

统计方法应用 国家标准汇编

术语符号和统计用表卷

(第二版)

中国标准出版社第四编辑室 编



中国质检出版社
中国标准出版社

统计方法应用国家标准汇编

术语符号和统计用表卷

机中若

(第 二 版)

中国标准出版社第四编辑室 编

中国质检出版社
中国标准出版社

北 京

图书在版编目 (CIP) 数据

统计方法应用国家标准汇编. 术语符号和统计用表卷
/中国标准出版社第四编辑室编. —2 版. —北京: 中国
标准出版社, 2011

ISBN 978-7-5066-6314-4

I. ①统… II. ①中… III. ①统计-方法-国家标准
-汇编-中国②统计-方法术语-国家标准-汇编-中国
IV. ①C81-65

中国版本图书馆 CIP 数据核字 (2011) 第 094632 号

中国质检出版社 出版发行
中国标准出版社
北京市朝阳区和平里西街甲 2 号(100013)
北京市西城区复外三里河北街 16 号(100045)

网址 www.spc.net.cn
电话:(010)64275360 68523946
中国标准出版社秦皇岛印刷厂印刷
各地新华书店经销

*

开本 880×1230 1/16 印张 28 字数 853 千字
2011 年 6 月第二版 2011 年 6 月第二次印刷

*

定价 145.00 元

如有印装差错 由本社发行中心调换
版权专有 侵权必究
举报电话:(010)68510107

前 言

GB/T 3358《统计学词汇及符号》分为以下部分：

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

本部分为 GB/T 3358 的第1部分，等同采用 ISO 3534-1:2006《统计学 词汇及符号 第1部分：一般统计术语与用于概率的术语》。与 ISO 3534-1:2006 相比，订正了原文的错误，修正原文中概念表述不够准确的部分，主要变化如下：

- 删去了 1.24 原文中的注 1；
- 2.38 示例中变异系数的计算式“ $0.99/0.995=0.994\ 97$ ”更正为“ $0.995/0.9=1.105\ 56$ ”；
- 2.69 中“[事件] σ 代数 $\mathcal{S}$ ”中，要求满足的性质 a)“属于 $\mathcal{S}$ ”修订为“ Ω 属于 $\mathcal{S}$ ”。

为便于使用，本部分作了下列编辑性修改：

- 删去了 ISO 前言；
- 为术语的简练起见，在少数术语中，使用中括号表示其中可省略部分。例如：2.5 中，[事件 A 的] 概率(probability [of an event A])，表示此术语实际定义的是“概率(probability)”，其中“事件 A 的”在许多场合可省略。又如 2.34“ r 阶[原点]矩 (moment of order r)”表示原文的“ r 阶矩 (moment of order r)”也称为“ r 阶原点矩”。

本部分代替 GB/T 3358.1—1993《统计学术语 第一部分 一般统计术语》，与 GB/T 3358.1—1993 相比，主要变化如下：

- 名称改为《统计学词汇及符号 第1部分：一般统计术语与用于概率的术语》；
- 对术语条目作了较大的调整：增加了一般统计术语及用于概率的术语；将 GB/T 3358.1—1993 中第4章“观测和测试结果的一般术语”及第5章“抽样方法的一般术语”中的内容移至 GB/T 3358 的第2部分；
- 增加了大量的示例及注释；
- 增加了术语概念图(附录 B、附录 C)及定义标准中的术语所使用的方法的附录 D，并将关于符号的附录 A 改为资料性附录。

本部分的附录 A、附录 B、附录 C 和附录 D 均为资料性附录。

本部分由全国统计方法应用标准化技术委员会提出并归口。

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

本部分主要起草人：冯士雍、陈敏、于丹、崔恒建、吴耀华、丁文兴、汪仁官、于振凡。

本部分于 1993 年首次发布，本次为第一次修订。

引 言

目前版本的 GB/T 3358.1 和 GB/T 3358.2 是兼容的,其共同目标是在一致、准确而简洁的前提下,将定义所需的数学程度限制在最低水平。由于 GB/T 3358.1 是概率和统计的基础术语,所以有必要用相对严格而复杂的数学语言来表述。考虑到 GB/T 3358.2 及其他统计方法应用标准的使用者有时需要查询 GB/T 3358.1 中术语的定义,因此本部分的术语尽可能用通俗的方式来描述,并辅以注释及示例。尽管这些非正式的描述并不能取代正式的定义,但为统计专业以外的人员提供了有效的概念性的定义,能满足这些术语标准的大多数用户的需要。为了进一步适应经常使用 GB/T 3358.2 或 GB/T 6379 等标准的用户,通过注释和示例使 GB/T 3358.1 更易于理解。

