

第3章

03

试验数据处理

3.1 概述

数据处理是数理统计学中的一部分重要内容，它主要研究试验测量或观察数据分析计算的处理方法，从而得出可靠和规律性的结果，并依据这个规律和结果对工业生产、农业生产、天气、地震等进行预报和控制，进而掌握客观事物的发展规律。同时，材料试验数据处理是整理试验数据、精准评价试验结果的关键。

数据处理的具体方法包括以下几种。

- (1) 参数估计：主要对某些重要参数表进行点估计和区间估计。
- (2) 假设检验：判断各种数据处理结果的可靠性程度。
- (3) 方差分析：分析各影响因子对考察指标影响的显著性程度。
- (4) 回归分析：分析如何获得反映事物客观规律的数学表达式。

3.2 参数估计

由中心极限定理和大数定理可知，只要抽样为大样本，无论总体是否服从正态分布，其样本平均数都近似服从 $N(\mu, \sigma_y^2)$ 的正态分布，故在给定某一概率水准 α （置信

度 $p=1-\alpha$) 下, 有

$$p(\bar{y} - u_{\alpha/2}\sigma_{\bar{y}} \leq \mu \leq \bar{y} + u_{\alpha/2}\sigma_{\bar{y}}) = 1 - \alpha$$

式中, $u_{\alpha/2}$ 为正态分布下置信度 $p=1-\alpha$ 时的 u 临界值。

结果表明: 尽管只知样本均数 $\bar{y}$ 而未知总体均数 μ , 但可知区间 $(\bar{y} - u_{\alpha/2}\sigma_{\bar{y}}, \bar{y} + u_{\alpha/2}\sigma_{\bar{y}})$ 包含在内的可靠程度为 $1-\alpha$, 因而将其称为 μ 的 $1-\alpha$ 置信区间 (confidence interval)。其中, $\bar{y} - u_{\alpha/2}\sigma_{\bar{y}}$ 称为 μ 的 $1-\alpha$ 置信区间的下限, 记为 “ L_L ”; $\bar{y} + u_{\alpha/2}\sigma_{\bar{y}}$ 称为 μ 的 $1-\alpha$ 置信区间的上限, 记为 “ L_H ”。

区间 (L_L, L_H) 称为采样样本均数 $\bar{y}$ 对总体均数 μ 的置信度为 $p=1-\alpha$ 的区间估计 (interval estimation)。

$L = \bar{y} \pm u_{\alpha/2}\sigma_{\bar{y}}$ 称为采用样本均数 $\bar{y}$ 对总体均数 μ 的置信度为 $p=1-\alpha$ 的点估计 (point estimation)。

参数估计按其估计方式的不同可分为两类: 点估计与区间估计。所谓点估计, 就是以统计量的值作为母体参数的估计值。所谓区间估计, 则是估计母体参数以某一概率包含在一个什么样的区间中。

3.2.1 点估计

将样本统计量直接作为总体相应参数的估计值。

点估计只给出了未知参数估计值的大小, 没有考虑试验误差的影响, 也没有指出估计的可靠程度。

主要介绍以下两种构造点估计常用的方法。

方法 1: 矩法, 即用样本矩估计总体矩。

(1) 用样本均数估计总体均数: $\hat{\mu} = \frac{1}{n}(y_1 + y_2 + \cdots + y_n) = \bar{y}$ 。

(2) 用样本方差估计总体方差: $\hat{\sigma}^2 = \frac{1}{n}[(x_1 - \bar{x})^2 + (x_2 - \bar{x})^2 + \cdots + (x_n - \bar{x})^2] = S^2$ 。

(3) 用相关系数 ρ 的估计: $\rho(\xi, \eta) = \frac{\text{Cov}(\xi, \eta)}{\sqrt{V(\xi)V(\eta)}}$ (试验数据的整理与分析)。

其中, $V(\xi)$ 与 $V(\eta)$ 可分别用相对应的子样方差来估计, 而协变方差 $\text{Cov}(\xi, \eta)$ 按定义为

$$\int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} [x - E(\xi)][y - E(\eta)]p(x, y) dx dy$$

可以设想用以下表达式来估计

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

这样, 可构造 $\rho(\xi, \eta)$ 的估计量为

$$\hat{\rho} = \frac{\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{S_x S_y} = r$$

此等式的右部即子样相关系数。因此, 可以用子样相关系数 r 作为母体相关系数的估计量。

再把上面所得的结果归纳为如下描述。

设母体平均为 μ , 母体方差为 σ^2 , 自母体中抽取的随机子样为 $(x_1, x_2, \dots, x_n)$, 则 μ 与 σ^2 的估计量分别为

$$\hat{\mu} = \bar{x} = \frac{1}{n} (x_1 + x_2 + \dots + x_n)$$

$$\hat{\sigma}^2 = S^2 = \frac{1}{n} [(x_1 - \bar{x})^2 + (x_2 - \bar{x})^2 + \dots + (x_n - \bar{x})^2]$$

设二元分布母体的母体相关系数为 ρ , 自母体中随机抽取的子样为 $(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)$, 则 ρ 的估计量为

$$\hat{\rho} = r = \frac{\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{S_x S_y}$$

其中

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i, \quad \bar{y} = \frac{1}{n} \sum_{i=1}^n y_i, \quad S_x = \sqrt{\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2}, \quad S_y = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \bar{y})^2}$$

方法 2: 极大似然法, 于 1912 年由英国统计学家 R.A. Fisher 提出, 利用样本分布密度构造似然函数, 从而求出参数的最大似然估计。

(1) 用样本均数估计总体均数: $\hat{\mu} = \bar{y}$ 。

(2) $\hat{\sigma}^2 = \frac{n-1}{n} S^2$ 。

显然, 均值 μ 的矩估计量与极大似然估计量一致, 而 σ^2 的极大似然估计量与矩估计量不同, 多了一个系数 $(n-1)/n$ 。当样本容量 n 较大时, 两者相差甚微, 而 S^2 是总体方差的无偏估计量, 这就是为何用 $\frac{1}{n-1} \sum_{i=1}^n (y_i - \bar{y})^2$ 表示样本方差。

3.2.2 区间估计

区间估计能够在一定概率保证下指出总体参数的可能范围。其中, 所给出的可能

范围称为置信区间 (confidence interval), 所给出的概率保证称为置信度或置信概率 (confidence probability)。

1. 一个正态总体的区间估计

(1) 总体均数 μ 的区间估计。

σ^2 已知: μ 的 $100(1-\alpha)\%$ 的置信区间为

$$\left[\bar{y} \pm u_{\alpha/2} \cdot \frac{\sigma}{\sqrt{n}} \right]$$

σ^2 未知: 对大样本, μ 的 $100(1-\alpha)\%$ 的置信区间为

$$\left[\bar{y} \pm u_{\alpha/2} \cdot \frac{S}{\sqrt{n}} \right]$$

对小样本, μ 的 $100(1-\alpha)\%$ 的置信区间为

$$\left[\bar{y} \pm t_{\alpha/2, n-1} \cdot \frac{S}{\sqrt{n}} \right]$$

(2) 方差 σ^2 的区间估计。

σ^2 的 $100(1-\alpha)\%$ 的置信区间为

$$\left[\frac{(n-1)S^2}{X_{\alpha/2, (n-1)}^2}, \frac{(n-1)S^2}{X_{1-\alpha/2, (n-1)}^2} \right]$$

2. 两个正态总体的区间估计

(1) 总体均数差 $\mu_I - \mu_{II}$ 的区间估计。

σ_I^2 与 σ_{II}^2 已知: $\mu_I - \mu_{II}$ 的 $100(1-\alpha)\%$ 的置信区间为

$$\left[\bar{y}_I - \bar{y}_{II} \pm u_{\alpha/2} \cdot \sqrt{\frac{\sigma_I^2}{n_I} + \frac{\sigma_{II}^2}{n_{II}}} \right]$$

σ_I^2 与 σ_{II}^2 未知: 对大样本 (n_I 、 n_{II} 很大), $\mu_I - \mu_{II}$ 的 $100(1-\alpha)\%$ 的置信区间为

$$\left[\bar{y}_I - \bar{y}_{II} \pm u_{\alpha/2} \cdot \sqrt{\frac{S_I^2}{n_I} + \frac{S_{II}^2}{n_{II}}} \right]$$

对小样本 (n_I 、 n_{II} 不是很大, 且 $n_I \neq n_{II}$), $\mu_I - \mu_{II}$ 的 $100(1-\alpha)\%$ 的置信区间为

$$\left[\bar{y}_I - \bar{y}_{II} \pm t_{\alpha/2, (n_I+n_{II}-2)} \cdot \sqrt{\frac{(n_I-1)S_I^2 + (n_{II}-1)S_{II}^2}{n_I+n_{II}-2} \left(\frac{1}{n_I} + \frac{1}{n_{II}} \right)} \right]$$

简记为

$$\left[\bar{y}_I - \bar{y}_{II} \pm t_{\alpha/2, (n_I+n_{II}-2)} \cdot S_{\bar{y}_I - \bar{y}_{II}} \right]$$

(2) 方差比 $\sigma_I^2 / \sigma_{II}^2$ 的区间估计。

$\sigma_I^2 / \sigma_{II}^2$ 的 $100(1-\alpha)\%$ 的置信区间为

$$\left[\frac{S_I^2}{S_{II}^2} \cdot \frac{1}{F_{\alpha/2, (n_I-1, n_{II}-1)}}, \frac{S_I^2}{S_{II}^2} \cdot \frac{1}{F_{1-\alpha/2, (n_I-1, n_{II}-1)}} \right]$$

3.3 假设检验

3.3.1 概述

1. 意义

【例 3.1】 随机各抽取 10 个甲加工件和乙加工件，测得寿命（年）如下。

甲加工件 (y_I): 11、11、9、12、10、13、13、8、10、13。

乙加工件 (y_{II}): 8、11、12、10、9、8、8、9、10、7。

甲加工件的产胶量的均数 $\bar{y}_I = 11$ ，标准差 $S_I = 1.760$ 。

乙加工件的产胶量的均数 $\bar{y}_{II} = 9.2$ ，标准差 $S_{II} = 1.549$ 。

能否仅凭这两个样本均数之差 $\bar{y}_I - \bar{y}_{II} = 1.8$ ，立刻得出甲加工件与乙加工件两种寿命不同的结论呢？

统计学认为：立刻得出的结论是不可靠的。若再分别随机抽测 10 个甲加工件和 10 个乙加工件的寿命，又可得到两个样本资料。因抽样误差的随机性，这两个样本均数就不一定是 11 和 9.2，其均数之差也就未必是 1.8 了。

造成上述差异可能的原因包括两个方面：一是确实由品系造成，即因甲加工件与乙加工件品质不同所致；二是试验误差（或抽样误差）造成。

对两个样本的比较，必须判断样本间差异是抽样误差造成的，还是本质不同所致。如何区分两类性质的差异，怎样通过样本来推断总体，这正是显著性检验要解决的问题。

两个总体间差异的比较可以采用以下两种方法进行。

方法 1：研究总体，即由总体中的所有个体数据计算出总体参数进行比较。这种研究总体的方法是很准确的，但通常是不可能进行的，因为总体往往是无限总体，或者是包含个体很多的有限总体。

方法 2: 研究局部样本, 通过所抽取样本研究其所代表的总体。例如, 设甲加工件寿命的总体均数为 μ_I , 乙加工件寿命的总体均数为 μ_{II} , 试验研究的目的, 就是要给 μ_I 、 μ_{II} 是否相同做出判断。由于总体均数 μ_I 、 μ_{II} 未知, 在进行显著性检验时只能以样本均数 $\bar{y}_I$ 、 $\bar{y}_{II}$ 作为检验对象, 更确切地说, 是以 $\bar{y}_I - \bar{y}_{II}$ 作为检验对象。

为什么可以用样本均数作为检验对象呢? 这是由样本均数所具有的以下三个特征决定的。

(1) 离均差的平方和最小。因样本均数与样本各个观测值最接近, 算数平均数是资料的代表数。

(2) 样本均数是总体均数的无偏估计值, 即 $E(\bar{y}) = \mu$ 。

(3) 根据统计学中心极限定理, 样本均数服从或逼近正态分布。

因此, 以样本均数作为检验对象, 由两个样本均数差异的大小推断样本所属总体均数是否相同是有依据的。

2. 基本思路

(1) 假设检验的根据——小概率原理。

统计学上, 把小概率事件在一次试验中看成实际不可能发生的事件称为“小概率事件实际不可能性原理”, 简称“小概率原理”。

小概率事件虽然不是不可能事件, 但在一次试验中出现的可能性很小, 不出现的可能性很大, 以至于实际上可以看成不可能发生的。

小概率原理是统计学上进行假设检验(显著性检验)的基本依据。

通常, 概率为 0.05、0.01 或 0.001 的事件称为“小概率事件”。

【例 3.2】设箱子中有红球、蓝球共 100 个, 但不知红球、蓝球到底各有多少个。现提出无效假设 H_0 : “箱子中有 99 个蓝球”(或“箱子中不只 1 个红球”作为备择假设 H_A), 暂时设 H_0 正确, 那么从箱子中任抽 1 球, 得红球的概率为 0.01, 是一个小概率事件。进行随机抽球 1 次, 若抽到红球, 则自然会使人对 H_0 的正确性产生怀疑, 从而否定 H_0 而接受 H_A , 也就是说, 箱子中不应该“不只 1 个红球”。

(2) 假设检验的思想方法——概率反证法。

当要判断无效假设 H_0 是否成立时, 先假设无效假设 H_0 成立, 然后根据成立的条件进行推导和运算。若最后导致小概率事件发生, 即出现了不合理的结果, 此时拒绝无效假设 H_0 , 否则接受无效假设 H_0 。

3. 步骤

【第一步】根据研究目的, 对试验样本所在总体先进行无效假设 H_0 (与备择假设 H_A)。

H_0 代表无效假设、原假设、零假设 (null hypothesis), 是被检验的假设, 通过检验可能被接受, 也可能被否定。

H_A 代表备择假设 (alternative hypothesis), 是与 H_0 对应的假设, 是只有在无效假设 H_0 被否定后才可接受的假设, 无充分理由是不能轻率接受的。

【第二步】在无效假设 H_0 成立的前提下, 营造适合的统计量, 并研究试验所得统计量的抽样分布, 计算无效假设正确的概率。

【第三步】据“小概率事件实际不可能性原理”否定或接受无效假设 H_0 , 并得出判断。

当试验表面效应是试验误差的概率小于 0.05、0.01 或 0.001 时, 可以认为在一次试验中试验表面效应是试验误差实际上是不可能的, 因而否定原先所做的无效假设 H_0 , 接受备择假设 H_A , 即认为试验的处理效应是存在的。

当试验表面效应是试验误差的概率大于 0.05、0.01 或 0.001 时, 则说明无效假设成立的可能性大, 不能被否定, 因而也就不能接受备择假设。

4. 两类错误

第一类错误 (I 型错误或 α 型错误): 无效假设应成立, 却否定了它, 犯了“弃真”错误, 即把非真实差异错判为真实差异, 即 $H_0(\mu_1 = \mu_2)$ 为真, 却接受 $H_A(\mu_1 \neq \mu_2)$ 。

第二类错误 (II 型错误或 β 型错误): 无效假设不成立, 却接受了它, 犯了“存伪”错误, 即把真实差异错判为非真实差异, 即 $H_A(\mu_1 \neq \mu_2)$ 为真, 却未能否定 $H_0(\mu_1 = \mu_2)$ 。

第一类错误只有在否定 H_0 时才会发生, 而第二类错误只有在接受 H_0 时才会发生, 两者不会同时发生。

在样本容量相同条件下, 若犯第一类错误的概率减小, 则犯第二类错误的概率会增加; 反之, 若犯第二类错误的概率减小, 则犯第一类错误的概率就会增加。

如将检验水准 α 从 0.05 提高到 0.01, 因分位数 ($u_{\alpha/2}$ 或 $t_{\alpha/2, n-1}$) 提高而导致统计量 (u 或 t) 更难大于该分位数, 就更容易接受 H_0 , 因此犯第一类错误的概率就减小, 相应地增加了犯第二类错误的概率。所以, 检验水准 α 的选取, 也并非越高越好, 应根据实际要求而定。

5. 双侧检验与单侧 (右 / 左) 检验

(1) 双侧检验: 利用两尾概率进行的检验, 如图 3-1 所示。

(2) 单侧检验: 利用一尾概率进行的检验, 如图 3-2 所示。

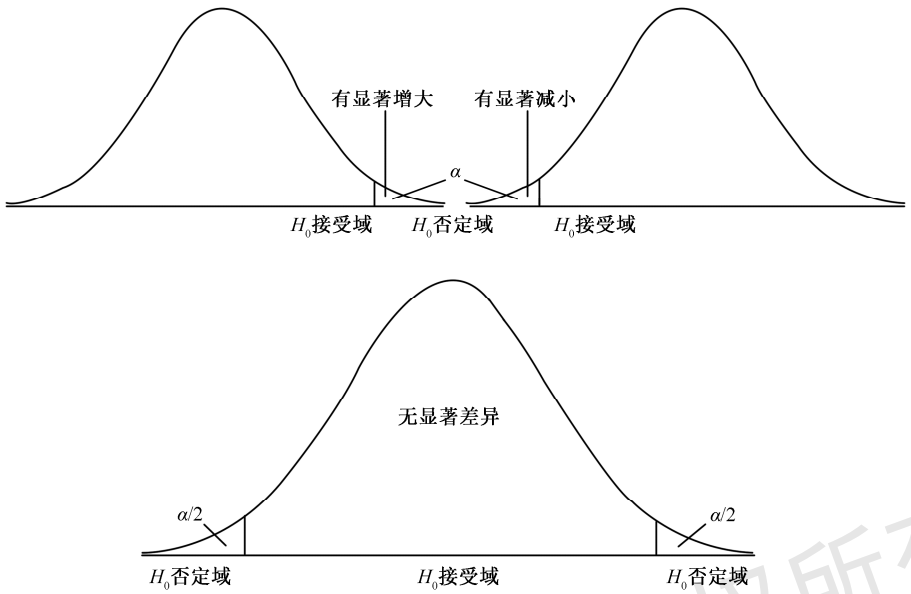


图 3-1 双侧检验

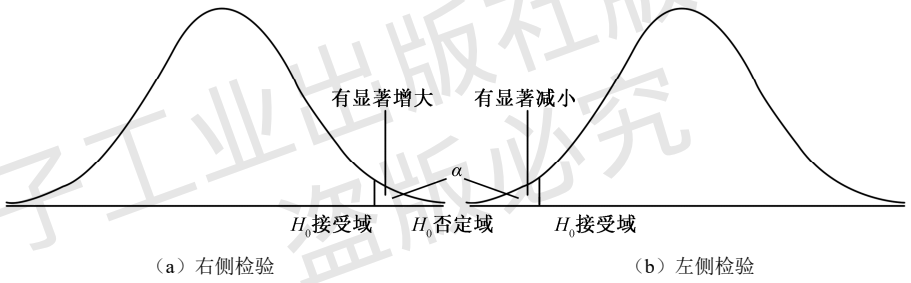


图 3-2 单侧检验

若对同一资料进行单侧检验也进行双侧检验，那么在水平上单侧检验显著，只相当于双侧检验在水平上显著。双侧检验显著，单侧检验一定显著；但单侧检验显著，双侧检验未必显著。所以，同一资料双侧检验与单侧检验所得的结论不一定相同。

选用单侧检验还是双侧检验应根据专业知识及问题的要求在试验设计时就确定。一般若事先不知道所比较的两个处理效果谁好谁坏，分析的目的在于推断两个处理效果有无差别，则选用双侧检验；若根据理论知识或实践经验判断甲处理的效果不会比乙处理的效果差（或好），分析的目的在于推断甲处理是否比乙处理好（或差），则用单侧检验。一般情况下，如没有特殊说明均指双侧检验。

6. 显著性检验应注意的问题

- (1) 要有严密合格的试验或抽样设计，且处理间要有可比性。
- (2) 选用的显著性检验方法应符合其应用条件。

- (3) 要正确理解差异显著或极显著的统计意义。
- (4) 合理建立统计假设, 正确营造出检验统计量。
- (5) 结论不能绝对化。

3.3.2 一个正态总体的假设检验

设一个总体 $Y \sim N(\mu_0, \sigma_0^2)$, μ_0 与 σ_0^2 均为已知, $Y_1, Y_2, \dots, Y_n$ 是 Y 的样本, $y_1, y_2, \dots, y_n$ 是 Y 的样本值。

- (1) 检验无效假设 $H_0: \mu = \mu_0$, 采用双侧检验。

【第一步】假定无效假设 $H_0: \mu = \mu_0$ 成立, 并计算统计量 $u = \frac{\bar{y} - \mu_0}{\sigma_0 / \sqrt{n}}$ 。

【第二步】对给定显著性检验水准 α , 查附表 A-6, 得 $u_{\alpha/2}$ 。

【第三步】若 $|u| > u_{\alpha/2}$, 拒绝 H_0 (接受 $H_A: \mu \neq \mu_0$), 即有 $100(1-\alpha)\%$ 把握认为有 (极) 显著差异性; 若 $|u| \leq u_{\alpha/2}$, 接受 H_0 (拒绝 $H_A: \mu \neq \mu_0$), 即有 $100(1-\alpha)\%$ 把握认为无差异显著性。

- (2) 检验无效假设 $H_0: \mu \leq \mu_0$, 采用右侧检验 (适用于 $u > 0$)。

总的检验过程与 (1) 相似, 不同的是查附表 A-6, 得 u_α , 然后按以下判断规则: 若 $u \leq u_\alpha$, 无显著增大; 若 $u > u_\alpha$, 有显著增大。

- (3) 检验无效假设 $H_0: \mu \geq \mu_0$, 采用左侧检验 (适用于 $u < 0$)。

总的检验过程与 (1) 相似, 不同的是查附表 A-6, 得 u_α , 然后按以下判断规则: 若 $|u| \leq u_\alpha$, 无显著减小; 若 $|u| > u_\alpha$, 有显著减小。

【例 3.3】某地小麦株产 $Y \sim N(\mu_0, \sigma_0^2)$, $\mu_0 = 33.5g$, $\sigma_0 = 1.6g$ 。今从国外引进一高产品种, 在 8 个小区种植, 得株产 (y, g): 35.6、37.6、33.4、35.1、32.7、36.8、35.9、34.6。问新引进品种与该地原有品种在株产上是否有显著差异? 若有显著差异, 是否显著高于该地原有品种? (取 $\alpha = 0.01$)

【解】(1) 新引进品种与该地原有品种在株产上是否有显著差异?

