



论文

What 和 Where: 视觉背腹侧通路在汉字组块破解过程中连接的变化

吴齐元^{①②}, 吴丽丽^{①②}, 罗劲^{①*}^① 中国科学院心理研究所, 北京 100101;^② 中国科学院研究生院, 北京 100049

* 联系人, Email: luoj@psych.ac.cn

收稿日期: 2009-09-14; 接受日期: 2009-12-25

中国科学院知识创新工程重要方向性项目(批准号: KSCX2-YW-R-28)、国家自然科学基金面上项目(批准号: 30770708, 30970890)、国家高技术研究发展计划(批准号: 2008AA021204)和国家重点基础研究发展计划(批准号: 2010CB833900)项目资助

摘要 组块破解是指将熟悉的知觉整体或对象分解为其组成部分, 再以新的方式组成其他项目. 以往对于组块破解的眼动和脑成像研究提示: 组块破解具有视觉-空间信息加工的特点, 包含“在哪里”和“是什么”两个基本的认知侧面, 可能同时激活视觉的背侧(what)和腹侧(where)通路. 本研究应用动态因果模型探讨不同难度的汉字组块破解任务对于视知觉信息加工中的背、腹侧通路的连接强度的影响. 实验操纵了阻碍组块破解实现的两种不同因素——熟悉度和结构紧密性. 其中, 熟悉度变量通过真字和假字两个水平实现, 而结构紧密程度变量则通过松散结构和紧密结构两个水平(如将“旧”拆解为“日”, 或将“四”拆解为“匹”)体现. 结果表明, 熟悉度的增加使 where 通路的调控连接增强; 空间结构紧密性的增加使得 what 通路和 where 通路的调控连接均增强; 而上述两个因素的同时增强不但使 what 通路和 where 通路的调控连接增强, 而且使 what 通路终端到 where 通路终端之间的调控连接也增强.

关键词有效连接
调控连接
动态因果模型
视觉通路
熟悉度
空间结构紧密性

组块(chunk)一方面能极大地提高人们的信息加工能力, 另一方面也可能阻碍问题解决过程中创造性思维的发生. 顿悟的“表征变换理论”认为, 在问题解决过程中, 问题知觉或理解过程所引起的不适合的表征会阻碍人们有效地解决问题, 而成功的问题解决取决于问题表征方式的有效变换^[1]. 表征方式的变换有时需要一个组块破解的过程, 例如古人在印刷和复制书本的过程中, 会自然地将每一页书看作一个整体, 只有当人们将这种认识上的整体组块破

解为一个个的字或字母并意识到其重组的可能性时, 活字印刷术才可能出现. 同样, 古人在用工具制造过程中也倾向于将这个过程(如制造一根针)看作是一个整体, 而当人们将这个看似整体的过程分解为各个不同的环节, 并用不同的工艺生产流程加以处理时, 生产效率就会数十倍地提高.

1999年, Knoblich等人^[2]用实验证明了顿悟中的组块破解学说, 应用火柴摆成的罗马数字算式, 要求被试通过挪动一根火柴, 使得原来不能成立的等式

变得能够成立. 例如对于 $VI=VII+I$ ($6=7+1$), 解决办法是将等式右边的 $VII(7)$ 中的最右边的一根火柴挪走, 使之变为 $VI(6)$, 而将挪出来的火柴放在等式左边的 $VI(6)$ 的最右边, 使之变为 $VII(7)$. 因为 $VII(7)$ 由一个 V (5) 和 2 个 I (1) 所组成, 是比较松散的组块, 因此, 将这个组块破解成一个 $VI(6)$ 和一个 I (1) 相对较容易. 但若解决 $XI=III+III$ ($11=3+3$) 就比较难了, 人们必须将紧密程度较高的 $X(10)$ 分解成“ $\backslash$ ”和“ $/$ ”并重构成 V (5). 行为实验表明, 需要破解的组块的紧密程度越大, 人们解决这个问题可能性越小, 需要的时间越长.

虽然火柴罗马数字算式问题被成功地应用于组块破解研究并在一定程度上揭示了其相应的认知机制, 但这个方法还存在明显的缺陷: 除了难以控制被试对罗马数字的熟悉程度之外, 罗马数字算式的变化方式和数量十分有限, 很难满足 fMRI 和 ERP 等脑成像实验多个数据叠加的需要. 因此, 人们还无法用它研究相关的脑认知过程. 与罗马数字不同, 汉字是一种典型的知觉组块, 利用汉字的拆字任务研究组块破解过程具有很多明显的优势, 除了汉字数量多、熟悉度便于控制、拆解方法变式丰富之外, 汉字由笔画、部首等不同层次的结构组成, 利用汉字的结构特点, 可以系统地变化组块的紧密程度, 操纵和控制组块破解的难度, 研究组块破解的认知与脑过程^[1,3]. 实验中, 研究者向人们呈现两个汉字, 如“白-排”, 要求人们从右面的字中取出一部分放入左面的字中, 使左右两边的新字仍然为真字, 如将右边的“排”字中的“扌”取出, 与左边的“白”字相结合, 形成“拍-非”. 部首水平的拆字相对容易(如将“青-话”变成“请-舌”, “女-宝”变成“安-玉”, “矢-匡”变成“医-王”), 而笔画水平的拆解则相对较难(如将“三-四”变成“王-匹”; “干-学”变成“平-字”, “三-兴”变成“兰-六”, “大-不”变成“天-个”等), 这是因为笔画水平的拆字涉及较为紧密的组块破解的缘故.

火柴罗马数字算式问题的眼动研究显示, 对算式各组成部分的注视时间会因火柴罗马数字算式问题的类型而异, 有关注视点的转移与注视时间的长短的实验也提示, what 与 where 过程可能同时参与组块破解^[4]. 后来, 应用汉字研究组块破解过程的认知与脑过程, 取得了一些重要进展. Luo 等人^[1]发现, 组块破解过程会同时激活视觉的背侧(where)和腹侧(what)通路的一些典型区域, 并且初级视皮层与高级视皮层在激活上存在不一致性, 揭示了汉字“拆装”

过程的视觉加工特征. 火柴罗马数字算式问题与汉字拆字问题中所包含的组块破解过程在更大的程度上属于空间信息加工. 可见, 组块破解过程包含“在哪里”和“是什么”两个基本的认知侧面, 但目前对组块破解过程中上述两个过程的参与情况及相互作用缺乏基本的了解.

