

全国高等林业院校试用教材

林业试验设计

续九如 黄智慧 编著

中国林业出版社

PDG

全国高等林业院校试用教材

林业试验设计

续九如 黄智慧 编著

中国林业出版社



(京) 新登字 033 号

图书在版编目 (CIP) 数据

林业试验设计/续九如, 黄智慧编著. —北京: 中国林业出版社, 1995. 6

全国高等林业院校试用教材

ISBN 7-5038-1361-X

I. 林… II. ①续… ②黄… III. 林业-试验设计 (数学) -高等学校-教材 IV. S7-3

中国版本图书馆 CIP 数据核字 (95) 第 02848 号

中国林业出版社出版

(100009 北京西城区刘海胡同 7 号)

北京龙华印刷厂印刷 新华书店北京发行所发行

1995 年 3 月第一版 1995 年 3 月 第一次印刷

开本: 787×1092 毫米 1/16 印张: 12.5

字数: 300 千字 印数: 1500 册

定价: 6.55 元

前 言

20 世纪以来,由于数理统计学的发展,使生物科学,包括林业科学逐渐成为可以用数学方法来研究的科学。农、林业的科学试验常常与田间试验方法密切联系在一起,借助于数理统计学的原理,试验的设计和分析方法不断得到发展和完善。农业和畜牧业的试验设计早已成为一门重要的专业课在高等学校开设,而且已有大量的教材和专著问世。相比之下,林业比较落后,直到现在国内还未见到一本正式出版的林业试验设计方面的高等学校教材。一般而言,林业试验比农业试验周期更长,地域更广,因而影响也就更为深远。为此,教学、科研、生产各方面都迫切需要一本既能结合林业实践,又能详细阐述各种试验设计方法的专门著作,本书就是适应这种形势而产生的。其主要目的是使读者掌握林业试验的基本原理,能够正确地设计试验和正确地分析试验结果,从而得出科学的结论,用以指导和发展生产。

本书是在本校两次印刷讲义基础上改写而成的。改写过程中,广泛借鉴了其他学科和国外有关教材。本书的主要特点是从数学模型出发,尽量反映本学科的现有先进水平,比较详尽地讨论各种适合于林业试验的设计方法与统计分析技术。内容侧重于应用,对于数学原理,强调理解,而不强调详细的推导过程。这样,既能使本书具有明确的实用性,可以成为林业科技工作者的工具,又在较少的篇幅内保持了学科的先进性和系统性。

全书共分 12 章。前 8 章属基本原理和常用设计方法,供本科生教学和科技工作者安排试验时选择使用;后 4 章为比较复杂的设计方法,供研究生和科研人员参考选用。为方便,各章附有习题,书末附有必要的统计用表。

本书第 7、8、9 章由黄智慧执笔,其余各章均由续九如执笔。全书经华南热带学院林德光教授和北京师范大学刘来福教授审稿和指导,东北林业大学张培呆教授和本校符伍儒教授、贾乃光教授审阅了书稿,并提出了许多宝贵的修改意见,谨此一并致谢!

由于水平所限,书中仍难免有错误。欢迎读者和专家批评指正。

编著者

1992 年 10 月于北京林业大学



目 录

前言

第1章 概论	(1)
第1节 林业试验设计的任务	(1)
第2节 林业试验设计的原则	(1)
第3节 林业试验设计的误差控制	(2)
第2章 林业试验设计的数学模型和方差分析	(5)
第1节 设计种类及其线性数学模型	(5)
第2节 不同试验设计的方差分析	(8)
第3章 简单林业试验设计	(18)
第1节 试验布置	(18)
第2节 统计分析	(21)
第3节 几种简单设计方法的评价	(27)
第4章 完全随机区组设计	(29)
第1节 设计方法	(29)
第2节 统计分析	(32)
第3节 评价	(40)
第5章 平衡不完全区组设计	(42)
第1节 设计方法	(42)
第2节 统计分析	(43)
第3节 评价	(46)
第6章 拉丁方设计	(48)
第1节 设计方法	(48)
第2节 统计分析	(51)
第3节 评价	(60)
第7章 裂区设计	(62)
第1节 设计方法	(62)
第2节 统计分析	(64)
第3节 评价及应用	(69)
第8章 正交设计	(71)
第1节 正交设计的原理和步骤	(71)
第2节 统计分析方法	(74)
第3节 正交设计的应用和正交表的构造	(80)
第9章 格子设计	(86)
第1节 格子设计的种类和方法	(86)