- ① 假定无效假设 $H_0: \mu = \mu_0$ 成立 (而备择假设 $H_A: \mu \neq \mu_0$)。

计算统计量: $u = \frac{\bar{y} - \mu_0}{\sigma_0 / \sqrt{n}} = \frac{35.2 - 33.5}{1.6 / \sqrt{8}} = 3.005$, 其中

$$\bar{y} = \sum_{i=1}^n y_i = (35.6 + 37.6 + \dots + 34.6) / 8 = 35.2$$

② 对给定显著性检验水准 $\alpha = 0.01$, 查附表 A-6 标准正态分布双侧分位数表得 $u_{\alpha/2} = u_{0.01/2} = 2.576$ 。

③ 判断：因 $u = 3.005 > u_{\alpha/2} = 2.576$ ，故拒绝 H_0 ，即有 99% 把握认为新引进品种的株产与该地原有品种有显著差异。

(2) 若有显著差异，是否显著高于该地原有品种？

① 假定 $H_0: \mu \leq \mu_0$ 成立，而 $H_A: \mu > \mu_0$ （因引进新品种的目的 $\mu > \mu_0$ ），计算统计量（与问题一完全相同）。

② 对给定的 $\alpha = 0.01$ ，查附表 A-6 标准正态分布表（右侧）得 $u_\alpha = u_{0.01} = 2.326$ 。

③ 判断：因 $u = 3.005 > u_\alpha = 2.326$ ，故拒绝 H_0 ，即有 99% 把握认为新引进品种的株产显著高于该地原有品种。

【推论】若 $H_A: \mu \neq \mu_0$ 成立，则根据样本表现， $H_A: \mu > \mu_0$ 一定成立，因为检验统计量完全相同，而始终有 $u_{\alpha/2}$ （双侧） $> u_\alpha$ （单侧）。

设一个总体 $Y \sim N(\mu_0, \sigma_0^2)$ ， μ_0 已知而 σ^2 未知， $Y_1, Y_2, \dots, Y_n$ 是 Y 的样本， $y_1, y_2, \dots, y_n$ 是 Y 的样本值。

(1) 检验无效假设 $H_0: \mu = \mu_0$ ，采用双侧检验。

【第一步】假定无效假设 $H_0: \mu = \mu_0$ 成立，并计算统计量 $t = \frac{\bar{y} - \mu_0}{S/\sqrt{n}}$ 。

【第二步】对给定显著性检验水准 α ，查附表 A-4 的 t 分布表，得 $t_{\alpha/2, (n-1)}$ 。

【第三步】判断：若 $|t| > t_{\alpha/2, (n-1)}$ ，则拒绝 H_0 （接受 $H_A: \mu \neq \mu_0$ ），即有 $100(1-\alpha)\%$ 把握认为有（极）显著差异；若 $|t| \leq t_{\alpha/2, (n-1)}$ ，则接受 H_0 ，即有 $100(1-\alpha)\%$ 把握认为无差异显著性。

(2) 检验无效假设 $H_0: \mu \leq \mu_0$ 采用右侧检验（适用于 $t > 0$ ）。

总的检验过程与（1）相似，不同的是查附表 A-4 的 t 分布表，得 $t_{\alpha, (n-1)}$ ，然后按以下判断规则：若 $t \leq t_{\alpha, n-1}$ ，无显著增大；若 $t > t_{\alpha, n-1}$ ，有显著增大。

(3) 检验无效假设 $H_0: \mu \geq \mu_0$ ，采用右侧检验（适用于 $t < 0$ ）。

总的检验过程与（1）相似，不同的是查附表 A-4 的 t 分布表，得 $t_{\alpha, (n-1)}$ ，然后按以下判断规则：若 $|t| \leq t_{\alpha, n-1}$ ，无显著减小；若 $|t| > t_{\alpha, n-1}$ ，有显著减小。

【例 3.4】某地小麦株产 $Y \sim N(\mu_0, \sigma^2)$ ， $\mu_0 = 33.5g$ ， σ 未知。现从国外引进一高产品种，在 8 个小区种植，得株产 (y, g): 35.6、37.6、33.4、35.1、32.7、36.8、35.9、34.6。问新引进品种与该地原有品种在株产上是否有显著差异？若有显著差异，是否显著高于该地原有品种？（取 $\alpha = 0.05$ ）

【解】(1) 新引进品种与该地原有品种在株产上是否有显著差异？

① 假定无效假设 $H_0: \mu = \mu_0$ 成立（而备择假设 $H_A: \mu \neq \mu_0$ ）。

计算统计量: $t = \frac{\bar{y} - \mu_0}{S/\sqrt{n}} = \frac{35.2 - 33.5}{1.64/\sqrt{8}} = 2.932$ 。其中

$$\begin{aligned}\bar{y} &= \sum_{i=1}^n y_i / n = (35.6 + 37.6 + \cdots + 34.6) / 8 = 35.2 \\ S &= \sqrt{\frac{\sum_{i=1}^n (y_i - \bar{y})^2}{n-1}} \\ &= \sqrt{\frac{\sum_{i=1}^n y_i^2 - \left(\sum_{i=1}^n y_i\right)^2 / n}{n-1}} \\ &= \sqrt{\frac{(35.6^2 + 37.6^2 + \cdots + 34.6^2) - (35.6 + 37.6 + \cdots + 34.6)^2 / 8}{8-1}} \\ &= 1.64\end{aligned}$$

② 对给定 $\alpha = 0.05$, 查 t 分布表 (双侧) 得: $t_{\alpha/2, (n-1)} = t_{0.05/2, 7} = 2.365$ 。

③ 判断: 因 $t = 2.932 > t_{\alpha/2, (n-1)} = 2.365$, 拒绝 H_0 , 即有 95% 把握认为新引进品种与该地原有品种在株产上有显著差异。

(2) 若有显著差异, 其株产是否显著高于该地原有品种?

① 假定 $H_0: \mu \leq \mu_0$ 成立, 而 $H_A: \mu > \mu_0$ (因引进新品种的目的 $\mu > \mu_0$), 计算统计量 (与问题一完全相同)。

② 对给定的 $\alpha = 0.05$, 查 t 分布表 (右侧) 得: $t_{\alpha, (n-1)} = t_{0.05, 7} = 1.895$ 。

③ 判断: 因 $t = 2.932 > t_{\alpha, (n-1)} = 1.895$, 故拒绝 H_0 , 即有 95% 的把握认为新引进品种的株产显著高于该地原有品种。

设一个总体 $Y \sim N(\mu, \sigma_0^2)$, σ_0^2 已知而 μ 未知, $Y_1, Y_2, \dots, Y_n$ 是 Y 的样本, $y_1, y_2, \dots, y_n$ 是 Y 的样本值。

(1) 检验无效假设 $H_0: \sigma^2 = \sigma_0^2$, 采用双侧检验。

【第一步】假定无效假设 $H_0: \sigma^2 = \sigma_0^2$ 成立, 并计算统计量 $X^2 = \frac{(n-1)S^2}{\sigma_0^2}$ 。

【第二步】对给定显著性检验水准 α , 查附表 A-8, 得 $X_{1-\alpha/2, (n-1)}^2$ 及 $X_{\alpha/2, (n-1)}^2$ 。

【第三步】判断: 若 $X^2 < X_{1-\alpha/2, (n-1)}^2$ 或 $X^2 > X_{\alpha/2, (n-1)}^2$, 则拒绝 H_0 , 即有 $100(1-\alpha)\%$ 把握认为总体方差并不是所给出的方差; 若 $X_{1-\alpha/2, (n-1)}^2 \leq X^2 \leq X_{\alpha/2, (n-1)}^2$, 则接受 H_0 , 即有 $100(1-\alpha)\%$ 把握认为总体方差正好是所给出的方差。

(2) 检验无效假设 $H_0: \sigma^2 \leq \sigma_0^2$, 采用右侧检验。

总的检验过程与(1)相似,不同的是查附表 A-8,得 $X_{\alpha,(n-1)}^2$, 然后按以下判断规则: 若 $X^2 \leq X_{\alpha,(n-1)}^2$, 无显著增大; 若 $X^2 > X_{\alpha,(n-1)}^2$, 有显著增大。

(3) 检验无效假设 $H_0: \sigma^2 \geq \sigma_0^2$, 采用左侧检验。

总的检验过程与(1)相似,不同的是查附表 A-8,得 $X_{1-\alpha,(n-1)}^2$, 然后按以下判断规则: 若 $X^2 \leq X_{1-\alpha,(n-1)}^2$, 无显著减小; 若 $X^2 > X_{1-\alpha,(n-1)}^2$, 有显著减小。

【例 3.5】 正常情况下, 某厂生产的尼龙纤度 $Y \sim N(\mu, 0.048^2)$ 。某日随机抽取 5 根纤维, 测得其纤度(y)分别为 1.32、1.36、1.55、1.44、1.40。试问: 能否认为这一天尼龙纤度的标准差 $\sigma = 0.048$? (取 $\alpha = 0.1$)

【解】 (1) 假定无效假设 $H_0: \sigma^2 = 0.048^2$ 成立。

计算统计量 $X^2 = \frac{(n-1)S^2}{\sigma_0^2} = \frac{(5-1) \times 0.00775}{0.048^2} = 13.51$ 。其中

$$\bar{y} = \sum_{i=1}^n y_i / n = (1.32 + 1.36 + \cdots + 1.40) / 5 = 1.414$$

$$\begin{aligned} S^2 &= \sum_{i=1}^n (y_i - \bar{y})^2 / (n-1) \\ &= [(1.32 - 1.414)^2 + (1.36 - 1.414)^2 + \cdots + (1.40 - 1.414)^2] / (5-1) \\ &= 0.00775 \end{aligned}$$

(2) 对给定的 $\alpha = 0.1$, 查附表 A-8, 得

$$X_{1-\alpha/2,(n-1)}^2 = X_{1-0.1/2,(5-1)}^2 = X_{0.95,4}^2 = 0.711$$

$$X_{\alpha/2,(n-1)}^2 = X_{0.1/2,(5-1)}^2 = X_{0.05,4}^2 = 9.488$$

(3) 判断: 因 $X^2 > X_{\alpha/2,(n-1)}^2$, 故拒绝 H_0 , 即有 90% 的把握认为这一天尼龙纤度的标准差 σ 并不等于 0.048。

3.3.3 两个正态总体的假设检验

设某一总体 $Y_I \sim N(\mu_I, \sigma_I^2)$, $Y_{I_1}, Y_{I_2}, \cdots, Y_{I_{m_1}}$ 是 Y_I 的样本, $y_{I_1}, y_{I_2}, \cdots, y_{I_{m_1}}$ 是 Y_I 的样本值; 另一总体 $Y_{II} \sim N(\mu_{II}, \sigma_{II}^2)$, $Y_{II_1}, Y_{II_2}, \cdots, Y_{II_{m_2}}$ 是 Y_{II} 的样本, $y_{II_1}, y_{II_2}, \cdots, y_{II_{m_2}}$ 是 Y_{II} 的样本值。已知 σ_I^2 、 σ_{II}^2 、 Y_I 与 Y_{II} 相互独立。

(1) 检验无效假设 $H_0: \mu_I = \mu_{II}$, 采用双侧检验。

【第一步】假定无效假设 $H_0: \mu_I = \mu_{II}$ 成立，计算统计量

$$U = \frac{(\bar{y}_I - \bar{y}_{II}) - (\mu_I - \mu_{II})}{\sqrt{\frac{\sigma_I^2}{n_I} + \frac{\sigma_{II}^2}{n_{II}}}} = \frac{(\bar{y}_I - \bar{y}_{II})}{\sqrt{\frac{\sigma_I^2}{n_I} + \frac{\sigma_{II}^2}{n_{II}}}}$$

【第二步】对给定的显著性检验水准 α ，查附表 A-6，得 $u_{\alpha/2}$ 。

【第三步】判断：若 $|U| > u_{\alpha/2}$ ，拒绝 H_0 ，即有 $100(1-\alpha)\%$ 把握认为两总体之间有（极）显著差异；若 $|U| \leq u_{\alpha/2}$ ，接受 H_0 ，即有 $100(1-\alpha)\%$ 把握认为两总体之间无显著差异。

(2) 检验无效假设 $H_0: \mu_I \leq \mu_{II}$ ，采用右侧检验（适用于 $U > 0$ ）。

总的检验过程与（1）相似，不同的是查附表 A-6，得 u_α ，然后按以下判断规则：若 $U \leq u_\alpha$ ，无显著增大；若 $U > u_\alpha$ ，有显著增大。

(3) 检验无效假设 $H_0: \mu_I \geq \mu_{II}$ ，采用左侧检验（适用于 $U < 0$ ）。

总的检验过程与（1）相似，不同的是查附表 A-6，得 u_α ，然后按以下判断规则：若 $|U| \leq u_\alpha$ ，无显著增大；若 $|U| > u_\alpha$ ，有显著增大。

【例 3.6】甲、乙两台车床加工同一种工件，现要测量工件的同轴度，设甲加工工件同轴度 $Y_I \sim N(\mu_I, 0.025^2)$ ，乙加工工件同轴度 $Y_{II} \sim N(\mu_{II}, 0.06^2)$ 。今从甲、乙两台车床加工件中分别测量 $n_I = 200$ 个， $n_{II} = 150$ 个，并求得 $\bar{y}_I = 0.081$ ， $\bar{y}_{II} = 0.060$ 。试问：这两台车床加工工件的同轴度是否有显著差异？（取 $\alpha = 0.05$ ）

【解】(1) 假定无效假设 $H_0: \mu_I = \mu_{II}$ 成立。

$$\text{计算统计量 } U = \frac{(\bar{y}_I - \bar{y}_{II})}{\sqrt{\frac{\sigma_I^2}{n_I} + \frac{\sigma_{II}^2}{n_{II}}}} = \frac{0.081 - 0.060}{\sqrt{\frac{0.025^2}{200} + \frac{0.06^2}{150}}} = 3.92。$$

(2) 对给定的 $\alpha = 0.05$ ，查附表 A-6，得 $u_{\alpha/2} = u_{0.05/2} = 1.96$ 。

(3) 判断：因 $U = 3.92 > u_{\alpha/2} = 1.96$ ，故拒绝 H_0 ，即有 95% 把握认为这两台车床加工工件的同轴度有显著差异。

设某一总体 $Y_I \sim N(\mu_I, \sigma_I^2)$ ， $Y_{I_1}, Y_{I_2}, \dots, Y_{I_{n_I}}$ 是 Y_I 的样本， $y_{I_1}, y_{I_2}, \dots, y_{I_{n_I}}$ 是 Y_I 的样本值；另一总体 $Y_{II} \sim N(\mu_{II}, \sigma_{II}^2)$ ， $Y_{II_1}, Y_{II_2}, \dots, Y_{II_{n_{II}}}$ 是 Y_{II} 的样本， $y_{II_1}, y_{II_2}, \dots, y_{II_{n_{II}}}$ 是 Y_{II} 的样本值。 σ_I^2 、 σ_{II}^2 未知，但 $\sigma_I^2 = \sigma_{II}^2$ ， Y_I 与 Y_{II} 相互独立。

【情形一】成组法——非配对设计两样本均数的差异显著性检验。

表 3-1 为非配对设计资料的一般形式。

表 3-1 非配对设计资料的一般形式

处 理	观测值 y_{ij}	样本含量	样本均数	总体均数
I	$y_{11}, y_{12}, \dots, y_{1n_1}$	n_1	$\bar{y}_I = \sum y_{1j} / n_1$	μ_I
II	$y_{II1}, y_{II2}, \dots, y_{II n_2}$	n_2	$\bar{y}_{II} = \sum y_{II j} / n_2$	μ_{II}

(1) 检验无效假设 $H_0: \mu_I = \mu_{II}$ ，采用双侧检验。

【第一步】假定无效假设 $H_0: \mu_I = \mu_{II}$ 成立，计算统计量

$$T = \frac{(\bar{y}_I - \bar{y}_{II})}{\sqrt{\frac{(n_1 - 1)S_I^2 + (n_2 - 1)S_{II}^2}{n_1 + n_2 - 1} \left(\frac{1}{n_1} + \frac{1}{n_2} \right)}}$$

简记为 $T = \frac{(\bar{y}_I - \bar{y}_{II})}{S_{\bar{y}_I - \bar{y}_{II}}}$ 。其中， $S_I^2 = \frac{1}{n_1 - 1} \sum_{i=1}^{n_1} (y_{1i} - \bar{y}_I)^2$ ， $S_{II}^2 = \frac{1}{n_2 - 1} \sum_{i=1}^{n_2} (y_{II i} - \bar{y}_{II})^2$ 。

【第二步】对给定显著性检验水准 α ，查附表 A-7，得 $t_{\alpha/2, (n_1 + n_2 - 2)}$ 。

【第三步】判断：若 $|T| > t_{\alpha/2, (n_1 + n_2 - 2)}$ ，拒绝 H_0 ，即有 $100(1 - \alpha)\%$ 把握认为有(极)显著差异；若 $|T| \leq t_{\alpha/2, (n_1 + n_2 - 2)}$ ，接受 H_0 ，即有 $100(1 - \alpha)\%$ 把握认为无显著差异。

(2) 检验无效假设 $H_0: \mu_I \leq \mu_{II}$ ，采用右侧检验（适用于 $T > 0$ ）。

总的检验过程与(1)相似，不同的是查附表 A-7 t 分布单侧分位数表，得 $t_{\alpha, (n_1 + n_2 - 2)}$ ，然后按以下判断规则：若 $T \leq t_{\alpha, (n_1 + n_2 - 2)}$ ，无显著增大；若 $T > t_{\alpha, (n_1 + n_2 - 2)}$ ，有显著增大。

(3) 检验无效假设 $H_0: \mu_I \geq \mu_{II}$ ，采用左侧检验（适用于 $T < 0$ ）。

总的检验过程与(1)相似，不同的是查附表 A-7 t 分布单侧分位数表，得 $t_{\alpha, (n_1 + n_2 - 2)}$ ，然后按以下判断规则：若 $|T| \leq t_{\alpha, (n_1 + n_2 - 2)}$ ，无显著减小；若 $|T| > t_{\alpha, (n_1 + n_2 - 2)}$ ，有显著减小。

【例 3.7】在某厂生产工件的工艺过程中，要考察温度对工件强度的影响，为了比较 70°C 与 80°C 的影响有无显著性差异，在这两种温度下分别做 8 次重复试验，所得工件强度如下。

70°C 时的工件强度 (y_I): 20.5、18.8、19.8、20.9、21.5、19.5、21.0、21.2。

80°C 时的工件强度 (y_{II}): 17.7、20.3、20.0、18.8、19.0、20.1、20.2、19.1。

据以往经验，可认为工件强度服从正态分布，且在不同温度下的方差相同，取 $\alpha = 0.05$ 。

【解】(1) 假定无效假设 $H_0: \mu_I = \mu_{II}$ 成立。

计算统计量 $T = \frac{(\bar{y}_I - \bar{y}_{II})}{S_{\bar{y}_I - \bar{y}_{II}}} = \frac{20.4 - 19.4}{0.4629} = 2.160$ 。其中

$$\bar{y}_I = \sum_{i=1}^{n_1} y_{I_i} / n_1 = (20.5 + 18.8 + \cdots + 21.2) / 8 = 20.4$$

$$\bar{y}_{II} = \sum_{i=1}^{n_2} y_{II_i} / n_2 = (17.7 + 20.3 + \cdots + 19.1) / 8 = 19.4$$

$$S_I^2 = \frac{1}{n_1 - 1} \sum_{i=1}^{n_1} (y_{I_i} - \bar{y}_I)^2$$

$$= \frac{1}{8 - 1} [(20.5 - 20.4)^2 + (18.8 - 20.4)^2 + \cdots + (21.2 - 20.4)^2] = 0.8857$$

$$S_{II}^2 = \frac{1}{n_2 - 1} \sum_{i=1}^{n_2} (y_{II_i} - \bar{y}_{II})^2$$

$$= \frac{1}{8 - 1} [(17.7 - 19.4)^2 + (20.3 - 19.4)^2 + \cdots + (19.1 - 19.4)^2] = 0.8286$$

$$S_{\bar{y}_I - \bar{y}_{II}} = \sqrt{\frac{(n_1 - 1)S_I^2 + (n_2 - 1)S_{II}^2}{n_1 + n_2 - 2} \left(\frac{1}{n_1} + \frac{1}{n_2} \right)}$$

$$= \sqrt{\frac{(8 - 1) \times 0.8857 + (8 - 1) \times 0.8286}{8 + 8 - 2} \left(\frac{1}{8} + \frac{1}{8} \right)} = \sqrt{0.2143} = 0.4629$$

