

全国科学技术名词审定委员会
征求意见稿



卫生统计学名词
HEALTH STATISTICS

2025
全国公共卫生与预防医学名词审定委员会
卫生统计学名词审定分委员会

公开征求意见期限：
2025年7月8日——2025年10月8日

内 容 简 介

本次公开征求意见的是第一版卫生统计学名词，内容包括：学科核心概念、统计学原理、统计学方法、研究设计和卫生统计指标等 5 部分，共 1816 条。每条词都提供了定义或注释。



公开征求意见期限：
2025年7月8日——2025年10月8日

全国公共卫生与预防医学名词审定委员会委员名单

主任：高 福

常务副主任：刘剑君

副主任：李立明 梁晓峰 唐金陵

委员（以姓名笔画为序）：

么鸿雁 王 辰 冯子健 任 涛 刘起勇 刘雅文
孙全富 孙 新 邬堂春 陈君石 何 纳 沈洪兵
吴 凡 吴息凤 张玉森 张 本 金 曦 林东昕
罗会明 周晓农 郝元涛 胡国清 施小明 赵文华
顾东风 郭中平 夏彦恺 徐建国 曹务春

秘书长：张玉森

副秘书长：罗会明 任 涛

秘 书：元 晓 马 静 刘梦冉 郑文静

全国公共卫生与预防医学名词编写委员会委员名单

总主编：刘剑君

委员（以姓名笔画为序）：

丁钢强 马 军 么鸿雁 刘起勇 吕 军 孙全富
孙 新 孙殿军 李 涛 吴永宁 张流波 邹 飞
孟庆跃 周晓农 郝元涛 胡国清 施小明 郭 岩
钱 序 夏彦恺

秘书长：么鸿雁

副秘书长：元 晓

秘 书：马 静 刘梦冉 王琦琦 董 哲

公开征求意见期限：

2025年7月8日——2025年10月8日

卫生统计学名词审定分委员会委员名单

顾问：方积乾 孙振球 徐勇勇 颜虹

主任：郝元涛

副主任：王彤 陈峰 夏结来

委员（以姓氏笔画为序）：

马骏 万崇华 方亚尹 平 刘美娜 刘玉秀 许传志 宇传华

杨土保 张菊英 易东 周晓华 贺佳 胡国清 施学忠 赵耐青

钟晓妮 郭海强 郭秀花 贾红凌 莉 薛付忠

秘书：李菁华 张王剑

卫生统计学名词编写分委员会委员名单

主编：郝元涛

副主编：于石成 朱彩蓉 赵杨

委员（以姓氏笔画为序）：

邓丹 王超龙 王陵 刘海燕 刘红波 伍亚舟 陈雯 吴骋

吴思英 言方荣 杨永利 余红梅 张晋昕 张涛 欧春泉 尚磊

郝春 侯艳 胡跃华 赵星 顾菁 秦国友 彭志行 曾令霞

黎燕宁 魏永越

秘书：王莹 赖颖斯 廖婧

公开征求意见期限：

2025年7月8日——2025年10月8日

前言

卫生统计学是应用统计学的基本原理和方法，研究健康领域中数量现象的学科，通常包括资料的搜集、整理、分析和结果呈现等过程。近年来，随着人工智能以及新兴算法的快速进展，卫生统计学所涵盖的方法体系不断扩展，应用范围日益广泛。作为一门理论与应用并重的综合性学科，卫生统计学涉及医学、公共卫生、数据科学、统计学、信息学等多领域交叉，建立完善的术语规范与概念统一的学科名词体系，不仅为学科发展、学术传播和国际合作交流奠定了坚实基础，也有助于提升研究的规范性、促进成果传播，并为跨领域交流提供有力支撑。

2021年6月，卫生统计学名词编写委员会与审定委员会正式成立，编审专家来自全国20余所高校，并邀请多位国内权威学者担任顾问，协力推进学科名词体系的规范化建设。卫生统计学名词编审工作也随之正式启动。全体编审专家通过多轮专题研讨会、反复论证与集中修订，不断统一术语标准、细化表述规范，旨在确保学科术语体系的科学性、规范性与前瞻性，形成涵盖学科核心概念、统计学原理、统计学方法、研究设计和卫生统计指标五个一级条目的术语体系，共收录核心术语1816条。

谨向卫生统计学名词编审分委员会全体委员、顾问专家及积极参与编审工作的各位学者致以衷心感谢，感谢大家为名词体系建设所付出的辛勤努力与专业支持。对全国科技名词委的全程指导和中国疾病预防控制中心等协作单位的大力支持表示衷心感谢。在学科名词编审过程中，尽管已多次反复审校，由于学科发展迅速、当前水平仍有局限，部分疏漏错误恐难以完全避免。恳请各界同仁积极就名词分类框架、定名规范及释义准确性提出宝贵意见，助力我国卫生统计学名词体系不断优化与提升。

卫生统计学名词审定分委员会

2025年6月30日

编排说明

- 一、本书征求意见稿是热带医学名词，共 1816 条。
- 二、全书分 5 部分：学科核心概念、统计学原理、统计学方法、研究设计和卫生统计指标。
- 三、正文按汉文名所属学科的相关概念体系排列。汉文名后给出了与该词概念相对应的英文名。
- 四、一个汉文名对应几个英文同义词时，英文词之间用“,” 分开。
- 五、凡英文词的首字母大、小写均可时，一律小写；英文除必须用复数者，一般用单数形式。
- 六、“[]” 中的字为可省略的部分。
- 七、异名包括：“全称”“简称”是与正名等效使用的名词；“又称”为非推荐名，只在一定范围内使用；“俗称”为非学术用语；“曾称”为被淘汰的旧名。