第2节	统计分析	(89)
第3节	评价与应用	(99)
7	第10章 析因设计	(100)
	第1节 设计方法	(100)
	第2节 统计分析	(102)
8	第11章 混杂设计	(115)
	第1节 完全混杂试验设计	(115)
	第2节 部分混杂试验设计	(124)
	第3节 评价	(128)
9	第12章 回归设计	(130)
	第1节 一次回归正交设计	(130)
	第2节 二次回归正交设计	(133)
	第3节 二次回归旋转设计	(142)
	附表	
	附表1 t 分布的临界值(t_{α})表	(145)
	附表2 F 分布的临界值(F_{α})表	(146)
	附表3 多重比较中的 q_{α} 表	(151)
	附表4 新复全距检验的 SSR_{α} 表	(153)
	附表5 BIB 设计表	(155)
	附表6 正交拉丁方表	(165)
	附表7 正交表	(167)
	附表8 正交表构造的运算法则	(188)
	附表9 百分数的反正弦转换表	(189)
	主要参考文献	(194)



第 1 章 概 论

第 1 节 林业试验设计的任务

广义的试验设计是指在科研工作之前对整个试验课题的设计，包括方案拟定、试验单位选择、资料收集、统计与分析等；狭义的试验设计是指对试验和分析方法的专门考虑。试验设计是否合理完善，应以它能否用最少的劳动和时间获得科学而有用的结论作为衡量准则。科学地设计一个试验时，应充分考虑到如何减少使试验结果偏离实际情况的各种干扰因素；判断试验结果的差异是由于处理效应造成的，还是由于误差造成的。要做到这一点，必须按照数理统计的原理来进行试验设计。数理统计学是制定科学的试验设计方案和合理地分析试验结果的必要基础。

林业试验设计与其它试验设计一样，都是为了更有效地研究和解决实际问题。由于树木生长周期长，地域广，易受环境影响，一次设计不当就会贻误多年，因此制订科学的林业试验设计方案就显得更为重要和迫切。

一个好的林业试验设计应该具备以下特性：①**典型性**：试验场地和试验材料都具有代表性，使得试验成果能够在生产上推广。②**精确性**：试验材料来源清楚，试验数据可靠，试验结果可信。③**重演性**：试验结果能够经得起重复，别人在相同情况下也能得到相同的结果。④**坚韧性**：遭受自然或人为的局部破坏之后仍然能够分析。在满足上述要求的前提下，林业试验设计还应尽可能做到简便易行。

第 2 节 林业试验设计的原则

与一切其它试验设计一样，林业试验设计也必须遵守以下 3 条原则。

一、设置重复

在田间试验中，试验误差是客观存在和不可避免的，只有设置重复的试验才能减少误差和正确地估计试验误差。试验误差是根据同一处理内的变异而测定的，如果每个处理只有一个观测值，就无法计算差异，也就无法估计误差的大小。同时，由于土壤肥力、水分以及小气候等环境条件的影响，一个观测值往往不能精确反映某个处理的效果，只有用几次重复的平均值才能代表这个处理的真实效果而避免偶然性影响。

试验证明，在同样面积的条件下，采用扩大小区面积的方法所提高的试验精确度不如增加重复次数来的好。这是因为增加重复次数可以使试验地各种肥力条件比较全面地包括在每一处理内，而且标准误（样本平均数的标准差）也随着重复数的增加而减小。

二、随机化

在试验中,各个区组的排列以及在每个区组内各个处理小区的排列都应该实现随机化,而不能按一定顺序或者试验者的主观愿望来排列。随机排列可以使各处理在试验中占居任何一个小区的机会均等,以消除系统误差,确保试验得到的处理平均数和试验误差是无偏估计值。

