

高中数学必修3知识点

一：算法初步

1：算法的概念

(1) 算法概念：在数学上，现代意义上的“算法”通常是指可以用计算机来解决的某一类问题是程序或步骤，这些程序或步骤必须是明确和有效的，而且能够在有限步之内完成.

(2) 算法的特点:

① 有限性：一个算法的步骤序列是有限的，必须在有限操作之后停止，不能是无限的.

② 确定性：算法中的每一步应该是确定的并且能有效地执行且得到确定的结果，而不应当是模棱两可.

③ 顺序性与正确性：算法从初始步骤开始，分为若干明确的步骤，每一个步骤只能有一个确定的后继步骤，前一步是后一步的前提，只有执行完前一步才能进行下一步，并且每一步都准确无误，才能完成问题.

④ 不唯一性：求解某一个问题的解法不一定是唯一的，对于一个问题可以有不同的算法.

⑤ 普遍性：很多具体的问题，都可以设计合理的算法去解决，如心算、计算器计算都要经过有限、事先设计好的步骤加以解决.

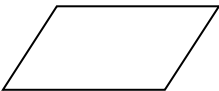
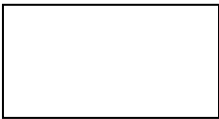
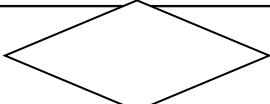
2：程序框图

(1) 程序框图基本概念:

① 程序构图的概念：程序框图又称流程图，是一种用规定的图形、指向线及文字说明来准确、直观地表示算法的图形。

一个程序框图包括以下几部分：表示相应操作的程序框；带箭头的流程线；程序框外必要文字说明。

② 构成程序框的图形符号及其作用

程序框	名称	功能
	起止框	表示一个算法的起始和结束，是任何流程图不可少的。
	输入、输出框	表示一个算法输入和输出的信息，可用在算法中任何需要输入、输出的位置。
	处理框	赋值、计算，算法中处理数据需要的算式、公式等分别写在不同的用以处理数据的处理框内。
	判断框	判断某一条件是否成立，成立时在出口处标明“是”或“Y”；不成立时标明“否”或“N”。

		“N°。
--	--	------

学习这部分知识的时候，要掌握各个图形的形状、作用及使用规则，画程序框图的规则如下：

- 1、使用标准的图形符号。
- 2、框图一般按从上到下、从左到右的方向画。
- 3、除判断框外，大多数流程图符号只有一个进入点和一个退出点。判断框具有超过一个退出点的唯一符号。
- 4、判断框分两大类，一类判断框“是”与“否”两分支的判断，而且有且仅有两个结果；另一类是多分支判断，有几种不同的结果。
- 5、在图形符号内描述的语言要非常简练清楚。

3： 算法的三种基本逻辑结构：顺序结构、条件结构、循环结构。

（1）顺序结构：顺序结构是最简单的算法结构，语句与语句之间，框与框之间是按从上到下的顺序进行的，它是由若干个依次执行的处理步骤组成的，它是任何一个算法都离不开的一种基本算法结构。

顺序结构在程序框图中的体现就是用流程线将程序框自上而下地连接起来，按顺序执行算法步骤。如在示意图中，A 框和 B 框是依次执行的，只有在执行完 A 框指定的操作后，才能接着执行 B 框所指定的操作。

（2）条件结构：条件结构是指在算法中通过对条件的判断根据条件是否成立而选择不同流向的算法结构。

条件 P 是否成立而选择执行 A 框或 B 框。无论 P 条件是否成立，只能执行 A 框或 B 框之一，不可能同时执行

A 框和 B 框，也不可能 A 框、B 框都不执行。一个判断结构可以有多个判断框。

（3）循环结构：在一些算法中，经常会出现从某处开始，按照一定条件，反复执行某一处理步骤的情况，这就是循环结构，反复执行的处理步骤为循环体，显然，循环结构中一定包含条件结构。循环结构又称重复结构，循环结构可细分为两类：