目 录

前言
编排说明

正文

1	卫生统计学.....	5
1.1	基本概念	5
2	统计学原理.....	6
2.1	统计描述	6
2.2	统计推断	9
3	统计学方法.....	11
3.1	t 检验	11
3.2	方差分析	12
3.3	卡方检验	13
3.4	秩和检验	14
3.5	相关分析	15
3.6	对应分析	16
3.7	线性回归	16
3.8	广义线性模型	22
3.9	生存分析	24
3.10	判别分析	27
3.11	聚类分析	29
3.12	主成分分析	30
3.13	荟萃分析	31
3.14	时间序列分析	33
3.15	空间分析	35
3.16	结构方程模型	40
3.17	多水平模型	43
3.18	广义估计方程	45
3.19	遗传统计	46
3.20	贝叶斯统计	48
3.21	缺失数据	50
3.22	量表评价	52
3.23	因果推断	54
4	研究设计	57
4.1	实验研究设计	57
4.2	临床试验	60
4.3	调查设计	70
4.4	现实世界研究	74
5	卫生统计指标	75
5.1	居民健康状况指标	75
5.2	健康影响因素指标	77

5.3 疾病控制指标..... 78

5.4 妇幼保健指标..... 79

5.5 卫生监督指标..... 80

5.6 医疗服务指标..... 81

5.7 药品与卫生材料供应保障指标..... 84

5.8 卫生资源指标..... 86

5.9 医疗保障统计指标..... 90



1 卫生统计学

1 卫生统计学 health statistics

应用统计学的基本原理和方法,研究健康领域中数量

现象的学科。通常包括资料的搜集、整理、分析和研究等过程。

1.1 基本概念

1.1.1 总体 population

研究对象的全体,具有同质性和变异性的个体组成的集合。

1.1.1.1 有限总体 finite population

在所规定的时间、空间、人群范围内观察单位数量有限的总体。

1.1.1.2 无限总体 infinite population

研究对象的数量不受限制,或在时间、空间、人群范围上无限制,难以准确计数全部单位的总体,通常用于理论分析或假设场景。

1.1.2 个体 individual

又称“观察单位(observed unit)”。研究的基本单位,被观察、测量或研究的单个独立单位或对象。

1.1.3 抽样 sampling

从研究总体中抽取一定数量的个体的过程。

1.1.3.1 随机抽样 random sampling

从总体中抽取样本时,遵循随机原则,确保每个个体或单位具有相同的、已知的非零概率被抽中,包括完全随机抽样、系统抽样、整群抽样、分层抽样等。

1.1.4 误差 error

研究对象的实际测量值与真实值之间的差异。

1.1.4.1 系统误差 systematic error

由测量工具、方法或环境等因素引起的测量结果与真实值之差,具有一致性和方向性,影响数据的准确性,无法通过增加样本量消除。

1.1.4.2 随机误差 random error

由不可控因素引起的测量结果与真实值之差,呈现无规律波动,具有随机性和不可预测性,可通过增加样本量或重复测量减小其影响。

1.1.4.2.1 抽样误差 sampling error

由于个体之间存在变异导致的样本统计量与总体参数之间的差异,随样本量增大而减小,反映样本代表性的不确定性。

1.1.5 样本 sample

按一定规则从总体中抽取的一定数量的个体。

1.1.5.1 样本量 sample size

按一定规则从总体中抽取的样本所包含研究对象的数量。

1.1.6 同质 homogeneity

研究对象具有的相同特征或属性等的共性。

1.1.7 变异 variatio;variability

同一总体中不同个体间存在的差异。

1.1.8 变量 variable

反映个体特征或属性的量。

1.1.8.1 定量变量 quantitative variable

又称“数值变量”。具有数值特征,并可以被量化或者测量的变量。

1.1.8.1.1 连续变量 continuous variable

在某一区间可以取任何数值的定量变量。

1.1.8.1.2 离散变量 discrete variable

在某一区间具有有限或可数无限取值的变量,每个取值之间存在明确间隔,常用于描述计数数据或事件次数。

1.1.8.2 定性变量 qualitative variable

又称“分类变量”。描述类别或属性的变量,取值为不同的类别或标签,无法进行数量比较,常用于分类或标识,如性别、职业类型等。

1.1.8.2.1 有序变量 ordinal variable

又称“等级变量”。各取值之间存在顺序或等级关系的分类变量。

1.1.8.2.2 名义变量 nominal variable

各取值之间不存在顺序或等级关系的分类变量。

1.1.9 参数 parameter

描述总体统计特征的指标。

1.1.10 统计量 statistic

描述样本统计特征的指标。

1.1.11 概率分布 probability distribution

对随机变量所有可能取值或取值区间,以及相应取值或取值区间出现可能性大小的一种数学描述,以函数或列表形式呈现可能性规律。

1.1.11.1 分布特征

1.1.11.1.1 对称分布 symmetric distribution

随机变量的取值分布形态关于某一中心轴（或中心点）两边完全相同，即数据在中心轴（或中心点）两侧的频数或频率分布呈现对等、均衡的状态。

1.1.11.1.2 偏峰分布 skewed distribution

数据分布不对称，其形态偏离标准的对称分布，表现为数据集中在某一侧，而另一侧有更长的“拖尾”，正偏峰右尾长，负偏峰左尾长。

1.1.11.1.3 偏度系数 coefficient of skewness

用于衡量随机变量概率分布偏态方向和程度的数值，通过对分布的三阶矩进行标准化计算得出，值为正表示右偏，为负表示左偏，绝对值越大偏斜程度越高。

1.1.11.1.4 峰度系数 coefficient of kurtosis

衡量数据分布峰态陡缓程度的统计量，通过对随机变量概率分布的四阶矩进行标准化计算得到的数值。值越大峰越尖，值越小峰越平，相对于正态分布基准，正态分布的峰度系数为3。

