

Visual Numerics

IMSL®

C 数字库

用户指南

第 2 卷 (共 2 卷): C-Stat 库

7.0 版



IMSL® C 数字 Stat 库

IMSL C Stat 库是科学编程中非常有用的 C 函数库。技术专家还将每个函数设计和记录为在研究活动中使用。

© 1970-2010 Visual Numerics、IMSL 和 PV-WAVE 是 Visual Numerics, Inc. 在美国和其它国家/地区的注册商标。PyIMSL、JMSL、JWAVE、TS-WAVE 和 Knowledge in Motion 是 Visual Numerics, Inc. 的商标。所有其它公司、产品或品牌名称均归其各自所有者拥有。

重要说明：本文档中包含的信息如有更改，恕不另行通知。对本文档的使用应受 Visual Numerics 软件许可协议的条款和条件（包括但不限于有限担保和责任限制）的制约。如果不接受许可协议的条款，则不能使用本文档，并且应立即返还本产品，以获得全额退款。未经 Visual Numerics 的明确书面许可，不得以任何形式复制或分发本文档。



C、C#、Java™、Java™ 和 Fortran
应用程序开发工具

目录

简介

5

目录	5
IMSL C Stat 库	5
文档的组织结构	5
查找正确的函数	6
命名约定	7
错误处理、下溢和上溢	8

时序和预测

9

例程	9
用法说明	10
arma	11
max_arma	18
auto_uni_ar	22
ts_outlier_identification	25
auto_arima	30
difference	39
box_cox_transform	42
autocorrelation	45
partial_autocorrelation	49
lack_of_fit	52
estimate_missing	54
garch	58

参考材料	62
用户错误	62
哪些因素决定错误严重级	62
错误和缺省操作的种类	62
产品支持	64
与 Visual Numerics 支持联系	64
附录 A	65
参考	65

简介

目录

IMSL C Stat 库.....5

文档的组织结构.....5

查找正确的函数.....6

命名约定.....7

错误处理、下溢和上溢.....8

IMSL C Stat 库

IMSL C Stat 库是科学编程中非常有用的 C 函数库。技术专家还将每个函数设计和记录为在研究活动中使用。

文档的组织结构

本手册包含每个函数的简要说明。有关特定函数的所有信息将在其中的某个章节中介绍。

每章的开头是简介，然后是列出该章中包含的函数的目录。函数的文档包含以下信息：

- 章节名称：通常为 *float* 和 *double* 类型版本的函数的公用根。
- 用途：函数用途的说明。
- 对照表：引用子程序的表单，其中列出必选参数。

必选参数：必选参数的说明，按其出现的顺序列出。

输入：必须初始化参数；函数无法对其进行更改。

输入/输出：必须初始化参数；函数通过此参数返回输出。参数不能为常量或表达式。

输出：无需进行初始化。参数不能为常量或表达式；函数通过此参数返回输出。

- **返回值：**函数返回的值。
- **具有可选参数的对照表：**引用函数的表单，其中列出必选参数和可选参数。
- **可选参数：**可选参数的说明，按其出现的顺序列出。
- **说明：**算法的说明以及对详细信息的引用。在许多情形下，还会记录其它 IMSL 函数以及类似或互补的函数。
- **错误：**列出特定函数可能发生的所有错误。参考资料的“用户错误”部分提供了有关错误类型的讨论。下面按类型列出这些错误：

信息性错误：列出函数可能发生的信息性错误。

警报错误：列出函数可能发生的警报错误。

警告错误：列出函数可能发生的警告错误。

致命错误：列出函数可能发生的致命错误。

参考：按作者姓名的字母顺序列出参考材料。

查找正确的函数

C Stat 库文档分为多个章节；每章都包含具有类似计算或分析功能的函数。若要查找适用于给定问题的函数，请使用每章简介中的目录。

命名约定

多数函数都提供有 *float* 和 *double* 两种类型的版本，两个版本的名称共享一个公用根。一些函数还可以是 *int* 类型。以下列表针对的是每种类型和多个类型版本所在的函数名称的相应前缀：

类型	前缀
<i>float</i>	imsls_f_
<i>double</i>	imsls_d_
<i>int</i>	imsls_i_

函数对应的章节名称仅包含公用根，从而使函数查找起来更容易。

适当时，将在整个 C Stat 库中使用相同的变量名称以保持一致性。例如，`anova_table` 表示包含方差统计信息的分析的数组，`y` 表示相关变量的响应矢量。

当编写访问 C Stat 库的程序时，请选择与 IMSL 外部名称不冲突的 C 名称。如果在选择名称时遵守以下规则，则细心的用户可避免与 IMSL 名称的任何冲突：

- 不要选择以任意大写或小写字符组合形式的 “imsls_” 开头的名称。

错误处理、下溢和上溢

C Stat 库中的函数会尝试检测和报告错误及无效输入。此错误处理功能提供对用户的自动保护，并不要求用户为错误条件的处理进行任何特定配置。错误根据严重性进行分类，并分配有代码编号。缺省情况下，出现严重性为中或更高的错误时，函数会自动输出消息。而且，严重性最高时会导致程序执行停止。严重级与错误的一般属性一样，由具有符号名称 `IMSL_FATAL`、`IMSL_WARNING` 等的“错误类型”指定。有关详细信息，请参见参考资料中的“[用户错误](#)”。

通常情况下，如果系统（硬件或软件）将下溢替换为值 0，则可以编写 C Stat 库代码，以便计算不会受下溢的影响。正常情况下，可忽略指示下溢的系统错误消息。

还可以编写 IMSL 代码来避免上溢。应对生成指出上溢问题的系统错误消息的程序进行检查，以发现其中是否存在编程错误，如输入数据不正确、参数类型不匹配或维度不正确。

在许多情形下，函数的文档会指出可导致算法失败的常见缺陷。

时序和预测

例程

ARIMA 模型

计算参数的最小二乘或矩量法估计值 **arma** 11

计算参数的最大似然估计值 **max_arma** 18

单变量自回归时序模型的自动选择和拟合 **auto_uni_ar** 22

检测和确定离群值，并同时估计时序中的
模型参数 **ts_outlier_identification** 25

在可能存在离群值的情况下，
进行的自动 ARIMA 建模和预测 **auto_arima** 30

对时序执行求差 **difference** 39

模型构建和评估实用程序

执行 Box-Cox 转换 **box_cox_transform** 42

示例自相关函数 **autocorrelation** 45

示例偏自相关函数 **partial_autocorrelation** 49

基于相关函数的失拟检测 **lack_of_fit** 52

估计时序中的缺失值 **estimate_missing** 54

异方差建模

计算异方差 (p,q) 模型的参数的估计值 **garch** 58

用法说明

本章中的函数假定时序不包含任何缺失的观测值。如果存在缺失值，则应将它们设置为 NaN，并且例程将返回相应的错误消息。要启用模型拟合，缺失值必须替换为相应的估计值。

时间域方法

将数据转换为平稳性后，通常推荐时间域中的临时模型，并执行参数估计、诊断检查和预测。

ARIMA 模型（自回归整合移动平均）

小但全面的平稳时序模型类由非季节性 ARMA 过程组成，这些过程由以下公式定义：

$$\phi(B) (W_t - \mu) = \theta(B)A_t, \quad t \in Z$$

其中 $Z = \{\dots, -2, -1, 0, 1, 2, \dots\}$ 表示一个整数集， B 是 $B_k W_t = W_{t-k}$ 定义的后移算子， μ 是 W_t 的平均值，并且以下等式成立：

$$\phi(B) = 1 - \phi_1 B - \phi_2 B^2 - \dots - \phi_p B^p, \quad p \geq 0$$

$$\theta(B) = 1 - \theta_1 B - \theta_2 B^2 - \dots - \theta_q B^q, \quad q \geq 0$$

模型是 (p, q) 阶的，称为 ARMA (p, q) 模型。

以下公式提供了等效版本的 ARMA (p, q) 模型：

$$\phi(B) W_t = \theta_0 + \theta(B)A_t, \quad t \in Z$$

其中 θ_0 是以下公式定义的总常量：

$$\theta_0 = \mu \left(1 - \sum_{i=1}^p \phi_i \right)$$

有关总常量的含义和作用的讨论，请参见 Box and Jenkins（1976 年，第 92–93 页）。

如果“原始”数据 $\{z_t\}$ 是同构且不平稳的，则使用 `imsls_f_difference` 的求差会产生平稳性，并且该模型称为 ARIMA（自回归整合移动平均）。参数估计在平稳时序 $w_t = \nabla d z_t$ 上执行，其中 $\nabla d = (1 - B)d$ 是后向差分算子，时间段为 1，阶数为 $d, d > 0$ 。

通常情况下，在调用 `imsls_f_arma` 函数以获取初步参数估计值时，矩量法会包括参数 `IMSLS_METHOD_OF_MOMENTS`。这些估计值可作为初始值用于最小二乘过程，方法是将参数 `IMSLS_LEAST_SQUARES` 包括在对函数 `imsls_f_arma` 的调用中。可使用用户提供的其它初始估计值。最小二乘过程可用于计算参数的条件或无条件最小二乘估计值，具体取决于选择的返溯长度。用于进行初步参数估计、最小二乘参数估计和预测的函数遵循 Box and Jenkins（1976 年，Programs 2–4，第 498–509 页）的方法。

arma

计算 ARMA 模型的参数的最小二乘估计值。

对照表

```
float *imsls_f_arma(int n_observations, float z[], int p, int q, ..., 0)
```

`double` 类型函数为 `imsls_d_arma`。

必选参数

`int n_observations` (输入)

观测值数。

`float z[]` (输入)

包含观测值的长度为 `n_observations` 的数组。

`int p` (输入)

自回归参数的数目。

`int q` (输入)

移动平均参数的数目。

返回值

指向长度 $1 + p + q$ 的数组的指针，具有估计常量、AR 和 MA 参数。如果指定 `IMSLS_NO_CONSTANT`，则此数组的第 0 个元素为 0.0。

说明

函数 `imsls_f_arma` 可根据观测样本 $\{w_t\}$ ($t=1, 2, \dots, n$, 其中 $n=n_{\text{observations}}$) 计算非季节性 ARMA 模型的参数估计值。有以下两种方法可供选择：矩量法和最小二乘法。缺省方法是矩量法。

提供两种参数估计方法：矩量法和最小二乘法。用户可以选择带有可选参数 `IMSLS_METHOD_OF_MOMENTS` 的矩量法。如果用户指定了 `IMSLS_LEAST_SQUARES`，则使用最小二乘法。如果用户希望使用最小二乘法，则初步估计值缺省为矩量法估计值。用户也可以通过指定可选参数 `IMSLS_INITIAL_ESTIMATES` 来输入初始估计值。下表列出了矩量法和最小二乘法的相应可选参数：

仅限于矩量法	仅限于最小二乘法	矩量法和最小二乘法
IMSLS_METHOD_OF_MOMENTS	IMSLS_LEAST_SQUARES	IMSLS_RELATIVE_ERROR
	IMSLS_CONSTANT (或 IMSLS_NO_CONSTANT)	IMSLS_MAX_ITERATIONS
	IMSLS_AR_LAGS	IMSLS_MEAN_ESTIMATE
	IMSLS_MA_LAGS	IMSLS_AUTOCOV(_USER)
	IMSLS_BACKCASTING	IMSLS_RETURN_USER
	IMSLS_CONVERGENCE_TOLERANCE	IMSLS_ARMA_INFO
	IMSLS_INITIAL_ESTIMATES	
	IMSLS_RESIDUAL(_USER)	
	IMSLS_PARAM_EST_COV(_USER)	
	IMSLS_SS_RESIDUAL	

矩量法估计

假设时序 $\{Z_t\}$ 是用以下形式的 ARMA (p, q) 模型生成的:

$$\phi(B)Z_t = \theta_0 + \theta(B)A_t$$

其中 $t \in \{0, \pm 1, \pm 2, \dots\}$

$\hat{\mu} = z_mean$ 是时序 $\{Z_t\}$ 的平均值 μ 的估计值, 其中 $\hat{\mu}$ 等于以下值:

$$\hat{\mu} = \begin{cases} \mu & \mu \text{ 已知} \\ \frac{1}{n} \sum_{t=1}^n Z_t & \mu \text{ 未知} \end{cases}$$

自协方差函数的估计方程如下

$$\hat{\sigma}(k) = \frac{1}{n} \sum_{t=1}^{n-k} (Z_t - \hat{\mu})(Z_{t+k} - \hat{\mu})$$

$k = 0, 1, \dots, K$, 其中 $K = p + q$ 。注意, $\hat{\sigma}(0)$ 是样本方差的估计值。

在给定样本自协方差的情况下, 函数将使用如下所示的扩展 Yule-Walker 方程计算自回归参数的矩量法估计值:

$$\hat{\Sigma} \hat{\phi} = \hat{\sigma}$$

其中

$$\begin{aligned} \hat{\phi} &= (\hat{\phi}_1, \dots, \hat{\phi}_p)^T \\ \hat{\Sigma}_{ij} &= \hat{\sigma}(|q+i-j|), \quad i, j = 1, \dots, p \\ \hat{\sigma}_i &= \hat{\sigma}(q+i), \quad i = 1, \dots, p \end{aligned}$$

总常量 θ_0 可通过以下方程来估计：

$$\hat{\theta}_0 = \begin{cases} \hat{\mu} & p = 0 \\ \hat{\mu} \left(1 - \sum_{i=1}^p \hat{\phi}_i \right) & p > 0 \end{cases}$$

移动平均参数是基于非线性方程组在给定 $K = p + q + 1$ 自协方差 $\sigma(k)$ (其中, $k = 1, \dots, K$) 以及 p 个自回归参数 ϕ_i (其中 $i = 1, \dots, p$) 的情况下估计出来的。

假定 $Z_t = \phi(B)Z_t$ 。可通过以下关系来估计派生的移动平均过程 $Z_t = \theta(B)A_t$ 的自协方差：

$$\hat{\sigma}'(k) = \begin{cases} \hat{\sigma}(k) & p = 0 \\ \sum_{i=0}^p \sum_{j=0}^p \hat{\phi}_i \hat{\phi}_j (\hat{\sigma}(|k+i-j|)) & p \geq 1, \hat{\phi}_0 = -1 \end{cases}$$

确定移动平均参数的迭代过程基于以下关系

$$\sigma(k) = \begin{cases} (1 + \theta_1^2 + \dots + \theta_q^2) \sigma_A^2 & k = 0 \\ (-\theta_k + \theta_1 \theta_{k+1} + \dots + \theta_{q-k} \theta_q) \sigma_A^2 & k \geq 1 \end{cases}$$

其中, $\sigma(k)$ 表示原始 Z_t 过程的自协方差函数。

假定 $\tau = (\tau_0, \tau_1, \dots, \tau_q)T$, $f = (f_0, f_1, \dots, f_q)T$, 其中

$$\tau_j = \begin{cases} \sigma_A & j = 0 \\ -\theta_j / \tau_0 & j = 1, \dots, q \end{cases}$$

并且

$$f_j = \sum_{i=0}^{q-j} \tau_i \tau_{i+j} - \hat{\sigma}'(j) \quad j = 0, 1, \dots, q$$

然后，可通过以下方程来确定第 $(i + 1)$ 次迭代时的 τ 值：

$$\tau^{i+1} = \tau^i - (T^i)^{-1} f^i$$

估计过程开始时使用以下初始值

$$\tau^0 = (\sqrt{\hat{\sigma}^2(0)}, 0, \dots, 0)^T$$

并在第 i 次迭代时当 $\|f^i\|$ 小于 `relative_error` 或 i 等于 `max_iterations` 时终止。
通过设置以下方程可从 τ 的最终估计值获取移动平均参数估计值

$$\hat{\theta}_j = -\tau_j / \tau_0 \quad j = 1, \dots, q$$

随机冲击方差可通过以下方程来估计：

$$\hat{\sigma}_A^2 = \begin{cases} \hat{\sigma}(0) - \sum_{i=1}^p \hat{\phi}_i \hat{\sigma}(i) & q = 0 \\ \tau_0^2 & q \geq 0 \end{cases}$$

有关执行类似计算的函数的说明，请参见 Box and Jenkins（1976 年，第 498–500 页）。

最小二乘估计

假设时序 $\{Z_t\}$ 是用以下形式的非季节性 ARMA 模型生成的,

$$\phi(B)(Z_t - \mu) = \theta(B)A_t \text{ 其中, } t \in \{0, \pm 1, \pm 2, \dots\}$$

其中, B 是后移算子, μ 是 Z_t 的平均值, 并且

$$\begin{aligned} \phi(B) &= 1 - \phi_1 B^{l_\phi(1)} - \phi_2 B^{l_\phi(2)} - \dots - \phi_p B^{l_\phi(p)} & p \geq 0 \\ \theta(B) &= 1 - \theta_1 B^{l_\theta(1)} - \theta_2 B^{l_\theta(2)} - \dots - \theta_q B^{l_\theta(q)} & q \geq 0 \end{aligned}$$

具有 p 个自回归参数和 q 个移动平均参数。在不丧失普遍性的情况下, 假定以下方程:

$$1 \leq l_f(1) \leq l_f(2) \leq \dots \leq l_f(p)$$

$$1 \leq l_q(1) \leq l_q(2) \leq \dots \leq l_q(q)$$

因此, 非季节性 ARMA 模型是 (p', q') 阶的, 其中 $p' = l_f(p)$, $q' = l_q(q)$ 。请注意, 常规的层次结构模型假定以下条件:

$$l_f(i) = i, 1 \leq i \leq p$$

$$l_q(j) = j, 1 \leq j \leq q$$

请考虑平方和函数

$$S_T(\mu, \phi, \theta) = \sum_{t=-T+1}^n [A_t]^2$$

其中

$$[A_t] = E[A_t | (\mu, \phi, \theta, Z)]$$

T 是向后原点。假定随机冲击 $\{A_t\}$ 是独立且均匀分布的

$$N(0, \sigma_A^2)$$

的随机变量。因此，对数似然函数为

$$l(\mu, \phi, \theta, \sigma_A) = f(\mu, \phi, \theta) - n \ln(\sigma_A) - \frac{S_T(\mu, \phi, \theta)}{2\sigma_A^2}$$