视觉系统包含整个枕叶、大部分颞叶和顶叶区域^[5]. 1982 年, Ungerleider 和 Mishkin^[6]首次提出视觉系统有背侧和腹侧两条视觉通路, 其中背侧通路从枕叶至后顶叶, 负责空间信息加工, 是 where 通路; 腹侧通路从枕叶至下颞叶, 负责物体识别, 是 what 通路. 目前比较认同的是背腹侧通路均是从初级视皮层 $V1$ 出发, 再经次级视皮层 $V2$, 最后分别到达后顶叶和下颞叶^[7], 其中 $V1$ 和 $V2$ 分别对应 Broadman(BA) 17 区和 18 区^[8]. 在恒河猴(*Rhesus monkey*)中, 背侧通路的终端是后顶叶的顶下小叶(BA 7). 在人类大脑中, 后顶叶包括顶下小叶(BA 40)和顶上小叶(BA 7), 其中顶上小叶(BA 7)在组织细胞学上与恒河猴的顶下小叶(BA 7)相同, 因此认为人类的背侧通路终点是顶上小叶(BA 7), 腹侧通路的终端在恒河猴和人类大脑中均为下颞回(BA 19)^[9]. 但对于背腹侧通路涉及的枕叶、颞叶和顶叶的许多其他脑区, 以及背腹侧通路之间的连接, 仍然存在着很多争议. 有些研究认为, 背腹侧通路相互独立^[10], 另有研究认为背腹侧通路相互连接^[11], Hagmann 等人^[12]证明了顶叶与颞叶之间存在着相互连接.

本实验根据 Wu 等人未发表研究中的 fMRI 数据做进一步的 DCM 分析. 拟以视觉通路中背侧 where 通路和腹侧 what 通路为研究对象, 应用基于脑功能成像数据的动态因果模型(dynamic causal modeling, DCM)研究汉字组块破解过程中背腹侧通路的连接变化情况. DCM 是研究大脑区域之间有效连接的方法之一, 与多元回归模型及结构方程模型等假设各区域的血氧输入是由随机的内源性噪声引起不同, 动态因果模型假设输入是确定的实验操作. 在过去的数年里, 动态因果模型被应用在语言、注意、运动和情绪等研究中, 为研究人类认知提供了一种新方法^[13-16]. Mechelli 等人^[17]应用动态因果模型研究了物体类别效应在视觉通路的有效连接, 发现初级视皮层对枕叶和颞叶的物体类别效应区有调控作用, 而从顶叶至物体类别效应区的有效连接与实验刺激无关, 表明枕叶和颞叶的物体类别效应区有自下而上的调控

机制。

本研究系统地探讨了组块破解过程中两种重要的阻碍因素——熟悉度和空间紧密程度,对组块破解过程的影响。其中,熟悉度是指人们对于将要被破解的特定知觉组块的熟悉程度,越是熟悉的组块破解起来就会越困难;而空间紧密程度则是指被破解的部分本身是否是小的完整的组块,若是(如VI中的 I),则组块的紧密程度相对较为松散,破解比较容易,若不是(如 X 中的“\”和“/”),则组块的紧密程度就相对较强,破解更难^[2]。本实验中,熟悉度包括人们常见常用的真字和人为合成的符合汉字构字规则的假字两个水平;空间紧密性则由汉字拆解发生的方式决定,若拆解发生在部首水平,仅是将一个现成的部首移出的话,被看作是“松散”水平的拆解,若拆解发生在比部首更为基本的笔画水平,则被看作是“紧密”水平的拆解(图 1)。

鉴于松散结构拆出部分和剩余部分都能构成比较明显的独立组块,因此,相对于假字条件下,真字条件下的松散结构拆解所带来的额外困难并不在于“是什么”(what)而在于“在哪里”(where),即由于组块原有熟悉度所造成的干扰作用使部件拆除过程需要更多的方位确定与控制过程;与松散结构拆解不同,紧密结构拆解发生之前,人们往往不能将拆解后的结果知觉为一个独立的组块(如人们通常不会自然地意识到“四”字中间包含了“匹”字)。因此,紧密结构的拆解过程在包含“在哪里”(where)的空间信息加工过程的同时,也必然包含“是什么”(what)的加工过程。尤其是当熟悉度与空间紧密性两方面的困难同时较大时,what 和 where 两个过程的相互作用可能也会增加。总之,根据以往眼动和脑成像的实验观察^[1,4]与本实验的条件设置,预期熟悉性影响背侧的 where 通路,空间结构的紧密性影响腹侧 what 通路和背侧 where 通路,而且高熟悉度与高空间结构紧密性的叠加可能使 what 通路和 where 通路连接增强,同时还增强了这两条加工通路间的交互作用。

1 材料与方法

1.1 被试

16 名某高校大学生,其中 8 名女生,年龄 19~25 岁。所有被试均身体健康,右利手,视力正常或矫正

后正常,无神经或心理疾病史。由于汉字是具有空间结构的表意文字系统,只选取平时习惯使用拼音输入法而不使用五笔输入法的大学生作为被试。所有被试均仔细阅读并签署了知情同意书。其中 2 名被试 fMRI 数据由于记录过程中头动超过 1°而被剔除。

1.2 实验设计和刺激材料

本研究应用 2(熟悉度:两种条件)×2(空间紧密性:两种水平)的实验设计,熟悉度有真字和人为合成的符合汉字构字规则的假字两个水平;空间紧密性则由汉字拆解的方式决定,若拆解发生在部首水平,则是“松散”的,若拆解发生在比部首更为基本的笔画水平,则是“紧密”的(图 1)。

4 个条件共 176 个刺激,其中假字_松散、假字_紧密、真字_松散、真字_紧密 4 种条件各 44 个。每种条件的探测词 4 个,实验用词 40 个。另请 60 名被试(有效评估人数为 58)对所有汉字刺激材料进行主观熟悉度评定(采用 5 点量表,1 表示非常不熟悉,5 表示非常熟悉)。结果表明,真字的熟悉度显著大于假字($F(1,57)=1870.599, P<0.001$),说明真、假字变量可有效操纵改变目标刺激的主观熟悉程度。而对各种条件下所使用汉字所含笔画数的统计表明,松散结构的笔画数显著大于紧密结构的笔画数($F(1,43)=252.146, P<0.001$)。

1.3 实验程序

共包含 4 个实验 Block,每个 Block 包含 44 个 trial,其中 40 个实验 trial(各实验条件 10 个),4 个探测 trial(各实验条件 1 个)。Block 顺序随机,每个 Block 中刺激的顺序随机。每个 trial 包含两个阶段:真假字判断阶段和组块破解阶段。fMRI 的 BOLD 信号变化所反映的是脑在特定时间范围内的所有信息加工活动的累积整体效应,因此,为了避免刺激项目的“识别编码过程”与“组块破解过程”混在一起,影响所获观察的单纯性,故对“编码”和“拆分”两个阶段加以区分。单个 trial 的实验过程见图 1。每个 trial 中,第一阶段是真假字判断阶段,屏幕上出现刺激 1200 ms,被试需要按键判断刺激是否为真字;然后出现注视点 2800 ms;注视点消失后是组块破解阶段,屏幕上出现新刺激 3000 ms,其中刺激的左边是真假字判断阶段中出现过的刺激,右边是它的一部分部首或笔