(2) 对给定的 $\alpha = 0.05$, 查附表 A-7 t 分布双侧分位数表, 得 $t_{0.05/2, (8+8-2)} = 2.145$ 。

(3) 判断: 因 $T = 2.16 > t_{0.05/2, (8+8-2)} = 2.145$, 故拒绝 H_0 , 即有 95% 的把握认为 70°C 与 80°C 下工件强度有显著差异。

【情形二】成对法——配对设计两样本均数得差异显著性检验。

表 3-2 为配对设计试验资料的一般形式。

表 3-2 配对设计试验资料的一般形式

处 理	观测值 y_{ij}				样本含量	样本平均数	总体均数
I	y_{I1}	y_{I2}	...	y_{In}	n	$\bar{y}_I = \sum y_{Ij} / n$	μ_I
II	y_{II1}	y_{II2}	...	y_{II_n}	n	$\bar{y}_{II} = \sum y_{II_j} / n$	μ_{II}
$d_j = y_{Ij} - y_{II_j}$	d_1	d_2	...	d_n	n	$\bar{d} = \bar{y}_I - \bar{y}_{II}$	$\mu_d = \mu_I - \mu_{II}$

(1) 检验无效假设 $H_0: \mu_I = \mu_{II}$, 采用双侧检验。

【第一步】假定无效假设 $H_0: \mu_d = 0$ ($\mu_I = \mu_{II}$) 成立, 计算统计量 $t = \frac{\bar{d}}{S_{\bar{d}}}$, 其中

$$\bar{d} = \sum_{j=1}^n d_j / n, d_j = y_{I_j} - y_{II_j} (j=1, 2, \cdots, n)$$

$$S_{\bar{d}} = \frac{S_d}{\sqrt{n}}$$

$$= \sqrt{\frac{\sum_{j=1}^n (d_j - \bar{d})^2}{n(n-1)}} = \sqrt{\frac{\sum_{j=1}^n d_j^2 - \left(\sum_{j=1}^n d_j\right)^2 / n}{n(n-1)}}$$

【第二步】对给定显著性检验水准 α ，查附表 A-7 t 分布双侧分位数表，得 $t_{\alpha/2, (n-1)}$ 。

【第三步】判断：若 $|t| \leq t_{\alpha/2, (n-1)}$ ，接受 $H_0: \mu_I = \mu_{II}$ ，表明有 $100(1-\alpha)\%$ 把握认为两处理均数无显著差异；若 $|t| > t_{\alpha/2, (n-1)}$ ，拒绝 $H_0: \mu_I = \mu_{II}$ ，表明有 $100(1-\alpha)\%$ 把握认为两处理均数有（极）显著差异。

(2) 检验无效假设 $H_0: \mu_I \leq \mu_{II}$ ，采用右侧检验（适用于 $\bar{d} > 0$ ）。

总的检验过程与 (1) 相似，不同的是查附表 A-7 t 分布单侧分位数表，得 $t_{\alpha, (n-1)}$ ，然后按以下判断规则：若 $t < t_{\alpha, (n-1)}$ ，无显著增大；若 $t > t_{\alpha, (n-1)}$ ，有显著增大。

(3) 检验无效假设 $H_0: \mu_I \geq \mu_{II}$ ，采用左侧检验（适用于 $\bar{d} < 0$ ）。

总的检验过程与 (1) 相似，不同的是查附表 A-7 t 分布单侧分位数表，得 $t_{\alpha, (n-1)}$ ，然后按以下判断规则：若 $|t| < t_{\alpha, (n-1)}$ ，无显著减小；若 $|t| > t_{\alpha, (n-1)}$ ，有显著减小。

【例 3.8】为测定甲、乙两种病毒对橡胶树的致病力，取 8 株橡胶树，每株皆半叶接种甲病毒，另半叶接种乙病毒，分别以叶面出现枯斑数多少作为致病力强弱的指标，试验结果如表 3-3 所示。试检验两种病毒致病力的差异显著性，取 $\alpha = 0.05$ 。

表 3-3 试验结果

株号 (j)	1	2	3	4	5	6	7	8
y_{Ij} (甲病毒)	9	17	31	18	7	8	20	10
y_{IIj} (乙病毒)	10	11	18	14	6	7	17	5
$d_j = y_{Ij} - y_{IIj}$	-1	6	13	4	1	1	3	5

【解】(1) 进行无效假设 $H_0: \mu_d = 0$ （或备择假设 $H_A: \mu_d \neq 0$ ）成立。

计算统计量： $t = \frac{\bar{d}}{S_{\bar{d}}} = \frac{\bar{d}}{S_d / \sqrt{n}} = \frac{4}{4.31 / \sqrt{8}} = 2.625$ ，其中

$$\bar{d} = \sum d_j / n = (-1 + 6 + \cdots + 5) / 8 = 32 / 8 = 4$$

$$S_d = \sqrt{\frac{\sum d_j^2 - (\sum d_j)^2 / n}{n-1}} = \sqrt{\frac{[(-1)^2 + 6^2 + \cdots + 5^2] - 32^2 / 8}{8-1}} = 4.31$$

(2) 对给定的 $\alpha = 0.05$ ，查附表 A-7 t 分布双侧分位数表，得 $t_{0.05/2, (8-1)} = t_{0.05/2, 7} = 2.365$ 。

(3) 判断: 因 $t = 2.625 > t_{0.05/2, (8-1)} = 2.365$, 故拒绝 H_0 , 即有 95% 的把握认为甲乙两种病毒致病力有显著差异。

其实, 在双侧假设 H_0 被否定后, 根据单侧假设两者必取其一的原则, 因 $d > 0$, 故可以判断甲病毒的致病力比乙病毒强。

【总结 1】成组法与成对法的比较。

(1) 在进行试验设计时, 如将两处理完全随机排列, 而处理间 (组间) 的各供试单位彼此独立, 无论两个样本的容量是否相同, 所获数据均为成组数据。

(2) 在科学研究中, 为减小试验误差, 提高试验结果的精确度, 在比较两个样本 (或处理效应) 差异是否显著时, 经常采用对比设计, 即将性质相同的两个供试单位配成一对, 并设有多个配对, 然后对每个配对的两个供试单位分别随机地给予不同处理, 以这种设计方法获得的数据称为成对数据。

(3) 与成组资料相比较, 成对比较的优点: ① 成对法因加强试验条件的控制, 每对观察值所处的试验条件一致性较强, 所得观察值的可比性提高, 故随机误差减小, 可发现较小的真实差异; ② 成对法不受两个样本的总体方差 $\sigma_1^2 \neq \sigma_{II}^2$ 的干扰, 分析时无须考虑 σ_1^2 和 σ_{II}^2 是否相等。

需要指出的是, 成组资料的样本容量即使相等, 即 $n_1 = n_2$, 也不能采用成对资料的分析方法, 因为成组资料的每个观察值都是独立的, 没有配对的基础。

设某一总体 $Y_I \sim N(\mu_I, \sigma_I^2)$, $Y_{I_1}, Y_{I_2}, \dots, Y_{I_{n_1}}$ 是 Y_I 的样本, $y_{I_1}, y_{I_2}, \dots, y_{I_{n_1}}$ 是 Y_I 的样本值; 另一总体 $Y_{II} \sim N(\mu_{II}, \sigma_{II}^2)$, $Y_{II_1}, Y_{II_2}, \dots, Y_{II_{n_2}}$ 是 Y_{II} 的样本, $y_{II_1}, y_{II_2}, \dots, y_{II_{n_2}}$ 是 Y_{II} 的样本值。 μ_I 、 μ_{II} 未知, 且 Y_I 与 Y_{II} 相互独立。

(1) 检验无效假设 $H_0: \sigma_I^2 = \sigma_{II}^2$, 采用双侧检验。

【第一步】假定无效假设 $H_0: \sigma_I^2 = \sigma_{II}^2$ 成立, 计算统计量 $F = \frac{S_I^2 / \sigma_I^2}{S_{II}^2 / \sigma_{II}^2} = \frac{S_I^2}{S_{II}^2}$ 。

【第二步】对给定显著性检验水准 α , 查附表 A-9 得

$$F_{\alpha/2, (n_1-1, n_2-1)}, F_{1-\alpha/2, (n_1-1, n_2-1)} \left(= \frac{1}{F_{\alpha/2, (n_2-1, n_1-1)}} \right)$$

【第三步】判断: 若 $F > F_{\alpha/2, (n_1-1, n_2-1)}$ 或 $F < F_{1-\alpha/2, (n_1-1, n_2-1)}$, 拒绝 H_0 , 即有 $100(1-\alpha)\%$ 的把握认为两总体方差有 (极) 显著差异; 若 $F_{1-\alpha/2, (n_1-1, n_2-1)} \leq F \leq F_{\alpha/2, (n_1-1, n_2-1)}$, 接受 H_0 , 即有 $100(1-\alpha)\%$ 的把握认为两总体方差没有 (极) 显著差异。

(2) 检验无效假设 $H_0: \sigma_I^2 \leq \sigma_{II}^2$, 采用右侧检验 (适用于 $F > 1$)。

总的检验过程与 (1) 相似, 不同的是查附表 A-9 得 $F_{\alpha, (n_1-1, n_2-1)}$, 然后按以下判断规则: 若 $F \leq F_{\alpha, (n_1-1, n_2-1)}$, 无显著增大; 若 $F > F_{\alpha, (n_1-1, n_2-1)}$, 有显著增大。

(3) 检验无效假设 $H_0: \sigma_I^2 \geq \sigma_{II}^2$, 采用左侧检验 (适用于 $F < 1$)。

总的检验过程与 (1) 相似, 不同的是查附表 A-9 得 $F_{1-\alpha, (n_1-1, n_2-1)}$, 然后按以下判断规则: 若 $F \geq F_{1-\alpha, (n_1-1, n_2-1)}$, 无显著减小; 若 $F < F_{1-\alpha, (n_1-1, n_2-1)}$, 有显著减小。

【例 3.9】用原子吸收光谱法 (新法) 和 EDTA (旧法) 测定某废水中 Al^{3+} 的含量 (%), 试验结果如下。

新法 (y_I): 0.163、0.175、0.159、0.168、0.169、0.161、0.166、0.179、0.174、0.173。

旧法 (y_{II}): 0.153、0.181、0.165、0.155、0.156、0.161、0.176、0.174、0.164、0.183、0.179。

问: 两种方法的精密度是否有显著差异? 新法是否比旧法精密度有显著提高?

【解】(1) 两种方法的精密度显著差异, 采用 F 双侧检验。

① 进行无效检验 $H_0: \sigma_I^2 = \sigma_{II}^2$ (备择假设 $H_A: \sigma_I^2 \neq \sigma_{II}^2$) 成立。

根据给出的相关数据, 可分别算出新法与旧法的样本方差

$$S_I^2 = 4.29 \times 10^{-5}, S_{II}^2 = 1.23 \times 10^{-4}$$

计算统计量: $F = S_I^2 / S_{II}^2 = 4.29 \times 10^{-5} / 1.23 \times 10^{-4} = 0.350$ 。

② 查附表 A-9 及进行相关计算, 得 $F_{0.025(9,10)} = 3.779$, $F_{0.975(9,10)} = 0.252$ 。

③ 判断: 因 $F_{0.975(9,10)} < F < F_{0.025(9,10)}$, 接受 H_0 , 故有 95% 把握认为两种方法的方差无显著差异, 即两种方法的精密度是一致的。

(2) 新法是否比旧法的精密度有显著提高, 即只要检验新法比旧法的方差有显著减小就可以, 采用 F 左侧检验 (因 $F < 1$)。

查附表 A-9 并进行相关计算, 得 $F_{0.95(9,10)} = 0.319$ 。

因 $F_{0.95(9,10)} = F$, 故新法比旧法的方差没有显著减小, 即新法比旧法的精密度无显著提高。

【总结 2】区间估计与假设检验的比较。

(1) 主要区别: ① 参数估计以样本资料估计总体参数的真值, 假设检验以样本资料检验对总体参数的先验假设是否成立; ② 区间估计求得的是以样本估计值为中心的双侧置信区间, 假设检验既有双侧检验, 又有单侧检验; ③ 区间估计立足于大概率, 假设检验立足于小概率。

(2) 主要联系: ① 都根据样本信息推断总体参数; ② 都以抽样分布为理论依据, 都是建立在概率论基础之上的推断; ③ 二者可相互转换, 形成对偶性。

总之, 区间估计与假设检验两者表示结果的形式不同, 但本质是一样的。

3.4 方差分析

3.4.1 概述

1. t 检验法的适应性及局限性

检验法适用于样本均数与总体均数及两样本均数之间的差异显著性检验，但在科学研究和生产实践中，经常会遇到比较多个处理优劣的问题，即需要进行多个均数之间的差异显著性检验。这时，若仍采用检验法就不适宜了，主要存在以下问题。

- (1) 检验过程烦琐；
- (2) 无统一的试验误差，误差估计的精确性和检验的灵敏性低；
- (3) 推断的可靠性低，检验的 I 型（“弃真”）错误率大。

2. 方差分析法的引入

方差分析（analysis of variance）是英国统计学家 R. A. Fisher 在 1923 年提出的，其主要内容是将 n 个处理的观测值作为一个整体看待，把观测值总变异的偏差平方和及其自由度分解为相应于不同变异来源的偏差平方和及其自由度，进而获得不同变异来源总体方差估计值；通过计算这些总体方差估计值的适当比值，采用检验法，检验各样本所属总体均数是否相等而做出显著性判断。

方差分析实质上是关于观测值变异原因的数量分析，它在科学试验研究和实际生产实践中得到了十分广泛的应用，特别适合于三个或三个以上处理的正态总体的有关参数估计和均值比较。

3. 完全随机设计

完全随机含有两层意思：一是试验处理中试验顺序的随机安排；二是试验材料的随机分组。

3.4.2 单因子试验设计及其方差分析

1. 完全随机设计

1) 线性模型与基本假定

设单一因子 A ，取 a 种水平（有 a 个处理），每个处理均重复观察 k 次，具体试验

设计安排、观察结果及相关数据(包括观察值和及其均数)计算的模式如表 3-4 所示。

表 3-4 试验相关数据

因子 A (处理)	重复观测值 (y_{ij})						y_{i*}	y_{**}	$\bar{y}_{i*}$	$\bar{y}_{**}$
	1	2	...	j	...	k				
A_1	y_{11}	y_{12}	...	y_{1j}	...	y_{1k}	y_{1*}		$\bar{y}_{1*}$	
A_2	y_{21}	y_{22}	...	y_{2j}	...	y_{2k}	y_{2*}		$\bar{y}_{2*}$	
$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$		$\vdots$	
A_i	y_{i1}	y_{i2}	...	y_{ij}	...	y_{ik}	y_{i*}		$\bar{y}_{i*}$	
$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$		$\vdots$	
A_a	y_{a1}	y_{a2}	...	y_{aj}	...	y_{ak}	y_{a*}		$\bar{y}_{a*}$	

y_{ij} 表示第 i 个处理的第 j 个观测值 ($i=1,2,\dots,a; j=1,2,\dots,k$): $y_{i*} = \sum_{j=1}^k y_{ij}$ 表示第 i 个处理 k 个观测值之和; $\bar{y}_{i*} = \sum_{j=1}^k y_{ij} / a$ 表示第 i 个处理的平均数; $y_{**} = \sum_{i=1}^a \sum_{j=1}^k y_{ij} = \sum_{i=1}^a y_{i*}$ 表示全部观测值的总和; $\bar{y}_{**} = \sum_{i=1}^a \sum_{j=1}^k y_{ij} / (a \times k)$ 表示全部观测值的总平均数。

在单因子完全随机重复试验中, 第 i 个处理的第 j 个观测值可表示为

$$y_{ij} = \mu + \alpha_i + \varepsilon_{ij}$$

式中, μ 为全部试验观测值总体的平均数, $\mu = \sum_{i=1}^a \mu_i / a$; α_i 为第 i 个的处理效应, 表示处理 i 对试验结果产生的影响, $\sum_{i=1}^a \alpha_i = 0$; ε_{ij} 是试验误差, 相互独立, 且服从正态分布 $N(0, \sigma^2)$ 。

2) 总偏差平方和与总自由度及其分解

(1) 总偏差平方和及其分解。

用于反映全部观测值总变异的总偏差平方和是各观测值 y_{ij} 与总均数 $\bar{y}_{**}$ 的离均差平方和, 即

$$SS_T = \sum_{i=1}^a \sum_{j=1}^k (y_{ij} - \bar{y}_{**})^2$$

进一步展开, 有

$$\begin{aligned} SS_T &= \sum_{i=1}^a \sum_{j=1}^k (y_{ij} - \bar{y}_{**})^2 = \sum_{i=1}^a \sum_{j=1}^k [(\bar{y}_{i*} - \bar{y}_{**}) + (y_{ij} - \bar{y}_{i*})]^2 \\ &= \sum_{i=1}^a \sum_{j=1}^k [(\bar{y}_{i*} - \bar{y}_{**})^2 + 2 \times (\bar{y}_{i*} - \bar{y}_{**}) \times (y_{ij} - \bar{y}_{i*}) + (y_{ij} - \bar{y}_{i*})^2] \end{aligned}$$

$$\begin{aligned}
&= k \sum_{i=1}^a (\bar{y}_{i*} - \bar{y}_{**})^2 + 2 \times \sum_{i=1}^a \left[(\bar{y}_{i*} - \bar{y}_{**}) \times \sum_{j=1}^k (y_{ij} - \bar{y}_{i*}) \right] + \sum_{i=1}^a \sum_{j=1}^k (y_{ij} - \bar{y}_{i*})^2 \\
&= k \sum_{i=1}^a (\bar{y}_{i*} - \bar{y}_{**})^2 + \sum_{i=1}^a \sum_{j=1}^k (y_{ij} - \bar{y}_{i*})^2
\end{aligned}$$

式中, $k \sum_{i=1}^a (\bar{y}_{i*} - \bar{y}_{**})^2$ 为各处理均数 $\bar{y}_{i*}$ 与总均数 $\bar{y}_{**}$ 的离均差平方和与重复数 k 的乘

积, 反映重复 k 次的处理间变异, 称为“处理间偏差平方和” (SS_t): $\sum_{i=1}^a \sum_{j=1}^k (y_{ij} - \bar{y}_{i*})^2$

为各处理内偏差平方和之和, 反映各处理内的差异, 即误差, 称为“处理内偏差平方和”或“误差平方和” SS_e , 即

$$SS_T = SS_t + SS_e$$

因此, 上述公式可简化表示为以下形式。

总偏差平方和: $SS_T = \sum_{i=1}^a \sum_{j=1}^k y_{ij}^2 - CT$, 其中 $CT = \left(\sum_{i=1}^a \sum_{j=1}^k y_{ij} \right)^2 / N = y_{**}^2 / (a \times k)$ 称为

“矫正数”。

处理间偏差平方和: $SS_t = \sum_{i=1}^a y_{i*}^2 / k - CT$

处理内偏差 (误差) 平方和: $SS_e = SS_T - SS_t$

(2) 总自由度及其分解。

总自由度: $df_T = a \times k - 1 = N - 1$

处理间自由度: $df_t = a - 1$

处理内 (误差) 自由度: $df_e = df_T - df_t = a \times (k - 1)$

(3) 均方及其计算。

总均方: $MS_T = SS_T / df_T$

处理间均方: $MS_t = SS_t / df_t$

处理内 (误差) 均方: $MS_e = SS_e / df_e$

3) 列方差分析

列方差分析如表 3-5 所示。

表 3-5 列方差分析

变异来源	平方和 SS	自由度 df	均方 MS	F 值	$F_{0.05}$	$F_{0.01}$
处理间 (因子 A)	$SS_t = \sum_{i=1}^a y_{i*}^2 / k - CT$	$df_t = a - 1$	$MS_t = SS_t / df_t$	MS_t / MS_e		
误差	$SS_e = SS_T - SS_t$	$df_e = a \times (k - 1)$	$MS_e = SS_e / df_e$			
总变异	$SS_T = \sum_{i=1}^a \sum_{j=1}^k y_{ij}^2 - CT$	$df_T = a \times k - 1$	$MS_T = SS_T / df_T$			

4) 检验方法 (F 检验)

采用 F 值出现概率的大小推断两个总体方差是否相等。

(1) 进行无效假设 $H_0: \mu_1 = \mu_2 = \dots = \mu_a$ (或 $H_0: \sigma^2 = 0$) [或备择假设 $H_A: \mu_i$ 不全相等 (或 $H_A: \sigma^2 \neq 0$)]。

(2) 计算统计量 $F = MS_t / MS_e$ 。

(3) 查附表 A-9, 得 $F_{\alpha(df_t, df_e)}$ 。

(4) 判断: 若 $F > F_{\alpha(df_t, df_e)}$, 拒绝 H_0 ; 若 $F \leq F_{\alpha(df_t, df_e)}$, 接受 H_0 。具体情况见以下表述。

若 $F \leq F_{0.05(df_t, df_e)}$, 即 $p \geq 0.05$, 接受 H_0 , 统计学上, 把这一检验结果表述为: 各处理间差异不显著, 在 F 值的右上方标记 “ns”, 或不进行任何标记。

若 $F_{0.05(df_t, df_e)} < F \leq F_{0.01(df_t, df_e)}$, 即 $0.01 \leq p < 0.05$, 拒绝 H_0 , 而接受 H_A , 统计学上把这一检验结果表述为: 各处理间差异达到显著, 在 F 值的右上方标记 “*”。

