去选择。成熟的做法是把否决权推迟到足以看见关系之后，使评价仍然严格，却不抢占生成的第一时刻。

（三）情绪调节。 「我必须停止焦虑」是一条内容级指令：它规定了下一刻应该出现什么感觉。越急于让它消失，越会把注意锁在「它还在不在」的连续监测上——这正是元接管。框架级的改写是：「我需要辨认强度、维持安全、让波峰完成，并在超过阈值时求助。」前者控制内容，后者控制关系与边界。**心理自主不是体验完全受控，而是拥有进入、退出、命名、求助与恢复的权力，同时不必逐项制造正确的体验。**

（四）康复。 受损的回路充分性下降，身体内部提示、镜像反馈与分段动作重新变得有益——这是丁格的正当用法。随着患者能在无提示条件下识别偏差并自我修复，治疗师应逐步把内容权交还给目标效果与环境约束。**若整个康复过程都坚持「盯住外面」，就是把一条阶段性规律当成了永久法则。**

（五）教学。 例题、步骤提示、概念分解对应内容支架；开放问题、迁移任务、要求解释对应框架控制。学生还没有误差识别能力时直接「自主探究」，容易漂移；已经掌握方法后仍被要求逐步照抄，则产生接管损耗。教学质量不在支架多少，而在**支架能否按能力证据撤出，并留一个返回入口。**

（六）管理与人机协作。 战略目标、资源边界、风险阈值、复盘窗口属于框架权；逐封邮件、逐句文案、逐项沟通的审批属于内容接管。团队能力低时完全放权会放大错误，能力高时持续微管会同时降低速度、责任感与局部创新。至于人机：一个系统可以成为内容级的接管者（实时补全每一句、推荐每一步、不断纠正风格，使人得到漂亮结果却失去局部生成与修复能力），也可以成为框架工具（帮助澄清目标、列出约束、设置检查点、在关键异常时提醒）。本文对这两者的差别有一条可裁决的预测，见第八节第三小节。

6.2 让步一：高风险领域——内容权不能「交下去就不管」

外科、航空、核电这类领域迫使本文收窄。

按本文的严格表述，充分性高就该把内容权交给执行的那一侧。可在高风险领域，即使术者的回路完全充分，制度也不允许完全撤除监督：核对清单、双人确认、异常上报，一样都不能少。若坚持严格表述，本文就与这些领域的安全实践正面冲突，而那些实践是有代价换来的。

正确的收窄是：**本文主张的从来不是「内容权下放」，而是「内容权的默认归属加上可撤回」**。默认归属决定了不出事时谁定下一步——是执行的那一侧；可撤回决定了什么信号一出现，上层就立刻收回。这两件事是分开的，而且第二件属于框架权（「什么信号触发中断」本来就写在框架权的定义里）。

由此得到一条更准确的表述：高风险领域需要的不是更深的内容控制，而是**更密的触发条件与更快的收回通道**。术中每一步都被口令打断，与术中一旦出现特定读数就立即停下，是两件后果相反的事。前者是乙格，后者是甲格加上一条硬的中断线。

6.3 让步二：创伤与惊恐——「交给局部」有时等于把人丢下

第二条让步来自临床。

对某些人而言，把注意转向内部会触发创伤记忆或惊恐循环。此时按本文的框架推理会出错：他们的问题看起来像「充分性不足」，于是似乎应当上收内容权、逐项规定该有什么感觉。但临床经验相反——逐项规定感觉恰恰会加重对抗。

正确的收窄是：**这类情形下下降的不是局部回路充分性，而是框架里的安全项**。需要补的是外部定向、身体接触边界、节律呼吸、停止信号与他人的共同调节——全都在框架权那一侧。这条让步反过来给本文加了一个硬性成分：**框架权必须包含中断权与退出权，否则「把内容权交给局部」在某些人身上等于把人独自留在过强的刺激里**。

还有一层：若训练把「持续觉察」设成道德义务，练习者失去退出权，那么框架本身就变成了接管。**自主不是坚持面对所有内容，而是拥有调节接近距离的权限**。

6.4 让步三：母语习得——丁格不是普遍的发展阶段

第三条让步最贵，它来自一个本文差点漏掉的反例。

没有任何人的母语是靠逐词指定学会的。抚养者不会规定下一个音节，不会逐句纠正语法，甚至绝大多数纠正都被证明无效或根本没发生；而每一个正常发育的孩子都长出了极其充分的语言回路。**照本文的表，这是一个「低充分性 × 内容权在执行侧」的长期状态，也就是丙格——可它产生的不是漂移，而是人类学习史上最可靠的成功**。

