



中华人民共和国国家标准

GB/T 3358.3—2009/ISO 3534-3:1999
代替 GB/T 3358.3—1993

统计学词汇及符号 第3部分:实验设计

Statistics—Vocabulary and symbols—
Part 3: Design of experiments

(ISO 3534-3:1999, IDT)

2009-10-15 发布

2010-02-01 实施

中华人民共和国国家质量监督检验检疫总局
中国国家标准化管理委员会 发布

目 次

前言 I

引言 II

范围..... 1

1 一般术语 1

2 实验安排术语 7

3 分析方法术语..... 19

参考文献 24

索引 25

汉语拼音索引 25

英文对应词索引 27

前 言

GB/T 3358《统计学词汇及符号》分为以下部分：

- 第1部分：一般统计术语与用于概率的术语
- 第2部分：应用统计
- 第3部分：实验设计

本部分为 GB/T 3358 的第3部分，等同采用 ISO 3534-3:1999《统计学 词汇及符号 第3部分：实验设计》。

GB/T 3358 的本部分与 ISO 3534-3:1999 相比，订正了原文的错误，修正了原文中概念表述不够准确的部分，主要变化如下：

- 将 1.10“处理(treatment)”的定义改为“每个因子的特定水平或不同因子水平的组合”；
- 将 1.27“重复(replication)”的定义改为“对给定的处理实施多于一次的实验”，并删去了注；
- 在 1.21“残差(residual)”的定义中在“(响应变量的)预测值”前增加“基于假定模型的”一词，以与 1.22“剩余误差(residual error)”的定义相对应；
- 在 2.1.2.1“2^k 析因实验”中增加了关于用“1”，“2”分别替代“+”“—”，表示因子两个水平的注；
- 将 2.3“区组设计(block design)”的定义改为“将全部实验单元分成若干个区组的实验设计”等。

与 ISO 3534-3:1999 相比，本部分作了必要的编辑性的修改，例如：

- 为与第1部分和第2部分相一致，在术语定义中增加了其他术语(包括条目编号)的引用；
- 对 ISO 3534-3:1999 引用 ISO 3534-1:1993 与 ISO 3534-2:1993 的条目，按等同采用 ISO 3534-1:2006 与 ISO 3534-2:2006 的 GB/T 3358.1—2009 与 GB/T 3358.2—2009 的相应内容，作适当的更改。

本部分代替 GB/T 3358.3—1993《统计学术语 第三部分 试验设计术语》，与 GB/T 3358.3—1993 相比，主要变化如下：

- 名称改为《统计学词汇及符号 第3部分：实验设计》；
- 在全文中用“实验”替代“试验”，作为相应英文词“experiment”的优先选用词，但将“试验”一词保留为“实验”的同义词；
- 调整了术语条目设置；
- 增加了大量的示例及注释。

本部分由全国统计方法应用标准化技术委员会提出并归口。

本部分主要起草单位：中国科学院数学与系统科学研究院、北京大学、中国标准化研究院、苏州大学。

本部分主要起草人：冯士雍、陈敏、石坚、艾明要、丁文兴、汪仁官、于振凡。

本部分于 1993 年首次发布，本次为第一次修订。

引言

实验设计实质上即是对实验的策划,以便有效和经济地得到正确和相关的结论。选择具体的实验方案依赖于所涉及问题的类型、结论的普遍性程度,以及可利用的资源(实验材料、人员与时间)。一个经过恰当设计和实施的实验,常导致相对简明的统计分析和对结果的解释。

近年来,实验设计的应用得到蓬勃发展,主要是由于认识到实验设计对于提高产品和服务的质量非常重要。虽然统计质量控制、管理目标(menagement resolve)、检验和其他质量工具也有此功能,实验设计代表了一种在复杂的、变化的和交互的环境中进行选择的方法。在历史上,实验设计在农业领域得到发展与繁荣,医学领域也经历了悠久历史的精心实验设计。目前,工业环境中目睹了实验设计带来的可观效益,因为实验设计便于开展工作(界面友好的软件),改进了培训,得到有影响力的倡导推广,也积累了许多成功的案例。

析因实验(见 2.1)为实验者提供了研究所关心的多因子之间相互关系的方法。这些类型的实验,远比简单的一次仅分析一个因子的实验更为有效。析因实验特别适用于确定在其他因子取不同水平时有不同响应的因子。通常,分析质量的“突破”来自研究交互作用(见 1.17)所揭示的内在联系。如果考虑的因子数量比较多,析因实验可能要占用过多资源而难以实施。不过,部分析因设计(见 2.1.1)提供了一种可能的折衷办法。实际上,如果最初的目标是找出哪些因素需要进一步研究,筛选设计(见 2.2)就比较可行。

在计划一个实验中,有必要对由于实验条件或实验单元处理的配置造成的偏倚进行控制。随机化(见 1.29)和分区组(见 1.28)之类的技术即是用来最小化讨厌的、外来的因素的影响。具体分区组技术包括随机化区组设计(见 2.3.1)、拉丁方设计(见 2.3.2)、平衡不完全区组设计(见 2.3.4.1)等。

实验设计是一个渐进的过程,以不断完善为目标,响应曲面设计(见 2.4)扮演了举足轻重的角色。通过对关键因子的不同水平的考察,响应曲面设计方法巧妙地解释了最优点附近的曲线效应。

混料设计(2.5)处理各因子在整体中比例的情况,例如合金中的成分。嵌套设计(见 2.6)尤其适用于多个实验室间进行的实验。

如果实验完全按方案实施,对实验数据的分析方法将是直接的。图方法(见 3.1)对揭示大体结论尤为有效。根据模型进行参数估计(见 1.1 及其后)常使用回归分析(见 3.3)。回归分析方法也可用来处理缺失数据,识别离群值,以及其他问题所带来的困难。

优良的实验设计应该:

- a) 结合在因子及其水平的选择、描述假定条件等方面的先验知识和经验;
- b) 以最少的精力处理相关信息;
- c) 实验前能确保该实验的设计可以实现实验的目标及所需的精度;
- d) 体现调查的连续性;
- e) 明确实验处理的安排及其次序,以避免实验过程中的误解。

统计学词汇及符号

第3部分：实验设计

范围

GB/T 3358 的本部分规定了实验设计领域和起草其他标准中常用的术语。

1 一般术语

1.1

模型 model

关于响应变量(1.2)与预测变量(1.3)关系及其附带假定的描述。

注1：模型由三个部分组成：第一部分是建模的响应(1.2)，第二部分是包含预测变量(1.3)的模型的不确定性或系统性的部分，第三部分是模型的随机部分或随机误差，其描述可以十分详细。例如，误差项可以结合成散度效应(1.14)，使响应值的变异随着响应值的增大而增加。

示例1：一个零件的寿命与它所处的环境条件有关。

示例2：一个典型模型：

$$y_{ij} = \mu + \alpha_i + \beta_j + \epsilon_{ij}$$

其中 y_{ij} 是在因子 A 的 i 水平和因子 B 的 j 水平时的响应， μ 是响应的总平均， α_i 是因子 A 在 i 水平时的附加效应， β_j 是因子 B 在 j 水平时的附加效应， ϵ_{ij} 是误差项。

模型的响应部分仅是 y_{ij} ；模型的预测部分是 $\mu + \alpha_i + \beta_j$ ，由一项响应的总平均和两项因子效应组成；模型的随机或误差部分是 ϵ_{ij} ，它表示产生该响应的过程的固有变异。

示例3：一个常用模型：

$$y_{ijk} = \alpha_i + \beta_j + \tau_{ij} + \epsilon_{ijk}$$

其中 y_{ijk} 是第 k 次重复(1.27)的响应， α_i 是由因子 1 造成的调整， β_j 是由因子 2 造成的调整， τ_{ij} 是由两因子的交互效应(1.16)所造成的调整， ϵ_{ijk} 是误差项。

此处的典型模型中不包含总平均项，所以用术语“调整”代替示例2中的“附加效应”；此外，采用 $y_{ijk}(\epsilon_{ijk})$ 而不是 $y_{ij}(\epsilon_{ij})$ ，表明可能存在的重复。

示例4：另一个典型模型：

$$y_i = e^{\beta_0 + \beta_1 x_i + \beta_2 x_i^2} + \epsilon_i$$

其中 y_i 是对应于 x_i 的响应， $e^{\beta_0 + \beta_1 x_i + \beta_2 x_i^2}$ 表示对应于 x_i 的平均响应， ϵ_i 是误差项。

注2：模型的上述描述不仅适用于带有可加误差的经典线性模型，也适用于误差可以用各种分布描述的广义线性模型，这些分布包括二项分布、泊松分布、指数分布、伽玛分布和正态分布。

1.2

响应变量 response variable

表示实验结果的变量。

注1：同义词是“输出变量”。

注2：不推荐把术语“因变量”也作为一个同义词，避免可能与“自变量”混淆(见 GB/T 3358.1—2009 的 1.11)。

注3：从每一个实验单元(1.9)记录多个响应时，响应变量可为向量。

1.3

预测变量 predictor variable

可用来解释实验结果的变量。

注1：常用同义词有：“输入变量”、“描述变量”和“解释变量”。

注2：在一个经设计的实验中，预测变量的可控程度表明了它在其中所起的潜在作用。预测变量可以是可控的(固定的)、可修正的(仅在短期内或在花费昂贵代价后可控的)或不可控的(随机的)。

注 3: 预测变量可以包含一个随机元素,也可以是没有随机误差的可观测或设置的一些定性组。

注 4: 不推荐把术语“自变量”作为一个同义词,避免可能与“因变量”混淆(见 GB/T 3358.1—2009,1.11)。

1.4

设计区域 design region

设计空间 design space

预测变量(1.3)容许值的集合。

1.5

因子 factor

为评估对响应变量(1.2)的效应而在实验中需变化的预测变量(1.3)。

注 1: 因子对实验的结果可能提供一种可解释的原因。

注 2: 与“预测变量(1.3)”的使用更为一般相比,此处“因子”的使用则有特定含义。

注 3: 因子可与“区组(1.11)”的建立有关。

1.6

水平 level

因子(1.5)可能的设置、取值或安排。

注 1: 同义词是预测变量的取值。

注 2: 术语“水平”形式上与定量特性有关,但它也用来描述定性特征的状态或设置。

示例: 表示催化剂的水平可能是“有”或“没有”;表示加热处理的有序尺度温度的四个水平可能是:100 °C、120 °C、140 °C 和 160 °C;表示实验室的名义尺度变量可能为对应于不同设备的三种水平:A,B 和 C。

注 3: 在因子不同水平观测的响应提供了实验水平范围内决定因子效应的信息。如果没有模型关系假定的可靠根据,在水平范围以外作外推通常是不合适的。在水平范围以内的插值可能依赖于水平个数和水平间距。内插通常认为是合理的,尽管在不连续或多峰的情况下会造成实验范围内的突然变化。水平的选取可能仅限于某些特定的固定值(已知或未知),也可以在研究范围内随机选择。

1.7

实验误差 experimental error

实验响应变量(1.2)在实施中,除因子(1.5)、区组(1.11)或其他可归属原因外,不能解释的变异。

注 1: 实验误差是实验的普遍特征。当实验重复时,尽管非常精细地控制了实验材料、环境条件和实验操作,每次实验结果仍会变化。因此实验误差是普遍存在的。实验误差给从实验结果得出的结论带来了一定程度的不确定性,从而在下结论时应予以考虑。

