

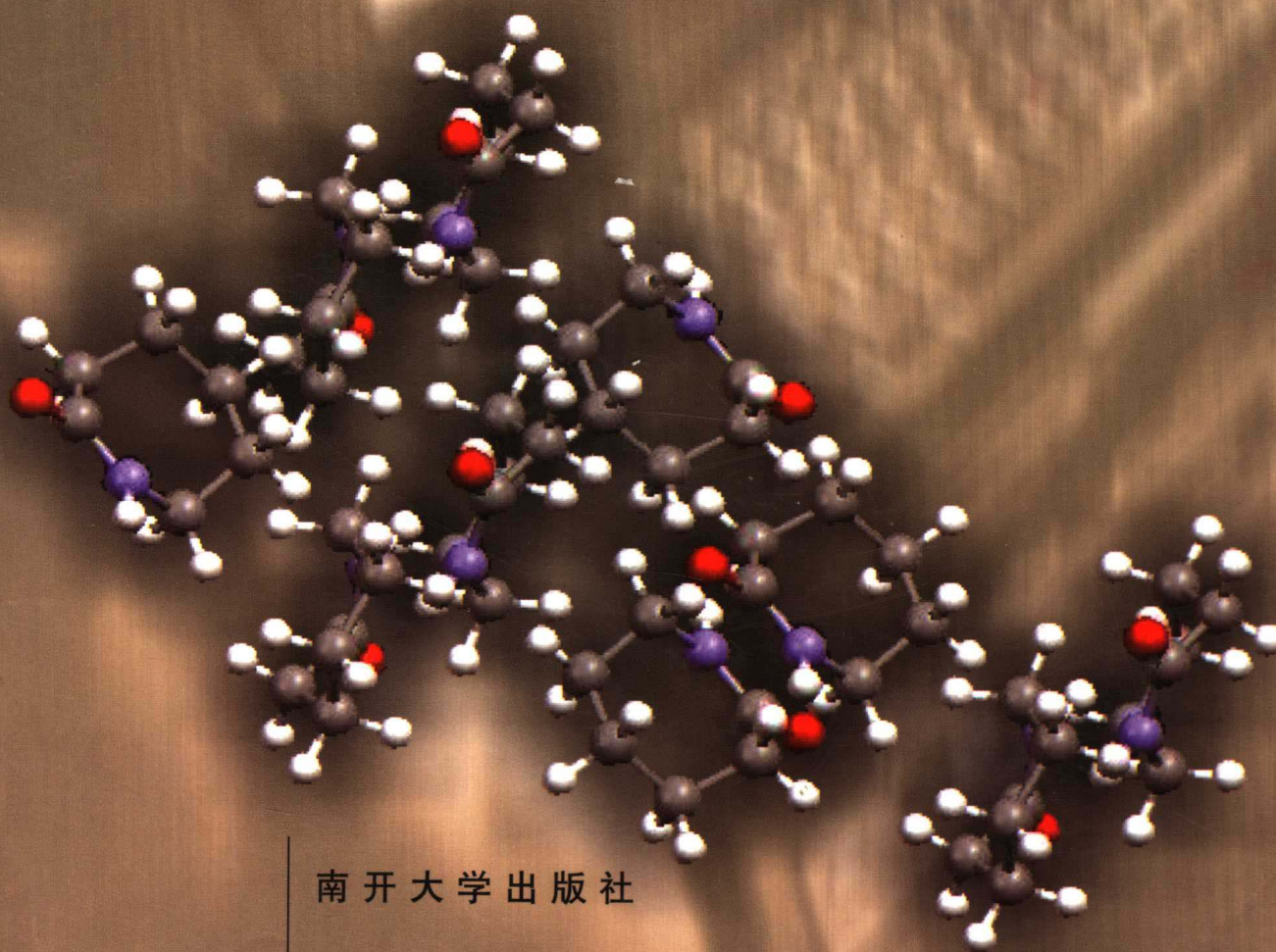
国家基础化学人才培养基地
南开大学近代化学教材丛书 申泮文院士 主编

简明

Jianming jisuanji huaxue jiaocheng

计算机化学 教程

乔园园 张明涛 主编



南开大学出版社

国家基础化学人才培养基地
南开大学近代化学教材丛书

申泮文院士 主编

简明计算机化学教程

主编 乔园园 张明涛

南开大学出版社
天 津

内容提要

计算机化学 (Computer chemistry) 是应用计算机技术来进行化学研究的科学。本书围绕着“化学信息”这一核心,从计算机化学的基本概念、方法入手,将实验数据、量测信号、谱图、分子结构、化学知识和 Internet 这一系列看似零散但却有着内在联系的内容串联起来,展现了计算机化学在数据处理、信号分析、结构解析、有机合成、分子设计等诸多方面的应用。