本文不能靠一句「语言是特例」把它推开。正确的收窄是把一条主张削掉：

本文**不主张「低充分性时必须上收内容权」**。低充分性只说明**上收内容权可能具有校准价值**，而它是否必要，取决于环境能不能在不接管内容的情况下提供足够密的反馈。

母语环境恰恰能：它提供海量的、即时的、自然的成败反馈（说错了没被听懂、说对了拿到了东西），而这些反馈全部落在**效果上**，不落在内容上。于是回路可以在没有任何内容接管的条件下自己长起来。

这条让步反过来给出了本文最有用的一条实践推论：**丁格是反馈稀缺的替代品，不是学习的必经阶段。**

当你发现必须逐项指定内容，先别急着做支架，先问一句：**是不是这个环境让误差变得看不见了？** 掷界外球之所以适合效果提示，正因为球的落点天然可见；而调整呼吸肌协同、术后重新招募某块肌肉、矫正一个危险的代偿动作，误差本身不可见——此时身体提示的真正作用不是「接管」，而是**临时把误差变得可见**。这也让本文与「向外优于向内」那条实践指引分开了：外部提示之所以常常更好，很可能是因为它顺带解决了可见性，而不是因为方向本身。

6.5 多人情形：内容权可以落在第三处

合奏、球队、戏剧与舞蹈还引出一件本文原来的两格装不下的事。在四人乐队里，一段乐句的内容权常常不在演奏者的意识里，也不在他的手指回路里，**而在别人手上**——他正在等对方那一句落下来再决定自己怎么接。轮流独奏时，内容权在几个人之间按小节流动。

因此更准确的说法是：内容权是一项**可以被托管的权限**，托管方可以是意识、局部回路，也可以是同伴、乐谱、清单或一台机器。本文的两格是把托管方分成「上层」与「执行侧」的一个简化；在多主体情形下，判据要改成：**下一个内容单元的决定，是否需要跨越一次等待**。需要等待的地方，就是权限所在。

七、判据：一个提示逐字相同的实验

7.1 核心可裁决判据

在信息量、句长、认知负荷与自主感全部匹配的条件下，仅改变「下一个内容单元能否在不经一次确认的情况下发生」，表现仍会出现方向相反的变化：局部回路充分者变差，不足者变好。

这条判据的全部价值在于它**不能被信息量解释**。这是本文与教育心理学里那条最接近的规律（见第八节第一小节）分歧的地方，也是本文能不能作为一条独立主张站住的唯一支点。

7.2 实验一：键盘序列，最便宜的一个

被试与操纵。 学习一段固定的按键序列。一组练习二十次，一组练习两百次，以此制造低、高两种局部回路充分性（并用独立指标验证：无提示保持、扰动恢复、双任务代价）。随后两组各随机分配到两个条件：

- **甲（内容权在执行侧）**：提示语「按顺序完成，注意整段的节奏均匀」，完成整段后再评价。
- **乙（内容权在上层）**：提示语完全相同，但要求每按一键后在屏幕上确认一次「这一键是否正确」，确认后才进入下一键。

为什么这个设计干净。 两条件的文字一字不差；乙条件多出来的不是信息，而是一次等待。若有人主张乙的额外操作本身带来负荷，可加第三个条件**丙**：同样每按一键后做一次与任务无关的按钮点击（等量的动作与中断，但不含确认）。**丙控制掉了「中断」与「额外动作」，只有乙含有权限的转移。**

主要结果。 撤除提示后的迁移表现、扰动恢复时间、按键间隔的变异结构、双任务代价。

预测。 高练习组：乙显著差于甲与丙，且丙与甲接近。低练习组：乙不差于甲，可能优于甲。**核心读数是两条曲线交叉，而不是某一组的平均值更高。**

什么结果会让本文错。 若乙与丙的差异不存在（即全部效应都由中断与额外动作解释），本文的自变量不成立；若两组都表现出同方向的乙劣于甲，只说明确认动作有普遍成本，与充分性无关，本文的第二根轴也不成立。

