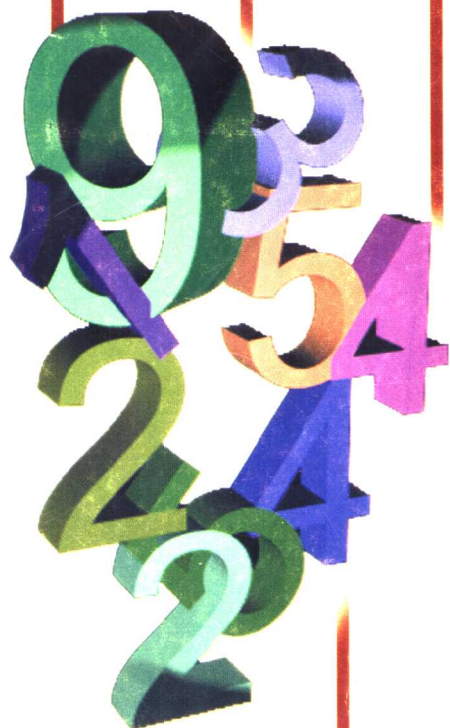


Ecological Mathematics Method for Studies on Forest Pathology

森林病理学研究的生态数学方法

张星耀 编著



中国林业出版社

国家林业局图书出版基金 资助出版
中国林业科学研究院人才培养基金

森林病理学研究的 生态数学方法

张星耀 编著

BBA16/25

中国林业出版社

图书在版编目(CIP)数据

森林病理学研究的生态数学方法/张星耀编著. —北京:
中国林业出版社, 1999. 6

ISBN 7-5038-2285-6

I. 森… II. 张… III. 生态学:数学-应用-森林-病理
学-研究 N. S763.1

中国版本图书馆 CIP 数据核字(1999)第 12426 号

出版 中国林业出版社

(100009 北京市西城区刘海胡同 7 号)

发行 新华书店北京发行所

印刷 北京地质印刷厂

版次 1999 年 6 月第 1 版

印次 1999 年 6 月第 1 次印刷

开本 850mm×1168mm 1/32

印张 7.625

字数 194 千字

印数 1~800 册

定价 18.00 元

前 言

1982 年从北京林业大学前身北京林学院毕业后，分配至西北林学院任教，并师从景耀先生和杨俊秀先生开始森林病理学的研究工作。1988 年出国留学前，初成了这本书的一稿，此后又进行了二稿，但由于无出版费而始终未能在国内或国外出版。1996 年回国后就职于中国林业科学研究院森林保护研究所，重拮已有些发黄的手稿，感到陈旧了许多。感谢当时的所长周淑芷先生和副所长张锡津先生鼓励我修改出版此书，使我有勇气又进行了三稿，并作为讲稿为北京林业大学的研究生课程进行了试讲。忙忙碌碌加之出版经费的筹集，拖延至今，这本书终于出版了。如果这本书对读者有所意义，作者则深感欣慰。

这本书分四章。第一章，是基础部分，叙述森林病理学研究的生态数学方法的基本技术和基本分析方法；第二章，以作者对杨树烂皮病和杨树溃疡病以及松材线虫病的研究结果为基础，叙述森林病害空间分布的研究方法，其中包括空间分布格局的研究方法和生态地理分布图式的研究方法；第三章，以作者对华山松疱锈病的研究结果以及对其它一些森林病害生态学数据的分析为基础，叙述森林病害时间动态的研究方法，其中包括时间动态的模拟及预测和数量动态相对静止状态的模拟及预测；第四章，以作者对杨树烂皮病和溃疡病损失估测和损失过程模拟及允许水平的研究结果为基础，叙述森林病害允许水平和管理决策的研究方法。书中的数学方法涉及概率论和数理统计以及模糊数学和灰色系统理论。

有很多人关心这本书的出版，其中有原林业部副部长刘于鹤

先生，全国政协委员、森林保护学专家陈昌洁先生，我永远的所长周淑芷先生和张锡津先生以及许多同事们，我从心里感谢他们。出版，得到了国家林业局出版基金和中国林业科学研究院人才培养基金的资助，谨致诚挚的感谢。