本书可供化学、化工、生物等相关专业人员和高等院校相关专业师生参考。

图书在版编目(CIP)数据

简明计算机化学教程 / 乔园园, 张明涛主编. — 天津: 南开大学出版社, 2005. 8

(国家基础化学人才培养基地南开大学近代化学教材丛书 / 申泮文主编)

ISBN 7-310-02327-7

I. 简... II. ①乔... ②张... III. 计算机应用—化学—高等学校—教材 IV. 06-39

中国版本图书馆 CIP 数据核字(2005)第 039180 号

版权所有 侵权必究

南开大学出版社出版发行

出版人: 肖占鹏

地址: 天津市南开区卫津路 94 号 邮政编码: 300071

营销部电话: (022)23508339 23500755

营销部传真: (022)23508542 邮购部电话: (022)23502200

*

河北省迁安万隆印刷有限责任公司印刷

全国各地新华书店经销

*

2005 年 8 月第 1 版 2005 年 8 月第 1 次印刷

787×1092 毫米 16 开本 12.75 印张 315 千字

定价: 22.00 元

如遇图书印装质量问题, 请与本社营销部联系调换, 电话: (022)23507125

化学要為中國的

經濟繁榮

學術進展

做出更大的貢獻

楊石先

一九八二年

南开大学近代化学教材丛书编委会

主 编 申泮文院士

副主编 王积涛教授

编 委 何锡文教授 朱志昂教授

袁满雪教授 程 鹏教授

刘靖疆教授 林少凡教授

张邦华教授 左育民教授

高如瑜教授 车云霞教授

邱晓航博士 (以上南开大学)

赵 康教授 (天津大学)

通信编委 姚守拙院士 (湖南大学) 方肇伦院士 (东北大学)

宋俊峰教授 (西北大学) 徐辉碧教授 (华中科技大学)

庞代文教授 (武汉大学) 王延吉教授 (河北工业大学)

李玉生高工 (航空航天工业部)

王华正高工 (航空航天工业部)

陈 亮教授 (山西大学)

《南开大学近代化学教材丛书》序

自 1997 年,南开大学化学学院开始进行面向 21 世纪的化学教育改革试点。我们参考了国内外高校的先进化学教育方案,设计了一套创新的教学计划和课程体系。优化后的化学课程设置体例如下:

第一类课程 必修基础课(本科一年级至三年级)	
化学概论 物理化学(含结构化学) 无机化学 有机化学 近代分析科学 实验化学——基础实验化学 中级实验化学 综合实验化学	
第二类课程 副修课(三年级下学期开设,每位学生人选 3 门)	
高等有机化学 高分子科学 量子化学与应用 计算机化学 化学生物学与生物技术 近代化学工程	
第三类课程 专业选修课及任选课(四年级)	
分支学科专业指定选修课程 学士毕业论文研究 特别任选课——绿色化学 纳米化学 组合化学 药物化学 材料化学 等	

大学本科一年级第一门启蒙课程“化学概论”,即国际高校通用的“General Chemistry”(过去此课程名称错译为“普通化学”,在当前改革之际,应加以纠正。经教育部高教司的同意,现正式定名为“化学概论”)。这门课程的教学目标是:以概论的形式向学生讲授化学学科的科学属性,它在科学体系中的地位及与其他相关学科的关系,它在人类社会中对人类生活与生产的作用与意义;化学学科的发展历史和它在当代的发展形势,特别是它的分支学科与边缘交叉学科在进入 21 世纪后的发展趋势,它对支持人类社会可持续发展中的重要作用;化学学科的教学计划和培养目标,对学生的要求等。本课程是一门学科概貌的引论课,是高中化学与大学化学沟通的桥梁课,既是通才教育课,又是素质教育课,集多重任务于一身。采用的主教材是申泮文主编的《近代化学导论》(高等教育出版社,2002)。