7.3 实验二：即兴生成中的评价时点

参与者先完成足量的技术熟悉，随后在生成阶段接受三种条件：连续在线评价；**总量相同**但集中于段落之间的评价；只保留主题与边界的框架条件。专家与新手分别分析；盲评者评新颖性、连贯性与风格适配，过程指标记录停顿、回撤与局部重复。

预测。 专家在「段间评价」与「框架条件」下新颖性更高且连贯性不降；新手在纯框架条件下连贯性下降。**关键是评价的总量相同而时点不同**——若两种评价条件没有差别，则「否决权分时」这一条不成立。

7.4 实验三：内部经验的两种指令

先用独立任务测量个体的局部回路充分性（走神发现延迟、情绪峰值后的恢复、身体信号辨识），再比较三种条件：改变感觉内容（「让这种感觉平静下来」）；维持观察框架（「留意它的强度变化，若超过某个程度就停下」）；无指导。结果包括反刍时间、体验回避、主观强度、返回延迟与生理恢复。

预测。 高充分性者在观察框架条件下恢复最好；低充分性且威胁较高者需要更具体的外部安全支撑——注意这条预测已经含进了第六节第三小节的让步：他们需要的是框架里的安全项，不是内容级的感觉规定。**若观察框架对所有人、所有威胁水平都产生同等作用，第二根轴的必要性下降。**

7.5 三条通用纪律

三组实验都必须把**信息量、认知负荷、评价威胁、提示频率**作为独立或匹配变量，否则「内容级提示较差」可能只是因为字更多、压力更大或工作记忆负荷更高。

操纵检验不能只由研究者按句子分类：要问参与者在行动中是否逐项指定内容、是否保有局部选择、评价频率如何变化。局部回路充分性必须由独立任务测量——这一条已在第四节第四小节写成自缚。

样本量应按交互效应而非主效应规划，并采用多实验室重复。运动注意研究已经出现过发表偏倚的争议，新研究需要注册报告、公开材料、报告零结果与模型比较。**若这条主张只能在小样本的事后分组里成立，它就不配跨领域。**

7.6 会让本文为假的八条

1. 在功效充分、信息量匹配的研究中，局部回路充分性与内容权归属长期不出现交互，某一侧始终以固定主效应获胜。
2. 逐字相同的提示下，甲乙差异完全被丙（等量中断与额外动作）解释——权限不是自变量。
3. 非评价性的身体觉察与逐关节的动作微管产生同样的负效应——说明「向内」本身才是原因，四条轴不必拆开。
4. 随着练习增加，最佳的内容权归属不移动：同一个人从新手到熟练始终需要同样的权限安排。
5. 接管损耗可被焦虑、工作记忆负荷或提示长度完全解释，权限归属不再提供独立预测。
6. 元接管无法与「越压制越反弹」的既有机制分开（见第八节第四小节的判决性对照）。
7. 脑成像中，前额叶的总活动量比「哪一项功能被减掉」更稳定地预测表现。
8. 在只需一步、无可分层结构的任务，或全部内容都必须经过显性符号计算的任务上，本文不适用——这是边界，不是证伪，但若这类任务其实占了绝大多数，本文的适用范围小到没有意义。

八、与最近的六种说法的分界

本节是本文最该被挑剔的地方。一条关于「指导该不该随能力退出」的主张，在好几个领域里都已经有人说过，而且说得比本文早、比本文有证据。**本文若不能给出可裁决的差别，就应当被吸收，而不是被当作新的。** 以下六条按接近程度排列，最近的排在最前。

8.1 与教学设计里的「专长逆转」——最近的一条

教育心理学里有一条被反复重复过的规律：对新手极其有效的教学支持（详细的解题范例、图文整合、逐步指引），用在已有相当知识的学习者身上会失效，甚至产生负作用；因此指导应当随专长的增长而淡出。这条规律的正式形式，就是「先验知识水平 × 教学方式」的显著交互。

这与本文第七节的核心预测在形状上几乎相同，而它有二十多年的实验积累。若本文只是把它换个说法搬到运动与情绪上，那本文就是多余的。

分界在**机制**，不在形状。那条规律的机制是**信息与负荷**：专家已经有了自己的组织结构，外来的详细讲解与之重复，于是必须同时处理两套信息，反而占用了工作记忆。按这个机制，**把冗余的信息去掉，问题就解决了。**

本文的机制是**权限**：损害不来自多出来的信息，而来自「下一个单元必须等一次批准」这件事本身。按这个机制，**即使一个字都不多，只要每一步都要等，损害照样发生。**