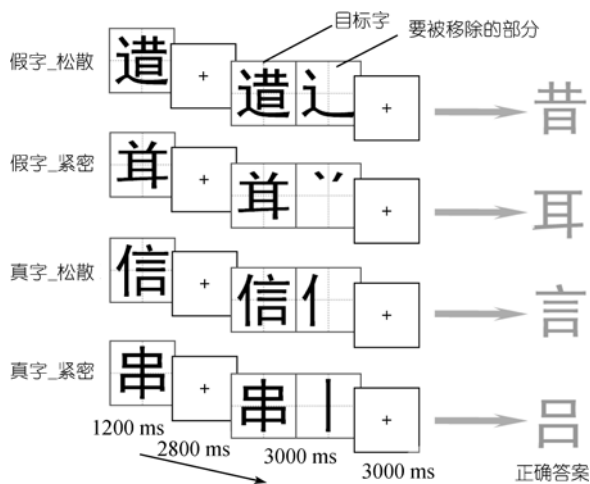


图 1 单个 trial 的实验过程

从上到下分别为假字_松散、假字_紧密、真字_松散和真字_紧密
4 个条件下单个 trial 的实验过程示例

画, 被试需要判断从左边刺激中移除右边刺激后, 剩下部分是否为真字。

1.4 MRI 数据采集

所有 MRI 数据均由 3.0T 场强共振成像系统(西门子)使用标准无线电频率线圈采集得到。首先采用 T1 加权的快速自旋回波(SE)脉冲序列采集三维结构像, 每个结构像包含 128 张切片(层厚=1.33 mm), 参数为: TR=2530 ms, TE=3.39 ms, FOV=256 mm×256 mm, FA=7°, VS=1.33 mm×1 mm×1.33 mm。再采用梯度回波平面成像脉冲序列(EPI)进行功能像扫描, 33 层连续轴位扫描, 成像参数为: TR=2000 ms, TE=30 ms, FOV=200 mm×200 mm, Matrix=64×64, FA=90°, VS=3.1 mm×3.1 mm×4.0 mm。每个被试数据包含 4 个 session, 每个 session 包含 222 张扫描图像。采集过程中, 被试的头被固定在泡沫块上, 以减小头动。

1.5 fMRI 数据分析

应用 SPM5(Welcome Department of Imaging Neuroscience, London, UK; <http://www.fil.ion.ucl.ac.uk/spm>)对所获得的 fMRI 数据进行分析。每个 session 的前 5 张 EPI 图像被剔除。预处理按以下顺序进行: 时间校正(参考图片为第 33 层切片), 空间重排, 标准化, 平滑(FWHM=8 mm)。预处理后剔除两个头动超过 1°的被试。

在确定单个被试的一般线性模型(GLM)时, 应

用截止频率为 1/128 Hz 的高通滤波器对所有被试的时间序列进行滤波, 不执行全局标准化。每个被试的数据包含 4 个 session, 每个 session 包含 9 种事件类型, 4 种实验条件(假字_松散、假字_紧密、真字_松散和真字_紧密)在真假字判断阶段和组块破解阶段分别有 1 种事件类型, 分别包含各阶段该实验条件下正确反应且反应时在平均反应时 3 个标准差之内的事件, 共 8 种事件类型, 最后一种事件类型是所有错误反应或反应时在平均反应时 3 个标准差之外的事件。在被试水平分别得到每个被试在组块破解阶段 4 种实验条件的 contrast, 以进行群组水平的分析。

在群组水平应用每个被试在组块破解阶段 4 种实验条件的 contrast 进行方差分析(ANOVA, $N=14$), 分别定义了以下 5 个 contrast(其中前 4 个为 t 对比(t -contrast), 最后一个为 F 对比(F -contrast)): (1) 熟悉 vs. 不熟悉: [(真字_松散+真字_紧密)-(假字_松散+假字_紧密)]; (2) 不熟悉 vs. 熟悉: [(假字_松散+假字_紧密)-(真字_松散+真字_紧密)]; (3) 紧密 vs. 松散: [(假字_紧密+真字_紧密)-(假字_松散+真字_松散)]; (4) 松散 vs. 紧密: [(假字_松散+真字_松散)-(假字_紧密+真字_紧密)]; (5) 所有刺激的主效应: [假字_松散+假字_紧密+真字_松散+真字_紧密]。最后报告的结果都是基于 FWE(family wise error)错误类型校正($P<0.05$)统计。

1.6 动态因果模型

动态因果模型假设大脑是一个确定的非线性动态系统, 接受来自外部的实验刺激并输出反应, 其输入是确定的实验操作。实验操作可通过输入直接作用于某些脑区, 也可作用于不同脑区之间的有效连接。根据各脑区在不同输入条件下的血氧动力学反应, 通过生物物理学方程与神经状态变量产生联系。最后给出 3 类模型参数: (1) 输入直接作用于某些区域的驱动强度; (2) 与实验操作无关的不同区域间的固有内在连接强度; (3) 实验操作对内在连接的调控强度的影响^[18]。这 3 个参数中, 最值得关注的是第 3 个, 即调控连接, 因为它是由实验操作直接引起的。动态因果模型的连接强度不是指区域本身的活动强度单位, 而是指区域间活动相互影响的速度。对于一个感兴趣的信息加工系统而言, 其连接方式可能有多个不同的假设, 而每个假设都有相应的 DCM 模型。贝叶斯模型选择方法可用来从数个可能的模型中挑

选最优模型, 它应用的是各模型的日志证据(log evidence), 通常被称为贝叶斯因子 (Bayesian factor, BF), 常用 AIC/BIC(Akaike's information criterion, AIC; Bayesian information criterion, BIC)或负自由能量(Negative free energy, NFE)来近似贝叶斯因子. 本研究应用中最新的 SPM8b(Wellcome Department of Imaging Neuroscience, London, UK)工具箱进行 DCM 分析, 采用负自由能量来近似贝叶斯因子. 负自由能量越大(绝对值越小), 则模型越好, 比较各模型的负自由能量, 从而可以得出最优模型^[19~22].

在具体的操作过程中, 对每个被试的 fMRI 数据定义了一个用于 DCM 计算的一般线性模型 (dcm-GLM), 该模型共包括 6 个事件类型, 其中 4 个分别是组块破解阶段的 4 个条件: 假字_松散, 假字_紧密, 真字_松散和真字_紧密, 第 5 个事件包括该以上 4 种条件(ALL), 第 6 个事件是所有真假字判断阶段的事件及组块破解阶段错误反应和反应时间有平均反应时间 3 个标准差之外的所有事件. 其中, 只有前 5 个事件被纳入动态因果模型分析. 同时, 为了减少 DCM 的一般线性模型中共线性的影响, 另外定义了提取感兴趣区域时间序列的一般线性模型(voi-GLM). 在 voi-GLM 中, 共包含 9 种事件类型, 与 1.5 小节中提到的事件类型完全一致.

在 dcm-GLM 和 voi-GLM 中, 都应用截止频率为 1/128 Hz 的高通滤波器对所有被试的时间序列进行滤波, 不执行全局标准化. 所有事件均锁时在刺激呈现时间点. 模型把 4 个 session 的数据串连成一个连接的 session, 因此模型中还包括 4 个额外的、均由 1 和 0 组成的回归量, 分别对应数据采集中的 4 个 session; 其中第一个 session 中的各 scan 与第一个回归量的 1 相对应, 而其他数据采集 session 的 scan 对应于第一个回归量中的 0, 依此类推. 基于单个被试 voi-GLM 的 t 检验结果, 用 SVC(small volume correction)定义动态因果模型中的单个被试的感兴趣区域.