一套明确定义的,且相对完整的概率统计术语对统计标准的编制及有效使用是必需的。定义必须足够准确、且具备数学意义上的严格性,使在编制其他统计标准时避免出现概念模糊。当然,对概念的更详细的解释、背景和应用领域可在初等概率统计教材中找到。

资料性附录 B 与附录 C 分别为一般统计术语与用于概率的术语提供了系列概念框图。其中一般统计术语包含六个概念图;用于概率的术语包含四个概念图。某些术语同时出现在几个不同的框图中,从而起到一组概念与另一组概念的联系作用。附录 D 提供了关于概念图的简要介绍及其解释。

这些框图有助于本次修订,因为它们有助于描述不同术语之间的相互联系。这些框图也有助于标准文本的翻译。

除非另有说明,本标准中大部分术语均在一维(单变量)场合下定义。这避免了许多术语在类似条件下进行重复定义。

目 录

GB/T 3358.1—2009	统计学词汇及符号 第1部分:一般统计术语与用于概率的术语	1
GB/T 3358.2—2009	统计学词汇及符号 第2部分:应用统计	61
GB/T 3358.3—2009	统计学词汇及符号 第3部分:实验设计	141
GB/T 4086.1—1983	统计分布数值表 正态分布	175
GB/T 4086.2—1983	统计分布数值表 χ^2 分布	185
GB/T 4086.3—1983	统计分布数值表 t 分布	211
GB/T 4086.4—1983	统计分布数值表 F 分布	227
GB/T 4086.5—1983	统计分布数值表 二项分布	332
GB/T 4086.6—1983	统计分布数值表 泊松分布	420
GB/T 4888—2009	故障树名词术语和符号	431



中华人民共和国国家标准

GB/T 3358.1—2009/ISO 3534-1:2006
代替 GB/T 3358.1—1993

统计学词汇及符号 第 1 部分：一般统计术语与 用于概率的术语

Statistics—Vocabulary and symbols—
Part 1: General statistical terms and terms used in probability

(ISO 3534-1:2006, IDT)

2009-10-15 发布

2010-02-01 实施

中华人民共和国国家质量监督检验检疫总局
中国国家标准化管理委员会

发布

出版说明

统计方法应用国家标准是用数理统计应用技术解决科研、设计、生产、贸易和管理中所遇到的某些实际问题必须遵循的依据,广泛应用于社会生活的各个领域,不仅为重大国家标准的研制提供重要的理论支持和实践指导,还直接应用在生产过程中产品抽样检验和流通领域产品质量监督等方面。因而,统计方法应用国家标准作为我国重要的基础性综合性标准,一直得到全社会的广泛关注。

为满足广大统计方法应用技术人员需要,向读者提供完整而有实用价值的技术资料,我们选编出版了《统计方法应用国家标准汇编》。此套系列汇编分为以下五卷:

- 术语符号和统计用表卷
- 统计分析与数据处理卷
- 抽样检验卷
- 过程控制与测量方法的准确度卷
- 可靠性统计方法卷

此套汇编的各卷根据其相关标准的制修订情况于1999年起陆续出版、再版。本卷为术语符号和统计用表卷(第二版),上一版于1999年出版。此次本卷第二版共收入截至2011年3月发布的术语符号和统计用表方面的国家标准10项。

由于本卷所收录标准的发布年代不同,请读者在使用时注意以下两点:

1. 本卷中与现行国家标准《量和单位》不统一之处及所收集标准在编排格式上的不统一之处未做改动。
2. 本汇编收集的国家标准的属性已在目录上标明(GB或GB/T),年号用四位数字表示。鉴于部分国家标准是在清理整顿前出版的,现尚未修订,故正文部分仍保持原样。读者在使用这些国家标准时,其属性以目录上标明的为准(标准正文“引用标准”中标准的属性请读者注意查对)。