1.1.11.2 分布类型

1.1.11.2.1 连续[型]分布 continuous distribution

描述连续随机变量取值范围内概率分布的函数，值域为实数区间。

1.1.11.2.1.1 正态分布 Normal distribution

又称“高斯分布(Gaussian distribution)”。一种对称的连续型概率分布，形状呈钟形，由均值和标准差完全确定。均值决定中心位置，标准差决定宽度。标准差越小，分布形状越窄，观察值越集中在均数附近。

1.1.11.2.1.2 t 分布 t distribution

全称“学生 t 分布 (Student's t distribution)”。概率密度曲线由自由度决定，呈现单峰，以0为中心的对称分布。形状类似正态分布但尾部较厚，随着样本量的增加，该分布逐渐趋近于正态分布。

1.1.11.2.1.3 z 分布 z distribution

总体均数为0，总体标准差为1的正态分布。

1.1.11.2.1.4 F 分布 F distribution

从一个正态总体中随机抽取样本量为 n_1 和 n_2 的两个样本，其方差的比率分布。

1.1.11.2.1.5 卡方分布 Chi-square distribution

又称“ X^2 分布”。 n 个相互独立且各自服从标准正态分布的随机变量的平方和的分布。

1.1.11.2.1.6 威布尔分布 Weibull distribution

其概率密度曲线由形状参数和尺度参数决定，形状参数决定了分布的偏度，尺度参数则影响曲线的伸缩和平移，常用于描述与时间相关的随机事件的分布，如随机事件发生的时间点或持续时间。

1.1.11.2.1.7 伽马分布 Gamma distribution

用于描述一系列独立同分布的随机事件发生所需时间的连续型概率密度分布，其中指数分布与卡方分布均为特例。

1.1.11.2.2 离散[型]分布 discrete distribution

描述离散随机变量取值的概率分布，随机变量仅能取有限个或可列无限个孤立值，以概率分布列描述其取值及对应概率，每个取值有确定概率，所有取值概率之和为1。典型例子包括二项分布和泊松分布。

1.1.11.2.2.1 二项分布 binomial distribution

对于 n 次独立、重复试验，如果每次试验只出现两种对立的结果（对立事件 A 与 B ），在每次试验中事件 A 发生的概率都是 π ，且各个观察对象的结果是相互独立的，出现事件 A 的次数 X 的概率分布。

1.1.11.2.2.2 泊松分布 Poisson distribution

用以描述单位时间、空间、面积等的罕见事件发生次数的概率分布。

1.1.11.2.2.3 负二项分布 negative binomial distribution

描述在一系列独立伯努利试验中，成功次数达到固定值之前失败次数的概率分布。概率密度函数包含两个参数：成功的次数和每次试验的成功概率。

2 统计学原理

2.1 统计描述

2.1 统计描述 statistical description

运用统计指标、统计表、统计图等方法，对数据的特征、分布规律、变量间关系等进行收集、整理和展示，以概括和呈现数据基本信息。

2.1.1 统计表 statistical table

将统计分析的指标及其数量结果用表格列出的形式。一般由标题、标目、线条、数字和脚注五部分组成，可以分为简单表和复合表。

2.1.1.1 简单表 simple table

纵标目只按一个变量分组的统计表。包括横标目和纵

标目。横标目位于表的左侧，用来说明各横行数字的含义；纵标目位于表的上方，只按一个变量分组，用来说明纵栏数字的含义。

2.1.1.2 复合表 combinative table

纵标目按两个或两个以上变量分组的统计表。包括横标目和纵标目。横标目位于表的左侧，用来说明各横行数字的含义；纵标目位于表的上方，用两个或两个以上变量联合说明纵栏数字的含义。

2.1.1.3 频率分布表 frequency distribution table

简称“频数表(frequency table)”。展示数据在各个类别或区间内出现次数的表格。列出每个类别的频数和频率，帮助直观理解数据分布情况。常用于分类数据和数值数据的初步分析，便于识别模式和异常值。

