

什么是智能

在电光石火间做出判断

竹小竹8167

费曼式写法 · 由多个 AI 代理撰写与互相审校

导读

在电光石火间做出判断

先别急着背定义。先想象一位将军。

地图是旧的，情报互相矛盾，左翼失联两个小时了。他没有时间开会，必须现在就说：打，还是撤。认知科学家侯世达说，这就是智能的样子——**在瞬息万变的局面里，电光石火之间做出判断，然后动手。**

这本书从一个反直觉的对比开始：一个能把整本书背下来的学生，遇到一道没见过的活题却愣住了；而一个只看了五秒钟棋盘的象棋大师，能立刻说出该怎么走。差别不在记忆里装了多少，而在装的是什麼——大师装的是「长成这样的局面，该这么办」的字条。

顺着这条线，我们会一路拆开判断这台机器：为什么所有判断的第一步都是「认出这是哪一类情况」；为什么心理学家发现把题混在一起做，比一类类刷题效果好几倍——因为它偷偷练了「该用哪一招」；为什么听过堡垒攻坚故事的人，解不出同结构的肿瘤射线题，直到有人提醒一句「这两个是一回事」。

最后，所有的线索会收拢成一句话：**在不同中发现相同，在相同中发现不同。**前半句是归类，把新问题认成老问题；后半句是分辨，在两个看起来一样的局面里，揪出那个决定成败的差异。将军靠它，医生靠它，品酒师靠它——而它，是可以练的。

这本书写给每个天天在做「小战场判断」的人：项目要不要砍、这条消息是不是谣言、孩子这次哭是不是真出了事。你没有参谋部，但你可以练格子。

第一章 · 战场上的电光石火

先别急着查字典，也别背任何定义。先想象一个人。

一位将军，站在战场上。地图是半小时前的，侦察兵的报告彼此矛盾，通信兵刚跑进来喊了一句什么就被打断。敌人在动，地形在变，自己这边的左翼已经两个小时没有消息了。他没有时间开会，没有时间把所有信息核对一遍，他必须**现在**就说出一句命令：打，还是撤；从左翼压上去，还是收缩防线。

认知科学家侯世达（Douglas Hofstadter）在《表象与本质》里用的大概就是这样一个例子来说明什么是智能——大白话说，就是**在瞬息万变的局面里，电光石火之间对形势做出判断、并立刻采取行动的能力**。

大白话

大白话：智能不是你脑子里存了多少东西，而是局面扑到你脸上的那一秒，你能不能判断出「现在是怎么回事、该怎么办」，并且动手。

注意这四个词

这个例子里藏着四个关键词，值得一个个掰开看。

瞬息万变。局面是活的，不是考卷上印好的题。你思考的每一分钟，局面本身都在变。一道静态的谜题，再难也不是智能的真正考场；真正的考场是流动的。

电光石火。时间是有成本的。将军可以花三个月做研究吗？可以，但那时候仗已经打完了。判断的价值和它的速度绑在一起——迟到的正确判断，约等于错误判断。

判断。注意，不是「计算」，不是「回忆」，是判断。信息永远不全，选项永远不清，没有一个公式能代入。将军做的不是求解，是在模糊中认出一个形状：「这个局面，像是要崩。」

行动。判断不出手的将军，和没有判断的将军，结果一样。智能最终要落到动作上，指挥棒挥下去才算数。

一个反例：考试状元

现在换一个场景。一个学生，把整本教科书背得滚瓜烂熟，你随便翻一页他都能接着往下背。考试来了，题目是书里原话的填空——满分。但换一道题：给一段他没见过的生活场景，问他该用书里的哪个原理。他愣住了。

他缺的不是知识，是判断——大白话说，就是「认出眼前这个情况到底属于哪一类、该调用哪个办法」的能力。书里的东西他都有，但「此刻该用哪一条」这个开关，他没练过。

这就是这本书要反复敲打的一件事：**知道**和**会判断**是两种东西，而我们日常说的「聪明」，常常把这两样混在一起。

这个问题为什么值得写一本书

你可能会说：将军的例子太极端了，我又不带兵。

但你每天都在做小号的战场判断：老板突然问「这个项目还要不要投人」，你有三十秒；孩子哭着说不想去上学，你要判断这是撒娇还是真的出了事；股市跳水，你要判断这是噪音还是趋势。这些场景共享着和那四个字一模一样的结构——**局面在变、时间有限、信息不全、必须出手**。

所以这本书要回答的贯穿问题是：**这种电光石火间的判断力，到底是什么东西？它能不能练？**

接下来的章节会一层层拆开它。我们会先看清「记住一切」为什么不够（第二章），再发现判断的第一步其实是一个你几乎意识不到的动作（第三章），然后借《认知天性》里的一个著名学习实验，看到这个动作原来是可以专门练的（第四章）。最后，所有的线索会收拢到一句话上——一句话说清智能最值得练的内核是什么。

将军没有时间读完这本书。但你还有。

第二章 · 记住一切不等于智能

上一章立了个靶子：知道和会判断是两种东西。这一章来拆一个更顽固的直觉——「聪明的人就是记性好、知道得多」。

这个直觉太自然了。我们从小到大的考试几乎都在考记忆，谁背得多谁就是「学霸」。于是「智能=大容量存储器」这个等式就悄悄长在脑子里了。要拆掉它，最好的工具是一个著名的心理学实验。