在化学概论课之后,继之以物理化学与结构化学大课,在物理化学与结构化学原理的指导下,后面并列开设无机化学、有机化学、近代分析科学三门化学主干课。化学实验课分年度独立设课。以上安排构成了基础必修课程体系(第一类课程)。这种课程设置模式,体现了 21 世纪学科发展特点的多学科知识交叉与渗透。化学学科的继往开来,适应了当前社会经济建设发展趋势,提高了学生的理论水平,拓宽了学生的知识面。

化学学科课程体系独有的特色是设立了副修课(第二类课程),目前暂设 6 门课:高等有机化学、高分子科学、量子化学与应用、计算机化学、化学生物学与生物技术和近代化学工程。学生在读完了必修基础课程之后,可以在自己感觉有兴趣的化学领域自主选择(或在教

师的指导下)选修其中的3门课程,构成一个未来发展方向的知识准备。例如,某学生在基础课学习中,对有机化学结构理论和反应机理产生了浓厚的兴趣,准备将来在这方面做深入研究,他可以在副修课程中选读高等有机化学、量子化学和计算机化学3门课,作为未来深入研究的准备。又如另一个学生对生命科学和生物技术感兴趣,他可以选读化学生物学、高分子科学和近代化学工程3门课,作为发展的准备。把选课自主权给予学生自己掌握,为的是培养学生的自主选择、自主发展意识,发挥他们自主的创新才能。副修课程准备常年开设,学生如果愿意增选其中1门或2门课程,可以在四年级有余力的时候加选。

到四年级,学生进入学科分支专业,按专业需要选修专业选修课和毕业论文。另外,鼓励全院教师发挥积极性,百花齐放、百家争鸣,开设当今前沿知识课,繁荣化学教坛。

千里之行,始于足下。欲使教学改革方案得到具体实施,必须把它落实到新课程的教材上。没有实体教材,任何教学改革只能是一句空话。为此,我们为教学改革投入了最大的力量,组织校内外专家学者成立了“南开大学近代化学教材丛书编辑委员会”,群策群力,按照教学改革方针和课程设计,编撰一整套新教材,迎接教育改革新时代的到来。

新教材的编撰原则是:(1)以百年诺贝尔自然科学奖为背景,展示化学未来发展趋势;(2)收列文献新颖度水平达到今日国际前沿;(3)重视我国科学家的科学成就;(4)重视培养学生的自主学习能力和发展创新思维。

编撰系统的整套教学改革教材的举措,在高等学校也是一项创举,感谢南开大学校院领导、天津市教育委员会、教育部高教司等领导单位给予的大力支持,也感谢高等教育出版社、科学出版社、化学化工出版社和南开大学出版社给予的慷慨帮助,分担了全套教材的出版任务。我们也对参加本套丛书编撰工作的专家学者(通信编委和顾问)表示衷心的感谢,没有他们的参与,本套丛书是难以顺利完成的。

本编委会将继续努力,编撰高年级选修课和研究生课程的教材,希望兄弟院校的专家学者们,将你们的化学专著加入到我们的教材系列中来,我们一定会全心全意地做好服务工作。

本套教学改革课程方案和教材丛书适合于大学本科化学各类专业使用,希望得到广泛的批评和指正,谢谢!

南开大学近代化学教材丛书编委会

2003年10月

前 言

计算机化学(Computer chemistry)的兴起与发展是与化学知识创新的迫切需要紧密联系的。从 20 世纪 60 年代起,随着计算机在工业发达国家的迅速普及,使得许多原来可望而不可及的科学计算任务得以实现。这些成就为人类探索发现新知识打开了新的大门。化学是一门需要化学家经验和灵感的科学,而计算机的高速度、高可靠性等优势,使化学家可以方便地进行数据处理、实验模拟等,因此计算机化学应运而生,并在 20 世纪七、八十年代得到了较快发展,从而成为一门独立的学科,从早期单一的数值分析、文献查询等到现在综合了信息学、统计学、数据库、图论、人工智能、控制论等多学科的理论与实践经验,有机地与化学相结合,形成了一个多学科交叉的具有特定内涵而又不断发展的新兴学科。

近年来,计算机化学在许多方面都取得了很大的进步:

- 数据处理 如数据采集、数值分析、化学数据库等;
- 谱图解析 如谱图模拟、结构解析等;
- 体系计算 如量子化学从头计算和半经验计算法等;
- 分子模拟 如分子力学、分子动力学计算、构象分析等;
- 合成设计 如合成路线选择、化学反应知识的发现、推理和预测等;
- 分子设计 如药物和农药、新材料的构效关系研究与分子设计等;
- 化学教学 如多媒体教学、智能题库等。