2.1.1.4 列联表 contingency table

一种将两个或多个分类变量的不同取值组合及其频数进行呈现的统计表，用于分析变量之间是否存在关联及关联程度。

2.1.1.4.1 2×2 列联表 2×2 contingency table

用于分析两个二分类变量关系的表格，由两行两列组成，四个格子分别为两种变量不同水平组合下的频数。

2.1.2 统计图 statistical chart

用点的位置、曲线的走势、直条的长短或面积的大小等直观地呈现所研究事物数量关系的图形。

2.1.2.1 直方图 histogram

横坐标表示定量变量的组段、纵坐标表示各组段变量值的频数、频率或频率密度（频率/组距）的统计图。用于描述定量变量的分布情况。

2.1.2.2 累积直方图 cumulative histogram

横坐标表示定量变量的组段、纵坐标表示各组段变量值的累计频率的统计图。用于描述定量变量的累计频率分布情况。

2.1.2.3 箱式图 box plot

又称“箱线图”、“盒须图”、“盒式图”。用于展示数据分布特征的统计图。由五个关键数值点构成：最小值、第一四分位数（ Q_1 ）、中位数（ Q_2 ）、第三四分位数（ Q_3 ）和最大值。箱体表示数据的中间 50%（ Q_1 至 Q_3 ），箱内横线为中位数，延伸的“须”显示数据范围，异常值通常单独标出。其优势在于直观呈现数据的离散程度、偏态及异常值，适用于多组数据对比分析。

2.1.2.4 误差条图 error bar chart

用于展示数据变异性的统计图，通常在柱状图或折线图的基础上添加误差条。误差条表示数据的离散程度，可反映标准差、标准误或置信区间等统计量。其长度代表数据的波动范围，有助于直观比较不同组间的差异显著性，能有效评估数据的可靠性和稳定性。

2.1.2.5 直条图 bar chart

又称“条形图”。用等宽长条表示分类数据的统计图。长条高度（或长度）反映数值大小，适用于比较不同类别间的离散数据。通常纵轴为数值，横轴为分类变量，可垂直或水平排列。

2.1.2.6 百分条图 percent bar chart

将总长度设为 100%，并按构成比例分段表示的统计图。各分段长度代表不同类别在整体中的占比，通常以不同颜色或纹理区分。适用于展示构成比或多组数据的相对比例关系，便于直观比较各部分在整体中的分布情况。

2.1.2.7 圆图 pie chart

又称“饼图”。用圆的总面积表示事物的全部，圆内各扇形面积表示各组成部分所占构成比的统计图。

2.1.2.8 线图 line chart

用线段的升降表示统计指标随时间的变化趋势，或某现象随另一现象的变化而变化情况的统计图。坐标轴一般均为算术尺度

2.1.2.9 半对数线图 semi-logarithmic line chart

横轴为算术尺度，纵轴为对数尺度的线图。所呈现数量关系为对数关系，通常用来反映事物的发展速度。

2.1.2.10 散点图 scatter plot

展示两个变量之间关系的图形，以一个变量为横坐标，另一变量为纵坐标，每个数据点代表一对变量值，利用数据点的分布形态来反映定量变量之间的关系的。

2.1.2.11 茎叶图 stem-leaf plot

数据被分离成整数部分和尾数部分后，由整数部分作为茎、尾数部分作为叶的统计图。每一行由一个整数的茎和若干叶构成。左边是茎的数值，茎的数量级标在图的下方；右边是叶，显示每个叶的尾数数值，在图的下方标示每叶代表几个实际观察值。用于直观地展示数据的分布范围和形态。

2.1.2.12 统计地图 statistical map

根据不同地方某种现象的数值大小，采用不同密度的线条或不同颜色绘制的地图。用于表示某种现象在地域空间上的分布。

2.1.2.13 热图 heat map

用不同的颜色（或者深浅）表示观测值大小的统计图。常用于表示疾病的时间与空间分布、基因表达谱等。

2.1.2.14 列线图 nomogram

将多个预测因素整合为一个可视化图形的工具，通常由多条带有刻度的线条组成，可用于直观地展示各因素对结果的影响程度，并通过简单计算预测特定事件发生的概率或风险。

2.1.3 百分位数 percentile

将一组数据从小到大排序后，把数据划分为 100 等份，

位于特定等分位置上的数值，反映数据在某一百分比位置的水平，可用于描述数据的分布特征，如第 25、50、75 个等分位置的数值分别为下四分位数、中位数、上四分位数。

2.1.3.1 四分位数 quartile

把一组数据从小到大排序后，平均分成四等份，处于三个分割点位置的数值。分别为第一四分位数（下四分位数），将数据的前 25% 与后 75% 分开；第二四分位数即中位数；第三四分位数（上四分位数），分开前 75% 和后 25% 的数据。

2.1.3.1.1 下四分位数 lower quartile

将所有观测值从小到大依次排序后，位置处于 25% 的数值。表示全部观测值中有四分之一的数值比它小。符号为 Q_1 或 P_{25} 。

2.1.3.1.2 上四分位数 upper quartile

将所有观测值从小到大依次排序后，位置处于 75% 的数值。表示全部观测值中有四分之一的数值比它大。符号为 Q_3 或 P_{75} 。

2.1.4 集中趋势 central tendency

又称“平均数（average）”。一组数据向某一中心值靠拢的趋势，反映了一组数据中心点的位置。可用于描述变量的平均水平。

2.1.4.1 算术[平]均数 arithmetic mean

简称“均数(mean)”。一组数值变量的观测值之和除以观测值个数所得的商。适用于服从对称分布变量的集中趋势描述。

2.1.4.2 调和均数 harmonic mean

又称“倒数平均数”、“调和平均(harmonic mean)”。一个变量所有观测值倒数的算术均数的倒数。常用于计算平均速率。

2.1.4.3 几何均数 geometric mean

所有观测值乘积开观测值个数的次方根。常用于描述呈明显正偏峰分布或观测值之间呈倍数关系数据的集中趋势。符号 G 。

2.1.4.4 截尾均值 trimmed mean

一个变量所有观测值去掉两端极端值后所计算的算术平均数。能够有效避免极端值对整体数据的影响。

2.1.4.5 中位数 median

一种描述数据集中趋势的统计量，将数据集按大小顺序排列后，位于中间位置的数。若数据量为偶数，则取中间两个数的平均值。对极端值不敏感，常用于表示偏态分布的中心位置。符号 M 。