五秒钟，你能记住一盘棋吗

上世纪六七十年代，心理学家德格鲁特（Adriaan de Groot）以及后来的蔡斯和西蒙（Chase & Simon）做过这样一组实验：给国际象棋大师和初学者看一个棋盘，只看**五秒钟**，然后撤走，让他们凭记忆把棋子摆回去。

结果是：如果棋盘上摆的是一局真实对局里出现过的局面，大师能摆回二十多个棋子，初学者只能摆回四五个。大师的记忆碾压初学者——看起来「记性好=水平高」的等式成立了。

但实验还有第二半。这次棋盘上摆的是**随机乱放**的棋子，同样五秒钟。结果大师的优势消失了——他们也只能摆回四五个，和初学者差不多。

大白话

大白话：大师的「好记性」只在棋局有意义的时候才有。棋子一乱摆，大师立刻被打回原形。所以他记住的根本不是一个个棋子的位置，而是别的东西。

大师记住的是「模式」

那个别的东西，心理学家叫它组块（chunk）——大白话说，就是「一个有意义的小团体」。初学者眼里是二十四个棋子，大师眼里是四五个「局面短语」：这是王翼堡垒，那是悬兵结构，这两个车在同一列要出事了。就像你读一句话，记住的不是二十六字母，而是几个词。

关键在下一步。大师认出「这是王翼堡垒」这个模式的同时，这个模式自带了**判断**：王翼堡垒意味着王比较安全、意味着该考虑从另一侧进攻。模式不是仓库里的存货，是一张写着「该怎么办」的字条。大师看一眼棋盘，五秒内完成的不是记忆表演，而是**归类**——把这个局面认成「我见过的那一类」，于是对策自动浮出来。

回头看第一章的将军：他脑子里也没有整本《战争论》在滚动播放，有的是成百上千个这样的「局面短语」。电光石火之间，他把眼前这一团混乱认成了其中某一张字条。

背书为什么会失灵

现在可以回答「学霸为什么会在活题面前发愣」了。记忆和判断的差别，在于你存进去的是什么：

- 如果存的是**孤立的事实**——一条条互不相关的句子，那你只能在「原题重现」时把它调出来。局面稍微变个样子，你就不知道哪条句子管用了。
- 如果存的是**带判断的模式**——「长成这样的局面，大概是这样的一类问题，该这么办」，那局面再怎么变形，只要骨架还在，你就认得出来。

同样是「记住了很多东西」，一种存的是词典，一种存的是字条。词典在考场上无敌，字条在战场上救命。

这也解释了为什么「读书很多却处理不了新情况」的人那么多：他们把书读成了词典。每页都是孤立的事实，从没问过自己「这一页讲的，是哪一类局面的字条」。书合上的那一刻，知识就死在了书页之间。

一个提醒

但别走极端，以为记忆不重要。大师能认出模式，前提是他下了成千上万盘棋，那些组块是实打实存进去的。**没有存货就没有模式**——巧妇难为无米之炊，空脑壳做不出判断。这一章真正要说的只是：存货的形式决定它能不能在电光石火间被用上。孤立的事实是死库存，带判断的模式才是活资产。

那么问题来了：「认出一类局面」这个动作，到底是怎么发生的？它发生得有多快、多自动？下一章我们凑近了看——你会发现，判断的第一步几乎快到你意识不到它存在。

第三章 · 判断的第一步是归类

上一章的结尾说，大师看一眼棋盘，五秒内完成的动作是「归类」——把眼前这团混乱认成「我见过的某一类」。这一章凑近了看这个动作，你会发现两件事：它快得惊人，而且它决定了后面的一切。

草丛动了一下

你走在山路上，脚边的草丛突然动了一下。在你「思考」之前，你的身体已经跳开了半米。这时候你的大脑完成了一个什么动作？

它完成了一个归类——大白话说，就是把眼前这个模糊的刺激，扔进脑子里某个已有的格子里。草丛动了，格子只有两个候选：「蛇」或者「风」。你的大脑在零点几秒内选了「蛇」，于是肾上腺素的指令先发出去了，腿先动了，至于到底对不对，回头再说。

注意这个顺序：**先归类，后行动**。归类成「蛇」，行动是跳开；归类成「风」，行动是无视。同样一个「草丛动了」，归进不同的格子，引出的是完全不同的动作。判断的全部后果，在归类落格的那一瞬间就已经定了。

医生看病的真相

再换一个安静的场景。病人描述了一堆症状：胸口闷、出汗、左臂发麻。有经验的医生听到第三句，心里已经冒出一个格子：「心绞痛？先按心梗排查。」接下来她问的问题、开的检查，全是在验证或推翻这个归类。

医学生和医生的差别就在这里。医学生把症状背得很熟，但面对一个真人，他要从几百种病里挨个对——像查词典。医生脑子里装的则是格子：「长成这样的一簇症状，大概率是这一类。」诊断这个词听起来高深，拆开了就是**归类，然后按这一类该做的做**。

大白话

大白话：所谓判断，第一步永远是回答「这是哪一类情况」。这个答案一下，后面查哪本书、用哪一招、防哪一手，全都跟着定了。归类错了，后面越努力错得越远。

归类是自动的，这正是问题所在

最要命的是：归类这个动作**快到你意识不到**。你不会觉得「我刚才归类了一下」，你觉得「我直接看到了事实」。草丛动了，你觉得「我看见了危险」，而不是「我把一个模糊刺激归进了蛇这个格子」。