计算机化学是应用计算机技术来进行化学研究的科学,在有机化学、无机化学、分析化学、物理化学等方面都有着广泛的应用。总的说来,计算机化学所面临的问题有两类:一类是对化学信息的采集、处理等,另一类是对化学体系的模拟、化学知识的推理等。据目前的总体情况看,计算机化学在处理前一类“信息型”问题上有着比较长的历史和相对成熟的方法,也积累了很多成功经验。例如:仪器分析数据的采集与处理、大量化学文献的存储与查询等。而在处理后一类“智能型”问题上尽管已经取得良好的进展,仍面临着巨大的挑战。例如:各类谱图的智能解析、有机合成路线设计、药物分子设计等。造成这种现状的原因是多方面的,其中最主要的困难就在于如何把化学家的知识与经验以及他们对化学体系的描述转化为适合用计算机来解决的形式。正因为如此,计算机化学才是一门处在不断发展中的学科,在可以预见的未来,还会有更多的创新领域。

现在,随着功能完备、性能强大的化学软件的不断涌现和计算机与网络技术的高速发展,越来越多的化学家通过计算机进行化学研究。纵观历史,化学研究早已从“实验、再实验”的“炼金术”演变到了“设计—实验—改进设计—再实验”的科学,现在又在逐步向“设计—计算—改进设计—改进计算—实验”的新模式发展。正如因对量子化学计算方法的贡献而于 1998 年获 Nobel 化学奖的 J. Pople 教授提出的模型化学(Model chemistry)所预言的那样,融合计算机技术、数学、化学及其相关学科的最新理论成就的计算机化学,通过对化学概念

的新表述、运用数字化工具对化学体系进行新观察，从而发现化学新知识，将是未来化学研究的新模式。

化学信息按照其特点大致可以分为四个类型：

- 数字信息 科学测量和计算的结果，如原子和分子的物性数据和谱图数据、化学反应的速率和产率；
- 化学文献 化学研究、化学实践的文字记录及其他媒体记录等；
- 结构信息 化学结构的分布、联接情况，如分子的立体结构与构象、晶体的空间结构等；
- 化学知识 有关化学研究的经验、感知与推理。

其中，数字信息和化学文献都可以比较直接地用计算机处理，而结构信息和化学知识则需要经过一定的转化步骤才可以用计算机处理。无论是进行化学结构数据库的查询，还是进行各类谱图的解析；无论是研究分子结构特征与其性能表现的关系，还是探索合成路线、进行分子设计，都需要适当的信息转化过程。当然，针对不同类型的信息，转化方式也是多种多样的。因此，针对不同类型的化学信息，处理的方式与过程也是不同的。

本书围绕着“化学信息”这一核心，从计算机化学的基本概念、方法入手，将实验数据、量测信号、谱图、分子结构、化学知识和 Internet 这一系列看似零散但却有着内在联系的内容串联起来，展现了计算机化学在数据处理、信号分析、结构解析、有机合成、分子设计等诸多方面的应用。全书大致分为三个部分：第一至第四章是数据统计分析基础；第五章到第九章是非数值计算及其应用；第十章介绍 Internet 上的化学资源。

全书由乔园园和张明涛共同编著。在此，编者感谢南开大学申泮文院士、林少凡教授、赖城明教授等多年来的教育与指导，感谢天津大学马沛生教授对本书提出的宝贵意见，感谢各届博士、硕士研究生和本科生的贡献。本书参考了国内外许多相关的教科书和期刊，还有 Internet 上的大量信息，在此对这些参考资料的作者一并致谢。

编 者

2005 年 2 月

目 录

第一章 实验数据的统计分布	1
第一节 误差	1
第二节 统计分布	1
第三节 置信区间	3
第四节 统计检验	3
第二章 实验数据的回归分析	7
第一节 一元线性回归分析	7
第二节 多元线性回归分析	10
第三章 化学信号的数据处理	16
第一节 有用信号与噪声信号	16
第二节 滤波	17
第三节 平滑	18
第四节 微分	24
第五节 拟合	26
第六节 积分	29
第四章 实验数据的主成分分析	32
第一节 数据的预处理	32
第二节 随机向量的数字特征	34
第三节 矩阵特征值、特征向量及其 Jacobi 解法	36
第四节 主成分分析	40
第五章 分子拓扑结构	49
第一节 分子拓扑结构的表示	50
第二节 分子通式结构	60
第三节 分子结构片段	64
第六章 分子结构参数	72
第一节 物理化学参数	73
第二节 拓扑指数	75
第三节 量子化学参数	81
第四节 构效关系模型	83