2.1.4.6 众数 mode

一组数据中出现次数最多的数据值。用于描述变量的集中趋势。

2.1.5 离散趋势 dispersion tendency

描述一组数据中各数据值远离其中心值（如均值、中位数）的程度，反映数据的分散状况，常用极差、方差、标准差、四分位间距等指标衡量，可展示数据的分布范围和变异性。

2.1.5.1 极差 range

又称“全距”。所有观测值中最大值与最小值的差值。其越大说明数据变异程度越大，或数据越离散。符号 R 。

2.1.5.2 四分位数间距 inter-quartile range, IQR

上四分位数与下四分位数之差。其越大说明数据变异程度越大，可用于各种分布资料离散趋势的描述。

2.1.5.3 方差 variance

所有观测值离均差平方和的平均值。表示所有观测值相对于均数的平均偏离程度。

2.1.5.4 标准差 standard deviation

所有观测值离均差平方和取平均后的算术平方根。表示所有观测值相对于均数的平均偏离程度。符号 s 。

2.1.5.5 变异系数 coefficient of variation, CV

一种相对离散程度的度量指标，为标准差与均值的比值，通常以百分比表示。消除了数据量纲的影响，用于比较不同均值、不同量纲数据的离散程度，数值越大表明数据相对于均值的离散程度越大，越小则离散程度越小。

2.1.6 相对数 relative number

两个有联系的指标数值通过对比得到的比值，以倍数、百分数、千分数等形式表示，用于反映事物间的数量对比关系、发展变化程度等。

2.1.6.1 频率 frequency

在一定范围内，某一事件或现象在多次重复试验或观测中出现的次数与总试验、观测次数的比值，通常用百分数或小数表示，反映该事件或现象出现的频繁程度，随着试验或观测次数增多，会趋近于理论概率。

2.1.6.1.1 频率分布 distribution of frequency

一种统计方法，将数据集按数值范围或类别分组，统计各组出现的次数或比例。通过表格或直方图展示数据分布特征，揭示数据的集中趋势、离散程度和分布形态，为数据分析提供基础信息。

2.1.6.2 发病率 incidence

单位时间内某事件的发生率。可用该单位时间内某事件的发生数除以可能发生该事件的总人数，或某时间段内的事件发生数除以该时间段内的观察总人数计算。

2.1.6.3 比 ratio

全称“相对比”，又称“比值”。两个有关指标之比，两个指标可以性质相同或不同。表示两个指标的相对大小。

2.1.6.4 标准化率 standardized rate

在对不同人群或样本的率进行比较时，消除内部构成（如年龄、性别等）不同的影响，按照统一的标准人口构成对原始率进行调整计算得到的率，便于更合理地进行率之间的对比分析。

2.1.6.4.1 直接标准化法 direct standardization

在比较不同人群某率时，先选定标准人口构成，将各人群按年龄、性别等分组的实际率乘以标准人口中相应组别人数，得到预期发生数，汇总预期发生数并除以标准人口总数，计算出消除构成差异的率以作比较。

2.1.6.4.2 间接标准化法 indirect standardization

比较不同人群某率时，已知标准人群各分组率，用待比较人群各分组观察单位数乘以标准人群对应分组率，得预期发生数，汇总得预期总数，再用待比较人群实际发生总数除以预期总数，计算出消除构成差异的率来进行比较。

2.1.6.4.3 标准化死亡比 standard mortality ratio

实际上死亡数除以理论上死亡数的商，其中理论上死亡数为各人口构成组中实际人口数与标准人群中这个组死亡率乘积的总和。用来刻画实际人群与标准人群间死亡率的相对大小。

2.1.6.5 动态数列 dynamic series

按时间顺序排列的一系列统计指标数值，如绝对数、相对数或平均数等，通过展示这些指标在不同时间点的变化，反映事物在一段时间内的发展过程、方向和速度，用于趋势分析和预测。

2.1.6.5.1 发展水平 development level

动态数列中，在一定时间条件下所达到的规模或水平，表现为某一时期或时点上的统计指标数值，可通过绝对数、相对数或平均数等形式呈现，是描述和分析事物发展变化的基础数据，反映事物在特定时间的状态。

2.1.6.5.1.1 平均发展水平 average development level

将动态数列中不同时期的发展水平加以平均，消除各期数据波动影响，根据数列类型（时期数列或时点数列），采用不同计算方法（如简单算术平均、加权平均等）得出的代表一定时期内发展的一般水平的数值。

2.1.6.5.2 增长量 increment

动态数列中报告期发展水平与基期发展水平之差。用于表示现象发展状态，正数表明呈上升（正增长）状态，负数表明呈下降（负增长）状态。

2.1.6.5.2.1 逐期增长量 periodic increment

动态数列中各个时间发展水平与前一时间发展水平之差。用来反映所关心指标在给定时间间隔上增加或减少的数量。

2.1.6.5.2.2 平均增长量 average increment

在动态数列中，用逐期增长量之和除以逐期增长量的个数，或用累计增长量除以时间数列项数减1，得到的在一定时期内平均每期增长的数量，反映现象在一定时期内平均增长的水平 and 幅度。