这就是为什么归类错了很难察觉。一个项目经理看到进度延期，自动归进「团队在偷懒」的格子，于是加人、催工、盯日报——但如果真实原因是需求方一直在改需求呢？他从第一步就归错了格子，后面所有动作都在让事情更糟，而他觉得自己「明明一直在处理问题」。

顺便回答第一章留下的一个悬念：将军的「电光石火」为什么可能？因为归类一旦被练成自动的，就不再占用思考时间。慢的是从零分析，快的是认格子。电光石火，是归类的速度。

格子的质量

既然一切取决于归类，那问题就收敛成两个：你的脑子里**有没有足够的格子**，以及你的格子**画得准不准**。

- 格子不够，遇到新情况你只能硬塞进一个不相干的旧格子——手里拿着锤子，看什么都像钉子。
- 格子画得不准，两个本质不同的局面被你当成一类，或者同一类局面被你拆成两堆——招招都打偏。

这两个问题，正好就是这本书的第五、第六章要分别处理的：怎么在不同里认出相同（格子从哪来），怎么在相同里认出不同（格子的边界画在哪）。

但在去那里之前，要先回答一个更实际的问题：格子是怎么练出来的？靠多做题吗？下一章会介绍一个有点反直觉的学习实验——它发现，把题混在一起做，竟然比一类一类刷，更能练出归类的能力。

第四章 · 交叉练习的秘密

上一章留了个问题：归类的格子怎么练出来？直觉的答案是把同一类情况反复刷——今天专练「蛇」，明天专练「风」，一类练熟了再换下一类。学校里就是这么教的：学完一个公式，做二十道一模一样的题。

《认知天性》(Make It Stick，作者罗迪格、麦克丹尼尔和布朗，几位研究记忆的心理学家)整本书几乎都在跟这个直觉对着干。他们的核心发现之一叫交叉练习——大白话说，就是**把不同类型的问题混在一起练，而不是一类刷完再刷下一类**。

一个让数学老师尴尬的实验

心理学家罗勒和泰勒 (Rohrer & Taylor) 找了一群学生学解四种不同几何体的体积题。学生分成两组：

- **集中组**：把第一种题型的四道题做完，再做第二种的四道……一类一类来，规规矩矩。
- **交叉组**：同样十六道题，但四种题型随机打乱混着来。做每一道之前，都不知道它属于哪一型。

练习过程中，集中组的成绩漂亮得多——这很合理，刚讲完公式立刻套，当然顺。交叉组练得磕磕绊绊，错得多，学生自己也觉得「没学好」。

但一周之后突击测验，结果反过来了，而且不是小反超：**交叉组的正确率大约是集中组的三倍**（实验中约六成对约两成）。练的时候感觉好的，考的时候垮了；练的时候感觉差的，真正学会了。

大白话

大白话：一类一类刷题，每道题都提前告诉你「我是哪一型」，你练的只是套公式的手。混在一起练，每道题都得先自己认「这是哪一型」——你练的是判断该用哪一招。而考场上、战场上，恰恰没有人在题目前面标好题型。

关键在这一句话

注意交叉练习真正多练出来的东西。两组人练的题目总量一样，四种公式的用法也都练了。唯一的区别是：**交叉组每做一道题，都多做了一个动作——先归类**。「这道题是哪种几何体？该套哪个公式？」这个动作，集中组一次都没练过，因为答案被题号的顺序免费剧透了。

这和第三章对上了：真实世界的判断，第一步永远是归类。集中练习把这个第一步外包给了练习册的编排顺序，于是练出来的招式全是「离了拐棍就不会走」的。交叉练习则逼着你每一道题都走一遍完整链路：**认类型 → 选招式 → 出招**。它练的不只是招，是选招。

类似的结果在别处也复现过。心理学家科内尔和比约克（Kornell & Bjork）让受试者学认不同画家的风格：一组一个画家的画集中看，一组混着看。结果混着看的组后来在认「没见过的新画」时表现明显更好——认新画，本质就是归类。还有棒球击球手的实验：训练时随机混合投快球、曲线球、变速球的打者，实战成绩优于一种球练饱再换下一种的打者。

为什么大家都选错的方法

如果交叉练习这么好，为什么学校和自己刷题的人都爱用集中练习？因为**感觉会骗人**。集中练习当下顺滑、正确率高，给你一种「我掌握了」的错觉；交叉练习当下痛苦、错误多，给你一种「我还没学会」的错觉。偏偏学习的效果要用延迟的、变化的考验来测，而我们的感觉只看当下。心理学家比约克给这类现象起了个名字：合意困难（desirable difficulty）——大白话说，就是「练的时候制造点恰当的麻烦，学的时候难一点，用的时候顺得多」。

所以这一章的结论可以写得很直白：**练技能时别把题型标好再练，要混着练、逼自己每次都先认类型**。学数学，混合着做题；学外语，混合着抽单词卡而不是按字母表背；学编程，别照着教程敲，去找不知道属于哪一章的真实问题。

但「认出这是哪一类」到底是怎么做到的？凭什么你能在一个没见过的问题里，看出它和一个老问题的相同之处？这就是下一章的事——智能的第一块基石：**在不同中发现相同**。

第五章 · 在不同中发现相同

前面四章把结论推到了这里：判断的第一步是归类，归类靠脑子里的格子，格子靠练。这一章和下一章分别拆格子的两面。先看第一面：**在不同中发现相同**。

这句话听着玄，其实每个人都做过。你第一次骑电动车就会骑，因为你骑过自行车；你第一次用一个新的打车软件不用看说明书，因为你用过别的打车软件。自行车和电动车、两个长得完全不同的软件界面，表面处处不同，你却在它们里面看见了同一个骨架。这个动作叫抽象——大白话说，就是**扒掉表面的细节，只留里面的结构**。