判决性对照就是第七节第二小节那个设计：提示逐字相同的甲与乙。负荷机制预测二者无差异（信息量、句长、冗余度完全一致）；权限机制预测熟练者在乙条件下变差。这一条是本文全部独立性的所在——**若这个对照做出来没有差异，本文应当被并入那条规律，不再单列。**

还有一处次要但值得说的差别：那条规律讲的是**教学材料的设计**（外部给什么），本文讲的是**权限的归属**（谁决定下一步），因此本文可以处理一件那条规律处理不了的事——**一个人对自己的接管**。没有任何外部指导的运动员，在压力下开始逐关节检查自己，信息量没有增加一个比特，表现却塌了。这一情形在负荷框架里只能被解释为「自我监控占用了资源」，而那已经是在借用另一套机制了。

8.2 与管理学里的「按成熟度调整领导方式」

管理学里有一个流传极广的模型：领导行为按「指导」与「支持」两维分成四种风格，下属的就绪度（能力与意愿）分四级，风格应当与就绪度匹配——低就绪度用高指导，高就绪度用授权；而且就绪度是**按具体任务**判断的，不是人格层面的标签。

这与本文第六节的管理部分几乎是同一张表，且早了几十年。分界有两条。

其一，那个模型调的是「行为的量」（多说还是少说），本文调的是「权限的归属」（谁定下一步）。

二者可以分离，而分离处正是最常见的一种管理病：**少说话，但每一步都要报批**。按那个模型，这属于低指导，应当接近授权；按本文，这是最纯粹的乙格——上层放弃了给信息，却没有放弃决定权，于是执行侧既得不到帮助又不能自己动。**判决性对照：把「指导量」与「决定权」正交操纵四格，那个模型预测结果主要由指导量决定，本文预测主要由决定权决定。**

其二，那个模型是他人对他人的，本文包含同一个人对自己。 运动员在压力下自己收回内容权，创作者在生成时自己启动否决权，练习者要求自己「现在必须平静」——这些没有第二个人在场，管理模型无从适用，而本文预测它们与被别人微管的后果同形。

8.3 与人因工程里的「自动化水平」

人因工程有一条成熟的结论：自动化程度越高，常规表现越好、负荷越低，但操作者会脱离回路——情境意识下降、技能退化，一旦自动化失效，接管的速度与质量都明显变差。

这一族与本文第六节的人机部分高度重叠，而且它有几十年的实证与事故记录。分界在于：**那一族的自变量是「机器替人做了多少」，本文的自变量是「谁决定下一个内容单元」。** 这两件事可以正交，而且正交之后能给出一条方向相反的预测。

把人机分工切成两格：

- **甲 替你执行、你来决定：**你说做什么，它去做（工具型）。自动化程度**高**。
- **乙 替你决定、你来执行：**它逐句给出下一句，你负责敲下去或点接受（补全型）。自动化程度**低**——毕竟动作还是你做的。

按自动化水平那一族排序，甲的脱离回路风险更高；按本文，乙的损害更大。 因为退化的从来不是手指，而是「下一步该是什么」这个判断——它正是局部回路里唯一不能被代做的部分。判决性预测：在总产出质量相当的前提下，**长期使用乙型辅助的人，在撤去辅助后的独立表现与迁移，会比长期使用甲型辅助的人差得更多，**尽管乙型的自动化程度更低。

这条预测有直接的实践含义，也很容易被做：把同一件工作分给两组，一组用逐句补全，一组用「先说清楚要做什么、由系统去执行」，三个月后撤掉工具做同一份任务。

8.4 与心理学里的「越压制越反弹」

本文的「元接管」需要与一条经典机制分开：刻意不去想某件事时，一个负责搜寻「不该出现的东西」的监控过程会被启动；在负荷或压力之下，这个监控过程反而把被压制的内容推到前台。

分界在机制与解法。那条机制的核心是**监控进程在负荷下反向启动**，因此解法是降低负荷、放弃压制。本文的元接管说的是另一件事：**「不控制」被当成了一个内容目标**，于是**内容权被重新收了上去**——执行者开始逐时刻检查自己是否足够放松、是否已经进入状态。

判决性对照。 给一个可执行的**内容替代目标**（把注意放到某个具体的外部单元上，例如球的落点、节拍器的第二拍、脚底与地面的接触），在不降低任何负荷的前提下：反弹机制预测无效甚至更糟（多了一件事要做），本文预测有效——因为内容权重新有了去处，不再空转在「检查自己有没有在控制」上。