编者

2011年4月

统计学词汇及符号

第 1 部分：一般统计术语与 用于概率的术语

范围

GB/T 3358 的本部分规定了用于标准起草的一般统计术语、用于概率的术语的定义及部分术语的符号。

本部分中的术语分为：

- a) 一般统计术语(第 1 章)；
- b) 用于概率的术语(第 2 章)。

附录 A 列出了本部分推荐使用的符号。

附录 B 和附录 C 是本部分所有术语条目的概念框图。

1 一般统计术语

1.1

总体 population

所考虑对象的全体。

注 1：总体可是真实有限或无限的，也可能是完全虚构的。有时，特别是在调查抽样中也使用“有限总体”；在一些流程性物质抽样中也使用“无限总体”。在第 2 章中，从概率的角度，总体在一定意义上可看作是样本空间(2.1)。

注 2：对于虚构的总体，允许人们想象在不同假定条件下的数据所具有的属性。因此，虚构总体在统计研究的设计阶段，特别是确定适宜样本量时非常有用。虚构总体所含对象数目可以是有限的也可以是无限的。在统计推断中，这是一个对评价统计研究证据强度特别有用的概念。

注 3：下面的例子能帮助理解总体这一概念：若有三个村庄被选中作人口统计或健康研究，总体即由这三个村庄的全体居民构成；若这三个村庄是从某个特定区域中的所有村庄中随机挑选出来的，则总体由该区域中的所有居民构成。

1.2

抽样单元 sampling unit

总体(1.1)划分成若干部分中的每一部分。

注：抽样单元依赖于具体问题中所感兴趣的最小部分。抽样单元可以是个人、一个家庭、一个学校或一个行政单位等。

1.3

样本 sample

由一个或者多个抽样单元(1.2)组成的总体(1.1)的子集。

注 1：根据所研究总体的情况，样本中的每个单元可是真实或抽象的个体，也可能是具体的数值。

注 2：在 GB/T 3358.2 关于样本的定义中，包括一个抽样框的示例。抽样框在从有限总体中抽取随机样本时是必须的。

1.4

观测值 observed value

由样本(1.3)中每个单元获得的相关特性的值。

注 1：常用的同义词是“实现”和“数据”。

注2：本定义并没有指明值的来源或如何被获得。观测值可表示某**随机变量**(2.10)的一次实现，但并不一定如此。它可以是相继用于统计分析的若干值中的一个。正确的推断需要一定的统计假定，但首先要做的是对观测值的计算概括或图形描述。仅当需要解决进一步的问题，如确定观测值落入某一指定集合的概率，统计机制才是重要而本质的。观测值分析的初始阶段通常称为数据分析。

1.5

描述性统计量 descriptive statistics

观测值(1.4)的图形、数值或其他概括性描述。

示例1：数值描述包括样本均值(1.15)、样本极差(1.10)、样本标准差(1.17)等。

示例2：图形描述包括箱线图、示意图、Q-Q图、正态分位图、散点图、多元散点图和直方图等。

1.6

随机样本 random sample

由随机抽取的方法获得的**样本**(1.3)。

注1：本定义比GB/T 3358.2给出的定义限制要少，样本允许来自无限总体。

注2：当从有限**样本空间**(2.1)中抽取 n 个抽样单元组成样本时， n 个抽样单元的任意一种组合都会以特定的概率(2.5)被抽中。对于调查抽样方案而言，每一种可能组合被抽中的概率可事先计算。

注3：对有限样本空间的调查抽样，随机样本可以通过不同的抽样方法得到，如分层随机抽样、随机起点的系统抽样、整群抽样、与辅助变量的大小成比例的概率抽样以及其他可能的抽样。

注4：本定义一般是指实际**观测值**(1.4)。这些观测值被认为是**随机变量**(2.10)的实现，其中每个观测值都对应于一个随机变量。当由随机样本构造**估计量**(1.12)、**统计检验**(1.48)的检验统计量或**置信区间**(1.28)时，本定义是指从样本中的抽象个体得到的随机变量而不是这些随机变量的实际观测值。