随机化是应用数理统计学方法分析试验数据的前提。所以,严格地讲,如果在试验设计中没有执行随机化的原则,那么将来对试验结果进行类似方差分析之类的统计学分析是没有意义的。

三、局部控制

试验表明,在一定范围内相邻小区间的土壤肥力差异往往较小,超过这个范围则差异较大。所以进行田间设计时,与其把所有处理的所有重复都完全随机地排列在试验地上,不如把试验地按照土壤肥力的高低分为几个区组,在每个区组内安排所有处理的一次重复。这样,由于在区组内土壤肥力差异较小,各个处理互相比较的准确性较高,误差较小。这种把要比较的处理设置在土壤均匀的局部地段,以便减少试验误差的原则叫作局部控制。

第3节 林业试验设计的误差控制

广义的试验误差包括错误、系统误差和偶然误差三类。所谓错误,是指量错了土地面积、算错了产量、记错了数字、混杂了品种等等。这些问题往往降低试验的精确性,甚至使整个试验报废。在试验过程中,应该十分细心谨慎,反复核对,避免错误的发生。所谓系统误差,是指数学期望有偏的误差,它是由于试验材料或量测方法不一致,以及处理没有随机化等原因造成的误差。一般来说,系统误差也可以通过合理的设计和精细的工作得到消除。偶然误差即狭义的试验误差,它是指除去处理本身之外,由于环境条件的微小变异、个体间的差异、量测时不可避免的观测误差、抽样误差等偶然性因素造成的对真实效应的影响。这一类误差的数学期望是无偏的,它具有随机性质,无论怎样小心谨慎也不可能完全消除。这种误差的大小是衡量试验精确性的依据,偶然误差越大,说明试验精确性越低。如果偶然误差大到与真实效应不相上下的程度,就很难通过试验得出正确的结论。因此,应尽可能采取各种措施减少这部分试验误差,提高试验的精确性。这些措施,主要可以归纳为试验地选择和小区技术两个方面。

一、试验地选择

为了使试验结果具有典型性,以便将来推广,试验地应该选在气候、土壤等条件有代表性的地段,即试验地应具有试验地区的典型气候、典型地势、典型土壤肥力和理化性状、典型的地下水位等。

试验地一般应设在同一坡向上,坡度也应基本一致,而不应该把试验地选在沟谷中或山顶上。为了避免自然条件的影响和人畜损坏,试验地不宜太靠近村庄、道路和建筑物。

二、小区技术

1. 小区的面积与形状

小区面积的大小与试验误差的大小直接相关。一般来说，小区面积较大时，小区间土壤肥力较均匀，试验较精确。但小区面积大小与区组大小又有关系。根据局部控制原则，区组面积不宜过大，因而小区面积也不能太大。

林业试验小区的大小，目前尚无统一规定。应根据试验种类、栽培年限、重复次数、土壤肥力差异程度等综合考虑决定。一般来说，树木试验以每小区栽植 1—10 株效率最高，在欧洲也有人主张安排更大的小区。

小区的形状一般采用长方形，有时也用正方形。当试验地肥力呈方向性变化时，采用长方形小区更为适宜。此时，区组长边的方向应与肥力变化方向相垂直，而区组内小区的长边方向应与肥力变化方向相平行。在山地进行试验时，区组的长边应与坡向垂直即平行于等高线，小区的长边则与坡向平行即垂直于等高线。这样安排，可以使区组内各小区间的肥力基本一致。一般来说，小区的长边与宽边的比例以 3:1—10:1 为宜。

2. 小区的重复数

设置重复是试验设计的基本原则之一，有人将重复原理称为现代田间试验的基本原理。因为只有设置重复才能客观地比较各个处理的效果，提高试验的精确度。

试验所需要的小区重复数，主要决定因素是土壤差异程度和试验所要求的精度。其次，还要综合考虑处理的数目、试验地面积、小区面积、小区排列方式等。一般来说，在土壤差异不大的地区安排随机区组设计需要 4—6 次重复；土壤差异较大或要求精度较高的试验可安排 8—10 次重复。比较粗放的试验或小区面积较大时也可安排 2—3 次重复。从统计学的观点看，进行方差分析时误差项的自由度最好不小于 12，否则 F 值很大，难于检验出处理的显著性。

