

逻辑斯谛曲线K值的四点 式平均值估计法*

王振中 林孔勋

(华南农业大学植保系, 广州)

摘 要

本文提出了逻辑斯谛曲线的一种新的K值估算方法——四点式平均值法。利用在时间序列 $t_0, t_1, t_2, \dots, t_n$ 内所得的 $n+1$ 个观察值中具有相同的时间间距的两对点 t_i, t_j, t_k, t_l ($t_1 - t_k = t_j - t_i = \Delta t, k \neq i$) 所得的4个观察值, 可对逻辑斯谛方程的上界K值进行估算。在全序列中举能满足此一条件的所有数组, 求得多个K的估算值, 取数学期望作为K的无偏估计值。

本法较以往的目测法、三点法和平均值法为好, 但所得结果可能稍逊于枚举选优法, 但是, 四点式平均值法可以应用于非自然数离散型分布、K值很大或可能出现的范围较大、以及在种群未达平衡时便终止试验等类型的数据。在这些情况下, 枚举选优法的应用则受到一定的限制。

逻辑斯谛方程在研究有限空间内种群的增长动态方面, 有着重要的作用 (Andrewartha and Birch, 1954; Pearl and Reed, 1920)。在植物病害流行病学中, 对流行事件如产孢动态和显症概率等的描述。也是一种有用的工具 (肖悦岩等, 1983; Jowett *et al.*, 1974)。但是, 由于其上界K值的估算方法不很理想, 导致它的应用受到一定的限制 (万昌秀等, 1983), 并引起数据变换方面的一些失误 (Jowett *et al.*, 1974)。已有不少人¹对K值的估算方法进行了探讨。目前常用的方法有三点法 (Pearl and Reed, 1920)、目测法 (郭祖超等, 1965) 和平均值法 (Andrewartha and Birch, 1954)。

三点法是在横轴上取等距离的三个点, 以相应的纵坐标求K值:

$$K = \frac{2N_1N_2N_3 - N_2^2(N_1 + N_3)}{N_1N_3 - N_2^2}$$

式中, N_1, N_2, N_3 为等时间距离的3个观察值。由于在同一试验中, 满足横轴等距离条件的3个点不止一组, 因而取法很多, 这就在很大程度上受主观因素的影响。

目测法则是预选一K值, 利用观察值的线性转换值与时间关系作图, 若图形不够理想, 便改变K值, 直到自己较满意为止。同样, 这种方法也受主观因素的影响。

平均值法是用群体生长达到平衡以后所取得的观察值的平均值代替K值。这似乎比较合理, 因为平衡后的种群数目总是在K值上下波动。但是, 正由于它的波动, 很难在观察值中判断从哪一个观察点开始, 群体生长便达到平衡。若试验数据不很完整, 即在种群生长未达

* 华南农业大学范怀忠教授、李郁治老师、广东省农科院周亮高研究员审阅全文, 均此谨致谢忱。

本文于1985年7月30日收到。

平衡便终止试验时,此法便难于应用。

万昌秀等(1983)提出用枚举选优法拟合逻辑斯谛方程。此法追踪了目标函数(剩余平方和)在所有可能的 K 值和用于组配方程的数对数 n 对应点上的数值,然后一一对比,找出最佳点。该法提出者用 Gause (1934) 的草履虫的试验数据对该法进行验证,结果都较理想。无疑,这种方法较以前的3种方法都好。但是,此法追踪了多种 $K-n$ 组合的剩余平方和数值,然后再进行比较,其计算量是相当大的,尤其在种群数量很大时尤其如此。此外,对于用多次试验的平均值作量度标准时,由于数据不是整数型,枚举选优法便难于应用了。

因此,对于逻辑斯谛方程这样一个在种群增长动态方面有重要作用的数学模型,寻找一种简单客观且能适用于各种类型的数据(整型和实型,完整和不完整)的 K 值估算方法,显然是有重要意义的。