侯世达：类比是思考的燃料

侯世达在《表象与本质》里把这件事推到了极致。他和桑德尔（Emmanuel Sander）论证：类比不只是诗人修辞的工具，它是**认知的发动机**。你做的每一次归类，本质上都是一次类比——「眼前这个东西，和我见过的那一类东西是一回事」。没有类比能力，每一次经历都是孤立的，世界上没有任何新情况能被认成老情况，你也就永远只能从零开始。

大白话

大白话：脑子里的每个格子，都是用「这和那个一样」的连线织出来的。见过的局面越多、连的线越多，新情况一出现就越容易被某根线钩住——钩住了，就叫「我见过这个」。

一个著名实验：堡垒与肿瘤

「在不同中发现相同」有多难、又多重要，心理学家吉克和霍利约克（Gick & Holyoak）的一个经典实验说得最清楚。他们先给受试者讲一个故事：

一位将军要攻打一座堡垒，通往堡垒的路有很多条，但每条路上都埋了地雷，大队人马通过会引爆。将军于是把军队拆成许多小队，从各条路同时进发，最后在堡垒前会师，一举拿下。

听完故事，研究者请受试者解另一道题（心理学家邓克尔的经典难题）：病人胃里有个肿瘤，可以用射线杀死，但高强度的射线会把沿途的健康组织一起杀死，低强度又杀不死肿瘤。怎么办？

解法就藏在将军的故事里：把射线拆成许多束低强度的，从不同方向射入，在肿瘤处会聚成高强度。表面上一个讲打仗一个讲治病，骨架一模一样：**分散进入、会聚发力**。

结果呢？听过故事的受试者里，只有约三成自己想到了这个解法。但如果研究者**提示一句**「刚才那个故事可能对解题有用」，比例立刻跳到七成以上。

这个落差太说明问题了：答案就在他们脑子里，骨架他们也刚刚读过，缺的只是那根「这和那个一样」的连线。在不同中发现相同，不是知识问题，是**连线问题**。

抽象能力是怎么长的

怎么让自己更容易连上这根线？研究给出的路子和第四章一脉相承：

- **同一个骨架，看多个表面**。吉克和霍利约克后来发现，让受试者读**两个**同构的故事（比如将军故事加一个消防员救火的同构故事）并比较它们，迁移率大增。一个例子容易让人记住表面，两个例子逼着大脑问「这两个到底有什么共同点」——那个共同点，就是骨架。
- **自己说出骨架**。被动地看相似，不如主动地问「这两件事的共同点是什么」。问完用自己的话写出来，格子才算真的画进了脑子。
- **换领域见面**。同一个结构在数学里见一次、在厨房里见一次、在管理里见一次，它就从「某章的知识」变成「你的格子」。这也是为什么交叉练习有效——它天然让不同表面混在一起。

小心：这一面只是半个格子

但「在不同中发现相同」有个天生的危险：连线一旦连上，人就倾向于**只看相同，不看不同**。把所有新情况都塞进老格子，是会出大事的——2008 年把房贷打包成证券的人，以为自己发现了「和以前所有证券一样」的东西，结果那个「不同」炸掉了全球金融系统。

格子不光要有，还要有边界。两个局面看起来一样，到底是什么让它们不一样？哪个差异是无关紧要的噪声，哪个差异决定成败？这就是格子的第二面，也是下一章：**在相同中发现不同**。

第六章 · 在相同中发现不同

上一章说了格子的一面：在不同里认出相同。这一章说另一面，也是更容易被忽略的一面：**在相同中发现不同**。

侯世达那句话的完整版应该这样读：格子用来认出「这和我见过的某类一样」，但格子的边界用来发现「等等，这一个和我见过的那些，差在了一个要命的地方」。智能有一半藏在「它们一样」里，另一半藏在「但它们不一样」里。

新手看相同，专家看差异

品酒师刚入门时，红酒只分红酒；练上几年，他能从两杯颜色几乎一样的酒里说出年份、产区、橡木桶。放射科医生看胸片，新手看到「一片阴影」，老手看到「边缘毛刺状的阴影」——边缘的毛刺，就是良性结节和恶性肿瘤之间那个决定生死的差异。

注意专家做的事：不是把「阴影」这个格子画得更细这么简单，而是知道**哪些差异算数**。一片阴影有一百个维度可以比较——大小、位置、形状、颜色深浅——其中九十几个维度上的差异是噪声，只有「边缘形态」这类少数维度上的差异会改变诊断。分辨，大白话说，就是**在一堆看似一样的局面里，找出那个会改变结论的差异维度**。

大白话

大白话：格子告诉你「往哪格扔」，分辨告诉你「这次别急着扔——先看看那个要命的细节上，这次是不是不一样」。

过度类比的坑

上一章结尾提了 2008 年的金融系统。这里再举一个更近的：过去十几年，创业者最爱用的句式是「某某领域的 Uber」——打车的 Uber 成了格子，凡是「把闲置资源和需求对接」的生意都往里扔。上门美甲的 Uber、家政的 Uber、共享充电宝的 Uber。骨架确实一样：平台、匹配、抽成。