如果事先了解了试验的变异系数，则可根据精确度所要求的标准误，参阅表 1—1 查出应有的重复次数。

表 1—1 不同重复次数和变异系数下的标准误 (表内数字: $\frac{\text{标准误}}{\text{平均数}} \times 100$)

变异系数(%) \ 重复次数	8	10	12	14
2	5.7	7.1	8.5	9.9
3	4.6	5.8	6.9	8.1
4	4.0	5.0	6.0	7.0
5	3.6	4.5	5.4	6.3
6	3.3	3.8	4.5	5.3
7	3.0	3.8	4.5	5.3
8	2.8	3.5	4.2	5.0

(引自马育华著《试验统计》P22)

例如，某试验具有 10% 的变异系数，要求处理平均数的标准误不大于 5%，则重复 4 次即可满足要求。

3. 设置对照小区

对照小区又称标准区，用来安排对照处理。设置对照的目的，一方面是使参加试验的各个处理都能以对照为标准进行比较，以鉴定其优劣；另一方面是利用对照矫正和估计试验地的土壤差异。

选择适宜的对照是试验设计的关键之一。如果对照设置不当或者不设对照，则试验结果的应用价值往往不大。例如进行引种试验，必须有当地的优良树种作对照才能评价引种的效果；进行不同施肥量的对比试验，必须有不施肥的小区为对照以鉴别施肥的效果。

4. 设置保护行

靠近试验地外侧的植株，通风透光条件好，长势旺，有明显的边际效应。为了减少边际效应所引起的试验误差，同时也为了防止人畜破坏践踏，应在试验地周围设置保护行。区组或小区之间一般不设保护行。

对保护行中的树木，应采取与试验地相同的品种和管理措施，必要时可以从保护行中采集标本和选取解析木。

习 题

1. 何为试验设计？林业试验设计的特点是什么？
2. 试验为什么要设置重复？确定重复数要考虑哪些因子？
3. 什么叫系统误差？为什么说没有经过随机化的试验设计其结果不能进行方差分析？
4. 当试验地的肥力有一个方向的变化时，区组和小区应如何安排？为什么？



第2章 林业试验设计的 数学模型和方差分析

第1节 设计种类及其线性数学模型

一、林业试验设计的种类

根据参加试验因子的多少,可以将所有林业试验设计分为单因子和多因子两大类,在多因子试验中根据因子间的相互组合方式,又可分为交叉式、巢式和混合式三种。无论哪一种试验设计,每一个观测值都可以写成该试验的有关因子的不同效应值及其交互效应值的线性函数式。

1. 单因子试验

只含有一个因子的试验称为单因子试验。该因子的若干个处理叫作水平。比如,为了研究油松半同胞家系在抗寒性方面有无显著差异,我们从10株油松上采种,分别培育出10个半同胞家系,在每个半同胞家系中随机抽取6株进行抗寒性试验,以作出抗寒性强弱的判断。这里,半同胞家系是该试验的因子,10个半同胞家系是该试验的10个水平,而每个家系观测6株是该试验的重复数。

对于单因子试验,任一观测值都可以用下列线性模型加以描述:

$$x_{ij} = \mu + \alpha_i + e_{ij}$$

式中, x_{ij} 表示第 i 个处理的第 j 个观测值, μ 为观测值的数学期望, α_i 表示第 i 个处理的效应值(即 i 处理的平均值距总平均值的离差), e_{ij} 为随机误差(即第 i 个处理第 j 个个体的观测值与第 i 处理平均值的离差)。

2. 多因子交叉式分组试验

如果在一个试验中同时考查两个以上的因子,即为多因子试验。比如,我们可以通过一次试验考查某一地区最适宜的毛白杨品种和最合理的栽植密度。这里,品种和密度就是本试验要考查的两个因子,它们各有若干个水平。如果我们在试验中使这两个因子的各个水平都能两两相遇,使这两个待考查的因子处于完全平等的地位上,这样进行分组的试验就叫作交叉式分组试验。