注 2: 将比本术语较为广泛的概念用于某单个响应变量时,术语“残差(1.21)”、“残余误差(1.22)”和“纯误差(1.23)”提供了更为精细的描述。

注 3: 与实验误差有关的术语有“重复性标准差”(GB/T 3358.2—2009,3.3.7)和“再现性标准差”(GB/T 3358.2—2009,3.3.12)。如果实验的实际设计符合重复性条件(GB/T 3358.2—2009,3.3.6)或再现性条件(GB/T 3358.2—2009,3.3.11),这些术语可直接用于实验设计中。

1.8

方差分量 variance component

描述因子(1.5)效应或实验误差(1.7)的随机变量的方差。

注 1: 考虑模型

$$y_{ij} = \mu + \tau_i + \epsilon_{ij}$$

其中 τ_i 是从可能取值的无限集中随机选择的一个水平(的效应), τ_i 和 ϵ_{ij} 的分布独立, τ_i 和 ϵ_{ij} 都是随机变量。一旦从可能水平的无限集中作出了随机选择,就根据 τ_i 的实现继续进行分析。鉴于这种概率结构,考虑以下方差等式是合理的: $Var(y_{ij}) = Var(\tau_i) + Var(\epsilon_{ij})$ 。右端也可表示成 $\sigma_\tau^2 + \sigma_\epsilon^2$,其中 σ_τ^2 和 σ_ϵ^2 是 y_{ij} 的方差分量。

1.9

实验单元 experimental unit

接收一个特定处理(1.10)、随后产生响应变量(1.2)值的个体。

注: 一个实验单元对应于一次实验。

1.10

处理 treatment

每个因子(1.5)的特定水平(1.6)或不同因子水平的组合。

1.11

区组 block

比实验单元(1.9)的全体更具齐性的实验单元的集合(见1.28)。

注1:术语“区组”来源于农业实验。作实验的田地被分为若干块每块有相同条件的小区域,如有相同或相似的风的位向、地下水状况及可耕层的深度等。其他情况下,区组可根据原材料的批、操作员、一天内考察的单元数等来确定。

注2:通常,识别出区组的存在会影响处理(1.10)安排到实验单元的方式。

1.12

单因子实验 one-factor experiment

只考察一个因子(1.5)对响应变量(1.2)的效应(如有的话)的实验。

示例:考虑模型:

$$y = \mu_i + \epsilon$$

其中 y 是响应变量, μ_i 是在因子的 i 水平时的平均响应, ϵ 是一个包涵所有其他效应和变异来源的随机变量。该模型将响应变量 y 与效应 μ_i (依赖于因子水平)和误差项 ϵ 联系起来。 μ_i 中的变化反映了因子对响应变量的影响(此时平均响应值是因子水平的函数)。

该模型的另一种表示形式是:

$$y = \mu + \alpha_i + \epsilon$$

其中 y 是响应变量, μ 是总平均响应, α_i 是因子的 i 水平所造成的附加效应, ϵ 是一个包涵所有其他效应和变异来源的随机变量。

1.13

主效应 main effect

单个因子(1.5)对响应变量(1.2)均值的影响。

注:对于两水平(1.6)因子,主效应描述了从一个水平到另一个水平时响应的变化。如果水平设置为-1(低水平)和+1(高水平),那么因子的主效应可估计成因子水平为+1时的平均响应减去因子水平为-1时的平均响应。

考虑模型:

$$y = \mu + \beta X + \epsilon$$

其中 y , μ 和 ϵ 与1.12中定义一致, X 是前述的+1或-1, β 表示因子 X 的调整。注意 β 的估计是因子 X 主效应(1.13)的一半。如果 $\beta=0$, 那么 X 不影响响应变量的均值(不管 X 的水平是+1或-1都一样), 从而 X 的主效应为0。

1.14

散度效应 dispersion effect

单个因子(1.5)对响应变量(1.2)方差的影响。

注:认识到对平均响应没有太大影响,但可能对响应的变异有显著效应的因子是很重要的。在此情形,使响应变异较小、数值比较一致的因子的某个特定水平(1.6)正是最想要的。当然对响应变量的均值和方差都有影响的因子也是可能的。

1.15

两因子实验 two-factor experiment

同时考察两个不同因子(1.5)对响应变量(1.2)的可能效应的实验。

注:如果两个因子的作用互不影响,术语“主效应(1.13)”仍可适用。即每个因子的主效应是它对响应变量均值的贡献。

1.16

k 因子实验 k-factor experiment

同时考察 $k(k \geq 2)$ 个不同因子(1.5)对响应变量(1.2)的可能效应的实验。

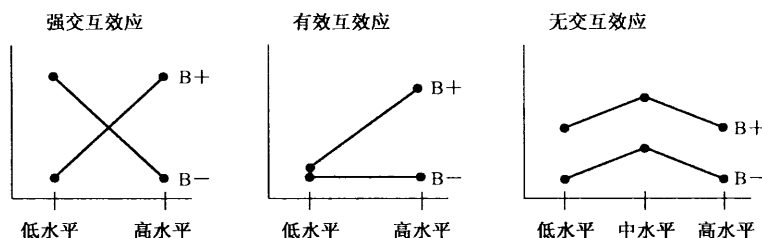
注:同义词是“多因子实验”。

1.17

交互效应 interaction

一个因子(1.5)对响应变量(1.2)的影响依赖于其他一个或多个因子的效应。

注1:交互效应表明一个因子对响应的主效应依赖于另一个因子水平而表现出来的不一致性。下面图形表示了这些现象。



注2:最常见的情况是考虑仅包含两个因子的交互效应,这种交互效应更准确地称为两向交互效应或一阶交互效应。也可考虑三个因子如A、B和C中,可能存在A、B的一阶交互效应依赖于因子C的水平,这称为二阶交互效应。类似地,可以考虑三阶、四阶或更高阶的交互效应。

注3:1.1中示例3提供了包含两个因子和它们的两向或一阶交互效应 τ_{ij} 实验的典型模型表示。

1.18

混杂 confounding

有意将两个或多个效应(主效应(1.13)或交互效应(1.16))相混合,使其不可区分。

注:混杂是一项重要的技术,比如在一些实验设计中,它可有效地利用指定的区组(1.11)。为应用这种技术,在设计时应有意将那些没有兴趣的效应(主效应或交互效应)与区组效应混杂,而同时避免与其他较重要的效应混杂。利用混杂技术可以减少“实验方案(1.30)”的实验次数。然而,混杂有时是因为实验过程中对设计的无意识改变或对设计的不完善造成的,从而程度不等地降低了实验的效率。

1.19

别名 alias

〈统计学〉由实验本质造成的与其他主效应(1.13)或交互效应(1.16)完全混杂(1.18)的效应(主效应或交互效应)。

1.20

曲性 curvature

对响应变量(1.2)与预测变量(1.3)之间直线关系的偏离。

注1:曲性只对于定量预测变量有意义,而对于分类的(名义的)或定性的(有序的)预测变量没有意义。检测曲性要求因子有多于两个水平。有些情况下,在中心点(因子最高水平和最低水平之间的中间点)处进行重复(1.27)可以检测和评估曲性。另一种观测曲性的方法是,将因子水平的范围作必要的扩展。

注2:对1.12示例中给出的模型,曲性容易通过下述形式建模:

$$Y = \mu + \beta X + \gamma X^2 + \epsilon$$

如果 γ 不为0,则有证据表明模型偏离了简单线性关系。

1.21

残差 residual

响应变量(1.2)的观测值与相应的基于假定模型(1.1)的预测值之差。

注1:响应变量的预测值所基于模型的参数要由数据来估计。

注2:残差是剩余误差(1.22)的观测值。

示例1:在1.1示例2的模型中, $y_{ij} - \hat{\mu} - \hat{\alpha}_i - \hat{\beta}_j$ 是对应于因子A的水平*i*、因子B的水平*j*的实验单元的残差。

示例2: $y_{ijk} - \hat{\alpha}_i - \hat{\beta}_j - \hat{\tau}_{ij}$ 是对应于1.1示例3中模型的残差。

示例3: $y_i - \hat{\epsilon}^2_0 + \hat{\beta}_1 x_i + \hat{\beta}_2 x_i^2$ 是对应于1.1示例4中模型的残差。

1.22

剩余误差 residual error

表示响应变量(1.2)的结果与相应的基于假定模型(1.1)的预测响应值之差的随机变量。

注1: 在本定义中,术语“预测响应值”可理解为利用假定模型、从实验数据导出的经验模型确定的估计响应值。

示例: 如果 $\hat{\mu}$ 和 $\hat{\beta}$ 分别是 1.13 注中 μ 和 β 的估计量,则 $y - \hat{\mu} - \hat{\beta}x$ 是在预测变量取值 x 时,给定 y 的观测值的剩余误差。

注2: 剩余误差包含实验误差和模型没有考虑的可查明变异来源。

注3: 实验中,剩余误差的方差通常采用从总平方和中减去假定模型所包含项的平方和,再除以相应自由度之差,即“残差方差”来估计(见 3.3 中示例 1 及其注和 3.4 中的示例)。

1.23

纯误差 pure error

反映一固定处理(1.10)组合上的重复(1.27)观测之变异的随机变量。

注1: 若仅在设计的中心点有重复,则响应在中心点的样本方差即是纯误差方差的一个估计。如果在多个处理组合上都有重复,则将这些处理组合上的估计合在一起即可得到纯误差方差的一个总估计。

示例: 回到 1.1 中示例 3,对固定的 (i, j) ,纯误差方差的一个估计为 $\frac{1}{n_{ij} - 1} \sum_{k=1}^{n_{ij}} (y_{ijk} - \bar{y}_{ij})^2$, 其中 $\bar{y}_{ij} = \frac{1}{n_{ij} - 1} \sum_{k=1}^{n_{ij}} y_{ijk}$;

若在每一个 (i, j) 组合都有重复,纯误差方差的一个合并估计为: $\frac{1}{N - IJ} \sum_{i,j,k} (y_{ijk} - \bar{y}_{ij})^2$, 其中 $i = 1, \dots, I; j = 1, \dots, J; k = 1, \dots, n_{ij}$ 。

注2: 纯误差在实际应用中,有两种不同的表示:一是指数学模型的总体方差 σ^2 ;二是指“样本”或“经验”纯误差,它和估计的剩余误差(1.22)方差(残差方差)一起是检验模型拟合不足的基础。在 1.1 关于模型列举的示例中,只有带有重复的示例 3 才容易直接估计纯误差。从数学的观点看,纯误差可构建为示例 2 中的 $\text{Var}(\epsilon_{ij})$ 、示例 3 中的 $\text{Var}(\epsilon_{ijk})$ 和示例 4 中的 $\text{Var}(\epsilon_i)$ 。

1.24

对照 contrast

(统计学)系数不全为 0,而系数和为 0 的响应值的线性函数。

注: 对观测值 $y_1, y_2, \dots, y_n$, 线性函数 $a_1 y_1 + a_2 y_2 + \dots + a_n y_n$ 是一个对照,当且仅当 $a_1 + a_2 + \dots + a_n = 0$,且所有的 a_i 不全为 0。