若 $F > F_{0.01(df_t, df_e)}$, 即 $p < 0.01$, 高度拒绝 H_0 (接受 H_A), 统计学上把这一检验结果表述为: 各处理间差异达到极显著, 在 F 值的右上方标记 “**”。

5) 方差分析的基本步骤

【第一步】相关数据计算, 包括数据求和、各项偏差平方和及其自由度等。

【第二步】列方差分析表, 进行 F 检验, 推断显著性。

【第三步】若检验结果显著, 再进行“多重比较”, 确定不同水平彼此之间的差异显著性。

6) 方差分析的应用条件

一是各观测值相互独立, 并服从正态分布; 二是各组总体方差相等, 即方差齐性。

【例 3.10】男衬衣制造商欲了解所有人造纤维的拉力强度 ($y/\text{kgf} \cdot \text{cm}^{-2}$), 认为人造纤维中棉花百分率 $A(x, \%)$ 是其影响的主因子。今取 5 种棉花百分率 (5 种处理), 分别为 15%、20%、25%、30%、35%, 且每种处理均取 5 个观察值, 按随机办法取得先后顺序, 测试结果如表 3-6 所示。

表 3-6 测试结果

试验号 (i)	棉花百分率 $A(x_i)/\%$	重复观察值[人造纤维的拉力强度 $y_{ij}/\text{kgf} \cdot \text{cm}^{-2}$]				
		1	2	3	4	5
1	15	7	7	15	11	9
2	20	12	17	12	18	18
3	25	14	18	18	19	19
4	30	19	25	22	19	23
5	35	7	10	15	15	11

(1) 试推断人造纤维中棉花百分率是否对其拉力强度有显著影响。
 (2) 试估计当置信度为 95% 时, 棉花百分率为 30% 的均数 (μ_i) 及棉花百分率为 20% 与棉花百分率为 30% 两均数之差 ($\mu_i - \mu_{i'}$) 的置信区间。

(3) 采用“多重比较”, 对不同棉花百分率彼此之间的差异进行显著性检验。

(4) 配合反应曲线, 试建立人造纤维中棉花百分率与其拉力强度的数学模型。

【解】(1) 推断人造纤维中棉花百分率是否对其拉力强度有显著影响。

相关数据计算 (包括观察总值、平均值) 如表 3-7 所示。

表 3-7 数据计算

试验号 (i)	棉花百分率 (x_i)/%	重复观察值[人造纤维的拉力强度 y_{ij} / kgf·cm ⁻²]					y_{i*}	y_{**}	$\bar{y}_{i*}$	$\bar{y}_{**}$
		1	2	3	4	5				
1	15	7	7	15	11	9	49	376	9.8	15.04
2	20	12	17	12	18	18	77		15.4	
3	25	14	18	18	19	19	88		17.6	
4	30	19	25	22	19	23	108		21.6	
5	35	7	10	15	15	11	54		10.8	

各偏差平方和与其自由度计算如下所示。

$$CT = y_{**}^2 / (a \times k) = 376^2 / (5 \times 5) = 5\,655.04$$

$$SS_T = \sum_{i=1}^a \sum_{j=1}^k y_{ij}^2 - CT = (7^2 + 7^2 + \cdots + 11^2) - 5\,655.04 = 636.96$$

$$df_T = a \times k - 1 = 5 \times 5 - 1 = 24$$

$$SS_t = \sum_{i=1}^a y_{i*}^2 / k - CT = (49^2 + 77^2 + 88^2 + 108^2 + 54^2) / 5 - 5\,655.04 = 475.76$$

$$df_t = a - 1 = 5 - 1 = 4$$

$$SS_e = SS_T - SS_t = 636.96 - 475.76 = 161.20$$

$$df_e = df_T - df_t = 24 - 4 = 20, \text{ 或 } df_e = a \times (k - 1) = 5 \times (5 - 1) = 20$$

列出方差分析表, 进行 F 检验, 如表 3-8 所示。

表 3-8 方差分析与 F 检验

变异来源	SS	df	MS	F	$F_{0.05}$	$F_{0.01}$
处理间 (棉花百分率 A)	475.76	4	118.94	14.76	2.87	4.43
误差	161.20	20	8.06			
总变异	639.96	24				

结果表明：处理间的影响达到极显著，即有 99%把握认为人造纤维中棉花百分率对其拉力强度有极显著影响。

(2) 试估计当置信度为 95%时，棉花百分率为 30%的均数 (μ_i) 及棉花百分率为 20%与棉花百分率为 30%两均数之差 ($\mu_i - \mu_{i'}$) 的置信区间。

首先进行模型参数估计。

总平均： $\hat{\mu} = \bar{y}_{**}$ 。

处理效应： $\bar{\tau}_i = \bar{y}_{i*} - \bar{y}_{**}$ 。

第 i 个处理均数 μ 的 $100(1-\alpha)\%$ 置信区间： $\bar{y}_{i*} \pm t_{\alpha/2, a \times (k-1)} \sqrt{MS_e / k}$ 。

任意两个处理均数之差 ($\mu_i - \mu_{i'}$) 的 $100(1-\alpha)\%$ 置信区间： $\bar{y}_{i*} - \bar{y}_{i'*} \pm t_{\alpha/2, a \times (k-1)} \sqrt{2MS_e / k}$ 。

查附表 A-7 知， $t_{\alpha/2, a \times (k-1)} = t_{0.05/2, 5 \times (5-1)} = t_{0.05(\text{双侧}), 20} = 2.086$ 。

当置信度为 95%时，棉花百分率为 30%的均数 (μ_i) 的置信区间： $[21.60 \pm 2.086 \times \sqrt{8.06/5}]$ ，即 $[21.60 \pm 2.65]$ 或 $[18.95, 24.25]$ 。

当置信度为 95%时，棉花百分率 20%与棉花百分率 30%两均数之差 ($\mu_i - \mu_{i'}$) 的置信区间： $[(15.40 - 21.60) \pm 2.086 \times \sqrt{2 \times 8.06/5}]$ ，即 $[-6.20 \pm 3.75]$ 或 $[-9.95, -2.45]$ 。

(3) 采用新复极差法进行“多重比较”，进行不同棉花百分率彼此之间的差异性检验。

对不同试验号，按所得均数 $\bar{y}_{i*}$ 的大小顺序，依次由上至下顺序排列，具体多重比较设计如表 3-9 所示。

表 3-9 多重比较设计

试验号 (i)	$\bar{y}_{i*}$	$\bar{y}_{i*} - 9.8$	$\bar{y}_{i*} - 10.8$	$\bar{y}_{i*} - 15.4$	$\bar{y}_{i*} - 17.6$
4	21.6	11.8	10.8	6.2	4.0
3	17.6	7.8	6.8	2.2	
2	15.4	5.6	4.6		
5	10.8	1.0			
1	9.8				

根据 $MS_e = 8.06$ ， $k = 5$ ，得样本标准误差： $S_y = \sqrt{MS_e / k} = \sqrt{8.06/5} = 1.27$ 。

并根据 $df_e = 20$ ，秩次距 $k = 2, 3, 4, 5$ ，查出 $\alpha = 0.05$ 和 $\alpha = 0.01$ 的各临界 SSR 值，并乘以样本标准误差 (1.27)，即各最小显著极差 LSR，相关试验结果如表 3-10 所示。

表 3-10 试验结果

df_e	秩次距 k	$SSR_{0.05}$	$SSR_{0.01}$	$LSR_{0.05}$	$LSR_{0.01}$
20	2	2.95	4.02	3.75	5.10
	3	3.10	4.22	3.94	5.36
	4	3.18	4.33	4.04	5.50
	5	3.25	4.40	4.13	5.59

将表中的差数与表中相应的最小显著极差比较并标记检验结果，如表 3-11 所示。

表 3-11 检验结果

试验号 (i)	$\bar{y}_{i*}$	$\bar{y}_{i*} - 9.8$	$\bar{y}_{i*} - 10.8$	$\bar{y}_{i*} - 15.4$	$\bar{y}_{i*} - 17.6$
4	21.6	11.8**	10.8**	6.2**	4.0*
3	17.6	7.8**	6.8**	2.2	
2	15.4	5.6**	4.6*		
5	10.8	1.0			
1	9.8				

结果表明：除试验号 3 与试验号 2、试验号 5 与试验号 1 外，其他各试验号（处理）彼此之间的差异均达到（极）显著。

（4）配合反应曲线，建立纤维棉花百分率与衬衣拉力强度的数学模型。

根据处理数，选取合适的正交多项式系数，并进行相关数据计算，如表 3-12 所示。

表 3-12 相关数据计算

试验号 (i)	棉花百分率 $A(x_i)/\%$	观察值和 $\bar{y}_{i*}$	正交对比系数 C_{ij}			
			一次	二次	三次	四次
1	15	49	-2	2	-1	1
2	20	77	-1	-1	2	-4
3	25	88	0	-2	0	6
4	30	108	1	-1	-2	-4
5	35	54	2	2	1	1
正交对比系数平方和 $Q_C = \sum_{i=1}^a C_{ij}^2$			10	14	10	70
常数项 λ_j			1	1	5/6	35/12
效应 $E = \sum_{i=1}^a C_{ij} y_{i*}$			41	-155	-57	-109
平方和 $SS = E^2 / (k \times Q_C)$			33.62	343.21	64.98	33.95
模型参数 $\alpha_j = E / (k \times Q_C)$			0.820 0	-2.214 3	-1.140 0	-0.311 4

列拓展的方差分析如表 3-13 所示。

表 3-13 列拓展的方差分析

变异来源	SS	df	MS	<i>F</i>	<i>F</i> _{0.05}	<i>F</i> _{0.01}
处理间 (棉花百分率 <i>A</i>)	475.76	4	118.94	14.76**	2.87	4.43
一次	33.62	1	33.62	4.17	4.35	8.10
二次	343.21	1	343.21	42.58**		
三次	64.98	1	64.98	8.06*		
四次	33.95	1	33.95	4.21		
误差	161.20	20	8.06			
总变异	639.96	24				

拓展的方差分析结果表明, 虽然四次、一次效应并未达到显著, 但鉴于二次、三次效应已经分别达到极显著、显著, 故应将衬衣拉力强度 (y) 与人造纤维中棉花百分率 (x) 的关系配合成三次多项式曲线方程

$$y = \alpha_0 p_0(x) + \alpha_1 p_1(x) + \alpha_2 p_2(x) + \alpha_3 p_3(x) + \varepsilon$$

这里, $p_u(x)$ 是 u 次正交多项式, 其中

$$p_0(x) = 1$$

$$p_1(x) = \lambda_1 \left[\frac{x - \bar{x}}{d} \right]$$

$$p_2(x) = \lambda_2 \left[\left(\frac{x - \bar{x}}{d} \right)^2 - \frac{m^2 - 1}{12} \right]$$

$$p_3(x) = \lambda_3 \left[\left(\frac{x - \bar{x}}{d} \right)^3 - \left(\frac{x - \bar{x}}{d} \right) \left(\frac{3m^2 - 7}{20} \right) \right]$$

$$p_4(x) = \lambda_4 \left[\left(\frac{x - \bar{x}}{d} \right)^4 - \left(\frac{x - \bar{x}}{d} \right)^2 \left(\frac{3m - 13}{14} \right) + \frac{3(m^2 - 1)(m^2 - m)}{560} \right]$$

⋮

这里, d 为 x 的水平间距; m 为水平数; λ_i 为多项式所固有的常数值; α 为正交多项式中模型参数估计, 且有

$$\begin{aligned} \hat{\alpha}_0 &= \bar{y}_{**} \\ \bar{\alpha}_j &= \frac{\sum_{i=1}^a C_{ij} y_{i*}}{k \sum_{i=1}^a C_{ij}^2} = \alpha_j \end{aligned}$$

针对本例, $m=5$, $\bar{x}=(15+35)/2=25$, $d=5$, 将相关数据代入, 有

$$\hat{y} = 15.04 + 0.8200 \times 1 \times \left(\frac{x-25}{5} \right) - 2.2143 \times 1 \times \left[\left(\frac{x-25}{5} \right)^2 - \frac{5^2-1}{12} \right] - 1.1400 \times \frac{5}{6} \times \left[\left(\frac{x-25}{5} \right)^3 - \left(\frac{x-25}{5} \right) \left(\frac{3 \times 5^2 - 7}{20} \right) \right]$$

整理得

$$\hat{y} = 64.5936 - 9.0100x + 0.4814x^2 - 0.0076x^3$$

因此, 配合成衬衣拉力强度 y 与棉花百分率 x 关系的曲线方程

$$\hat{y} = 64.5936 - 9.0100x + 0.4814x^2 - 0.0076x^3$$

据此方程, 可预测在 15%~35% 内的任意棉花百分率的纤维拉力强度。如当棉花百分率 $x=30\%$ 时的纤维拉力强度的预测值为

$$\hat{y} = 64.5936 - 9.0100 \times 30 + 0.4814 \times 30^2 - 0.0076 \times 30^3 = 22.35 \text{ kgf/cm}^2$$

实际(平均)为 21.6 kgf/cm^2 。两者有一定的接近程度, 表明预测估计具有一定的价值。

设单一因子 A 的处理数为 a , 各处理重复数分别为 $k_1, k_2, \dots, k_a$, 则

试验观测值总数: $N = \sum k_i$

矫正数: $CT = y_{**}^2 / N$

总偏差平方和: $SS_T = \sum_{i=1}^a \sum_{j=1}^{k_i} y_{ij}^2 - CT$

总自由度: $df_T = N - 1$

处理偏差平方和: $SS_t = \sum_{i=1}^a y_{i*}^2 / k_i - CT$

处理自由度: $df_t = a - 1$

误差平方和: $SS_e = SS_T - SS_t$

误差自由度: $df_e = df_T - df_t = N - a$

【例 3.11】五种不同配方(五个处理)橡胶材料的拉伸强度试验结果如表 3-14 所示, 试比较不同配方之间对拉伸强度有无差异显著性。

表 3-14 五种不同配方橡胶材料的拉伸强度试验结果

试验号 (i)	配方 (B)	重复观察值 y_{ij} (拉伸强度 / MPa)					
		1	2	3	4	5	6
1	B_1	21.5	19.5	20.0	22.0	18.0	20.0
2	B_2	16.0	18.5	17.0	15.5	20.0	16.0
3	B_3	19.0	17.5	20.0	18.0	17.0	
4	B_4	21.0	18.5	19.0	20.0		
5	B_5	15.5	18.0	17.0	16.0		

【解】相关数据求和，如表 3-15 所示。

表 3-15 相关数据求和

试验号 (<i>i</i>)	配方 (<i>B</i>)	重复观察值 y_{ij} (拉伸强度/MPa)						k_i	N	y_{i*}	y_{**}	$\bar{y}_{i*}$	$\bar{y}_{**}$
		1	2	3	4	5	6						
1	B_1	21.5	19.5	20.0	22.0	18.0	20.0	6	25	121.0	460.5	20.2	18.42
2	B_2	16.0	18.5	17.0	15.5	20.0	16.0	6		103.0		17.2	
3	B_3	19.0	17.5	20.0	18.0	17.0		5		91.5		18.3	
4	B_4	21.0	18.5	19.0	20.0			4		78.5		19.6	
5	B_5	15.5	18.0	17.0	16.0			4		66.5		16.6	

计算各偏差平方和与其自由度

$$CT = y_{**}^2 / N = 460.5^2 / 25 = 8\,482.21$$

$$SS_T = \sum_{i=1}^a \sum_{j=1}^{k_i} y_{ij}^2 - CT = (21.5^2 + 19.5^2 + \cdots + 16.0^2) - 8\,482.41 = 85.34$$

$$df_T = N - 1 = 25 - 1 = 24$$

$$SS_t = \sum_{i=1}^a y_{i*}^2 / k_i - CT = \left(121.0^2 / 6 + 103.0^2 / 6 + 91.5^2 / 5 + 78.8^2 / 4 + 66.5^2 / 4 \right) - 8\,482.41 = 46.50$$

$$df_t = a - 1 = 5 - 1 = 4$$

$$SS_e = SS_T - SS_t = 85.34 - 46.50 = 38.84$$

$$df_e = df_T - df_t = 24 - 4 = 20$$

列出方差分析表，进行 F 检验，如表 3-16 所示。

表 3-16 方差分析与 F 检验

变异来源	SS	df	MS	F	$F_{0.05}$	$F_{0.01}$
处理间 (配方)	46.50	4	11.63	5.99	2.87	4.43
误差	38.84	20	1.94			
总变异	85.34	24				

结果表明：配方间的影响达到极显著，即有 99%把握认为配方对所得橡胶材料的拉伸强度有极显著的影响。

采用新复极差法作多重比较，进一步检验各配方之间差异显著性。

【第一步】对不同配方号，按所得均数 $\bar{y}_{i*}$ 的大小顺序，依次由上至下顺序排列，具体多重比较数据如表 3-17 所示。

表 3-17 多重比较数据

配方 (B)	$\bar{y}_{i*}$	$\bar{y}_{i*} - 16.6$	$\bar{y}_{i*} - 17.2$	$\bar{y}_{i*} - 18.3$	$\bar{y}_{i*} - 19.6$
B_1	20.2	3.6	3.0	1.9	0.6
B_2	19.6	3.0	2.4	1.3	
B_3	18.3	1.7	1.1		
B_4	17.2	0.6			
B_5	16.6				

【第二步】因各处理重复数不等，应先计算出平均重复次数 k_0 来代替标准误 $S_{\bar{y}} = \sqrt{MS_e / k}$ 中的 k 。

针对此例

$$k_0 = \frac{1}{a-1} \left[\sum k_i - \frac{\sum k_i^2}{\sum k_i} \right] = \frac{1}{5-1} \left[25 - \frac{6^2 + 6^2 + 5^2 + 4^2 + 4^2}{6+6+5+4+4} \right] = 4.96$$

因此，标准误为

$$S_{\bar{y}} = \sqrt{MS_e / k} = \sqrt{1.94 / 4.96} = 0.625$$

并根据 $df_e = 20$ ，秩次距 $k = 2, 3, 4, 5$ ，查出 $\alpha = 0.05$ 与 $\alpha = 0.01$ 的临界 SSR 值，再乘以 0.625，即得各最小显极差 LSR，所得结果如表 3-18 所示。

表 3-18 各最小显极差结果

df_e	秩次距 k	$SSR_{0.05}$	$SSR_{0.01}$	$LSR_{0.05}$	$LSR_{0.01}$
20	2	2.95	4.02	1.844	2.513
	3	3.10	4.22	1.938	2.638
	4	3.18	4.33	1.988	2.706
	5	3.25	4.40	2.031	2.750

【第三步】将差数与相应的最小显著极差比较并标记检验结果，如表 3-19 所示。

表 3-19 比较与检验结果

配方 (B)	$\bar{y}_{i*}$	$\bar{y}_{i*} - 16.6$	$\bar{y}_{i*} - 17.2$	$\bar{y}_{i*} - 18.3$	$\bar{y}_{i*} - 19.6$
B_1	20.2	3.6**	3.0**	1.9	0.6
B_2	19.6	3.0**	2.4*	1.3	
B_3	18.3	1.7	1.1		
B_4	17.2	0.6			
B_5	16.6				

结果表明： B_1 、 B_4 得平均拉伸强度极显著或显著高于 B_2 、 B_5 得平均拉伸强度，其余不同配方彼此之间得差异并没有达到显著。

2. 随机区组设计

1) 数据模式

设单一因子 A 有 a 个处理, q 个区组, 共有 $a \times q$ 个观察值, 其数据模式如表 3-20 所示。

表 3-20 因子 A 数据模式

因子 A (处理)	区组观测值 (y_{ij})					
	I	II	...	j	...	q
1	y_{11}	y_{12}	...	y_{1j}	...	y_{1q}
2	y_{21}	y_{22}	...	y_{2j}	...	y_{2q}
$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$
i	y_{i1}	y_{i2}	...	y_{ij}	...	y_{iq}
$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$
a	y_{a1}	y_{a2}	...	y_{aj}	...	y_{aq}

2) 数学模型

单因子随机区组试验中, 每一区组观察值 y_{ij} 的数学模型可表示为 $y_{ij} = \mu + \alpha_i + \beta_j + \varepsilon_{ij}$, 其中 μ 为全部试验观测值总体的平均数, $\mu = \sum_{i=1}^a \mu_i / a$; α_i 为第 i 个处理效应, 表示处理 i 对试验结果产生的影响, $\sum_{i=1}^a \alpha_i = 0$; β_j 为第 j 个区组效应, 表示区组 j 对试验结果产生的影响, $\sum_{j=1}^q \beta_j = 0$; ε_{ij} 是试验误差, 相互独立, 且服从正态分布 $N(0, \sigma^2)$ 。

3) 总变异分解

随机区组设计的方差分析关键仍在于总变异分解, 因随机区组设计可将区组间变异从完全随机设计的组内变异中分离出来以反映不同区组对结果的影响, 故随机区组设计全部测量值总变异 (T) 相应地分成处理 (t)、区组 (q) 与误差 (e) 三个部分, 即有

$$\begin{aligned} SS_T &= SS_t + SS_q + SS_e \\ df_T &= df_t + df_q + df_e \end{aligned}$$