双因素交叉式分组试验的线性模型为:

$$x_{ijk} = \mu + \alpha_i + \beta_j + \alpha\beta_{ij} + e_{ijk}$$

式中, α_i 为 A 因子第 i 水平的效应值, β_j 为 B 因子第 j 水平的效应值, $\alpha\beta_{ij}$ 为它们的交互效应值。

类似上式,我们也可以写出三因子以至更多因子试验的线性模型。

3. 多因子巢式分组试验

在多因子试验中,如果因子间的各个水平不能两两相遇,而是一个因子从属于另一个

因子,处于下级的因子的各个水平因上级因子的不同水平而变化,这样进行分组的试验则叫作巢式分组试验,也叫套式或系统分组试验。

比如,当我们进行某一树种的种源试验时,先在该树种全分布区选定若干个种源,在每个种源中各选取若干个林分,分别在这些选定的林分中采种,再将由这些种子育出的后代苗木栽于某地点进行比较试验,以鉴别不同种源间以及同一种源内不同林分间有无显著差异。在这个试验中,林分这一因子是从属于种源这一因子的,第1个种源的第1个林分并不等于第2个种源的第1个林分。种源属于上一级因子,林分属于下一级因子。

双因素巢式分组试验的线性模型为:

$$x_{ijk} = \mu + \alpha_i + \beta_{j(i)} + e_{ijk}$$

这里, $\beta_{j(i)}$ 表示从属于 A 因子第 i 水平的 B 因子第 j 水平的效应值。因为二者是从属关系,所以不存在交互效应。类似上式,也可以写出三因子乃至更多因子巢式分组试验的线性模型。

以上,我们写出了几种设计的最简单模型,但没有给出任何限制条件。实际上,任何数学模型都是有条件的,这些条件又因各因素效应值的类型不同而异。下面就来讨论这些类型及其限制条件。

二、方差分析模型

依据各种试验设计中不同因子的效应值类型,可以将方差分析模型分为三类,即固定模型、随机模型和混合模型。

1. 固定模型

如果在试验中要考察的某因素的各个水平是特意选择的,对试验结果进行分析所得到的结论也只限于原来设计的那几个水平,并不扩展到未加试验的其它水平上,那么我们要考察的这个因素就叫作固定因素,处理固定因素的数学模型叫作固定效应模型,简称固定模型。固定模型试验的目的是找出各因素各水平的最佳组合。比如,我们在某地进行白榆 50 个无性系的对比试验,如果我们的目的就是通过试验寻找这 50 个无性系中有哪几个最适宜在当地栽培推广,而不涉及白榆的其它无性系,则应将白榆无性系视作固定因素,分析它的数学模型就是固定模型。

固定模型的限制条件是处理效应值的总和为零。

对于单因素试验,线性模型为:

$$x_{ij} = \mu + \alpha_i + e_{ij}$$

其限制条件为处理效应值总和 $\sum_i \alpha_i = 0$, 进行方差分析时的零假设为:

$$H_0: \alpha_1 = \alpha_2 = \cdots = 0。$$

对于双因素交叉分组试验,线性模型为:

$$x_{ijk} = \mu + \alpha_i + \beta_j + \alpha\beta_{ij} + e_{ijk}$$

其限制条件为 $\sum_i \alpha_i = 0$, $\sum_j \beta_j = 0$, $\sum_i \alpha\beta_{ij} = \sum_j \alpha\beta_{ij} = 0$ 。

进行方差分析时的零假设为:

$$H_0: \alpha_1 = \alpha_2 = \cdots = 0$$

$$\beta_1 = \beta_2 = \cdots = 0$$

$$\alpha\beta_{11} = \alpha\beta_{12} = \cdots = \alpha\beta_{ij} = 0$$

除了上述限制条件之外，无论采用哪种设计，都要求随机误差是独立的随机变量，并遵从正态分布 $N(0, \sigma^2)$ 。只有满足了所有这些限制条件，才能够进行方差分析。