2.1.6.5.3 发展速度 development rate

动态数列中，报告期发展水平与基期发展水平之比，以倍数或百分数表示，反映事物在不同时期发展变化的相对程度，分为定基发展速度和环比发展速度，定基是与某一固定基期比，环比是与前一期比。

2.1.6.5.3.1 环比发展速度 chained development rate

动态数列中某指定时间发展水平与前一相邻时间发展水平之比。用来反映所关心指标在前后时间的发展变化情况。

2.1.6.5.3.2 平均发展速度 average development rate

动态数列中一定时间段内各期环比发展速度（即报告期水平与前一期水平之比）的几何平均数，通过开 n 次方根计算得出。反映指标在较长时间内逐期发展的平均变化程度。

2.1.6.5.4 增长速度 growth rate

动态数列中某指定时期相对于基期的增长量除以基期发展水平所得的商。用以反映所关心指标在指定时期的水平相对于基期水平的增长程度。

2.1.6.5.4.1 平均增长速度 average growth rate

动态数列中一定时间段内各个时间发展水平与前一时间发展水平之比的几何均数。用以反映所关心指标在较长一段时间内逐期平均增长变化的程度。

2.2 统计推断

2.2 统计推断 statistical inference

根据样本数据对总体的特征从统计学上进行的推断。

2.2.1 参数估计 parameter estimation

根据样本数据估算总体中所关注特征的取值大小。

2.2.1.1 抽样分布 sampling distribution

从同一总体中随机抽取例数相同的若干样本，可以获得不完全相同的样本统计量，这些样本统计量的概率

分布。

2.2.1.1.1 标准误 standard error

描述样本统计量（如均值）抽样分布离散程度的指标，反映样本估计值与总体参数的接近程度。计算公式为总体标准差除以样本量的平方根。标准误越小，表明样本统计量越稳定，对总体参数的估计越精确。广泛应用于假设检验、置信区间构建等统计分析，是衡量

估计可靠性的关键指标。

2.2.1.1.2 均方误差 mean squared error, MSE

样本的离均差平方和除以对应的自由度所得到相应的平均变异指标。

2.2.1.1.3 t 变换 t transformation

将随机变量转换为正态分布的方法，与标准正态变换不同的是，总体标准差常常未知，用样本标准差代替。

2.2.1.1.4 自由度 degree of freedom

计算某一统计量时，取值不受约束的个体数目。通常取值为 $n-k$ 。其中， n 为样本含量； k 为被限制的条件数或变量个数，或计算某一统计量时用到其他独立统计量的个数。

2.2.1.1.5 t 界值 critical value of t

在给定自由度和 t 分布曲线下的尾部面积（即概率）时，与横轴对应的分界值。

2.2.1.1.6 单侧概率 one-sided probability

又称“单尾概率(one-tailed probability)”。在统计假设检验中，只考虑可能产生一个方向的极端结果的概率，即概率分布下的一侧尾部面积。可分为左侧概率和右侧概率。