相关数据计算如表 3-21 所示。

表 3-21 相关数据计算

处理	重复观测值 (y_{ij})						y_{i*}	$\bar{y}_{i*}$	处理平方之和 $\left(\sum_{j=1}^q y_{ij}^2\right)$	处理和之平方 (y_{i*}^2)
	I	II	...	j	...	q				
1	y_{11}	y_{12}	...	y_{1j}	...	y_{1q}	y_{1*}	$\bar{y}_{1*}$	$\sum_{j=1}^q y_{1j}^2$	y_{1*}^2
2	y_{21}	y_{22}	...	y_{2j}	...	y_{2q}	y_{2*}	$\bar{y}_{2*}$	$\sum_{j=1}^q y_{2j}^2$	y_{2*}^2
$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
i	y_{i1}	y_{i2}	...	y_{ij}	...	y_{iq}	y_{i*}	$\bar{y}_{i*}$	$\sum_{j=1}^q y_{ij}^2$	y_{i*}^2
$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
a	y_{a1}	y_{a2}	...	y_{aj}	...	y_{aq}	y_{a*}	$\bar{y}_{a*}$	$\sum_{j=1}^q y_{aj}^2$	y_{a*}^2
y_{*j}	y_{*1}	y_{*2}	...	y_{*j}	...	y_{*q}	y_{**}	$\bar{y}_{**}$	$\sum_{i=1}^a \sum_{j=1}^q y_{ij}^2$	$\sum_{i=1}^a y_{i*}^2$
y_{*j}^2	y_{*1}^2	y_{*2}^2	...	y_{*j}^2	...	y_{*q}^2	—	—		$\sum_{j=1}^q y_{*j}^2$

(1) 总变异：反映全部试验数据间大小不等的状况。

$$SS_T = \sum_{i=1}^a \sum_{j=1}^q y_{ij}^2 - CT$$

其中

$$CT = \left(\sum_{i=1}^a \sum_{j=1}^q y_{ij} \right)^2 / N = y_{**}^2 / (a \times q)$$

$$df_T = a \times q - 1$$

(2) 处理组间变异：各处理组间测量值的均数大小不等。

$$SS_t = \sum_{i=1}^a y_{i*}^2 / q - CT$$

$$df_t = a - 1$$

(3) 区组间变异：各区组间测量值的均数大小不等。

$$SS_q = \sum_{j=1}^q y_{*j}^2 / a - CT$$

$$df_q = q - 1$$

(4) 误差变异：反映随机误差产生的变异。

$$SS_e = SS_T - SS_t - SS_q$$

$$df_e = df_T - df_t - df_q = (a - 1) \times (q - 1)$$

4) 列方差分析

列方差分析如表 3-22 所示。

表 3-22 列方差分析

变异来源	SS	df	MS	F	$F_{0.05}$	$F_{0.01}$
区组	$SS_q = \sum_{j=1}^q y_{*j}^2 / a - CT$	$df_q = q - 1$	$MS_q = \frac{SS_q}{df_q}$	MS_q / MS_e		
处理间 (因子 A)	$SS_t = \sum_{i=1}^q y_{i*}^2 / q - CT$	$df_t = a - 1$	$MS_t = \frac{SS_t}{df_t}$	MS_t / MS_e		
误差	$SS_e = SS_T - SS_t - SS_q$	$df_e = (a - 1) \times (q - 1)$	$MS_e = \frac{SS_e}{df_e}$			
总变异	$SS_T = \sum_{i=1}^a \sum_{j=1}^q y_{ij}^2 - CT$	$df_T = a \times q - 1$				

【例 3.12】研究甲、乙、丙三种改性剂对高聚物性能的影响，假定原料批次也设置为影响因子，以区组加以考虑。拟用 6 批原料，每批 3 份（500g/份），随机安排试验，研究甲、乙、丙三种改性剂对高聚物拉伸强度的影响，结果如表 3-23 所示。

- (1) 不同改性剂之间对高聚物拉伸强度提高是否相同？
(2) 不同原料批次之间对高聚物拉伸强度提高是否相同？

表 3-23 三种改性剂对高聚物拉伸强度的影响

处理（改性剂）	区组（原料批次号）观察值（ y_{ij} ）					
	I	II	III	IV	V	VI
甲	64	53	71	41	50	42
乙	65	54	68	46	58	40
丙	73	59	79	38	65	46

【解】建立检验假设

H_0 : $\mu_{\text{甲}} = \mu_{\text{乙}} = \mu_{\text{丙}}$ （三种改性剂对高聚物拉伸强度提高作用相同）

H_A : $\mu_{\text{甲}}$ 、 $\mu_{\text{乙}}$ 、 $\mu_{\text{丙}}$ 不全相等（三种改性剂对高聚物拉伸强度提高作用不全相同）

H_0 : $\mu_1 = \mu_2 = \cdots = \mu_6$ （原料批次对高聚物拉伸强度提高无影响）

H_A : $\mu_1, \mu_2, \cdots, \mu_6$ 不全相等（原料批次对高聚物拉伸强度提高有影响）

数据求和，如表 3-24 所示。

表 3-24 数据求和

处理（改性剂）	区组（原料批次号）观察值（ y_{ij} ）						y_{i*}	$\sum_{j=1}^q y_{ij}^2$	$\sum_{i=1}^a \sum_{j=1}^q y_{ij}^2$
	I	II	III	IV	V	VI			
甲	64	53	71	41	50	42	321	17 891	59 572
乙	65	54	68	46	58	40	331	18 845	
丙	73	59	79	38	65	46	360	228 36	
y_{*j}	202	166	218	125	173	128	$1012(y_{**})$		

计算各偏差平方和与其自由度

$$CT = y_{**}^2 / (a \times q) = 1012^2 / (3 \times 6) = 56\,896.89$$

$$SS_T = \sum_{i=1}^a \sum_{j=1}^q y_{ij}^2 - CT = 59\,572 - 56\,896.89 = 2\,675.111$$

$$df_T = a \times q - 1 = 3 \times 6 - 1 = 17$$

$$SS_t = \sum_{i=1}^a y_{i*}^2 / q - CT = (321^2 + 331^2 + 360^2) / 6 - 56\,896.89 = 136.778$$

$$df_t = a - 1 = 3 - 1 = 2$$

$$SS_q = \sum_{j=1}^q y_{*j}^2 / a - CT = (202^2 + 166^2 + \cdots + 128^2) / 3 - 56\,896.89 = 2\,377.111$$

$$df_q = q - 1 = 6 - 1 = 5$$

$$SS_e = SS_T - SS_t - SS_q = 2\,675.111 - 136.778 - 2\,377.111 = 161.222$$

$$df_e = (a - 1) \times (q - 1) = (3 - 1) \times (6 - 1) = 10$$

列方差分析如表 3-25 所示。

表 3-25 列方差分析

变异来源	SS	df	MS	F	$F_{0.05}$	$F_{0.01}$
区组（原料批次）间	2 377.111	5	475.422	29.49**	3.33	5.64
处理组（改性剂种类）间	136.778	2	68.389	4.24*	4.10	7.56
误差	161.222	10	16.122			
总变异	2 675.111	17				

处理因子查附表 A-9, $F_{0.05,(2,10)} = 4.10$, 因 $F = 4.24 > F_{0.05,(2,10)} = 4.10$, 故 $p < 0.05$ 。

【推断】按 $\alpha = 0.05$ 检验水准, 拒绝 H_0 (即接受 H_A) 差别有统计学意义, 可认为改性剂种类对高聚物拉伸强度提高有显著影响。

区组因子查附表 A-9, $F_{0.01,(5,10)} = 5.64$, 因 $F = 29.49 > F_{0.01,(5,10)} = 5.64$, 故 $p < 0.01$ 。

【推断】按 $\alpha = 0.01$ 检验水准, 高度拒绝 H_0 (即高度接受 H_A), 差别有统计学意义, 可认为原料批次对高聚物拉伸强度提高有极显著影响。

值得注意的是, 若不把批次设计看成是完全随机区组化设计, 而是当作完全随机单因子设计, 其方差分析如表 3-26 所示。

表 3-26 方差分析

变异来源	SS	df	MS	F	$F_{0.05}$	$F_{0.01}$
处理组（改性剂）间	136.778	2	68.389	0.404 2	3.49	5.95
误差	2 538.22	15	169.215			
总变异	2 675.111	17				

结果表明：改性剂种类对高聚物拉伸强度提高并无显著影响。这一结论与事实并不相符，因而是错误的。

采用完全随机区组试验设计，减少了干扰成分，使得 3 种改性剂间的高聚物拉伸强度提高差异能够充分显示出来。

3. 拉丁方设计

完全随机设计只涉及一个处理因子，随机区组设计涉及一个处理因子和一个区组因子。若试验涉及一个处理因子和两个控制因子，且每个因子的水平数相等，此时可采用拉丁方设计来安排试验，将两个控制因子分别安排在拉丁方的行和列上。

p 阶拉丁方是由 p 个拉丁字母排成的 $p \times p$ 方阵，每行或每列中每个字母都只出现一次。下面给出了几个拉丁方设计。

(1) 3×3

A	B	C
C	A	B
B	C	A

(2) 4×4

A	B	C	D
D	A	B	C
C	D	A	B
B	C	D	A

(3) 5×5

A	B	C	D	E
E	A	B	C	D
D	E	A	B	C
C	D	E	A	B
B	C	D	E	A

应用时，根据水平数 p 来选定拉丁方大小。

拉丁方设计是在随机区组设计基础上发展的，可多安排一个已知的对试验结果有影响的非处理因子来提高效率。

拉丁方设计的要求：一定是 p 因子，且 p 因子水平数相等；行间、列间、处理间均无交互作用；各行、列处理的方差齐。

拉丁方设计的优点：可同时研究 p 个因子，减少试验次数；从组内变异中不但分

离出行区组变异,而且还分离出列区组变异,使误差变异进一步减小,因而精确度高,适用有两种较明显误差的试验。

拉丁方设计的缺点:要求处理组数与所要控制的 $(p-1)$ 个因子水平数相等(即处理数=横行数=直行数=重复数);试验空间很难伸缩,且一般试验不容易满足此条件,而且数据缺失会增加统计分析的难度。(一般处理数只限于 5~8 个。整个试验操作要在短期内完成,工作不好安排。)

【例 3.13】研究 A、B、C、D 四种改性剂,以及甲、乙、丙、丁四种制备方法对高聚物性能的影响。拟用四批原料,每批四份(500g/份),每份高聚物随机搭配一种改性剂、随机采用一种制备方法,研究改性结果,试验结果如表 3-27 所示。

- (1) 改性剂种类是否影响高聚物撕裂强度?
- (2) 制备方法是否影响高聚物撕裂强度?
- (3) 不同原料批次对高聚物改性效果是否相同?

表 3-27 试验结果

原料批次	制备方法			
	甲	乙	丙	丁
1	80 (D)	70 (B)	51 (C)	48 (A)
2	47 (A)	75 (C)	78 (D)	45 (B)
3	48 (B)	80 (D)	47 (A)	52 (C)
4	46 (C)	81 (A)	49 (B)	77 (D)

【解】建立检验假设

$H_{\text{处理}0}$: $\mu_A = \mu_B = \mu_C = \mu_D$, 即不同改性剂种类对高聚物撕裂强度提高相同

$H_{\text{处理}A}$: μ_A 、 μ_B 、 μ_C 、 μ_D 不全相等,即不同改性剂种类对高聚物撕裂强度提高不全相同

$H_{\text{行}0}$: $\mu_1 = \mu_2 = \mu_3 = \mu_4$, 即不同原料批次对高聚物撕裂强度提高相同

$H_{\text{行}A}$: μ_1 、 μ_2 、 μ_3 、 μ_4 不全相等,即不同原料批次对高聚物撕裂强度提高不全相同

$H_{\text{列}0}$: $\mu_{\text{甲}} = \mu_{\text{乙}} = \mu_{\text{丙}} = \mu_{\text{丁}}$, 即不同制备方法对高聚物撕裂强度提高相同

$H_{\text{列}A}$: $\mu_{\text{甲}}$ 、 $\mu_{\text{乙}}$ 、 $\mu_{\text{丙}}$ 、 $\mu_{\text{丁}}$ 不全相等,即不同制备方法对高聚物撕裂强度提高不全相同

数据求和,如表 3-28 所示。

表 3-28 数据求和

原料批次	制备方法				y_{i*} (横行)
	甲	乙	丙	丁	
1	80 (D)	70 (B)	51 (C)	48 (A)	249
2	47 (A)	75 (C)	78 (D)	45 (B)	245
3	48 (B)	80 (D)	47 (A)	52 (C)	227
4	46 (C)	81 (A)	49 (B)	77 (D)	253
y_{*j} (直列)	221	306	225	222	
处理和 (改性剂种类) $\sum y_{ij}$	A	B	C	D	974(y_{**})
	223	212	224	315	
各种类平方之和 $\sum y_{ij}^2$	$\sum A^2$	$\sum B^2$	$\sum C^2$	$\sum D^2$	$62\,772\left(\sum_{i=1}^p\sum_{j=1}^py_{ij}^2\right)$
	13\,283	11\,630	13\,046	24\,813	

计算各偏差平方和与其自由度

$$CT = y_{**}^2 / (p \times p) = 974^2 / (4 \times 4) = 59\,292.25$$

$$SS_T = \sum_{i=1}^p \sum_{j=1}^p y_{ij}^2 - CT = 62\,772 - 59\,292.25 = 3\,479.75$$

$$df_T = p \times p - 1 = 4 \times 4 - 1 = 15$$

$$SS_t = \sum_{k=A}^D y_{**k}^2 / p - CT = (223^2 + 212^2 + 224^2 + 315^2) / 4 - 59\,292.25 = 1\,726.25$$

$$df_t = p - 1 = 4 - 1 = 3$$

$$SS_{\text{横行}} = \sum_{i=1}^p y_{i*}^2 / p - CT = (249^2 + 245^2 + 227^2 + 253^2) / 4 - 59\,292.25 = 98.75$$

$$df_{\text{横行}} = p - 1 = 4 - 1 = 3$$

$$SS_{\text{直列}} = \sum_{j=1}^p y_{*j}^2 / p - CT = (221^2 + 306^2 + 225^2 + 222^2) / 4 - 59\,292.25 = 1\,304.25$$

$$df_{\text{直列}} = p - 1 = 4 - 1 = 3$$

$$SS_e = SS_T - SS_t - SS_{\text{横行}} - SS_{\text{直列}} = 3\,479.75 - 1\,726.25 - 98.75 - 1\,304.25 = 350.5$$

$$df_e = df_T - df_t - df_{\text{横行}} - df_{\text{直列}} = 15 - 3 - 3 - 3 = 6$$

列方差分析如表 3-29 所示。

表 3-29 列方差分析

变异来源	SS	df	MS	F	$F_{0.05}$	$F_{0.01}$
处理间 (改性剂种类)	1\,726.25	3	575.417	9.85**	4.76	9.78
行区组 (原料批次)	98.75	3	32.917	0.56	4.76	9.78
列区组 (制备方法)	1\,304.25	3	434.750	7.44*	4.76	9.78
误差	350.50	6	58.417			
总变异	3\,479.75	15				

对改性剂种类: 由 $df_{\text{处理}} = 3$, $df_{\text{误差}} = 6$, 查附表 A-9, $F_{0.05,(3,6)} = 4.76$, $F_{0.01,(3,6)} = 9.78$, 得 $p < 0.01$, 按 $\alpha = 0.01$ 检验水准, 高度拒绝 $H_{\text{处理}0}$ (高度接受 $H_{\text{处理}A}$), 差别有统计学意义, 可认为改性剂种类对高聚物撕裂强度有极显著影响。

对原料批次: 由 $df_{\text{行}} = 3$, $df_{\text{误差}} = 6$, 查附表 A-9, $F_{0.05,(3,6)} = 4.76$, $F_{0.01,(3,6)} = 9.78$, 得 $p > 0.05$, 按 $\alpha = 0.05$ 检验水准, 无法拒绝 $H_{\text{行}0}$ (应该接受 $H_{\text{行}A}$), 差别无统计学意义, 尚不能认为原料批次对高聚物撕裂强度造成影响。

对制备方法: 由 $df_{\text{列}} = 3$, $df_{\text{误差}} = 6$, 查附表 A-9, $F_{0.05,(3,6)} = 4.76$, $F_{0.01,(3,6)} = 9.78$, 得 $p < 0.05$, 按 $\alpha = 0.05$ 检验水准, 拒绝 $H_{\text{列}0}$ (接受 $H_{\text{列}A}$), 差别有统计学意义, 可认为制备方法会显著影响高聚物撕裂强度。

4. 希腊拉丁方设计

在一个 $p \times p$ 的拉丁方 (采用拉丁字母表示) 上重叠第二个 $p \times p$ 拉丁方 (采用希腊字母表示), 若这两个拉丁方中, 每一个拉丁字母与每一个希腊字母当且仅当相遇一次, 则这两个拉丁方是正交的, 由这两个拉丁方所确定的设计称为希腊拉丁方设计。其特点有: 可以控制用于 3 个方面的外部变异, 即区组的 3 个方向控制; 可容纳 4 个因子 (即横行、直列、拉丁字母、希腊字母), 每一个因子有 p 种水平, 共有 p^2 个试验单元; 除 $p = 6$ 之外, 所有 $p \geq 3$ 的希腊拉丁方均存在。

以 4×4 希腊拉丁方设计为例, 其数据模式及其相关数据计算见表 3-30, 方差分析见表 3-31, 其中 $CT = y_{****}^2 / (p \times p)$ 。

表 3-30 4×4 希腊拉丁方设计数据模式及其相关数据计算

横行	直列				y_{i***} (横行)
	1	2	3	4	
1	A_α	B_β	C_γ	D_δ	y_{1***}
2	B_β	A_γ	D_δ	C_α	y_{2***}
3	C_γ	D_α	A_δ	B_γ	y_{3***}
4	D_δ	C_δ	B_α	A_β	y_{4***}
y_{**j} (直列)	y_{***1}	y_{***2}	y_{***3}	y_{***4}	y_{***}
y_{*j**} (处理-拉丁字母)	y_{*A**}	y_{*B**}	y_{*C**}	y_{*D**}	
y_{**k*} (处理-希腊字母)	$y_{**\alpha*}$	$y_{**\beta*}$	$y_{**\gamma*}$	$y_{**\delta*}$	

表 3-31 4×4 希腊拉丁方设计方差分析

变异来源	SS	df	MS	F	$F_{0.05}$	$F_{0.01}$
拉丁字母处理	$SS_{\text{拉丁}} = \sum_{j=1}^p y_{*j**}^2 / p - CT$	$df_{\text{拉丁}} = p - 1$	$MS_{\text{拉丁}} = \frac{SS_{\text{拉丁}}}{df_{\text{拉丁}}}$	$\frac{MS_{\text{拉丁}}}{MS_e}$		
希腊字母处理	$SS_{\text{希腊}} = \sum_{k=1}^p y_{**k*}^2 / p - CT$	$df_{\text{希腊}} = p - 1$	$MS_{\text{希腊}} = \frac{SS_{\text{希腊}}}{df_{\text{希腊}}}$	$\frac{MS_{\text{希腊}}}{MS_e}$		
横行	$SS_{\text{横行}} = \sum_{i=1}^p y_{i**k*}^2 / p - CT$	$df_{\text{横行}} = p - 1$	$MS_{\text{横行}} = \frac{SS_{\text{横行}}}{df_{\text{横行}}}$	$\frac{MS_{\text{横行}}}{MS_e}$		
直列	$SS_{\text{直列}} = \sum_{l=1}^p y_{i***l}^2 / p - CT$	$df_{\text{直列}} = p - 1$	$MS_{\text{直列}} = \frac{SS_{\text{直列}}}{df_{\text{直列}}}$	$\frac{MS_{\text{直列}}}{MS_e}$		
误差	$SS_e = SS_T - SS_{\text{拉丁}} - SS_{\text{希腊}} - SS_{\text{横行}} - SS_{\text{直列}}$	$df_e = (p-1) \times (p-1)$	$MS_e = \frac{SS_e}{df_e}$			
总变异	$SS_T = \sum_{i=1}^p \sum_{j=1}^p \sum_{k=1}^p \sum_{l=1}^p y_{ijkl}^2 - CT$	$df_T = p^2 - 1$				

【例 3.14】炸药试验。在拉丁方设计基础上增加一个测定装配的因子，也有五种水平，以 α 、 β 、 γ 、 δ 、 ε 表示。具体希腊拉丁方设计及其结果如表 3-32 所示。试进行方差分析。

表 3-32 希腊拉丁方设计及其结果

原料批次	操作员				
	1	2	3	4	5
1	$A_\alpha = 24$	$B_\gamma = 20$	$C_\varepsilon = 19$	$D_\beta = 24$	$E_\delta = 24$
2	$B_\beta = 17$	$C_\delta = 24$	$D_\alpha = 30$	$E_\gamma = 27$	$A_\varepsilon = 36$
3	$C_\gamma = 18$	$D_\varepsilon = 38$	$E_\beta = 26$	$A_\delta = 27$	$B_\alpha = 21$
4	$D_\delta = 26$	$E_\alpha = 31$	$A_\gamma = 26$	$B_\varepsilon = 23$	$C_\beta = 22$
5	$E_\varepsilon = 22$	$A_\beta = 30$	$B_\delta = 20$	$C_\alpha = 29$	$D_\gamma = 31$

