

文字的视觉形象初论

——认知过程、区别特征结构及其模糊性

揭 春 雨

文字，是用来记录人类自然语言的书面符号系统，是语音体系之外的又一种传递语言信息的信息载体。语音，是耳治的；而文字，则是目治的，是用来视读的：视，然后读。视，是识别文字的图象；读，是发出声音和获得意义。在视读中，人们的视觉器官（或者触觉器官，如果是盲文的话）最先接触到的，是文字符号具体的图象信息，即文字的视觉形象。

因此，一种文字符号体系，不管它是通过所谓的表音方式，还是表意方式，或者其它别的什么方式来记录语言，它和语言无非是载体和信息的关系，衡量其优劣，视觉形象的好坏应当成为最为重要的标准之一。考

查和评价一种文字体系，还要看其符号图形是否适宜于我们的视觉器官的视读——文字符号的宜读程度。一种文字，如果符号的数量及其内部结构方式都恰到好处，但是符号与符号之间由于图形过分相似或者别的什么原因，导致人们在阅读中对它们的辨识发生困难，那么，这种文字的视觉形象便是不够理想的。

文字符号的视觉形象越是符合人们的视读认知心理规律，则它的宜读程度越高，这是显然的。分析文字的视觉形象，可以从视读的认知过程着手，这一认知过程如图1所示。

这是一个动态的心理过程，但图中的语

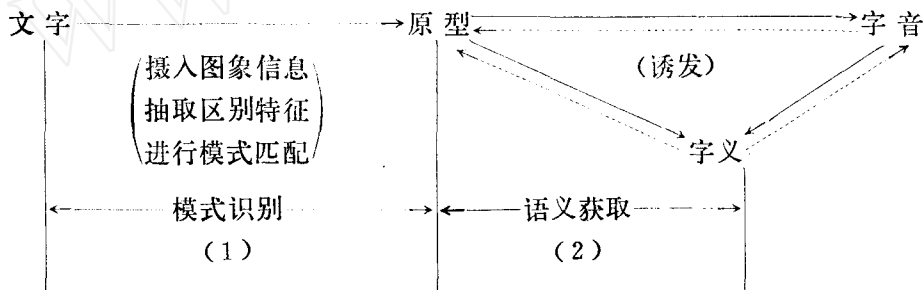


图1 视读认知过程

义获取部分，同时也表示了文字系统中字形、字义和字音三者之间的静态关系，这种关系是在语言文字习得时建立于人脑记忆里的。习得一种文字，简而言之，就是在字形、字义和字音三者间习得这种由传统确定下来的静态对应关系：每见到一个字形，即可唤醒对相应的字音和字形记忆；每听到一个字音，就知道它相应的字形和字义；每要表达一个语义单位，就知道相应需要哪些字，知道这些字的字形和字音是什么。建立

这种静态的对应关系，首先要完成的，是记住一个个字的图形，在对字形没有记忆以前，字形与字义、字音之间的对应无从谈起。

那么，习得文字时，我们又是怎样记忆文字的图形的呢？我们记忆的，当然不是一个个字的具体图形，这一点现代心理学已经给出过结论。同一个字（或字母、词），写得大一点或小一点，正一点或斜一点，它的具体图形都各不相同，直接记住这些千差万

别的图形相当困难。汉字,尤其是手写体,书写方式更是灵活不拘,恣意挥洒,五花八门,没有必要把这些无穷无尽的具体图形一个个都细微具体地记住。

我们记忆文字,就象我们记忆其它事物一样,记的是文字系统中每个字的原型。所谓原型,纯粹从心理学的概念上说,是表示某一类(范畴)客体的基本成分的一种抽象形式,它是从这类客体的各个具体经验中抽取出来的。一个原型,或者说某一个范畴的客体的原型,代表了这个特定范畴中所有客体的集中趋势。例如飞机,它的原型就是两个扁圆形翅膀和一个长筒形机身。各种型号的飞机,无论大小、形状差别如何,但作为飞机,它们在我们的认知中的集中趋势却是一样的:两个扁圆翅膀,一个长筒机身,三者构成了飞机区别于其它物体的基本成分。这些基本成分称为区别特征。同样地,我们不难理解,一个字的原型,就是从这个字(相当于前面所述的一个范畴)的千姿百态的无数个具体图形中抽取出来的、代表其具体图形的集中趋势的若干基本成分的一个有机组合。原型是一个字的所有区别特征的一个有序集。

虽然迄今为止,现代科学尚未能向我们直接宣示人脑记忆事物时所特有的区别特征的真实细节,但是,区别特征概念的提出,却大大加深了我们对人脑认知的了解。区别特征的存在已为一些心理实验所证实,在一些科技领域如模式识别,尤其是文字自动识别中,区别特征也得到卓有成效的运用。科技应用中,文字的区别特征的选取,主要是根据经验而人为确定的。