- ① 一类是当型循环结构，如下左图所示，它的功能是当给定的条件 P 成立时，执行 A 框，A 框执行完毕后，再判断条件 P 是否成立，如果仍然成立，再执行 A 框，如此反复执行 A 框，直到某一次条件 P 不成立为止，此时不再执行 A 框，离开循环结构。
- ② 另一类是直到型循环结构，如下右图所示，它的功能是先执行，然后判断给定的条件 P 是否成立，如果 P 仍然不成立，则继续执行 A 框，直到某一次给定的条件 P 成立为止，此时不再执行 A 框，离开循环结构。

p
当型循环结构

直到型循环结构

注意：1 循环结构要在某个条件下终止循环，这就需要条件结构来判断。因此，循环结构中一定包含条件结构，但不允许“死循环”。2 在循环结构中都有一个计数变量和累加变量。计数变量用于记录循环次数，累加变量用于输出结果。计数变量和累加变量一般是同步执行的，累加一次，计数一次。

4： 输入、输出语句和赋值语句

（1） 输入语句

①输入语句的一般格式

②输入语句的作用是实现在算法的输入信息功能；③“提示内容”提示用户输入什么样的信息，变量是指程序在运行时其值是可以变化的量；④输入语句要求输入的值只能是具体的常数，不能是函数、变量或表达式；⑤提示内容与变量之间用分号“；”隔开，若输入多个变量，变量与变量之间用逗号“，”隔开。

（2） 输出语句

①输出语句的一般格式

②输出语句的作用是实现在算法的输出结果功能；③“提示内容”提示用户输入什么样的信息，表达式是指程序要输出的数据；④输出语句可以输出常量、变量或表达式的值以及字符。

（3） 赋值语句

①赋值语句的一般格式

②赋值语句的作用是将表达式所代表的值赋给变量；③赋值语句中的“=”称作赋值号，与数学中的等号的意义是不同的。赋值号的左右两边不能对换，它将赋值号右边的表达式的值赋给赋值号左边的变量；④赋值语句左边只能是变量名字，而不是表达式，右边表达式可以是一个数据、常量或算式；⑤对于一个变量可以多次赋值。

注意：①赋值号左边只能是变量名字，而不能是表达式。如：**2=X**是错误的。②赋值号左右不能对换。如“**A=B**”“**B=A**”的含义运行结果是不同的。③不能利用赋值语句进行代数式的演算。（如化简、因式分解、解方程等）④赋值号“=”与数学中的等号意义不同。

5： 条件语句

（1） 条件语句的一般格式有两种：① **IFD THEND ELSE** 语句；② **IFD THEN** 语句。

①**IFD THEND ELSE** 语句 **IFD THEND ELSE** 语句的一般格式为图 1，对应的程序框图为图 2。

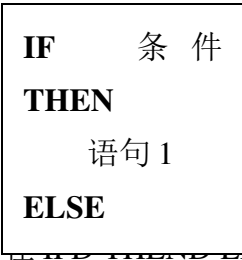


图 2

分析：**IFD THEND ELSE** 语句中，“条件”表示判断的条件，“语句 1”表示满足条件时执行的操作

内容：“语句 2”表示不满足条件时执行的操作内容；END IF 表示条件语句的结束。计算机在执行时，首先对 IF 后的条件进行判断，如果条件符合，则执行 THEN 后面的语句 1；若条件不符合，则执行 ELSE 后面的语句 2。

②IFD THEN 语句

IFD THEN 语句的一般格式为图 3, 对应的程序框图为图 4。

注意：“条件”表示判断的条件；“语句”表示满足条件时执行的操作内容，条件不满足时，结束程序；END IF 表示条件语句的结束。计算机在执行时首先对 IF 后的条件进行判断，如果条件符合就执行 THEN 后边的语句，若条件不符合则直接结束该条件语句，转而执行其它语句。

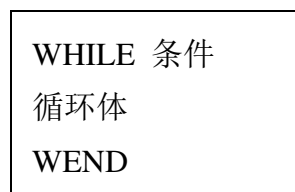
6: 循环语句

循环结构是由循环语句来实现的。对应于程序框图中的两种循环结构，一般程序设计语言中也有当型（WHILE 型）和直到型（UNTIL 型）两种语句结构。即 WHILE 语句和 UNTIL 语句。