【解】数据求和，如表 3-33 所示。

表 3-33 数据求和

原料批次	操作员					y_{i***} (横行)
	1	2	3	4	5	
1	$A_\alpha = 24$	$B_\gamma = 20$	$C_\varepsilon = 19$	$D_\beta = 24$	$E_\delta = 24$	111
2	$B_\beta = 17$	$C_\delta = 24$	$D_\alpha = 30$	$E_\gamma = 27$	$A_\varepsilon = 36$	134
3	$C_\gamma = 18$	$D_\varepsilon = 38$	$E_\beta = 26$	$A_\delta = 27$	$B_\alpha = 21$	130
4	$D_\delta = 26$	$E_\alpha = 31$	$A_\gamma = 26$	$B_\varepsilon = 23$	$C_\beta = 22$	128
5	$E_\varepsilon = 22$	$A_\beta = 30$	$B_\delta = 20$	$C_\alpha = 29$	$D_\gamma = 31$	132
y_{i***l} (直列)	107	143	121	130	134	635(y_{****})
处理-拉丁字母 y_{*j**}	143	101	112	149	130	
处理-希腊字母 y_{**k*}	135	119	122	121	138	

计算各偏差平方和与其自由度

$$CT = y_{***}^2 / (p \times p) = 635^2 / (5 \times 5) = 16129$$

$$SS_T = \sum_{i=1}^p \sum_{j=1}^p \sum_{k=1}^p \sum_{l=1}^p y_{ijkl}^2 - CT = (24^2 + 20^2 + \cdots + 31^2) - 16129 = 676$$

$$df_T = p \times p - 1 = 5 \times 5 - 1 = 24$$

$$SS_{\text{拉丁}} = \sum_{j=1}^p y_{*j**}^2 / p - CT = (143^2 + 101^2 + \cdots + 130^2) / 5 - 16129 = 330$$

$$df_{\text{拉丁}} = p - 1 = 5 - 1 = 4$$

$$SS_{\text{希腊}} = \sum_{k=1}^p y_{**k*}^2 / p - CT = (135^2 + 119^2 + \cdots + 138^2) / 5 - 16129 = 62$$

$$df_{\text{希腊}} = p - 1 = 5 - 1 = 4$$

$$SS_{\text{横行}} = \sum_{i=1}^p y_{i***}^2 / p - CT = (111^2 + 134^2 + \cdots + 132^2) / 5 - 16129 = 68$$

$$df_{\text{横行}} = p - 1 = 5 - 1 = 4$$

$$SS_{\text{直列}} = \sum_{l=1}^p y_{***l}^2 / p - CT = (107^2 + 143^2 + \cdots + 134^2) / 5 - 16129 = 150$$

$$df_{\text{直列}} = p - 1 = 5 - 1 = 4$$

$$SS_e = SS_T - SS_{\text{拉丁}} - SS_{\text{希腊}} - SS_{\text{横行}} - SS_{\text{直列}} = 66$$

$$df_e = df_T - df_{\text{拉丁}} - df_{\text{希腊}} - df_{\text{横行}} - df_{\text{直列}} = 24 - 4 - 4 - 4 - 4 = 8$$

列方差分析如表 3-34 所示。

表 3-34 列方差分析

变异来源	SS	df	MS	F	$F_{0.05}$	$F_{0.01}$
成分（处理-拉丁）	330	4	82.5	10**	3.84	7.01
测验装配（处理-希腊）	62	4	15.5	1.878 788	3.84	7.01
原料（横行）	68	4	17	2.060 606	3.84	7.01
操作员（直列）	150	4	37.5	4.545 455*	3.84	7.01
误差	66	8	8.25			
总变异	676	24				

结果表明：混合成分及操作员对炸药的爆炸力有（极）显著的影响，而测验装配及原料批次则对炸药的爆炸力无显著影响。

对比前例，同等条件下，希腊拉丁方设计因分离出测验装配所引起的变异，试验误差将减小。然而，试验误差减少的同时，其自由度也相应减少（从 $df_e = 12$ 减少到 $df_e = 8$ ）。因而，估计误差的自由度减少，测验的灵敏度也相应降低。

3.4.3 两因子试验设计及其方差分析

1. 完全随机设计

【情形一】两因子无重复观测试验资料的方差分析

1) 试验数据模型

两因子完全随机试验中，设 A 因子有 a 种水平， B 因子有 b 种水平，共有 $a \times b$ 种完全组合，即有 $a \times b$ 个处理。不考虑重复观察，其试验方案设计及结果如表 3-35 所示。

表 3-35 试验方案设计及结果

因子 A	因子 B	观察值 (y_{ih})	y_{i*}	y_{*h}	$\bar{y}_{i*}$	$\bar{y}_{*h}$	
A_1	B_1	y_{11}	y_{11}	y_{*1}	$\bar{y}_{1*}$	y_{*1}	
	B_2	y_{12}		y_{*2}		y_{*2}	
	$\vdots$	$\vdots$		$\vdots$		$\vdots$	
	B_h	y_{1h}		y_{*h}		y_{*h}	
	$\vdots$	$\vdots$		$\vdots$		$\vdots$	
	B_b	y_{1b}		y_{*b}		y_{*b}	
A_2	B_1	y_{21}	y_2	—	$\bar{y}_{2*}$		
	B_2	y_{22}					
	$\vdots$	$\vdots$					
	B_h	y_{2h}					
	$\vdots$	$\vdots$					
	B_b	y_{2b}					
$\vdots$	B_1	$\vdots$	$\vdots$	—	$\vdots$	—	
A_i	B_2	y_{i1}	y_{ia}		$\bar{y}_{i*}$		
	$\vdots$	y_{i2}					
	B_h	$\vdots$					
	$\vdots$	y_{ih}					
	$\vdots$	$\vdots$					
	B_b	$\vdots$					
$\vdots$	B_1	y_{ib}	$\vdots$		$\vdots$		
$\vdots$	$\vdots$	$\vdots$	$\vdots$		$\bar{y}_{a*}$		
A_a	B_1	y_{a1}	y_a				$\bar{y}_{a*}$
	B_2	y_{a1}					
	$\vdots$	$\vdots$					
	$\vdots$	$\vdots$					
	B_h	y_{ah}					
	$\vdots$	$\vdots$					
$\vdots$	B_b	y_{ab}					
总计			y_{**}	$\bar{y}_{**}$			

表 3-35 中， $y_{i*} = \sum_{j=1}^b y_{ij}$ ， $\bar{y}_{i*} = \sum_{j=1}^b y_{ij} / b$ ， $y_{*h} = \sum_{i=1}^a y_{ih}$ ， $\bar{y}_{*h} = \sum_{i=1}^a y_{ih} / a$ ， $y_{**} = \sum_{i=1}^a \sum_{j=1}^b y_{ij}$ ，

$$\bar{y}_{**} = \sum_{i=1}^a \sum_{h=1}^b y_{ih} / (a \times b)。$$

2) 数学模型

两因子无重复观测试验中，任一试验结果 y_{ih} 的数学模型可表示为

$$y_{ih} = \mu + \alpha_i + \beta_h + \varepsilon_{ih} (i=1,2,\cdots,a; h=1,2,\cdots,b)$$

其中， μ 为总体均数； α_i 、 β_h 分别为 A_i 、 B_h 的效应（ $\alpha_i = \mu_i - \mu$ ， $\beta_h = \mu_h - \mu$ ， μ_i 、 μ_h 分别为 A_i 、 B_h 观测值总体均数）， $\sum \alpha_i = 0$ ， $\sum \beta_h = 0$ ； ε_{ih} 为随机误差，相互独立，且服从 $N(0, \sigma^2)$ 。

3) 列方差分析

列方差分析如表 3-36 所示。

表 3-36 列方差分析

变异来源	SS	df	MS	F	$F_{0.05}$	$F_{0.01}$
因子(A)	$SS_A = \sum_{i=1}^a y_{i*}^2 / b - CT$	$df_A = a - 1$	$MS_A = \frac{SS_A}{df_A}$	$\frac{MS_A}{MS_e}$		
因子(B)	$SS_B = \sum_{h=1}^b y_{*h}^2 / a - CT$	$df_B = b - 1$	$MS_B = \frac{SS_B}{df_B}$	$\frac{MS_B}{MS_e}$		
误差(e)	$SS_e = SS_T - SS_A - SS_B$	$df_e = (a-1) \times (b-1)$	$MS_e = \frac{SS_e}{df_e}$			
总变异	$SS_T = \sum_{i=1}^a \sum_{h=1}^b y_{ih}^2 - CT$	$df_T = a \times b - 1$				

注：CT = $y_{**}^2 / (a \times b)$ 。

两因子无重复观测值试验只适用于两个因子间无交互作用的情形。

若两因子间有交互作用，则每个水平组合中只设一个试验单位（观察单位）的试验设计是不正确或不完善的。因为此种情况下，式中 SS_e 、 df_e 实际上是 A 、 B 两因子交互作用平方和与自由度，所算得的 MS_e 是交互作用均方，主要反映由交互作用引起的变异。这时若仍按情形一所用方法进行方差分析，因误差均方值大（包含交互作用在内），有可能掩盖试验因子的显著性，从而增大犯 II 型（存伪）错误的概率。

因每个水平组合只有一个观测值，故无法估计真正的试验误差，也不可能对因子之间的交互作用进行研究。

进行两因子或多因子试验时，一般应设置重复，以便正确估计试验误差。

进行两因子或多因子试验时，除了研究每一因子对试验指标的影响外，往往更希望研究因子与因子之间的交互作用。

【例 3.15】 化学生产中杂质的出现与压力及温度两个因子有关。为确认压力与温度之间是否存在交互作用，采用无重复的两因子完全随机试验，方案设计及结果如

表 3-37 所示，试加以分析。

表 3-37 方案设计及结果

温度 ($A/^{\circ}\text{C}$)	压力 (B/MPa)	观察值 (y_{ih})
100	25	5
	30	4
	35	6
	40	3
	45	5
125	25	3
	30	1
	35	4
	40	2
	45	3
150	25	1
	30	1
	35	3
	40	1
	45	2

【解】数据求和，如表 3-38 所示。

表 3-38 数据求和

温度 ($A/^{\circ}\text{C}$)	压力 (B/MPa)	观察值 (y_{ih})	y_{i*}	y_{*h}	y_{**}			
100	25	5	23	9	44			
	30	4		6				
	35	6		13				
	40	3		6				
	45	5		10				
125	25	3	13	—		44		
	30	1						
	35	4						
	40	2						
	45	3						
150	25	1	8				—	44
	30	1						
	35	3						
	40	1						
	45	2						

计算各偏差平方和与其自由度

$$CT = y_{**}^2 / (a \times b) = 44^2 / (3 \times 5) = 129.07$$

$$SS_T = \sum_{i=1}^a \sum_{j=1}^b y_{i*}^2 - CT = (5^2 + 4^2 + \cdots + 2^2) - 129.07 = 166 - 129.07 = 36.93$$

$$df_T = a \times b - 1 = 3 \times 5 - 1 = 14$$

$$SS_A = \sum_{i=1}^a y_{i*}^2 / b - CT = (23^2 + 13^2 + 8^2) / 5 - 129.07 = 23.33$$

$$df_A = a - 1 = 3 - 1 = 2$$

$$SS_B = \sum_{h=1}^b y_{*h}^2 / a - CT = (9^2 + 6^2 + 13^2 + 6^2 + 10^2) / 3 - 129.07 = 11.60$$

$$df_B = b - 1 = 5 - 1 = 4$$

$$SS_e = SS_T - SS_A - SS_B = 36.93 - 23.33 - 11.60 = 2.00$$

$$SS_N = \frac{\left[\sum_{i=1}^a \sum_{h=1}^b (y_{ih} \times y_{i*} \times y_{*h}) - y_{**} \times (SS_A + SS_B + C) \right]^2}{a \times b \times SS_A \times SS_B}$$

$$= \frac{[7\,236 - 44 \times (23.33 + 11.60 + 129.07)]^2}{3 \times 5 \times 23.33 \times 11.60} = 0.098\,5$$

其中, $\sum_{i=1}^3 \sum_{h=1}^5 (y_{ih} \times y_{i*} \times y_{*h}) = 5 \times 23 \times 9 + 4 \times 23 \times 6 + \cdots + 2 \times 8 \times 10 = 7\,236$, $df_N = 1$,

$$SS_{e*} = SS_e - SS_N = 2.00 - 0.098\,5 = 1.901\,5, \quad df_{e*} = (a-1)(b-1) - 1 = (3-1) \times (5-1) - 1 = 7.$$

列方差分析如表 3-39 所示。

表 3-39 列方差分析

变异来源	SS	df	MS	F	$F_{0.05}$	$F_{0.01}$
温度(A)	23.333 3	2	11.666 7	42.949 1**	4.71	9.55
压力(B)	11.600 0	4	2.900 0	10.675 9**	4.12	7.85
非可加性	0.098 5	1	0.098 5	0.362 7	5.59	12.25
新误差 e*	1.901 5	7	0.271 6			
总变异	36.933 3	14				

结果表明: 温度与压力之间并不存在交互作用, 而温度与压力对杂质的影响均达到极显著。

特别注意的是, 每单元一个观察值的两因子试验模型, 看上去与单因子随机区组化试验模型颇为相似。然而, 导出这两个模型的试验状态截然不同。每单元一个观察值的两因子试验模型, 即 $a \times b$ 个处理组合按随机顺序实施。单因子随机区组化试验模型, 虽有 b 个区组, 但供试因子 a 种水平中的任一种水平仅仅在区组内实行随机。这两种试验模型的资料收集及其方差分析与结果解释是完全不同的。

【情形二】两因子有重复观测试验资料的方差分析

1) 试验数据模式

两因子完全随机试验中，设 A 因子有 a 种水平， B 因子有 b 种水平，含有 $a \times b$ 种完全组合，即有 $a \times b$ 个处理。每一处理均重复观察 k 次，其试验方案设计及结果、相关计算如表 3-40 所示。

表 3-40 试验方案设计及结果、相关计算

因子 A	因子 B	重复观察值 (y_{ij})						y_{ij*}	y_{i**}	y_{*h*}	$\bar{y}_{ij*}$	$\bar{y}_{i**}$	$\bar{y}_{*h*}$
		1	2	...	j	...	k						
A_1	B_1	y_{111}	y_{112}	...	y_{11j}	...	y_{11k}	y_{11*}	y_{1**}	y_{*1*}	$\bar{y}_{11*}$	$\bar{y}_{1**}$	$\bar{y}_{*1*}$
	B_2	y_{121}	y_{122}	...	y_{12j}	...	y_{12k}	y_{12*}		y_{*2*}	$\bar{y}_{12*}$		$\bar{y}_{*2*}$
	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$		$\vdots$	$\vdots$		$\vdots$
	B_h	y_{1h1}	y_{1h2}	...	y_{1hj}	...	y_{1hk}	y_{1h*}		y_{*h*}	$\bar{y}_{1j*}$		$\bar{y}_{*h*}$
	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$		$\vdots$	$\vdots$		$\vdots$
	B_b	y_{1b1}	y_{1b2}	...	y_{1bj}	...	y_{1bk}	y_{1b*}		y_{*b*}	$\bar{y}_{1b*}$		$\bar{y}_{*b*}$
A_2	B_1	y_{211}	y_{212}	...	y_{21j}	...	y_{21k}	y_{21*}	y_{2**}		$\bar{y}_{21*}$	$\bar{y}_{2**}$	
	B_2	y_{221}	y_{222}	...	y_{22j}	...	y_{22k}	y_{22*}			$\bar{y}_{22*}$		
	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$			$\vdots$		
	B_h	y_{2h1}	y_{2h2}	...	y_{2hj}	...	y_{2hk}	y_{2h*}			$\bar{y}_{2j*}$		
	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$			$\vdots$		
	B_b	y_{2b1}	y_{2b2}	...	y_{2bj}	...	y_{2bk}	y_{2b*}			$\bar{y}_{2b*}$		
$\vdots$	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$	$\vdots$		$\vdots$	$\vdots$	
A_i	B_1	y_{i11}	y_{i12}	...	y_{i1j}	...	y_{i1k}	y_{i1*}	y_{i**}	—	$\bar{y}_{i1*}$	$\bar{y}_{i**}$	—
	B_2	y_{i21}	y_{i22}	...	y_{i2j}	...	y_{i2k}	y_{i2*}			$\bar{y}_{i2*}$		
	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$			$\vdots$		
	B_h	y_{ih1}	y_{ih2}	...	y_{ihj}	...	y_{ihk}	y_{ih*}			$\bar{y}_{ij*}$		
	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$			$\vdots$		
	B_b	y_{ib1}	y_{ib2}	...	y_{ibj}	...	y_{ibk}	y_{ib*}			$\bar{y}_{ib*}$		
$\vdots$	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$	$\vdots$		$\vdots$	$\vdots$	
A_a	B_1	y_{a11}	y_{a12}	...	y_{a1j}	...	y_{a1k}	y_{a1*}	y_{a**}		$\bar{y}_{a1*}$	$\bar{y}_{a**}$	
	B_2	y_{a21}	y_{a22}	...	y_{a2j}	...	y_{a2k}	y_{a2*}			$\bar{y}_{a2*}$		
	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$			$\vdots$		
	B_h	y_{ah1}	y_{ah2}	...	y_{ahj}	...	y_{ahk}	y_{ah*}			$\bar{y}_{aj*}$		
	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$			$\vdots$		
	B_b	y_{ab1}	y_{ab2}	...	y_{abj}	...	y_{abk}	y_{ab*}			$\bar{y}_{ab*}$		
合计									y_{***}		$\bar{y}_{***}$		

任一组合所有重复观察值之和 $y_{ih*} = \sum_{j=1}^k y_{ihj}$ ，相应的平均值 $\bar{y}_{ihj} = \sum_{j=1}^k y_{ihj} / k$ ；因子 A 某一水平所有观察值之和 $y_{i**} = \sum_{h=1}^b \sum_{j=1}^k y_{ihj}$ ，相应的平均值 $\bar{y}_{i**} = \sum_{h=1}^b \sum_{j=1}^k y_{ihj} / (b \times k)$ ；因子 B 某一水平所有观察值之和 $y_{*h*} = \sum_{i=1}^a \sum_{j=1}^k y_{ihj}$ ，相应的平均值 $\bar{y}_{*h*} = \sum_{i=1}^a \sum_{j=1}^k y_{ihj} / (a \times k)$ ；所有观察值之和 $y_{***} = \sum_{i=1}^a \sum_{h=1}^b \sum_{j=1}^k y_{ihj}$ ，相应的平均值 $\bar{y}_{***} = \sum_{i=1}^a \sum_{h=1}^b \sum_{j=1}^k y_{ihj} / (a \times b \times k)$ 。

2) 数学模型

两因子有重复观测试验中，任一试验结果 y_{ijl} 的数学模型可表示为

$$y_{ijl} = \mu + \alpha_i + \beta_h + (\alpha\beta)_{ih} + \varepsilon_{ijl} \quad (i=1,2,\dots,a; h=1,2,\dots,b; j=1,2,\dots,k)$$

μ 为总体均数， α_i 为 A_i 的效应， β_h 为 B_h 的效应， $(\alpha\beta)_{ih}$ 为 A_i 与 β_h 交互效应，且有

$$\alpha_i = \mu_{i*} - \mu$$

$$\beta_h = \mu_{*h} - \mu$$

$$(\alpha\beta)_{ih} = \mu_{ih} - \mu_{i*} - \mu_{*h} + \mu$$

μ_i 、 μ_h 、 μ_{ih} 分别为 A_i 、 B_h 、 $A_i B_h$ 观测值总体均数，且满足

$$\sum_{i=1}^a \alpha_i = 0, \sum_{h=1}^b \beta_h = 0, \sum_{i=1}^a (\alpha\beta)_{ih} = \sum_{h=1}^b (\alpha\beta)_{ih} = \sum_{i=1}^a \sum_{h=1}^b (\alpha\beta)_{ih} = 0$$

ε_{ijl} 为随机误差，相互独立，且都服从 $N(0, \sigma^2)$ 。

3) 列方差分析

列方差分析如表 3-41 所示。

表 3-41 列方差分析

变异来源	SS	df	MS	F	$F_{0.05}$	$F_{0.01}$
因子(A)	$SS_A = \sum_{i=1}^a y_{i**}^2 / (b \times k) - CT$	$df_A = a - 1$	$MS_A = \frac{SS_A}{df_A}$	$\frac{MS_A}{MS_e}$		
因子(B)	$SS_B = \sum_{h=1}^b y_{*h*}^2 / (a \times k) - CT$	$df_B = b - 1$	$MS_B = \frac{SS_B}{df_B}$	$\frac{MS_B}{MS_e}$		
$A \times B$	$SS_{A \times B} = SS_{AB} - SS_A - SS_B$	$df_{A \times B} = (a-1) \times (b-1)$	$MS_{A \times B} = \frac{SS_{A \times B}}{df_{A \times B}}$	$\frac{MS_{A \times B}}{MS_e}$		
误差(e)	$SS_e = SS_T - SS_{AB}$	$df_e = a \times b \times (k-1)$	$MS_e = \frac{SS_e}{df_e}$			
总变异	$SS_T = \sum_{i=1}^a \sum_{h=1}^b \sum_{j=1}^k y_{ijl}^2 - CT$	$df_T = a \times b \times k - 1$				

其中， $CT = y_{***}^2 / (a \times b \times k)$ ； $SS_{AB} = \sum_{i=1}^a \sum_{h=1}^b y_{ih*}^2 / k - CT$ ，代表 AB 水平组合的偏

差平方和，称为“次总(AB)”。

2. 随机区组设计

1) 试验数据模式

两因子随机区组试验中，设 A 因子有 a 种水平， B 因子有 b 种水平，区组 R 有 q 个，含有 $a \times b$ 种完全组合，即有 $a \times b$ 个处理。每一处理随机分配在 k 个区组内，共有 $a \times b \times q$ 个观察值，其试验方案设计及结果、相关计算如表 3-42 所示。