示例 1: 一个三水平因子的实验结果为 y_1, y_2 和 y_3 。在许多可能问题中,问题 1 关注实验在水平 1 和水平 3 上的响应之差(暂时忽略中间水平)。回答此问题的对照系数可取为 $(-1, 0, +1)$,它只需要 y_1 和 y_3 的值。如果因子水平是等间隔的,问题 2 关注是否有迹象表明响应变化模式具有(二次的)曲性(1.20),而不是一个简单的线性关系。此时需将 y_1 和 y_3 的平均值与 y_2 比较。(如果没有曲性, y_2 会靠近连接 y_1 和 y_3 的直线,即近似地等于它们的平均值。)

响应	y_1	y_2	y_3
问题 1 的对照系数	-1	0	+1
对照 1	$-y_1$		$+y_3$
问题 2 的对照系数	-1/2	+1	-1/2
对照 2	$-y_1/2$	$+y_2$	$-y_3/2$

本例说明了对连续变量的一种回归型研究方法。通常使用整数作为对照系数比使用分数更加方便,因而本例中,对照 2 的系数可采用 $(-1, +2, -1)$ 。

示例 2: 这个处理因子离散水平的示例提出了与示例 1 不同的两个问题。假定有三个供方: A_1, A_2, A_3 , 其中 A_1 采用一种新的制造技术,而 A_2 和 A_3 采用常规技术。问题 1: 采用新技术的供方 A_1 确实与采用常规技术的 A_2 和 A_3 不同吗? 此时要对照的是 y_1 与 y_2 和 y_3 的平均值。问题 2: 采用常规技术的两个供方有什么不同吗? 此时要对照 y_2 与 y_3 。尽管结果的解释方式不同,但对照系数的模式与前一个示例中的相似。

响应	y_1	y_2	y_3
问题 1 的对照系数	-2	+1	+1
对照 1	$-2y_1$	$+y_2$	$+y_3$
问题 2 的对照系数	0	-1	+1
对照 2		$-y_2$	$+y_3$

1.25

正交对照 orthogonal contrast

对应系数乘积之和等于 0 的两个或一组对照(1.24)。

示例 1:

	y_1	y_2	y_3
a_{i1} 对照 1	-1	0	+1
a_{i2} 对照 2	0	-1	+1

$$a_{i1} a_{i2} \quad \quad \quad 0 \quad \quad 0 \quad \quad +1$$

$\sum a_{i1} a_{i2} = 1$, 因此对照 1 与对照 2 不正交。

示例 2:

	y_1	y_2	y_3
a_{i1} 对照 1	-1	0	+1
a_{i2} 对照 2	-1	+2	-1

$$a_{i1} a_{i2} \quad \quad \quad +1 \quad \quad 0 \quad \quad -1$$

$\sum a_{i1} a_{i2} = 0$, 因此对照 1 与对照 2 为正交对照。

1.26

正交表 orthogonal array

任何一对因子(1.5)的所有可能因子水平(1.6)对出现的次数都相等的处理(1.10)的组合。

注:与正交表有关的“强度”概念将与“筛选设计(2.2)”有关,筛选设计也是正交表的一种应用。一个强度 d 的设计是任何 d 个因子的完全析因设计。强度 1 意味着每个因子的水平出现相同次数(有时称作平衡因子)。本术语定义的正交表的强度等于 2,也可以定义强度大于 2 的正交表。

1.27

重复 replication

对给定的处理(1.10)实施多于一次的实验。

注:有时只做一次实验也说成一次重复。

1.28

分区组 blocking

实验单元(1.9)在相对齐性区组(1.11)内的一种安排,每个区组内的实验误差(1.7)可望小于将处理随机安排在数目相近的实验单元时的实验误差(见 1.11,2.3)。

注 1:如果在实验中,除需要研究的已作为因子(主因子)之外,还有另外一些可查明原因的效应会对实验产生影响,且这些因素在全部实验中很难甚至不可能做到对所有实验单元都保持不变。此时就应该采用分区组的技术。区组的选择应尽可能地使同区组内这些可查明原因的效应最小,从而得到一个比较齐性的实验子空间。分析实验结果时也必须考虑区组的效应。

注 2:一个区组若可容纳全部处理,则称为“完全区组”;而只能容纳全部处理中的一部分的区组称为“不完全区组”。当将处理成对安排时,“对”就是区组。

1.29

随机化 randomization

将处理(1.10)安排到实验单元(1.9),并使每个实验单元有同等机会被安排某特定处理的过程。

注:随机化的目的在于尽可能避免实验中没有明确考虑的因素所造成的偏倚。随机化也可缩小由于不同时间或空间对实验结果的潜在影响。

1.30

实验方案 experimental plan

将处理(1.10)安排每个实验单元(1.9),同时确定其实施时间次序的过程。

1.31

经设计的实验 **designed experiment**

为满足特定目标所选择的实验方案(1.30)。

注：设计实验的目的是使实验以最有效且经济的方法获得相关结论。为实验选择一个合适的设计基于诸多考虑，比如问题种类、所得结论的一般性程度、能以高概率(功效)检出的效应的大小、实验单元的齐性和实验的实施费用等。一个好的经设计的实验通常会导致相对简单的统计分析和对结果的解释。

1.32

调优操作 **evolutionary operation****EVOP**

在正常生产过程中实施的、对生产设施及生产条件的序列实验。

注：调优操作的主要目的是为了获得改善过程及其产品的知识，以最小的费用，采用因子水平相对较小的偏移(在生产容许条件内)来设计实验。调优操作实验的因子变动范围通常极小，以避免生产不合格的产品，且需要大量的重复(1.27)以减小随机变异的影响。

1.33

完全随机化设计 **completely randomized design**

将处理(1.10)随机地安排到所有实验单元(1.9)的设计。

注：完全随机化设计仅适用于可以认为所有实验单元是齐性的，没有系统差异的情形。

1.34

顶点 **cube point**

因子(1.5)水平(1.6)形如向量 $(a_1, a_2, \dots, a_k)$ 的实验点，其中 $a_i = +1$ 或 -1 。

注：这些点正是二水平完全或部分析因设计(参见2.1)中点的类型。在中心复合设计中，顶点数多达 2^k 个(见2.4中示例1)。

1.35

星点 **star point**

因子(1.5)水平(1.6)形如向量 $(a_1, a_2, \dots, a_k)$ 的实验点，其中仅有一个 a_i 等于非零的 $+\alpha$ 或 $-\alpha$ ，其余皆为0。

注：所有星点仅有一个等于 $+\alpha$ 或 $-\alpha$ 的非零分量。在典型的中心复合设计中，一般设置 $2k$ 个星点(见2.4中示例1)。

1.36

中心点 **center point**

因子(1.5)水平(1.6)形如向量 $(a_1, a_2, \dots, a_k)$ 的实验点，其中所有 a_i 皆为0。

注：中心点的因子水平设置向量具有形式 $(0, 0, \dots, 0)$ ，对应于编码变量设计的中心点。中心点的重复次数 n_0 ，根据响应曲面设计的目的选定。在中心点重复有时是为估计所研究过程的纯误差。

1.37

可旋转性 **rotatability**

从拟合模型(1.1)预测的响应在与中心点等距离的点上都有相同方差的设计特征。

2 实验安排术语

2.1

完全析因实验 **full factorial experiment****析因实验** **factorial experiment**

包含有两个或多个因子(1.5)，每个因子考虑两个或多个水平(1.6)所有可能处理(1.10)的实验。

注1：从完全析因实验中可估计所有交互效应(1.16)和主效应(1.13)。

注2：析因实验通常从符号上表示为每个因子水平数的乘积。例如，基于三水平的因子A、二水平的因子B和四水平的因子C的实验记为 $3 \times 2 \times 4$ 析因实验。其乘积表示处理数。

注 3: 当析因实验中所有因子有相同的水平数时,通常记为水平数的因子数的 k 次幂。如有两个因子,每个因子有三个水平的实验记为 3^2 析因实验($k=2$),它需要 9 个用于安排不同处理的实验单元。

注 4: 完全析因设计有时也称为交叉设计。

2.1.1

部分析因实验 fractional factorial experiment

由完全析因实验(2.1)的一个子集组成的实验。

注: 术语“部分析因实验”中的“部分”通常是处理(1.10)组合数的一个简单比例。如: $1/2$, $1/4$ 等。

2.1.2

二水平实验 two-level experiment

所有因子(1.5)都为两个水平(1.6)的实验。

2.1.2.1

2^k 析因实验 2^k factorial experiment

有 k 个因子(1.5)且每个因子都为两个水平(1.6)的析因实验。

示例: 以下 2^4 析因实验适宜考察 4 个因子: 压力、温度、催化剂和操作人员对过程产生的效应。A 表示压力(水平取为低或高), B 表示温度(水平取为低或高), C 表示催化剂(水平取为有或没有), D 对应于操作人员(水平取为两个操作员中的某一个)。

实验单元	处理	A	B	C	D
1	(1)	—	—	—	—
2	a	+	—	—	—
3	b	—	+	—	—
4	ab	+	+	—	—
5	c	—	—	+	—
6	ac	+	—	+	—
7	bc	—	+	+	—
8	abc	+	+	+	—
9	d	—	—	—	+
10	ad	+	—	—	+
11	bd	—	+	—	+
12	abd	+	+	—	+
13	cd	—	—	+	+
14	acd	+	—	+	+
15	bcd	—	+	+	+
16	abcd	+	+	+	+

一个 2^4 实验由上表所列的 16 个不同的处理组成。符号“—”和“+”表示每个因子两个水平。通常用“—”号表示因子的低水平,“+”号表示高水平。当然这种规定是任意的。

上述表中的排列次序称作“标准叶茨(Yates)次序”,在分析阶段很有用。实验中实施这些处理的实际次序应经“随机化(1.29)”。因子 A 的“—”“+”号交替排列(—, +, —, +, …); 因子 B 是 2 个“—”号和 2 个“+”号交替排列; 因子 C 是 4 个“—”号和 4 个“+”号交替排列; 最后因子 D 是 1 至 8 实验单元(1.9)为“—”号, 9 至 16 单元为“+”号。在 GB/T 3358 的本部分的其后部分中,“—”也记为“—1”,“+”也记为“+1”。

上表第 2 列表示了处理的缩写记号。小写字母出现表示对应的大写因子的水平设置在高水平; 字母不出现意味着对应因子设置在低水平。所有因子都设置在低水平的情形记为“(1)”。

一个完全析因实验能对所有主效应(1.13)和交互效应(1.16)进行估计。如对 2^4 析因实验,有 4 个主效应(A,B,C,D)、6 个两因子(一阶)交互效应(AB,AC,AD,BC,BD,CD)、4 个三因子(二阶)交互效应(ABC,ABD,ACD,BCD)和 1 个四因子(三阶)交互效应(ABCD)。

每个效应(例如,A 的主效应、A 和 B 的交互效应,甚至 A、B、C 和 D 的四因子交互效应)都可用对照系数来估计。这将在 3.3 中讨论。

注:在实际应用中,也可用“1”与“2”分别代替示例表中的“-”(或“-1”)与“+”(或“+1”),它们分别表示二水平实验中因子的两个水平。这样 A_1 、 A_2 分别表示因子 A 取水平 1 与水平 2; B_1 、 B_2 分别表示因子 B 取水平 1 与水平 2,……。用这种表示方法,示例中的表有以下形式:

实验单元	处理	A	B	C	D
1	$A_1 B_1 C_1 D_1$	1	1	1	1
2	$A_2 B_1 C_1 D_1$	2	1	1	1
3	$A_1 B_2 C_1 D_1$	1	2	1	1
4	$A_2 B_2 C_1 D_1$	2	2	1	1
5	$A_1 B_1 C_2 D_1$	1	1	2	1
6	$A_2 B_1 C_2 D_1$	2	1	2	1
7	$A_1 B_2 C_2 D_1$	1	2	2	1
8	$A_2 B_2 C_2 D_1$	2	2	2	1
9	$A_1 B_1 C_1 D_2$	1	1	1	2
10	$A_2 B_1 C_1 D_2$	2	1	1	2
11	$A_1 B_2 C_1 D_2$	1	2	1	2
12	$A_2 B_2 C_1 D_2$	2	2	1	2
13	$A_1 B_1 C_2 D_2$	1	1	2	2
14	$A_2 B_1 C_2 D_2$	2	1	2	2
15	$A_1 B_2 C_2 D_2$	1	2	2	2
16	$A_2 B_2 C_2 D_2$	2	2	2	2

这种表示方法的优点是对处理的表示更为直观,而且容易推广到三个或三个以上水平的情形;缺点是当由两列生成另一列时的运算规则的定义不如用符号“-”和“+”方便(见 2.1.2.2 的注)。

2.1.2.2

2^{k-p} 部分析因实验 2^{k-p} fractional factorial experiment

从 2^k 完全析因实验(2.1)中精心选择的包含全部实验的 2^{k-p} 的实验。

注 1: 因子数很大时, 2^k 析因实验需要的处理(1.10)数非常大,实际中不可行。通过精心选择,从部分析因实验中也可获得与从完全析因实验中几乎等量的信息。特别注意,在作选择时应使认为实际上重要的主效应和交互效应仅与那些可忽略的交互效应相混杂。

注 2: 当 $p=1$ 时,相应的部分析因实验是“ $1/2$ 析因实验”;当 $p=2$ 时,相应的部分析因实验是“ $1/4$ 析因实验”,等。

注 3: 构造 2^{k-p} 部分析因实验可通过将 k 个因子分成两个组来实现。第一组有 $k-p$ 个因子,第二组有 p 个因子。第一组中的 $k-p$ 个因子分配到有 2^{k-p} 个实验单元的完全析因实验, 2^{k-p} 也是设计的实验单元数。对于每个实验单元,第二组中每个因子的水平由第一组因子水平定义。用第一组因子定义第二组因子的 p 个等式称作“生成关系”,因为由它可生成设计。由 p 个生成关系可导出 $2^{k-p}-1$ 个“定义关系”,它们确定了设计的性质。

示例: 考虑有 6 个因子和 16 个处理的实验,建议使用一个 2^{6-2} 部分析因设计。先设计一个 4 个因子(A,B,C 和 D)的完全析因实验,另两个因子 E 和 F 可用 A,B,C 和 D 的水平设置。生成关系的一种可能形式是: $E=ABC$, $F=BCD$ 。(注意,由此构造方法产生的 4 个字母链或字符串 ABCE 和 BCDF 称作“字”。例如,ABC 是一个 3 字母链的字,ACDEG 是一个 5 字母链的字,等等。)用 +1, -1 记因子的水平, A, B 和 C 的水平决定了 E 的对应水平(通过乘积 ABC);而 B, C 和 D 的水平决定了 F 的水平(通过乘积 BCD)。例如,对于实验单元 1, A 至 D 已经显示在 2.1.2.1 中示例的表中, E 和 F 也设置在低水平。主效应 E 与三因子交互效应 ABC 别名,主效应 F 与三因子交互效应 BCD 别名。完整的别名与混杂结构可由定义关系 $I=ABCE=BCDF=ADEF$ 导出。其中 I 是全由“+”(或 +1)构成的列。

2.1.3

设计分辨率 design resolution

定义关系中最短字母链的长度。

注1: 设计分辨率提供了特定设计中别名程度的一种描述,通常用罗马数字来表示。在实际中,三种最常见的分辨率是Ⅲ,Ⅳ和Ⅴ。

对于“分辨度为Ⅲ”的设计,最短的字母链(“I”除外)是3,这意味着设计的主效应不与其他主效应别名,但至少有一个主效应与一个两因子交互效应别名。对于“分辨度Ⅳ”设计,主效应不与其他主效应或任何两因子交互效应别名,但至少有一个两因子交互效应与另一个两因子交互效应别名。对于“分辨度为Ⅴ”的设计,主效应和两因子交互效应不与其他任何主效应或任何交互效应别名。

注2: 分辨率越高,可清晰估计的效应(主效应或交互效应)就越多。如果存在两个包含相同因子和实验单元(1.9)数的可选设计,应该选择分辨率较高的设计。幸运的是,对于实际感兴趣的大多数的 k 和 p ,已给出了最合适的定义关系(例如,参见参考文献[5] 410页)。

示例:续2.1.2.2中的示例。这个 2^{5-2} 部分析因设计的分辨率可由其定义关系得到。准确地说,设计分辨率是定义关系中最短的字或字母链(除“I”之外)的长度。利用运算规则: $IA=AI=A$, $IB=BI=B$, $I=A^2=B^2=C^2, \dots$, 生成关系 $E=ABC$ 等价于 $EE=ABCE$, 它也等价于 $I=ABCE$ 。类似地,由 $F=BCD$,可导出 $I=BCDF$ 。再考虑“广义交互效应” $ABCE \times BCDF = ADEF$ 便可得到所有定义关系。于是,定义关系是 $I=ABCE=BCDF=ADEF$ 。最短的字母链(除I之外)的长度为4,因此设计分辨率是Ⅳ。

注3: 设计生成元通常称为 Box-Hunter 生成元,尽管这一概念可以追溯到更早的工作。

2.2

筛选设计 screening design

从全部因子(1.5)中识别并选出其中的一部分,以作为后续研究的实验。

注1: 这样的设计通常重点只研究主效应(1.13)。如果考虑交互效应会造成分析复杂化,且可能需要增加实验次数。

示例1:2.1.2.2中的 2^{4-p} 部分析因设计(特别是 p 较大的设计)可看作筛选设计。

示例2:Plackett 和 Burman^[4]给出了一类处理数是4的整数倍的二水平筛选设计。当所研究的主效应数接近容许的处理个数时(无重复(1.27)),通常可选用这些设计。例如,下面给出的12个处理的 Plackett-Burman 设计可用来筛选多达11个因子的主效应。对于这个设计,如果存在两因子(一阶)交互效应,如A与B的交互效应AB,则会影响对因子C,D, $\dots$,K的主效应估计。

处理	A	B	C	D	E	F	G	H	I	J	K
1	+	-	+	-	-	-	+	+	+	-	+
2	+	+	-	+	-	-	-	+	+	+	-
3	-	+	+	-	+	-	-	-	+	+	+
4	+	-	+	+	-	+	-	-	-	+	+
5	+	+	-	+	+	-	+	-	-	-	+
6	+	+	+	-	+	+	-	+	-	+	-
7	-	+	+	+	-	+	+	-	+	-	-
8	-	-	+	+	+	-	+	+	-	+	-
9	-	-	-	+	+	+	-	+	+	-	+
10	+	-	-	-	+	+	+	-	+	+	-
11	-	+	-	-	-	+	+	+	-	+	+
12	-	-	-	-	-	-	-	-	-	-	-

注2: 许多 Plackett-Burman 设计与阿达玛(Hadamard)矩阵有关,阿达玛矩阵最初仅限于理论研究,后来才被认识到在实验设计中的用途。只要知道其中任意一列(或等价地,任意一行),整个阿达玛矩阵就完全确定。一种可行的构造方法是,令最后一行全为负号(-),其余的列元素可通过已知列向右移动一列,将列中的每个元素下移一个位置,而将第11行的元素移到第一行,即得到一个新列,继续这个过程直至构造出整个矩阵。下面给出了这类矩阵的一些例子。对于每种构造,只要标明第1列中正号(+)的位置就够了。

n	第 1 列中包含(+)的行
12	1,2,4,5,6,10
20	1,2,5,6,7,8,10,12,17,18
24	1,2,3,4,5,7,9,10,13,14,17,19

注意,上述 $n=12$ 时标示的行与示例 2 中完全列出的设计一致。许多 Plackett-Burman 设计可以按照这种一般方法以单个列的元素为基础构造出来。但 $n=28,52,76,92$ 和 100 的情形构造较为困难。有关细节参见参考文献[7]。

注 3: 田口玄一(G. Taguchi)采用一些缩写规则普及了 Plackett-Burman 设计的应用: L_{12} 表等同于上面给出的 12 个处理的 Plackett-Burman 设计, L_{20} 等同于 20 个处理的 Plackett-Burman 设计。需要注意的是, L 表给出设计矩阵元素的顺序通常不同于阿达玛矩阵元素的顺序。

注 4: Plackett-Burman 设计可适用于超饱和设计(因子比实验处理多)。有关细节可参见参考文献[8]和[9]。

2.3

区组设计 block design

将全部实验单元(1.9)分成若干个区组(1.11)的实验设计。

注:如果在实验设计中忽略了实验单元之间的非齐性,由于观测变异的增加减少了从实验中得到的信息量。采用区组设计可以提高实验满足设计目标的能力。

2.3.1

随机化区组设计 randomized block design

每个区组(1.11)内的处理(1.10)通过随机化(1.29)安排到各实验单元(1.9)的区组设计(2.3)。

注:随机化区组设计是将实验单元分为区组,同一区组内的单元比不同区组内的单元更加齐性的一种设计。每个区组内的处理随机分配到区组内的实验单元,处理效应可不受区组效应的干扰而有效地估计。

2.3.2

拉丁方设计 Latin square design

包含三个因子(1.5),每个因子有 h 个水平(1.6)的设计,其中任何一个因子的水平与其他两个因子水平的组合在 h^2 次实验中都出现且仅出现一次。

注 1:一个拉丁方设计包含三个因子,一个与处理(1.10)有关的主要因子和两个与区组效应有关的次要因子(行因子和列因子),所有因子有相同水平数 h 。该设计中包含由两个区组因子划分的 h^2 个实验单元,主要因子的 h 个水平随机分配到 h^2 个实验单元中使得每行和每列包含主要因子的每个水平(处理水平)恰好一次。因此,拉丁方设计是有两个区组变量(或者外部变异来源)的随机化区组设计。这种结构的一个约束条件是主要因子和区组因子的水平数必须相同。

示例:下面给出了三个 4×4 拉丁方,每一个都可变成一个拉丁方设计。4 个行可作为第一个区组因子的水平,4 个列可作为第二个区组因子的水平,主要因子的 4 个处理水平是 A,B,C 和 D。

ABCD	ABCD	ABCD
BADC	DCBA	CDAB
CDAB	BADC	DCBA
DCBA	CDAB	BADC

注 2:通过“相互抵消”,拉丁方设计一般可用来消除实验中不是主要关注的两个区组的效应。见“分区组(1.28)”。通常用方阵的行和列来表示区组。比如,行可以是天,列可以是操作员。主要因子和每个区组因子的水平数 h 必须相同。随机化(1.29)要求随机地分配主要因子的水平到字母,从列表中随机地选择或者按照参考文献[10]中描述的程序选择一个拉丁方,并且随机地安排区组因子的水平到方阵的行和列。[有 1 个 2×2 、12 个 3×3 、576 个 4×4 和 161 280 个 5×5 拉丁方,其中只有 1 个 2×2 、1 个 2×2 、4 个 4×4 和 56 个 5×5 “标准”拉丁方,它们的第一行和第一列以字母次序排列,通过行和列置换可以得到其他拉丁方。]