(1) WHILE 语句

①WHILE 语句的一般格式是

对应的程序框图是



②当计算机遇到 **WHILE** 语句时，先判断条件的真假，如果条件符合，就执行 **WHILE** 与 **WEND** 之间的循环体；然后再检查上述条件，如果条件仍符合，再次执行循环体，这个过程反复进行，直到某一次条件不符合为止。这时，计算机将不执行循环体，直接跳到 **WEND** 语句后，接着执行 **WEND** 之后的语句。因此，当型循环有时也称为“前测试型”循环。

(2) UNTIL 语句

①UNTIL 语句的一般格式是

对应的程序框图是



② 直到型循环又称为“后测试型”循环，从 UNTIL 型循环结构分析，计算机执行该语句时，先执行一次循环体，然后进行条件的判断，如果条件不满足，继续返回执行循环体，然后再进行条件的判断，这个过程反复进行，直到某一次条件满足时，不再执行循环体，跳到 LOOP UNTIL 语句后执行其他语句，是先执行循环体后进行条件判断的循环语句。

分析：当型循环与直到型循环的区别：（先由学生讨论再归纳）

(1) 当型循环先判断后执行，直到型循环先执行后判断：

在 WHILE 语句中, 是当条件满足时执行循环体, 在 UNTIL 语句中, 是当条件不满足时执行循环

7: 辗转相除法与更相减损术

(1) 辗转相除法。也叫欧几里德算法，用辗转相除法求最大公约数的步骤如下：

①用较大的数 m 除以较小的数 n 得到一个商 S_0 和一个余数 R_0 ； ②若 $R_0 = 0$ ，则 n 为 m ， n 的最大公约数；若 $R_0 \neq 0$ ，则用除数 n 除以余数 R_0 得到一个商 S_1 和一个余数 R_1 ； ③若 $R_1 = 0$ ，则 R_1 为 m ， n 的最大公约数；若 $R_1 \neq 0$ ，则用除数 R_0 除以余数 R_1 得到一个商 S_2 和一个余数 R_2 ；…… 依次计算直至 $R_n = 0$ ，此时所得到的 R_{n-1} 即为所求的最大公约数。

(2) 更相减损术

我国早期也有求最大公约数问题的算法，就是更相减损术。在《九章算术》中有更相减损术求最大公约数的步骤：可半者半之，不可半者，副置分母·子之数，以少减多，更相减损，求其等也，以等数约之。翻译为：①任意给出两个正数；判断它们是否都是偶数。若是，用 2 约简；若不是，执行第二步。②以较大的数减去较小的数，接着把较小的数与所得的差比较，并以大数减小数。继续这个操作，直到所得的数相等为止，则这个数（等数）就是所求的最大公约数。

(3) 辗转相除法与更相减损术的区别：

- ①都是求最大公约数的方法，计算上辗转相除法以除法为主，更相减损术以减法为主，计算次数上辗转相除法计算次数相对较少，特别当两个数字大小区别较大时计算次数的区别较明显。
- ②从结果体现形式来看，辗转相除法体现结果是以相除余数为 0 则得到，而更相减损术则以减数与差相等而得到

8: 秦九韶算法与排序

(1) 秦九韶算法概念：

$f(x) = a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0$ 求值问题

$$\begin{aligned} f(x) &= a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0 = (a_n x^{n-1} + a_{n-1} x^{n-2} + \dots + a_1) x + a_0 \\ &= ((a_n x^{n-2} + a_{n-1} x^{n-3} + \dots + a_2) x + a_1) x + a_0 \\ &= \dots = (\dots (a_n x + a_{n-1}) x + a_{n-2}) x + \dots + a_1) x + a_0 \end{aligned}$$

求多项式的值时，首先计算最内层括号内依次多项式的值，即 $v_1 = a_n x + a_{n-1}$ 然后由内向外逐层计算一次多项式的值，即 $v_2 = v_1 x + a_{n-2}$ $v_3 = v_2 x + a_{n-3}$ $v_n = v_{n-1} x + a_0$
这样，把 n 次多项式的求值问题转化成求 n 个一次多项式的值的问题。