其中， $f(\mu, \phi, \theta)$ 是 μ 、 ϕ 、和 θ 的函数。

当 $T=0$ 时，对数似然函数取决于初始化模型所需的 z_t 和 A_t 的过去值。选择这些初始值的方法通常会向模型中引入瞬态偏差（Box and Jenkins, 1976 年, 210–211 页）。当 $T=\infty$ 时，此依赖性消失，估计问题涉及无条件对数似然函数的最大化。Box and Jenkins（1976 年, 213 页）论述到

$$S_\infty(\mu, \phi, \theta) / (2\sigma_A^2)$$

控制

$$l(\mu, \phi, \theta, \sigma_A^2)$$

使平方和函数最小化的参数估计值被称为最小二乘估计值。当 n 较大时，无条件最小二乘估计值近似等于最大似然估计值。

实际上， T 的有限值允许无条件平方和函数实现足够接近。将使用通过后向预测获得的 z_t 初始值以迭代方式计算 $[AT]$ 的值，计算无条件平方和时需要这些值。还会计算最终参数估计值的残差（包括回测）、随机冲击方差和协方差矩阵。通过将 `imsls_f_arma` 与 `imsls_f_difference` 结合使用可以计算 ARIMA 参数。

警告错误

`IMSLS_LEAST_SQUARES_FAILED`

参数的最小二乘估计无法收敛。增大“maxbc”和/或“tolerance”以及/或“convergence_tolerance”。可将上次迭代中参数的估计值用作新的起始值。

致命错误

IMSLT_TOO_MANY_CALLS	函数的调用次数超出“itmax”*(“n”+1) = %(i1)。用户可以尝试新的初始猜测值。
IMSLT_INCREASE_ERRREL	相对误差的范围 “errrel”= %(r1) 太小。近似解中无法做进一步改进。用户应增大“errrel”。
IMSLT_NEW_INITIAL_GUESS	迭代尚未取得重大进展。用户可以尝试新的初始猜测值。

max_arma

单变量 ARMA（自回归移动平均）时序模型中参数的精确最大似然估计值。

对照表

float *imslt_f_max_arma (*int* n_obs, *float* w[], *int* p, *int* q, ..., 0)

double 类型函数为 imslt_d_max_arma。

必选参数

- int* n_obs (输入)
时序中的观测值数。
- float* w[] (输入)
包含时序的长度为 n_obs 的数组。
- int* p (输入)
自回归参数的数量，非负。
- int* q (输入)
移动平均参数的非负数量。

返回值

指向长度为 1+p+q 的数组的指针，具有估计常量、AR 和 MA 参数。如果无法计算值，则返回 NULL。

具有可选参数的对照表

```
Float *imsls_f_max_arma (int n_obs, float w[], int p, int q,  
    IMSLS_INITIAL_ESTIMATES, float init_ar[], float init_ma[],  
    IMSLS_PRINT_LEVEL, int iprint,  
    IMSLS_MAX_ITERATIONS, int maxit,  
    IMSLS_LOG_LIKELIHOOD, float *log_likeli,  
    IMSLS_VAR_NOISE, float *avar,  
    IMSLS_ARMA_INFO, Imsls_f_arma **arma_info,  
    IMSLS_MEAN_ESTIMATE, float *w_mean,  
    IMSLS_RETURN_USER, float *constant, float ar[], float ma[],  
    0)
```

可选参数

IMSLS_INITIAL_ESTIMATES, float init_ar[], float init_ma[] (输入)

如果指定，init_ar 是长度为 p 的数组，它包含自回归参数的初步估计值，init_ma 是长度为 q 的数组，它包含移动平均参数的初步估计值；否则，将在内部进行计算。如果 p=0 或 q=0，则忽略对应参数。

IMSLS_PRINT_LEVEL, int iprint (输入)

输出选项：

0 — 无输出。

1 — 只输出最终结果。

2 — 输出中间和最终结果。

缺省设置：iprint = 0

IMSLS_MAX_ITERATIONS, int maxit (输入)

最大估计值迭代次数。

缺省设置：maxit = 300

IMSLS_VAR_NOISE, float *avar (输出)

创新方差的估计值。

IMSLS_LOG_LIKELIHOOD, float *log_likeli (输出)

拟合模型的 $-2*(\ln(\text{likelihood}))$ 值。

IMSLS_ARMA_INFO, *lmsls_f_arma* **arma_info (输出)

指向内部分配结构的指针的地址, 该结构类型为 *lmsls_f_arma*, 包含调用 *lmsls_f_arma_forecast* 时必需的信息。

IMSLS_MEAN_ESTIMATE, *float* *w_mean (输入/输出)

时序 *w* 的平均值估计。返回时, *w_mean* 会包含该平均值的更新。

缺省设置: 时序 *w* 以其样本平均值为中心。

IMSLS_RETURN_USER, *float* *constant, *float* ar[], *float* ma[] (输出)

如果指定, *constant* 将为常量参数估计值, *ar* 是长度为 *p* 的数组, 包含最终的自回归参数估计值, *ma* 是长度为 *q* 的数组, 包含最终的移动平均参数估计值。

说明

函数 *lmsls_f_max_arma* 派生自 Akaike, Kitagawa, Arahata and Tada (1979) 介绍的最大似然估计值算法和 TIMSAC-78 库中发布的 XSARMA 例程。

使用本章开头用时间域方法开发的表示法, 可通过以下关系用非季节性自回归移动平均 (ARMA) 模型来表示平均值为 μ 的平稳时序 W_t :

$$\phi(B)(W_t - \mu) = \theta(B)a_t$$

其中

$$t \in ZZ = \{\dots, -2, -1, 0, 1, 2, \dots\},$$

B 是由 $B^k W_t = W_{t-k}$ 定义的后移算子,

$$\phi(B) = 1 - \phi_1 B - \phi_2 B^2 - \dots - \phi_p B^p, \quad p \geq 0,$$

并且

$$\theta(B) = 1 - \theta_1 B - \theta_2 B^2 - \dots - \theta_q B^q, \quad q \geq 0.$$

函数 *lmsls_f_max_arma* 使用最大似然估计值来估计 AR 系数 $\phi_1, \phi_2, \dots, \phi_p$ 和 MA 系数 $\theta_1, \theta_2, \dots, \theta_q$ 。

函数 `imsls_f_max_arma` 可检查自回归和移动平均系数的初始估计值，以确保它们分别表示平稳和可逆序列。

如果

$$\phi_1, \phi_2, \dots, \phi_p$$

是平稳序列的初始估计值，则以下多项式的所有（复数）根都将处于单位圆的范围之外：

$$1 - \phi_1 z - \phi_2 z^2 - \dots - \phi_p z^p.$$

如果

$$\theta_1, \theta_2, \dots, \theta_q$$

是可逆序列的初始估计值，则该多项式的所有（复数）根

$$1 - \theta_1 z - \theta_2 z^2 - \dots - \theta_q z^q$$

都将处于单位圆的范围之外。

AR 和 MA 系数的初始值可由矢量 `init_ar` 和 `init_ma` 提供。否则，将通过矩量法在内部计算估计值。`imsls_f_max_arma` 用于计算关联多项式的根。如果 AR 估计值表示非平稳序列，`imsls_f_max_arma` 将发出警告消息并将 `init_ar` 替换为平稳的初始值。如果 MA 估计值表示不可逆序列，`imsls_f_max_arma` 将发出终止性错误，并且必须寻找新的初始值。

`imsls_f_max_arma` 还会验证 AR 系数的最终估计值，以确保它们也表示平稳序列。这样是为了保证内部对数似然优化程序不会收敛到非平稳解。如果遇到非平稳估计值，`imsls_f_max_arma` 将发出致命错误消息。

在选择模型时，可能会优先选择 AIC 值最小的 ARMA 模型，

$$AIC = \log_likeli + 2(p+q)$$

函数 `imsls_f_max_arma` 还可以处理白噪音过程，即 ARMA(0,0) 过程。

auto_uni_ar

单变量自回归时序模型的自动选择和拟合。使用 Akaike 的信息准则 (AIC) 自动选择模型的滞后。使用矩量法、最小二乘法或最大似然法计算具有最小 AIC 的模型的自回归参数估计值。

对照表

```
float *imsls_f_auto_uni_ar(int n_obs, float z[], int maxlag,  
                           int *p, ..., 0)
```

double 类型函数为 imsls_d_auto_uni_ar。

必选参数

int n_obs (输入)

时序中的观测值数。

float z[] (输入)

包含平稳时序的长度为 n_obs 的数组。

int maxlag (输入)

请求的自回归参数的最大数目。要求

$1 \leq \text{maxlag} \leq n_obs/2$ 。

int *p (输出)

具有最小 AIC 的模型中自回归参数的数量。

返回值

长度为 $1 + \text{maxlag}$ 的矢量，包含具有最小 AIC 的模型中常量和自回归参数的估计值。估计值位于此数组的第一个 $1 + p$ 位置。

具有可选参数的对照表

```
float *imsls_f_auto_uni_ar(int n_obs, float z[], int maxlag, int *p,  
    IMSLS_PRINT_LEVEL, int iprint,  
    IMSLS_MAX_ITERATIONS, int maxit,  
    IMSLS_METHOD, int method,  
    IMSLS_VAR_NOISE, float *avar,  
    IMSLS_AIC, float *aic,  
    IMSLS_MEAN_ESTIMATE, float *z_mean,  
    IMSLS_RETURN_USER, float *constant, float ar[],  
    0)
```

可选参数

IMSLS_PRINT_LEVEL, *int* iprint (输入)

输出选项:

0 — 无输出。

1 — 只输出最终结果。

2 — 输出中间和最终结果。

缺省设置: iprint = 0

IMSLS_MAX_ITERATIONS, *int* maxit (输入)

最大估计值迭代次数。

缺省设置: maxit = 300

IMSLS_METHOD, *int* method (输入)

估计方法选项:

0 — 矩量法

1 — 通过 Householder 变换实现的最小二乘法

2 — 最大似然法

缺省设置: method = 1

IMSLS_VAR_NOISE, *float* *avar (输出)

创新方差的估计值。

IMSLS_AIC, *float* *aic (输出)

最小 AIC。

IMSL_S_MEAN_ESTIMATE, *float* *z_mean (输入/输出)

时序 z 的平均值估计。返回时, z_mean 包含平均值的更新。

缺省设置: 时序 z 以其样本平均值为中心。

IMSL_S_RETURN_USER, *float* *constant, *float* ar[] (输出)

如果指定, 则 $constant$ 是常量参数估计值, ar 是长度为 $maxlag$ 的数组, 其中在其第一个 p 位置包含最终自回归参数估计值。

说明

函数 `auto_uni_ar` 会自动选择拟合数据效果最好的 AR 模型的阶数, 然后计算 AR 系数。
`auto_uni_ar` 中所用的算法来自 Akaike, H., et. al (1979) 和 Kitagawa and Akaike (1978) 的著作。
此代码是根据在 TIMSAC-78 库中发布的 UNIMAR 过程改编的。

用 $0, 1, 2, \dots, maxlag$ 自回归系数连续拟合 AR 模型可确定拟合效果最好的 AR 模型。对于每个模型, 均会根据以下公式计算 Akaike 的信息准则 (AIC)

$$AIC = -2 \ln(likelihood) + 2(p+1)$$

函数 `auto_uni_ar` 使用由 Ozaki and Oda (1979) 开发的此公式的近似,

$$AIC = (n_obs - maxlag) \ln(\hat{\sigma}^2) + 2(p+1) + (n_obs - maxlag)(\ln(2\pi) + 1),$$

其中 $\hat{\sigma}^2$ 是序列残差的估计值, 在时序分析中通常称为创新方差。通过删除常量

$$(n_obs - maxlag)(\ln(2\pi) + 1),$$

计算可简化为

$$AIC = (n_obs - maxlag) \ln(\hat{\sigma}^2) + 2(p+1)$$

最佳拟合模型是具有最小 AIC 的模型。如果此模型中的参数数量等于拟合的最高阶自回归模型, 即 $p=maxlag$, 则对于较大的 $maxlag$ 值, 可能存在 AIC 较小的模型。在此情况下, 增大 $maxlag$ 可以保证用更多自回归参数展开 AR 模型。

如果 $method=0$, 则使用矩量法计算具有最小 AIC 的模型的自回归系数的估计值。如果 $method=1$, 则通过以 Kitagawa and Akaike (1978) 所述形式应用的最小二乘法确定系数。否则, 如果 $method=2$, 则使用最大似然法估计系数。

ts_outlier_identification

检测并确定离群值，同时估计某个时序中的模型参数的值，该时序的基础无离群值序列遵循常规季节性或非季节性 ARMA 模型。

对照表

```
float *imsls_f_ts_outlier_identification(int n_obs, int model[],
                                         float w[], ..., 0)
```

double 类型的函数是 `imsls_d_ts_outlier_identification`。

必选参数

int `n_obs` (输入)

时序中的观测值数。

int `model[]` (输入)

长度为 4 的矢量，包含无离群值序列遵循的 $ARIMA(p, 0, q) \times (0, d, 0)$ 模型的数值 p 、 q 、 s 、 d 。

float `w[]` (输入)

包含时序的长度为 `n_obs` 的数组。

返回值

指向长度为 `n_obs` 的数组的指针，该数组包含无离群值时序。
如果出现错误，则返回 `NULL`。

具有可选参数的对照表

```
float *imsls_f_ts_outlier_identification(int n_obs,
                                         int model[], float w[],
                                         IMSLS_RETURN_USER, float x[],
                                         IMSLS_DELTA, float delta,
                                         IMSLS_CRITICAL, float critical,
                                         IMSLS_EPSILON, float epsilon,
                                         IMSLS_RELATIVE_ERROR, float relative_error,
                                         IMSLS_RESIDUAL, float **residual,
                                         IMSLS_RESIDUAL_USER, float residual[],
```

```

IMSL_RESIDUAL_SIGMA, float *res_sigma,
IMSL_NUM_OUTLIERS, int *num_outliers,
IMSL_OUTLIER_STATISTICS, int **outlier_stat,
IMSL_OUTLIER_STATISTICS_USER, int outlier_stat[],
IMSL_TAU_STATISTICS, float **tau_stat,
IMSL_TAU_STATISTICS_USER, float tau_stat[],
IMSL_OMEGA_WEIGHTS, float **omega,
IMSL_OMEGA_WEIGHTS_USER, float omega[],
IMSL_ARMA_PARAM, float **parameters,
IMSL_ARMA_PARAM_USER, float parameters[],
IMSL_AIC, float *aic,
0)

```

可选参数

IMSL_RETURN_USER, *float* x[] (输出)

长度为 `n_obs` 的用户提供的数组，其中包含无离群值序列。

IMSL_DELTA, *float* delta (输入)

检测临时变化离群值 (TC) 时使用的阻尼影响参数， $0 < \text{delta} < 1$ 。

缺省设置: `delta = 0.7`

IMSL_CRITICAL, *float* critical (输入)

用作离群值检测阈值的临界值，`critical > 0`。

缺省设置: `critical = 3.0`

IMSL_EPSILON, *float* epsilon (输入)

正公差值在离群值检测期间控制参数估计值的精度。

缺省设置: `epsilon = 0.001`

IMSL_RELATIVE_ERROR, *float* relative_error (输入)

函数 `imsls_f_arma` 中使用的非线性方程求解器的停止准则。

缺省设置: `relative_error = 10-10`。

IMSL_RESIDUAL, *float* **residual (输出)

指向内部分配的长度为 `n_obs` 的数组的指针地址，该数组包含无离群值序列的残差。

IMSL_RESIDUAL_USER, *float* residual[] (输出)

用户提供数组 `residual` 的存储。请参阅 `IMSL_RESIDUAL`。

IMSL_RESIDUAL_SIGMA, *float* *res_sigma (输出)

无离群值序列的残差标准差。

IMSL_NUM_OUTLIERS, *int* *num_outliers (输出)

检测到的离群值的数量。

IMSL_OUTLIER_STATISTICS, *int* **outlier_stat (输出)

指向内部分配的长度为 $\text{num_outliers} \times 2$ 的数组的指针地址，该数组包含离群值统计信息。第一列包含观测离群值的时间 ($t=1,2,\dots,n_{\text{obs}}$)，第二列包含指示所观测的离群值类型的标识符。

离群值类型分为以下五个类别：

- 0 创新离群值 (IO)
- 1 附加离群值 (AO)
- 2 水平平移离群值 (LS)
- 3 临时变化离群值 (TC)
- 4 无法识别 (UI)。

使用 `IMSL_NUM_OUTLIERS` 可获取检测到的离群值数量 `IMSL_NUM_OUTLIERS`。

如果 `num_outliers=0`，将返回 `NULL`。

IMSL_OUTLIER_STATISTICS_USER, *int* outlier_stat[] (输出)

用户分配的长度为 $n_{\text{obs}} \times 2$ 的数组，该数组在第一个 `num_outliers` 位置包含离群值统计信息。请参见 `IMSL_OUTLIER_STATISTICS`。

如果 `num_outliers=0`，则 `outlier_stat` 保持不变。

IMSL_TAU_STATISTICS, *float* **tau_stat (输出)

指向内部分配的长度为 `num_outliers` 的数组的指针地址，该数组包含每个检测的离群值的 t 值。

如果 `num_outliers=0`，则返回 `NULL`。

IMSLS_TAU_STATISTICS_USER, *float* tau_stat[] (输出)

用户分配的长度为 `n_obs` 的数组, 在其第一个 `num_outliers` 位置包含每个检测的离群值的 `t` 值。

如果 `num_outliers=0`, `tau_stat` 将保持不变。

IMSLS_OMEGA_WEIGHTS, *float* **omega (输出)

指向内部分配数组的指针的地址, 该数组长度为 `num_outliers`, 包含为检测到的离群值计算的 ω 权重。

如果 `num_outliers=0`, 将返回 `NULL`。

IMSLS_OMEGA_WEIGHTS_USER *float* omega[] (输出)

用户分配的长度为 `n_obs` 的数组, 其中在第一个 `num_outliers` 位置包含为检测到的离群值计算的 ω 权重。

如果 `num_outliers=0`, `omega` 将保持不变。

IMSLS_ARMA_PARAM, *float* **parameters (输出)

指向内部分配数组的指针的地址, 该数组长度为 $1+p+q$, 包含估计的常量、AR 和 MA 参数。

IMSLS_ARMA_PARAM_USER *float* parameters[] (输出)

长度为 $1+p+q$ 的用户分配数组, 其中包含估计的常量、AR 和 MA 参数。

IMSLS_AIC, *float* *aic (输出)