注5：无限总体中的随机样本一般是从样本空间中重复抽取产生的。根据注4的解释，此时样本由独立同分布的随机变量组成。

1.7

简单随机样本 simple random sample

〈有限总体〉给定样本量的每个子集都有相等的被抽选概率的**随机样本**(1.6)。

注：此处的定义与GB/T 3358.2中的定义是一致的，仅在措辞上稍有不同。

1.8

统计量 statistic

由**随机变量**(2.10)完全确定的函数。

注1：在1.6注4的意义下，统计量是**随机样本**(1.6)中随机变量的函数。

注2：按注1，若 $\{X_1, X_2, \dots, X_n\}$ 是来自未知均值(2.35) μ 和未知标准差(2.37) σ 的正态分布(2.50)的随机样本，则**样本均值**(1.15) $(X_1 + X_2 + \dots + X_n)/n$ 是一个统计量；而 $[(X_1 + X_2 + \dots + X_n)/n] - \mu$ 不是统计量，因为它包含了未知**参数**(2.9) μ 。

注3：相应于数理统计中的表述，此处给出的是统计量的一种技术性定义。英语中，统计量(statistic)的复数形式就是统计学(statistics)，它是一门包括了统计方法应用标准中所叙述的分析方法的技术学科。

1.9

次序统计量 order statistic

由**随机样本**(1.6)中的**随机变量**(2.10)的值，依非降次序排列所确定的**统计量**(1.8)。

示例：假设样本观测值为9,13,7,6,13,7,19,6,10,7，则次序统计量的观测值为：6,6,7,7,7,9,10,13,13,19。这些值是 $X_{(1)}, \dots, X_{(10)}$ 的一次实现。

注1：假设**随机样本**(1.6)的观测值(1.4)为 $\{x_1, x_2, \dots, x_n\}$ ，按非降的次序排列为 $x_{(1)} \leq \dots \leq x_{(k)} \leq \dots \leq x_{(n)}$ ，则 $(x_{(1)}, \dots, x_{(k)}, \dots, x_{(n)})$ 是次序统计量 $(X_{(1)}, \dots, X_{(k)}, \dots, X_{(n)})$ 的观测值， $x_{(k)}$ 为第 k 个次序统计量的观测值。

注2：在实际应用中，为获得一组数据的次序统计量，即是将数据按照注1中所述方式进行排序。将一组数据按上述方法排序后，还可获得其他几个术语定义的有用的统计量，如1.10、1.11等。

注3：次序统计量涉及按照非降次序排列后的位置来识别的样本值。正如示例所示，将样本值(随机变量的实现)排序比将未观测的随机变量排序更容易理解。它可以通过按照非降次序排列的**随机样本**(1.6)来理解随机变

量。比如, n 个随机变量的最大值可以先于它的实现值来研究。

注 4: 单个次序统计量是随机变量的一个特定函数。这个函数可以简单地由其在随机变量排序集合中的位置或序次(称为秩)来确定。

注 5: 结点值会引起一些潜在的问题,特别是对于离散随机变量或者是低分辨的实现。用“非降”而不是“递增”的说法可解决这个问题。需要强调的是结点值都要保留而不能合并成一个。在上面的示例中,“6”有两个实现,所以“6”是结点值。

注 6: 排序按照随机变量的实数值进行,而不是按照其绝对值进行。

注 7: 次序统计量 $(X_{(1)}, \dots, X_{(k)}, \dots, X_{(n)})$ 组成 n 维随机变量, n 是样本中观测值的个数。

注 8: 次序统计量的分量也是次序统计量,而且保持其在原样本排序中的位置标识。

注 9: 最小值,最大值以及样本量为奇数时的样本中位数(1.13)都是特殊的次序统计量。比如样本量为 11,那么 $X_{(1)}$ 是最小值, $X_{(11)}$ 是最大值, $X_{(6)}$ 是样本中位数。