在此，我应该感谢我的妻子贾秀贞同志。从一稿的手写稿至终稿的计算机稿，大量的计算、稿件的誊写和计算机输入、校对、编排，都是由她完成的。

这本书耽搁 10 年，显有陈旧，书中的错误和不足肯定不少，敬请读者批评指正。

张星耀

1999 年 2 月 15 日于北京

目 录

前言

第一章 基本技术和基本分析	(1)
1.1 抽样技术	(1)
1.1.1 典型抽样	(1)
1.1.2 机械抽样	(1)
1.1.3 随机抽样	(1)
1.2 试验设计的原则和技术	(3)
1.2.1 试验设计的原则	(3)
1.2.2 试验设计技术	(4)
1.3 原始数据的整理与分析	(5)
1.3.1 平均数	(5)
1.3.2 变异数	(7)
1.4 差异分析	(12)
1.4.1 t 检验和 u 检验	(12)
1.4.2 F 检验	(14)
1.4.3 χ^2 检验	(20)
1.4.4 多重差异分析	(20)
1.5 正交和 BIB 设计及分析	(25)
1.5.1 正交设计的森林病理学意义	(25)
1.5.2 正交表	(26)
1.5.3 正交试验的分析	(28)
1.5.4 BIB 设计及分析	(37)

1.6 因子及关系分析	(44)
1.6.1 相关系数	(44)
1.6.2 灰色关联度分析	(46)
1.6.3 列联表分析	(48)
1.6.4 判别分析	(53)
1.6.5 通径分析	(55)
第二章 森林病害空间分布的研究方法	(59)
2.1 空间分布格局的研究方法	(59)
2.1.1 空间分布格局的测定方法	(59)
2.1.2 测定空间分布格局的森林病理学意义	(62)
2.1.3 空间分布格局研究的程序和实例	(65)
2.1.4 森林病害传播空间图式的识别和传播梯度的模拟	(70)
2.2 地理分布图式的研究方法	(75)
2.2.1 生态地理分布环境的识别	(75)
2.2.2 生态地理分布环境梯度的识别	(88)
2.2.3 地理分布图式研究的实例	(101)
第三章 森林病害时间动态的研究方法	(105)
3.1 时间动态的模拟及预测	(105)
3.1.1 病害发展阶段的模拟及预测	(106)
3.1.2 病害累积发展过程的模拟及预测	(117)
3.1.3 模型的检验	(122)
3.1.4 病害多波动过程的模拟及预测	(123)
3.2 数量动态相对静止状态的模拟及预测	(153)
3.2.1 回归分析	(153)
3.2.2 隶属函数模型	(171)

3.2.3 模糊状态聚类分析	(183)
第四章 森林病害允许水平和管理决策的研究方法	(187)
4.1 森林病害造成损失的定义和模式及其测定	(187)
4.1.1 损失的定义	(187)
4.1.2 损失的模式	(189)
4.1.3 损失估测的数学模型	(190)
4.1.4 损失估测的野外测定	(192)
4.2 允许水平的有关定义及其数学模型	(193)
4.2.1 允许水平的有关定义	(193)
4.2.2 允许水平的有关数学模型	(195)
4.2.3 研究实例：杨树烂皮病和溃疡病	(198)
4.3 森林病害控制管理的决策方法	(211)
4.3.1 森林病害控制决策的数学原理	(211)
4.3.2 森林病害管理的灰色多目标局势决策	(214)
4.3.3 森林病害管理的模糊多准则决策	(225)
参考文献	(235)

第一章 基本技术和基本分析

1.1 抽样技术

抽样技术可分为 3 大类：典型抽样、机械抽样和随机抽样。随机抽样又可分为简单随机抽样、分层随机抽样、整群随机抽样以及混合随机抽样。在森林病理学研究的抽样过程中，常将几种抽样方法混合使用。

1.1.1 典型抽样

典型抽样是指根据研究目的，从抽样总体内有意识、有目的地选取有代表性的典型抽样单位或单位群。所选取的单位应具有总体的代表性。典型抽样获得的样本称为典型样本。森林病虫害生态学观测设立标准木或标准地时，往往采用这种抽样技术。但典型抽样技术不具有随机性，所以无法估计抽样误差。

1.1.2 机械抽样

机械抽样是指根据某种即定的顺序抽取一定数量抽样单位构成样本。如在林地内按对角线抽取样木、按每隔一行或数行抽取一行样木等抽样方式都属于机械抽样。机械抽样获得的样本称为机械样本，其抽样误差也难以估计。

1.1.3 随机抽样

【1】简单随机抽样 是指抽样单位是直接从总体中随机抽取的，样本容量固定，且每一个样本都有被抽取的同等概率。这种

抽样技术可以得到抽样误差的无偏估计,同时可得到总体平均数、总和或成数的估计值。

【2】分层随机抽样

(1) 根据变异原因将被研究总体分为较为均匀同质的若干部分,这些部分称为区层;