在提取单个被试感兴趣区域的时间序列时, 先定义概率阈值 $P=1$, 然后以群组水平的峰值点坐标为球心, 在半径 $r=12$ mm 的球内寻找单个被试在感兴趣区域的峰值点, 并取最大峰值点作为单个被试的感兴趣区域. 对于每个被试的 3 个感兴趣区域, 经确认用该方法获得的最大峰值点与相应的群组水平峰

值点属于同一解剖位置. 最后, 提取以最大峰值点为球心, 半径 $r=6$ mm 的球内所有体素的时间序列, 同时应用 4 个条件的 F 检验对提取的时间序列进行矫正.

同样, DCM 也有两个水平的随机效应分析. 首先在单个被试水平, 基于每个被试的 dcm-GLM 建立 3 个动态因果模型, 在群组水平应用贝叶斯模型选择方法得到的最优模型为模型 3. 最后, SPSS 对每个连接进行双尾检验($n=14$), t 检验其均值是否显著地不等于 0. Shapiro-Wilk 正态检验确定每个连接服从正态分布. 对于显著非正态分布的连接强度, 用非参数统计方法.

2 结果

2.1 行为结果

在真假字判断阶段和组块破解阶段, 被试都需要按键做出反应. 用 SPSS 对被试在各种实验条件下组块破解阶段的正确率和反应时做了统计分析(表 1 和图 2). 应用 2×2 ANOVA(熟悉度 $\times$ 结构紧密性)分析发现, 在正确率中, 熟悉度的主效应显著($P<0.001$, $F(1,13)=24.202$), 假字的正确率显著地大于真字的正确率, 松散结构的正确率显著大于紧密结构的正确率($P<0.001$, $F(1,13)=41.566$), 交互作用也显著($P<0.001$, $F(1,13)=14.610$); 在反应时中, 熟悉度的主效应显著($P<0.001$, $F(1,13)=104.289$), 真字的反应时显著大于假字的反应时, 紧密结构的反应时显著长于松散结构的反应时($P<0.001$, $F(1,13)=278.297$), 交互作用也显著($P<0.001$, $F(1,13)=43.974$).

2.2 fMRI 分析结果

本实验虽然设置了两个实验任务, 真假字判断任务和组块破解任务, 但设置真假字判断任务的目的在于减小编码过程对于组块破解过程的影响, 从而获得较为单纯的组块破解过程. 从 fMRI 结果来看, 两个实验任务激活的脑区分布与激活程度存在较大差异. 本研究的主要目的是应用动态因果模型研究汉字组块破解过程中背腹侧通路的连接变化情况, 因此只报告组块破解任务中的 fMRI 结果.

(1) 所有刺激的主效应. 为了考查所有刺激的主效应激活区, 对 4 种实验条件作联合(conjunction)分析(FEW, $P<0.05$)时发现, 激活最强的区域是左侧

枕下回(left inferior occipital gyrus, LIOG, BA18, MNI 坐标: $x=33, y=-84, z=-9$), 说明该区域在 4 种实验条件下均有较强的激活. 在定义的所有刺激的 F 对比中也发现, 激活最显著的峰值点是左侧枕下回(LIOG, MNI 坐标: $x=-33, y=-84, z=-9$). 该区域属于 V2 区域, 也是背腹侧视觉通路的共有区域.

(2) 熟悉度的主效应. 为了考查熟悉度对大脑活动的影响, 分析了[(真字_松散+真字_紧密)-(假字_松散+假字_紧密)]. 发现[(真字_松散+真字_紧密)-(假字_松散+假字_紧密)]的激活区域主要包含双侧前扣带回(BA 32)和左侧梭状回(BA 37). 进一步分析得出, (真字_松散-假字_松散)无显著激活区域, 熟悉度所造成的差别主要来自(真字_紧密-假字_紧密).

(3) 空间紧密性的主效应. 为了考查空间紧密性对大脑活动的影响, 分析了[(假字_紧密+真字_紧密)-(假字_松散+真字_松散)], 观察到的激活区域包括双侧额下回(BA 47/9/10/46), 双侧额中回(BA 6/9/10/46), 双侧颞下回(BA 19/20), 左侧枕中回(BA 19)右侧枕上回(BA 39), 双侧顶上小叶(BA 7), 双侧纺锤状回(BA 20/37), 双侧楔前叶(BA 31/7), 右侧前扣带回(BA 32), 左侧顶下小叶(BA 40), 右侧后中央回(BA 2/40)及右侧前中央回(BA 9). 再分析(假字_紧密-假字_松散)与(真字_紧密-真字_松散), 发现其激活区域与[(假字_紧密+真字_紧密)-(假字_松散+真字_松散)]相似, 其中属于视觉通路中终端的两个最显著的峰值点分别如下: 左侧颞下回(MNI 坐标: $x=-48, y=-63, z=-9$, LITG, BA 19), 解剖学上属于左侧颞下

回(IT), 是腹部视觉通路的终端; 左侧顶上小叶(MNI 坐标: $x=-24, y=-69, z=45$, LSPL, BA 7), 解剖学上属于后顶叶, 是人类背侧视觉通路的终端.

2.3 动态因果模型

大量研究证明, 左侧枕上小叶(left superior posterior lobule, LSPL, BA7)和左侧后部下颞叶(left posterior inferior temporal cortex, LPIT, BA 19/37)在处理汉字视觉信息过程中有重要作用^[23~28]. 研究者认为, 左侧枕上小叶的激活是负责汉字的亚结构(如部首)加工^[28], 而左侧后部下颞叶负责从长时记忆中提取汉字的视觉图像信息^[29]. fMRI 结果发现, 所有刺激的主效应激活区在左侧枕下回, 属于 V2 区域, 在背腹侧视觉通路的共有区域上. 结构紧密性的主效应区域包含视觉系统的大部分区域, 其中激活最强的左侧大脑的视觉系统区域有左侧顶下小叶(BA 40)、左侧顶下小叶(BA 7)和左侧颞下回(BA 19)等. 由于语言认知的左侧大脑优势, 同时为了简化模型^[22], 借鉴其他的动态因果模型研究^[13], 只选取了左侧枕

表 1 组块破解阶段的正确率和反应时

	松散		紧密	
	假字	真字	假字	真字
正确率	0.9929	0.9875	0.9393	0.8750
标准误	0.00313	0.00435	0.01404	0.01777
平均反应时 (ms)	701.6907	815.1471	1119.7843	1454.8214
标准误	44.03068	50.89998	56.74077	43.65215