1.10

样本极差 sample range

最大次序统计量(1.9)与最小次序统计量的差。

示例:在 1.9 中的示例中,样本极差的观测值为 $19-6=13$ 。

注:在统计过程控制中,尤其当样本量相对比较小时,样本极差通常用来监测过程的离散程度随时间的变化。

1.11

中程数 mid-range

最大和最小次序统计量(1.9)的平均值(1.15)。

示例:1.9 的示例中,中程数的观测值为 $(6+19)/2=12.5$ 。

注:中程数能够对较小数据集的中心提供一种快捷而简单的估计。

1.12

估计量 estimator

$\hat{\theta}$

用于对参数 θ 估计(1.36)的统计量(1.8)。

注 1: 样本均值(1.15)是总体均值(2.35) μ 的一个估计量。例如,对于正态分布(2.50),样本均值是总体均值 μ 的估计量。

注 2: 要估计总体的特征(如一维(元)分布(2.16)的众数(2.27)),一个合适的估计量可以是分布参数估计量的函数,也可以是随机样本(1.6)的复杂函数。

注 3: 此处所讲的“估计量”是一个宽泛的概念。它包括某参数的点估计,也包括用于预测的区间估计。估计量也包括该估计量和其他特殊形式的统计量。另见 1.36 注的讨论。

1.13

样本中位数 sample median

若样本量(见 GB/T 3358.2—2009,1.2.26) n 为奇数,则是第 $(n+1)/2$ 个次序统计量(1.9);若样本量 n 是偶数,则是第 $n/2$ 与第 $(n/2)+1$ 个次序统计量之和除以 2。

示例:续 1.9 的示例,8 为样本中位数的一个实现,此时样本量为 10(偶数),第 5 和第 6 个次序统计量分别为 7 和 9,其平均值为 8。尽管严格来说样本中位数是作为一个随机变量来定义的,但在实际中也说“样本中位数为 8”。

注 1: 对于样本量为 n 的随机样本(1.6),其随机变量(2.10)按照非降顺序从 1 到 n 排列,如果样本量为奇数,则样本中位数为第 $(n+1)/2$ 个随机变量,如果样本量为偶数,则样本中位数为第 $(n/2)$ 个与第 $(n+1)/2$ 个随机变量的平均值。

注 2: 从概念上讲,对一个没有观测到的随机变量进行排序似乎是不可能的。但不经观测也可理解次序统计量的结构。在实际中,通过获得观测值并对其进行排序,从而得到次序统计量的实现。这些实现值可用于解释次序统计量的结构。

注 3: 样本中位数是分布中间位置的一个估计,各有一半的样本单元大于等于或小于等于它。

注 4: 样本中位数在实际问题中是有用的,它提供了一个对数据极端值不敏感的估计量。例如,中位收入和中位房价都是常用的统计指标。

1. 14

k 阶样本矩 sample moment of order k随机样本(1.6)中随机变量(2.10)的 k 次幂的和除以和中的项数。注1: 对于样本量为 n 的随机样本 $\{X_1, X_2, \dots, X_n\}$, k 阶样本矩为:

$$\frac{1}{n} \sum_{i=1}^n X_i^k$$

注2: 本术语也称为 k 阶样本原点矩。

注3: 一阶样本矩即为样本均值(1.15)。

注4: 虽然本定义中 k 可取任意值,但在实际中常用的是 $k=1$ [样本均值(1.15)], $k=2$ [与样本方差(1.16)和样本标准差(1.17)有关], $k=3$ [与样本偏度系数(1.20)有关]和 $k=4$ [与样本峰度系数(1.21)有关]的情形。

1. 15

样本均值 sample mean**平均数 average****算术平均值 arithmetic mean**

随机样本(1.6)中随机变量(2.10)的和除以和中的项数。

示例:续 1.9 中的示例,观测值的和为 97,样本量为 10,样本均值的实现为 9.7。

注1: 在 1.8 中注3的意义下,样本均值作为统计量是随机样本中随机变量的函数。必须区分统计量与由随机样本中观测值(1.4)计算得出的样本均值的数值。

注2: 样本均值作为统计量,常用作总体均值(2.35)的估计量。算术平均值是它的同义词。

注3: 对样本量为 n 的随机样本 $\{X_1, X_2, \dots, X_n\}$, 样本均值为: $\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$ 。