(2) 独立地从每一区层内随机抽取所确定的抽样单位数目,计算各区层的样本平均数、总和或成数;

(3) 根据各区层的估计值采用加权法估计总体真值。

各区层的抽样单位数目,可参照下述两类方法配置。①相等配置法:各区层的抽样单位数目相等;②比例配置法:即按区层所包含的抽样单位数目,用同一比例计算应抽取的单位数目。

例如,有3个面积不相等的区层,抽样单位为1株林木,第一区层包含120个抽样单位,第二区层包含200个抽样单位,第三区层包含80个抽样单位。若按5%的比例抽样,则第一区层应抽 $120 \times 5\% = 6$ (株)林木,第二区层应抽 $200 \times 5\% = 10$ (株)林木,第三区层应抽 $80 \times 5\% = 4$ (株)林木。

森林病害在野外的分布,绝大多数是非随机性的,往往具有一定的分布梯度。因此,在不同的梯度区层,根据比例与变异性进行抽样,其结果将比简单随机抽样精确。

【3】整群随机抽样 如果所抽样本不是从总体直接抽选各个抽样单位,而是随机抽选单位整群,然后在每一群内进行全部抽样单位的调查观测,这种抽样方法则称为整群抽样,这是森林病害野外调查时常用的一种抽样技术。如随机抽取100株林木构成的林分作为样地(整群)、各样地(群)内每木检尺测其树高、胸径、发病等级。

和简单随机抽样相比,整群抽样能提供更为精确的总体估计值。当病害在林间分布为非随机性时,采用扩大调查单位面积的整群抽样技术,可使调查结果更接近真实。若各群包含相等数目

的抽样单位,则此法可获得总体平均数、总和或成数的无偏估计。

【4】混合随机抽样 在森林病害的野外生态学观测、调查过程中,常将几种抽样方法混合使用。

笔者等(1988)在调查杨树烂皮病(*Valsa sordida* Nit)在张掖地区的病情和材积损失时,曾采用了下述的混合抽样技术。首先在地图上采取机械抽样,等距离地抽取样区,然后在各样区采取典型整群抽样技术设置样地,样地内每木实测树高、胸径、病斑长、宽度。

观测、调查叶部病害时,可采用如下的混合抽样技术。笔者等(1983)在观测、调查秦岭南坡引种的落叶松早期落叶病(*Mycosphaerella larici-leptolepis* Ito et al.)时,首先采用典型整群抽样技术在林区设置样地,在样地内采用五点式机械抽样抽取5个样株。由于受害针叶在整个树冠分布的不均匀性,进一步采取分层抽样将样株树冠按上、中、下、东、南、西、北划分为12个抽样区层,在各区层内随机抽取相等数量的若干簇针叶,调查针叶的发病情况。

1.2 试验设计的原则和技术

1.2.1 试验设计的原则

为了控制差异、降低试验、观测、调查误差,在进行试验、观测、调查设计时应遵循如下原则。

重复 即同一处理必须有两个以上的重复,其目的是估计试验、观测、调查误差,提高试验、观测、调查的精度。

随机 设置重复虽然提供了估计误差的条件,但是为了获得无偏的试验、观测、调查误差估计值,则要求每一个处理都有相等的机会设置在任何一个试验小区上。因此,必须采取随机设计来满足这个要求。

局部控制 即分范围地控制非处理因素和环境,使其对各处理的影响趋于最大程度的一致。

设置保护行 为了使野外试验、观测、调查能在较为均匀、稳定的环境条件下进行,应在试验、观测、调查点周围设置保护行。其目的是保护试验、观测、调查体不受外来因素的干扰,防止边缘效应,使各处理间能合理客观地比较。设置保护行的意义和局部控制的意义是一致的。

设置对照 为了比较和校正试验、观测、调查的结果,必须设置对照。常以 CK 表示。

1.2.2 试验设计技术

对比设计 对比设计常用于少数处理的比较试验,其排列特点是每一处理均直接排列于对照区旁边,即每隔 2 个处理设 1 个对照,使每一个处理可与其邻旁的对照直接比较。图 1-1 是 3 个处理 2 次重复的对比试验设计。

保 护 行	CK			CK			CK			CK	保 护 行
	1	A	B	2	C	A	3	B	C	4	

图 1-1 对比设计

随机设计 随机设计根据局部控制的原则,将试验地按异质性程度划分为等于重复次数的区组,在各区组内,各处理小区借助随机数字发生器或其它方法随机排列。如施用 A、B、C,3 种杀菌剂防治苗圃杨树溃疡病,3 个重复的试验,可采用图 1-2 的设计。