注 3:一个基本的假定是,区组因子与要研究的主要因子之间不存在交互效应,它们之间也不存在交互效应。如果这个假定不成立,剩余误差就会增加,因子的主效应就会与这些交互效应混杂。如果假定成立,这类设计对于最小化实验次数是特别有用。有时在区组的位置使用其他的主要因子,这样会有三个主要因子而没有区组因

子,这等价于一个在没有交互效应假定下的部分析因实验(2.1.1)。一些部分析因实验的设计可以形成拉丁方,为了更充分地理解对交互效应所作的假定,更希望从部分析因实验的角度来阐述这个问题。

2.3.3

正交拉丁方设计 Graeco-Latin square design

包含四个因子(1.5),每个因子有 h 个水平(1.6)的设计,其中任何一个因子的水平与其他三个因子水平的组合在 h^2 次实验中出现且仅出现一次。

注1:一个正交拉丁方设计包含四个因子,有 h^2 个根据三个区组因子(例如行因子、列因子和希腊字母表示的)划分成的实验单元,每个区组因子都有 h 个水平。主要因子的 h 个水平随机分配到 h^2 个实验单元中,使得每个处理在每行和每列中恰好出现一次,而且每个处理与每一个希腊字母也恰好同时出现一次。

注2:两个拉丁方称作正交的,如果一个方阵中的每个字母与另一个方阵中的每个字母恰好同时出现一次。两个正交的拉丁方可以结合生成一个正交拉丁方设计。

示例:下面是一个 4×4 正交拉丁方的示例:

A α	B β	C γ	D δ
B δ	A γ	D β	C α
C β	D α	A δ	B γ
D γ	C δ	B α	A β

行表示因子1,列表示因子2,希腊字母表示因子3,拉丁字母表示因子4,为主要因子。

2.3.4

不完全区组设计 incomplete block design

每个区组(1.11)不足以安排全部实验处理(1.10)的区组设计(2.3)。

注:事实上,一个随机化区组设计(2.3.1)可以解释成一个“完全”区组设计,旨在强调每个区组内有足够的实验单元来安排全部处理。

2.3.4.1

平衡不完全区组设计 balanced incomplete block design

BIBD

每个区组(1.11)包含主要因子(1.5)全部 l 个水平(1.6)中的 k 个水平,且每个由两个不同水平组成的水平对在全部 b 个区组中的 λ 个区组中出现的不完全区组设计(2.3.4)。

注1:本术语中,“平衡”是指相同的配对数,“不完全”是指在每个区组内不能安排主要因子的全部水平,“区组”是指在实验单元(1.9)的齐性集合上实施实验的策略。

示例1:考虑有4个处理、6个区组、每个区组包含2个处理的情形。例如要研究主要因子的四个水平(T_1, T_2, T_3, T_4),然而一天内只能考察两个水平。如果有6天可作实验,那么下面的方案是合适的:

天	T_1	T_2	T_3	T_4
1	*	*		
2			*	*
3	*		*	
4		*		*
5	*			*
6		*	*	

本示例中所有可能的处理对在同一区组内都出现一次。本示例的 BIB 设计的参数为: ($l=4, k=2, b=6, \lambda=1$)。

示例2:考虑主要因子有6个水平、10个区组、每个区组有3个水平的情形。由于此种情况下6个水平有20个可能的三元组,人们自然会怀疑实际中20个区组是否有必要。考虑下面的处理集合,每个三元组对应于一个区组:

$$(T_1, T_2, T_3), (T_1, T_2, T_4), (T_1, T_3, T_5), (T_1, T_4, T_6), (T_1, T_5, T_6) \\ (T_2, T_3, T_6), (T_2, T_4, T_5), (T_2, T_5, T_6), (T_3, T_4, T_5), (T_3, T_4, T_6)$$

其中每个水平对都恰好在两个区组中出现,这表明10个区组就足够。本示例的 BIB 设计的参数为: ($l=6, k=3, b=10, \lambda=2$)。

示例 3: 考虑有 7 个水平、7 个区组、每个区组内有 4 个水平的 BIBD ($l=7, k=4, b=7, \lambda=2$)。

区组	主要因子的水平			
1	1	2	3	4
2	2	3	4	7
3	3	4	5	1
4	4	5	6	2
5	5	6	7	3
6	6	7	1	4
7	7	1	2	5

注 2: 平衡不完全区组设计意味着, 主要因子的每个水平在实验中出现相同的次数 h , 且下述关系成立:

$$b \cdot k = l \cdot h, b \geq l, \text{ 且 } h(k-1) = \lambda(l-1)。$$

由于上述等式中的每个字母表示一个整数, 所以很清楚只有满足上述约束的组合 (l, k, b, h, λ) 才可能构造平衡不完全区组设计。但满足上述三个条件的 5 个整数 (l, k, b, h, λ) 的 BIBD 不一定存在。

注 3: 随机化(1.29)要求独立随机地安排区组和区组内的水平。

2.3.4.2

部分平衡不完全区组设计 partially balanced incomplete block design

PBIB

每个区组(1.11)包含主要因子(1.5)全部 l 个水平(1.6)中的 k 个水平, 但不是每个水平对在全部 b 个区组中出现次数都相同的不完全区组设计(2.3.4)。

注 1: 一个 l 个水平、 b 个区组的不完全区组设计称为是具有 $m(\geq 2)$ 种结合类的 PBIB 设计, 如果满足下述条件:

- 每个区组包含 $k(< l)$ 个不同的水平;
- 每个水平在 h 个区组内出现;
- 水平之间有满足下列条件的关系:
 - 任何两个水平必具有 $m(\geq 2)$ 种结合关系的一种, 这种结合关系是对称的(即: 如果水平 α 是水平 β 的第 i 种结合, 则水平 β 也是水平 α 的第 i 种结合);
 - 每个水平有 n_i 个第 i 种结合, 且 n_i 与选择的水平无关, $i=1, 2, \dots, m$;
 - 给定一对第 i 结合类的水平 α 和 β , 同时是 α 的第 j 结合类和 β 的第 k 结合类的水平数是 $p_{jk}^i, i, j, k=1, 2, \dots, m, p_{jk}^i$ 与第 i 结合类的水平对 (α, β) 无关;
- 互为第 i 种结合的水平对在 λ_i 个区组中出现($i=1, 2, \dots, m$), 不是所有的 λ_i 都相同。

注 2: PBIB 设计的参数 $l, b, h, k, \lambda_1, \lambda_2, \dots, \lambda_m, n_1, n_2, \dots, n_m, p_{jk}^i (i, j, k=1, 2, \dots, m)$ 满足以下关系:

$$lh=bk$$

$$n_1\lambda_1+n_2\lambda_2+\dots+n_m\lambda_m=h(k-1)$$

$$n_1+n_2+\dots+n_m=l-1$$

$$n_i p_{jk}^i = n_j p_{ik}^j = n_k p_{ij}^k$$

示例: 考虑下表参数为 $l=6, k=4, b=6, h=4, n_1=1, n_2=4, \lambda_1=4, \lambda_2=2$ 的 PBIB 设计:

区组	主要因子的水平			
1	1	4	2	5
2	2	5	3	6
3	3	6	1	4
4	4	1	5	2
5	5	2	6	3
6	6	3	4	1

本设计中每个水平出现 $h=4$ 次。从任意一个水平(比如, 水平 1)出发, 有 $n_1=1$ 个水平(比如, 水平 4)与水平 1 一起出现在 $\lambda_1=4$ 个区组中; 有 $n_2=4$ 个水平(水平 2, 3, 5, 6)与水平 1 一起出现在 $\lambda_2=2$ 个区组中。无论从哪个水平出发, 上述 n_1, n_2, λ_1 和 λ_2 都保持不变。

2.3.5

尧敦方 Youden square

通过删除或者添加拉丁方的行(或列)获得的,对于其中一个区组(1.11)因子(1.5)为随机化区组设计(2.3.1),而对于另一个区组因子为平衡不完全区组设计(2.3.4.1)的区组设计(2.3)。

注:尧敦方可以看作一个带有两个区组因子的设计,矩阵的行和列对应于区组因子,矩阵的元素表示主要因子的水平。例如,假定设计的列数与水平数相同,行数小于列数。每个水平在每行出现一次,生成一个关于行区组因子的随机化区组设计。然而,把注意力集中在列区组因子上就会生成一个平衡不完全区组设计。删除 4×4 拉丁方的第4行得到一个 3×4 尧敦方。

示例1:

		区组因子2 (列)			
		1	2	3	4
区组因子1 (行)	1	A	D	C	B
	2	B	A	D	C
	3	C	B	A	D
从拉丁方中删除		D	C	B	A

A, B, C 和 D 表示主要因子的四个水平。

示例2:下面的列表是一个 4×7 尧敦方:

3	4	5	6	7	1	2
5	6	7	1	2	3	4
6	7	1	2	3	4	5
7	1	2	3	4	5	6

本例中,行组成了一个随机化区组设计,列组成了一个参数为 $l=b=7, h=k=4, \lambda=2$ 的 BIB 设计。

2.3.6

裂区设计 split-plot design

将全部实验单元(1.9)按所安排的主要因子(1.5)相同水平(1.6)划分为称为裂区的小组,使在每个因子水平内可以考察一个或多个额外主要因子的设计。

示例:在两组重复(1.27)中测试因子 A 的三个水平。在 A 的每个水平(对应于一个裂区)内,考察因子 B 的两个水平。

裂区	重复 1		重复 2	
A ₁	A ₁ B ₂	A ₁ B ₁	A ₁ B ₂	A ₁ B ₁
A ₂	A ₂ B ₁	A ₂ B ₂	A ₂ B ₁	A ₂ B ₂
A ₃	A ₃ B ₁	A ₃ B ₂	A ₃ B ₂	A ₃ B ₁

注1:本例中的重复对第一阶主要因子(A)起着区组作用,安排 A 的某个水平的裂区对 A 内(裂区因子内)考察的第二阶主要因子 B 起着区组作用。这样,裂区内因子 B 的剩余误差(1.22)应该小于整个实验的剩余误差。在裂区设计中,对裂区内和裂区间效应会得到剩余误差的不同度量。裂区设计可以进一步推广到在第二阶因子水平内引入第三阶因子的情况。裂区设计常用来研究某一个因子,该因子在一个很长系列或很大区域内不易改变,而其他因子则容易改变的情形。

注2:裂区设计不仅普遍应用于农业,而且由此而得名,在工业实验中也常见。通常,一组因子水平需要大的实验单元(1.9),而另一组因子水平可以在较小的实验单元内比较。例如,与比较合金注入的不同类型的铸模所需要的合金量相比,比较预备合金的不同类型的熔炉所需要的合金量要多些。熔炉类型可看作第一阶因子的水平,铸模类型可看作第二阶(裂区内)因子的水平。另一个例子是一台大型机器,要改变其速度只能通过更换齿轮,这要耗费时间而且费用昂贵,因而希望不要频繁改变这个因子。在每种速度下生产的材料可以通过几种不同的技术热处理、在各种压力下成型、使用不同的打磨设备抛光。这些因子从一个水平滑动到另一个水平相对容易,它们组成裂区内因子(或第二阶因子),而速度变异组成裂区间因子(或第一阶因子)。