(2) 两种排序方法：直接插入排序和冒泡排序

①直接插入排序

基本思想：插入排序的思想就是读一个，排一个。将第 1 个数放入数组的第 1 个元素中，以后读入的数与已存入数组的数进行比较，确定它在从大到小的排列中应处的位置。将该位置以及以后的元素向后推移一个位置，将读入的新数填入空出的位置中。（由于算法简单，可以举例说明）

②冒泡排序

基本思想：依次比较相邻的两个数，把大的放前面，小的放后面。即首先比较第 1 个数和第 2 个数，大数放前，小数放后。然后比较第 2 个数和第 3 个数……直到比较最后两个数。第一趟结束，最小的一定沉到最后。重复上过程，仍从第 1 个数开始，到最后第 2 个数……由于在排序过程中总是大数往前，小数往后，相当气泡上升，所以叫冒泡排序。

9：进位制

（1）概念：进位制是一种记数方式，用有限的数字在不同的位置表示不同的数值。可使用数字符号的个数称为基数，基数为 n ，即可称 n 进位制，简称 n 进制。现在最常用的是十进制，通常使用 10 个阿拉伯数字 0-9 进行记数。对于任何一个数，我们可以用不同的进位制来表示。比如：十进数 57，可以用二进制表示为 111001，也可以用八进制表示为 71、用十六进制表示为 39，它们所代表的数值都是一样的。

一般地，若 k 是一个大于 1 的整数，那么以 k 为基数的 k 进制可以表示为：

$$a_n a_{n-1} \dots a_1 a_0 (k) \quad (0 < a_n < k, 0 \leq a_{n-1}, \dots, a_1, a_0 < k),$$

而表示各种进位制数一般在数字右下脚加注来表示，如 $111001_{(2)}$ 表示二进制数， $34_{(5)}$ 表示 5 进制数

二：统计

1：简单随机抽样

（1）总体和样本

① 在统计学中，把研究对象的全体叫做总体。② 把每个研究对象叫做个体。③ 把总体中个体的总数叫做总体容量。

④ 为了研究总体 x 的有关性质，一般从总体中随机抽取一部分： $x_1, x_2, \dots, x_n$ 研究，我们称它为样本。其中个体的个数称为样本容量。

（2）简单随机抽样，也叫纯随机抽样。就是从总体中不加任何分组、划类、排队等，完全随机地抽取调查单位。特点是：每个样本单位被抽中的可能性相同（概率相等），样本的每个单位完全独立，彼此间无一定的关联性和排斥性。简单随机抽样是其它各种抽样形式的基础。通常只是在总体单位之间差异程度较小和数目较少时，才采用这种方法。

（3）简单随机抽样常用的方法：

① 抽签法 ② 随机数表法 ③ 计算机模拟法 ④ 使用统计软件直接抽取。

在简单随机抽样的样本容量设计中，主要考虑：① 总体变异情况；② 允许误差范围；③ 概率保证程度。

(4) 抽签法:

- ①给调查对象群体中的每一个对象编号; ②准备抽签的工具, 实施抽签;
- ③对样本中的每一个个体进行测量或调查

(5) 随机数表法:

2: 系统抽样

(1) 系统抽样(等距抽样或机械抽样):

把总体的单位进行排序, 再计算出抽样距离, 然后按照这一固定的抽样距离抽取样本。第一个样本采用简单随机抽样的办法抽取。 $K(\text{抽样距离}) = N(\text{总体规模}) / n(\text{样本规模})$

前提条件: 总体中个体的排列对于研究的变量来说, 应是随机的, 即不存在某种与研究变量相关的规则分布。可以在调查允许的条件下, 从不同的样本开始抽样, 对比几次样本的特点。如果有明显差别, 说明样本在总体中的分布承某种循环性规律, 且这种循环和抽样距离重合。

(2) 系统抽样, 即等距抽样是实际中最为常用的抽样方法之一。因为它对抽样框的要求较低, 实施也比较简单。更为重要的是, 如果有某种与调查指标相关的辅助变量可供使用, 总体单元按辅助变量的大小顺序排队的话, 使用系统抽样可以大大提高估计精度。

3: 分层抽样

(1) 分层抽样(类型抽样):