8.5 与运动学习里的「挑战点」

运动学习内部还有一条与本文形状接近的框架：最佳的信息量与任务难度随技能水平移动，太易与太难都损害学习，因此存在一个随能力移动的最优点。

这条框架与本文共享「最优点会移动」这个结构，分界在**移动的是什么**：那里移动的是**难度与信息量**，本文移动的是**权限的归属**。二者可以分开：一个难度恰当、信息量恰当的任务，仍然可以要求执行者每一步等一次确认（乙格）；一个难度过低的任务，也可以完全把内容权交下去。**判决性对照：把难度与权限正交**，若表现完全由**难度—能力匹配**解释，本文应当被并入那条框架。

需要说明的是，这一条本来是本文最容易漏掉的一位——它在同一个领域、用同一种「随能力移动」的结构，却因为讲的是难度而不是权限，很容易被当成不相干的邻居而略过。**把它放进来，本文才算真的清过场。**

8.6 与三家自身的分界

与干扰论。 二者都预言熟练动作会被身体微管破坏。差别在于本文不把「向外」当成普遍原因：外部提示若逐项规定微观时序，同样是接管；内部提示若只维持一个非评价的整体锚点，或用于低充分性的康复，则可能有益。这可以用「外部微管」与「内部框架」两个条件直接裁决。

与来源论。 二者都承认自我表达可以增强而在线审查可以减弱。差别在于本文给出条件而不只是状态：低充分性者撤去内容控制不会得到创造，只会得到重复或随机。因此「关掉审查」不是处方，**能不能在偏离之后回到结构、能不能识别失去连贯的时刻，才是可以放手的证据。**

与关系论。 二者共享「想法与感觉可以作为事件被观察」。差别在于本文把同一结构推到了动作与作品上，并加上了充分性这根轴，因而能预测它在什么时候失败——回路不足、威胁过高、安全框架不足时，「只是观察」不但不够，还可能有害。这一条已在第六节第三小节兑现为让步。

九、三家各自为何到不了这一条

运动学习到不了，因为它的因变量长在可见的一侧。这门学问的标准任务——投掷、平衡、击球——误差天然可见，于是「把注意放到效果上」既解决了权限问题，也顺带解决了可见性问题，两件事在它的实验里从不分开。一门自变量与另一个变量长期共变的学问，不会觉得需要把它们拆开。

音乐即兴研究到不了，因为它只有一个高水平样本。被扫描的都是职业演奏者，也就是充分性一端的人。在这一端，内容权本来就在局部，研究者看到的只能是「审查退场」这半张图；没有低端样本的领域，看不见一条会移动的线，只能看见一个状态。

临床心理学到不了，因为它的两种模式是按内容分的，不是按权限分的。它区分的是「与自我有关的叙述」与「当下的感觉」——这是所有权那条轴。至于「谁决定下一刻出现什么」，在它的框架里不是一个问题，因为感觉本来就不由人指定。一个默认内容不可能被指定的领域，不会把「指定权」当成变量。

三家的边界都不是能力问题：一个的自变量与可见性共变，一个只有一端的样本，一个的对象天然不可指定。三者并排放着，那个共同的变量才显出来。

十、推论、适用边界与反噬预言

10.1 推论

推论一。 任何以「他还不够熟练」为由收回内容权的做法，都必须同时说出撤出条件；说不出撤出条件的支架，与依赖没有区别。反过来，任何以「让他自由发挥」为由交出内容权的做法，都必须同时说出返回入口。

推论二。 想让一个人学会，第一件该做的不是设计支架，而是让误差变得可见。可见性够高时，回路可以在完全没有内容接管的条件下长起来（第六节第四小节的母语）；可见性不足时，内容提示的真正作用是临时充当可见性，一旦可见性被别的办法补上，它就该退。

推论三。 评价不是越少越好，而应从连续在线转为间歇、分段与事件触发。一次完整动作、一个乐句、一轮呼吸可以先运行完再检查；若评价持续侵入每个微单元，局部回路就变成了一个等待上级批准的系统。