但活下来的没几个。因为决定成败的差异不在骨架上，在细节上：打车的供需双方都是高频、标准化、即时决策，而上门美甲是低频、非标、重信任。格子的骨架相同，要命的维度全不同。只练了「发现相同」、没练「发现不同」的人，会把每个新情况都看成「又一个老机会」。

反过来，真正的高手问的问题总是反着的：「这次和上次，哪里不一样？」老医生听到「和上个月那个病人一样的症状」，第一反应是找反例；老将军看到「这和当年那场仗一样的地形」，第一反应是问敌人的指挥官换没换人。

找不同，恰恰也是混着练出来的

这里有个漂亮的回环。还记得第四章认画家风格的实验吗？科内尔和比约克发现，混着看不同画家的画，比集中看一个画家的画，学得好得多。为什么？现在可以给出更深一层的答案了：集中看莫奈，你学会的是「莫奈长什么样」；莫奈和马奈混着看，你被迫每一幅都回答「这幅是莫奈还是马奈」——你练的正是在最相似的两个东西之间找那个决定性的差异。

交叉练习练的不只是归类，还有归类的边界。它天然地把你按在两个最像的格子之间摩擦：这道题到底用公式 A 还是公式 B？这幅画到底是莫奈还是马奈？每一次摩擦，都是在给格子的边界描线。

所以练「找不同」有个很具体的方法：专挑最像的两个东西对比着学。学英语就把 affect 和 effect、complement 和 compliment 放在一起比；学医学就把症状最像的两种病放在一起比；做投资就把「上一次这么涨的之后崩了」的案例和「这么涨的之后继续涨」的案例放在一起比。差异维度，只在对比中显形。

两面合一

把第五、六章合起来，一个合格的格子长这样：中间是骨架（「这些局面是一类」），边缘是带警戒线的边界（「但只要这个维度上不一样，就归另一类」）。只会发现相同的人，格子有中间没边界，处处生搬硬套；只会发现不同的人，格子全是碎边界，看什么都是「这次不一样」，永远从零开始。将军的电光石火，是两者同时在电光石火间完成：先认出「像哪一类」，再扫一眼「要命的维度上有没有不一样」，然后下令。

但还有一个现实问题没解决：战场上的将军根本没时间把所有信息收集齐再判断。信息永远不全，时间永远不够——那专家到底是怎么在缺信息的情况下敢下手的？下一章，去看一群真正在火场里做生死判断的人。

第七章 · 信息不全时怎么办

教科书里的决策长这样：列出所有选项，给每个选项列出利弊，打分，选总分最高的。管理学课本、决策树软件、「用 Excel 做选择」的人生建议，全是这个套路。它假设你有完整的信息、充足的时间、清晰的选项。

但火场里的消防队长三样都没有。他们是怎么决策的？心理学家加里·克莱因（Gary Klein）真的去研究了他们——这一章基本是他《力量的源泉》（Sources of Power）里讲的事。

厨房火灾里的「第六感」

克莱因收集过一个著名的案例。一位消防队长带队扑灭一处厨房火灾：火在厨房烧，水枪喷上去，火势却不退，而且他发现两件事不太对——**屋里比预想的热，火却比预想的安静**（真正的火源在别处闷烧时，现场反而会显得安静）。他说不清为什么，就是觉得不对劲，于是下令全员撤出。他们刚退到街上，地板塌了——火根本不在厨房，在厨房底下的地下室里，已经把地板烧空了。

事后别人问他怎么知道的，他说了类似「直觉」「第六感」的话。克莱因不同意这个解释。他一层层追问，发现所谓第六感拆得开：队长干了十几年，脑子里存了几百个火灾模式；「热但不安静」这个组合不在任何一个正常厨房火灾的模式里——**现场和格子对不上，警报就响了**。这不是超能力，是第六章说的「在相同中发现不同」在高速运转：表面看这就是一次普通厨房火灾，但「热」和「静」这两个维度对不上，而要命的恰恰在这两个维度上。

专家根本不比较选项

克莱因访谈了大量消防队长、护理主任、棋手之后，发现了一个让决策理论难堪的事实：**这些专家几乎从不同时比较多个选项**。他们不列利弊表。他们做的事是：认出这是什么局面 → 脑子里浮现出第一个可行方案 → 在脑子里把这个方案「放一遍电影」推演一下 → 行就干，不行再想下一个。

他给这个模式起了个名字：识别启动决策（recognition-primed decision，RPD）——大白话说，就是「先认局面，方案跟着认出来，一次只验一个」。为什么不比较？因为**比较是**

在没有格子的前提下才需要干的事。新手分不清哪个方案好，才需要把所有方案摆在一起打分；专家的格子本身已经按「这类局面通常什么管用」排好了序，第一个冒出来的方案，平均而言就是最好的那个。

大白话

大白话：利弊表是给没经验的人用的拐杖。专家不比较选项，是因为他的经验已经替他把选项排过序了——认出局面，答案自己浮上来，然后用脑子里的模拟器快速过一遍：这招在这儿行得通吗？

信息不全？专家要的本来就不是「全」

这回答了这一章标题的问题。新手觉得信息不全就没法判断，因为他判断的方式是「收集信息→逐项分析」，信息一缺，机器就卡死。专家判断的方式是「认出模式→模式补全缺失的信息」。消防队长不知道地下室有火，但「厨房火灾」这个模式自带预测：「应该越喷越小、应该嘈杂」。预测和现实对不上的那一刻，缺失的信息反而以「哪里不对劲」的形式暴露了出来。