注4: 样本均值就是一阶样本矩。

注5: 样本量为 2 时,样本均值、样本中位数(1.13)和中程数(1.11)皆相同。

1. 16

样本方差 sample variance S^2

随机样本(1.6)中随机变量(2.10)与样本均值(1.15)差的平方和用和中项数减 1 除。

示例:续 1.9 中的示例,样本观测值与样本均值差的平方和为 158.10,样本量 10 减 1 为 9,计算得样本方差为 17.57。

注1: 样本方差 S^2 作为统计量(1.8),是随机样本中随机变量的函数。必须区分这个统计量与根据随机样本观测值(1.4)计算得出的样本方差的数值,该值称为经验样本方差或观测样本方差,通常记作 s^2 。注2: 对样本量为 n 的随机样本 $\{X_1, X_2, \dots, X_n\}$, 样本均值为 $\bar{X}$, 则

$$S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$$

注3: 样本方差作为一个统计量“差不多”等于该随机变量(2.10)与样本均值(1.15)差的平方的平均数(其中“差不多”是指这里平均用 $n-1$ 而不是用 n 作分母),用 $n-1$ 作分母是为总体方差(2.36)提供一个无偏估计量(1.34)。注4: $n-1$ 称为自由度(2.54)。注5: 样本方差可以近似认为是中心化样本随机变量(2.31)的二阶样本矩(仅以 $n-1$ 代替 n)。

1. 17

样本标准差 sample standard deviation S

样本方差(1.16)的非负平方根。

示例:续 1.9 中的示例,观测样本方差为 17.57,观测样本标准差为 4.192。

注1: 实际中样本标准差用来估计总体标准差(2.37)。再次强调 S 也是一个随机变量(2.10),而并不是随机样本(1.6)的实现。

注2: 样本标准差是分布(2.11)离散程度的一个度量。

1.18

样本变异系数 sample coefficient of variation

样本标准差(1.17)除以非零样本均值(1.15)的绝对值。

注：变异系数通常表示成百分数。

1.19

标准化样本随机变量 standardized sample random variable

随机变量(2.10)与其样本均值(1.15)的差除以样本标准差(1.17)。

示例：续1.9中的示例，观测样本均值为9.7，观测样本标准差为4.192，观测标准化随机变量(表示为两位小数)为：-0.17;0.79;-0.64;-0.88;0.79;-0.64;2.22;-0.88;0.07;-0.62。

注1：标准化样本随机变量应区别于理论上的标准化随机变量(2.33)。将随机变量标准化的目的在于使得其均值为0、标准差为1，便于解释和比较。

注2：标准化样本观测值的观测样本均值为0，观测样本标准差为1。

1.20

样本偏度系数 sample coefficient of skewness

随机样本(1.6)的标准化样本随机变量(1.19)三次幂的算术平均值。

示例：续1.9中的示例，观测样本偏度系数的计算结果为0.97188。如本例中的样本量为10的情形，样本偏度系数不够稳定，因此应谨慎使用。根据注1给出的另一公式计算出的值为1.34983。

注1：对应于定义中公式为：

$$\frac{1}{n} \sum_{i=1}^n \left(\frac{X_i - \bar{X}}{S} \right)^3$$

有些统计软件里使用下面的公式修正样本偏度系数的偏倚(1.33)：

$$\frac{n}{(n-1)(n-2)} \sum_{i=1}^n Z_i^3$$

其中：

$$Z_i = \frac{X_i - \bar{X}}{S}$$

当样本量很大时，两个公式的差别可以忽略。当 $n=10, 100, 1\,000$ 时，修正估计值与定义中的估计值之比分别为1.389, 1.031, 1.003。

注2：偏度系数是对分布不对称性的度量，如果偏度系数接近0意味着真实分布近似对称。偏度系数不为零时意味着在某一侧尾部可能有极端值。有偏的数据也会在样本均值(1.15)与样本中位数(1.13)的差异上体现出来。