先将总体中的所有单位按照某种特征或标志(性别、年龄等)划分成若干类型或层次, 然后再在各个类型或层次中采用简单随机抽样或系用抽样的办法抽取一个子样本, 最后, 将这些子样本合起来构成总体的样本。

两种方法:

- ①先以分层变量将总体划分为若干层, 再按照各层在总体中的比例从各层中抽取。
- ②先以分层变量将总体划分为若干层, 再将各层中的元素按分层的顺序整齐排列, 最后用系统抽样的方法抽取样本。

(2) 分层抽样是把异质性较强的总体分成一个个同质性较强的子总体, 再抽取不同的子总体中的样本分别代表该子总体, 所有的样本进而代表总体。

分层标准:

- ①以调查所要分析和研究的主要变量或相关的变量作为分层的标准。
- ②以保证各层内部同质性强、各层之间异质性强、突出总体内在结构的变量作为分层变量。
- ③以那些有明显分层区分的变量作为分层变量。

(3) 分层的比例问题: 抽样比 = $\frac{\text{样本容量}}{\text{个体容量}} = \frac{\text{各层样本容量}}{\text{各层个体容量}}$

- ① 按比例分层抽样：根据各种类型或层次中的单位数目占总体单位数目的比重来抽取子样本的方法。
- ② 不按比例分层抽样：有的层次在总体中的比重太小，其样本量就会非常少，此时采用该方法，主要是便于对不同层次的子总体进行专门研究或进行相互比较。如果要用样本资料推断总体时，则需要先对各层的数据资料进行加权处理，调整样本中各层的比例，使数据恢复到总体中各层实际的比例结构。

类别	共同点	各自特点	相互关系	适用范围
简单随机抽样	抽样过程中每个个体被抽取的机会相等	从总体中逐个抽取		总体中的个体数较少
系统抽样		将总体均匀分成几部分，按事先确定的规则在各部分抽取	再起时部分抽样时采用简单随机抽样	总体中的个数较多
分成抽样		经总体分成几层，分层进行抽取	各层抽样时采用简单随机抽样	总体由差异明显的几部分组成

4：用样本的数字特征估计总体的数字特征

(1) 样本均值：

$$\bar{x} = \frac{x_1 + x_2 + \cdots + x_n}{n}$$

(2) 样本标准差：

$$s = \sqrt{s^2} = \sqrt{\frac{(x_1 - \bar{x})^2 + (x_2 - \bar{x})^2 + \cdots + (x_n - \bar{x})^2}{n}}$$

- 用样本估计总体时，如果抽样的方法比较合理，那么样本可以反映总体的信息，但从样本得到的信息会有偏差。在随机抽样中，这种偏差是不可避免的。
- 虽然我们用样本数据得到的分布、均值和标准差并不是总体的真正的分布、均值和标准差，而只是一个估计，但这种估计是合理的，特别是当样本量很大时，它们确实反映了总体的信息。
- (3) 众数：在样本数据中，频率分布最大值所对应的样本数据（可以是多个）。
 - (4) 中位数：在样本数据中，累计频率为 1.5 时所对应的样本数据值（只有一个）。

注意：

① 如果把一组数据中的每一个数据都加上或减去同一个共同的常数，标准差不变
 ② 如果把一组数据中的每一个数据乘以一个共同的常数 k，标准差变为原来的 k 倍
 ③ 一组数据中的最大值和最小值对标准差的影响，区间 $(\bar{x} - 3s, \bar{x} + 3s)$ 的应用；

“去掉一个最高分，去掉一个最低分” 中的科学道理

5:用样本的频率分布估计总体分布

- 1：频率分布表与频率分布直方图
- 频率分布表盒频率分布直方图，是从各个小组数据在样本容量中所占比例大小的角度，来表示数据分布规律，它可以帮助我们看到整个样本数据的频率分布情况。
- 具体步骤如下：

- 第一步：求极差，即计算最大值与最小值的差.
- 第二步：决定组距和组数：组距与组数的确定没有固定标准，需要尝试、选择，力求有合适的组数，以能把数据的规律较清楚地呈现为准.太多或太少都不好，不利对数据规律的发现.组数应与样本的容量有关，样本容量越大组数越多.一般来说，容量不超过 100 的组数在 5 至 12 之间.组距应最好“取整”，它与 $\frac{\text{极差}}{\text{组距}}$ 有关.