这就是模式最厉害的用法：**它不只告诉你该怎么办，还告诉你正常情况下应该看到什么。**于是信息不全不再可怕——你带着模式进场，让现实和模式对表，对不上的地方就是你需要补看的地方。将军看战场也是这样：他不是把战场「全部看清」，而是带着「这种地形通常怎么打」的模式，专看和模式对不上的那几处。

那直觉是不是可以无脑信

读到这儿容易产生一个危险的结论：专家直觉这么神，那我跟着感觉走就行了。慢着。消防队长的直觉灵，是因为他所在的领域给他送了两样礼物——而很多领域一样都不送。是哪两样，缺了它们直觉会变成什么？这是下一章的事：直觉的边界。

第八章 · 快思考什么时候会骗你

上一章把专家直觉夸了一通。这一章来泼冷水，而且泼冷水的人分量足够——丹尼尔·卡尼曼（Daniel Kahneman），研究人类判断偏误拿了诺贝尔奖的那位，写过《思考，快与慢》。有趣的是，他和上一章的主角克莱因曾经是学术界著名的对头：卡尼曼认为直觉多半是偏见，克莱因认为直觉是专家的本事。两人吵了很多年，后来做了一件学术界罕见的事：坐下来合作写一篇论文，把分歧解决掉。

解决的结果，一句话就能说清：**直觉靠不靠得住，不取决于人，取决于环境。**

直觉有效的两个条件

卡尼曼和克莱因在联合论文（《直觉专长的条件》，2009）里达成了一致：专家的直觉值得信任，当且仅当两个条件同时满足。

条件一：环境本身有规律。大白话说，就是这个领域里「同样的因」大概率导致「同样的果」。火灾的物理规律是稳定的，棋的规则是死的，同样的心电图模式背后是同样的病——在这种高有效性环境（大白话：有规律、可预测的环境）里，模式是真的存在，认模式才有意义。

条件二：你有机会快速、清楚地得到反馈。消防队长每次行动，几分钟内就知道对错；护士的判断，病人的变化马上给答案。反馈快而且不含糊，脑子里的格子才能一次次被修正校准。

大白话

大白话：直觉是「用过去练出来的格子套现在」。如果那个领域根本没什么稳定规律可练，或者你做完一个判断十年才知道对错，那你的格子就是拿噪声练出来的——看起来很专业，其实不灵。

反例：股市和预言家

现在看两个条件都不满足的领域。心理学家泰特洛克（Philip Tetlock）做过一个长达二十年的研究：收集了几百位政治、经济专家的几万个预测——苏联会不会解体、某国会不会打仗、经济会不会衰退。结果：专家的平均准确率，和「随机瞎猜」差不多。而且越自信、越有名的专家，并没有更准。

股市择时也一样：短期股价基本是噪声，你今天买明天涨，这个「反馈」跟你的判断对不对几乎没有关系——但感觉上它像反馈。于是人拿着假反馈练了二十年，练出一套无比自信的假格子。

注意这和第四章「感觉会骗人」是同一个病根的升级版：第四章是练习方式骗人（集中练习感觉好），这里是**自信本身骗人**。卡尼曼指出过一个扎心的事实：人的自信程度，取决于他脑子里那个故事讲得顺不顺，而不取决于那个故事对不对。故事越连贯，感觉越笃定——哪怕信息是缺的、环境是没规律的。他给这种倾向起了个名字叫 WYSIATI（What You See Is All There Is）——大白话说，就是「你只看见眼前这些料，就以为这些就是全部」。

一张自检清单

所以「该不该信直觉」不该凭感觉回答，该用两个条件检查。下次你想跟着感觉走，先问两句：

- **这个领域有规律吗？** 火灾、下棋、看病、写代码调试——有。明天股价、明年哪款产品会爆、十年后世界格局——基本没有。规律越少，直觉越接近抛硬币。
- **我收到过干净快速的反馈吗？** 注意「干净」这个词：结果要明确地由你的判断造成，而且你要真的去对了答案。很多人所谓的「二十年经验」，其实是「一年经验重复二十次」，或者「从没回头核对过对错的二十年」。

两句都答「是」，放心电光石火；有一句答「否」，就该慢下来，把直觉降级为「一个待验证的假设」，老老实实收集信息。

将军怎么办

你可能发现一个问题：第一章的将军，处境的恰恰是规律模糊、反馈迟缓的环境——战争里同样的战术这次赢下次输，打完一仗要很久才知道当初判断对不对。那将军的电光石火还算

数吗？

算数，但原因微妙：战场的**底层规律**（地形、后勤、士气、火力对比）是稳定的，真正不可预测的是敌人那个人。所以优秀指挥官的直觉锚定在稳定的那一半上，同时为不稳定的那一半留后手——这正是「判断」和「赌博」的区别。至于普通人怎么在真假反馈里把格子练好，下一章给出具体的练法。

第九章 · 判断力是可以练的

前面八章把判断拆开了：它是归类，归类靠格子，格子靠经验加反馈喂出来，而反馈有真假。这一章把零件装回去，回答全书问题的一半：**具体怎么练？**

原则只有一句：把「归类」这个动作从幕后拉到台前，单独练它。下面四个方法，每一个前面都已经见过原理，这里只说怎么做。

方法一：混合着练，别一类类刷

这是第四章交叉练习的直接应用，但要点值得再拧紧一圈：练习的关键变量不是「练了多少题」，而是**每道题是否逼你先认类型**。

- 学数学：别按章节刷，把整本书的题打乱做成卡片随机抽。抽出来先不看题，先答一句「这是哪一型」。
- 学编程：别顺着教程走，去真实问题里泡——真实 bug 从不标注「本题属于第三章」。
- 学任何技能：只要发现自己在「同一招连续用十次」，就主动打断，换个场景。