第七章 计算机辅助结构解析	87
第一节 谱图数据库.....	87
第二节 计算机辅助结构解析.....	99
第八章 计算机辅助有机合成设计	106
第一节 有机合成反应的表示方法.....	107
第二节 计算机辅助有机合成设计系统.....	112
第九章 计算机辅助分子设计	119
第一节 药物作用的基本理论.....	119
第二节 药物分子设计的基本过程.....	121
第三节 药物分子设计的有关技术.....	123
第四节 药物分子设计的基本方法.....	129
第十章 化学信息与 Internet	138
第一节 Internet 服务.....	138
第二节 搜索引擎.....	141
第三节 网上化学信息.....	145
参考文献	179
附录	181
索引	185

第一章 实验数据的统计分布

对大量化学实验数据进行统计分布研究,是获取这些数据所蕴含的丰富信息必需的基础。

第一节 误差

误差(Error)是实验数据的基本特点。如果对系统中的某个量进行多次量测,每次的结果不会完全相同,即存在误差。误差按照性质可以分为随机误差(Random error)、系统误差(Systematic error)和过失误差(Accidental error)。

随机误差是由于在测定过程中一系列因素微小的随机波动而形成的具有相互抵偿性的误差。在一次测定中,随机误差的大小及其符号是无法预言的,没有任何规律性,但在多次测定中,随机误差的出现还是有规律的,它具有统计规律性。由于随机误差有大有小,可正可负,随着测定次数的增加,正、负误差相互抵偿,误差平均值趋向于零。因此,多次测定平均值的随机误差比单次测定值的随机误差小。随机误差的形成取决于测定过程中一系列无法避免的随机因素,可以设法将它大大减小。随机误差的影响一般可以采用统计方法来表征。

系统误差是指在一定试验条件下由某个或某些因素按照某一确定的规律起作用而形成的误差。系统误差的大小及其符号在同一试验中是恒定的,或在试验条件改变时按照某一确定的规律变化,重复测定不能发现和减小系统误差。只有改变试验条件才能发现系统误差。一旦发现了系统误差产生的原因,是可以设法避免和校正的。系统误差的影响往往可以采用标准物质来校正。

过失误差是指一种显然与事实不符的误差,没有一定的规律。不管造成过失误差的具体原因如何,只要确知存在过失误差,就应将含有过失误差的测定值作为异常值从量测数据中舍弃。

第二节 统计分布

化学量测可以看作是随机的、相互独立的事件,因此化学量测数据符合一定的统计规律,能够运用统计分布(Statistical distribution)等方法对量测数据进行描述。

随机数据的分布一般用它在一定区域内出现的频率来表示,那么频率随数据进行变化的函数称为分布函数(Distribution function)。对于理想的无穷多次随机量测来说,随机数据的分布应该符合正态分布(Normal distribution),相应的正态分布曲线可以用如下函数表示:

$$f(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x-\mu)^2}{2\sigma^2}}$$

其中:

- $f(x)$ 是概率密度函数 (Probability density function);
- x 是量测值 (Measured value);
- μ 是量测值所对应的真值 (True value): $\mu = \int_{-\infty}^{+\infty} xf(x)dx$;
- σ 是标准偏差 (Standard deviation): $\sigma^2 = \int_{-\infty}^{+\infty} (x-\mu)^2 f(x)dx$ 。

绝大多数情况下的化学量测都符合正态分布。正态分布具有很多特点, 例如:

- 曲线 $f(x)$ 由真值 μ 和标准偏差 σ 两个参数决定;
- 曲线关于真值 μ 对称, 极大值出现在 μ 处, 在 $\mu \pm \sigma$ 处各有一个拐点;
- 真值 μ 的大小不影响曲线形状;
- 标准偏差 σ 越大, 数据的离散程度越大;
- 量测数据 x 出现在确定范围内的几率是有规律的: 在 $\mu \pm \sigma$ 范围内约为 68.3%; 在 $\mu \pm 2\sigma$ 范围内约为 95.5%; 在 $\mu \pm 3\sigma$ 范围内约为 99.7%。