注意：组数的“取舍”不依据四舍五入，而是当 $\frac{\text{极差}}{\text{组距}}$ 不是整数时，组数= $\lceil \frac{\text{极差}}{\text{组距}} \rceil +1$.

- ② 频率分布折线图：连接频率分布直方图中各个小长方形上端的重点，就得到频率分布折线图。
- ③ 总体密度曲线：总体密度曲线反映了总体在各个范围内取值的半分比，它能给我们提供更加精细的信息。

2：茎叶图：茎是指中间的一列数，叶是指从茎旁边生长出来的数。

例：例如：为了了解某地区高三学生的身体发育情况，抽查了地区内 100 名年龄为 17.5~18 岁的男生的体重情况，结果如下（单位：kg）.

56.5	69.5	65	61.5	64.5	76	71	66	63.5	56
66.5	64	64.5	76	58.5	59.5	63.5	65	70	74.5
72	73.5	56	67	70	68.5	64	55.5	72.5	66.5
57.5	65.5	68	71	75	68	76	57.5	60	71.5
62	68.5	62.5	66	59.5	57	69.5	74	64.5	59
63.5	64.5	67.5	73	68	61.5	67	68	63.5	58
55	72	66.5	74	63	59	65.5	62.5	69.5	72
60	55.5	70	64.5	58	64.5	75.5	68.5	64	62
64	70.5	57	62.5	65	65.5	58.5	67.5	70.5	65
69	71.5	73	62	58	66	66.5	70	63	59.5

试根据上述数据画出样本的频率分布直方图，并对相应的总体分布作出估计.

解：按照下列值的差

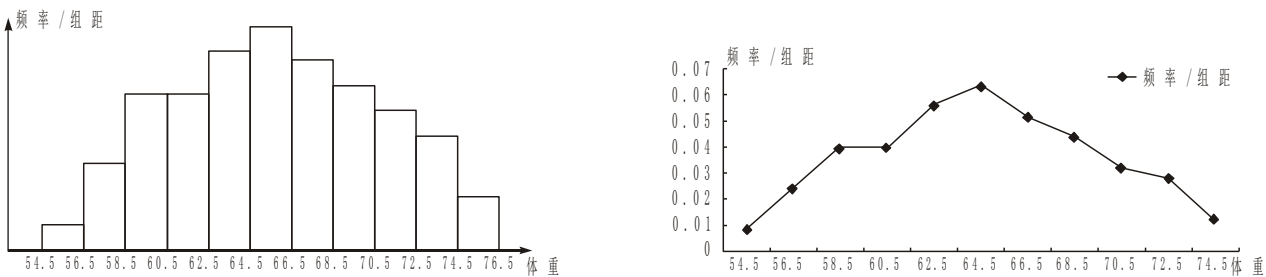
- (1) 求最大值与最小计.在上述数据中，最大值是 76，最小值是 55，极差是 76-55=21.
- (2) 确定组距与组数.如果将组距定为 2，那么由 21÷2=10.5，组数为 11，这个组数适合的.于是组距为 2，组数为 11.
- (3) 决定分点.根据本例中数据的特点，第 1 小组的起点可取为 54.5，第 1 小组的终点可取为 56.5，为了避免一个数据既是起点，又是终点而造成重复计算，我们规定分组的区间是“左闭右开”的.这样，所得到的分组是
- [54.5, 56.5)， [56.5, 58.5)， ..., [74.5, 76.5) .

(4) 列频率分布表.

分组	频数	频率	累计频率
[54.5, 56.5)	2	0.02	0.02
[56.5, 58.5)	6	0.06	0.08
[58.5, 60.5)	10	0.10	0.18
[60.5, 62.5)	10	0.10	0.28
[62.5, 64.5)	14	0.14	0.42
[64.5, 66.5)	16	0.16	0.58
[66.5, 68.5)	13	0.13	0.71
[68.5, 70.5)	11	0.11	0.82
[70.5, 72.5)	8	0.08	0.90
[72.5, 74.5)	7	0.07	0.97
[74.5, 76.5)	3	0.03	1.00

合计	100	1.00	
----	-----	------	--

(5) 绘制频率分布直方图.
 频率分布直方如图 2—2—3 所示.