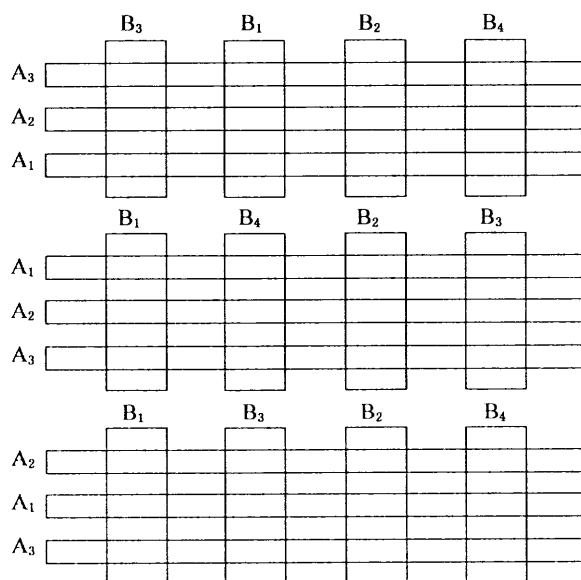
2.3.7

两向裂区设计 two-way split-plot design

裂区组设计 split-block design

第二阶因子(1.5)的水平(1.6)不是在每个区组(1.11)内独立地进行随机化(1.29)安排,而是在每次重复(1.27)内跨区带状安排的裂区设计(2.3.6)。

示例:对一个 3×4 两向裂区设计,(随机化后)合适的安排如下所示:



注1:两向裂区设计会使A和B的主效应(平均效应)的估计精度降低,但能提高交互效应的估计精度。与随机化区组或普通的裂区设计相比,两向裂区设计确定的交互效应通常更准确。

注2:在工业实验中,实际的限制有时要求必须使用两向裂区设计。例如,在纺织品工业中,因子A可能是不同的过氧化氯漂白过程,而因子B是冷却过程中不同的过氧化氢用量。

2.4

响应曲面设计 response surface design

旨在研究响应变量(1.2)和一组预测变量(1.3)之间函数关系的实验设计。

注1:响应曲面设计的一个明显好处是通过建议对预测变量(假定是连续的)的调节,来“改进”响应。

示例1:本例给出一个中心复合设计。它由顶点(1.34)、星点(1.35)和中心点(1.36)组成的一组处理(1.10),选择这些点以生成一个有效的(如可旋转的)设计。以下即是一个3个预测变量(1.3)的中心复合设计。

实验单位	x_1	x_2	x_3
1	-1	-1	-1
2	1	-1	-1
3	-1	1	-1
4	1	1	-1
5	-1	-1	1
6	1	-1	1
7	-1	1	1
8	1	1	1
9	0	0	0
10	0	0	0
11	-2	0	0

实验单位	x_1	x_2	x_3
12	2	0	0
13	0	-2	0
14	0	2	0
15	0	0	-2
16	0	0	2

实验单元 1-8 是设计的顶点,对应于一个 2^3 析因设计。预测变量的水平以其代码值给出。实验单元 9 和 10 是中心点,实验单元 11~16 是星点。先将前 8 个实验单元看作一组,然后再考虑安排其他处理。实际实施次序应随机化(1.29)。中心复合设计对这类设计单元的系列配置是很方便的。根据该设计得到的数据可拟合包含线性的(x_1, x_2, x_3)、二次的(x_1^2, x_2^2, x_3^2)和两个因子(1.5)的交互效应($x_1 x_2, x_1 x_3, x_2 x_3$)的模型。

注 2: 中心复合设计的一种常用变形是,对所有星点令 $\alpha=1$,此时考虑的因子水平较少,称为面中心的中心复合设计。由于因子水平的减少,设计牺牲了可旋转性(1.37)。

示例 2: 2^k 析因设计和平衡不完全区组设计的巧妙结合可以构造 Box-Behnken 设计。下面的设置组成三个变量的 Box-Behnken 设计。

实验单位	x_1	x_2	x_3
1	0	-1	-1
2	0	1	-1
3	0	-1	1
4	0	1	1
5	-1	0	-1
6	1	0	-1
7	-1	0	1
8	1	0	1
9	-1	-1	0
10	1	-1	0
11	-1	1	0
12	1	1	0
13	0	0	0
14	0	0	0
15	0	0	0

示例 3: 五边形设计是一个两因子设计,其设计点由单位圆(使用代码变量的单位)上等间距的五个点和可能重复的中心点组成。满足定义的一个五点集合是: $(1, 0), (0.309, 0.951), (-0.809, 0.588), (-0.809, -0.588), (0.309, -0.951)$ 。注意, $\cos 72^\circ = 0.309, \sin 72^\circ = 0.951$, 等等。

示例 4: 六边形设计是一个两因子设计,其设计点由单位圆(使用代码变量的单位)上等间距的六个点和可能重复的中心点组成。满足定义的一个六点集合是: $(1, 0), (0.5, 0.866), (-0.5, 0.866), (-1, 0), (-0.5, -0.866), (0.5, -0.866)$ 。注意, $\cos 60^\circ = 0.5, \sin 60^\circ = 0.866$, 等等。

注 3: 单位圆上任何规则几何图形都可作为响应曲面设计类中可旋转设计的基础。

2.5

混料设计 mixture design

在预测变量(1.3)之和等于固定常数的约束下构造的实验设计。

注: 表示合金中不同金属比例的因子是混料设计的典型例子,其设计空间必须满足 $X_1 + X_2 + \dots + X_k = 1$ 。如果还有进一步的约束条件,比如关于选定因子的最小比例,则可得到特殊的混料设计。有关混料设计的全面阐述参见文献[11]。

2.6

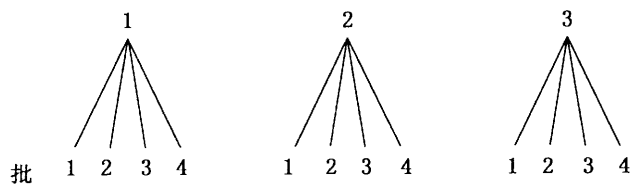
嵌套设计 nested design**层次设计 hierarchical design**

一个给定因子(1.5)的每个水平(1.6)仅在某其他因子的单个水平中出现的实验设计。

注1: 本设计主要用于评估嵌入因子的方差分量(1.8)。对于三个因子 A,B 和 C 的情形, 因子 B 的每个水平仅和因子 A 的单个水平一起出现; 类似地, 因子 C 的每个水平仅和因子 B 的单个水平一起出现。 k 因子嵌套设计有时称作 k 阶嵌套设计($k \geq 2$)。

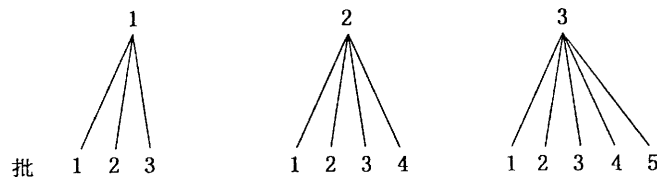
示例: 考虑三个不同的供方给一个公司提供了四批原材料的情形, 公司随后对其纯度进行测定。

供方



如图中所描绘的, 批嵌入在每个供方中, 由于来自供方 1 的批 1 不同于来自供方 2 的批 1。尽管批“标记”相同, 但是因子批和供方不交叉。当对每个供方提供不同批数的情形, 也可构造一个嵌套设计或层次设计, 如下面的设计:

供方



然而, 如果每个供方提供的批数相同, 分析就会简单得多。

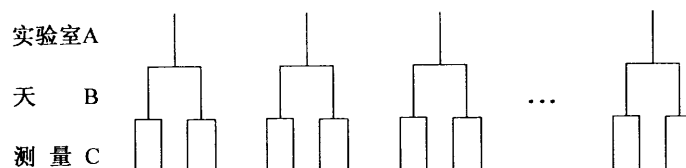
注2: 一般地, 嵌套设计用于评估方差分量方面的结果, 而不是评估有关响应水平或预测模型中差异的结果。

2.6.1

平衡嵌套设计 balanced nested design**完全嵌套设计 full nested design**

嵌入因子(1.5)的水平(1.6)为常数的嵌套设计(2.6)。

示例: 下图描述了一个平衡嵌套设计:



该设计是一个平衡嵌套设计, 因为每个实验室花费两天(因子 B 的水平数是 2), 每个实验室每一天得到两个测量结果(因子 C 的水平个数也是 2)。实验室所选定的日子可能是不同的, 因为测量大概是在给定测试时限内随机选择的。

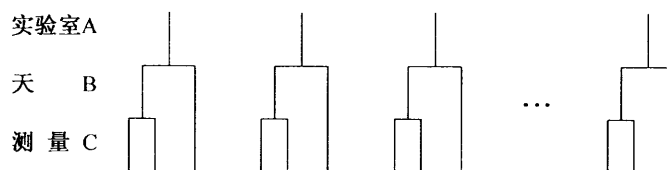
注: 有时可能调整一个因子水平的定义, 使之可与其他因子比较, 以便获得更有意义的结果。如 B_1 分配到星期一, B_2 分配到星期五。因此, 星期一得到的结果可以与星期五得到的结果比较。这样, 所有实验室都有共同的水平, 不像前述的情形, 它们(在实验室之间)是两个无关的天(日子)分配。这个结构是一个交叉的(即一个因子的每个水平与其他因子的每个水平都相遇), 而不是嵌套的分类, 故可看作是一个析因实验。

2.6.2

错层嵌套设计 staggered nested design

第二个嵌入因子(1.5)在第一个嵌入因子的一个水平(1.6)内有两个水平, 而在第一个嵌入因子的另一个水平内仅有一个水平的嵌套设计(2.6)。

示例: 下图描述了一个平衡嵌套设计:



注：对于错层嵌套设计，可以用近似相同的自由度数估计每个因子效应。

2.7

最优设计 optimal design

对某个特定目标函数，通常是设计矩阵(2.7.1)的一个函数，通过优化来确定因子(1.5)水平(1.6)设置的实验设计。

注：注意，最优设计的基础是模型正确。如果假定模型不正确，那么最优设计可能只是理论(数学上)最优的，而实际上可能是无用的。2.7.1.1~2.7.1.3提到的几类设计是常用的最优设计。

2.7.1

设计矩阵 design matrix

行表示单个处理(1.10)(根据假定模型(1.1)可能经过变换)的矩阵，可以由因子(1.5)水平(1.6)其他函数的导出水平(交互效应、二次项等)扩展，不过依赖于假定的模型。

注1：对给定的实验方案，可以根据假定模型设想多个设计矩阵。

注2：设计矩阵也称为模型矩阵，一般用 X 表示。 X 的每行对应于一个处理。如果模型包括总平均项 μ ，则 X 的第1列就全为1，其他列则表示因子或预测变量的函数。

2.7.1.1

D 最优设计 D-optimal design

最大化 $X'X$ 的行列式的设计。

注：D 最优设计准则与设计矩阵 X 对应系数的置信椭球体积有关。2.2 中的 Plackett-Burman 设计关于主效应模型是 D 最优的。

2.7.1.2

A 最优设计 A-optimal design

最大化 $X'X$ 的迹的设计。

注1： $X'X$ 的迹即是矩阵 $X'X$ 对角线上元素之和。

注2：A 最优准则结合了椭球体积和椭球球形的程度。

2.7.1.3

G 最优设计 G-optimal design

最小化设计空间上预测的最大方差的设计。

注：G 最优性不直接或明显包含 X 。但数学上，可以证明 D 和 G 最优性准则是等价的，从而可以用 G 最优准则得到 D 最优设计，因为前者的优化过程较为容易。