Akaike 的信息准则 (AIC)。

说明

考虑单变量时序 $\{Y_t\}$, 可通过下面的 $(p,0,q) \times (0,d,0)_s$ 阶乘积季节性 ARIMA 模型来对其进行描述:

$$Y_t - \mu = \frac{\theta(B)}{\Delta_s^d \phi(B)} a_t, t = 1, \dots, n.$$

这里, $\Delta_s^d = (1 - B^s)^d$, $\theta(B) = 1 - \theta_1 B - \dots - \theta_q B^q$, $\phi(B) = 1 - \phi_1 B - \dots - \phi_p B^p$ 。 B 是滞后算子, $B^k Y_t = Y_{t-k}$, $\{a_t\}$ 是白噪音过程, μ 表示序列 $\{Y_t\}$ 的平均值。

通常，由于离群值的影响，无法直接观测 $\{Y_t\}$ 。Chen and Liu (1993) 对四种类型的离群值进行了区分：创新离群值 (IO)、附加离群值 (AO)、临时变化 (TC) 和水平平移 (LS)。如果离群值作为序列的最后一个观测值出现，则 Chen and Liu 的算法无法确定该离群值的分类。在 `imsls_f_ts_outlier_identification` 中，将此类离群值称为 UI（无法识别），并将其作为创新离群值对待。

为了考虑在时间点 $t_1, t_2, \dots, t_m$ 出现的多个离群值的效应，Chen and Liu 考虑了以下模型：

$$Y_t^* - \mu = \sum_{j=1}^m \omega_j L_j(B) I_t(t_j) + \frac{\theta(B)}{\Delta_s^d \phi(B)} a_t.$$

这里， $\{Y_t^*\}$ 是观测到的离群值污染序列， ω_j 和 $L_j(B)$ 分别表示离群值 j 的幅值和动态模式。 $I_t(t_j)$ 是用于确定离群值效应的临时过程的指示器函数， $I_{t_j}(t_j)=1$ ，否则， $I_t(t_j)=0$ 。注意， $L_j(B)$ 通过 $B^k I_t = I_{t-k}, k=0,1,\dots$ 作用于 I_t 。

最后一个公式表明，可通过从原始序列 $\{Y_t^*\}$ 中删除所有出现的离群值效应来获得无离群值序列 $\{Y_t\}$ ：

$$Y_t = Y_t^* - \sum_{j=1}^m \omega_j L_j(B) I_t(t_j).$$

不同类型的离群值具有不同的 $L_j(B)$ 值：

1. $L_j(B) = \frac{\theta(B)}{\Delta_s^d \phi(B)}$ ，用于创新离群值，
2. $L_j(B) = 1$ ，用于附加离群值，
3. $L_j(B) = (1-B)^{-1}$ 用于水平平移离群值，以及
4. $L_j(B) = (1-\delta B)^{-1}, 0 < \delta < 1$ ，用于临时变化离群值。

函数 `imsls_f_ts_outlier_identification` 是 Chen and Liu 算法的实现。它通过三个阶段共同确定 $\phi(B), \theta(B)$ 中的系数以及观测序列模型中的离群值效应。离群值效应的幅值由最小二乘估计值确定。离群值检测本身则通过检查离群值效应的标准化统计信息的最大值来实现。有关详细描述，请参见 Chen and Liu 的原稿 (1993)。

$\phi(B)$ 和 $\theta(B)$ 中系数的中间和最终估计值由函数 `imsls_f_arma` 和 `imsls_f_max_arma` 来计算。如果 $\phi(B)$ 或 $\theta(B)$ 的根位于单位圆之上或之内，算法将停止，并显示相应的错误消息。在此情况下，应对 p 和 q 尝试不同的值。

auto_arima

自动标识时序离群值、确定乘积季节性 $ARIMA(p, 0, q) \times (0, d, 0)$ 模型的参数，并生成合并了其效应持续至时序结束之后的离群值的预测。

对照表

`float *imsls_f_auto_arima(int n_obs, int tpoints[], float x[], ..., 0)`

`double` 类型函数为 `imsls_d_auto_arima`。

必选参数

`int n_obs` (输入)

原始时序中的观测值数。假定在时间点 $t_1, \dots, t_{n_obs}$ 定义序列，则序列（包括缺失的观测值）的实际长度为 $n = t_{n_obs} - t_1 + 1$ 。

`int tpoints[]` (输入)

观测长度为 `n_obs` 的矢量，包含时间点 $t_1, t_2, \dots, t_{n_obs}$ 。要求 $t_1, t_2, \dots, t_{n_obs}$ 严格按升序排列。

float x[] (输入)

长度为 `n_obs` 的矢量，包含观测的时序值 $Y_1^*, Y_2^*, \dots, Y_{n_obs}^*$ 。此序列可以包含离群值和缺失的观测值。离群值由此例程标识，缺失值由矢量 `tpoints` 中的时间值标识。如果两个连续时间点之间的时间间隔大于 1，即 $t_{i+1} - t_i = m > 1$ ，则假定在 $m-1$ 和 t_i 之间在时间 t_{i+1} 处存在 $t_i + 1, t_i + 2, \dots, t_{i+1} - 1$ 个缺失值。因此，假定为等距时间点定义无间隔序列。在确定离群值并生成预测之前，会自动估计缺失值。将同时为缺失值和观测值生成预测。

返回值

指向长度为 $1 + p + q$ 的数组的指针，该数组中包含估计的常量，以及使用 $\text{ARIMA}(p, 0, q) \times (0, d, 0)_s$ 模型拟合无离群值序列时要使用的 AR 和 MA 参数。完成后，如果 $d = \text{model}[3] = 0$ ，则将 $\text{ARMA}(p, q)$ 模型或 $\text{AR}(p)$ 模型拟合为观测序列 Y_t^* 的无离群值版本。如果 $d = \text{model}[3] > 0$ ，则为季节性调整序列 $Z_t^* = \Delta_s^d \cdot Y_t^* = (1 - B_s)^d \cdot Y_t^*$ 的 $\text{ARMA}(p, q)$ 表示计算这些参数，其中 $B_s Y_t^* = Y_{t-s}^*$ ， $s = \text{model}[2] \geq 1$ 。
如果出现错误，将返回 `NULL`。

具有可选参数的对照表

```
float *imsls_f_auto_arima(int n_obs, int tpoints[], float x[],  
    IMSLS_METHOD, int method,  
    IMSLS_MAX_LAG, int maxlag,  
    IMSLS_MODEL, int model[],  
    IMSLS_DELTA, float delta,  
    IMSLS_CRITICAL, float critical,  
    IMSLS_EPSILON, float epsilon,  
    IMSLS_RESIDUAL_SIGMA, float *res_sigma,  
    IMSLS_P_INITIAL, int n_p_initial, int p_initial[],  
    IMSLS_Q_INITIAL, int n_q_initial, int q_initial[],  
    IMSLS_S_INITIAL, int n_s_initial, int s_initial[],  
    IMSLS_D_INITIAL, int n_d_initial, int d_initial[],  
    IMSLS_OUTLIER_STATISTICS, int **outlier_stat,  
    IMSLS_OUTLIER_STATISTICS_USER, int outlier_stat[],  
    IMSLS_AIC, float *aic,  
    IMSLS_AICC, float *aicc,
```

```

IMSL_S_BIC, float *bic,
IMSL_MODEL_SELECTION_CRITERION, int criterion,
IMSL_CONFIDENCE, float confidence,
IMSL_NUM_PREDICT, int n_predict,
IMSL_OUTLIER_FORECAST, float **outlier_forecast,
IMSL_OUTLIER_FORECAST_USER, float outlier_forecast[],
0)

```

可选参数

IMSL_METHOD, int method (输入)

选择模型时使用的方法:

1 — 自动选择 $ARIMA(p, 0, 0) \times (0, d, 0)_s$

3 — 指定的 $ARIMA(p, 0, q) \times (0, d, 0)_s$ 模型

需要参数 `IMSL_MODEL`。

缺省设置: `method = 1`

IMSL_MAX_LAG, int maxlag (输入)

拟合 $AR(p)$ 模型时允许的最大滞后。

缺省设置: `maxlag = 10`

IMSL_MODEL, int model[] (输入/输出)

长度为 4 的数组, 其中包含 p 、 q 、 s 、 d 的值。如果选择 `method=3`, 则必须定义 p 和 q 的值。如果未定义 `IMSL_S_INITIAL` 和 `IMSL_D_INITIAL`, 则还必须给定 s 和 d 。如果 `method=1`, 则模型在作为输入数组时将被忽略。输出时, 模型分别包含 `model[0]`、`model[1]` 和 `model[3]` 中 p 、 q 、 s 、 d 的最佳值。

IMSL_DELTA, float delta (输入)

检测临时变化离群值 (TC) 时使用的阻尼影响参数, $0 < \text{delta} < 1$ 。

缺省设置: `delta = 0.7`

IMSL_CRITICAL, float critical (输入)

用作离群值检测阈值的临界值, $\text{critical} > 0$ 。

缺省设置: `critical = 3.0`

IMSL_EPSILON, float epsilon (输入)

正公差值在离群值检测期间控制参数估计值的精度。

缺省设置: `epsilon = 0.001`

IMSL_RESIDUAL_USER, *float* residual[] (输出)

用户提供数组 residual 的存储。

IMSL_RESIDUAL_SIGMA, *float* *res_sigma (输出)

无离群值原始序列的残差标准差 (RSE)。

IMSL_NUM_OUTLIERS, *int* *num_outliers (输出)

检测到的离群值的数量。

IMSL_P_INITIAL, *int* n_p_initial, *int* p_initial[] (输入)

具有 n_p_initial 个元素的数组，其中包含 p 的候选值，可从中选择最佳值。

p_initial[] 中的所有候选值都必须是非负数，并且 $n_p_initial \geq 1$ 。如果 method=2，则必须定义 IMSL_P_INITIAL。否则，将忽略 n_p_initial 和 p_initial。

IMSL_Q_INITIAL, *int* n_q_initial, *int* q_initial[] (输入)

具有 n_q_initial 个元素的数组，其中包含 q 的候选值，可从中选择最佳值。

q_initial[] 中的所有候选值都必须是非负数，并且 $n_q_initial \geq 1$ 。如果 method=2，则必须定义 IMSL_Q_INITIAL。否则，将忽略 n_q_initial 和 q_initial。

IMSL_S_INITIAL, *int* n_s_initial, *int* s_initial[] (输入)

长度为 n_s_initial 的矢量，其中包含 s 的候选值，可从中选择最佳值。

s_initial[] 中的所有候选值都必须是非负数，并且 $n_s_initial \geq 1$ 。

缺省设置: n_s_initial=1, s_initial={1}

IMSL_D_INITIAL, *int* n_d_initial, *int* d_initial[] (输入)

长度为 n_d_initial 的矢量，其中包含 d 的候选值，可从中选择最佳值。

d_initial[] 中的所有候选值都必须是非负数，并且 $n_d_initial \geq 1$ 。

缺省设置: n_d_initial=1, d_initial={0}

IMSLI_OUTLIER_STATISTICS, *int* **outlier_stat (输出)

指向内部分配的长度为 $\text{num_outliers} * 2$ 的数组的指针地址，该数组包含离群值统计信息。第一列包含观测离群值的时间 ($t = t_1, t_1 + 1, t_1 + 2, \dots, t_{n_obs}$)，第二列包含指示所观测的离群值类型的标识符。离群值类型分为以下五个类别：

- 0 创新离群值 (IO)
- 1 附加离群值 (AO)
- 2 水平平移离群值 (LS)
- 3 临时变化离群值 (TC)
- 4 无法识别 (UI)。

如果 $\text{num_outliers}=0$ ，则返回 NULL。

IMSLI_OUTLIER_STATISTICS_USER, *int* outlier_stat[] (输出)

用户分配的长度为 $n \times 2$ 的数组，该数组在第一个 num_outliers 行中包含离群值统计信息。这里， $n = t_{n_obs} - t_1 + 1 \geq n_obs$ 。

如果 $\text{num_outliers} = 0$ ，则 outlier_stat 保持不变。

IMSLI_AIC, *float* *aic (输出)

最佳模型的 AIC (Akaike 的信息准则) 值。使用基于序列的创新方差估计值的最大对数似然近似值。

IMSLI_AICC, *float* *aicc (输出)

最佳模型的 AICC (修正的 AIC) 值。使用基于序列的创新方差估计值的最大对数似然近似值。

IMSLI_BIC, *float* *bic (输出)

最佳模型的 BIC (Bayesian 信息准则) 值。使用基于序列的创新方差估计值的最大对数似然近似值。

IMSLI_MODEL_SELECTION_CRITERION, *int* criterion (输入)

用于最佳模型选择的信息准则。

准则	选择的信息准则
0	Akaike 的信息准则 (AIC)
1	Akaike 的修正信息准则 (AICC)
2	Bayesian 信息准则 (BIC)

缺省设置：准则 = 0。

IMSLI_OUT_FREE_SERIES_USER, *float* outfree_series[] (输出)

用户分配的长度为 $n * 2$ 的数组，其中 $n = t_{n_obs} - t_1 + 1$ 。

IMSLI_CONFIDENCE, *float* confidence (输入)

用于计算预测置信界限的置信度，从不包含区间 (0, 100) 中选择。confidence 的典型选择有 90.0、95.0 和 99.0。

缺省设置：confidence = 95.0

IMSLI_NUM_PREDICT, *int* n_predict (输入)

请求的预测数。从原点 t_{n_obs} （即序列的最后一个观测值）处进行的预测。

缺省设置：n_predict = 0

IMSLI_OUTLIER_FORECAST, *float* **outlier_forecast (输出)

指向内部分配的长度为 $n_predict * 3$ 的数组的指针地址。第一列包含 $t = t_{n_obs} + 1, t_{n_obs} + 2, \dots, t_{n_obs} + n_predict$ 的原始序列的预测值。第二列包含这些预测值的标准差，第三列包含模型的无限阶移动平均表单的 ψ 权重。

IMSLI_OUTLIER_FORECAST_USER, *float* outlier_forecast[] (输出)

用户分配的长度为 $n_predict * 3$ 的数组。

IMSLI_RETURN_USER, *float* x[] (输出)

用户分配的数组，该数组在其第一个 $1+p+q$ 位置中包含估计的常量、AR 和 MA 参数。可根据上限估计 p 和 q 的值：如果 method=1，则 p 的上限将为 maxlag，并且 $q=0$ 。如果 method=2， p 和 q 的上限将分别为数组 p_initial 和 q_initial 中的最大值。如果 method=3，则 $p = \text{model}[0]$ ， $q = \text{model}[1]$ 。

说明

函数 `imsls_f_auto_arima` 可确定乘积季节性 $\text{ARIMA}(p, 0, q) \times (0, d, 0)_s$ 模型的参数, 然后使用拟合模型识别离群值并准备预测值。可以指定也可以自动确定此模型的阶数。

`imsls_f_auto_arima` 处理的 $\text{ARIMA}(p, 0, q) \times (0, d, 0)_s$ 模型具有如下形式:

$$\phi(B)\Delta_s^d(Y_t - \mu) = \theta(B)a_t, \quad t = 1, 2, \dots, n,$$

其中

$$\phi(B) = 1 - \phi_1 B - \phi_2 B^2 - \dots - \phi_p B^p, \quad \theta(B) = 1 - \theta_1 B - \theta_2 B^2 - \dots - \theta_q B^q, \quad \Delta_s^d = (1 - B^s)^d$$

并且

$$B^k Y_t = Y_{t-k}.$$

假定 $\phi(B)$ 和 $\theta(B)$ 的所有根都位于单位圆之外。很显然, 如果 $s=1$, 此模型将简化为传统 $\text{ARIMA}(p, d, q)$ 模型。

Y_t 是未观测且无离群值的时序, 其平均值为 μ , 白噪音为 a_t 。此模型称为基本的无离群值模型。函数 `imsls_f_auto_arima` 不假定此序列是可观测序列。它假定观测值可能被一个或多个离群值污染, 因此将其影响添加到基本无离群值序列中:

$$Y_t^* = Y_t + \text{outlier_effect}_t.$$

离群值识别使用 Chen and Liu (1993) 提出的算法。离群值分为以下 5 种类型:

1. 创新
2. 附加
3. 水平平移
4. 临时变化和
5. 无法识别

确定离群值后, `imsls_f_auto_arima` 将通过删除估计的离群值效应来估计数据的无离群值序列表示 Y_t 。

使用有关调整的 $ARIMA(p,0,q) \times (0,d,0)_s$ 模型和已删除离群值的信息，准备无离群值序列的预测值。向这些预测值中添加离群值效应，以生成观测序列 Y_t^* 的预测值。如果没有离群值，则无离群值序列与观测序列的预测值将相同。

模型选择

用户可以选择为 p 、 q 、 s 和 d 指定特定值或让 `imsls_f_auto_arima` 自动选择最佳拟合值。根据变量 `method` 的值，可通过下列三种方法之一执行模型选择。

方法 1: 自动的 $ARIMA(p,0,0) \times (0,d,0)_s$ 选择

此方法先搜索具有最小 AIC 的 $AR(p)$ 表示形式中的噪音数据，其中 $p=0,...,maxlag$ 。

如果已定义 `IMSLD_INITIAL`，则在搜索中包括 `s_initial` 和 `d_initial` 中的值，以查找序列的最佳 $ARIMA(p,0,0) \times (0,d,0)_s$ 表示形式。这里，将对 `s_initial` 中的 p 、 s 和 `d_initial` 中的 d 的每个可能值组合进行检查。然后，将找到的最佳 $ARIMA(p,0,0) \times (0,d,0)_s$ 表示形式用作离群值检测例程的输入。

将分别在 `model[0]`、`model[1]`、`model[2]` 和 `model[3]` 中返回 p 、 q 、 s 和 d 的最优值。

方法 3: 指定的 $ARIMA(p,0,q) \times (0,d,0)_s$ 模型

在第三种方法中，会为 p 、 q 、 s 和 d 提供特定值。 p 和 q 的值必须分别在 `model[0]` 和 `model[1]` 中定义。如果未定义 `IMSLD_S_INITIAL` 和 `IMSLD_D_INITIAL`，则必须在 `model[2]` 和 `model[3]` 中指定值 $s > 0$ 和 $d \geq 0$ 。如果已定义 `IMSLD_S_INITIAL` 和 `IMSLD_D_INITIAL`，将使用 `s_initial` 和 `d_initial` 中输入值的所有可能组合来执行网格搜索，以查找 s 和 d 的最优值。 s 和 d 的最优值可分别在 `model[2]` 和 `model[3]` 中找到。