对于文字,我们也未能确定人脑记忆里文字的区别特征具体是什么内容,但是,我们可以一般性地认为:拼音文字中,一个字(或一个词)里,字母、字母的个数以及排

列情况,都是视觉信息的基本成分,属于区别特征的范畴;而对于方块汉字,则基本笔画、构字部件如偏旁、部首等,以及它们的数目、长短、大小、方向、位置,乃至相互间的配置方式,即布局,都是种种区别特征。拼音文字中作为区别特征元素的字母是从左到右线形排列的,可看作是一维图形。汉字的区别特征元素则是在每个汉字的方形框框内平面分布的,是二维图形。从复杂度上看,当元素数量均等时,平面分布的二维图形的复杂度是线性一维图形的复杂度的平方倍。汉字的构字部件多达数百,比之一般拼音文字的三、五十个字母,在数量上高出数十倍,再加上汉字区别特征在平面上的组合比之拼音文字的线形排列显得更为灵活,更为复杂,因此,在视觉上,汉字字形相互区别的性能,比之拼音文字存在一定程度的优越性是可信的。国内外的一些实验也直接或间接的支持或证实这一结论。

既然字形是由一些称之为区别特征的成分构成的,那么,习得文字时,自然就要从文字的一个个具体图形中抽取、总结出这些区别特征,再将它们构成一个个字的原型并且记住,之后,才能在字的原型和字音、字义之间建立确定的约定俗成的对应关系。这时,习得者才有能力对文字进行视读。

对文字的视读,正如视读认知过程所示,第一步,是模式识别,这是辨认出一个字是这个字而不是那个字的过程,也即是把一个字的图形和它的原型连接起来的过程;一种文字的视觉形象的优劣和宜读程度,体现在完成这一过程的难易快慢之中。第二步,是语义获取,这是获得一个字的原型之后,再从该原型联想或对应到它的字义和字音的过程;这是一个诱发过程,是模式识别这种心理过程完成之后必然发生的后续事件,大脑一旦捕捉到一个字的原型,就能自

动地根据记忆中已有的对应关系获得该原型所属的字义和字音（相类似地，字音也能诱发出字义和字形）。文字视觉形象的好坏、宜读程度的高低，对诱发过程的影响虽然较小，但对模式识别过程的影响却相当大，这种影响对于阅读是非常关键的。

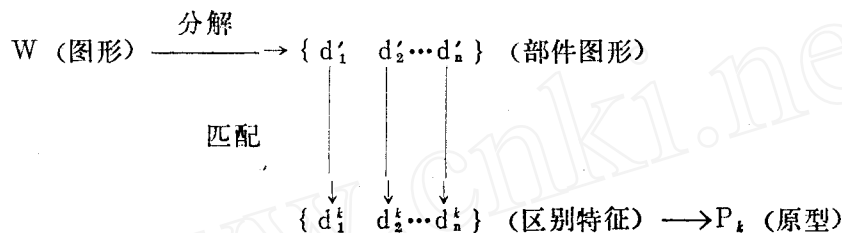
模式识别的完整过程是这样的：每一个字都有自己的原型，原型是人脑记忆文字图形的抽象形式，它是一个字的所有区别特征组成的有序集合；当视读到某一个字时，大脑首先通过视觉器官接受到这个字的具体图形信息，得到视觉形象，而视觉形象可以分解为一个一个区别特征图形，这些图形和人脑记忆装置中各个原型的区别特征进行匹配比

较，即可辨认出这个字的原型。如果用W表示文字的视觉形象，P表示原型，d表示区别特征，则记忆一个拥有m个字的文字符号系统，其记忆内容可用公式表示为：

$$\text{文字符号系统 } S = \sum_{i=1}^m P_i$$

$$P_i = \prod_{j=1}^{n_i} d_j^i$$

公式中， Σ 是“和”的数学符号，这里表示文字系统S由 P_1 、 $P_2 \dots P_m$ 等m个字符组成； Π 是“积”的数学符号，这里表示字符 P_i 是由 d_1^i 、 $d_2^i \dots d_{n_i}^i$ 等 n_i 个区别特征构成的一个有序集。而模式识别过程则如下所示：



这样，就识别出W的原型是 P_k ，从而得知W是系统S中的第k个字符。

至此可知，文字的识别、原型的获得，都依恃于区别特征。那么，文字符号的区别特征的多寡强弱，无疑就成了影响宜读程度和视觉形象的好坏的最直接的因素。视读中，常常需要小心辨认两个字的异同，特别是当遇到图形相当近似的字时，辨认时更需集中注意力。这时，文字的视觉形象的好坏就显示出至关重要的作用。