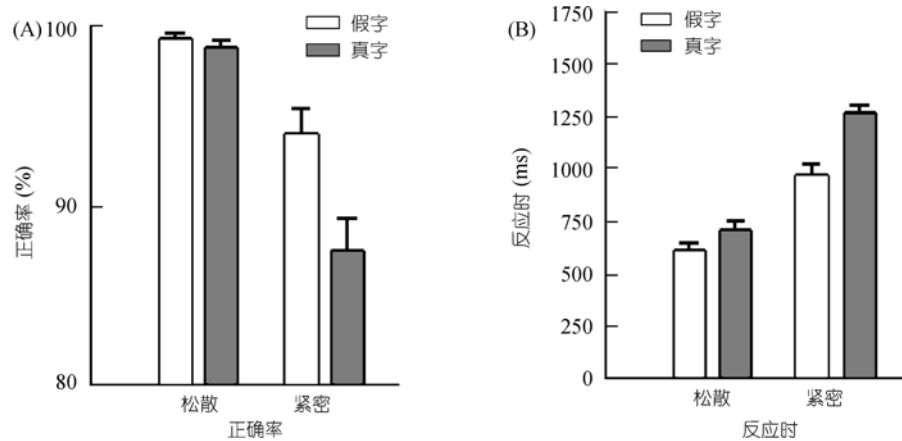


图 2 行为结果: 组块破解阶段的正确率与反应时

下回、左侧顶上小叶和左侧颞下回 3 个区域作为动态因果模型的感兴趣区域. 其中左侧颞下回(LITG)和左侧顶上小叶(LSPL)是从空间紧密结构紧密性的主效应(紧密 vs.松散: [(假字_紧密+真字_紧密)-(假字_松散+真字_松散)])中提取的, 其坐标是基于空间紧密性的主效应((紧密 vs.松散: [(假字_紧密+真字_紧密)-(假字_松散+真字_松散)])在群组水平的结果. 左侧枕下回(LIOG)是基于群组水平 ANOVA 的所有组块破解阶段实验条件 F 检验(F 对比: [假字_松散+假字_紧密+真字_松散+真字_紧密])下的主效应区, 并且作为所有刺激作用于该视觉通路网络的输入点.

在单个被试水平, 基于每个被试的 dcm-GLM, ALL(包含组块破解阶段的所有实验条件)直接作用于左侧枕下回, 作为整个模型的输入. 4 种实验条件可以作用于该视觉信息加工网络中的任何连接. 假设 3 个 DCM 的内在连接模型(图 3): (1) 背腹侧视觉通路分别只有自下而上的单向连接, 即 $LIOG \rightarrow LITG$, $LIOG \rightarrow LSPL$; (2) 背侧视觉通路存在着双向连接, 腹侧视觉通路只有自下而上的单向连接, 同时背侧视觉通路的终端左侧顶上小叶至左侧颞下回也有连接, 即 $LIOG \rightarrow LITG$, $LIOG \leftrightarrow LSPL$, $LSPL \rightarrow LITG$; (3) 由于有研究认为, DCM 的内在连接不必拘泥于已知的解剖连接^[13], 定义了全连接模型, 即 3 个区域之间均有双向连接, 即 $LIOG \leftrightarrow LITG$, $LIOG \leftrightarrow LSPL$, $LSPL \leftrightarrow LITG$. 最后, 在群组水平应用贝叶斯模型选择方法得到最优模型为模型 3.

为了得到最优模型, 在群组水平上应用随机效应方法比较了 3 个模型的负自由能量, 发现对于每个被试数据, 模型 3 的负自由能量都是最大的, 模型 3 的各被试的负自由能量之和也是最大的(绝对值最小), 由此推断最优模型都是模型 3, 即假设 3 个区域

之间均有双向连接的全连接模型(图 3(C)). 负自由能量结果详见附加材料.

然后, 应用 SPSS 的 t 检验统计分析了所有被试在模型 3(全连接模型)中的所有连接参数(多重比较, 采用 Bonferroni 矫正). 结果表明, 所有刺激作用于左侧枕下回的连接强度显著(0.4319 ± 0.0653 , $P < 0.001$). 对于不因任务刺激输入的不同而改变的“内在连接”而言, 只有 where 通路($LIOG \rightarrow LSPL$)以及 where 通路终端与 what 通路终端之间的双向连接($LITG \rightarrow LSPL$, $LSPL \rightarrow LITG$)显著大于 0(图 4 和表 2).

而对于因任务刺激输入的不同而改变的“调控连接”而言, 假字_松散条件下无任何显著的调控连接, 真字_松散条件下只有 what 通路($LIOG \rightarrow LITG$)是显著的, 假字_紧密条件下 where 通路和 what 通路均显著, 真字_紧密条件下不仅 where 通路和 what 通路均显著, 同时还有 $LITG \rightarrow LSPL$ 的调控连接是显著的(图 4 和表 2).

3 讨论

本研究应用汉字的熟悉度和空间结构紧密程度双因素设计研究汉字组块破解过程. 在刺激材料视觉复杂度统计中发现, 松散结构的笔画数显著大于紧密结构的笔画数, 而行为表明, 松散结构的正确率显著大于紧密结构的正确率, 松散结构的反应时显著小于紧密结构的反应时. 说明组块破解的难度不是由笔画数多少的视觉复杂度造成的. 58 名被试对于刺激材料的熟悉度评分的统计表明, 真字的熟悉度显著大于假字的熟悉度. 说明应用真字和人为合成的假字表示刺激材料熟悉度的两个水平是合理的.