2.8

正交设计 orthogonal design

每对因子(1.5)都正交的设计。

注：一对因子是正交的，如果对任意两列的每个水平组合 (i, j) ，满足：

$$n_{ij} = n_i n_j / N$$

其中， n_{ij} 是在这两列中水平组合 (i, j) 出现次数， n_i 是水平 i 在其中一列中的出现次数， n_j 是水平 j 在另一列中出现次数， N 是实验单元总数。

2.9

饱和设计 saturated design

设计矩阵的列数与处理(1.10)数相同的设计。

注：当参数个数多于设计中的实验单元(1.9)时，不能得到合适的估计。

3 分析方法术语

3.1

图方法 graphical method

基于实验结果的图示性描述的分析。

注：简单的图能够对经设计的实验(1.31)的结果提供初步而有效的评价。

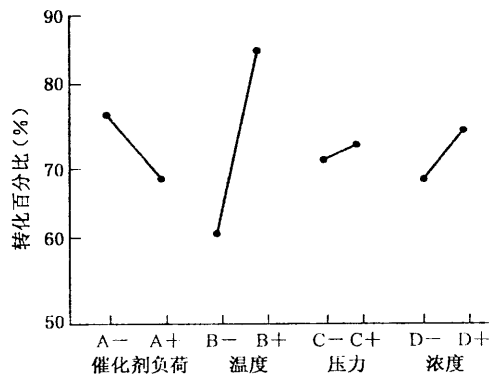
3.1.1

主效应图 main effects plot

表示单个因子(1.5)在不同水平(1.6)下的平均响应图。

注1：下图是参考文献[5]10.8例子的主效应图。响应变量是转化率，预测变量分别是催化剂用量(A)、温度(B)、压力(C)和浓度(D)。每个预测变量有两个水平，其中“-”表示低水平，“+”表示高水平。做了 2^4 析因实验。从图中可以看出，温度显然对转化起本质的作用，催化剂次之，其他两个因子的作用相当。为评判图中直线的斜率是否显著不等于零，还需进一步作分析。

注2：主效应图给出每个因子在不同水平下的平均响应，每个因子对响应所起的作用及其大小从图中可清楚地表达出来。不过交互作用的存在可能掩盖了因子的主效应(1.13)。



3.1.2

交互效应图 interaction plot

表示两个不同因子(1.5)水平下的平均响应图。

[见交互效应(1.17)]

注：交互效应图提供了一种检测交互效应是否存在的图示工具，若图中的线段不平行则表明存在交互效应。

3.1.3

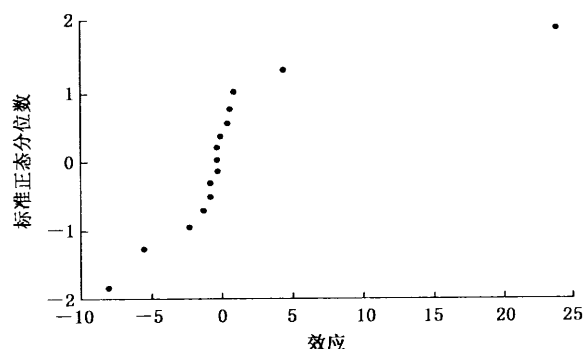
效应分位图 quantile plot of effects

在完全或部分析因设计下，标准正态分布的分位数对应于估计效应的图

示例：下图是一个效应分位图，数据取自3.1.1中的示例。

注：对无重复的实验，此图说明可能存在支配效应(即那些远离由图中主要点所形成的“主”线的左端或右端的点)。

图中右边最上面的点对应的温度效应为24。

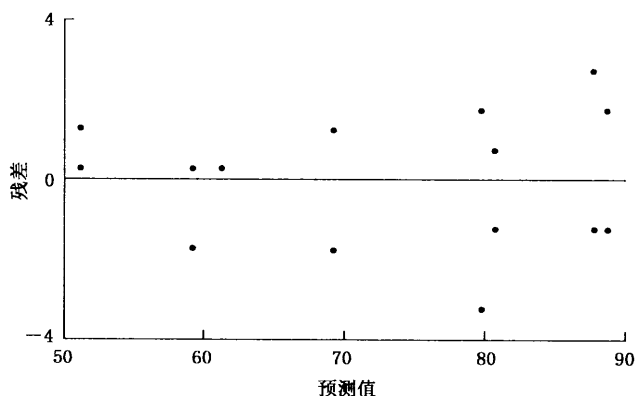


3.1.4

残差图 residual plot

残差(1.21)对应于相应的预测值或特定因子(1.5)水平(1.6)的图。

示例：续 3.1.1 中的示例，使用模型(1.1)为 4 个主效应加上交互效应 BD。



3.2

最小二乘法 method of least squares

通过最小化观测值与假定模型(1.1)的预测值的差值的平方和，进行参数估计的技术。

注 1：记 e 为观测值与假定模型的预测值的差值，则最小二乘法是通过最小化 $\sum e^2$ ，其中求和是对所有处理的。

注 2：虽然推断方法经适当调整后，也可以处理彼此相关的实验误差，但通常假定个体观测的实验误差是独立的。

常用的方差分析、回归分析和协方差分析都基于最小二乘法的，它们在计算和解释上有不同的优势，源于为便于数据分组的实验安排的某种平衡。

3.3

回归分析 regression analysis

评价响应变量(1.2)与预测变量(1.3)关系模型(1.1)的技术。

注 1：回归分析通常包括假定模型下的进行参数估计的过程，它通过优化目标函数值来达到(例如，最小化观测响应与模型预测响应之差的平方和)。已有的统计软件为获得参数估计、估计量的标准差，以及进行模型诊断提供了方便。回归分析同样可以考虑响应的其他度量。例如，如果在有重复的析因实验设计中，对散度效应(1.14)感兴趣，将 S^2 (S^2 是重复点的样本方差)的对数作为响应，比原响应更容易进行分析与解释。

注 2：回归分析扮演着与方差分析相似的角色，特别适用于当因子的水平(1.6)是连续的而且强调的显式的预测模型时。回归分析也可用于带缺失数据的经设计的实验(1.31)，而不象方差分析那样要求数据间的平衡。不过，不平衡会增加假设检验的序相依性(公共的元素包含在第一相关项中，而不在后面的项中)并失去平衡实验的其他优良性。对平衡实验而言，这两种方法都是最小二乘法的变形，并给出可比较的结果。

示例 1：考虑一有 3 个定量因子的 2^3 析因实验的正交设计，无重复，实验单元的假定模型为：

$$Y = \beta_0 x_0 + \beta_1 x_1 + \beta_2 x_2 + \beta_3 x_3 + e$$

其中：

$x_0 = 1$ ；

x_1 是因子 A 的水平；

x_2 是因子 B 的水平；

x_3 是因子 C 的水平；

e 是随机误差。

此模型可应用于具有水平(1.6)代码为 -1 和 +1 的 3 因子的实验。

示例 1 的回归分析表

(记号: $x_{ji} = X_{ji} - \bar{X}_i$, 其中 $\bar{X}_i = \sum_{j=1}^4 X_{ji}/4$)

变异来源	回归系数	平方和(SS)	自由度(DF)	均方(MS)
总和	—	$S_T = \sum Y_i^2$	8	—
常数项(X_0)	$\beta_0 = \frac{\sum x_{0i} Y_i}{\sum x_{0i}^2}$	$Sx_0 = \beta_0 \sum x_{0i} Y_i$	1	Sx_0
X_1 (因子 A)	$\beta_1 = \frac{\sum x_{1i} Y_i}{\sum x_{1i}^2}$	$Sx_1 = \beta_1 \sum x_{1i} Y_i$	1	Sx_1
X_2 (因子 B)	$\beta_2 = \frac{\sum x_{2i} Y_i}{\sum x_{2i}^2}$	$Sx_2 = \beta_2 \sum x_{2i} Y_i$	1	Sx_2
X_3 (因子 C)	$\beta_3 = \frac{\sum x_{3i} Y_i}{\sum x_{3i}^2}$	$Sx_3 = \beta_3 \sum x_{3i} Y_i$	1	Sx_3
残差	—	$S_E = S_T - Sx_0 - Sx_1 - Sx_2 - Sx_3$	4	$S_E/4$

注 3: 若在同一区组里重复做 2^3 次实验,“总和”(第一行)的自由度为 16,“残差”的自由度为 12。“残差平方和”可分解为自由度分别是 8 和 4 的“重复”和“拟合不足”两部分。

示例 1 的回归分析表-有重复实验的附加部分

变异来源	平方和(SS)	自由度(DF)	均方(MS)	F 值	均方的期望
残差	S_E	12	$S_E/12$		
重复	$S_R = \sum_j (Y_{.j} - \bar{Y}_i)^2$	8	$S_R/8$		
拟合不足	$S_L = S_E - S_R$	4	$S_L/4$		

在适当的正态性假定下,可以通过变异来源的均方与适当的误差项的 F 统计量来检验变异来源的统计显著性。对一次重复实验的情况,应该对“回归”项相对“残差”项进行检验。对两次重复实验的情况,应该对“拟合不足”项相对于“重复(实验误差(1.7))”项进行检验,以确定模型(1.1)是否恰当。“重复”项是对实验误差的度量,它不受模型(1.1)不当的影响,应纳入“残差”项。

3.4

方差分析 analysis of variance

ANOVA

将响应变量(1.2)的总变异分解为若干有意义的分量以确定变异来源的技术。

注: 方差分析可作方差分量(1.8)估计以及模型参数的假设检验。

方差分析表通常包含如下各列:

- 变异来源;
- 平方和(SS);
- 自由度(DF);
- 均方(MS)(平方和除以自由度);
- F 统计量(行中的均方与误差均方的比值);
- 均方的期望(由模型参数确定的均方的数学期望)。

表中的行表示特定的因子效应或交互效应、区组(若实验设计中进行了分区组(1.28))效应或误差(模型或区组没有考虑残差(1.21)效应)。“总和”行通常对应于对总平均的总平方和,自由度为样本量减 1。

示例: 考虑一随机化区组设计,记 Y_{ij} 为在 h 个区组中第 j 区组里因子 A 的 l 个水平(1.6)中第 i 水平下的观测值, ($i=1,2,\dots,l; j=1,2,\dots,h$),主因子 A 表示一个固定的处理效应;因子 B 表示一个固定的区组效应。方差分析表计算如下:

方差分析表

变异来源	平方和(SS)	自由度(DF)	均方(MS)	F 值	均方的期望 [E(MS)]
总和	$S_T = \sum_i \sum_j (Y_{ij} - \bar{Y})^2$	$v_T = hl - 1$	—	—	—
因子 A (处理)	$S_A = h \sum_i (Y_{ij} - \bar{Y}_i)^2$	$v_A = l - 1$	$MS_A = \frac{S_A}{v_A}$	$F(v_A, v_e) = \frac{MS_A}{MS_e}$	$\sigma^2 + hK_A^2$
因子 B (区组)	$S_B = l \sum_j (Y_{ij} - \bar{Y}_j)^2$	$v_B = h - 1$	$MS_B = \frac{S_B}{v_B}$	$F(v_B, v_e) = \frac{MS_B}{MS_e}$	$\sigma^2 + lK_B^2$
误差	$S_e = \sum_{i,j} (Y_{ij} - \bar{Y}_i - \bar{Y}_j + \bar{Y})^2$	$v_e = (l-1)(h-1)$	$MS_e = \frac{S_e}{v_e}$	—	σ^2