两个字的辨别，在认知上是根据它们的区别特征差做出心理决策的。区别特征差又分为绝对差和相对差两种。区别特征绝对差，并不是两个字的区别特征集作减法运算所得到的结果，而是两个区别特征集的两个差集的并集，用 ΔP 表示，则有：

$$\Delta P = (P_i - P_j) \cup (P_j - P_i)$$

如图2所示，区别特征绝对差 ΔP 是两个阴

影部分的总和：

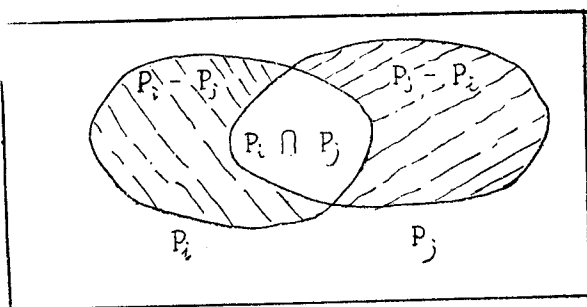


图2 区别特征集和区别特征差

如果两个字符的区别特征绝对差 ΔP 为空集，则说明它们区别特征集重叠，两字相同；否则，是不同的两个字。如果两个区别特征集的交集 $P_i \cap P_j$ 为空集，则两个区别特征集没有重叠部分，两个字完全不同，没有任何共同的区别特征。这里，两个字也即两个区别特征集的比较，是以一个文字符号系统的区别特征总集E为背景进行的。所谓区别特征总集，就是这种文字体系内部所有

的区别特征的总和。例如,汉字体系内部的所有构字部件的总和,构成汉字的区别特征总集。

区别特征相对差,则是两个区别特征集的绝对差 ΔP 和它们的并集 $P_1 \cup P_2$ 之比,用公式表示即:

$$\frac{\Delta P}{P_1 \cup P_2} = \frac{(P_1 - P_2) \cup (P_2 - P_1)}{P_1 \cup P_2}$$

这个公式主要是为了在形式上说明问题,并不旨在计算。如果要计算区别特征相对差的值,必须先根据模糊数学的原理给每个相关的区别特征按其区别作用的大小予以赋值。但是,应该说进行人为的赋值是带有很大的主观因素在内的,也是比较困难的,各个区别特征的区别作用互不相同,而且即使是同一个区别特征,在不同的字形环境里,它的区别作用也不尽相同,给它们赋值时很难找到同一的、前后完全一致的依据。这是模糊数学应用于文字研究的一个难点,有待于有志者的尝试和努力。

对文字的区别特征予以赋值之后,计算

$$\begin{aligned} \frac{|\Delta P|}{|P_1 \cup P_2|} &= \frac{V(\text{乚}) + V(\text{丿})}{|V(\text{丿}) + V(\text{儿}) + V(\text{又})| + |V(\text{乚}) + V(\text{儿}) + V(\text{又})|} \\ &= \frac{1 + 1}{(1 + 1 + 1) + (1 + 1 + 1)} \\ &= 1/3 \end{aligned}$$

这两个字的字形比较相似,因而它们的区别特征相对差的值比较小。对于大部分汉字来说,字形差别一般都比较大或完全不同,它们相互间的区别特征相对差的值多为接近于或等于1。

区别特征相对差决定着把两个字辨别开来的难易程度,差值大,辨别容易,差值越小,辨别难越。静态地看,两个字只要存在区别特征差,就能够把它们辨别开来。但是,书面语的阅读,总是在动态中,在一定的速度下进行的,相对于某个特定的阅读速度,两个之间的区别特征相对差值低于相应阈限,则视读认知过程中的模式识别将会发

公式中的分子、分母分别加上绝对值符号表示有值,公式变为:

$$\frac{|\Delta P|}{|P_1 \cup P_2|} = \frac{|(P_1 - P_2) \cup (P_2 - P_1)|}{|P_1 \cup P_2|}$$

例如,对于汉字中的“没”、“设”两个相似的字,我们有:

“没”字的区别特征集 $P_1 = \{\text{丿}, \text{儿}, \text{又}\}$

“设”字的区别特征集 $P_2 = \{\text{乚}, \text{儿}, \text{又}\}$

两个字的区别特征绝对差

$$\begin{aligned} \Delta P &= (P_1 - P_2) \cup (P_2 - P_1) \\ &= \{\text{丿}\} \cup \{\text{乚}\} \\ &= \{\text{丿}, \text{乚}\} \end{aligned}$$

两个字的区别特征相对差:

$$\frac{\Delta P}{P_1 \cup P_2} = \frac{\{\text{丿}, \text{乚}\}}{\{\text{丿}, \text{乚}, \text{儿}, \text{又}\}}$$