如果令正态分布的真值 μ 为零、标准偏差 σ 为 1, 定义 $z = \frac{x-\mu}{\sigma}$, 这样得到的正态分布称为标准正态分布, 相应的标准正态分布曲线可以表示为:

$$f(z) = \frac{1}{\sqrt{2\pi}} e^{-\frac{z^2}{2}}$$

当 z 取不同的数值时, 标准正态分布曲线下的面积就表示量测值出现的概率, 可以通过查表获得。

这种理想的正态分布是针对无限次量测事件的, 但是在量测次数比较多 (大于 100 次) 的情况下, 量测数据的分布一般仍然可以采用正态分布来描述。对于 N ($N > 100$) 次量测, 其相应的正态分布参数可以按照如下的公式进行计算:

$$\mu = \frac{1}{N} \sum_{i=1}^N x_i$$

$$\sigma^2 = \frac{1}{N} \sum_{i=1}^N (x_i - \mu)^2$$

然而, 在实际应用中由于量测次数是有限的, 因此不可能获取真值 μ 和标准偏差 σ , 通常是采用如下的估计值:

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$$

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

其中:

- n 是重复量测的次数;
- x_i 是第 i 次量测的数值;

- $\bar{x}$ 是 n 次重复量测的平均值, 相当于真值的估计值;

- S 是标准偏差 σ 的估计值, S^2 也称为方差。

有时也采取其他一些参数来表述量测数据的统计分布情况, 例如:

- 变异系数: $cv = S / \bar{x}$;

- 相对标准偏差: $RSD(\%) = 100S / \bar{x}$;

- 均值标准偏差: $S_{\bar{x}} = S / \sqrt{n}$ 。

第三节 置信区间

有多大的把握相信真值 μ 就处于多次量测值的范围之内? 为了描述这种情况, 根据正态分布的性质可以把以真值 μ 为中心的某个范围作为置信区间 (Confidence interval), 这个置信区间范围的大小是由置信水平 (Confidence degree) 来决定的, 置信水平越高, 相应的置信区间就越大。有时也用显著性水平 (Significance level), 即 1 减去置信水平来表示。正如前面正态分布特点所显示的那样, n 次量测平均值 $\bar{x}$ 在置信水平为 95% (显著性水平为 5%) 时的可能取值范围是:

$$\mu - 1.96 \left(\frac{\sigma}{\sqrt{n}} \right) < \bar{x} < \mu + 1.96 \left(\frac{\sigma}{\sqrt{n}} \right)$$

可以改写为如下形式, 即得置信水平为 95% 时量测均值的置信区间:

$$\bar{x} - 1.96 \left(\frac{\sigma}{\sqrt{n}} \right) < \mu < \bar{x} + 1.96 \left(\frac{\sigma}{\sqrt{n}} \right)$$

常用的是置信水平为 99% (显著性水平为 1%) 时的置信区间:

$$\bar{x} - 2.58 \left(\frac{\sigma}{\sqrt{n}} \right) < \mu < \bar{x} + 2.58 \left(\frac{\sigma}{\sqrt{n}} \right)$$

因为 σ 是未知的, 往往在量测次数 n 足够大 (如 $n > 100$) 时, 用 S 来代替; 然而当 n 不够大的时候, 用 S 会引入比较大的误差, 此时往往采用如下方式计算:

$$\bar{x} - tS_{\bar{x}} < \mu < \bar{x} + tS_{\bar{x}}$$

其中, t 值来自于 t 分布, 是由置信区间和自由度 $n-1$ 所决定的。 t 分布是一种用于统计检验的重要统计分布。

第四节 统计检验

当获取了量测数据以后, 需要对这些数据进行检验, 看其是否符合随机分布规律。统计检验 (Statistical inference) 就是在一定的置信水平 (或显著性水平) 上, 检验量测数据是否符合某种假设的分布规律。一般不能证明假设是绝对正确或绝对错误, 所能做的就是当数据支持该假设时接受它, 当数据不支持该假设时舍弃它。

一般这种假设检验分为如下的步骤进行:

- 提出零假设 H_0 (Null hypothesis, H_0) 和备择假设 H_1 (Alternative hypothesis, H_1),

从而使检验结论有意义;

- 选取适当的显著性水平 (取 5% 或 1%, 相当于置信水平为 95% 或 99%);
- 在假设成立的条件下选择适当的统计量并确定其分布;
- 根据统计量的分布和显著性水平提出判断规则;
- 计算样本的统计量值;
- 根据统计量的分布查表得到对应显著性水平其下临界值, 然后根据判断规则做出判断。