在方差分析表中：

$$S_T = S_A + S_B + S_e$$

$$v_T = v_A + v_B + v_e$$

$F(v_1, v_2)$ 是自由度为 v_1, v_2 的 F 统计量。

关于观测值的模型为：

$$Y_{ij} = \mu + \alpha_i + \beta_j + e_{ij}; i = 1, 2, \dots, l; j = 1, 2, \dots, h$$

且

$$\sum \alpha_i = \sum \beta_j = 0; e_{ij} \sim N(0, \sigma^2);$$

$$K_A^2 = \frac{\sum \alpha_i^2}{l-1}; K_B^2 = \frac{\sum \beta_j^2}{h-1}$$

式中：

μ ——总平均；

α_i ——第 i 个处理的效应；

β_j ——第 j 个区组的效应；

e_{ij} ——实验误差。

在本例中，假定水平(1.6)是设定的。 μ, α_i, β_j 和 σ^2 的最小二乘估计为：

$$\hat{\mu} = \bar{Y} = \sum \sum \frac{Y_{ij}}{hl}$$

$$\hat{\alpha}_i = \bar{Y}_i - \bar{Y} = \sum_j \frac{Y_{ij}}{h} - \bar{Y}$$

$$\hat{\beta}_j = \bar{Y}_j - \bar{Y} = \sum_i \frac{Y_{ij}}{l} - \bar{Y}$$

$$\hat{\sigma}^2 = \sum_i \sum_j \frac{(Y_{ij} - \bar{Y}_i - \bar{Y}_j + \bar{Y})^2}{(l-1)(h-1)} = S_e$$

在本例中，由于在随机化区组设计中每个单元的观测数都相等，所以公式都得到简化。

注 2：方差分析的基本假定是效应是可加的，且实验误差都是独立、正态、零均值等方差(齐性)的。方差分析技术与 F 比一起可用来检验变异来源的效应的统计显著性，且可以估计这些变异来源的方差。只在检验统计显著性和构造置信区间时才需要正态分布的假定。若要考察平均效应和交互效应，通常将它们列成二向(或 k 向)表的形式。本示例采用模型 1(固定效应模型(3.4.1))。

当误差正态性的假定不成立时，可以对响应变量进行变换(例如，取对数)或用非参数方法来处理。

3.4.1

固定效应方差分析 fixed effects analysis of variance

每个因子(1.5)的水平(1.6)是在因子取值范围内被预先选定情形下的方差分析(3.4)。

注：在固定水平(1.6)下不能计算方差分量。此模型有时称为方差分析模型 1。

3.4.2

随机效应方差分析 random effects analysis of variance

每个因子(1.5)的水平(1.6)都是从该因子水平的总体中随机抽取情形下的方差分析(3.4)。

注：在随机水平下，主要兴趣通常是获取方差分量(1.8)的估计。此模型一般称为方差分析模型 2。

示例：考虑一个对原料进行批处理操作的情形。当少数批原料是从所有批的总体中随机抽出时，“批”可以看作是实验里的一个随机因子。

3.4.3

混合模型方差分析 mixed model analysis of variance

一些因子(1.5)的水平(1.6)是预先选定的,其他因子的水平是从这些因子水平的总体中抽样而得的情形下的方差分析(3.4)。

注:方差分量(1.8)只对随机因子有意义,对效应的估计只适用于固定因子。此模型也称为方差分析模型3。

3.5

协方差分析 analysis of covariance**ANCOVA**

在一个或多个伴随变量影响响应变量(1.2)情形,估计和检验处理(1.10)效应的技术。

注1:协方差分析可以看作是回归分析(3.3)和方差分析(3.4)的组合。

注2:通常在实验设计时不考虑伴随变量,而在分析时必须考虑其对结果的影响。例如,不同实验单元(1.9)中,每个单元的某些化学成分在量上可能有所不同,它们可以测量,但不能控制。

参 考 文 献

- [1] GB/T 3358.1—2009 统计学词汇及符号 第1部分:一般统计术语与用于概率的术语.
- [2] ISO 10241:1992 International terminology standards—Preparation and layout.
- [3] Box G E P, Hunter W G, Hunter J S. Statistics for Experimenters. An Introduction to Design, Data Analysis, and Model Building. John Wiley & Sons, New York, 1978.
- [4] Plackett R L, Burman J P. The design of optimum multifactorial experiments, *Biometrika*, 33, 1946, pp. 305-325.
- [5] John P W M. Statistical Design and Analysis of Experiments. The Macmillan Company, New York, 1971.
- [6] Lin D K J. A new class of supersaturated designs. *Technometrics*, 35, 1993, pp. 28-31.
- [7] Wu C F J. Construction of supersaturated designs through partially aliased interactions. *Biometrika*, 80, (3), 1993, pp. 661-669.
- [8] Fisher R A, Yates F. Statistical Tables for Biological, Agricultural, and Medical Research. Oliver and Boyd, Edinburgh, 4th edition, 1953.
- [9] Cornell J A. Experiments with Mixtures. 2nd ed. John Wiley, New York, 1990.
- [10] GB/T 3358.2—2009 统计学词汇及符号 第2部分:应用统计.
- [11] GB/T 6379.1—2004 测量方法与结果的准确度(正确度与精密度) 第1部分:总则与定义.
- [12] Box G E P, Draper N R. Empirical Model-Building and Response Surfaces. John Wiley & Sons, New York, 1987.

索引

汉语拼音索引

- B**
- 饱和设计····· 2.9
别名····· 1.19
不完全区组设计····· 2.3.4
部分析因实验····· 2.1.1
- C**
- 残差····· 1.21
残差图····· 3.1.4
重复····· 1.27
处理····· 1.10
纯误差····· 1.23
错层嵌套设计····· 2.6.2
- D**
- 单因子实验····· 1.12
等级嵌套设计····· 2.6
调优操作····· 1.32
顶点····· 1.34
对照····· 1.24
- F**
- 方差分量····· 1.8
方差分析····· 3.4
分区组····· 1.28
- G**
- 固定效应方差分析····· 3.4.1
- H**
- 回归分析····· 3.3
混合效应方差分析····· 3.4.3
混料设计····· 2.5
混杂····· 1.18
- J**
- 交互效应····· 1.17
- 交互效应图····· 3.1.2
经设计的实验····· 1.31
- K**
- 可旋转性····· 1.37
- L**
- 拉丁方设计····· 2.3.2
两水平实验····· 2.1.2
两向裂区设计····· 2.3.7
两因子实验····· 1.15
裂区设计····· 2.3.6
裂区组设计····· 2.3.7
- M**
- 模型····· 1.1
- P**
- 平衡不完全区组设计····· 2.3.4.1
平衡嵌套设计····· 2.6.1
- Q**
- 嵌套设计····· 2.6
区组····· 1.11
区组设计····· 2.3
曲性····· 1.20
- S**
- 散度效应····· 1.14
筛选设计····· 2.2
设计分辨率····· 2.1.3
设计矩阵····· 2.7.1
设计空间····· 1.4
设计区域····· 1.4
剩余误差····· 1.22
实验单元····· 1.9
实验方案····· 1.30
实验误差····· 1.7

水平·····	1.6
随机化·····	1.29
随机区组设计·····	2.3.1
随机效应方差分析·····	3.4.2

T

图方法·····	3.1
----------	-----

W

完全嵌套设计·····	2.6.1
完全随机化设计·····	1.33
完全析因实验·····	2.1

X

析因实验·····	2.1
响应变量·····	1.2
响应曲面设计·····	2.4
效应分位图·····	3.1.3
协方差分析·····	3.5
星点·····	1.35

Y

尧敦方·····	2.3.5
----------	-------

因子·····	1.5
预测变量·····	1.3

Z

正交表·····	1.26
正交对照·····	1.25
正交拉丁方设计·····	2.3.3
正交设计·····	2.8
中心点·····	1.36
主效应·····	1.13
主效应图·····	3.1.1
最小二乘法·····	3.2
最优设计·····	2.7

2^{k-p} 部分析因实验·····	2.1.2.2
2^k 析因实验·····	2.1.2.1
A 最优设计·····	2.7.1.2
D 最优设计·····	2.7.1.1
G 最优设计·····	2.7.1.3
k 因子实验·····	1.16

英文对应词索引

A

alias	1. 19
analysis of covariance	3. 5
analysis of variance	3. 4
ANCOVA	3. 5
ANOVA	3. 4
A-optimal design	2. 7. 1. 2

B

balanced incomplete block design	2. 3. 4. 1
balanced nested design	2. 6. 1
block	1. 11
block design	2. 3
blocking	1. 28

C

centre point	1. 36
completely randomized design	1. 33
confounding	1. 18
contrast	1. 24
cube point	1. 34
curvature	1. 20

D

design matrix	2. 7. 1
design region	1. 4
design resolution	2. 1. 3
design space	1. 4
designed experiment	1. 31
dispersion effect	1. 14
D-optimal design	2. 7. 1. 1

E

evolutionary operation	1. 32
experimental error	1. 7
experimental plan	1. 30
experimental unit	1. 9

F

factor	1. 5
--------------	------

factorial experiment	2. 1
fixed effects analysis of variance	3. 4. 1
fractional factorial experiment	2. 1. 1
full factorial experiment	2. 1
full nested design	2. 6. 1

G

G-optimal design	2. 7. 1. 3
Graeco-Latin square design	2. 3. 3
graphical method	3. 1

H

hierarchical design	2. 6
---------------------------	------

I

incomplete block design	2. 3. 4
interaction	1. 17
interaction plot	3. 1. 2

K

<i>k</i> -factor experiment	1. 16
-----------------------------------	-------

L

Latin square design	2. 3. 2
level	1. 6

M

main effect	1. 13
main effects plot	3. 1. 1
method of least squares	3. 2
mixed model analysis of variance	3. 4. 3
mixture design	2. 5
model	1. 1

N

nested design	2. 6
---------------------	------

O

one-factor experiment	1. 12
optimal design	2. 7
orthogonal array	1. 26

orthogonal contrast	1. 25
orthogonal design	2. 8

P

partial balanced incomplete block design	2. 3. 4. 2
predictor variable	1. 3
pure error	1. 23

Q

quantile plot of effects	3. 1. 3
--------------------------------	---------

R

random effects analysis of variance	3. 4. 2
randomization	1. 29
randomized block design	2. 3. 1
regression analysis	3. 3
replication	1. 27
residual	1. 21
residual error	1. 22
residual plot	3. 1. 4
response surface design	2. 4
response variable	1. 2
rotatability	1. 37

S

saturated design	2. 9
screening design	2. 2
split-block design	2. 3. 7
split-plot design	2. 3. 6
staggered nested design	2. 6. 2
star point	1. 35

T

treatment	1. 10
two factor experiment	1. 15
two-level experiment	2. 1. 2
two-way split-plot design	2. 3. 7

V

variance component	1. 8
--------------------------	------

Y

Youden square	2.3.5
2^k factorial experiment	2.1.2.1
2^{k-p} fractional factorial experiment	2.1.2.2