推论四。 对人机分工的选择应当按「谁决定下一步」而不是按「省了多少事」来定。省事最多的那种辅助，恰恰是替你决定的那种。

10.2 适用边界

其一，本文只处理有可分层结构的行为。纯反射、单步动作、以及全部内容都必须经过显性符号计算的任务，不适用。

其二，本文不含价值排序：内容权被上层握着不是罪，丁格是必要的过渡；有问题的只是**与充分性错配**，以及**没有撤出条件**。

其三，本文的两格是简化。多主体情形下内容权可以落在同伴、乐谱、清单或机器上（第六节第五小节），此时判据要改成「下一个内容单元的决定是否需要跨越一次等待」。

其四，「充分性」不能用身份标签代替。专家在新器械、新曲式、新情绪或伤后任务上照样可能处在低充分性；新手在极小的局部技能上也可能很快形成可靠回路。测量必须针对具体任务，且至少包含保持、迁移、扰动与恢复四项。

10.3 反噬预言

按本文设计干预，最自然的做法是要求教练、教师与管理者在每次给出提示前，先评估对方的局部回路充分性，再选择提示的层级。

这会适得其反，方式可以精确说出。

第一重反噬：评估本身会变成新的接管。 一旦「先评估再提示」成为流程，评估表就会变长，而填表的人是上层。于是每一次提示之前多出一轮上层的判断，执行侧被观察得更细、被记录得更全——以「**判断该不该放手**」的名义，实施了一次更完整的接管。

第二重反噬：这条主张会被执行者用在自己身上。 读完本文的人可能开始在动作中自问「我的控制是不是又太深了」。这恰恰是本文自己命名的元接管：一条要求「别在行动中逐项检查」的主张，被做成了一件行动中要逐项检查的事。

因此本文不建议在行动中使用它。**它的正确用法在行动之前与行动之后：**行动前只设定少量框架——目标、边界、一个锚点、一条中断线；行动中让局部先跑；行动后按录像、作品、体验与迁移结果判断该下沉还是该上移。

另有一条更要紧的自我限制：「**充分性**」这个说法可以被用来制造等级——把某些人长期定为「回路不足」，据此长期剥夺他们的决定权。为防这一点，充分性的评估必须任务化、可逆、透明，并允许当事人参与决定交接。**低充分性只说明某一个具体回路当前需要支撑，不说明一个人缺乏自主的资格。**

十一、本判断对它自己适用吗

本文自己落在哪一格？

写作也是分层的：题目、边界、结构、什么算写完，是框架；每一句怎么落，是内容。而本文的产出方式是：**框架由人给定，内容由机器逐段生成，人再审定。** 按本文自己的表，这是「内容权在执行侧、框架权在上层」——甲格。

但这个自评不能就这么收下，因为它正好是本文第八节第三小节警告过的那一格：**替你决定、你来执行的反面——替你写，你来审**。审定权是一种否决权，而本文在第六节第二小节说过，否决权推迟到能看见关系之后行使是成熟的做法。这听起来很像在给自己发一张合格证。

诚实的说法是：**本文无法用自己的判据判定自己**。判据要求用无提示迁移与扰动恢复来测充分性，而一篇文章不会被扰动，也不会被要求在没有工具的条件下重做一遍。**能被这套判据测量的从来不是作品，是产出作品的那个人或那套流程；而本文只交出了作品**。

还有第二重，比第一重更硬。本文主张「评价不该侵入尚未成形的内容」，而本文本身是**按一次外部评分意见重写的**——上一稿被独立评分之后，评分意见指出了三处最接近的既有说法未被处理，第八节前三小节正是照着那份意见补的。**这恰好是本文推荐的用法：评价发生在段落之间，不发生在句子之内**。但它同时说明一件不那么体面的事：**那三处最近的说法，是被别人指出来的，不是本文自己找到的**。一篇论证「要主动去找谁早已说过同样的话」的文章，自己没有做到这一点——这一句留在这里，不删。

结 论

运动学习说向内有害，即兴研究说内部驱动是创造的来源，临床心理学说向内可以不带自我。三条不能同时为真，除非「向内」这个词里压着四件不同的事：来源、对象、层级、所有权。

拆开之后剩下的那个变量，本文命名为**内容权**——「下一个动作、下一个音、下一个念头是什么」由谁决定。与它相对的是**框架权**：什么算完成、哪条线不许越、什么时候评、什么信号中断、失败后从哪里重来。

卓越不是自我消失，是自我换岗。主体必须保留框架权，却不能一条已经能自己纠错的局部回路里继续握着内容权。

把内容权与局部回路充分性交叉，得到四格：生成性掌握、接管损耗、放任性漂移、校准性支架。同一条内容级指令在低充分性时救人、在高充分性时害人；同一份自由在高充分性时是创造、在低充分性时是随机。