表 3-42 试验方案设计及结果、相关计算

因子 A	因子 B	区组观察值 (y_{ihj})						y_{ij*}	y_{i**}	y_{*h*}	$\bar{y}_{ij*}$	$\bar{y}_{i**}$	$\bar{y}_{*h*}$
		1	2	...	j	...	k						
A_1	B_1	y_{111}	y_{112}	...	y_{11j}	...	y_{11k}	y_{11*}	y_{1**}	y_{*1*}	$\bar{y}_{11*}$	$\bar{y}_{1**}$	$\bar{y}_{*1*}$
	B_2	y_{121}	y_{122}	...	y_{12j}	...	y_{12k}	y_{12*}		y_{*2*}	$\bar{y}_{12*}$		$\bar{y}_{*2*}$
	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$		$\vdots$	$\vdots$		$\vdots$
	B_h	y_{1h1}	y_{1h2}	...	y_{1hj}	...	y_{1hk}	y_{1h*}		y_{*h*}	$\bar{y}_{1j*}$		$\bar{y}_{*h*}$
	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$		$\vdots$	$\vdots$		$\vdots$
	B_b	y_{1b1}	y_{1b2}	...	y_{1bj}	...	y_{1bk}	y_{1b*}		y_{*b*}	$\bar{y}_{1b*}$		$\bar{y}_{*b*}$
A_2	B_1	y_{211}	y_{212}	...	y_{21j}	...	y_{21k}	y_{21*}	y_{2**}		$\bar{y}_{21*}$	$\bar{y}_{2**}$	
	B_2	y_{221}	y_{222}	...	y_{22j}	...	y_{22k}	y_{22*}			$\bar{y}_{22*}$		
	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$			$\vdots$		
	B_h	y_{2h1}	y_{2h2}	...	y_{2hj}	...	y_{2hk}	y_{2h*}			$\bar{y}_{2j*}$		
	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$			$\vdots$		
	B_b	y_{2b1}	y_{2b2}	...	y_{2bj}	...	y_{2bk}	y_{2b*}			$\bar{y}_{2b*}$		
$\vdots$	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$	$\vdots$		$\vdots$	$\vdots$	
A_i	B_1	y_{i11}	y_{i12}	...	y_{i1j}	...	y_{i1k}	y_{i1*}	y_{i**}		$\bar{y}_{i1*}$	$\bar{y}_{i**}$	
	B_2	y_{i21}	y_{i22}	...	y_{i2j}	...	y_{i2k}	y_{i2*}			$\bar{y}_{i2*}$		
	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$			$\vdots$		
	B_h	y_{ih1}	y_{ih2}	...	y_{ihj}	...	y_{ihk}	y_{ih*}			$\bar{y}_{ij*}$		
	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$			$\vdots$		
	B_b	y_{ib1}	y_{ib2}	...	y_{ibj}	...	y_{ibk}	y_{ib*}			$\bar{y}_{ib*}$		
$\vdots$	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$	$\vdots$		$\vdots$	$\vdots$	
A_a	B_1	y_{a11}	y_{a12}	...	y_{a1j}	...	y_{a1k}	y_{a1*}	y_{a**}		$\bar{y}_{a1*}$	$\bar{y}_{a**}$	
	B_2	y_{a21}	y_{a22}	...	y_{a2j}	...	y_{a2k}	y_{a2*}			$\bar{y}_{a2*}$		
	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$			$\vdots$		
	B_h	y_{ah1}	y_{ah2}	...	y_{ahj}	...	y_{ahk}	y_{ah*}			$\bar{y}_{aj*}$		
	$\vdots$	$\vdots$	$\vdots$	...	$\vdots$	...	$\vdots$	$\vdots$			$\vdots$		
	B_b	y_{ab1}	y_{ab2}	...	y_{abj}	...	y_{abk}	y_{ab*}			$\bar{y}_{ab*}$		
合计									y_{***}			$\bar{y}_{***}$	

2) 数学模型

两因子随机区组试验中,任一试验结果 y_{ijl} 的数学模型可表示为

$$y_{ijl} = \mu + \alpha_i + \beta_h + (\alpha\beta)_{ih} + k_i + \varepsilon_{ihj} \quad (i=1,2,\dots,a; h=1,2,\dots,b; j=1,2,\dots,k)$$

μ 为总体均数, α_i 为 A_i 的效应, β_h 为 B_h 的效应, $(\alpha\beta)_{ih}$ 为 A_i 与 β_h 交互效应, q_j 为区组效应, ε_{ihj} 为相互独立的随机误差,且都服从 $N(0, \sigma^2)$ 。

3) 列方差分析

列方差分析如表 3-43 所示。

表 3-43 列方差分析

变异来源	SS	df	MS	F	$F_{0.05}$	$F_{0.01}$
区组	$SS_k = \sum_{j=1}^k y_{**j}^2 / (a \times b) - CT$	$df_k = k - 1$	$MS_k = \frac{SS_k}{df_k}$	$\frac{MS_k}{MS_e}$		
因子(A)	$SS_A = \sum_{i=1}^a y_{i**}^2 / (b \times k) - CT$	$df_A = a - 1$	$MS_A = \frac{SS_A}{df_A}$	$\frac{MS_A}{MS_e}$		
因子(B)	$SS_B = \sum_{h=1}^b y_{*h*}^2 / (a \times k) - CT$	$df_B = b - 1$	$MS_B = \frac{SS_B}{df_B}$	$\frac{MS_B}{MS_e}$		
$A \times B$	$SS_{A \times B} = SS_{AB} - SS_A - SS_B$	$df_{A \times B} = (a-1) \times (b-1)$	$MS_{A \times B} = \frac{SS_{A \times B}}{df_{A \times B}}$	$\frac{MS_{A \times B}}{MS_e}$		
误差(e)	$SS_e = SS_T - SS_k - SS_{AB}$	$df_e = (a \times b - 1) \times (k - 1)$	$MS_e = \frac{SS_e}{df_e}$			
总变异	$SS_T = \sum_{i=1}^a \sum_{h=1}^b \sum_{j=1}^k y_{ihj}^2 - CT$	$df_T = a \times b \times k - 1$				

其中, $CT = y_{****}^2 / (a \times b \times k)$; $SS_{AB} = \sum_{i=1}^a \sum_{h=1}^b y_{ih*}^2 / k - CT$ 。

3.5 回归分析

在生产和科学试验中,测量与数据处理的目的有时并不在于求被测量的估计值,而是为了寻求两个变量或多个变量之间的内在关系。表达变量之间关系的方法有散点图、表格、曲线、数学表达式等,其中数学表达式能较好地反映事物的内在规律性,形式紧凑,且便于从理论上进一步分析研究,对认识自然界量与量之间的关系有着重要意义,而数学表达式可通过回归分析方法得到。

3.5.1 概述

一切客观事物都有其内部的规律性，而且每一事物的运动都与周围其他事物发生相互的联系和影响，事物间的联系和影响反映到数学上，就是变量和变量之间的相互关系。回归分析就是一种研究变量与变量之间关系的数学方法。

科学实践表明，变量之间的关系可以分成两大类，即确定性关系和相关关系。

1. 确定性关系

确定性关系即函数关系，它可以通过反复的精确试验或用严格的数学推导得到，在科学领域中这种关系是大量的。例如，通过具有一定电阻 R 的电路中的电流 I 与加在电路两端的电压 U 之间就遵循欧姆定律，即 $I = U/R$ 。这就是说，对一定的电压值，电流强度就可按前面的公式确定。再如，热力学中的气体状态方程式 $PV = RT$ ，流体力学中的运动方程式，这些都属于这种确定性的函数关系。

2. 相关关系

实际问题中，许多变量之间虽然有非常密切的关系，但是要找出它们之间的确切关系是困难的。造成这种情况的原因极其复杂，影响因素很多，其中包括尚未被发现的或者还不能控制的影响因素，而且各变量的测量总存在测量误差，因此所有这些因素的综合作用就造成了变量之间关系的不确定性。

3.5.2 最小二乘法原理

最小二乘法是一种在数据处理和误差估计等多学科领域得到广泛应用的数学工具，最初应用于天文测量领域。随着现代数学和计算机技术的发展，最小二乘法成为参数估计、数据处理、回归分析和经验公式拟合中必不可少的手段。此外，最小二乘法还是进行组合测量问题数据处理的重要工具。

假设 x 和 y 是具有某种相关关系的物理量，它们之间的关系可用式 (3-1) 给出

$$y = f(x, c_1, c_2, \dots, c_N) \quad (3-1)$$

式中， $c_1, c_2, \dots, c_N$ 是 N 个待定常数，即曲线的函数形式已经确定，而曲线的具体形状是未定的。为求得具体曲线，可同时测定 x 和 y 的数值。设共获得 m 对观测结果

$$(x_1, y_1), (x_2, y_2), \dots, (x_m, y_m) \quad (3-2)$$

接下来的任务就是根据这些观测值确定常数 $c_1, c_2, \dots, c_N$ 。设 x 和 y 关系的最佳形式为

$$\hat{y} = f(x, \hat{c}_1, \hat{c}_2, \dots, \hat{c}_N) \quad (3-3)$$

式中, $\hat{c}_1, \hat{c}_2, \dots, \hat{c}_N$ 是 $c_1, c_2, \dots, c_N$ 的最佳估计值。如若不存在测量误差, 则式 (3-2) 各观测值都应落在式 (3-1) 曲线上, 即

$$y_i = f(x_i, c_1, c_2, \dots, c_N) \quad (i = 1, 2, \dots, m) \quad (3-4)$$

但由于存在测量误差, 因而式 (3-4) 与式 (3-3) 不相重合, 即有

$$e_i = y_i - \bar{y}_i \quad (i = 1, 2, \dots, m) \quad (3-5)$$

通常称 e_i 为残差, 它是误差的实测值。式 (3-3) 中 x 变化时, y 也随之变化。如果 m 对观测值中有比较多的 y 值落到式 (3-3) 曲线上, 则所得曲线就能较为满意地反映被测物理量之间的关系。当 y 值落在曲线上的概率最大时, 式 (3-3) 曲线就是式 (3-1) 曲线的最佳形式。如果误差服从正态分布, 则概率 $P(e_1, e_2, \dots, e_m)$ 为

$$P(e_1, e_2, \dots, e_m) = \frac{1}{\sigma\sqrt{2\pi}} \exp \left[-\sum_{i=1}^m \frac{(y_i - \hat{y}_i)^2}{2\sigma^2} \right] \quad (3-6)$$

当 $P(e_1, e_2, \dots, e_m)$ 最大时, 求得的曲线是最佳形式。显然, 此时式 (3-7) 应最小

$$S = \sum_{i=1}^m (y_i - \hat{y}_i)^2 = \sum_{i=1}^m e_i^2 \quad (3-7)$$

即残差平方和最小, 这就是最小二乘法原理的由来。

这里, 实际上假定了 x_i 无误差, 或者虽然有误差, 但相对于 y_i 的误差来说很小, 可以忽略, 即 $S_x \ll S_y$ 。这一假设使得问题大为简化, 而且是合理的。因为在实际测量中, x 和 y 的测量精度一般是不相同的, 可取其中较为精确的一方作为 x 。严格地说, 最小二乘法仅在误差服从正态分布的情况下才是成立的, 然而在与正态分布相差不太大的误差分布, 以及所有误差都非常小的其他分布中, 也常采用最小二乘法进行处理。

式 (3-7) 可以写成

$$S = \sum_{i=1}^m [y_i - f(x_i, \hat{c}_1, \hat{c}_2, \dots, \hat{c}_N)]^2 \quad (3-8)$$

残差平方和最小, 就应有

$$\frac{\partial S}{\partial c_1} = 0, \quad \frac{\partial S}{\partial c_2} = 0, \quad \dots, \quad \frac{\partial S}{\partial c_N} = 0 \quad (3-9)$$

即要求求解如下方程组

$$\begin{cases} \sum_{i=1}^m [y_i - f(x_i, \hat{c}_1, \hat{c}_2, \dots, \hat{c}_N)] \left(\frac{\partial f}{\partial c_1} \right) = 0 \\ \sum_{i=1}^m [y_i - f(x_i, \hat{c}_1, \hat{c}_2, \dots, \hat{c}_N)] \left(\frac{\partial f}{\partial c_2} \right) = 0 \\ \vdots \\ \sum_{i=1}^m [y_i - f(x_i, \hat{c}_1, \hat{c}_2, \dots, \hat{c}_N)] \left(\frac{\partial f}{\partial c_N} \right) = 0 \end{cases} \quad (3-10)$$

该方程组称为正规方程（normal equation），解该方程组可得未定常数，通常称为最小二乘解。

3.5.3 直线的回归

找出描述变量之间相关关系定量表达式的最直观办法是作图。在回归分析中，最简单也是最基本的情况就是线性回归。对一元线性回归而言，就是配直线的问题，下面通过例题加以分析说明。

【例 3.16】研究腐蚀时间与腐蚀深度两个量之间的关系，可把腐蚀时间作为自变量 x ，把腐蚀深度作为因变量 y ，将试验数据记录在表 3-44 中。求出 x 、 y 之间的线性关系。

表 3-44 试验数据

时间 x/min	3	5	10	20	30	40	50	60	65	90	120
腐蚀深度 $y/\mu\text{m}$	40	60	80	130	160	170	190	250	250	290	460

【解】将表 3-44 中的数据，在直角坐标系中对应地作出一系列点，可得图 3-3，从图中可以直观看出两个变量之间的大致关系。从图中可看出两个变量大致成直线关系，但并不是确定性的线性关系，而是一种相关关系。这种相关关系可表示为

$$\hat{y} = a + bx \quad (3-11)$$

如图 3-3 中所示的方程，称为腐蚀深度与腐蚀时间的回归直线或回归方程。回归直线的斜率 b 称为回归系数，它表示当 x 增加一个单位时， y 平均增加地数量。

上一节指出，式（3-1）的最佳估计值应使其残差平方和最小。直线回归的残差可写为

$$e_i = y_i - (a + bx_i) \quad (3-12)$$

平方和

$$S = \sum_{i=1}^m e_i^2 = \sum_{i=1}^m [y_i - (a + bx_i)]^2 \quad (3-13)$$

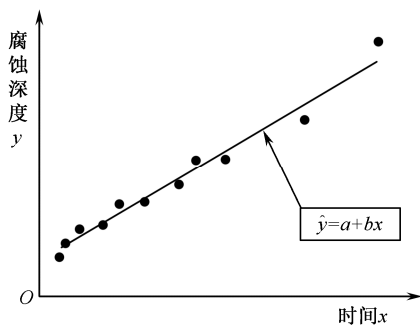


图 3-3 两变量之间的关系

平方和最小，即

$$\begin{cases} \frac{\partial S}{\partial a} = -2 \sum_{i=1}^m [y_i - (a + bx_i)] = 0 \\ \frac{\partial S}{\partial b} = -2 \sum_{i=1}^m x_i (y_i - a + bx_i) = 0 \end{cases} \quad (3-14)$$

因此可得正规方程组

$$\begin{cases} am + b \sum_{i=1}^m x_i = \sum_{i=1}^m y_i \\ a \sum_{i=1}^m x_i + b \sum_{i=1}^m x_i^2 = \sum_{i=1}^m x_i y_i \end{cases} \quad (3-15)$$

令平均值为

$$\begin{cases} \bar{x} = \frac{1}{m} \sum_{i=1}^m x_i \\ \bar{y} = \frac{1}{m} \sum_{i=1}^m y_i \end{cases} \quad (3-16)$$

则由式 (3-11) 可得

$$\begin{cases} a + b\bar{x} = \bar{y} \\ a = \bar{y} - b\bar{x} \end{cases} \quad (3-17)$$

$$(3-18)$$

同样，由式 (3-15) 可得

$$b = \frac{\sum_{i=1}^m x_i y_i - \frac{1}{m} \left(\sum_{i=1}^m x_i \right) \left(\sum_{i=1}^m y_i \right)}{\sum_{i=1}^m x_i^2 - \frac{1}{m} \left(\sum_{i=1}^m x_i \right)^2} = \frac{\sum_{i=1}^m (x_i - \bar{x})(y_i - \bar{y})}{\sum_{i=1}^m (x_i - \bar{x})^2} \quad (3-19)$$

式中

$$\begin{aligned}
 \sum_{i=1}^m (x_i - \bar{x})(y_i - \bar{y}) &= \sum_{i=1}^m x_i y_i - \bar{x} \left(\sum_{i=1}^m y_i \right) - \bar{y} \left(\sum_{i=1}^m x_i \right) + m \bar{x} \bar{y} \\
 &= \sum_{i=1}^m x_i y_i - m \bar{x} \bar{y} \\
 &= \sum_{i=1}^m x_i y_i - \frac{1}{m} \left(\sum_{i=1}^m x_i \right) \left(\sum_{i=1}^m y_i \right)
 \end{aligned} \quad (3-20)$$

由式 (3-18) 和式 (3-19) 可以求得回归直线方程式中的常数 a 及回归系数 b ，这样 $\hat{y} = a + bx$ 便可确定。

把式 (3-18) 代入 $\hat{y} = a + bx$ 中可得 $\hat{y} - \bar{y} = b(x - \bar{x})$ 。因此可知，回归直线通过点 $(\bar{x}, \bar{y})$ 。从力学的观点看， $(\bar{x}, \bar{y})$ 就是 m 各散点 (x_i, y_i) 的重心位置，因而回归直线必须通过这些散点的重心。这一结论对作回归直线是很有用的。

令

$$\begin{cases} l_{xx} = \sum_{i=1}^m (x_i - \bar{x})^2 = \sum_{i=1}^m x_i^2 - \frac{\left(\sum_{i=1}^m x_i \right)^2}{m} \\ l_{yy} = \sum_{i=1}^m (y_i - \bar{y})^2 = \sum_{i=1}^m y_i^2 - \frac{\left(\sum_{i=1}^m y_i \right)^2}{m} \\ l_{xy} = \sum_{i=1}^m (x_i - \bar{x})(y_i - \bar{y}) = \sum_{i=1}^m x_i y_i - \frac{\left(\sum_{i=1}^m x_i \right) \left(\sum_{i=1}^m y_i \right)}{m} \end{cases} \quad (3-21)$$

便可得到回归系数的另一种表达式

$$b = \frac{l_{xy}}{l_{xx}} \quad (3-22)$$

并且，习惯上称 $\sum_{i=1}^m x_i^2$ 为 x 的平方和， $\frac{\left(\sum_{i=1}^m x_i \right)^2}{m}$ 为平方和的修正项， $\sum_{i=1}^m x_i y_i$ 为 x 与 y 的乘积和， $\frac{\left(\sum_{i=1}^m x_i \right) \left(\sum_{i=1}^m y_i \right)}{m}$ 为乘积和的修正项。

上述回归直线的具体计算，通常都是列表进行的，本节的示例，具体计算如表 3-45 所示。完成表 3-45 的计算，就可得到回归直线方程

$$\hat{y} = 3.21x + 45.01 \quad (3-23)$$

表 3-45 具体计算

编号	x	y	x^2	y^2	xy
1	3	40	9	1 600	120
2	5	60	25	3 600	300
3	10	80	100	6 400	800
4	20	130	400	16 900	2 600
5	30	160	900	25 600	4 800
6	40	170	1 600	28 900	6 800
7	50	190	2 500	36 100	9 500
8	60	250	3 600	62 500	15 000
9	65	250	4 225	62 500	16 250
10	90	290	8 100	84 100	26 100
11	120	460	14 400	211 600	55 200
参数值	$\sum x$	$\sum y$	$\sum x^2$	$\sum y^2$	$\sum xy$
	493	2 080	35 859	539 800	137 470
	$\bar{x}$	$\bar{y}$	$\frac{(\sum x)^2}{m}$	$\frac{(\sum y)^2}{m}$	$\frac{(\sum x)(\sum y)}{m}$
	44.818 181 82	189.09	22 095.363 64	393 309	93 221.818 18
$l_{xx}=137\ 64$; $l_{yy}=146\ 491$; $l_{xy}=44\ 248.181\ 82$; $b=\frac{l_{xy}}{l_{xx}}=3.21$; $a=\bar{y}-b\bar{x}=45.01$					

3.5.4 多元回归

前面章节讨论的是只有两个变量的回归问题, 其中一个变量是自变量, 另一个变量是因变量。但在大多数情况下, 自变量不是一个而是多个, 这类问题称为多元回归问题。多元回归中最简单且最基本的是多元线性回归。如自变量 $x_i (i=1, 2, \dots, G)$, 进行 m 次试验所得的数据可以写成两个数组, 即两个矩阵

$$X = \begin{bmatrix} x_{11} & x_{21} & \cdots & x_{G1} \\ x_{12} & x_{22} & \cdots & x_{G2} \\ \vdots & \vdots & \ddots & \vdots \\ x_{1m} & x_{2m} & \cdots & x_{Gm} \end{bmatrix}_{m \times G}, \quad Y = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_m \end{bmatrix}_{m \times 1}$$

显然, 多元线性统计模型是

$$\hat{y} = a_0 + a_1 x_1 + a_2 x_2 + \cdots + a_G x_G \quad (3-24)$$

这种多元线性回归分析的原理, 与一元线性回归分析的原理完全相同, 只是计算上复杂得多。根据最小二乘法, 应令

$$\sum_{j=1}^m (y_j - \hat{y}_j)^2 = \sum_{j=1}^m \left[y_j - (a_0 + a_1 x_{1j} + a_2 x_{2j} + \cdots + a_G x_{Gj}) \right]^2$$