练的时候会 slower、更挫败——这是合意困难，是好事。记住实验里那组数据：练的时候感觉差的那组，一周后成绩是感觉好的那组的三倍。

方法二：写「预测日记」，给判断记账

第八章说直觉的敌人是假反馈和不核对的反馈。对策很笨但很硬：**把预测写下来，事后对账**。泰特洛克在后来的研究里发现，预测做得准的那批人（他称为「超级预测者」）有一个共同习惯：他们把预测写成具体的、可验证的形式，并且持续复盘记录，校准自己的把握程度。

做法极简：每次做重要判断时，在本子上写三行——

- **我判断会发生什么**（要具体到能证伪：「这个版本上线后一周内客服投诉会降」，而不是「应该会变好」）；
- **我有多大把握**（写个百分比，别写「挺有信心的」）；

- **依据是什么**（我把它归进了哪个格子？和我见过的哪次一样？哪里不一样？）。

然后，等结果出来，回去对账。这个本子干两件事：一是把模糊的感觉变成可核对的记录，让假反馈无处遁形；二是日积月累，你会得到一张别人给不了的图——**你自己的校准曲线**：你说「八成把握」的时候，实际中了几成。知道自己什么时候容易飘，是判断力最重要的元数据。

大白话

大白话：感觉会美化自己的历史。写下来，白纸黑字，才看得出你的格子哪些灵、哪些是拿噪声练出来的。

方法三：专练「找不同」

第六章说过，格子的边界是在对比里磨出来的。落地的方法就一条：**遇到任何「这和上次一样」的念头，强制自己补一句「哪里不一样」**。

更主动的做法是制造对比练习：把最像的两个案例并排摆开，一条条列差异，然后给每条差异标注「这条要紧吗」。医生的鉴别诊断就是这么练的——症状最像的两种病，逐项对。你可以把它用在任何地方：两份看起来一样的方案、两个看起来一样的候选人、两次看起来一样的下跌。

方法四：复盘时复盘「归类」，不只复盘结果

很多人复盘只问「成了没有」。成了就庆祝，败了就归因于运气或执行。这种复盘练不出判断力，因为判断力的主战场在归类那一步。有效的复盘要问的是：

- 我当时把这个局面归进了哪个格子？归对了吗？
- 如果归错了，是哪个差异维度被我当成了噪声？
- 如果归对了但结果不好，是格子本身不灵，还是这次属于小概率？

第三个问题尤其要紧。第八章说过，在噪声多的领域里，好的判断也常常带来坏结果，反过来也成立。只按结果复盘，等于用假反馈训练自己——你会把灵格子改掉，把不灵的死守住。

一张最小行动清单

如果这一章只能带走三句话：

1. 学东西，混着练，每次先认类型再出招；
2. 做判断，写下来，给把握打个分，事后对账；
3. 复盘时，先审归类再审结果。

格子不是一天画成的，但每一天它都在被画——问题只是你在用真反馈画，还是用笔感觉画。下一章换一个视角：在一个机器记忆碾压人类的时代，这套练判断的手艺，价值是涨了还是跌了？

第十章 · 当机器记住了全世界

这本书写到这里，有一个大象在房间里坐不住了：现在有大语言模型了。它「读」过的书比任何人类多几个数量级，随问随答，不知疲倦。第二章辛辛苦苦区分的「记忆」和「判断」，在机器面前还有意义吗？人练格子，还值吗？

值。而且这本书前面的每一章，恰好都能用来说明为什么。

记忆确实贬值了，这没什么可伤心的

先承认事实。第二章说「记住一切不等于智能」，当时说的是人；现在机器把这句话演到了极致：它真的几乎记住了一切，而我们确实不能说它因此就有了将军的判断力。反过来，对人来说，「背得出」这个能力的市价确实崩了——就像计算器普及后，心算多位数乘法从硬技能变成了才艺表演。

但回忆一下第二章真正的结论：重要的从来不是存货量，而是存货的**形式**——孤立的事实是死库存，带判断的模式才是活资产。这个区分在 AI 时代反而更锋利了：凡是能被检索到的东西都不值钱，值钱的是**把眼前这个具体局面归对格子**的能力。机器把「词典」的价格打到零，「字条」的价格就相对涨了。

AI 自己就是一台「找相同」的机器

更有意思的是看 AI 本身的长短板，简直是为第五、六章做的注释。大语言模型做的事，本质上是把人类写过的海量文本里的模式吃进去，然后在你说一件事时，调出「最像的那些文本里通常会怎么接」。这是「**在不同中发现相同**」的**极限形态**——它在百万量级的文本上连类比之线，所以它特别擅长给出「这类问题一般怎么答」。

但它的软肋也正在另一面：「**在相同中发现不同**」。两个局面的骨架相同、要命的维度不同——第六章说过，这是人最值钱的那半个格子。机器对「一般如何」很敏感，对「这一次的特殊之处」很迟钝；更麻烦的是，它分不清自己什么时候对，错了也常常说得一样流畅、一样自信——它内置了一个没有安装自检模块的 WYSIATI。第八章的两条自检（环境有规律吗？反馈干净吗？）对它同样适用：它见过棋谱，但没为一步棋付出过代价。