2. 随机模型

如果在试验中要考察的某因素的若干水平不是特意选定的，而是从该因素的水平总体中随机抽取的样本，将来对试验结果的分析也是通过这些样本去对该因素的水平总体作出推断，那么我们要考察的这个因素就叫作随机因素，处理随机因素的数学模型就叫作随机效应模型，简称随机模型。随机模型试验的目的是估计总体的变异度（判断总体内有无差异以及估算遗传参数等）。我们仍以在固定模型中讨论过的白榆 50 个无性系对比试验为例，如果我们的试验目的并不在于从这 50 个无性系中选取当地最适宜栽培的无性系，而在于探索白榆无性系这个总体内是否存在差异，或者说我们考察的对象不是无性系效应值本身，而是无性系效应值的变异程度即方差，那么，就应将白榆无性系视作随机因素，分析它的数学模型就属随机模型。所以，同一个试验中的同一个因素可以因试验目的和分析角度而改变数学模型。如果我们上述试验的目的既要作白榆无性系变异度的分析，又要同时选出最适宜的无性系，那么对试验观测值就应该同时从随机模型和固定模型两个角度去进行分析。

随机模型的限制条件是处理效应为遵从正态分布的随机变量。

对于单因素试验：

$$x_{ij} = \mu + \alpha_i + e_{ij}$$

其限制条件为：

$$\alpha_i \sim N(0, \sigma_a^2)$$

进行方差分析时的零假设为：

$$H_0: \sigma_a^2 = 0$$

对于双因素交叉分组试验：

$$x_{ijk} = \mu + \alpha_i + \beta_j + \alpha\beta_{ij} + e_{ijk}$$

其限制条件为：

$$\alpha_i \sim N(0, \sigma_a^2)$$

$$\beta_j \sim N(0, \sigma_b^2)$$

$$\alpha\beta_{ij} \sim N(0, \sigma_{a\beta}^2)$$

进行方差分析时的零假设为：

$$H_0: \sigma_a^2 = 0, \quad \sigma_b^2 = 0, \quad \sigma_{a\beta}^2 = 0$$

对于双因素巢式分组试验：

$$x_{ijk} = \mu + \alpha_i + \beta_{j(i)} + e_{ijk}$$

其限制条件为：

$$\alpha_i \sim N(0, \sigma_a^2)$$

$$\beta_{j(i)} \sim N(0, \sigma_{\beta(\alpha)}^2)$$

进行方差分析时的零假设为：

$$H_0: \sigma_a^2 = 0, \quad \sigma_{\beta(a)}^2 = 0$$

与固定模型一样，随机模型也要求试验的随机误差是独立的随机变量，并遵从正态分布 $N(0, \sigma^2)$ 。否则也不能进行方差分析。

3. 混合模型

在多因素试验中，如果既有固定因素，又有随机因素，那么处理这些因素的数学模型就叫作混合模型。

例如，在双因素交叉式分组试验中，如果一个因素 A 是固定因素，另一个因素 B 是随机因素，则应建立混合模型。其线性模型是：

$$x_{ijk} = \mu + \alpha_i + \beta_j + \alpha\beta_{ij} + e_{ijk}$$

其限制条件是：

$$\sum \alpha_i = 0$$

$$\beta_j \sim N(0, \sigma_\beta^2)$$

$$\alpha\beta_{ij} \sim N\left(0, \frac{a-1}{a} \sigma_{a\beta}^2\right)$$

进行方差分析时的零假设为：

$$H_0: \alpha_1 = \alpha_2 = \cdots = 0$$

$$\sigma_\beta^2 = 0$$

$$\sigma_{a\beta}^2 = 0$$

第 2 节 不同试验设计的方差分析

一、单因素试验

对于单因素试验，采用固定模型或随机模型进行方差分析的程序完全一样。设该因素有 k 个水平，每水平重复 n 次，则方差分析如表 2-1 所示。

表 2-1 单因素试验方差分析

变异来源	自由度	平方和	均方	F	EMS (固定、随机)
处 理 间	$k-1$	SS_A	V_A	V_A/V_e	$\sigma^2 + n\sigma_a^2$
机 误	$k(n-1)$	SS_e	V_e		σ^2
总 计	$kn-1$	SS_T			