保 护 行	A	B	C	CK	B	CK	A	C	C	A	B	CK	保 护 行
-------------	---	---	---	----	---	----	---	---	---	---	---	----	-------------

图 1-2 随机设计

拉丁方设计 拉丁方设计根据试验处理项目划分小区，处理项目和重复次数相等，每种处理无论在纵行或横行内只出现 1 次。图 1-3 是 3 个处理 A、B、C，1 个对照 D，4 个重复的试验设计。

保 护 行				保 护 行
保 护 行	A	B	C	
	B	A	D	
	D	C	B	
	C	D	A	
保 护 行				

图 1-3 拉丁方设计

其它试验设计技术 上述 3 种方法是最常用的试验设计技术，另外还有一些较为复杂的诸如正交设计、BIB 设计等，这些将在本章中结合基本分析叙述。

1.3 原始数据的整理与分析

1.3.1 平均数

平均数是试验、观测、调查变量中最有代表性的数据，是其结果的集中体现。

【1】算术平均数 ($\bar{X}$)

(1) 直接计算 将一组数据的各个数值逐个相加，再除以数值的总次数，即

$$\bar{X} = \frac{X_1 + X_2 + \cdots + X_n}{n} = \frac{\sum_{i=1}^n x_i}{n} \quad (1-1)$$

[例 1—1] 对渭河林业试验站 10 个标准地的北京杨溃疡病 (*Dothiorella gregaria* Sacc) 的病情指数进行测定，其结果如表 1-1，求其平均病情指数 $\bar{X}$ 。

表 1-1 测定结果

标准地	1	2	3	4	5	6	7	8	9	10
病情指数(x_i)	64.8	67.5	53.0	66.0	64.0	56.0	72.9	50.0	65.5	54.0

根据表 1-1 的原始数据,

$$\bar{X} = \frac{64.8 + 67.5 + \cdots + 65.5 + 54.0}{10} = \frac{\sum_{i=1}^{10} x_i}{10} = 61.37$$

(2) 加权计算 当试验、观测、调查的若干个结果在整个试验、观测、调查中具有不同程度的比重时, 我们称其为具有不同的权重。此时, 应采用加权计算求其平均数。公式如下:

$$\bar{X} = \frac{f_1 x_1 + f_2 x_2 + \cdots + f_n x_n}{f_1 + f_2 + \cdots + f_n} = \frac{\sum_{i=1}^n f_i x_i}{\sum_{i=1}^n f_i} \quad (1-2)$$

式中: f_i —— x_i 所具有的权重。

[例 1-2] 对渭河林业试验站杨树无性系试验对比林的溃疡病发病率进行测定, 其结果: 52 号杨发病率 94.0%, 57 号杨发病率 20.0%, 65 号杨发病率 26.5%。试验对比林内, 52 号杨 50 株, 57 号杨 200 株, 65 号杨 100 株。

若采用直接计算, 对比试验林的平均发病率

$$\bar{X} = \frac{94 + 20 + 26.5}{3} = 46.83(\%)$$

若采用加权计算, 对比试验林的平均发病率

$$\bar{X} = \frac{94 \times 50 + 20 \times 200 + 26.5 \times 100}{50 + 200 + 100} = 32.43(\%)$$

两者计算结果相差甚远。由于加权计算考虑了各无性系的权重, 即对比试验林内各无性系所占比例不同, 因此加权计算的结果更准确地反映了实际情况。

【2】几何平均数 (G) 当统计数值呈逐个递增或递减时, 可采用几何平均数表示平均增长 (递减) 率, 其公式如下:

$$G = \sqrt[n]{x_1 \cdot x_2 \cdots x_n} = (x_1 \cdot x_2 \cdots x_n)^{\frac{1}{n}} \quad (1-3)$$

$$\lg G = \frac{\lg x_1 + \lg x_2 + \cdots + \lg x_n}{n} = \frac{\sum_{i=1}^n \lg x_i}{n} \quad (1-4)$$

$$G = \text{arc lg} \left(\sum_{i=1}^n \lg x_i / n \right) \quad (1-5)$$

〔例 1—3〕 秦岭南坡火地塘林区的华山松疱锈病 (*Cronartium ribicola* J. C. Fischer ex Rab.) 1982~1986 年的病情指数如表 1-2 所示, 求该病害的平均增长率。