本研究探讨了伴随组块破解过程的脑认知机制,

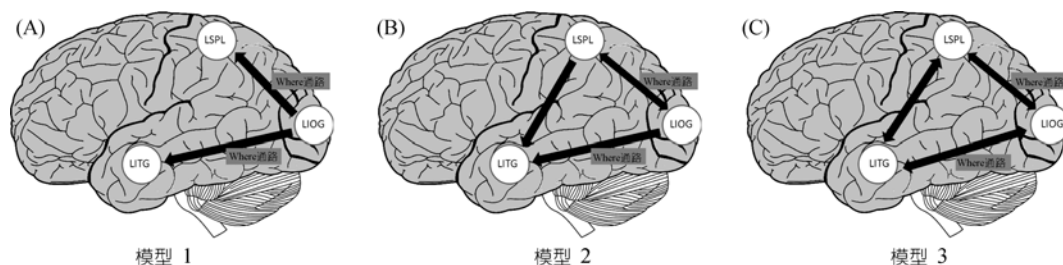


图 3 3 个 DCM 模型

LIOG: 左侧枕下回; LITG: 左侧颞下回; LSPL: 左侧顶上小叶

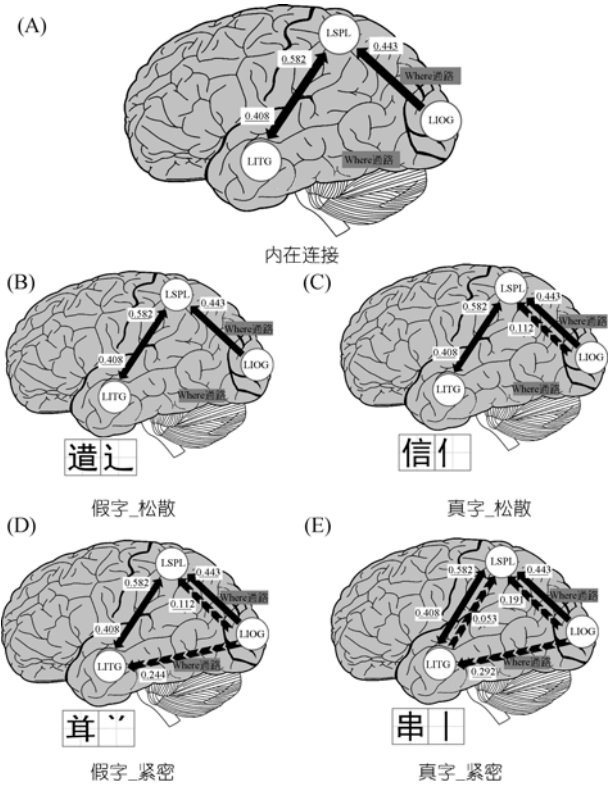


图 4 内在连接与 4 种条件下的调控连接

(A) 内在连接; (B)~(E) 条件下的调控连接. 实线箭头示内在连接; 虚线箭头示各条件下的调控连接. 每个连接图下方为刺激示例

研究结果不仅重复了本实验室^[1]汉字组块破解任务同时激活视觉的背、腹侧通路的发现, 而且揭示了不同任务难度下背、腹侧通路的连接强度变化情况. 行为结果发现, 熟悉性因素(即假字_松散条件相对于真字_松散条件, 或假字_紧密条件相对于真字_紧密条件)和空间结构紧密性因素(即假字_松散条件相对于假字_紧密条件, 或真字_松散条件相对于真字_紧密

条件)在反应时或正确率方面均有可能造成显著差异, 而且存在上述两种因素的交互作用, 意味着如果这两种影响任务难度的因素同时出现的话, 就有可能产生叠加放大效应. 更重要地, 本研究通过 DCM 分析, 揭示了与上述行为观察相一致的功能连接变化, 发现在空间结构较为松散(即部首水平拆字)的条件下, 高熟悉度会使背侧 where 通路的有效连接增强; 而在熟悉度较低(即假字)条件下, 高空间结构紧密性会使腹侧 what 通路和背侧 where 通路的连接均增强; 而在高熟悉度、高空间结构紧密性(即真字_紧密)条件下, 不但 what 通路和 where 通路的有效连接增强, 而且 what 通路终端到 where 通路终端之间的有效连接也增强. 上述结果进一步加深了人们对于组块破解过程脑认知过程的理解.

近年有关顿悟的一个重要认识是顿悟的多重障碍假说^[20]. 该假说认为, 尽管从表面上看人们往往在一个单一的思维步骤中就实现了顿悟, 这个单一的思维步骤却包含着多重障碍的同时克服. 如果人们想要成功地解决 9 点问题, 就必须同时克服两重思维障碍: 一是要把直线的拐点画在没有黑点的地方, 二是要能构思出目标状态的由 4 条直线构成的陌生的箭头状图形^[20]. 本实验为顿悟的多重障碍假说提供了脑科学方面的证据, 不仅观察到了高熟悉度和高空间结构紧密性在各自单独作用的条件下对于视觉的背、腹侧通路有效连接的影响, 而且证明了高熟悉度+高空间结构紧密性的共同作用会使一个新的、在原来各自作用条件下未被观察到的连接——what 通路终端到 where 通路终端之间的连接表现出明显的增强. 从一个侧面证实了阻碍顿悟组块破解的不同因素如果同时出现就会产生叠加放大效应, 使原本无需参与的脑-认知过程加入其中.

表 2 DCM 模型 3 的内在连接和调控连接^{a)}

通路连接	内在连接	松散		紧密	
		假字_松散	真字_松散	假字_紧密	真字_紧密
what 通路: LOIG→LITG	0.085	-0.019	0.021	0.244*	0.292*
where 通路: LOIG→LSPL	0.443*	0.102	0.112*	0.202*	0.191*
LITG→LOIG	0.137	0.023	0.021	-0.039	-0.010
LITG→LSPL	0.582*	-0.045	-0.007	0.041	0.053*
LSPL→LOIG	0.241	-0.021	0.016	-0.014	-0.027
LSPL→LITG	0.408*	-0.080	-0.087	-0.070	-0.047

a) 多重检验, Bonferroni 矫正, $P < 0.05$. * 示在 t 检验下显著的有效连接

本研究利用 DCM 方法研究了组块破解过程中视觉的背、腹侧通路的连接变化, 将 DCM 分析置于一个公认可靠的前提和背景下, 从而大大减少了这类分析所可能具有的随意性。但负作用是未能将一些可能发挥重要作用的脑区, 如额叶, 置于分析的框架之内。额叶在组块破解过程中的重要作用已经被脑损伤的研究所证实。研究发现, 如果在思维过程中负责问题空间的确定以及自上而下加工控制的额叶区域被病变所损伤, 那么这些人顿悟破解组块的能力不但不会降低, 反而较正常人有明显升高^[30]。但考虑到额叶所包含的区域较广, 而且与各个脑区之间都有十分复杂和广泛的联系。因此, 作为对于组块破解问题的初步探索, 有理由认为将讨论局限于一个有较高把握度的框架之下可能更加合适。同样, 对于视

觉的背、腹侧通路在组块破解中的连接变化的探讨, 即使是在不考虑额叶及其他可能发挥作用脑区的情况下, 也仍然是相对充分及有意义的。

4 结论

视觉的背、腹侧通路在不同难度的组块破解任务条件下, 其连接强度会发生变化, 在空间结构较为松散条件下, 高熟悉性会使背侧 where 通路的有效连接增强; 在熟悉性较低条件下, 高空间结构紧密性会使腹侧 what 通路和背侧 where 通路的连接同时增强; 而在高熟悉度加高空间结构紧密性条件下, 不但 what 通路和 where 通路的有效连接增强, 而且 what 通路终端到 where 通路终端之间的有效连接也随之增强。

参考文献

- 1 Luo J, Niki K, Knoblich G. Perceptual contributions to problem solving: chunk decomposition of Chinese characters. *Brain Res Bull*, 2006, 70: 430—443
- 2 Knoblich G, Ohlsson S, Haider H, et al. Constraint relaxation and chunk decomposition in insight problem solving. *JEP: LMC*, 1999, 25: 1534—1555
- 3 Luo J, Knoblich G. Studying insight problem solving with neuroscientific methods. *Methods*, 2007, 42: 77—86
- 4 Knoblich G, Ohlsson S, Raney G E. An eye movement study of insight problem solving. *Mem Cognit*, 2001, 29: 1000—1009
- 5 Wandell B A, Dumoulin O D, Brewer A A. Visual field maps in human cortex. *Neuron*, 2007, 56: 366—383
- 6 Ungerleider L G, Mishkin M. Two Cortical Visual Systems, in *Analysis of Visual Behavior*. Goodale M A, Mansfield R J W, eds. Cambridge: MIT Press, 1982. 549—586
- 7 Kaas J H, Lyon D C. Pulvinar contributions to the dorsal and ventral streams of visual processing in primates. *Brain Res Rev*, 2007, 55: 285—296
- 8 Wohlschlagel A M, Specht K, Lie C, et al. Linking retinotopic fMRI mapping and anatomical probability maps of human occipital areas V1 and V2. *Neuroimage*, 2005, 26: 73—82
- 9 Milner A D, Goodale M A. *The visual brain in action*. Oxford: Oxford University Press, 1995. 248
- 10 Laycock R, Crewther D, Crewther S. The advantage in being magnocellular: a few more remarks on attention and the magnocellular system. *Neurosci Biobehav Rev*, 2008, 32: 1409—1415
- 11 Zanon M, Busan P, Monti F, et al. Cortical Connections Between Dorsal and Ventral Visual Streams in Humans: Evidence by TMS/EEG Co-Registration. *Brain Topogr*, 2010, 22: 307—317
- 12 Hagmann P, Cammoun L, Gigandet X, et al. Mapping the structural core of human cerebral cortex. *PLoS Biol*, 2008, 6: 159
- 13 Heim S, Eickhoff S B, Ischebeck A K, et al. Effective connectivity of the left BA 44, BA 45, and inferior temporal gyrus during lexical and phonological decisions identified with DCM. *Hum Brain Mapp*, 2009, 30: 392—402
- 14 Siman-Tov T, Mendelsohn A, Schonberg T, et al. Bihemispheric leftward bias in a visuospatial attention-related network. *J Neurosci*, 2007, 27: 11271—11278
- 15 Grefkes C, Eickhoff S B, Nowak D A, et al. Dynamic intra- and interhemispheric interactions during unilateral and bilateral hand