本文与最接近的六种说法的分界，全部压在同一个设计上：**提示文本逐字相同，只改变「下一步能否在不经一次确认的情况下发生」**。若这个对照做不出差异，本文就该被并入教学设计里那条讲信息负荷的规律，不再单列。

三处让步各削掉了本文一块：高风险领域说明本文主张的是「默认归属加可撤回」，不是撤除监督；创伤情形说明框架权必须包含中断权与退出权；**而母语习得说明丁格不是发展的必经阶段——当环境能让误差自己变得可见，回路可以在完全没有内容接管的条件下长起来**。

最后是两条写死日期的赌注。第一条沿用本文所依据的原稿：

到 2030 年 12 月 31 日，若至少两类任务领域中的预注册研究均未观察到「局部回路充分性 × 控制深度」的方向性交互，本文的跨领域版本作废；若只在一个领域获得支持，降格为领域特定模型。

第二条是本文自己新加的，也是更严的一条：

到 2029 年 12 月 31 日，若第七节第二小节那个「提示逐字相同、只差一次确认」的对照被做过而未出现甲乙差异，或差异被等量中断条件完全解释，则本文的核心主张不成立，应当被并入既有的信息负荷解释——届时本文不再申辩，也不再要求「换个任务再试一次」。

这一次作废了多少条

本文在写作过程中作废掉的判断，按顺序记在这里：

1. 作废「外部优于内部」：外部提示同样可以逐项接管，内部提示也可以停在框架层；方向不是原因。
2. 作废「控制越少越好」：减掉的必须是内容权，框架权减掉就成了放任性漂移。
3. 作废「框架越清楚越好」：框架密到需要在行动中逐项核对，它就已经是内容接管；框架清晰度不能作为独立的加分项累加。
4. 作废「低充分性必须上收内容权」：母语习得这个反例迫使本文把「必须」削成「可能有校准价值，取决于误差可不可见」。
5. 作废「充分性高就该完全放手」：高风险领域迫使本文改成「默认归属加可撤回」，中断线本来就在框架权里。
6. 作废「内容权只有两个可能的持有者」：合奏与球队里它可以落在同伴身上，判据改为「是否需要跨越一次等待」。
7. 作废本文对自己的合格自评：本文无法用自己的判据判定自己——那套判据测的是流程，不是作品。

注 释

1. 本文所称「内容单元」按任务而定，其边界可争。这是本文最软的一处：单元划得越细，内容权看上去越容易被上收；划得越粗，两格越容易混。实验中应以「一次可被单独确认的操作」为准，并在报告中写明口径。
2. 第一节三项研究的数字与方向依据公开发表的原文；本文对它们的机制解释由本文负责，不应归给原作者。三项研究各有其限制，均已在第二节写明：动作学习一族存在发表偏倚争议，即兴研究样本为六人且用固定效应分析，正念研究不是严格的随机前后测。

3. 第七节的三组实验**均未执行**。本文交出的是设计与预测，不是结果。第七节第二小节的实验成本极低（一台电脑、两组被试、一段固定序列），本文希望它先被做。
4. 第八节所列六种说法的转述依据其公开表述的核心结论，未逐条核对原著页码。若某一条的原始表述与本文的转述不符，分界线应按原始表述重画，而不是按本文的方便重画。
5. 「局部回路」不指某个脑区，而指在特定时间尺度上完成感知、行动、误差检测与快速修正的一组耦合过程。这个概念刻意保持功能性，以避免把复杂状态还原成单一部位；代价是它的边界依赖「使用者」与「任务」的划法（第十节第二小节第三条）。