连接频率直方图中各小长方形上端的中点, 就得到频率分布折线图.如图 2—2—4 所示.

例 2: 某赛季甲、乙两名篮球运动员每场比赛得分情况如下
 甲的得分: 15, 21, 25, 31, 36, 39, 31, 45, 36, 48, 24, 50, 37;
 乙的得分: 13, 16, 23, 25, 28, 33, 38, 14, 8, 39, 51.
 上述的数据可以用下图来表示, 中间数字表示得分的十位数, 两边数字分别表示两个人各场比赛得分的个位数.

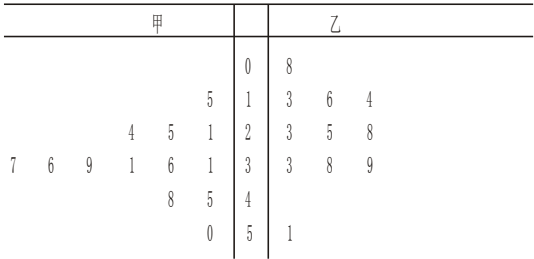


图 2—2—5

通常把这样的图叫做茎叶图.请根据上图对两名运动员的成绩进行比较.
 从这个茎叶图上可以看出, 甲运动员的得分情况是大致对称的, 中位数是 36; 乙运动员的得分情况除一个特殊得分外, 也大致对称, 中位数是 25.因此甲运动员发挥比较稳定, 总体得分情况比乙好.
 用茎叶图表示有两个突出的优点:其一, 从统计图上没有信息的损失, 所有的信息都可以从这个茎叶图中得到; 其二, 茎叶图可以在比赛时随时记录, 方便记录与表示.但茎叶图只能表示两位的整数, 虽然可以表示两个人以上的比赛结果 (或两个以上的记录), 但没有两个记录表示得那么直观, 清晰.
 6: 变量间的相关关系: 自变量取值一定时因变量的取值带有一定随机性的两个变量之间的关系交相关关系。对具有相关关系的两个变量进行统计分析的方法叫做回归分析。

(1) 回归直线: 根据变量的数据作出散点图, 如果各点大致分布在一条直线的附近, 就称这两个变量之间具有线性相关的关系, 这条直线叫做回归直线方程。如果这些点散布在从左下角到右上角的区域, 我们就成这两个变量呈正相关; 若从左上角到右下角的区域, 则称这两个变量呈负相关。

设已经得到具有线性相关关系的一组数据:

x	$\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})$	$\sum_{i=1}^n x_i y_i$	$n \bar{x} \bar{y}$
y	$\sum_{i=1}^n y_i$	$\sum_{i=1}^n (x_i - \bar{x})^2$	$\sum_{i=1}^n x_i^2$
	$\bar{y} = \frac{\sum_{i=1}^n y_i}{n}$	$\bar{x} = \frac{\sum_{i=1}^n x_i}{n}$	$\bar{y}_2$
	$a = \bar{y} - b \bar{x}$		

所要求的回归直线方程为： $\hat{y} = bx + a$ ，其中， a, b 是待定的系数。

(2) 回归直线过

x	x_1	$\circ \circ \circ$	x_n
y	y_1	$\circ \circ \circ$	y_n

的样本中心点 $(\bar{x}, \bar{y})$

三：概 率

1：随机事件的概率及概率的意义

- (1) 必然事件：在条件 S 下，一定会发生的事件，叫相对于条件 S 的必然事件；
- (2) 不可能事件：在条件 S 下，一定不会发生的事件，叫相对于条件 S 的不可能事件；
- (3) 确定事件：必然事件和不可能事件统称为相对于条件 S 的确定事件；
- (4) 随机事件：在条件 S 下可能发生也可能不发生的事件，叫相对于条件 S 的随机事件；
- (5) 频数与频率：在相同的条件 S 下重复 n 次试验，观察某一事件 A 是否出现，称 n 次试验中事件 A