一、四点式平均值法

逻辑斯谛方程是具有密度反馈作用的种群增长模型,其微分式为:

$$dN/dt = rN(1 - N/K) \quad (1)$$

式中, K 为环境容纳量, r 为内禀增长率, N 为种群在时间 t 时的数量。

式(1)的积分式可表示为:

$$N = \frac{K}{1 + \exp(a - rt)} \quad (2)$$

从式(2)可以看到, K 值为当时间为无穷大时种群增长的极限,从理论上说, K 只能逼近,永远无法达到。基于这一点,我们提出了 K 值的四点式平均值估计法,利用双等区间的观察值求 K 。

由初始条件: $t = t_0$ 时 $N = N_0$,则式(1)的积分式经整理后可用下式表示:

$$\frac{N}{K - N} = \frac{N_0}{K - N_0} e^{r(t - t_0)} \quad (3)$$

设在时间序列 $t_0, t_1, t_2, \dots, t_n$ 内得 $n+1$ 个观察值: $N_0, N_1, N_2, \dots, N_n$,若在序列内有四个观察点 t_i, t_j, t_k, t_l 使得下述条件成立:

$$t_l - t_k = t_j - t_i = \Delta t$$

$$\text{且:} \quad \Delta t \neq 0, \quad i \neq k$$

则可将相应的观察值 N_i, N_j, N_k, N_l 代入式(3)并整理得:

$$K = \frac{N_i N_j (1 - e^{\Delta t \cdot r})}{N_j - N_i \cdot e^{\Delta t \cdot r}} \quad (4)$$

$$K = \frac{N_k N_l (1 - e^{\Delta t \cdot r})}{N_l - N_k \cdot e^{\Delta t \cdot r}} \quad (5)$$

令 $\rho = e^{\Delta t \cdot r}$,联立(4)(5)二式解出 ρ :

$$\rho = \frac{N_j N_l (N_k - N_i)}{N_i N_k (N_l - N_j)}$$

将 ρ 代入(4)或(5)式,便可以得到计算 K 值的公式;

$$K = \frac{N_j N_k (N_i + N_l) - N_i N_l (N_j + N_k)}{N_j N_k - N_i N_l} \quad (6)$$

$$(i < j, k < l, k \neq i)$$

在 $n+1$ 个观察值中, 将所有具有相等时间间距的两对观察值作为一组(设有 m 组), 按公式(6)可以求得 m 个 K 的计算值。

此 m 个 K 计算值均由试验数据求得, 都包含了所研究的种群的极限值信息。不过, 由于试验误差等方面的影响, 它们是在 K 真值上下波动的。

在 m 个 K 计算值中, 进行 K 真值的点估计, 我们可以用这 m 个 K 计算值的数学期望值代替 K 真值。从统计学上可证明此值为 K 真值的无偏估计。

在根据公式(6)计算各 K 值计算值时, 可能会遇到分母为0, 或 K 值为0或小于0的情况, 这时应舍去这一计算值, 以使 K 的估算值更加合理(出现这类情况显然是试验误差引起)。

经估算 K 真值以后, 可用一般的线性化方法求方程的参数 a 和 r 。同时, 根据Andrewart-ha和Birch(1954)的经验, 以上升部分的数据拟合方程, 舍去第一次出现 $N \geq K$ 情况后的所有观察点(包括该点)的观察值。

二、方法的应用

作者在1983年研究花生锈病显症概率分布时进行了两次试验, 结果列于表1。肖悦岩等(1983)认为, 显症率的分布是逻辑斯蒂曲线型的。

表 1 花生锈病显症动态试验结果¹⁾

Table 1 Result of experiments on the apparition dynamic of peanut rust

接种后时间(天)	9	10	11	12	13	14	15	16	17	18	19	20	23
试验 I	79	126	257	625	1,856	4,000	7,389	12,133	16,057	18,898	22,361	23,528	25,919
试验 II	/	122	154	247	606	3,287	7,547	12,990	22,552	28,970	34,649	39,507	42,916