参考文献

- Beaty, R. E., Benedek, M., Kaufman, S. B., & Silvia, P. J. (2015). Default and executive network coupling supports creative idea production. *Scientific Reports*, 5, 10964.
- Beilock, S. L., & Carr, T. H. (2001). On the fragility of skilled performance: What governs choking under pressure? *Journal of Experimental Psychology: General*, 130(4), 701–725.
- Endsley, M. R., & Kiris, E. O. (1995). The out-of-the-loop performance problem and level of control in automation. *Human Factors*, 37(2), 381–394.
- Farb, N. A. S., Segal, Z. V., Mayberg, H., Bean, J., McKeon, D., Fatima, Z., & Anderson, A. K. (2007). Attending to the present: Mindfulness meditation reveals distinct neural modes of self-reference. *Social Cognitive and Affective Neuroscience*, 2(4), 313–322.
- Fitts, P. M., & Posner, M. I. (1967). *Human Performance*. Brooks/Cole.
- Guadagnoli, M. A., & Lee, T. D. (2004). Challenge point: A framework for conceptualizing the effects of various practice conditions in motor learning. *Journal of Motor Behavior*, 36(2), 212–224.
- Hersey, P., & Blanchard, K. H. (1969). Life cycle theory of leadership. *Training and Development Journal*, 23(5), 26–34.
- Kalyuga, S., Ayres, P., Chandler, P., & Sweller, J. (2003). The expertise reversal effect. *Educational Psychologist*, 38(1), 23–31.
- Limb, C. J., & Braun, A. R. (2008). Neural substrates of spontaneous musical performance: An fMRI study of jazz improvisation. *PLoS ONE*, 3(2), e1679.
- Masters, R. S. W. (1992). Knowledge, knerves and know-how: The role of explicit versus implicit knowledge in the breakdown of a complex motor skill under pressure. *British Journal of Psychology*, 83(3), 343–358.
- McKay, B., Corson, A. E., Seedu, J., De Faveri, C. S., Hasan, H., Arnold, K., Adams, F. C., & Carter, M. J. (2024). Reporting bias, not external focus: A robust Bayesian meta-analysis and systematic review of the external focus of attention literature. *Psychological Bulletin*, 150(11), 1347–1362.

- Onnasch, L., Wickens, C. D., Li, H., & Manzey, D. (2014). Human performance consequences of stages and levels of automation: An integrated meta-analysis. *Human Factors*, 56(3), 476–488.
- Parasuraman, R., Sheridan, T. B., & Wickens, C. D. (2000). A model for types and levels of human interaction with automation. *IEEE Transactions on Systems, Man, and Cybernetics — Part A*, 30(3), 286–297.
- Wegner, D. M. (1994). Ironic processes of mental control. *Psychological Review*, 101(1), 34–52.
- Wood, D., Bruner, J. S., & Ross, G. (1976). The role of tutoring in problem solving. *Journal of Child Psychology and Psychiatry*, 17(2), 89–100.
- Wulf, G., Chiviawsky, S., Schiller, E., & Ávila, L. T. G. (2010). Frequent external-focus feedback enhances motor learning. *Frontiers in Psychology*, 1, 190.
- Wulf, G., & Lewthwaite, R. (2016). Optimizing performance through intrinsic motivation and attention for learning: The OPTIMAL theory of motor learning. *Psychonomic Bulletin & Review*, 23(5), 1382–1414.

人机分工声明

本文由王德生与 Claude 共同署名。**承重判断、三个领域的选定与原始稿（控制层级迁移假说 V1.0, 2026-08-01）由王德生提出并完成**；本文是在此基础上按一次独立评分意见重写的加固版。

相对原稿，本文新增或改动的部分是：第三节把「内部注意」拆成四条轴并用它逐条消解三场冲突；**第四节第二小节把「内容权」与「指导量」分开，并给出信息量完全匹配的操纵**；第四节第四小节写出并接受了循环风险的两条自缚；第四节第六小节更正了原稿把框架清晰度当作独立加项的写法（改为与权限归属相乘）；**第六节第四小节的母语反例与由此削掉的一条主张**；第六节第五小节的多主体托管；**第八节前三小节与教学设计、管理学、人因工程三族的正面分界，以及各自的判决性对照**——这三处是评分意见指出的、原稿完全未处理的最接近的既有说法；第七节把核心实验改成提示逐字相同的设计并加了等量中断的控制条件；第十三节的作废清单；第十一节末段的自认。

Claude 负责上述新增部分的论证与全部文字。第七节的三组实验均未执行；第八节六种说法的转述依据其公开表述的核心结论，未逐条核对原著页码。**本文未印任何评分**——评过它的人不能给自己改过的稿子发合格证。

字数 纯汉字约 1.7 万 **版本** v2.0 · 2026-08-01（基于王德生原稿 V1.0 重写）

本文由三个分属不同学科、且互相冲突的理论体系碰撞而成。全部来源均为站外公开文献，可自行核对。 作者：王德生 + Claude · SDE Universes · 学科通融