离群值

Chen and Liu (1993) 的算法可用于识别离群值。在 `num_outliers` 中返回识别的离群值数量。在 `outlier_stat[]` 中返回这些离群值的时间和分类。根据每种离群值类型的标准化统计信息，离群值分为以下五个类别。`outlier_stat` 的第一列会给出出现离群值的时间。在第二列中返回的离群值标识符以下表中的描述为依据：

离群值 标识符	名称	一般性说明
0	(IO) 创新离群值	创新离群值将持久存在。也就是说，出现离群值时存在初始影响。对于将来的所有观测值，此影响将以滞后方式持续存在。滞后系数由基本 $ARIMA(p, 0, q) \times (0, d, 0)_s$ 模型的系数确定。
1	(AO) 附加离群值	附加离群值不持久存在。顾名思义，附加离群值只影响出现离群值时的观测值。因此，附加离群值对将来的预测值没有任何影响。
2	(LS) 水平平移	水平平移离群值将持久存在。它们的影响是从出现离群值之时起，增大或减小序列的平均值。这种平均值移动是突然性的，并且会持久存在。
3	(TC) 临时变化	临时变化离群值将持久存在，类似于水平平移离群值，但两者之间有一个主要区别。与水平平移离群值一样，序列的平均值在出现此离群值时会发生突然变化。但与水平平移离群值不同的是，此平移不是持久性的。TC 离群值逐渐衰减，最终使序列的平均值恢复到其原始值。此衰减速率使用参数 <code>delta</code> 来建模。缺省设置 <code>delta=0.7</code> 是 Chen and Liu (1993) 推荐用于一般情况的值。
4	(UI) 无法识别	如果离群值被识别为最后一个观测值，则该算法无法确定该离群值的分类。进行预测时，会将 UI 离群值视为 IO 离群值。也就是说，其影响会滞后到预测中。

除了附加离群值 (AO) 外，离群值的效应会持续到该离群值之后的观测值。`imsls_f_auto_arima` 生成的预测会考虑这一点。

difference

区分季节性或非季节性时序。

对照表

```
float *imsls_f_difference(int n_observations, float z[], int n_differences, int periods[], ..., 0)
```

double 类型函数为 `imsls_d_difference`。

必选参数

int `n_observations` (输入)

观测值数。

float `z[]` (输入)

包含时序的长度为 `n_observations` 的数组。

int `n_differences` (输入)

要执行的差分数。参数 `n_differences` 必须大于或等于 1。

int `periods[]` (输入)

长度为 `n_differences` 的数组，包含差分 `z` 所在的时间段。

返回值

指向长度为 `n_observations` 的数组的指针，该数组包含已差分的序列。

具有可选参数的对照表

```
float *imsls_f_difference(int n_observations, float z[], int n_differences, int periods[],
    IMSLS_ORDERS, int orders[],
    IMSLS_LOST, int v*n_lost,
    IMSLS_EXCLUDE_FIRST, or
    IMSLS_SET_FIRST_TO_NAN,
    IMSLS_RETURN_USER, float w[],
    0)
```

可选参数

IMSLS_ORDERS, *int* orders[] (输入)

长度为 `n_differences` 的数组，包含在时间段中给定的每个差分的阶数。阶数元素必须大于或等于 0。

IMSLS_LOST, *int* *n_lost (输出)

因为差分时序 `z` 而丢失的观测值数。

IMSLS_EXCLUDE_FIRST 或

IMSLS_SET_FIRST_TO_NAN

如果指定 `IMSLS_EXCLUDE_FIRST`，则会因为差分而从 `w` 中排除第一个 `n_lost`。

差分序列 `w` 的长度为 `n_observations - n_lost`。如果指定

`IMSLS_SET_FIRST_TO_NAN`，则会将前 `n_lost` 个观测值设置为 NaN（不是数字）。

如果 `IMSLS_EXCLUDE_FIRST` 和 `IMSLS_SET_FIRST_TO_NAN` 均未指定，则此设置为缺省设置。

IMSLS_RETURN_USER, *float* w[] (输出)

如果指定，`w` 将包含差分序列。如果同时还指定了 `IMSLS_EXCLUDE_FIRST`，则 `w` 的长度为 `n_observations`。如果指定了 `IMSLS_SET_FIRST_TO_NAN` 或者 `IMSLS_EXCLUDE_FIRST` 和 `IMSLS_SET_FIRST_TO_NAN` 均未指定，则 `w` 的长度为 `n_observations - n_lost`。

说明

函数 `imsls_f_difference` 可对 $n = n_{\text{observations}}$ 个观测值 $\{Z_t\}$ ($t = 1, 2, \dots, n$) 执行 $m = n_{\text{differences}}$ 次连续后向差分，其中时间段为 $si = \text{periods}[i - 1]$ ，阶数为 $di = \text{orders}[i - 1]$ ， $i = 1, \dots, m$ 。

考虑下式为所有 k 给定的后移算子 B

$$B^k Z_t = Z_{t-k}$$

然后，通过下式定义时间段为 s 的后向差分算子：

$$\Delta_s Z_t = (1 - B^s) Z_t = Z_t - Z_{t-s} \quad \text{对于} \quad s > 0。$$

注意，仅为 $t=(s+1), \dots, n$ 定义 $B^s Z_t$ 和 $\Delta^s Z_t$ 。时间段为 s 的重复差分可简化为

$$\Delta_s^d Z_t = (1 - B^s)^d Z_t = \sum_{j=0}^d \frac{d!}{j!(d-j)!} (-1)^j B^{sj} Z_t$$

其中 $d \geq 0$ 为差分阶数。注意

$$\Delta_s^d Z_t$$

仅为 $t = (sd+1), \dots, n$ 定义。

函数 `imsls_f_difference` 中使用的一般差分公式可通过下式给定

$$W_t = \begin{cases} \text{NaN} & t = 1, \dots, n_L \\ \Delta_{s_1}^{d_1} \Delta_{s_2}^{d_2} \dots \Delta_{s_m}^{d_m} Z_t & t = n_L + 1, \dots, n \end{cases}$$

其中， n_L 表示因差分而丢失的观测值数，NaN 表示缺失值代码。注意

$$n_L = \sum_j s_j d_j$$

通过对同构的非平稳时序进行相应差分处理可以得到同构的平稳时序（Box and Jenkins，1976 年，第 85 页）。在对序列应用适当的初步转换后进行差分处理，可在同构平稳的自回归移动平均模型的类中启用模型识别和参数估计。

致命错误

`IMSLS_PERIODS_LT_ZERO`

“period[#]” = #。“period”的所有元素都必须大于 0。

`IMSLS_ORDER_NEGATIVE`

“order[#]” = #。“order”的所有元素都必须是非负数。

`IMSLS_Z_CONTAINS_NAN`

“z[#]” = NaN；“z”不能包含缺失值。“z”的其它元素可能等于 NaN。

box_cox_transform

执行正向或反向 Box-Cox（幂）转换。

对照表

```
float*imsls_f_box_cox_transform(int n_observations,float z[],float power,...,0)
```

double 类型函数为 `imsls_d_box_cox_transform`。

必选参数

int `n_observations` (输入)

`z` 中的观测值数。

float `z[]` (输入)

包含观测值的长度为 `n_observations` 的数组。

float `power` (输入)

Box-Cox（幂）转换中的指数参数。

返回值

指向内部分配的长度为 `n_observations` 的数组的指针，该数组包含转换数据。若要释放此空间，请使用 `imsls_free`。如果无法计算任何值，则返回 `NULL`。

具有可选参数的对照表

```
float*imsls_f_box_cox_transform(int n_observations,float z[],float power,  
                                IMSLS_SHIFT,float shift,  
                                IMSLS_INVERSE_TRANSFORM,  
                                IMSLS_RETURN_USER,float x[],  
                                0)
```

可选参数

IMSL_SHIFT, *float* shift (输入)

将 Box-Cox (幂) 转换中的参数移位。参数移位必须满足关系

$\min(z(i)) + \text{shift} > 0$ 。

缺省设置: shift = 0.0。

IMSL_INVERSE_TRANSFORM

如果指定 IMSL_INVERSE_TRANSFORM, 则执行反向转换。

IMSL_RETURN_USER, *float* x[] (输出)

用户分配的长度为 n_observations 的数组, 其中包含转换数据。

说明

函数 imsls_f_box_cox_transform 可对 $n = n_observations$ 个观测值 $\{Z_t\} (t = 1, 2, \dots, n)$ 执行正向或反向 Box-Cox (幂) 转换。

正向转换在分析线性模型或具有非常规误差或非恒定方差的模型时很有用 (Draper and Smith, 1981 年, 第 222 页)。在时序设置中, 对序列应用适当转换再进行差分处理可在同构平稳的自回归移动平均模型的类中启用模型识别和参数估计。之后可对某些分析结果 (如预测值和预测值的预测限制) 应用反向转换, 从而以原始数据比例的形式表示结果。Box and Jenkins (1976 年, 第 328 页) 针对在时序模型中选择转换进行了简短说明。

Box and Cox (1964) 论述的幂转换类由以下公式定义

$$X_t = \begin{cases} \frac{(Z_t + \xi)^\lambda - 1}{\lambda} & \lambda \neq 0 \\ \ln(Z_t + \xi) & \lambda = 0 \end{cases}$$

其中, 对于所有 t , $Z_t + \xi > 0$ 。由于

$$\lim_{\lambda \rightarrow 0} \frac{(Z_t + \xi)^\lambda - 1}{\lambda} = \ln(Z_t + \xi)$$

因此，幂转换的系列是连续的。

假定 $\lambda = \text{power}$, $\xi = \text{shift}$; 那么, `imsls_f_box_cox_transform` 使用的计算公式为

$$X_t = \begin{cases} (Z_t + \xi)^\lambda & \lambda \neq 0 \\ \ln(Z_t + \xi) & \lambda = 0 \end{cases}$$

其中, 对于所有 t , $Z_t + \xi > 0$ 。计算和 Box-Cox 公式的差别仅在于转换数据的比例和原点。因此, 数据的常规分析不受影响 (Draper and Smith 1981, 第 225 页)。

反向转换的计算公式为

$$X_t = \begin{cases} Z_t^{1/\lambda} - \xi & \lambda \neq 0 \\ \exp(Z_t) - \xi & \lambda = 0 \end{cases}$$

其中 $\{Z_t\}$ 现在表示 `imsls_f_box_cox_transform` 使用参数 λ 和 ξ 针对原始数据的正向转换计算的结果。

致命错误

IMSL_ILLEGAL_SHIFT	“shift” = #, 并且 “z” 的最小元素为 “z[#]” = #。“shift” 并且 “z[#]” = #。“shift” + “z[i]” 必须大于 0, 适用于 $i = 1, \dots, \text{“n_observations”}$ 。“n_observations” = #。
IMSL_BCTR_CONTAINS_NAN	“z” 的一个或多个元素等于 NaN (不是数字)。不允许缺失值。等于 NaN 的 “z” 元素的最小索引为 #。
IMSL_BCTR_F_UNDERFLOW	正向转换。“power” = #。“shift” = #。“z” 的最小元素为 “z[#]” = #。(“z[#]” + “shift”) ^ “power” 将下溢。
IMSL_BCTR_F_OVERFLOW	正向转换。“power” = #。“shift” = #。“z” 的最大元素为 “z[#]” = #。(“z[#]” + “shift”) ^ “power” 将溢出。
IMSL_BCTR_I_UNDERFLOW	反向转换。“power” = #。“z” is “z[#]” = #. exp(“z[#]”) 的最小元素将下溢。
IMSL_BCTR_I_OVERFLOW	反向转换。“power” = #。“z[#]” = #. exp(“z[#]”) 的最大元素将溢出。

IMSLI_BCTR_I_ABS_UNDERFLOW	反向转换。“power” = #。“z” 的绝对值最小的元素为“z[#]” = #。“z[#]” ^ (1/ “power”) 将下溢。
IMSLI_BCTR_I_ABS_OVERFLOW	反向转换。“power” = #。“z” 的绝对值最大的元素为“z[#]” = #。“z[#]” ^ (1/ “power”) 将溢出。

autocorrelation

计算平稳时序的示例自相关函数。

对照表

```
float *imsls_f_autocorrelation (int n_observations, float x[],
                                int lagmax, ...0)
```

double 类型函数为 imsls_d_autocorrelation。

必选参数

int n_observations (输入)

时序 x 中的观测值数。n_observations 必须大于或等于 2。

float x[] (输入)

包含时序的长度为 n_observations 的数组。

int lagmax (输入)

要计算的自协方差、自相关以及自相关的标准误差的最大滞后。lagmax 必须大于或等于 1 且小于 n_observations。

返回值

指向长度为 lagmax + 1 的数组的指针，该数组包含时序 x 的自相关。该数组的第 0 个元素为 1。该数组的第 k 个元素包含滞后 k 的自相关，其中 $k = 1, \dots, \text{lagmax}$ 。

具有可选参数的对照表

```
float*imsls_f_autocorrelation (int n_observations,float x[],
    int lagmax,
    IMSLS_PRINT_LEVEL,int iprint,
    IMSLS_X_MEAN_IN,float x_mean_in,
    IMSLS_X_MEAN_OUT,float *x_mean_out,
    IMSLS_ACV,float **autocovariances,
    IMSLS_ACV_USER,float autocovariances[],
    IMSLS_SEAC,float **standard_errors,int se_option,
    IMSLS_SEAC_USER,float standard_errors[],int se_option,
    IMSLS_RETURN_USER, float autocorrelations[],
    0)
```

可选参数

IMSLS_PRINT_LEVEL, int iprint (输入)
输出选项。
缺省值 = 0。

Iprint	操作
0	不执行任何输出。
1	输出平均值和方差。
2	输出平均值、方差和自协方差。
3	输出平均值、方差、自协方差、自相关以及自相关的标准误差。

IMSLS_X_MEAN_IN, float x_mean_in (输入)
用户输入时序 x 的估计值。

IMSLS_X_MEAN_OUT, float *x_mean_out (输出)
如果指定，则 x_mean_out 为时序 x 的平均值估计。

IMSL5_ACV, *float* **autocovariances (输出)

指向长度为 $\text{lagmax} + 1$ 的数组的指针地址, 该数组包含时序 x 的方差和自协方差。该数组的第 0 个元素是时序 x 的方差。第 k 个元素包含滞后 k 的自协方差, 其中 $k = 1, \dots, \text{lagmax}$ 。

IMSL5_ACV_USER, *float* autocovariances[] (输出)

如果指定, 则 autocovariances 是长度为 $\text{lagmax} + 1$ 的数组, 其中包含时序 x 的方差和自协方差。

IMSL5_SEAC, *float* **standard_errors, *int* se_option (输出)

指向长度为 lagmax 的数组的指针地址, 该数组包含时序 x 的自相关的标准误差。自相关的标准误差的计算方法可通过 se_option 来选择。

se_option	操作
1	使用 Barlett 的公式计算 autocorrelations 的标准误差。
2	使用 Moran 的公式计算 autocorrelations 的标准误差。

IMSL5_SEAC_USER, *float* standard_errors[], *int* se_option (输出)

如果指定, 则 autocovariances 为长度为 lagmax 的数组, 其中包含时序 x 的自相关的标准误差。

IMSL5_RETURN_USER, *float* autocorrelations[] (输出)

如果指定, 则自相关是长度为 $\text{lagmax} + 1$ 的数组, 其中包含时序 x 的自相关。该数组的第 0 个元素是 1。该数组的第 k 个元素包含滞后 k 的自相关, 其中 $k = 1, \dots, \text{lagmax}$ 。

说明

函数 `imsls_f_autocorrelation` 可根据给定的 $n = n_observations$ 观测值样本 $\{x_t\}$ (其中 $t = 1, 2, \dots, n$) 估计平稳时序的自相关函数。

假定

$$\hat{\mu} = x_mean$$

是时序 $\{x_t\}$ 的平均值 μ 的估计值, 其中

$$\hat{\mu} = \begin{cases} \mu, & \mu \text{ 已知} \\ \frac{1}{n} \sum_{t=1}^n X_t, & \mu \text{ 未知} \end{cases}$$

自协方差函数 $\sigma(k)$ 的估计方程如下

$$\hat{\sigma}(k) = \frac{1}{n} \sum_{t=1}^{n-k} (X_t - \hat{\mu})(X_{t+k} - \hat{\mu}), \quad k = 0, 1, \dots, K$$

其中 $K = lagmax$ 。注意

$$\hat{\sigma}(0)$$

是样本方差的估计值。自相关函数 $\rho(k)$ 的估计方程如下

$$\hat{\rho}(k) = \frac{\hat{\sigma}(k)}{\hat{\sigma}(0)}, \quad k = 0, 1, \dots, K$$

注意, 根据定义

$$\hat{\rho}(0) \equiv 1$$

可以选择根据参数 `se_option` 计算可选参数 `IMSLs_SEAC` 的样本自相关的标准误差。一种方法 (Bartlett 1946) 基于平稳时序的样本自相关系数方差的常规渐进表达式，该时序具有独立且均匀分布的标准误差。理论公式为

$$\text{var}\{\hat{\rho}(k)\} = \frac{1}{n} \sum_{i=-\infty}^{\infty} [\rho^2(i) + \rho(i-k)\rho(i+k) - 4\rho(i)\rho(k)\rho(i-k) + 2\rho^2(i)\rho^2(k)]$$

其中

$$\hat{\rho}(k)$$

假定 μ 未知。为了便于计算，对于 $|k| \leq K$ ，将自相关 $r(k)$ 替换为估计值，

$$\hat{\rho}(k)$$

并在假定对于满足 $|k| > K$ 的所有 k ， $r(k) = 0$ 时，设定总和限制。

第二种方法 (Moran 1947) 针对具有独立且均匀分布标准误差的随机过程，利用其样本自相关系数方差的准确公式。理论上的公式为

$$\text{var}\{\hat{\rho}(k)\} = \frac{n-k}{n(n+2)}$$

其中，假定 μ 等于零。注意，此公式不依赖于自相关函数。

partial_autocorrelation

计算平稳时序的样本偏自相关函数。

对照表

`float *imsls_f_partial_autocorrelation(int lagmax, int cf[], ..., 0)`

`double` 类型函数为 `imsls_d_partial_autocorrelation`。

必选参数

int lagmax (输入)

要计算的偏自相关的最大滞后。

float cf[] (输入)

包含时序 *x* 的长度为 lagmax + 1 的数组。

返回值

指向长度为 lagmax 的数组的指针，该数组包含时序 *x* 的偏自相关。

具有可选参数的对照表

```
float *imsls_f_partial_autocorrelation(int lagmax, float cf[],  
                                       IMSLS_RETURN_USER, float partial_autocorrelations[],  
                                       0)
```