表中： $C = \frac{1}{kn} \sum x_i^2$

$$SS_A = \frac{1}{n} \sum x_i^2 - C$$

$$SS_T = \sum \sum x_{ij}^2 - C$$

$$SS_e = SS_T - SS_A$$

〔例 2.1〕 欲比较毛白杨 4 个无性系的生长量，每个无性系随机抽查 3 株，结果如表 2-2，试判断这 5 个无性系间是否存在差异。

表 2-2 4 个毛白杨无性系试验结果

无性系	x_{ij}			$x_{i\cdot}$	$x_{\cdot\cdot}$
A	2	4	9	15	120
B	6	7	11	24	
C	11	13	15	39	
D	12	12	18	42	

【解】 先按公式求各项离差平方和

$$C = \frac{1}{kn} x_{\cdot\cdot}^2 = \frac{1}{4 \times 3} \times 120^2 = 1200$$

$$SS_T = \sum_i \sum_j x_{ij}^2 - C = 2^2 + 4^2 + \cdots + 18^2 - 1200 = 234$$

$$SS_A = \frac{1}{n} \sum_i x_{i\cdot}^2 - C = \frac{1}{3} (15^2 + 24^2 + 39^2 + 42^2) - 1200 = 162$$

$$SS_e = SS_T - SS_A = 234 - 162 = 72$$

根据上述结果列成方差分析表 2-3。

表 2-3 毛白杨无性系试验方差分析

变异来源	自由度	平方和	均方	F
无性系	3	162	54	6*
机 误	8	72	9	
总 计	11	234		

$F_{0.05}(3,8) = 4.07$ 因 $F > F_{0.05}$, 所以无性系间差异显著。

二、双因素交叉式分组试验的方差分析

对于 A、B 两个因素的交叉式分组试验, 设 A 因素有 a 个水平, B 因素有 b 个水平, 每个交叉点有 n 次重复, 则因采用的数学模型不同而有不同的期望均方和 F 值计算式。其方差分析如表 2-4 所示。

表 2-4 双因素交叉式分组试验方差分析

变异来源	自由度	平方和	均方	固定模型		随机模型		混合模型 (A 固定、B 随机)	
				EMS	F	EMS	F	EMS	F
因素 A	$a-1$	SS_A	V_A	$\sigma^2 + nb\sigma_A^2$	$\frac{V_A}{V_e}$	$\sigma^2 + n\sigma_{AB}^2 + nb\sigma_A^2$	$\frac{V_A}{V_{AB}}$	$\sigma^2 + n\sigma_{AB}^2 + nb\sigma_A^2$	$\frac{V_A}{V_{AB}}$
因素 B	$b-1$	SS_B	V_B	$\sigma^2 + na\sigma_B^2$	$\frac{V_B}{V_e}$	$\sigma^2 + n\sigma_{AB}^2 + na\sigma_B^2$	$\frac{V_B}{V_{AB}}$	$\sigma^2 + na\sigma_B^2$	$\frac{V_B}{V_e}$
交互 AB	$(a-1)(b-1)$	SS_{AB}	V_{AB}	$\sigma^2 + n\sigma_{AB}^2$	$\frac{V_{AB}}{V_e}$	$\sigma^2 + n\sigma_{AB}^2$	$\frac{V_{AB}}{V_e}$	$\sigma^2 + n\sigma_{AB}^2$	$\frac{V_{AB}}{V_e}$
机 误	$ab(n-1)$	SS_e	V_e	σ^2		σ^2		σ^2	
总 计	$abn-1$	SS_T							

表中: $C = \frac{1}{abn} x_{...}^2$

$$SS_T = \sum_i \sum_j \sum_k x_{ijk}^2 - C$$

$$SS_A = \frac{1}{bn} \sum_i x_{i..}^2 - C$$

$$SS_B = \frac{1}{an} \sum_j x_{.j.}^2 - C$$

$$SS_{AB} = \frac{1}{n} \sum_i \sum_j x_{ij.}^2 - C - SS_A - SS_B$$

$$SS_e = SS_T - SS_A - SS_B - SS_{AB}$$

关于各种模型的期望均方，在数理统计学中都有严格的推导。从应用的角度出发，我们总想寻找一种简便而准确的方法直接写出期望均方。美籍华人孔繁浩教授找到了这种方法。他的方法适用于随机模型，而随机模型正是我们林业试验设计中使用最广泛的模型。