1) 表中试验结果为孢子堆数目。

在调查数据序列中以具有任何相同时间间距的两对调查点($\Delta t = 1, 2, \dots, 11$)按公式(6)求 K 的计算值, 并用其数学期望代替 K 真值, 试验 I、II 的 K 值分别为25,868和43,999。为了避免人为的选择, 且也全部利用数据中含有的 K 值信息, 在求 K 值的过程中, 要将数据中全部具有任何相同时间间距的成组的两对调查点都参与运算。

以 $N < K$ 的所有数据拟合方程, 求参数 a 、 r , 同时计算一个回归效果的参数: 决定系数 R^2 ($R^2 = (ss - Q)/ss$, ss 为总平方和, Q 为剩余平方和)。结果(见表2)表明: 两次试验的

表 2 花生锈病显症动态试验拟合结果

Table 2 Fitness of the two-paired points method to the apparition dynamic of peanut rust

参 数	K	a	r	R^2
试验 I	25,868	12.8356	0.7757	0.9932
试验 II	43,999	14.2707	0.8147	0.9832

r 值均较接近,可以较真实地反映孢子堆群体的增长情况,且都具有很高的决定系数。实际上。用本法估算 K 值以后,再用有效积温代替时间,则两次试验的 r 值相差极小 (1×10^{-4})。

三、结论和讨论

四点式平均值法消除了主观因素对估算 K 值的影响,推导过程逻辑性强,结果标准化,方法简单易于应用。

在用来计算 K 计算值的两对等时间间距的观察点中,若 $j=k$,则 K 值的计算公式 (6) 即与三点法的 K 值计算公式一致,因而,三点法所估算的 K 值仅是本法众多的 K 计算值之一,是本法的一种特殊情况。显然,本法在利用三点法可能用到的全部结果以外,还利用了三点法所未涉及的样本中的 K 值信息,因而结果会比三点法所得的结果更为合理。但无疑,本法估算 K 值的计算量要比三点法的计算量大一些。目测法和平均值法则由于受主观因素的影响较大,其结果的逻辑性和合理性显然比不上四点式平均值法。

为了与枚举选优法 (万昌秀等, 1983) 进行比较,作者特利用 Gause (1934) 的实验数据利用两种方法分别拟合曲线。结果 (见表 3) 表明,四点式平均值法除了在第 IV 组数据与枚举选优法的结果差异较大外,其余各组均较接近,在剩余平方和 Q 和决定系数 R^2 两列看来,

表 3 Gause 数据的四点式平均值法和枚举选优法的计算结果比较
Table 3 Comparison of the two-paired points method (A) and the optimization method (B) based on Gause (1934) data

数 据		方法 ^a	K	a	r	剩余平方和 Q^b	R^2 ^b
<i>P. aurelia</i> 每0.5毫升的数据	一环(Ⅰ)	A	452	4.9098	0.9576	14,927	0.9709
		B	443	5.0683	1.0451	12,488	0.9756
	半环(Ⅱ)	A	248	4.7058	1.1978	4,974	0.9624
		B	241	5.2021	1.5284	2,692	0.9797
<i>P. caudatum</i> 每0.5毫升的数据	一环(Ⅲ)	A	131	3.9223	1.0939	2,244	0.9418
		B	130	4.0101	1.1414	2,245	0.9418
	半环(Ⅳ)	A	69	2.6858	0.5706	1,697	0.7810
		B	59	3.2086	0.9379	971	0.8747

a. A 为四点式平均值法, B 为枚举选优法;

b. 枚举选优法的 Q 、 R^2 值是用 K 、 a 、 r 值和相应数据由作者计算。

枚举选优法的结果可能较四点式平均值法好些但差异也不很大,第 III 组数据的拟合效果几乎完全一致,且剩余平方和还较大些。

但是,枚举选优法的计算量要比四点式平均值法大得多。根据枚举选优法,本法中花生锈病显症动态试验 I (表 1) 的 K 值范围应定在 22,361—25,919 之间,按每个 K 值有 10 种 n 的取法,比较点将近 36,000 个,试验 II 的 K 值范围要在 34,649—42,916 之间, n 有 9 种取法,比较点将超过 74,000 个,计算量是相当大的,且后一个范围还不一定包含最佳 K 值点。