可选参数

IMSLS_RETURN_USER, *float* partial_autocorrelations[] (输出)

如果指定，则偏自相关存储在用户提供的长度为 lagmax 的数组中。

说明

在给定 $K = \text{lagmax}$ 样本自相关的情况下，函数 `imsls_f_partial_autocorrelation` 可估计平稳时序的偏自相关

$$\hat{\rho}(k)$$

其中 $k = 0, 1, \dots, K$ 。请考虑下式定义的 AR(*k*) 过程

$$X_t = \phi_{k1}X_{t-1} + \phi_{k2}X_{t-2} + \dots + \phi_{kk}X_{t-k} + A_t$$

其中， ϕ_{kj} 表示该过程中的第 *j* 个系数。估计值的集合

$$\{\hat{\phi}_{kk}\}$$

(其中 $k = 1, \dots, K$) 是样本偏自相关函数。自回归参数

$$\{\hat{\phi}_{kj}\}$$

(其中, $j = 1, \dots, k$) 近似于连续 AR(k) 模型的 Yule-Walker 估计值, 其中 $k = 1, \dots, K$ 。根据样本 Yule-Walker 方程

$$\hat{\rho}(j) = \hat{\phi}_{k1}\hat{\rho}(j-1) + \hat{\phi}_{k2}\hat{\rho}(j-2) + \dots + \hat{\phi}_{kk}\hat{\rho}(j-k), \quad j = 1, 2, \dots, k$$

Durbin (1960) 开发了递归关系, 其中 $k = 1, \dots, K$ 。方程由下式给定

$$\hat{\phi}_{kk} = \begin{cases} \hat{\rho}(1) & k = 1 \\ \frac{\hat{\rho}(k) - \sum_{j=1}^{k-1} \hat{\phi}_{k-1,j} \hat{\rho}(k-j)}{1 - \sum_{j=1}^{k-1} \hat{\phi}_{k-1,j} \hat{\rho}(j)} & k = 2, \dots, K \end{cases}$$

和

$$\hat{\phi}_{kk} = \begin{cases} \hat{\rho}(1) & k = 1 \\ \frac{\hat{\rho}(k) - \sum_{j=1}^{k-1} \hat{\phi}_{k-1,j} \hat{\rho}(k-j)}{1 - \sum_{j=1}^{k-1} \hat{\phi}_{k-1,j} \hat{\rho}(j)} & k = 2, \dots, K \end{cases}$$

此过程容易出现舍入错误, 不应在参数接近非平稳界限时使用。可能的替代方法是使用最小二乘或最大似然法计算连续 AR(k) 模型的估计值 $\{\phi_{kk}\}$ 。根据真实过程为 AR(p) 这一假设, Box and Jenkins (1976 年, 第 65 页) 得出

$$\text{var}\{\hat{\phi}_{kk}\} \simeq \frac{1}{n} \quad k \geq p+1$$

有关偏自相关函数的更多信息, 请参见 Box and Jenkins (1976 年, 第 82–84 页)。

lack_of_fit

在给定适当的相关函数的情况下针对单变量时序或传递函数执行失拟检验。

对照表

```
float *imsls_lack_of_fit(int n_observations, float cf[],
                        int lagmax, int npfree, ..., 0)
```

必选参数

int n_observations (输入)
平稳时序的观测值数。

float cf[] (输入)
包含相关函数的长度为 lagmax+1 的数组。

int lagmax (输入)
相关函数的最大滞后。

int npfree (输入)
时序模型公式中的自由参数数量。npfree 必须大于或等于 0 并小于 lagmax。
Woodfield (1990) 建议 npfree = p + q。

返回值

指向长度为 2 的数组的指针，该数组具有测试统计值 Q 及其 p 值 p。在空假设下，Q 服从具有 lagmax-lagmin+1-npfree 个自由度的近似卡方分布。

具有可选参数的对照表

```
#include <imsls.h>

float *imsls_f_lack_of_fit(int n_observations, float cf[], int lagmax,
                          int npfree,
                          IMSLS_LAGMIN, int lagmin,
                          IMSLS_RETURN_USER, float stat[],
                          0)
```

可选参数

IMSL_S_LAGMIN, *int* lagmin (输入)
相关函数的最小滞后。lagmin 对应于失拟测试统计值的总和和下限。缺省值为 1。

IMSL_S_RETURN_USER, *float* stat[] (输出)
用于存储失拟统计值的用户定义数组。

说明

例程 `imsls_f_lack_of_fit` 可用于在 ARMA 和传递函数模型中检测失拟。这些情况下的典型参数有：

模型	LAGMIN	LAGMAX	NPFREE
ARMA (<i>p</i> , <i>q</i>)	1	$\sqrt{\text{NOBS}}$	<i>p</i> + <i>q</i>
传递函数	0	$\sqrt{\text{NOBS}}$	<i>r</i> + <i>s</i>

在给定下面的适当样本相关函数

$$\hat{\rho}(k)$$

（其中 $k = L, L + 1, \dots, K$, $L = \text{lagmin}$, $K = \text{lagmax}$ ）时，函数 `imsls_f_lack_of_fit` 可对包含 *n* 个观测值的时序或传递函数执行多用途失拟测试。

测试统计值 *Q* 的基本形式为

$$Q = n(n + 2) \sum_{k=L}^K (n - k)^{-1} \hat{\rho}(k)$$

其中 $L = 1$ ，前提是

$$\hat{\rho}(k)$$

是自相关函数。假定模型是充分的, Q 服从具有 $K - L + 1 - m$ 个自由度的卡方分布, 其中 $m = \text{npfree}$ 是模型中估计的参数数量。如果时序的平均值是估计值, Woodfield (1990) 建议不要在模型的估计参数计数中包括此参数。因此, 对于 $\text{ARMA}(p, q)$ 模型, 无论平均值是否是估计值, 均设置 $\text{npfree} = p + q$ 。时序模型的最初形式是 Box and Pierce (1970) 对 Ljung and Box (1978) 讨论的上述版本进行修改而得出的。Box and Jenkins (1976 年, 第 394–395 页) 讨论了传递函数模型的测试范围。

estimate_missing

估计时序中的缺失值。

对照表

```
float *imsls_f_estimate_missing(int n_obs, int tpoints[],
                                float z[], ..., 0)
```

double 类型函数为 `imsls_d_estimate_missing`。

必选参数

int `n_obs` (输入)

时序中的非缺失观测值数。时序不得包含具有 3 个以上缺失值的间隔。

int `tpoints[]` (输入)

长度为 `n_obs` 的矢量, 包含观测时序值所在的时间点 $t_1, \dots, t_{n_obs}$ 。时间点必须严格按照升序排列。缺失值的时间点必须位于开区间 (t_1, t_{n_obs}) 中。

float `z[]` (输入)

包含时序值的长度为 `n_obs` 的矢量。值必须根据矢量 `tpoints` 中的值排序。假定估计缺失值之后的时序包含等距时间点处的值, 其中两个连续时间点之间的距离为 1。如果在时间点 $t_1, \dots, t_{n_obs}$ 观测非缺失时序值, 则当 $t_{i+1} - t_i > 1$ 时在 t_i 与 t_{i+1} $i = 1, \dots, n_obs - 1$ 之间存在缺失值。 t_i 与 t_{i+1} 之间的间隔大小为 $t_{i+1} - t_i - 1$ 。具有非缺失和估计缺失值的时序的总长度为 $t_{n_obs} - t_1 + 1$, 或 `tpoints[n_obs-1]-tpoints[0]+1`。

返回值

指向长度为 $(\text{tpoints}[\text{n_obs}-1]-\text{tpoints}[0]+1)$ 的数组的指针，该数组包含时序以及缺失值的估计。如果出现错误，则返回 `NULL`。

具有可选参数的对照表

```
float *imsls_f_estimate_missing(int n_obs, int tpoints[], float z[],
                                IMSLS_METHOD, int method,
                                IMSLS_MAX_LAG, int maxlag,
                                IMSLS_NTIMES, int *ntimes,
                                IMSLS_MEAN_ESTIMATE, float mean_estimate,
                                IMSLS_CONVERGENCE_TOLERANCE, float convergence_tolerance,
                                IMSLS_RELATIVE_ERROR, float relative_error,
                                IMSLS_MAX_ITERATIONS, int max_iterations,
                                IMSLS_TIMES_ARRAY, int **times,
                                IMSLS_TIMES_ARRAY_USER, int times[],
                                IMSLS_MISSING_INDEX, int **missing_index,
                                IMSLS_MISSING_INDEX_USER, int missing_index[],
                                IMSLS_RETURN_USER, float u_z[],
                                0)
```

可选参数

`IMSLS_METHOD, int method` (输入)

用于估计缺失值的方法：

- 0 — 使用中值。
- 1 — 使用三次样条内插。
- 2 — 使用 $\text{AR}(1)$ 模型的一步向前预测。
- 3 — 使用 $\text{AR}(p)$ 模型的一步向前预测。

缺省设置： `method = 3`

如果选择 `method = 2`，则使用 `method 0` 估计从时间点 t_1+1 或 t_1+2 开始的所有间隔值。如果选择 `method = 3`，并且第一个间隔从 t_1+1 开始，则也由 `method 0` 估计此间隔的值。如果在间隔之前的序列长度（表示为 `len`）大于 1 且小于 $2 \cdot \text{maxlag}$ ，则在此间隔内计算缺失值时，会将 `maxlag` 减小为 $\text{len}/2$ 。

IMSL_MAX_LAG, *int* maxlag (输入)

选择 method = 3 时的最大滞后数。

缺省设置: maxlag = 10

IMSL_NTIMES, *int* *ntimes (输出)

具有估计缺失值的时序中的元素数。注意, ntimes = tpoints[n_obs-1] - tpoints[0] + 1。

IMSL_MEAN_ESTIMATE, *float* mean_estimate (输入)

时序平均值的估计值。

IMSL_CONVERGENCE_TOLERANCE, *float* convergence_tolerance (输入)

用于确定方法 2 中所用的非线性最小二乘法收敛的公差级别。参数

convergence_tolerance 表示为了确定收敛需要两次迭代之间平方和的最小相对减少值。因此, convergence_tolerance 必须大于或等于 0。

缺省设置: $\max\{10^{-10}, \text{eps}^{2/3}\}$ (用于单精度) 和 $\max\{10^{-20}, \text{eps}^{2/3}\}$ (用于双精度), 其中 $\text{eps} = \text{imsls_f_machine}(4)$ 用于单精度, $\text{eps} = \text{imsls_d_machine}(4)$ 用于双精度。

IMSL_RELATIVE_ERROR, *float* relative_error (输入)

方法 2 所使用的非线性方程求解器中使用的停止准则。

缺省设置: 对于单精度, $\text{relative_error} = 100 \times \text{imsls_f_machine}(4)$; 对于双精度, $\text{relative_error} = 100 \times \text{imsls_d_machine}(4)$ 。

IMSL_MAX_ITERATIONS, *int* max_iterations (输入)

方法 2 所使用的非线性方程求解器中允许的最大迭代数。

缺省设置: max_iterations = 200。

IMSL_TIMES_ARRAY, *int* **times (输出)

指向内部分配的长度为

ntimes = tpoints[n_obs-1] - tpoints[0] + 1 的数组的指针地址, 该数组包含具有缺失值估计的时序的时间点。

IMSL_TIMES_ARRAY_USER, *int* times[] (输出)

用户提供数组时间的存储。请参见 IMSL_TIMES_ARRAY。

IMSL_MISSING_INDEX, *int* **missing_index (输出)

指向内部分配的长度为 (ntimes-n_obs) 数组的指针地址, 该数组包含数组时间中缺失值的索引。如果 ntimes-n_obs = 0, 则找不到缺失值并返回 NULL。

IMSL_MISSING_INDEX_USER, *int* missing_index[] (输出)

用户提供数组 missing_index 的存储。请参见 IMSLS_MISSING_INDEX。

IMSL_RETURN_USER, *float* u_z[] (输出)

如果指定, 则 u_z 是长度为 tpoints[n_obs-1]-tpoints[0]+1 的矢量, 其中包含时序值以及缺失值的估计值。

说明

Box、Jenkins 和 Reinsel (1994) 介绍的传统时序分析要求在等距时间点 $t_1, t_1+1, t_1+2, \dots, t_n$ 进行观测。如果缺少观测值, 在确定合适的估计值时将会出现问题。函数 `imsls_f_estimate_missing` 提供 4 种估计方法:

方法 0 用间隔之前的最后四个时序值和间隔之后的前四个值的中值估计间隔中的缺失观测值。如果间隔之前或之后没有足够的可用值, 则相应减少数量。此方法非常快捷简单, 但其使用范围仅限于没有离群值和水平平移的平稳遍历性序列。

方法 1 使用三次样条内插法估计缺失值。在此方法中, 同样对间隔之前的最后四个时序值和间隔之后的前四个值执行插值法。由生成的插值估计缺失值。此方法可使缺失值实现平滑过渡。

方法 2 假定 AR(1) 过程可以恰当描述间隔之前的时序。如果间隔之前的最后一个观测值是在时间点 t_m 观测到的, 则它使用 $t_1, t_1+1, \dots, t_m$ 处的时序值计算原点 t_m 处的一步向前预测。此值用作时间点 t_{m+1} 处的缺失值的估计值。如果还缺少 t_{m+1} 处的值, 则使用时间点 $t_1, t_1+1, \dots, t_{m+1}$ 处的值重新计算 AR(1) 模型, 估计 t_{m+2} 处的值, 依此类推。AR(1) 模型中的系数 ϕ_1 由最小二乘法根据例程 `imsls_f_arma` 在内部计算。

最后，方法 3 使用 AR(p) 模型的一步向前预测估计缺失值。首先，使用在缺失值之前应用于时序的函数 ... 从可能值集合 $\{0, 1, \dots, \max_lag\}$ 中确定最佳 p 并计算所生成 AR(p) 模型的参数 $\phi_1, \dots, \phi_p$ 。由最小二乘法根据 Kitagawa and Akaike (1978) 中介绍的 Householder 转换来估计这些参数。用 μ 表示序列 $y_{t_1}, y_{t_1+1}, \dots, y_{t_m}$ 的平均值后，可通过以下公式计算原点 t_m 处的一步向前预测值 $\hat{y}_{t_m}(1)$

$$\hat{y}_{t_m}(1) = \mu(1 - \sum_{j=1}^p \phi_j) + \sum_{j=1}^p \phi_j y_{t_m+1-j}.$$

此值用作缺失值的估计值。然后，将从 `imsls_f_auto_uni_ar` 开始，对间隔中的每个缺失值重复此过程。所有四种估计方法都按递增时间顺序处理缺失值间隔。

garch

计算 GARCH(p,q) 模型的参数估计值。

对照表

`float *imsls_f_garch(int p, int q, int m, float y[], float xguess[], ..., 0)`

`double` 类型函数为 `imsls_d_garch`。

必选参数

`int p` (输入)

GARCH 参数的数量。

`int q` (输入)

ARCH 参数的数量。

`int m` (输入)

观测的时序的长度。

`float y[]` (输入)

长度为 `m` 的数组，其中包含观测的时序数据。

`float xguess[]` (输入)

长度为 `p + q + 1` 的数组，包含参数数组 `x[]` 的初始值。

返回值

指向长度为 $p + q + 1$ 的参数数组 $x[]$ 的指针，该数组包含 σ^2 的估计值，后跟 q 个 ARCH 参数和 p 个 GARCH 参数。

具有可选参数的对照表

```
float *imsls_f_garch(int p, int q, int m, float y[], float xguess[],  
                    IMSLS_MAX_SIGMA, float max_sigma,  
                    IMSLS_A, float *a,  
                    IMSLS_AIC, float *aic,  
                    IMSLS_RETURN_USER, float x[],  
                    0)
```

可选参数

IMSLS_MAX_SIGMA, float max_sigma, (输入)

返回的估计系数数组中第一个元素 (σ) 的上限值。缺省值 = 10。

IMSLS_A, float *a, (输出)

在估计的参数数组 x 处计算的对数似然函数值。

IMSLS_AIC, float *aic, (输出)

在估计的参数数组 x 处计算的 Akaike 信息准则值。

IMSLS_RETURN_USER, float x[], (输出)

如果指定， x 将返回长度为 $p + q + 1$ 的数组，其中包含 σ^2 的估计值，后跟 q 个 ARCH 参数和 p 个 GARCH 参数。用户提供估计参数数组 x 的存储。

时序 $\{w_t\}$ 的广义自回归条件异方差 (GARCH) 定义为

$$w_t = z_t \sigma_t$$

$$\sigma_t^2 = \sigma^2 + \sum_{i=1}^p \beta_i \sigma_{t-i}^2 + \sum_{i=1}^q \alpha_i w_{t-i}^2,$$

其中, z_t 是独立且均匀分布的标准常规随机变量,

$$0 < \sigma^2 < \max_sigma, \beta_i \geq 0, \alpha_i \geq 0 \text{ 并且}$$

$$\sum_{i=2}^{p+q+1} x(i) = \sum_{i=1}^p \beta_i + \sum_{i=1}^q \alpha_i < 1。$$

上述模型表示为 GARCH(p, q)。 β_i 和 α_i 系数将分别称为 GARCH 和 ARCH 系数。当 $\beta_i = 0$ ($i = 1, 2, \dots, p$) 时, 上述模型将简化为 Engle (1982) 提出的 ARCH(q)。参数的非负性条件隐含非负方差, 需要对 β_i 和 α_i 之和设置条件, 以实现广义平稳性。

在观测数据的经验分析中, 通常发现 GARCH(1,1) 或 GARCH(1,2) 模型会适当考虑条件异方差 (Palm 1996)。此发现类似于基于 ARMA 模型的线性时序分析。

必须注意, 上述模型中的正值和负值对条件方差的影响是对称的, 这一点很重要。实际上, 许多序列都会对条件方差产生强大的非对称影响。考虑到这一现象, Nelson (1991) 提出了指数 GARCH (EGARCH)。Lai (1998) 提出并研究了常规模型类的某些属性, 该类可将 ARCH 和 GARCH 中条件方差的线性关系扩展为非线性方式。

在估计 GARCH(p, q) 中的参数时, 使用最大似然法。长度为 $m = \text{nobs}$ 的观测序列 $\{w_t\}$ 模型的对数似然法为

$$\log(L) = -\frac{m}{2} \log(2\pi) - \frac{1}{2} \sum_{t=1}^m y_t^2 / \sigma_t^2 - \frac{1}{2} \sum_{t=1}^m \log \sigma_t^2,$$

$$\text{其中 } \sigma_t^2 = \sigma^2 + \sum_{i=1}^p \beta_i \sigma_{t-i}^2 + \sum_{i=1}^q \alpha_i w_{t-i}^2.$$