为最小。

在多元线性回归中，回归平方和 U 为

$$U = \sum (\hat{y}_j - \bar{y})^2 \quad (3-25)$$

由式 (3-24) 可知，有 G 个自变量对因变量 y 有影响，所以回归平方和的自由度为

$$f_{\text{回}} = G \quad (3-26)$$

残差平方和 Q 为

$$Q = \sum (y_j - \hat{y}_j)^2 \quad (3-27)$$

自由度为

$$f_{\text{残}} = m - G - 1 \quad (3-28)$$

总平方和 T 为

$$T = \sum (y_j - \bar{y})^2 \quad (3-29)$$

自由度为

$$f_{\text{总}} = m - 1 \quad (3-30)$$

标准误差平方和为残差平方和除以它的自由度，即

$$S^2 = \frac{Q}{m - G - 1} \quad (3-31)$$

标准误差为

$$S = \sqrt{\frac{Q}{m - G - 1}} \quad (3-32)$$

由于线性多元回归的人工计算方法非常烦琐，而且计算量非常大，又容易出错，因此不再介绍。

3.6 试验数据处理常用软件

目前，用于试验设计的统计软件较多，常用的统计软件包括 SPSS (Statistical Package for the Social Science)、SAS (Statistical Analysis System)、DPS (Data Processing System)、MATLAB、Excel 和 Origin 等。为了便于试验设计的更好应用，本节简要地介绍这几种试验设计常用的统计软件，便于读者选用。

3.6.1 SPSS for Windows 软件

SPSS for Windows 是一个适用于自然科学、社会科学和工程技术各领域的统计分析软件包，它集数据录入、整理、分析功能于一身，是世界上流行的统计软件。用户可以根据实际需要和计算机的功能选择模块，以降低对系统硬盘容量的要求，有利于该软件的推广应用。

SPSS for Windows 的分析结果清晰、直观、易学易用，而且可以直接读取 Excel 及 DBF 数据文件，现已推广到多种计算机的操作系统上。在数据的统计分析方面，SPSS for Windows 具有以下特点。

(1) 界面非常友好，除了数据录入及部分命令程序等少数输入工作需要键盘键入外，大多数操作通过“菜单”“按钮”“对话框”来完成。操作简便，易于学习和使用。

(2) SPSS 的命令语句、子命令及选择项的选择绝大部分由“对话框”的操作完成。因此，用户无须花大量时间记忆大量的命令、过程、选择项。

(3) 具有第 4 代语言的特点，只要通过“对话框”的操作告诉系统要做什么，无须告诉怎样做。对于常见的统计方法，只要了解统计分析的原理，无须通晓统计方法的各种算法，即可得到需要的统计分析结果。

(4) 具有完整的数据输入、编辑、统计分析、报表、图形制作等功能。SPSS 提供了从简单的统计描述到复杂的多因素统计分析方法，比如数据的探索性分析、统计描述、列联表分析、二维相关、秩相关、偏相关、方差分析、非参数检验、多元回归、生存分析、协方差分析、判别分析、因子分析、聚类分析、非线性回归、Logistic 回归等。

(5) 该软件分为若干功能模块，用户可以根据自己的统计分析工作需要和计算机的实际配置情况灵活选择。

(6) 与其他软件有数据转换接口，其他软件生成的数据文件均可方便地转换成可供分析的 SPSS 的数据文件。

(7) 该软件针对初学者、熟练者及精通者都比较适用。很多只掌握简单操作分析的群体青睐于 SPSS，而那些熟练或精通者也较喜欢 SPSS，因为他们可以通过编程来实现更强大的功能。

3.6.2 SAS 软件

SAS (Statistical Analysis System) 软件于 20 世纪 70 年代由美国 SAS 研究所开发，经历了许多版本。经过多年的完善和发展，SAS 系统在国际上已被誉为统计分析

的标准软件，在各个领域得到广泛应用。

SAS 是一个模块化、集成化的大型应用软件系统。它由数十个专用模块构成，功能包括数据访问、数据储存及管理、应用开发、图形处理、数据分析、报告编制、运筹学方法、计量经济学与预测等。SAS 系统基本上可以分为四大部分：SAS 数据库部分，SAS 分析核心，SAS 开发呈现工具，SAS 对分布处理模式的支持及其数据仓库设计。SAS 系统主要完成以数据为中心的四大任务：数据访问、数据管理（SAS 的数据管理功能并不是很出色，但是数据分析能力强大。所以常常用微软的产品管理数据，再导成 SAS 数据格式，因此要注意与其他软件的配套使用）、数据呈现、数据分析。Base SAS 模块是 SAS 系统的核心，其他各模块均在 Base SAS 提供的环境中运行。用户可选择需要的模块与 Base SAS 一起构成一个用户化的 SAS 系统。

SAS 把数据存取、管理、分析和展现有机地融为一体，主要特点如下。

（1）SAS 以一个通用的数据（DATA）步产生数据集，而后以不同的过程调用完成各种数据分析。其编程语句简洁、短小，通常只需很小的几句语句即可完成一些复杂的运算，得到满意的结果，因此使用简便，操作灵活。

（2）SAS 提供了从基本统计数的计算到各种试验设计的方差分析，以及相关回归分析和多变数分析的多种统计分析过程，几乎囊括了所有最新分析方法，不仅功能十分强大，而且分析技术先进可靠。分析方法的实现通过过程调用完成。许多过程同时提供了多种算法和选项。例如回归分析中自变量选择，提供了包括 STEPWISE、BACKWARD、FORWARD、RSQUARE 等多种方法；方差分析中的多重比较，提供了包括 LSD、DUNCAN、TUKEY 等多种方法。

（3）结果输出以简明的英文给出提示，统计术语规范易懂，用户只需要具有初步英语和统计基础即可。用户只要告诉 SAS “做什么”，而不需要告诉其“怎么做”。SAS 的设计，使得任何 SAS 能够“猜”出的东西用户都不必告诉它（无须设定），并且能自动修正一些小的错误（例如将 DATA 语句的 DATA 拼写成 DATE，SAS 将假设为 DATA 继续运行，仅在 LOG 中给出注释说明）。

（4）用户按下功能键 F1，可提供联机帮助，随时获得帮助信息，得到简明的操作指导。

SAS 已在全球拥有近 3 万个客户群，直接用户超过 300 万人。SAS 已被广泛应用于经济管理、社会科学、生物医学、质量控制、试验设计等多个领域，并且发挥着越来越重要的作用。

3.6.3 Excel 软件

Excel 是第一款允许用户自定义界面的电子制表软件（包括字体、文字属性和单元格格式）。它还引进了“智能重算”的功能，当单元格数据变动时，只有与之相关的数据才会更新，而原先的制表软件只能重算全部数据或者等待下一个指令。同时，Excel 还有强大的图形功能，具有方便性和普遍性，很容易被用户掌握和使用。

建立一个新的 Excel 文件之后，便可以进行数据的输入操作，建立试验数据表格，这是 Excel 处理试验数据的基础。数据输入的方法很简单，只需要单击需要输入数据的单元格，使之成为活动单元格，然后从键盘上输入数据即可。

Excel 中的数据按类型有很多种，如数值型、字符型和逻辑型等。在输入数据时，需要注意不同类型数据的输入方法。输入有规律的数值或文本时，可以采用特殊的方法来完成。如果需要在相邻几个单元格中输入相同的数值或文本时，不必一个一个地输入，可以采用自动填充方法输入数据或采用数组输入方法输入数据。

Excel 提供了完整的算术运算符，如+（加）、-（减）、*（乘）、/（除）、%（百分比）、^（指数）等，以及丰富的内置函数（公式），如 SUM（求和）、AVERAGE（求算术平均值）、STDEV（求样本标准差）等，从而可以根据数据处理需要，建立各种公式，对数据执行计算操作，生成所需要的数据。

Microsoft Excel 提供了多种非常实用的数据分析工具，使用这些工具时，只需为每一个分析工具提供必要的数据和参数，该工具就会使用适宜的统计或工程函数，在输出表格中显示相应的结果，极大地简化运算过程。

3.6.4 DPS 软件

DPS 数据处理系统，英文全称为 Data Processing System。该系统采用多级下拉式菜单，用户使用时整个屏幕犹如一张工作平台，可被随意调整，因而形象地称其为 DPS 数据处理工作平台，简称 DPS 平台。与上文提及的三种软件相比，其优势如下。

（1）DPS 是当前国内唯一一款试验设计及统计分析功能齐全、价格适合于国内用户、具自主知识产权、技术上达到国际先进水平的国产多功能统计分析软件包，并且在某些方面已处于国际领先地位（如试验设计中大样本时的均匀试验设计、多元统计分析中的动态聚类分析）。

（2）该软件具有包括均匀设计、混料均匀设计在内的丰富试验设计功能，并在均匀设计中采用了独创算法，实现了大型均匀设计表构造的重大突破，且混料均匀设计可适合任意约束条件的情形。

(3) DPS 是目前国内统计分析功能最全的软件包,其完善的统计分析功能涵盖了所有统计分析内容。DPS 的一般线性模型 GLM 可以处理各种类型的试验设计方差分析,特别是一些用 SPSS 菜单操作解决不了、用 SAS 编程很难处理的多因素裂区混杂设计、格子设计等方差分析问题,用 DPS 菜单操作可轻松搞定。目前版本应用 GLM 进行分析,可处理因子数已不受限制。

(4) DPS 独特的非线性回归建模技术实现了“可想即可得”的用户需求,参数拟合精度高。

(5) 从 20 世纪 90 年代开始, DPS 一直在丰富专业统计分析模块,不断地完成了随机前沿面模型、数据包络分析、顾客满意指数模型、数学生态、生物测定、地理统计、遗传育种、生存分析、水文频率分析、量表分析、质量控制图、ROC 曲线分析等内容,并且还在不断地扩充。

(6) DPS 不仅实现了 SPSS 高级统计分析的计算,还具备 Excel 那样方便地在工作表里处理基础统计分析的功能。它提供了十分方便的可视化操作界面、可借助图形处理的数据建模功能,为处理复杂模型提供了最直观的途径,而这些功能是同类软件中所欠缺的。

(7) DPS 统计软件根据用户要求,不断吸纳新的统计方法,如一般线性模型的多元方差分析、小波分析、偏最小二乘回归、投影寻踪回归、投影寻踪综合评价、灰色系统方法、混合分布参数估计、含定性变量的多元逐步回归分析、三角模糊数分析、M 估计、随机前沿面模型及面板数据统计分析等,并在不断探索,吸纳更多功能,使系统更加完善。

3.6.5 MATLAB 软件

MATLAB 是 Matrix Laboratory 的简称,是美国 Mathwork 公司出品的商业数学软件,用于算法开发、数据可视化、数据分析及数值计算的高级技术计算语言和交互式环境。它将数值分析、矩阵计算、科学数据可视化及非线性动态系统的建模和仿真等诸多强大功能集成在一个易于使用的视窗环境中,为科学研究、工程设计及必须进行有效数值计算的众多科学领域提供了一种全面的解决方案,并在很大程度上摆脱了传统非交互式程序设计语言(如 C、Fortran)的编辑模式,代表了当今国际科学计算软件的先进水平。其优势和特点如下。

(1) MATLAB 用户界面非常友好,其接近数学表达式的自然化语言使用户易于学习和掌握。

(2) 该软件不仅具有高效的数值计算及符号计算功能,能使用户从繁杂的数学运

算分析中解脱出来,还有完备的图形处理功能,可实现计算结果和编程的可视化,极大地方便了用户的使用。

(3) 此外 MATLAB 还具备功能丰富的应用工具箱(如信号处理工具箱、通信工具箱等),为用户提供了大量方便实用的处理工具。

3.6.6 Origin 软件

Origin 是由 OriginLab 公司开发的专业函数绘图软件,由于其简单易学、操作方便、功能强大,很快就成为国际流行的分析软件之一。使用 Origin 就像使用 Excel 那样简单,只需点击鼠标,选择菜单命令就可以完成大部分工作,获得满意的结果。像 Excel 一样,Origin 是个多文档界面应用程序,它将所有工作都保存在 Project(*.OPJ)文件中。该文件可以包含多个子窗口,如 Worksheet、Graph、Matrix、Excel 等。各子窗口之间是相互关联的,可以实现数据的即时更新。子窗口可以随 Project 文件一起存盘,也可以单独存盘,以便其他程序调用。

Origin 的功能主要有数据分析和绘图。其数据分析主要包括统计、信号处理、图像处理、峰值分析和曲线拟合等各种完善的数学分析功能。准备好数据后,进行数据分析时,只需选择所要分析的数据,然后再选择相应的菜单命令即可。绘图是基于模板的,软件本身提供了几十种二维和三维绘图模板而且允许用户自己定制模板。绘图时,只要选择所需要的模板即可。用户可以自定义数学函数、图形样式和绘图模板,也可以和各种数据库软件、办公软件、图像处理软件等方便地连接。Origin 可以导入包括 ASCII、Excel、pClamp 在内的多种数据。除此之外,它可以把 Origin 图形输出到多种格式的图像文件,比如 JPEG、GIF、EPS、TIFF。该软件也支持编程,以方便拓展功能和执行批处理任务,其编程语言有两种:LabTalk 和 Origin C。用户可以通过编写 X-Function 来建立自己需要的特殊工具。X-Function 可以调用 Origin C 和 NAG 函数,而且可以很容易地生成交互界面。用户可以定制自己的菜单和命令按钮,把 X-Function 放到菜单和工具栏上,以后就可以非常方便地使用自己的定制工具。

3.7 试验报告的编写

试验报告是把试验的目的、方法、过程、结果等记录下来,经过整理,写成的书

面汇报。科技试验报告是描述、记录某个科研课题过程和结果的一种科技应用文体。撰写试验报告是科技试验工作不可缺少的重要环节。虽然试验报告与科技论文一样都以文字形式阐明了科学研究的成果,但二者在内容和表达方式上仍有所差别。科技论文一般是把成功的试验结果作为论证科学观点的根据。试验报告则客观地记录试验的过程和结果,着重告知一项科学事实,不带试验者的主观看法。

3.7.1 试验报告的分类

按科目分类:因科学试验的对象而异,如化学试验的报告叫化学试验报告,物理试验的报告就叫物理试验报告。随着科学事业的日益发展,试验的种类、项目等日见繁多,但其格式大同小异,比较固定。试验报告必须在科学试验的基础上进行。它主要的用途在于帮助试验者不断地积累研究资料,总结研究成果。

按专业分类:① 型式试验报告;② 拉伸试验报告;③ 盐雾试验报告;④ 土工试验报告;⑤ 电气试验报告;⑥ 水压试验报告;⑦ 变压器试验报告;⑧ 拉拔试验报告;⑨ 动力触探试验报告;⑩ 击实试验报告。

3.7.2 试验报告的特点

正确性:试验报告的写作对象是科学试验的客观事实,因此需要内容科学,表述真实、质朴,判断恰当。

客观性:试验报告以客观的科学研究的 facts 为写作对象,它是对科学试验的过程和结果的真实记录,虽然也要表明对某些问题的观点和意见,但这些观点和意见都是在客观事实的基础上提出的。

确证性:确证性是指试验报告中记载的试验结果能被任何人所重复和证实,也就是说,任何人按给定的条件去重复这项试验,都能观察到相同的科学现象,得到同样的结果。

可读性:可读性是指为使读者了解复杂的试验过程,试验报告的写作除了以文字叙述和说明以外,还常常借助画图像、列表格、作曲线图等方式,说明试验的基本原理和各步骤之间的关系,解释试验结果等。

3.7.3 试验报告的结构

试验报告的书写是一项重要的基本技能。它不仅是对每次试验的总结,更重要的

是它可以初步培养和训练学生的逻辑归纳能力、综合分析能力和文字表达能力,是科学论文写作的基础。因此,参加试验的每位学生,均应及时认真地书写试验报告。要求内容实事求是,分析全面具体,文字简练通顺,誊写清楚整洁。

试验报告内容与格式如上所述。

试验名称

要用最简练的语言反映试验的内容。如验证某程序、定律、算法,可写成“验证×××”“分析×××”。

学生姓名、学号及合作者

试验日期和地点(年、月、日)

试验目的

目的要明确,在理论上验证定理、公式、算法,并使试验者获得深刻和系统的理解。在实践上,掌握使用试验设备的技能技巧和程序的调试方法。一般需说明是验证型试验还是设计型试验,是创新型试验还是综合型试验。

试验设备(环境)及要求

主要包括在试验中需要用到的试验用物,药品及对环境的要求。

试验原理

在此阐述试验相关的主要原理。

试验内容

这是试验报告中极其重要的内容。要抓住重点,可以从理论和实践两个方面考虑。这部分要写明依据何种原理、定律算法、操作方法进行试验。应给出详细的计算过程。

试验步骤

只写主要操作步骤,不要照抄实习指导,要简明扼要。还应该画出试验流程图(试验装置的结构示意图),再配以相应的文字说明,这样既可以节省许多文字说明,又能使试验报告简明扼要。

试验结果

应包括试验现象的描述,试验数据的处理等。原始资料应附在本次试验主要操作者的试验报告上,同组的合作者要复制原始资料。

对于试验结果的表述,一般有以下三种方法。

(1) 文字叙述:根据试验目的将原始资料系统化、条理化,用准确的专业术语客观地描述试验现象和结果,要有时间顺序及各项指标在时间上的关系。

(2) 图表:用表格或坐标图的方式使试验结果突出、清晰,便于相互比较,尤其适合分组较多,且各组观察指标一致的试验,使组间异同一目了然。每一图表应有表

目和计量单位，应说明一定的中心问题。

(3) 曲线图：应用记录仪器标记出的曲线图，因为这些指标的变化趋势形象生动、直观明了。

在试验报告中，可任选其中一种或几种方法并用，以获得最佳效果。

讨论

根据相关的理论知识对所得的试验结果进行解释和分析。如果所得的试验结果和预期的结果一致，那么它可以验证什么理论？试验结果有什么意义？说明了什么问题？这些是试验报告应该讨论的。但是，不能用已知的理论或生活经验硬套在试验结果上；更不能由于所得到的试验结果与预期的结果或理论不符而随意取舍甚至修改试验结果，这时应该分析其异常的可能原因。如果本次试验失败了，应找出失败的原因及以后试验应注意的事项。不要简单地复述课本上的理论而缺乏自己主动思考的内容。

另外，也可以写一些本次试验的心得或者提出一些问题与建议等。

结论

结论不是具体试验结果的再次罗列，也不是对今后研究的展望，而是针对这一试验所能验证的概念、原则或理论的简明总结，是从试验结果中归纳出的一般性、概括性的判断，要简练、准确、严谨、客观。

本章习题

3-1 为了考察同学们对某一试验操作掌握的熟练程度，实验室老师们观察每位同学完成试验操作的情况，随机记录了 10 位同学完成试验操作的时间，测得平均时间为 15 分钟，样本标准差为 3.2 分钟，假定同学们完成该试验操作的时间服从正态分布，则

(1) 同学们完成该试验操作平均时间的置信度为 95% 的区间估计是什么？

(2) 若样本容量为 30，而数据的样本均值和样本标准差不变，则置信读为 95% 的置信区间是什么？

3-2 某新能源器件企业组装部门一次抽样调查表明，工人们平均完成组装一套新能源器件时间的 95% 置信区间为 [1.5, 2.6] 小时，问该次抽样样本平均完成时间 $\bar{x}$ 是多少？若样本容量为 100，则样本标准差是多少？

3-3 已知某高岭土公司生产的高岭土中含铁量服从正态分布 $N(1.07, 0.0272)$ ，现在抽检了 9 批次产品，其平均含铁量为 1.013。如果含铁量的方差没有变化，可否认为现在产品的平均含铁量仍为 1.07 ($\alpha=0.05$)？

3-4 试表 3-46 所列的试验数据，画出散点图，并求某物质在溶液中的浓度 $c(\%)$ 与其沸点温度 T 之间的函数关系，并检验所建立的函数是否有意义 ($\alpha=0.05$)。

表 3-46 习题 3-4 试验数据

$c/\%$	19.6	20.5	22.3	25.1	26.3	27.8	29.1
$T/^\circ\text{C}$	105.4	106.0	107.2	108.9	109.6	110.7	111.5

3-5 某材料生产企业采用新旧两种技术生产某种金属材料，分别抽检新旧两种技术生产的产品，测定产品中杂质的含量 ($\%$)，结果如表 3-47 所示。

表 3-47 习题 3-5 试验数据

旧技术/ $\%$	3.51	3.02	3.73	3.64	3.10	3.35	3.83	3.07	3.79
新技术/ $\%$	3.21	3.15	3.29	3.31	3.27	3.19	3.35	3.23	3.34

试分析新技术是否比旧技术生产更稳定，并检验两种技术之间是否存在系统误差。

3-6 在某提钒工艺中，选择酸浸时间、酸浸次数和加酸量三个因素进行考察，以样品中钒含量作为试验指标，试验数据列于表 3-48 中。试对试验数据进行线性回归分析，并检验线性回归方程的显著性，确定因素主次顺序 ($\alpha=0.05$)。

表 3-48 习题 3-6 试验数据

试 验 号	酸浸时间/min	酸浸次数	加酸量/倍	钒含量/ $\text{mg}\cdot\text{L}^{-1}$
1	30	1	2.5	180
2	50	2	3.5	444
3	70	3	2	552
4	90	1	3.5	312
5	110	2	2	408
6	130	3	3	684
7	150	3	4	687