表 1-2 华山松疱锈病病情增长情况

年代	病情指数	相对增长率 (x_i)	$\lg x_i$
1982	3.3		
1983	4.8	1.4545	0.1627
1984	7.8	1.6250	0.2109
1985	15.3	1.9615	0.2926
1986	24.0	1.5686	0.1955
Σ			0.8617

根据表 1-2

$$\lg G = \frac{0.8617}{4} = 0.2154 \quad G = \text{arc lg} 0.2154 = 1.6421$$

因此可以认为该林区华山松疱锈病 1982~1986 年 5 年间, 是以平均每年 1.6421 倍的速度增长的。

【3】中数和众数 将变数内所有变量从小到大依次排列, 中间位置的变量称为中数。若变量个数为偶数, 则以中间两个变量的算术平均数为中数。

众数是指变数中最常见的一个数值, 或次数最多一组的中点。这两个平均数使用较少。

1.3.2 变异数

一般将变量个数 $n \geq 30$ 的样本称为大样本, $n < 30$ 的样本称

为小样本，平均数是一个样本的代表值。由于样本中各变量的变异程度不同，样本平均数的代表性则有差异。因此，为了全面地描述一个样本，在平均数的基础上还必须度量其变异程度。用于度量变异程度的统计数，称为变异数。常用的变异数有极差、方差、标准差和变异系数。

【1】极差 (R) 极差亦称全距，是样本中最大变量 $X_{\max}$ 与最小变量 $X_{\min}$ 的差数，即

$$R = X_{\max} - X_{\min} \quad (1-6)$$

[例 1-4] 测定长武县杨树溃疡病的发病情况，一个样本为北京杨，另一个样本为大官杨，各抽取了 10 个样地，测定结果。

北京杨 10 个样地的病情指数： $X_i = [13, 14, 15, 17, 18, 18, 19, 21, 22, 23]$

大官杨 10 个样地的病情指数： $X_j = [16, 16, 17, 18, 18, 18, 18, 19, 20, 20]$

计算两个样本的平均数： $\bar{X}_i = \bar{X}_j = \frac{180}{10} = 18$

计算两个样本的极差： $R_i = X_{i,\max} - X_{i,\min} = 23 - 13 = 10$

$$R_j = X_{j,\max} - X_{j,\min} = 20 - 16 = 4$$

因此，两个样本的平均病情指数虽然相同，但北京杨样本的极差较大，其平均数的代表性较差，大官杨样本的极差较小，即变异幅度较小，其平均数的代表性较好。

极差只取决于样本的极端变量，并未充分利用样本的全部信息，因此，只适于 $n \leq 10$ 的样本，此时它具有计算简便明了的特点。

【2】方差 (S^2) 方差是充分利用了样本信息的变异数，也称均方。

$$S^2 = \frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n - 1} \quad (1-7)$$

式中, n 为样本变量个数, $n-1$ 为自由度, 方差在统计分析中用途很广, 后面还将讨论。

【3】标准差 (S) 方差的正根值称为标准差, 其单位同样本变量的度量单位。

$$S = \sqrt{S^2} = \sqrt{\frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n-1}} = \sqrt{\frac{\sum_{i=1}^n x_i^2 - \frac{(\sum_{i=1}^n x_i)^2}{n}}{n-1}} \quad (1-8)$$

【4】变异系数 (CV) 当比较两个以上样本的变异程度时, 若样本平均数、变量单位不同, 则应采用变异系数进行比较。变异系数为样本标准差对平均数的百分数, 即

$$CV = \frac{S}{\bar{X}} \times 100 \quad (1-9)$$

【5】平均数的标准差 ($S_{\bar{x}}$) 和变异系数 ($CV_{\bar{x}}$) 样本标准差 S 是表示观测值对其平均数 $\bar{X}$ 的平均离散程度, 并不反映样本平均数 $\bar{X}$ 本身的误差。样本平均数距总体平均数的平均误差用平均数标准差表示。因此, 平均数 $\bar{X}$ 的代表性常用 $\bar{X} \pm S_{\bar{x}}$ 来说明, 显然 $S_{\bar{x}}$ 越小, 平均数 $\bar{X}$ 的代表性越强。

$$S_{\bar{x}} = \frac{S}{\sqrt{n}} \quad (1-10)$$

或

$$S_{\bar{x}} = \sqrt{\frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n(n-1)}} \quad (1-11)$$

平均百分率的标准差

$$S_p = \sqrt{\frac{p(1-p)}{n}} \times 100\% \quad (1-12)$$