这样，根据 α_i 、 β_i 和 σ 的约束，可实现 $\log(L)$ 的最大化。

在此模型中，如果 $q = 0$ ，由于估计的 Hessian 矩阵是奇异矩阵，因此 GARCH 模型也是奇异的。

在矢量 `xguess` 中输入的参数矢量 `x` 的初始值必须满足某些约束。`xguess` 的第一个元素表示 σ^2 ，必须大于零并小于 `max_sigma`。其余的 $p+q$ 个初始值中的每个值都必须大于或等于 0，且总和必须小于 1。

为保证模型拟合的平稳性，对

$$\sum_{i=2}^{p+q+1} x(i) = \sum_{i=1}^p \beta_i + \sum_{i=1}^q \alpha_i < 1$$

进行内部检查。应在 0 与 1 之间选择初始值。

AIC 的计算方法如下：

$$-2 \log(L) + 2(p+q+1),$$

其中， $\log(L)$ 是对数似然函数的值。

可在例程 `GARCH` 外根据对数似然函数的输出 (A)、Akaike 信息准则 (AIC) 和方差-协方差矩阵 (VAR) 执行统计推断。

参考材料

用户错误

IMSL 尝试检测用户错误，并以可为用户提供尽可能多信息的方式处理这些错误。为此，将识别各种严重级别的错误，还会在函数用途上下文中考虑错误范围；在一种情况下的细小错误在其它情况下可能变得非常严重。IMSL 尝试报告尽量多的可以合理检测到的错误。由于在检测到第一个错误后，将在不确定的上下文中解释输入，因此多个错误会为错误检测带来很难解决的问题。

哪些因素决定错误严重级

有些情况下，用户的输入从数学角度看可能是正确的，但由于计算机算术和所用算法的限制，可能无法准确计算答案。在此情况下，评估的精确度决定错误的严重级。当函数计算多个输出数量时，有些不可计算，但大多数可以计算，此时存在错误情况。此错误的严重级取决于对错误的总体影响的评估。

错误和缺省操作的种类

IMSL C/Stat/Library 中定义了五个错误严重级。每个级别都有一个关联的 PRINT 属性和一个 STOP 属性。这些属性具有缺省设置（YES 或 NO），但也可由用户进行设置。具有多种错误类型的用途是针对不同严重级的错误所采取的操作提供独立控制。从 IMSL 函数返回后，只存在一种错误状态。（代码为 0 的“错误”表示没有错误。）即使出现多个信息性错误，也只会输出一条消息（如果 PRINT 属性为 YES）。不必或无需在调用程序内对其执行任何更正操作的多个错误会导致输出多条消息（如果其严重级的 PRINT 属性为 YES）。除 IMSLS_TERMINAL 之外的任何严重级错误都可能是信息性错误。包含文件 *imsls.h* 将每个 IMSLS_NOTE、IMSL alert、IMSL warning、IMSL fatal、IMSL_TERMINAL、IMSL_WARNING_IMMEDIATE 和 IMSL_FATAL_IMMEDIATE 都定义为枚举数据类型 *imsls_error*。

IMSL_NOTE。发出通知指示可能存在细小错误或只提供有关计算的信息。

缺省属性：PRINT=NO，STOP=NO

IMSL_ALERT。警报指示因为下溢已将函数值设置为 0。

缺省属性：PRINT=NO，STOP=NO

IMSL warning。警告指示存在一种情况，可能需要用户或调用函数执行更正操作。在以下情况下可能会发出警告错误：结果只精确到几个小数位；某些输出可能出现错误，但大部分输出都正确无误；或者违反了分析方法的某些基础假设。通常无需执行更正操作，可以忽略该情况。

缺省属性：PRINT=YES，STOP=NO

IMSL fatal。致命错误指示可能存在严重情况。大多数情况下，用户或调用函数必须执行更正操作才能恢复。

缺省属性：PRINT=YES，STOP=YES

IMSL terminal。终止性错误很严重。这通常是由错误指定造成的，例如将负数指定为方程数。这些错误也可能是由 C 中无法正确诊断的各种编程错误引起的。生成的错误消息可能会令用户难以理解。在此情况下，建议用户认真比较传递给函数的实际参数和文档中提供的虚拟参数说明。在检查参数顺序和数据类型时应特别注意。

终止性错误不是信息性错误，因为程序内的更正操作通常是不合理的。在正常使用中，出现终止性错误时，执行过程会立即终止。如果出现终止性错误，则会输出与多个终止性错误相关的消息。

缺省属性：PRINT=YES，STOP=YES

IMSL warning immediate。即时警告错误与警告错误相同，唯一不同的是它会立即输出。

缺省属性：PRINT=YES，STOP=NO

IMSL fatal immediate。即时致命错误与致命错误相同，唯一不同的是它会立即输出。

缺省属性：PRINT=YES，STOP=YES

用户可以通过调用函数 `imsl_error_options` 来设置 PRINT 和 STOP 属性。

产品支持

与 Visual Numerics 支持联系

在支持保修期内的用户可针对 IMSL C Numerical Library 使用事宜与 Visual Numerics 取得联系。Visual Numerics 可就以下主题提供咨询：

- 文档的清晰性
- 与 Visual Numerics 相关的可能的编程问题
- 选择用于特定问题的 IMSL 库函数或过程

此外，还有数学/统计咨询主题以及程序调试主题。

有关 Visual Numerics 产品支持的联系信息，请参见以下内容：

- <http://www.vni.com/tech/ims1/phone.php>

下面介绍向 Visual Numerics 咨询的过程：

1. 提供您的 Visual Numerics 许可证号
2. 提供产品名称和版本号：IMSL C Numerical Library 7.0 版
3. 提供编译器和操作系统版本号
4. 提供需要帮助的例程名称及问题描述

附录 A

参考

Abe

Abe, S. (2001) Pattern Classification:Neuro-Fuzzy Methods and their Comparison, Springer-Verlag.

Abramowitz and Stegun

Abramowitz, Milton and Irene A. Stegun (editors) (1964), *Handbook of Mathematical Functions with Formulas, Graphs, and Mathematical Tables*, National Bureau of Standards, Washington.

Afifi and Azen

Afifi, A.A. and S.P.Azen (1979), *Statistical Analysis:A Computer Oriented Approach*, 2d ed., Academic Press, New York.

Agresti, Wackerly, and Boyette

Agresti, Alan, Dennis Wackerly, and James M. Boyette (1979), Exact conditional tests for cross-classifications:Approximation of attained significance levels, *Psychometrika*, **44**, 75-83.

Aha

Aha, D. W. (1991).Incremental constructive induction:An instance-based approach.

In Proceedings of the Eighth International Workshop on Machine Learning (pp. 117--121).Evanston, ILL:Morgan Kaufmann.

Ahrens and Dieter

Ahrens, J.H. and U. Dieter (1974), Computer methods for sampling from gamma, beta, Poisson, and binomial distributions, *Computing*, **12**, 223–246.

Ahrens, J.H., and U. Dieter (1985), Sequential random sampling, *ACM Transactions on Mathematical Software*, **11**, 157–169.

Akaike

Akaike, H., (1978), *Covariance Matrix Computation of the State Variable of a Stationary Gaussian Process*, Ann.Inst.Statist.Math.30 , Part B, 499-504.

Akaike et al

Akaike, H. , Kitagawa, G., Arahata, E., Tada, F., (1979), Computer Science Monographs No. 13, The Institute of Statistical Mathematics, Tokyo.

Anderberg

Anderberg, Michael R. (1973), *Cluster Analysis for Applications*, Academic Press, New York.

Anderson

Anderson, T.W.(1971), *The Statistical Analysis of Time Series*, John Wiley & Sons, New York.

Anderson, T. W. (1994) *The Statistical Analysis of Time Series*, John Wiley & Sons, New York.

Anderson and Bancroft

Anderson, R.L. and T.A.Bancroft (1952), *Statistical Theory in Research*, McGraw-Hill Book Company, New York.

Asuncion and Newman

Asuncion, A. and Newman, D.J.(2007), UCI Machine Learning Repository,
<http://www.ics.uci.edu/~mllearn/MLRepository.html>.Irvine, CA:University of California, School of
Information and Computer Science.

Atkinson

Atkinson, A.C.(1979), A Family of Switching Algorithms for the Computer Generation of Beta Random
Variates, *Biometrika*, **66**, 141–145.

Atkinson, A.C.(1985), *Plots, Transformations, and Regression*, Claredon Press, Oxford.

Baker

Baker, J. E. (1987), Reducing Bias and Inefficiency in the Selection Algorithm.*Genetic Algorithms and
their Applications:Proceeding of the Second international Conference on Genetic Algorithms*, 14-21.

Barrodale and Roberts

Barrodale, I., and F.D.K.Roberts (1973), An improved algorithm for discrete L_1 approximation, *SIAM
Journal on Numerical Analysis*, **10**, 839–848.

Barrodale, I., and F.D.K.Roberts (1974), Solution of an overdetermined system of equations in the l_1
norm, *Communications of the ACM*, **17**, 319–320.

Barrodale, I., and C. Phillips (1975), Algorithm 495. Solution of an overdetermined system of linear
equations in the Chebyshev norm, *ACM Transactions on Mathematical Software*, **1**, 264–270.

Bartlett, M. S.

Bartlett, M.S.(1935), Contingency table interactions, *Journal of the Royal Statistics Society Supplement*, **2**, 248–252.

Bartlett, M. S. (1937) Some examples of statistical methods of research in agriculture and applied biology, *Supplement to the Journal of the Royal Statistical Society*, **4**, 137-183.

Bartlett, M. (1937), The statistical conception of mental factors, *British Journal of Psychology*, **28**, 97–104.

Bartlett, M.S.(1946), On the theoretical specification and sampling properties of autocorrelated time series, *Supplement to the Journal of the Royal Statistical Society*, **8**, 27–41.

Bartlett, M.S.(1978), *Stochastic Processes*, 3rd. ed., Cambridge University Press, Cambridge.

Bays and Durham

Bays, Carter and S.D.Durham (1976), Improving a poor random number generator, *ACM Transactions on Mathematical Software*, **2**, 59–64.

Bendel and Mickey

Bendel, Robert B., and M. Ray Mickey (1978), Population correlation matrices for sampling experiments, *Communications in Statistics*, **B7**, 163–182.

Berry

Berry, M. J. A. and Linoff, G. (1997) *Data Mining Techniques*, John Wiley & Sons, Inc.

Best and Fisher

Best, D.J., and N.I.Fisher (1979), Efficient simulation of the von Mises distribution, *Applied Statistics*, **28**, 152–157.

Bishop

Bishop, C. M. (1995) *Neural Networks for Pattern Recognition*, Oxford University Press.

Bishop et al

Bishop, Yvonne M.M., Stephen E. Feinberg, and Paul W. Holland (1975), *Discrete Multivariate Analysis: Theory and Practice*, MIT Press, Cambridge, Mass.

Bjorck and Golub

Bjorck, Ake, and Gene H. Golub (1973), Numerical Methods for Computing Angles Between Subspaces, *Mathematics of Computation*, **27**, 579–594.

Blom

Blom, Gunnar (1958), *Statistical Estimates and Transformed Beta-Variables*, John Wiley & Sons, New York.

Bosten and Battiste

Bosten, Nancy E. and E.L.Battiste (1974), Incomplete beta ratio, *Communications of the ACM*, **17**, 156s–157.

Box and Jenkins

Box, George E.P. and Gwilym M. Jenkins (1970) *Time Series Analysis: Forecasting and Control*, Holden-Day, Inc.

Box, George E.P. and Gwilym M. Jenkins (1976), *Time Series Analysis: Forecasting and Control*, revised ed., Holden-Day, Oakland.

Box and Pierce

Box, G.E.P., and David A. Pierce (1970), Distribution of residual autocorrelations in autoregressive-integrated moving average time series models, *Journal of the American Statistical Association*, **65**, 1509–1526.

Box and Tidwell

Box, G.E.P. and P.W.Tidwell (1962), Transformation of the Independent Variables, *Technometrics*, **4**, 531–550.

Box et al.

Box, George E.P., Jenkins,Gwilym M. and Reinsel G.C., (1994) *Time Series Analysis*, Third edition, Prentice Hall, Englewood Cliffs, New Jersey.

Boyette

Boyette, James M. (1979), Random RC tables with given row and column totals, *Applied Statistics*, **28**, 329–332.

Bradley

Bradley, J.V.(1968), *Distribution-Free Statistical Tests*, Prentice-Hall, New Jersey.

Breiman

Breiman, L., Friedman, J. H., Olshen, R. A. and Stone, C. J. (1984) *Classification and Regression Trees*, Chapman & Hall.For the latest information on CART visit:<http://www.salford-systems.com/cart.php>.

Breslow

Breslow, N.E.(1974), Covariance analysis of censored survival data, *Biometrics*, **30**, 89–99.

Bridel

Bridle, J. S. (1990) Probabilistic Interpretation of Feedforward Classification Network Outputs, with relationships to statistical pattern recognition, in F. Fogelman Soulie and J. Herault (Eds.), *Neuralcomputing:Algorithms, Architectures and Applications*, Springer-Verlag, 227-236.

Brown

Brown, Morton E. (1983), MCDP4F, two-way and multiway frequency tables-measures of association and the log-linear model (complete and incomplete tables), in *BMDP Statistical Software, 1983 Printing with Additions*, (edited by W.J.Dixon), University of California Press, Berkeley.

Brown and Benedetti

Brown, Morton B. and Jacqueline K. Benedetti (1977), Sampling behavior and tests for correlation in two-way contingency tables, *Journal of the American Statistical Association*, **42**, 309–315.

Calvo

Calvo, R. A. (2001) Classifying Financial News with Neural Networks, *Proceedings of the 6th Australasian Document Computing Symposium*.

Chen and Liu

Chen, C. and Liu, L., *Joint Estimation of Model Parameters and Outlier Effects in Time Series*, Journal of the American Statistical Association, Vol. 88, No.421, March 1993.

Cheng

Cheng, R.C.H.(1978), Generating beta variates with nonintegral shape parameters, *Communications of the ACM*, **21**, 317–322.

Chiang

Chiang, Chin Long (1968), *Introduction to Stochastic Processes in Statistics*, John Wiley & Sons, New York.

Clarkson and Jenrich

Clarkson, Douglas B. and Robert B Jenrich (1991), Computing extended maximum likelihood estimates for linear parameter models, submitted to *Journal of the Royal Statistical Society, Series B*, **53**, 417-426.

Coley

Coley, D. A. (1999), *An Introduction to Genetic Algorithms for Scientists and Engineers*, World Scientific Publishing Co.

Conover

Conover, W.J.(1980), *Practical Nonparametric Statistics*, 2d ed., John Wiley & Sons, New York.

Conover and Iman

Conover, W.J. and Ronald L. Iman (1983), *Introduction to Modern Business Statistics*, John Wiley & Sons, New York.

Conover, W. J., Johnson, M. E., and Johnson, M. M

Conover, W. J., Johnson, M. E., and Johnson, M. M. (1981) A comparative study of tests for homogeneity of variances, with applications to the outer continental shelf bidding data, *Technometrics*, **23**, 351-361.

Cook and Weisberg

Cook, R. Dennis and Sanford Weisberg (1982), *Residuals and Influence in Regression*, Chapman and Hall, New York.

Cooper

Cooper, B.E.(1968), Algorithm AS4, An auxiliary function for distribution integrals, *Applied Statistics*, **17**, 190–192.

Cox

Cox, David R. (1970), *The Analysis of Binary Data*, Methuen, London.

Cox, D.R.(1972), Regression models and life tables (with discussion), *Journal of the Royal Statistical Society, Series B, Methodology*, **34**, 187–220.

Cox and Lewis

Cox, D.R., and P.A.W.Lewis (1966), *The Statistical Analysis of Series of Events*, Methuen, London.

Cox and Oakes

Cox, D.R., and D. Oakes (1984), *Analysis of Survival Data*, Chapman and Hall, London.

Cox and Stuart

Cox, D.R., and A. Stuart (1955), Some quick sign tests for trend in location and dispersion, *Biometrika*, **42**, 80–95.

Cranley and Patterson

Cranley, R. and Patterson, T.N.L.(1976), Randomization of Number Theoretic Methods for Multiple Integration, *SIAM Journal of Numerical Analysis*, **13**, 904-914.

D'Agostino and Stevens

D'Agostino, Ralph B. and Michael A. Stevens (1986), *Goodness-of-Fit Techniques*, Marcel Dekker, New York.

Dallal and Wilkinson

Dallal, Gerald E. and Leland Wilkinson (1986), An analytic approximation to the distribution of Lilliefors's test statistic for normality, *The American Statistician*, **40**, 294–296.

Davis and Rabinowitz

Davis, P.J. and Rabinowitz, P. (1984), *Methods of Numerical Integration*, Academic Press, 482-483.

De Jong

De Jong, K. A. (1975), An Analysis of the Behavior of a Class of Genetic Adaptive Systems.(Doctorial dissertation, Univ. of Michigan).*Dissertation Abstracts International* 36(10), 5140B.(University Microfilms No. 76-9381).

Demiroz et al.

Demiroz, G., H. A. Govenir, and N. Ilter (1988), "Learning Differential Diagnosis of Eryhemato-Squamous Diseases using Voting Feature Intervals", *Artificial Intelligence in Medicine*.

Dennis and Schnabel

Dennis, J.E., Jr. and Robert B. Schnabel (1983), *Numerical Methods for Unconstrained Optimization and Nonlinear Equations*, Prentice-Hall, Englewood Cliffs, New Jersey.

Devore

Devore, Jay L (1982), *Probability and Statistics for Engineering and Sciences*, Brooks/Cole Publishing Company, Monterey, Calif.

Doornik

Doornik, J.A.(2005), An Improved Ziggurat Method to Generate Normal Random Samples, <http://www.doornik.com/research/ziggurat.pdf>., University of Oxford.

Draper and Smith

Draper, N.R. and H. Smith (1981), *Applied Regression Analysis*, 2d ed., John Wiley & Sons, New York.

Dunnett and Sobel

Dunnett, C. W. and Sobel, M. (1955), Approximations to the Probability Integral and Certain Percentage Points of a Multivariate Analogue of Student's *t*-distribution.*Biometrika*, **42**, 258-260.