这种方法大体可分为以下 5 步：

1. 第 1 列写出变异来源；
2. 第 1 行写出对应的均方成分；
3. 所有变异来源都含有误差项，系数为 1；
4. 在对角线上写出对应的方差成分，其系数为变异来源中没有出现的系数的乘积；
5. 从对角线往上看，若均方成分的下标全部包括了变异来源，就照抄；否则就划掉这一项。

例如，2 个母本各与 3 个父本交配，子代苗重复 4 次，则全部观测值个数即全部系数的乘积为： $\frac{F \times M \times n}{2 \times 3 \times 4} = 24$ ，按上述方法，可以直接写出该试验的期望均方如表 2-5。

表 2-5 随机模型期望均方例

变 异 来 源	EMS			
	σ_e^2	σ_{FM}^2	σ_M^2	σ_F^2
母本 (F)	σ_e^2	$4\sigma_{FM}^2$	—	$12\sigma_F^2$
父本 (M)	σ_e^2	$4\sigma_{FM}^2$	$8\sigma_M^2$	
交互 (FM)	σ_e^2	$4\sigma_{FM}^2$		
机误 (e)	σ_e^2			

这种方法不仅适用于双因素试验，也适用于更多因素的试验；不仅适用于交叉式分组设计，也同样适用于巢式设计以及交叉式与巢式复合分组的设计。这样，就使各种试验设计的方差分析工作大大简化了。

写出期望均方的主要目的是为了正确地进行 F 检验。计算某一变异来源的 F 值时，应以该变异来源的均方值为分子；分母项的期望均方应该比分子项少一种均方成分即要检验的均方成分。比如上例中母本效应的期望均方是 $\sigma_e^2 + 4\sigma_{FM}^2 + 12\sigma_F^2$ ，其中 $12\sigma_F^2$ 是要检验的均方成分，因此 F 检验的分母项的期望均方应为 $\sigma_e^2 + 4\sigma_{FM}^2$ ，对应于这一期望均方的变异来源是父母本交互作用项，故检验母本效应的 F 值计算式为：

$$F = \frac{MS_F}{MS_{FM}}$$

同理，检验父本效应的 F 值计算式为：

$$F = \frac{MS_M}{MS_{FM}}$$

检验父本、母本交互效应的 F 值计算式为：

$$F = \frac{MS_{FM}}{MS_e}$$

〔例 2.2〕 某杂交试验，设有 2 个母本、3 个父本，子代苗重复 2 次，结果如表 2-6，试判断母本、父本及交互作用的差异显著性。

表 2-6 杂交试验结果

x_{ijk} 父 本 \ 母 本	1	2	$x_{.j.}$
1	1, 3	9, 7	20
2	5, 3	3, 5	16
3	5, 7	5, 7	24
$x_{i..}$	24	36	$x_{...} = 60$

【解】 根据表 2-4 计算式计算离差平方和：

$$C = \frac{1}{abn} x_{...}^2 = \frac{1}{2 \times 3 \times 2} \times 60^2 = 300$$

$$SS_T = \sum_i \sum_j \sum_k x_{ijk}^2 - C = 1^2 + 3^2 + \cdots + 7^2 - 300 = 56$$

$$SS_A = \frac{1}{bn} \sum_i x_{i..}^2 - C = \frac{1}{3 \times 2} (24^2 + 36^2) - 300 = 12$$

$$SS_B = \frac{1}{an} \sum_j x_{.j.}^2 - C = \frac{1}{2 \times 2} (20^2 + 16^2 + 24^2) - 300 = 8$$

$$SS_{AB} = \frac{1}{n} \sum_i \sum_j x_{ij.}^2 - C - SS_A - SS_B = \frac{1}{2} (4^2 + 16^2 + \cdots + 12^2) - 300 - 12 - 8 = 24$$

$$SS_e = SS_T - SS_A - SS_B - SS_{AB} = 56 - 12 - 8 - 24 = 12$$

把计算结果列成方差分析表如表 2-7。

试验分析结果，固定模型时母本及母×父间差异显著，随机模型时只有母×父交互作