大白话：AI 是地球上最强的「这类问题通常怎么办」机器。但你的问题从来不是「通常」，而是「这一次」。「这一次」里那个决定成败的差异，要你自己盯着。

判断里还有一样东西机器没有：后果

再往深一层。第一章说智能的定义里有个词是「行动」。行动意味着有人承担后果。将军下令，输的要死人的；医生说「先按心梗排查」，错了要负责的。**判断之所以是判断而不是聊天，是因为判断者被后果绑着。**这份绑定带来两样机器没有的东西：一是真正的谨慎——知道哪里是要命的维度，因为代价你付；二是责任——「出了问题算谁的」这个问题，最终只能落在一个会疼、会被问责的人身上。

所以人机分工的自然形状是：让机器做它擅长的——穷举选项、提供「通常的做法」、当你的第二意见；人守住机器接不住的——认出这次哪里不一样、在不完整的信息里拍板、为结果负责。机器把你的「查词典」时间省下来，全部预算都可以投给练格子。

顺便说，AI 还是个很好的练习对手

最后给一个实操建议，把第九章的方法和 AI 接起来：大语言模型是历史上最便宜的「交叉练习出题机」。你可以让它一口气出二十道打乱类型的题，逼你先归类再解；可以让它扮演面试里的刁难者、谈判桌的对面，给你即时反馈；可以把你写的预测日记扔给它，让它挑你推理里的漏洞。它不知道答案对你是否正确，但它有的是耐心，而练习最难凑齐的条件——一个随叫随到、不厌其烦、随时给你出「没标题型」的题的陪练——突然就免费了。

只是别让它替你判断。陪练可以帮你练格子，上战场扛责任的，是你。

终章 · 回到战场

回到第一章的战场。现在再看那位将军，他身上的「电光石火」已经拆得开了。

情报涌进来，他不慌，因为他不需要把信息「看全」——他带着几百个局面模式进场，让现实和模式对表，专盯对不上的地方（第七章）。「左翼失联、正面压力加大、敌方补给线拉长」这一团混乱，在他眼里自动归入了一个格子：「敌人在孤注一掷」。这是归类，是在不同中发现相同（第三、五章）。但他没有立刻下令——他又扫了一眼那些要命的维度：这次的对手指挥官是谨慎型还是赌徒型？天气和上次那场仗差在哪？这是分辨，是在相同中发现不同（第六章）。两个动作都发生在电光石火之间，因为格子是现成的，边界是磨好的。然后他下令。行动，承担后果（第一章）。

把整本书压成三句话

第一句：智能不是存储，是判断。 记住一切不等于智能——象棋大师五秒钟复现棋局的本事，随机摆子就失效，说明他存的不是棋子位置，而是带判断的模式（第二章）。判断的第一步永远是归类：认出「这是哪一类情况」，而这一步归错，后面越努力错得越远（第三章）。

第二句：格子靠练，但要练对。 交叉练习的实验说明，把不同类型的问题混着练，效果远好过一类类刷——因为真实世界里没有人在题目前标好题型，混着练才真正练到「该用哪一招」（第四章）。而直觉能不能信，取决于两件事：这个领域有没有稳定规律，你收到的是不是干净快速的反馈——缺了这两样，练出来的自信只是拿噪声喂大的（第八章）。

第三句：格子有两面，一面是相同，一面是不同。 在不同中发现相同，是抽象、是类比、是把新问题认成老问题——堡垒和肿瘤共享同一个骨架，但没人提醒就只有三成的人看得见（第五章）。在相同中发现不同，是分辨、是不被表象骗——两个局面骨架一样，要命的差异藏在维度里，而差异只在对比中显形（第六章）。

三句话合起来，就是这本书的标题问题的答案：**智能，是在电光石火间做出判断；而判断的内核，是在不同中发现相同、在相同中发现不同。**

它能练吗

全书开头的问题，现在可以正面回答了。能，但有条件。

它不像背单词那样能靠堆时间速成，因为格子只能在真实的归类-反馈循环里长出来；它也不像身高那样由天赋定死，因为交叉练习、预测日记、对比训练、归类复盘（第九章），每一招都有实验或研究在后面撑着。比较准确的说法是：**判断力像肌肉，不像存款**——不能继承，不能转账，但任何人按正确的方法练，都会长。

而且时机恰好。第十章说过，当机器把「记住」和「通常怎么做」的价格打到接近零，人身上剩下的那块最贵的资产，恰好就是这个能练的东西。这不是坏消息，这是一张涨价通知单。

最后，一点诚实

这本书把判断讲得很利落，但要对读者诚实：真实的判断永远比书里写的狼狈。格子会过时，边界会画错，最好的将军也会打败仗——第八章的教训是，在规律稀薄的领域里，判断力强的人也只是从「抛硬币」升级到「抛一枚灌了铅的硬币」，不是从不失手。

但这恰恰是让这件事值得练的最后一条理由：既然失手无法消除，能优化的就只有格子的质量和校准的诚实。写下来的预测对账、复盘时对归类的拷问，都是在做同一件事——让自己比昨天的自己错得更明白一点。

智能不是「永不出错」，而是「每一个错都变成格子边界上一条新的刻线」。将军的电光石火不是天赋的闪电，是千万条刻线在一瞬间同时发光。

下一次局面扑到你脸上的那一秒——老板的突然发问、孩子的反常、市场的跳水——愿你的格子已经等在那里了。