Durbin

Durbin, J. (1960), The fitting of time series models, *Revue Institute Internationale de Statistics*, **28**, 233–243.

Efroymson

Efroymson, M.A.(1960), Multiple regression analysis, *Mathematical Methods for Digital Computers*, Volume 1, (edited by A. Ralston and H. Wilf), John Wiley & Sons, New York, 191–203.

Eklblom

Eklblom, Hakan (1973), Calculation of linear best L_p -approximations, *BIT*, **13**, 292–300.

Eklblom, Hakan (1987), The L_1 -estimate as limiting case of an L_p or Huber-estimate, in *Statistical Data Analysis Based on the L_1 -Norm and Related Methods* (edited by Yadolah Dodge), North-Holland, Amsterdam, 109–116.

Elandt-Johnson and Johnson

Elandt-Johnson, Regina C., and Norman L. Johnson (1980), *Survival Models and Data Analysis*, John Wiley & Sons, New York, 172–173.

Elman

Elman, J. L. (1990) Finding Structure in Time, *Cognitive Science*, 14, 179-211.

Emmett

Emmett, W.G.(1949), Factor analysis by Lawless method of maximum likelihood, *British Journal of Psychology, Statistical Section*, **2**, 90–97.

Engle

Engle, C. (1982), Autoregressive conditional heteroskedasticity with estimates of the variance of U.K. inflation, *Econometrica*, **50**, 987–1008.

Fisher

Fisher, R.A.(1936), The use of multiple measurements in taxonomic problems, *The Annals of Eugenics*, **7**, 179–188.

Fishman

Fishman, George S. (1978), *Principles of Discrete Event Simulation*, John Wiley & Sons, New York.

Fishman and Moore

Fishman, George S. and Louis R. Moore (1982), A statistical evaluation of multiplicative congruential random number generators with modulus , *Journal of the American Statistical Association*, **77**, 129–136.

Forsythe

Forsythe, G.E.(1957), Generation and use of orthogonal polynomials for fitting data with a digital computer, *SIAM Journal on Applied Mathematics*, **5**, 74–88.

Frey and Slate

Frey, P. W. and D. J. Slate.(1991), “Letter Recognition using Holland-style Adaptive Classifiers”.(Machine Learning Vol 6 #2).

Fuller

Fuller, Wayne A. (1976), *Introduction to Statistical Time Series*, John Wiley & Sons, New York.

Furnival and Wilson

Furnival, G.M. and R.W.Wilson, Jr. (1974), Regressions by leaps and bounds, *Technometrics*, **16**, 499–511.

Fushimi

Fushimi, Masanori (1990), Random number generation with the recursion $X_t = X_{t-3p} \oplus X_{t-3q}$, *Journal of Computational and Applied Mathematics*, **31**, 105–118.

Gentleman

Gentleman, W. Morven (1974), Basic procedures for large, sparse or weighted linear least squares problems, *Applied Statistics*, **23**, 448–454.

Genz

Genz, A. (1992), Numerical Computation of Multivariate Normal Probabilities. *J. Comp. Graph Stat.*, **1**, 141-149.

Gibbons

Gibbons, J.D.(1971), *Nonparametric Statistical Inference*, McGraw-Hill, New York.

Girschick

Girschick, M.A.(1939), On the Sampling Theory of Roots of Determinantal Equations, *Annals of Mathematical Statistics*, **10**, 203–224.

Goldberg

Goldberg, D. E. (1989), *Genetic Algorithms in Search, Optimization and Machine Learning*, Addison-Wesley Publishing Co.

Goldberg, D. E. and Deb, K. (1991), A Comparative Analysis of Selection Schemes Used in Genetic Algorithms. In G. Rawlins, Ed., *Foundations of Genetic Algorithms*. Morgan Kaufmann.

Golub and Van Loan

Golub, Gene H. and Charles F. Van Loan (1983), *Matrix Computations*, Johns Hopkins University Press, Baltimore, Md.

Gonin and Money

Gonin, Rene, and Arthur H. Money (1989), *Nonlinear L_p -Norm Estimation*, Marcel Dekker, New York.

Goodnight

Goodnight, James H. (1979), A tutorial on the SWEEP operator, *The American Statistician*, **33**, 149–158.

Graybill

Graybill, Franklin A. (1976), *Theory and Application of the Linear Model*, Duxbury Press, North Scituate, Mass.

Griffin and Redish

Griffin, R. and K.A.Redish (1970), Remark on Algorithm 347:An efficient algorithm for sorting with minimal storage, *Communications of the ACM*, **13**, 54.

Gross and Clark

Gross, Alan J., and Virginia A. Clark (1975), *Survival Distributions:Reliability Applications in the Biomedical Sciences*, John Wiley & Sons, New York.

Gruenberger and Mark

Gruenberger, F., and A.M.Mark (1951), The d^2 test of random digits, *Mathematical Tables and Other Aids in Computation*, **5**, 109–110.

Guerra et al.

Guerra, Victor O., Richard A. Tapia, and James R. Thompson (1976), A random number generator for continuous random variables based on an interpolation procedure of Akima, in *Proceedings of the Ninth Interface Symposium on Computer Science and Statistics*, (edited by David C. Hoaglin and Roy E. Welsch), Prindle, Weber & Schmidt, Boston, 228–230.

Giudici

Giudici, P. (2003) *Applied Data Mining:Statistical Methods for Business and Industry*, John Wiley & Sons, Inc.

Haldane

Haldane, J.B.S.(1939), The mean and variance of χ^2 when used as a test of homogeneity, when expectations are small, *Biometrika*, **31**, 346.

Hamilton

Hamilton, James D., *Time Series Analysis*, Princeton University Press, Princeton (NewJersey), 1994.

Harman

Harman, Harry H. (1976), *Modern Factor Analysis*, 3d ed. revised, University of Chicago Press, Chicago.

Hart et al

Hart, John F., E.W.Cheney, Charles L. Lawson, Hans J. Maehly, Charles K. Mesztenyi, John R. Rice, Henry G. Thacher, Jr., and Christoph Witzgall (1968), *Computer Approximations*, John Wiley & Sons, New York.

Hartigan

Hartigan, John A. (1975), *Clustering Algorithms*, John Wiley & Sons, New York.

Hartigan and Wong

Hartigan, J.A. and M.A.Wong (1979), Algorithm AS 136:A K-means clustering algorithm, *Applied Statistics*, **28**, 100–108.

Hayter

Hayter, Anthony J. (1984), A proof of the conjecture that the Tukey-Kramer multiple comparisons procedure is conservative, *Annals of Statistics*, **12**, 61–75.

Hebb

Hebb, D. O. (1949) *The Organization of Behaviour:A Neuropsychological Theory*, John Wiley.

Heiberger

Heiberger, Richard M. (1978), Generation of random orthogonal matrices, *Applied Statistics*, **27**, 199–206.

Hemmerle.

Hemmerle, William J. (1967), *Statistical Computations on a Digital Computer*, Blaisdell Publishing Company, Waltham, Mass.

Herraman

Herraman, C. (1968), Sums of squares and products matrix, *Applied Statistics*, **17**, 289–292.

Hill

Hill, G.W.(1970), Student's *t*-distribution, *Communications of the ACM*, **13**, 617–619.

Hill, G.W.(1970), Student's *t*-quantiles, *Communications of the ACM*, **13**, 619–620.

Hinkelmann, K and Kemthorne

Hinkelmann, K and Kemthorne, O (1994) *Design and Analysis of Experiments – Vol 1*, John Wiley.

Hinkley

Hinkley, David (1977), On quick choice of power transformation, *Applied Statistics*, **26**, 67–69.

Hoaglin and Welsch

Hoaglin, David C. and Roy E. Welsch (1978), The hat matrix in regression and ANOVA, *The American Statistician*, **32**, 17–22.

Hocking

Hocking, R.R.(1972), Criteria for selection of a subset regression:Which one should be used?, *Technometrics*, **14**, 967–970.

Hocking, R.R.(1973), A discussion of the two-way mixed model, *The American Statistician*, **27**, 148–152.

Hocking, R.R.(1985), *The Analysis of Linear Models*, Brooks/Cole Publishing Company, Monterey, California.

Hopfield

Hopfield, J. J. (1987) Learning Algorithms and Probability Distributions in Feed-Forward and Feed-Back Networks, *Proceedings of the National Academy of Sciences*, 84, 8429-8433.

Holland

Holland, J.H.(1975), *Adaptation in Natural and Artificial Systems*.Ann Arbor:The University of Michigan Press.

Huber

Huber, Peter J. (1981), *Robust Statistics*, John Wiley & Sons, New York.

Hutchinson

Hutchinson, J. M. (1994) *A Radial Basis Function Approach to Financial Timer Series Analysis*, Ph.D. dissertation, Massachusetts Institute of Technology.

Hughes and Saw

Hughes, David T., and John G. Saw (1972), Approximating the percentage points of Hotelling's generalized T_0^2 statistic, *Biometrika*, **59**, 224–226.

Hwang

Hwang, J. T. G. and Ding, A. A. (1997) Prediction Intervals for Artificial Neural Networks, *Journal of the American Statistical Society*, 92(438) 748-757.

Iman and Davenport

Iman, R.L., and J.M.Davenport (1980), Approximations of the critical region of the Friedman statistic, *Communications in Statistics*, **A9(6)**, 571–595.

Jacobs

Jacobs, R. A., Jorday, M. I., Nowlan, S. J., and Hinton, G. E. (1991) Adaptive Mixtures of Local Experts, *Neural Computation*, 3(1), 79-87.

Jennrich and Robinson

Jennrich, R.I. and S.M.Robinson (1969), A Newton-Raphson algorithm for maximum likelihood factor analysis, *Psychometrika*, **34**, 111–123.

Jennrich and Sampson

Jennrich, R.I. and P.F.Sampson (1966), Rotation for simple loadings, *Psychometrika*, **31**, 313–323.

John

John, Peter W.M.(1971), *Statistical Design and Analysis of Experiments*, Macmillan Company, New York.

Jöhnk

Jöhnk, M.D. (1964), Erzeugung von Betaverteilten und Gammaverteilten Zufallszahlen, *Metrika*, **8**, 5–15.

Johnson and Kotz

Johnson, Norman L., and Samuel Kotz (1969), *Discrete Distributions*, Houghton Mifflin Company, Boston.

Johnson, Norman L., and Samuel Kotz (1970a), *Continuous Univariate Distributions-1*, John Wiley & Sons, New York.

Johnson, Norman L., and Samuel Kotz (1970b), *Continuous Univariate Distributions-2*, John Wiley & Sons, New York.

Johnson and Kotz

Johnson, N.L. and Kotz, S. (1972), *Distributions in Statistics:Continuous Multivariate Distributions*, John Wiley & Sons, Inc., New York.

Johnson and Welch

Johnson, D.G., and W.J.Welch (1980), The generation of pseudo-random correlation matrices, *Journal of Statistical Computation and Simulation*, **11**, 55–69.

Jonckheere

Jonckheere, A.R.(1954), A distribution-free k -sample test against ordered alternatives, *Biometrika*, **41**, 133–143.

Jöreskog

Jöreskog, K.G.(1977), Factor analysis by least squares and maximum-likelihood methods, *Statistical Methods for Digital Computers*, (edited by Kurt Enslein, Anthony Ralston, and Herbert S. Wilf), John Wiley & Sons, New York, 125–153.

Kachitvichyanukul

Kachitvichyanukul, Voratas (1982), *Computer generation of Poisson, binomial, and hypergeometric random variates*, Ph.D. dissertation, Purdue University, West Lafayette, Indiana.

Kaiser

Kaiser, H.F.(1963), Image analysis, *Problems in Measuring Change*, (edited by C. Harris), University of Wisconsin Press, Madison, Wis.

Kaiser and Caffrey

Kaiser, H.F. and J. Caffrey (1965), Alpha factor analysis, *Psychometrika*, **30**, 1–14.

Kalbfleisch and Prentice

Kalbfleisch, John D., and Ross L. Prentice (1980), *The Statistical Analysis of Failure Time Data*, John Wiley & Sons, New York.

Keast

Keast, P. (1973) Optimal Parameters for Multidimensional Integration, *SIAM Journal of Numerical Analysis*, **10**, 831-838.

Kemp

Kemp, A.W., (1981), Efficient generation of logarithmically distributed pseudo-random variables, *Applied Statistics*, **30**, 249–253.

Kendall and Stuart

Kendall, Maurice G. and Alan Stuart (1973), *The Advanced Theory of Statistics*, Volume 2:*Inference and Relationship*, 3d ed., Charles Griffin & Company, London.

Kendall, Maurice G. and Alan Stuart (1979), *The Advanced Theory of Statistics*, Volume 2:*Inference and Relationship*, 4th ed., Oxford University Press, New York.

Kendall et al.

Kendall, Maurice G., Alan Stuart, and J. Keith Ord (1983), *The Advanced Theory of Statistics*, Volume 3:*Design and Analysis, and Time Series*, 4th. ed., Oxford University Press, New York.

Kennedy and Gentle

Kennedy, William J., Jr. and James E. Gentle (1980), *Statistical Computing*, Marcel Dekker, New York.

Kohonen

Kohonen, T. (1995), *Self-Organizing Maps*, Third Edition. Springer Series in Information Sciences., New York.

Kuehl, R. O.

Kuehl, R. O. (2000) *Design of Experiments: Statistical Principles of Research Design and Analysis*, 2nd edition, Duxbury Press.

Kim and Jennrich

Kim, P.J., and R.I. Jennrich (1973), Tables of the exact sampling distribution of the two sample Kolmogorov-Smirnov criterion D_{mn} ($m < n$), in *Selected Tables in Mathematical Statistics*, Volume 1, (edited by H. L. Harter and D.B. Owen), American Mathematical Society, Providence, Rhode Island.

Kinderman and Ramage

Kinderman, A.J., and J.G. Ramage (1976), Computer generation of normal random variables, *Journal of the American Statistical Association*, **71**, 893–896.

Kinderman et al.

Kinderman, A.J., J.F. Monahan, and J.G. Ramage (1977), Computer methods for sampling from Student's t distribution, *Mathematics of Computation* **31**, 1009–1018.

Kinnucan and Kuki

Kinnucan, P. and H. Kuki (1968), *A Single Precision INVERSE Error Function Subroutine*, Computation Center, University of Chicago.

Kirk

Kirk, Roger E. (1982), *Experimental Design: Procedures for the Behavioral Sciences*, 2d ed., Brooks/Cole Publishing Company, Monterey, Calif.

Kitagawa and Akaike

Kitagawa, G. and Akaike, H., A Procedure for the modeling of non-stationary time series, *Ann.Inst.Statist.Math.*30 (1978), Part B, 351-363.

Konishi and Kitagawa

Konishi, S. and Kitagawa, G (2008), *Information Criteria and Statistical Modeling*, Springer, New York.

Knuth

Knuth, Donald E. (1981), *The Art of Computer Programming*, Volume 2: *Seminumerical Algorithms*, 2d ed., Addison-Wesley, Reading, Mass.

Kshirsagar

Kshirsagar, Anant M. (1972), *Multivariate Analysis*, Marcel Dekker, New York.

Lachenbruch

Lachenbruch, Peter A. (1975), *Discriminant Analysis*, Hafner Press, London.

Lai

Lai, D. (1998a), Local asymptotic normality for location-scale type processes. *Far East Journal of Theoretical Statistics*, (in press).

Lai, D. (1998b), Asymptotic distributions of the correlation integral based statistics. *Journal of Nonparametric Statistics*, (in press).

Lai, D. (1998c), Asymptotic distributions of the estimated BDS statistic and residual analysis of AR Models on the Canadian lynx data. *Journal of Biological Systems*, (in press).

Laird and Oliver

Laird, N.M., and D. Fisher (1981), Covariance analysis of censored survival data using log-linear analysis techniques, *JASA* **76**, 1231–1240.

Lawless

Lawless, J.F.(1982), *Statistical Models and Methods for Lifetime Data*, John Wiley & Sons, New York.

Lawley and Maxwell

Lawley, D.N. and A.E.Maxwell (1971), *Factor Analysis as a Statistical Method*, 2d ed., Butterworth, London.

Lawrence et al

Lawrence, S., Giles, C. L, Tsoi, A. C., Back, A. D. (1997) Face Recognition:A Convolutional Neural Network Approach, *IEEE Transactions on Neural Networks, Special Issue on Neural Networks and Pattern Recognition*, 8(1), 98-113.

Learmonth and Lewis

Learmonth, G.P. and P.A.W.Lewis (1973), *Naval Postgraduate School Random Number Generator Package LLRANDOM, NPS55LW73061A*, Naval Postgraduate School, Monterey, Calif.

Lee

Lee, Elisa T. (1980), *Statistical Methods for Survival Data Analysis*, Lifetime Learning Publications, Belmont, Calif.

Lehmann

Lehmann, E.L.(1975), *Nonparametrics:Statistical Methods Based on Ranks*, Holden-Day, San Francisco.

Levenberg

Levenberg, K. (1944), A method for the solution of certain problems in least squares, *Quarterly of Applied Mathematics*, **2**, 164–168.

Levene, H.

Levene, H. (1960) In *Contributions to Probability and Statistics: Essays in Honor of Harold Hotelling*, I. Olkin et al. editors, Stanford University Press, 278-292.

Lewis et al.

Lewis, P.A.W., A.S.Goodman, and J.M.Miller (1969), A pseudorandom number generator for the System/360, *IBM Systems Journal*, **8**, 136–146.

Li

Li, L. K. (1992) Approximation Theory and Recurrent Networks, Proc.Int. Joint Conf. On Neural Networks, vol. II, 266-271.

Lilliefors

Lilliefors, H.W.(1967), On the Kolmogorov-Smirnov test for normality with mean and variance unknown, *Journal of the American Statistical Association*, **62**, 534–544.

Lippmann

Lippmann, R. P. (1989) Review of Neural Networks for Speech Recognition, *Neural Computation*, I, 1-38.

Ljung and Box

Ljung, G.M., and G.E.P.Box (1978), On a measure of lack of fit in time series models, *Biometrika*, **65**, 297–303.

Loh

Loh, W.-Y. and Shih, Y.-S.(1997) Split Selection Methods for Classification Trees, *Statistica Sinica*, 7, 815-840. For information on the latest version of QUEST see: <http://www.stat.wisc.edu/~loh/quest.html>.

Longley

Longley, James W. (1967), An appraisal of least-squares programs for the electronic computer from the point of view of the user, *Journal of the American Statistical Association*, **62**, 819–841.