四点式平均值法还适用于连续分布的观察值,这是枚举选优法所不能做到的。

在试验数据不很完整 (即对象种群未达到平衡便终止试验) 的情况下,平均值法是无法

进行的,枚举选优法更难于确定K值的选优范围,难于着手。但这种情况仍不影响四点式平均值法的应用,例如,在前述的花生锈病显症动态试验中(表1),若在第19天便停止观察,则群体仍在增长阶段,远未达到平衡,但应用四点式平均值法仍可较好地估算K值和拟合生长曲线,计算结果(见表4)与表2的较相近。这表明,在不很完整的种群增长数据中,仍

表4 花生锈病显症动态试验第9—19天数据的拟合结果

Table 4 Fitness of the two-paired points method to the apparition dynamic of peanut rust before the equilibrium of the population

参 数	K	a	r	R ²
试验 I	25,137	13.2194	0.8100	0.9941
试验 II	41,483	15.4457	0.9076	0.9932

含有K值的系统反馈信息,四点式平均值法便可以利用这些信息估计K值和拟合方程。

四点式平均值法适用范围广,计算量小,且方法标准化(不同的人在同一数据上可得到也只能得到同一结果),对逻辑斯谛曲线的拟合,有较大的使用价值。

参 考 文 献

- 万昌秀、梁中宇 1983 逻辑斯谛曲线的一种拟合方法。生态学报 3(3):288—296。
 肖悦岩、曾士迈等 1983 SIMYR——小麦条锈病流行的简要模拟模型。植物病理学报 13(1):1—13。
 郭祖超等 1965 医用数理统计方法。人民卫生出版社。
 Andrewartha, H. G. and Birch, L. G. 1954 The Distribution and Abundance of Animals, The University of Chicago Press.
 Gause, G. F. 1934 The struggle for existence. Baltimore. Wellians and Wilkins.
 Jowett, D., Browing, J. A. and Haning, B. C. 1974 Nonlinear disease progress curves. In Epidemics of Plant Disease (Kranz, J. ed.), pp. 115—136. Springer-Verlag, Berlin and New York pp. 170.
 Pearl, R. and Reed, L. J. 1920 On the rate of growth of the population of the United States since 1790 and its mathematical representation. *Proc. Nat. Acad. Sci.* 6:275—288.

TWO-PAIRED POINTS METHOD FOR ESTIMATING K VALUE OF LOGISTIC EQUATION

Wang Zhenzhong Lin Kunghsun

(Department of Plant Protection, South China Agricultural University, Guangzhou)

The traditional methods for fitting the logistic curve are found to have some weak points and can not be applied to all kinds of data. A new method, two-paired points method for estimating K value was therefore proposed.

A group of 4 observed points arranged in two pairs with a same time interval is used for calculating an estimate-value of K , the upper asymptotic value of the curve; many estimate-values of K are thus obtained with all of such 4 points and the expectation value of these K estimate-values is thus used as the true value of K . An equation may be obtained by fitting the curve with all observations with values less than the value of K . The equation for calculating an estimate-value of K with 4 points is as follows:

$$K = \frac{N_j N_k (N_i + N_l) - N_i N_l (N_j + N_k)}{N_j N_k - N_i N_l}$$

$$(t_j - t_i = t_l - t_k \neq 0, i < j, k < l, i \neq k)$$

A detail discussion was made of the two-paired points method with the three traditional methods—three points method, method of mean of the observations near asymptote, and stepwise adjustment method, and the former was found to be better than any of the three methods.

The fitness of the equation obtained with two-paired points method might probably be not so good as that obtained with the optimization method. However the former method can be applied to the data (a) of non-integral values, (b) with very high value of K and a large variable range of probable values of K , and (c) of the experiment ended before the equilibrium of the population, while the optimization method is difficult to be applied to such data.