如果认为这些区别特征的区别作用相同,赋值时可以都赋值为1,用 $V(x)$ 表示元素 x 的区别特征值,即是:

$$V(\text{丿}) = V(\text{乚}) = V(\text{儿}) = V(\text{又}) = 1$$

这时,两字的区别特征相对差的值为:

生困难,甚至误识。这时,为了保证识别的正确率,只好把阅读速度减慢。如果一个文字符号系统,它的字符与字符之间的区别特征相对差值总是低到让人减慢阅读速度,则这种文字的视觉形象显然是不能为人们接受的。

视觉形象不够理想的例子,无论在汉语拼音还是现行汉字中,都可以大量找到,尤以汉语拼音为多。例如,汉语拼音体系中较为常见的同形词(即同音词)和混形词(或称近形词,即近音词),视觉形象上都存在着或大或小的问题。同形词的问题最为突出,对于它们,除了依赖于语义,没有其它

任何办法可以区别,混形词,阅读中对它们的识别也颇费力气,有时也不得不借助语义手段。阅读书面语,其目的本来是为了获取语义,而在这种情形下,却反过来要依靠语义来识别文字,多少总给人一种本末倒置的感觉。汉语拼音中许多词之间的区别特征相对差的确实是比较小,显然对于阅读不利,例如,zhǎnzhǎng(站长)和zhǎngzhēng(战争)两个词,各有九个字母、两个调号,只有一个字母和一个调号不一样,区别特征相对差是很小的。如果认为每个字母的区别作用相同,都赋值为1,标调号的区别作用相对小一点,赋值为0.5,那么,可以算得这两个词的区别特征相对差的计算值为:

$$\frac{|P_1 - P_2|}{|P_1| + |P_2|} = \frac{(1+0.5) + (1+0.5)}{(9 \times 1 + 2 \times 0.5) + (9 \times 1 + 2 \times 0.5)} = 0.15$$

这个值不到“没”、“设”两个近形汉字的区别特征相对差值的一半。

现行汉字里,也存在一些非常容易混淆的近形字,例如,“戊、戌、戍、戍”,“己、巳、已”等,为数不少。它们的区别特征相对差都非常小,视觉形象不够理想,很容易引起误读,即便有上下文做参考,最终总能识别,但也经常使得正常阅读不时卡一下壳,影响思维,降低阅读速度和效率。

因此,无论是汉语拼音的完善,还是现行汉字的改革,都应当充分考虑到文字的视觉形象的优劣问题,应当尽可能地增加同形词和混形词(字)之间的区别性特征,增大区别特征差,使得汉语拼音和现行汉字的视

觉形象都尽可能地完美一些。这样的工作我们不妨称之为“增繁”,对于文字系统中一些字、词的混形现象,它是有益的,也是必要的。

我们知道,字与字之间的差异性越大,即区别特征差越大,辨识它们就越容易。但是,不能因此得出结论说区别特征差越大就越好。汉字的许多字体中,有些区别特征本来就是冗余的,区别特征差越大,则意味着冗余区别特征越多,超过一定限度,就会“过犹不及”。

冗余区别特征对于视读是有利的,在汉语书面语的阅读中,我们往往只看清一个字的大体轮廓就能认出这个字来,便是得益于它们。冗余区别特征可以分为两类:一类能明显改善文字的视觉形象而不过分增加记忆负担和书写负担,是良性的;另一类是对文字的视觉形象没有什么特别好处而只是增加记忆和书写负担的,这类冗余区别特征必须尽可能地减少,简化汉字所做的许多有益的工作都出自这样的目的,取得比较理想的效果。但是,如果汉字一味简化,唯简是好,也是片面的。根据前面的论述,我们知道,有些汉字适当“增繁”一下也是必要的。虽然汉字“增繁”一说在我国几十年的文字工作中未有先例,似有异端之嫌,但笔者认为,它和简化实质上是相辅相成的、希望本文能引起人们,尤其是相关人士的适当关注。汉字无论是简化还是“增繁”,其目的都是为了使之更为完善、美好,而不应是别的。

(上接52页)

法未能进一步论定,待借到该书后再行补苴,特此说明。

- ① 《中国古籍善本书目》(经部),上海古籍出版社出版,1986,60页。
- ② 陆志韦先生在《记〈三教经书文字根本〉,附〈谐声韵学〉》一文中抄引的这四句题记的“司马温公集二十一母”一句脱一“订”字,见《陆

志韦近代汉语音韵论集》,商务印书馆,1952,12页;赵荫棠先生在《〈谐声韵学〉跋》一文中抄引这四句题记时则无“司马温公订集二十一母”一句,见《中原音韵研究》,商务印书馆,1956,311页。笔者所据的《三教经书文字根本》一书系中国社会科学院语言研究所藏本。

- ③ 《陆志韦近代汉语音韵论集》,商务印书馆,1988,122页。

④⑤⑥ 《中原音韵研究》,商务印书馆,1956,311页。