Matsumoto and Nishimure

Makoto Matsumoto and Takuji Nishimura, ACM Transactions on Modeling and Computer Simulation, Vol. 8, No. 1, January 1998, Pages 3–30.

Mandic

Mandic, D. P. and Chambers, J. A. (2001) *Recurrent Neural Networks for Prediction*, John Wiley & Sons, LTD.

Manning

Manning, C. D. and Schütze, H. (1999) *Foundations of Statistical Natural Language Processing*, MIT Press.

Marsaglia

Marsaglia, George (1964), Generating a variable from the tail of a normal distribution, *Technometrics*, **6**, 101–102.

Marsaglia, G. (1968), Random numbers fall mainly in the planes, *Proceedings of the National Academy of Sciences*, **61**, 25–28.

Marsaglia, G. (1972), The structure of linear congruential sequences, in *Applications of Number Theory to Numerical Analysis*, (edited by S. K. Zaremba), Academic Press, New York, 249–286.

Marsaglia, George (1972), Choosing a point from the surface of a sphere, *The Annals of Mathematical Statistics*, **43**, 645–646.

Marsaglia and Tsang

Marsaglia, G. and Tsang, W. W. (2000), The Ziggurat Method for Generating Random Variables, *Journal of Statistical Software*, **5-8**, 1–7.

McCulloch

McCulloch, W. S. and Pitts, W. (1943), A Logical Calculus for Ideas Imminent in Nervous Activity, *Bulletin of Mathematical Biophysics*, **5**, 115-133.

McKean and Schrader

McKean, Joseph W., and Ronald M. Schrader (1987), Least absolute errors analysis of variance, in *Statistical Data Analysis Based on the L_1 -Norm and Related Methods* (edited by Yadolah Dodge), North-Holland, Amsterdam, 297–305.

McKeon

McKeon, James J. (1974), F approximations to the distribution of Hotelling's T_0^2 , *Biometrika*, **61**, 381–383.

McCullagh and Nelder

McCullagh, P., and J.A.Nelder, (1983), *Generalized Linear Models*, Chapman and Hall, London.

Maindonald

Maindonald, J.H.(1984), *Statistical Computation*, John Wiley & Sons, New York.

Marazzi

Marazzi, Alfio (1985), Robust affine invariant covariances in ROBETH, ROBETH-85 document No. 6, Division de Statistique et Informatique, Institut Universitaire de Medecine Sociale et Preventive, Lausanne.

Mardia et al.

Mardia, K.V.(1970), Measures of multivariate skewness and kurtosis with applications, *Biometrics*, **57**, 519–530.

Mardia, K.V., J.T.Kent, J.M.Bibby (1979), *Multivariate Analysis*, Academic Press, New York.

Mardia and Foster

Mardia, K.V. and K. Foster (1983), Omnibus tests of multinormality based on skewness and kurtosis, *Communications in Statistics A, Theory and Methods*, **12**, 207–221.

Marquardt

Marquardt, D. (1963), An algorithm for least-squares estimation of nonlinear parameters, *SIAM Journal on Applied Mathematics*, **11**, 431–441.

Marsaglia

Marsaglia, George (1964), Generating a variable from the tail of a normal distribution, *Technometrics*, **6**, 101–102.

Marsaglia and Bray

Marsaglia, G. and T.A.Bray (1964), A convenient method for generating normal variables, *SIAM Review*, **6**, 260–264.

Marsaglia et al.

Marsaglia, G., M.D. MacLaren, and T.A.Bray (1964), A fast procedure for generating normal random variables, *Communications of the ACM*, **7**, 4–10.

Merle and Spath

Merle, G., and H. Spath (1974), Computational experiences with discrete L_p approximation, *Computing*, **12**, 315–321.

Miller

Miller, Rupert G., Jr. (1980), *Simultaneous Statistical Inference*, 2d ed., Springer-Verlag, New York.

Milliken and Johnson

Milliken, George A. and Dallas E. Johnson (1984), *Analysis of Messy Data, Volume 1: Designed Experiments*, Van Nostrand Reinhold, New York.

Mitchell

Mitchell, M. (1996), *An Introduction to Genetic Algorithms*, MIT Press.

Moran

Moran, P.A.P.(1947), Some theorems on time series I, *Biometrika*, **34**, 281–291.

Moré et al.

Moré, Jorge, Burton Garbow, and Kenneth Hillstrom (1980), *User Guide for [4] MINPACK-1*, Argonne National Laboratory Report ANL-80_74, Argonne, Ill.

Morrison

Morrison, Donald F. (1976), *Multivariate Statistical Methods*, 2nd. ed. McGraw-Hill Book Company, New York.

Muller

Muller, M.E.(1959), A note on a method for generating points uniformly on N-dimensional spheres, *Communications of the ACM*, **2**, 19–20.

Nelson

Nelson, D. B. (1991), Conditional heteroskedasticity in asset returns:A new approach.*Econometrica*, , **59**, 347–370.

Nelson

Nelson, Peter (1989), Multiple Comparisons of Means Using Simultaneous Confidence Intervals, *Journal of Quality Technology*, **21**, 232–241.

Neter

Neter, John (1983), *Applied Linear Regression Models*, Richard D. Irwin, Homewood, Ill.

Neter and Wasserman

Neter, John and William Wasserman (1974), *Applied Linear Statistical Models*, Richard D. Irwin, Homewood, Ill.

Noether

Noether, G.E.(1956), Two sequential tests against trend, *Journal of the American Statistical Association*, **51**, 440–450.

Owen

Owen, D.B.(1962), *Handbook of Statistical Tables*, Addison-Wesley Publishing Company, Reading, Mass.

Owen, D.B.(1965), A Special Case of the Bivariate Non-central t Distribution, *Biometrika*, **52**, 437–446.

Ozaki and Oda

Ozaki, T and Oda H (1978) Nonlinear time series model identification by Akaike's information criterion. *Information and Systems*, Dubuisson eds, Pergamon Press.83-91.

Pao

Pao, Y. (1989) *Adaptive Pattern Recognition and Neural Networks*, Addison-Wesley Publishing.

Palm

Palm, F. C. (1996), GARCH models of volatility. In *Handbook of Statistics*, Vol. 14, 209-240. Eds: Maddala and Rao. Elsevier, New York.

Parker

Parker, D. B., (1985), *Learning Logic*. Technical Report TR-47, Cambridge, MA: MIT Center for Research in computational Economics and Management Science.

Patefield

Patefield, W.M.(1981), An efficient method of generating $R \times C$ tables with given row and column totals, *Applied Statistics*, **30**, 91–97.

Patefield and Tandy

Patefield, W.M.(1981), and Tandy D. (2000) Fast and Accurate Calculation of Owen’s T-Function, *J. Statistical Software*, **5**, Issue 5.

Peixoto

Peixoto, Julio L. (1986), Testable hypotheses in singular fixed linear models, *Communications in Statistics:Theory and Methods*, **15**, 1957–1973.

Petro

Petro, R. (1970), Remark on Algorithm 347:An efficient algorithm for sorting with minimal storage, *Communications of the ACM*, **13**, 624.

Pillai

Pillai, K.C.S.(1985), Pillai’s trace, in *Encyclopedia of Statistical Sciences, Volume 6*, (edited by Samuel Kotz and Norman L. Johnson), John Wiley & Sons, New York, 725–729.

Poli

Poli, I. and Jones, R. D. (1994) A Neural Net Model for Prediction, *Journal of the American Statistical Society*, 89(425) 117-121.

Pregibon

Pregibon, Daryl (1981), Logistic regression diagnostics, *The Annals of Statistics*, **9**, 705–724.

Prentice

Prentice, Ross L. (1976), A generalization of the probit and logit methods for dose response curves, *Biometrics*, **32**, 761–768.

Priestley

Priestley, M.B.(1981), *Spectral Analysis and Time Series*, Volumes 1 and 2, Academic Press, New York.

Quinlan

Quinlan, J. R. (1993).C4.5 Programs for Machine Learning, Morgan Kaufmann.For the latest information on Quinlan’s algorithms see <http://www.rulequest.com/>.

Quinlan (1987).*Simplifying Decision Trees*.Int J Man-Machine Studies 27, pp. 221-234.

Rao

Rao, C. Radhakrishna (1973), *Linear Statistical Inference and Its Applications*, 2d ed., John Wiley & Sons, New York.

Reed

Reed, R. D. and Marks, R. J. II (1999) *Neural Smithing:Supervised Learning in Feedforward Artificial Neural Networks*, The MIT Press, Cambridge, MA.

Ripley

Ripley, B. D. (1994) Neural Networks and Related Methods for Classification, *Journal of the Royal Statistical Society B*, 56(3), 409-456.

Ripley, B. D. (1996) *Pattern Recognition and Neural Networks*, Cambridge University Press.

Robinson

Robinson, Enders A. (1967), *Multichannel Time Series Analysis with Digital Computer Programs*, Holden-Day, San Francisco.

Rosenblatt

Rosenblatt, F. (1958) The Perceptron: A Probabilistic Model for Information Storage and Organization in the Brain, *Psychol.Rev.*, 65, 386-408.

Royston

Royston, J.P.(1982a), An extension of Shapiro and Wilk's W test for normality to large samples, *Applied Statistics*, **31**, 115–124.

Royston, J.P.(1982b), The W test for normality, *Applied Statistics*, **31**, 176–180.

Royston, J.P.(1982c), Expected normal order statistics (exact and approximate), *Applied Statistics*, **31**, 161–165.

Royston, J. P. (1991), Approximating the Shapiro-Wilk W -test for non-normality, *Statistics and Computing*, **2**, 117-119.

Rumelhart

Rumelhart, D. E., Hinton, G. E. and Williams, R. J. (1986) Learning Representations by Back-Propagating Errors, *Nature*, 323, 533-536.

Rumelhart, D. E. and McClelland, J. L. eds.(1986) *Parallel Distributed Processing: Explorations in the Microstructure of Cognition*, 1, 318-362, MIT Press.

Sallas

Sallas, William M. (1990), An algorithm for an L_p norm fit of a multiple linear regression model, *American Statistical Association 1990 Proceedings of the Statistical Computing Section*, 131–136.

Sallas and Lioni

Sallas, William M. and Abby M. Lioni (1988), *Some useful computing formulas for the nonfull rank linear model with linear equality restrictions*, IMSL Technical Report 8805, IMSL, Houston.

Savage

Savage, I. Richard (1956), Contributions to the theory of rank order statistics-the two-sample case, *Annals of Mathematical Statistics*, **27**, 590–615.

Scheffe

Scheffe, Henry (1959), *The Analysis of Variance*, John Wiley & Sons, New York.

Schmeiser

Schmeiser, Bruce (1983), Recent advances in generating observations from discrete random variates, *Computer Science and Statistics: Proceedings of the Fifteenth Symposium on the Interface*, (edited by James E. Gentle), North-Holland Publishing Company, Amsterdam, 154–160.

Schmeiser and Babu

Schmeiser, Bruce W. and A.J.G.Babu (1980), Beta variate generation via exponential majorizing functions, *Operations Research*, **28**, 917–926.

Schmeiser and Kachitvichyanukul

Schmeiser, Bruce and Voratas Kachitvichyanukul (1981), *Poisson Random Variate Generation*, Research Memorandum 81–4, School of Industrial Engineering, Purdue University, West Lafayette, Ind.

Schmeiser and Lal

Schmeiser, Bruce W. and Ram Lal (1980), Squeeze methods for generating gamma variates, *Journal of the American Statistical Association*, **75**, 679–682.

Searle

Searle, S.R.(1971), *Linear Models*, John Wiley & Sons, New York.

Seber

Seber, G.A.F.(1984), *Multivariate Observations*, John Wiley & Sons, New York.

Snedecor and Cochran

Snedecor and Cochran (1967) *Statistical Methods*, 6th edition, Iowa State University Press.

Snedecor and Cochran

Snedecor, George W. and Cochran, William G. (1967) *Statistical Methods*, 6th edition, Iowa State University Press, 296-298.

Snedecor, George W. and Cochran, William G. (1967) *Statistical Methods*, 6th edition, Iowa State University Press, 432–436.

Shampine

Shampine, L.F.(1975), Discrete least-squares polynomial fits, *Communications of the ACM*, **18**, 179–180.

Siegal

Siegal, Sidney (1956), *Nonparametric Statistics for the Behavioral Sciences*, McGraw-Hill, New York.

Singleton

Singleton, R.C.(1969), Algorithm 347:An efficient algorithm for sorting with minimal storage, *Communications of the ACM*, **12**, 185–187.

Smirnov

Smirnov, N.V.(1939), Estimate of deviation between empirical distribution functions in two independent samples (in Russian), *Bulletin of Moscow University*, **2**, 3–16.

Smith and Dubey

Smith, H., and S. D. Dubey (1964), "Some reliability problems in the chemical industry", *Industrial Quality Control*, 21 (2), 1964, 64-70.

Smith

Smith, M. (1993) *Neural Networks for Statistical Modeling*, New York:Van Nostrand Reinhold.

Snedecor and Cochran

Snedecor, George W. and William G. Cochran (1967), *Statistical Methods*, 6th ed., Iowa State University Press, Ames, Iowa.

Sposito

Sposito, Vincent A. (1989), Some properties of L_p -estimators, in *Robust Regression: Analysis and Applications* (edited by Kenneth D. Lawrence and Jeffrey L. Arthur), Marcel Dekker, New York, 23–58.

Spurrier and Isham

Spurrier, John D. and Steven P. Isham (1985), Exact simultaneous confidence intervals for pairwise comparisons of three normal means, *Journal of the American Statistical Association*, **80**, 438–442.

Stablein, Carter, and Novak

Stablein, D.M, W.H.Carter, and J.W.Novak (1981), Analysis of survival data with nonproportional hazard functions, *Controlled Clinical Trials*, **2**, 149–159.

Stahel

Stahel, W. (1981), Robuste Schatzugen: Infinitesimale Opimalitat und Schatzugen von Kovarianzmatrizen, Dissertation no. 6881, ETH, Zurich.

Steel and Torrie

Steel and Torrie (1960) *Principles and Procedures of Statistics*, McGraw-Hill.

Stephens

Stephens, M.A. (1974), EDF statistics for goodness of fit and some comparisons, *Journal of the American Statistical Association*, **69**, 730–737.

Stirling

Stirling, W.D. (1981), Least squares subject to linear constraints, *Applied Statistics*, **30**, 204–212. (See correction, p. 357.)

Stoline

Stoline, Michael R. (1981), The status of multiple comparisons: simultaneous estimation of all pairwise comparisons in one-way ANOVA designs, *The American Statistician*, **35**, 134–141.

Strecok

Strecok, Anthony J. (1968), On the calculation of the inverse of the error function, *Mathematics of Computation*, **22**, 144–158.

Studenmund

Studenmund, A. H. (1992) *Using Economics: A Practical Guide*, New York: Harper Collins.

Swingler

Swingler, K. (1996) *Applying Neural Networks: A Practical Guide*, Academic Press.

Tanner and Wong

Tanner, Martin A., and Wing H. Wong (1983), The estimation of the hazard function from randomly censored data by the kernel method, *Annals of Statistics*, **11**, 989–993.

Tanner, Martin A., and Wing H. Wong (1984), Data-based nonparametric estimation of the hazard function with applications to model diagnostics and exploratory analysis, *Journal of the American Statistical Association*, **79**, 123–456.

Taylor and Thompson

Taylor, Malcolm S., and James R. Thompson (1986), Data based random number generation for a multivariate distribution via stochastic simulation, *Computational Statistics & Data Analysis*, **4**, 93–101.

Tesauro

Tesauro, G. (1990) Neurogammon Wins Computer Olympiad, *Neural Computation*, **1**, 321–323.

Tezuka

Tezuka, S. (1995), *Uniform Random Numbers: Theory and Practice*. Academic Publishers, Boston.

Thompson

Thompson, James R, (1989), *Empirical Model Building*, John Wiley & Sons, New York.

Tong

Tong, Y. L. (1990), *The Multivariate Normal Distribution*, Springer-Verlag, New York.

Tucker and Lewis

Tucker, Ledyard and Charles Lewis (1973), A reliability coefficient for maximum likelihood factor analysis, *Psychometrika*, **38**, 1–10.

Tukey

Tukey, John W. (1962), The future of data analysis, *Annals of Mathematical Statistics*, **33**, 1–67.

Velleman and Hoaglin

Velleman, Paul F. and David C. Hoaglin (1981), *Applications, Basics, and Computing of Exploratory Data Analysis*, Duxbury Press, Boston.

Verdooren

Verdooren, L. R. (1963), Extended tables of critical values for Wilcoxon's test statistic, *Biometrika*, **50**, 177–186.

Wallace

Wallace, D.L.(1959), Simplified Beta-approximations to the Kruskal-Wallis H-test, *Journal of the American Statistical Association*, **54**, 225–230.

Warner

Warner, B. and Misra, M. (1996) Understanding Neural Networks as Statistical Tools, *The American Statistician*, 50(4) 284-293.

Weisberg

Weisberg, S. (1985), *Applied Linear Regression*, 2d ed., John Wiley & Sons, New York.

Werbos

Werbos, P. (1974) *Beyond Regression: New Tools for Prediction and Analysis in the Behavioral Science*, PhD thesis, Harvard University, Cambridge, MA.

Werbos, P. (1990) Backpropagation Through Time: What It Does and How to do It, *Proc. IEEE*, 78, 1550-1560.

Wetzel

Wetzel, A. (1983), *Evaluation of the Effectiveness of Genetic Algorithms in Combinatorial Optimization*, Unpublished manuscript, Univ. of Pittsburg, Pittsburg.

Williams

Williams, R. J. and Zipser, D. (1989) A Learning Algorithm for Continuously Running Fully Recurrent Neural Networks, *Neural Computation*, 1, 270-280.

Witten

Witten, I. H. and Frank, E. (2000) *Data Mining: Practical Machine Learning Tools and Techniques with Java Implementations*, Morgan Kaufmann Publishers.

Woodfield

Woodfield, Terry J. (1990), Some notes on the Ljung-Box portmanteau statistic, *American Statistical Association 1990 Proceedings of the Statistical Computing Section*, 155–160.

Wu

Wu, S-I (1995) Mirroring Our Thought Processes, *IEEE Potentials*, 14, 36-41.

Yates, F.

Yates, F. (1936) A new method of arranging variety trials involving a large number of varieties. *Journal of Agricultural Science*, **26**, 424-455.