正偏(右偏)数据表明可能有少数大的极端值。同样，负偏(左偏)数据表明可能有少数小的极端值。

注3：样本偏度系数也是标准化样本随机变量(1.19)的三阶样本矩。

1.21

样本峰度系数 sample coefficient of kurtosis

随机样本(1.6)的标准化样本随机变量(1.19)四次幂的算术平均值。

示例：续1.9中的示例，观测样本峰度系数的计算结果为2.67419。如本例中的样本量为10的情形，样本峰度系数极不稳定，因此应谨慎使用。统计软件包在计算样本峰度系数时常进行了各种修正(参见2.40中的注3)。应用注1中的另一公式计算的值为0.43605。不能直接比较2.67419和0.43605这两个数值。为此，应将2.67419减去3(正态分布的峰度系数为3)，其差为-0.32581，这个数值可与0.43605进行比较。

注1：与定义对应的公式是：

$$\frac{1}{n} \sum_{i=1}^n \left(\frac{X_i - \bar{X}}{S} \right)^4$$

一些统计软件包使用下面公式来修正样本峰度系数的偏倚(1.33)，它表示对正态分布峰度系数(等于3)的偏离：

$$\frac{n(n+1)}{(n-1)(n-2)(n-3)} \sum_{i=1}^n Z_i^4 - \frac{3(n-1)^2}{(n-2)(n-3)}$$

其中：

$$Z_i = \frac{X_i - \bar{X}}{S}$$

当 n 充分大时，上式第二项近似为 3。有时为了强调与正态分布的比较，峰度表示为如 2.40 中定义的值减去 3。显然，实际应用者需要注意到统计软件包中是否包含任何修正。

注 2：峰度描述了(单峰)分布的重尾程度。对正态分布(2.50)，由于抽样随机性，样本峰度系数一般只近似，而不是恰好为 3。在实际应用中正态的峰度提供了一个基准值：峰度值小于 3 的分布(2.11)有比正态轻的尾部；峰度值大于 3 的分布有比正态重的尾部。

注 3：对于峰度观测值大于 3 很多的情形，一种可能是因为真实分布的尾部比正态尾部重，另一可能是分布中存在潜在的离群值。

注 4：样本峰度系数可认为是标准化随机变量的四阶样本矩。

1.22

样本协方差 sample covariance

S_{XY}

随机样本(1.6)中两个随机变量(2.10)对各自样本均值(1.15)的离差的乘积之和被求和项数减 1 除。

示例 1：考虑下列三个变量的 10 组观测值。在这个示例中，只考虑 x 和 y 。

表 1 示例 1 的观测结果

i	1	2	3	4	5	6	7	8	9	10
x	38	41	24	60	41	51	58	50	65	33
y	73	74	43	107	65	73	99	72	100	48
z	34	31	40	28	35	28	32	27	27	31

X 的观测样本均值是 46.1， Y 的观测样本均值是 75.4， X 与 Y 的样本协方差等于：

$$[(38-46.1) \times (73-75.4) + (41-46.1) \times (74-75.4) + \dots + (33-46.1) \times (48-75.4)]/9 = 257.178$$

示例 2：在上例的表中，考虑 y 和 z ， Z 的观测样本均值是 31.3， Y 与 Z 的样本协方差等于：

$$[(73-75.4) \times (34-31.3) + (74-75.4) \times (31-31.3) + \dots + (48-75.4) \times (31-31.3)]/9 = -54.356$$

注 1：作为统计量(1.8)，样本协方差是样本量为 n 的随机变量对：(X_1, Y_1), (X_2, Y_2), ..., (X_n, Y_n) 在(1.6)注 3 意义下的函数。这个统计量需要与随机样本中由抽样单元(1.2) (x_1, y_1), (x_2, y_2), ..., (x_n, y_n) 的观测值计算得到样本协方差的数值相区别。后者称为经验样本协方差或观测样本协方差。