2.2.1.1.7 双侧概率 two-sided probability

又称“双尾概率(two-tailed probability)”。在统计假设检验中，考虑两个方向极端结果的概率，即概率分布下的两侧尾部面积之和。

2.2.1.2 点估计 point estimation

将相应的样本统计量直接作为总体参数的估计值。

2.2.1.2.1 无偏估计 unbiased estimation

在多次重复抽样的情况下，估计参数的期望值等于总体参数的真实值。

2.2.1.2.2 有偏估计 biased estimation

在多次重复抽样的情况下，估计参数的期望值不等于总体参数的真实值。

2.2.1.2.3 有效估计 effective estimation

相对于其他估计方法所得结果而言，具有较小方差的估计值。

2.2.1.2.4 无效估计 ineffective estimation

相对于其他估计方法所得结果而言，具有较大方差的估计值。

2.2.1.2.5 相合估计 consistent estimation

又称“一致估计”。当样本容量趋于无穷大时，以高概率趋于真实参数取值的估计值。

2.2.1.3 区间估计 interval estimation

根据样本对总体参数所在范围（区间）做出的估计，同时给出这一估计的置信度。

2.2.1.3.1 置信区间 confidence interval

由样本对某总体参数所做的区间估计，当重复抽样进

行多次估计时，其中一定比率的区间将成功包含总体参数的真值。

2.2.1.3.1.1 总体均数置信区间 confidence interval for population mean

n 次抽样结果得到的 n 个置信度为（1-检验水准）的样本均数的置信区间，这些区间理论上成功涵盖总体均数的可能性为（1-检验水准）。

2.2.1.3.1.2 总体概率置信区间 confidence interval for population proportion

n 次抽样结果得到的 n 个置信度为（1-检验水准）的样本率的置信区间，这些区间理论上成功涵盖总体概率的可能性为（1-检验水准）。

2.2.1.3.2 置信度 confidence level

置信区间理论上包含未知总体参数的可能性。

2.2.1.3.3 置信限 confidence limit

在给定置信水平下，通过样本统计量及其标准误差估计的总体参数所在的可能范围。

2.2.1.3.3.1 置信上限 upper confidence limit

置信区间上侧的边界值。

2.2.1.3.3.2 置信下限 lower confidence limit

置信区间下侧的边界值。

2.2.1.3.4 估计精度 precision of estimation

估计结果（均数、区间等）的准确程度或可信度的度量，表示估计值与真实值之间的接近程度，通常用误差范围或标准误差等指标反映。

2.2.2 假设检验 hypothesis testing

对所分析的总体情况首先提出一个假设，然后通过样本数据去判断是否拒绝这个假设而实施的过程。

2.2.2.1 零假设 null hypothesis

又称“原假设”。假设检验中提出的一种假设，通常表示无效应、无差异或无关联，统计推断阶段计算有关检验统计量时是以这个假设成立作为前提的。

2.2.2.2 备择假设 alternative hypothesis

又称“对立假设”。假设检验中提出的一种假设，与零假设相互对立，通常表示有效应、有差异或有关联。统计推断当样本信息不支持无效假设时将接受这种假设。

2.2.2.3 检验水准 significance level

假设检验中用以表示拒绝实际上成立的零假设的最大允许概率。

2.2.2.4 双侧检验 two-sided test

仅考虑有无差异，不考虑差异方向的假设检验方法。

2.2.2.5 单侧检验 one-sided test

不仅考虑有无差异，还考虑差异的方向的假设检验方法。

2.2.2.6 P 值 P value

在零假设成立的条件下,出现当前样本统计量以及更极端情况的概率大小。

2.2.2.7 统计学意义 statistical significance

又称“统计显著性”,指研究结果由随机误差导致的概率极小(通常 P 值 <0.05),表明观察到的效应或差异很可能真实存在。其核心是通过假设检验判断数据是否拒绝原假设。

2.2.2.8 I类错误 type I error

又称“I型错误”。假设检验中拒绝了实际上成立的零假设的一类错误。

2.2.2.9 II类错误 type II error

又称“II型错误”。假设检验中未拒绝实际上不成立的零假设的一类错误。

2.2.2.9.1 功效 power

又称“检验效能”。零假设实际不成立,按照设定的检验水准能够拒绝零假设的概率。

2.2.2.10 参数检验 parametric test

当样本来自理论上由若干参数规定的总体分布,且对未知总体参数做统计推断的假设检验方法。

2.2.2.11 非参数检验 nonparametric test

不依赖于总体分布类型,也不对参数进行推断,而是对总体分布进行比较的假设检验方法。

2.2.2.12 拒绝域 rejection area

给定检验水准下,由相应界值确定的得到拒绝零假设结论时,检验统计量所在的区间。

3 统计学方法

3.1 t 检验

3.1 t 检验 t test

又称“学生 t 检验(Student's t test)”。在总体方差未知且样本含量较小时以 t 分布为基础的检验方法。常用于两组及以下定量资料均数比较的假设检验。

3.1.1 t 检验基本方法

3.1.1.1 单样本 t 检验 t test for one sample

用于检验样本均数所代表的未知总体均数是否与某一给定数值有差别的 t 检验方法。

3.1.1.2 配对 t 检验 t test for paired samples

用于检验配对设计中两组样本均数分别代表的总体均数是否有差别的 t 检验方法。

3.1.1.3 两独立样本 t 检验 t test for two independent samples

用于检验两独立样本均数所代表的两总体均数是否有差别的 t 检验方法,要求两样本来自的总体具有相等的方差。

3.1.1.4 近似 t 检验 approximate t test

当两总体方差未知或不相等时,检验两独立样本均数所代表的两总体均数是否有差别的假设检验方法。

3.1.2 前提条件检验方法

3.1.2.1 方差齐性检验 test for homogeneity of variance

判断不同样本所来自的总体是否具有相等方差而实施的假设检验。

3.1.2.1.1 方差齐性 homogeneity of variance

不同组别或条件下的数据方差相等的统计特性,是方

差分析(ANOVA)和线性回归等重要统计方法的前提假设之一。若方差非齐(异方差),可能导致统计检验效能降低或结果偏差。常用 Levene 检验(莱文检验)、Bartlett 检验(巴特利特检验)等方法进行验证。满足方差齐性时,组间比较的结论更可靠,反之需采用校正方法(如 Welch ANOVA)进行分析。

3.1.2.1.2 F 检验 F test

基于两个样本方差的比值,判定所对应的总体方差是否相等的假设检验方法。

3.1.2.1.3 莱文检验 Levene's test

简称“Levene 检验”。将原始观测值转换为与所在组别均数之差的绝对值,再作方差分析,用于检验两个或多个总体方差是否相等的检验方法。

3.1.2.1.4 巴特利特检验 Bartlett's test

简称“Bartlett 检验”。对于本来自正态分布的样本,基于各样本方差和合并方差构造服从卡方分布的检验统计量,据此推断多个样本所来自总体的方差是否相等的检验方法。

3.1.2.1.5 布朗-福赛思检验 Brown-Forsythe test

简称“Brown-Forsythe 检验”。将原始观测值转换为与中位数之差的绝对值,再作方差分析,用于判断多个总体方差是否相等的检验方法。

3.1.2.2 正态性检验 normality test

判断样本所来自的总体是否服从正态分布而进行的假设检验。

3.1.2.2.1 P-P 图 probability-probability plot

以一份观测值从小到大的累计频率作为横坐标,以按照正态分布计算的相应累计概率作为纵坐标,把观测值表现为直角坐标系中的散点,通过判断各散点是否在第一象限紧邻 45 度线来判断资料是否服从正态分布的统计图。