- movements assessed with fMRI and DCM. *Neuroimage*, 2008, 41: 1382—1394
- 16 Schlosser R G, Wagner G, Koch K, et al. Fronto-cingulate effective connectivity in major depression: a study with fMRI and dynamic causal modeling. *Neuroimage*, 2008, 43: 645—655
 - 17 Mechelli A, Price C J, Noppeney U, et al. A dynamic causal modeling study on category effects: bottom-up or top-down mediation? *J Cogn Neurosci*, 2003, 15: 925—934
 - 18 Penny W D, Stephan K E, Mechelli A, et al. Modeling functional integration: a comparison of structural equation and dynamic causal models. *Neuroimage*, 2004, 23: 264—274
 - 19 Ashbuner J, Friston K J, Penny W D, Dynamical Causal Modeling. In: Frackowiak R.S J, Friston K J, Frith C, eds. *Human Brain Function*, 2nd ed. London: Academic Press, 2003. 1063—1090
 - 20 Kershaw T C, Ohlsson S. Multiple causes of difficulty in insight: the case of the nine-dot problem. *J Exp Psychol Learn Mem Cogn*, 2004, 30: 3—13
 - 21 Penny W D, Stephan K E, Mechelli A, et al. Friston Comparing dynamic causal models. *Neuroimage*, 2004, 22: 1157—1172
 - 22 Stephan K E, Penny W D, Daybuzeau J, et al. Bayesian model selection for group studies. *Neuroimage*, 2009, 46: 1004—1017
 - 23 Chee M W, Tan E W, Thiel T. Mandarin and English single word processing studied with functional magnetic resonance imaging. *J Neurosci*, 1999, 19: 3050—3056
 - 24 Chee M W, Weekes B, Lee K M, et al. Overlap and dissociation of semantic processing of Chinese characters, English words, and pictures: evidence from fMRI. *Neuroimage*, 2000, 12: 392—403
 - 25 Lee C Y, Tsai J L, Kuo W J, et al. Neuronal correlates of consistency and frequency effects on Chinese character naming: an event-related fMRI study. *Neuroimage*, 2004, 23: 1235—1245
 - 26 Tan L H, Feng C M, Foxy P T, et al. An fMRI study with written Chinese. *Neuroreport*, 2001, 12: 83—88
 - 27 Tan L H, Liu H L, Perfetti C A, et al. The neural system underlying Chinese logograph reading. *Neuroimage*, 2001, 13: 836—846
 - 28 Deng Y, Booth J R, Chou T L, et al. Item-specific and generalization effects on brain activation when learning Chinese characters. *Neuropsychologia*, 2008, 46: 1864—1876
 - 29 Nakamura K, Honda M, Okada T, et al. Participation of the left posterior inferior temporal cortex in writing and mental recall of kanji orthography: a functional MRI study. *Brain*, 2000, 5: 954—967
 - 30 Reverberi C, Toraldo A, Serena D, et al. Better without (lateral) frontal cortex? Insight problems solved by frontal patients. *Brain*, 2005, 128: 2882—2890
 - 31 Tsai P S, Yu B H, Lee C Y, et al. An event-related potential study of the concreteness effect between Chinese nouns and verbs. *Brain Res*, 2009, 1253: 149—160
 - 32 王湘, 程灶火, 姚树桥. 汉字再认 ERP 新旧效应的性别差异. *中国心理卫生杂志*, 2005, 19: 769—773

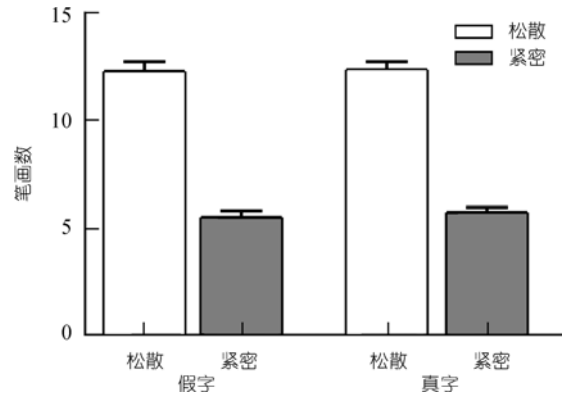
附加材料:

1 刺激材料的视觉复杂度和熟悉度评估

为了评估刺激的视觉复杂度, 参考文献[30], 使用汉字笔画为指标, 分别统计了 4 种实验条件下刺激的笔画数. 应用 SPSS 对 4 种条件刺激的笔画数进行了统计分析(附表 1 和附图 1. 结果表明, 松散结构的笔画数显著大于紧密结构的笔画数($F(1,43)=252.146, P<0.001$). 真字和假字的笔画数没有显著差异.

附表 1 4 种条件刺激材料的笔画数

	假字_松散	假字_紧密	真字_松散	真字_紧密
$\bar{x} \pm SD$	12.1364 $\pm$ 2.66407	5.6364 $\pm$ 1.60074	12.2045 $\pm$ 2.66397	5.7955 $\pm$ 1.37383

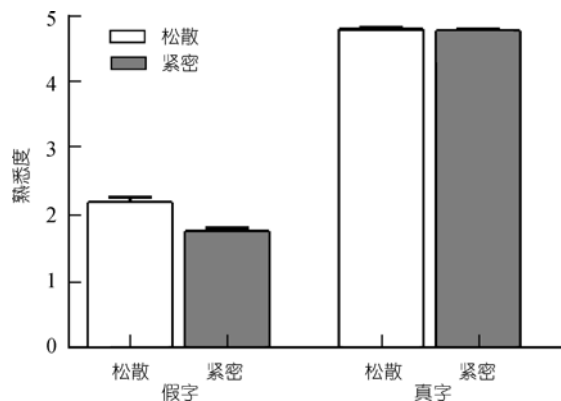


附图 1 4 种条件刺激材料的笔画数

根据以往的汉字研究^[31],采用被试评分的方法来考察汉字的熟悉度.请 60 名普通大学生(男生 29 名,女生 31 名,18~30 岁)对所有刺激材料进行熟悉度评定.把 176 个刺激材料(2 $\times$ 2 设计,共 4 种条件,每种条件 44 个刺激材料)顺序随机,做成纸质问卷,请被试根据自己的经验评价每个刺激材料的熟悉度.采用 5 点量表,1 示非常不熟悉,2 示不熟悉,3 示不确定是否熟悉,4 示熟悉,5 示非常熟悉.被试只需在每个刺激材料的右侧表格中相应熟悉度分数上打钩.共收回 60 份问卷(其中有 2 个男性被试由于答题不认真,整个页面的评分分数完全一样,数据被剔除),有效评定份数为 58 份.应用 SPSS 对 58 份有效问卷进行统计分析(附表 2 和附图 2).结果表明,真字的熟悉度得分显著地高于假字的熟悉度得分($F(1,57)=1870.599, P<0.001$),松散结构的熟悉度显著高于紧密结构的熟悉度得分($F(1,57)=80.335, P<0.001$),真假字和结构紧密程度的交互作用显著($F(1,57)=67.080, P<0.001$).