注 2：样本协方差 S_{XY} 由下式给出：

$$\frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})(Y_i - \bar{Y})$$

注 3：用 $n-1$ 除是为总体协方差(2.43)提供一个无偏估计量(1.34)。

注 4：表 1 的示例包含 3 个变量，而协方差定义中只涉及 2 个变量。在实际应用中经常会遇到多个变量的情况。

1.23

样本相关系数 sample correlation coefficient

r_{xy}

样本协方差(1.22)用相应样本标准差(1.17)的乘积来除。

示例 1：续 1.22 中的示例 1。 X 的观测标准差为 12.945， Y 的观测标准差为 21.329。从而 X 和 Y 的观测样本相关系数为：

$$\frac{257.118}{12.948 \times 21.329} = 0.9312$$

示例 2：继续 1.22 的示例 2， Y 的观测标准差为 21.329， Z 的观测标准差为 4.165。从而 Y 和 Z 的观测样本相关系数为：

$$\frac{-54.356}{21.329 \times 4.165} = -0.612$$

注 1: 样本相关系数的计算公式如下:

$$\frac{\sum_{i=1}^n (X_i - \bar{X})(Y_i - \bar{Y})}{\sqrt{\sum_{i=1}^n (X_i - \bar{X})^2 \sum_{i=1}^n (Y_i - \bar{Y})^2}}$$

这个表达式等价于样本协方差与两方差乘积的平方根的比。有时用 r_{xy} 表示样本相关系数。观测样本相关系数是基于实现值 $(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)$ 的。

注 2: 观测样本相关系数取值在 $[-1, 1]$ 之间。取值接近于 1 表示强的正相关; 取值接近于 -1 表示强的负相关。取值接近于 1 或 -1 表明数据点近似在一条直线上。

1.24

标准误差 standard error

$\sigma_{\hat{\theta}}$

估计量(1.12) $\hat{\theta}$ 的标准差(2.37)。

示例: 如果以样本均值(1.15)作为总体均值(2.35)的一个估计, 且随机变量(2.10)的标准差为 σ , 则样本均值的标准误差为 $\sigma/\sqrt{n}$, 其中 n 是样本中观测值的个数。标准误差的一个估计是 $S/\sqrt{n}$, 其中 S 是样本标准差(1.17)。

注: 不存在反义词“非标准”误差。通常应用中, 标准误差特指样本均值的标准差, 记为 $\sigma_{\bar{X}}$, 此时也常简称为“标准误差”。

1.25

区间估计 interval estimator

由一个上限统计量和一个下限统计量(1.8)所界定的区间。

注 1: 区间的一个端点可以是 $+\infty, -\infty$ 或是参数值的一个自然界限。如“0”是总体方差(2.36)区间估计的一个自然下限。在此情形, 区间称为是单侧的。

注 2: 区间估计可结合参数(2.9)估计(1.36)给出。区间估计通常是以假定在重复抽样下, 区间包含所估计的参数确定比例或其他某种概率意义下给出的。

注 3: 区间估计通常有三种: 参数的置信区间(1.28), 对未来观测的预测区间(1.30)和分布(2.11)被包含一个确定比例的统计容忍区间(1.26)。

1.26

统计容忍区间 statistical tolerance interval

在规定置信水平下, 由随机样本(1.6)确定的至少覆盖抽样总体(1.1)的指定比例的区间。

注: 这里“置信”一词是指在大量重复意义下, 所构造区间应至少包含抽样总体的指定比例。

1.27

统计容忍限 statistical tolerance limit

表示统计容忍区间(1.26)端点的统计量(1.8)。

注: 统计容忍限可为以下两种情况的一种:

- 单侧容忍限, 即单侧的统计容忍上限或单侧的统计容忍下限, 此时另一个容忍限为随机变量的自然界限;
- 双侧容忍限, 此时有两个统计容忍限。

1.28

置信区间 confidence interval

参数(2.9) θ 的区间估计(1.25) (T_0, T_1) , 其中作为区间限的统计量(1.8) T_0, T_1 , 满足 $P[T_0 < \theta < T_1] \geq 1 - \alpha$ 。

注 1: 置信度反映了在同一条件下大量重复随机抽样(1.6)中, 置信区间包含参数真值的比例。置信区间并不能反映观测到的区间包含参数真值的概率(2.5)(观测到的区间只能是包含要么不包含参数真值)。

注 2: 一个与置信区间相关的量是 $100(1 - \alpha)\%$, 称为置信系数或置信水平, 其中 α 是一个小的数。对任意确定但未知的总体 θ 值, $P[T_0 < \theta < T_1] \geq 1 - \alpha$ 。置信系数通常取为 95% 或 99%。