出现的次数 n_A 为事件 A 出现的频数；称事件 A 出现的比例 $f_n(A) = \frac{n_A}{n}$ 为事件 A 出现的概率：对于

给定的随机事件 A，如果随着试验次数的增加，事件 A 发生的频率 $f_n(A)$ 稳定在某个常数上，把这个常数记作 $P(A)$ ，称为事件 A 的概率。

- (6) 频率与概率的区别与联系：随机事件的频率，指此事件发生的次数 n_A 与试验总次数 n 的比值 $\frac{n_A}{n}$ ，

它具有一定的稳定性，总在某个常数附近摆动，且随着试验次数的不断增多，这种摆动幅度越来越小。我们把这个常数叫做随机事件的概率，概率从数量上反映了随机事件发生的可能性的。频率在大量重复试验的前提下可以近似地作为这个事件的概率

2: 概率的基本性质

- (1) 必然事件概率为 1, 不可能事件概率为 0, 因此 $0 \leq P(A) \leq 1$
- (2) 事件的包含、并事件、交事件、相等事件
- (3) 若 $A \cap B$ 为不可能事件, 即 $A \cap B = \emptyset$, 那么称事件 A 与事件 B 互斥;
- (4) 若 $A \cap B$ 为不可能事件, $A \cup B$ 为必然事件, 那么称事件 A 与事件 B 互为对立事件;
- (5) 当事件 A 与 B 互斥时, 满足加法公式: $P(A \cup B) = P(A) + P(B)$;

若事件 A 与 B 为对立事件, 则 $A \cup B$ 为必然事件, 所以 $P(A \cup B) = P(A) + P(B) = 1$, 于是有 $P(A) = 1 - P(B)$

(6) 互斥事件与对立事件的区别与联系, 互斥事件是指事件 A 与事件 B 在一次试验中不会同时发生, 其具体包括三种不同的情形: ① 事件 A 发生且事件 B 不发生; ② 事件 A 不发生且事件 B 发生; ③ 事件 A 与事件 B 同时不发生, 而对立事件是指事件 A 与事件 B 有且仅有一个发生, 其包括两种情形: ④ 事件 A 发生 B 不发生; ⑤ 事件 B 发生事件 A 不发生, 对立事件互斥事件的特殊情形。

3: 基本事件

- (1) 基本事件: 基本事件是在一次试验中所有可能发生的基本结果中的一个, 它是试验中不能再分的最简单的随机事件。
- (2) 基本事件的特点: ①任何两个基本事件是互斥的②任何事件 (除不可能事件外) 都可以表示成基本事件的和。

4: 古典概型:

- (1) 古典概型的条件: 古典概型是一种特殊的数学模型, 这种模型满足两个条件:
 - ① 试验结果的有限性和所有结果的等可能性。
 - ② 所有基本事件必须是有限个。
- (2) 古典概型的解题步骤:
 - ① 求出总的基本事件数;

② 求出事件 A 所包含的基本事件数, 然后利用公式 $p(A) = \frac{A \text{ 所包含的基本事件的个数}}{\text{总的基本事件个数}}$

5: 几何概型

- (1) 几何概率模型: 如果每个事件发生的概率只与构成该事件区域的长度 (面积或体积) 成比例, 则称这样的概率模型为几何概率模型;
- (2) 几何概型的概率公式: $p(A) = \frac{\text{构成事件A的区域长度 (面积或体积)}}{\text{试验的全部结果所构成的区域长度 (面积或体积)}}$;
- (3) 几何概型的特点: ①试验中所有可能出现的结果 (基本事件) 有无限多个; ②每个基本事件出现的可能性相等。

注意: 几何概型也是一种概率模型, 它与古典概型的区别是试验的可能结果不是有限个。其特点是在一个区域内均匀分布, 所以随机事件的概率大小与随机事件所在区域的形状位置无关, 值域该区域的大小有关。如果随即事件所在区域是一个单点, 由于单点的长度、面积、体积均为 0, 则它出现的概率为 0, 但它不是不可能事件; 如果一个随机事件所在区域是全部区域扣除一个单点, 则

它出现的概率为 1，但他不是必然事件。

综上可得：必然事件的概率为 1；不可能事件的概率为 0。

概率为 1 的事件不一定为必然事件；概率为 0 的事件不一定为不可能事件。