零假设 H_0 的含义是量测数据与真值一致 (并非没有随机误差), 而备择假设 H_1 的含义就是量测数据与真值不一致。如果 H_0 被接受, 那么备择假设 H_1 就被舍弃; 否则, 就只有舍弃 H_0 而接受备择假设 H_1 。一般是当量测数据与真值差异的概率小于所选取的显著性水平时, 舍弃 H_0 而取 H_1 。

进行这样的检验不论结论如何, 总是不可避免地存在发生错误的风险。在实际操作中用的是“反证法”: 虚设一个零假设进行检验, 目的是看当这个假设成立时, 检验结论是否合理。若合理就接受这个假设; 否则就舍弃这个假设。应该注意的是, 这里的合理与否不是绝对的, 而是根据科学实践中广泛采用的“概率接近零的事件实际上是不可能发生的”这一原则来考虑的。

一、 t 检验

在进行量测数据均值与真值的比较和量测数据均值之间的比较时, 一般采用 t 检验。

1. 量测数据均值与真值的比较: 单边 t 检验 (One-sided t test)

单边 t 检验的主要方法如下:

- 提出零假设 H_0 为 $\bar{x} \leq \mu$, 则备择假设 H_1 为 $\bar{x} > \mu$;
- 选取显著性水平 α 为 5% (或 1%);
- 计算统计量 $t = \frac{|\bar{x} - \mu|}{S} \sqrt{n}$, 其中 S 是标准偏差 (估计值), n 是量测次数;
- 通过查表获得 $t(\alpha, f)$, 其中 $f = n - 1$;
- 如果 $t > t(\alpha, f)$, 拒绝 H_0 而接受 H_1 ; 否则接受 H_0 而拒绝 H_1 。

例 4 次测定同一饮用水样品中硝酸根离子的浓度 (ppm), 分别得到: 51.0、51.3、51.6、50.9。若法定值为 50.0 ppm, 问该样品的硝酸根离子浓度是否合格。

按照单边 t 检验的方法进行处理。

$$\mu = 50.0;$$

设 H_0 为 $\bar{x} \leq \mu$, H_1 为 $\bar{x} > \mu$;

$$\bar{x} = 51.2, S = 0.316, n = 4, t = 7.59;$$

取 α 为 5%, 查表, $t(0.05, 3) = 2.353$;

$$t > t(0.05, 3), \text{接受 } H_1。$$

结论: 样品的硝酸根离子浓度不合格。

2. 量测数据均值之间的比较: 双边 t 检验 (Two-sided t test)

设有两组量测数据, 均值分别是 $\bar{x}_1$ 和 $\bar{x}_2$, 方差分别是 S_1^2 和 S_2^2 。双边 t 检验的主要方法如下:

- 提出零假设 H_0 为两组数据之间的差异是随机的, 有统计学意义; 备择假设 H_1 为两

组数据之间的差异不是随机的，没有统计学意义；

- 选取显著性水平 α 为 5%（或 1%）；

- 计算统计量 $t = \frac{|\bar{x}_1 - \bar{x}_2|}{S_d \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}}$ ，其中 n_1 和 n_2 分别是两均值的量测次数，而

$$S_d = \sqrt{\frac{(n_1 - 1)S_1^2 + (n_2 - 1)S_2^2}{n_1 + n_2 - 2}}$$
 是加权标准偏差；

- 通过查表获得 $t(1 - \alpha/2, f)$ ，其中 $f = n_1 + n_2 - 2$ ；
- 如果 $t > t(1 - \alpha/2, f)$ ，拒绝 H_0 而接受 H_1 ；否则接受 H_0 而拒绝 H_1 。

例 通过在盐酸介质中进行蒸馏来测定食物中锡的含量（ppm），在 30 分钟和 75 分钟的蒸馏时间下得到的结果分别是：55.0、57.0、59.0、56.0、56.0、59.0 和 57.0、55.0、58.0、59.0、59.0、59.0。问蒸馏时间对于测定结果有无统计学影响？