附表 2 4 种条件刺激材料的熟悉度评分

	假字_松散	假字_紧密	真字_松散	真字_紧密
$\bar{x} \pm SD$	2.1776 $\pm$ 0.55575	1.7474 $\pm$ 0.49443	4.7812 $\pm$ 0.24534	4.7602 $\pm$ 0.15132

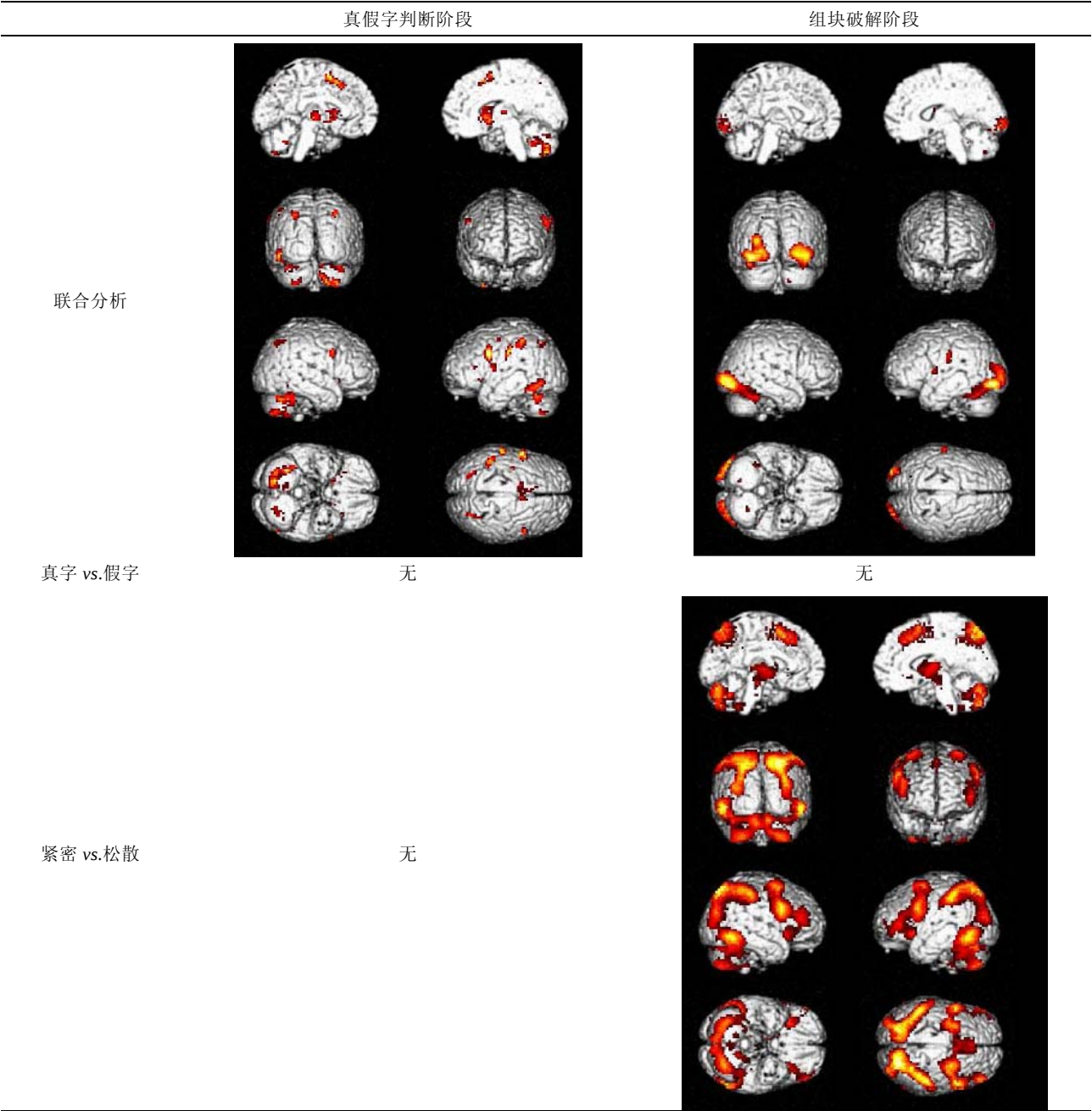


附图 2 4 种条件刺激材料的熟悉度评分结果

2 真假字判断任务与组块破解任务 fMRI 结果比对

根据两个实验任务的 fMRI 结果, 分别对比了 4 种实验条件在两个实验任务下的联合分析、熟悉度的主效应和结构紧密性的主效应(附表 3).

附表 3 2 种实验任务下的 fMRI 结果分析比对(FEW, $P<0.05$)



由附表 3 可见, 同样应用 FEW 矫正方法($P<0.05$), 两个实验任务中激活脑区分布和激活程度存在很大差异. 在 4 种实验条件的联合分析中, 组块破解阶段在左侧枕下回有显著的激活区域, 而真假字判断阶段的显著

激活区域较分散, 包括枕叶、顶叶和额叶等, 在左侧枕下回虽然也有激活区域, 但是激活强度不如组块破解阶段. 在紧密结构与松散结构的对比, 在组块破解阶段视觉通路的背腹侧通路的多个脑区均有激活, 而真假字判断阶段无任何显著激活.

3 3 个 DCM 模型的负自由能量(附表 4)

附表 4 3 个 DCM 模型的负自由能量

被试	模型 1	模型 2	模型 3
01	-4649.41	-4648.86	-4645.81
02	-4028.98	-4028.95	-3980.44
03	-4439.51	-4436.58	-4384.74
04	-4121.14	-4120.81	-4118.75
05	-4535.67	-4533.64	-4519.72
06	-3910.62	-3907.34	-3851.60
07	-4445.75	-4445.92	-4429.94
08	-3447.15	-3442.93	-3392.63
09	-4415.70	-4415.62	-4363.33
10	-4316.67	-4271.07	-4138.05
11	-4006.52	-4005.89	-3939.24
12	-3549.56	-3546.10	-3543.82
13	-3798.87	-3791.50	-3785.73
14	-3498.83	-3480.54	-3445.43
总计	-57164.36	-57075.75	-56539.23

因此, 最优模型是模型 3.