按照双边 t 检验的方法进行处理。

H_0 为 $\mu_1 = \mu_2$ ，即蒸馏时间无影响； H_1 为 $\mu_1 \neq \mu_2$ ；

$$n_1 = 6, n_2 = 6, f = 10;$$

$$\bar{x}_1 = 57.0, \bar{x}_2 = 57.8;$$

$$S_1^2 = 2.80, S_2^2 = 2.57;$$

$$S_d = 1.64;$$

$$t = 0.84;$$

取 $\alpha = 0.05$ ，查表， $t(1 - 0.05/2, 10) = 2.228$ ；

$$t < t(1 - 0.05/2, 10), \text{ 接受 } H_0。$$

结论：蒸馏时间对于测定结果无统计学影响。

如果认为两组量测数据的方差 S_1^2 和 S_2^2 之间的差异不是随机的，不具有统计学意义，则可用如下方法计算 t 和 f 值：

$$t = \frac{|\bar{x}_1 - \bar{x}_2|}{\sqrt{\frac{S_1^2}{n_1} + \frac{S_2^2}{n_2}}}$$

$$f = \frac{\left(\frac{S_1^2}{n_1} + \frac{S_2^2}{n_2}\right)^2}{\frac{(S_1^2/n_1)^2}{n_1 - 1} + \frac{(S_2^2/n_2)^2}{n_2 - 1}}$$

例 风湿病患者与正常人对照者的血内硫醇含量（mmol）分别是 2.81、4.06、3.62、3.27、3.27、3.76 和 1.84、1.92、1.94、1.92、1.85、1.91、2.07。问两者的血内硫醇含量是否有统计学意义？

按照双边 t 检验的方法进行处理。

设 H_0 为 $\mu_1 = \mu_2$ ； H_1 为 $\mu_1 \neq \mu_2$ ；

$$n_1 = 7, \bar{x}_1 = 1.921, S_1^2 = 0.005776;$$

$$n_2 = 6, \bar{x}_2 = 3.465; S_2^2 = 0.1936;$$

$$f = 5;$$

$$t = 8.5;$$

取 $\alpha=0.01$, 查表, $t(1-0.01/2, 5) = 4.032$;

$$t > t(1-0.01/2, 10), \text{ 接受 } H_1。$$

结论: 两者的血内硫醇含量之间的差异不是随机的, 不具有统计学意义。

二、 F 检验

设有两组量测数据, 方差分别是 S_1^2 和 S_2^2 ($S_1^2 \geq S_2^2$)。为了比较其方差, 一般采用 F 检验, 步骤如下:

- 提出零假设 H_0 和备择假设 H_1 , 前者为 S_1^2 和 S_2^2 之间无显著性差异, 有统计学意义, 后者为 S_1^2 和 S_2^2 之间有显著性差异, 无统计学意义;
- 选取显著性水平 α 为 5% (或 1%);
- 计算统计量 $F = \frac{S_1^2}{S_2^2}$;
- 如果 $F > F(1 - \alpha/2; f_1, f_2)$, 拒绝 H_0 而接受 H_1 ; 否则接受 H_0 而拒绝 H_1 。

例 两个实验室各自采用原子吸收光谱法测定钢中的钛含量。其中一个实验室的结果是: 0.470、0.448、0.463、0.449、0.482、0.452、0.477、0.409; 另一个实验室的结果是: 0.529、0.490、0.489、0.521、0.486、0.502。问两者的量测结果是否有统计学意义?

首先进行 F 检验, 看其方差是否有统计学意义。

设 H_0 为 S_1^2 和 S_2^2 有统计学意义; H_1 为 S_1^2 和 S_2^2 无统计学意义;

$$S_1^2 = 0.000\ 522\ 4;$$

$$S_2^2 = 0.000\ 331\ 2;$$

$$F = 1.58;$$

取 $\alpha=0.05$, 查表, $F(1 - 0.05/2; 7, 5) = 6.85$;

$$F > F(1 - 0.05/2; 7, 5), \text{ 接受 } H_0。$$

然后进行双边 t 检验, 看两组数据之间的差异是否有统计学意义。

设 H_0 为 $\mu_1 = \mu_2$, 即在两个实验室中进行实验对结果无影响; H_1 为 $\mu_1 \neq \mu_2$;

$$\bar{x}_1 = 0.467, \quad \bar{x}_2 = 0.503; \quad S_d = 0.022\ 1;$$

$$t = 4.07;$$

取 $\alpha=0.05$, 查表, $t(1 - 0.05/2, 12) = 2.179$;

$$t < t(1 - 0.05/2, 12), \text{ 接受 } H_0。$$

结论: 两者的方差具有统计学意义; 在两个实验室所进行的量测对钢中钛的含量无统计学影响。