3.1.2.2.2 Q-Q 图 quantile-quantile plot

全称“分位数-分位数图”。以样本的分位数作为横坐标,以按照正态分布计算所得相应分位数作为纵坐标,把观测值表现为直角坐标系中的散点,通过判断各散点是否在第一象限紧邻 45 度线来判断资料是否服从正态分布的统计图。

3.1.2.2.3 *W* 检验 Shapiro-Wilk *W* test

基于样本顺序统计量的正态性检验方法,适用于小样本数据(通常 $n \leq 50$)。其通过计算观测值与预期正态值间的相关性得到统计量 *W*,*W* 值越接近 1 表明数据越符合正态分布。该检验具有较高的统计效能,是统计学中最常用的正态性检验之一,常用于参数检验前对数据分布假设的验证。

3.1.2.2.4 *D* 检验 D'Agostino's test

用于评估数据是否服从正态分布的一种假设检验方法。该检验结合偏度和峰度指标,通过计算综合统计量 *D*,并与临界值比较判断正态性。相较于单指标检验(如 *W* 检验),该检验对非正态特征更敏感,尤其适用于中等样本量($50 \leq n \leq 1000$)的数据分析,是参数检验前重要的正态验证方法之一。

3.1.2.2.5 K-S 检验 Kolmogorov-Smirnov test

一种非参数检验方法,用于检验一个样本是否来自某一特定分布,或比较两个样本是否来自同一分布。通过比较经验分布函数和理论分布函数,或两个经验分布函数间的差异来判断,输出检验统计量和相应的 *p* 值。

3.1.2.2.6 矩法 method of moment

分别对样本分布的偏度和峰度做检验,在两个检验结论均不拒绝零假设时,可作出该样本来自正态分布总体的推论。

3.1.3 Z 检验 Z test

又称“*u* 检验(*u* test)”。在总体标准差已知时以正态分布为基础的检验方法,常用于两组及以下定量资料均数比较的假设检验。

3.2 方差分析

3.2 方差分析 analysis of variance, ANOVA

一种统计方法,通过分解总变异为组间变异和组内变异,检验多个组别均值是否存在显著差异。适用于比较三个及以上独立或相关样本的均值,广泛应用于实验设计和多组数据比较分析。

3.2.1 单因素方差分析 one-way analysis of variance

一种统计分析方法,用于检验单一因素不同水平下多个独立样本均值的差别是否有统计学意义。通过比较组间变异与组内变异,判断因素对观测变量的影响是否有统计学意义,适用于单变量多组数据的比较研究。

3.2.2 多因素方差分析 multi-factor analysis of variance

考虑多个研究因素对实验效应的影响,按照研究设计将总变异分解为多个部分,并确定各因素对研究效应的影响的统计方法。

3.2.2.1 随机区组设计的方差分析 analysis of variance for randomized block design

将所有观察值的总变异按照研究设计分解为处理组间、区组间和误差三部分变异,推断各处理组及各区组间的总体均数是否存在差异的统计方法。

3.2.2.2 重复测量的方差分析 analysis of variance for repeated measures design

将所有观察值的总变异按照研究设计分解为受试对

象间与受试对象内两大部分,推断各处理组及各时间点的总体均数是否存在差异的统计方法。

3.2.2.2.1 球形检验 Mauchly's test of sphericity

用于检验重复测量数据的方差-协方差矩阵是否符合球形假设的统计方法。通过比较不同时间点或条件下的测量值,评估是差异是否具有统计学意义。若不足球形假设,需调整自由度或使用其他方法。广泛应用于心理学和医学的重复测量设计中,确保分析结果的准确性。

3.2.2.3 析因设计的方差分析 analysis of variance for factorial design

所有观察值的总变异按照研究设计分解为主效应变异、交互效应变异与误差变异,推断各处理因素的主效应是否为零及各处理因素是否存在交互效应的统计方法。

3.2.2.4 交叉设计的方差分析 analysis of variance for crossover design

用于分析交叉试验数据的统计方法。在该设计中,受试者按特定顺序接受不同处理,通过比较阶段、处理和个体差异的变异,评估处理效应及残留效应。其优势在于控制个体变异,提高检验效能,广泛应用于医学、心理学等领域的干预研究,但需满足无携带效应和周期效应等前提假设。

3.2.3 变异分解 decomposition of variation

在方差分析中,将全部观测值的总变异按研究设计和需要分解成若干个组成部分的统计方法。

3.2.3.1 总变异 total variation

衡量数据集中所有观测值之间变异程度的统计量,其大小可以用各观察值与总均数的离均差平方和表示。

3.2.3.1.1 离均差平方和 sum of squares of deviations from mean

每个观测值与均值之差的平方和,反映观测值的变异大小。

3.2.3.1.2 均方 mean square

各部分的离均差平方和除以其自由度得到的平均变异指标。

3.2.3.1.3 F 值 F value

一种检验统计量,常用于方差分析或线性回归分析中,数值上等于组间均方、区组均方等均方值除以误差均方。

3.2.3.2 组间变异 variation between groups

每组均数与总均数的离均差平方和,反映不同处理及随机误差的效应。

3.2.3.3 组内变异 variation within group

组内每个个体与组内均数的离均差平方和,反映随机误差的效应。

3.2.3.4 区组变异 variation between blocks

各区组均数与总体的离均差平方和,反映区组的平均效应和随机误差效应。

3.2.3.5 个体内变异 within-subject variation

各处理组下所有个体不同时间点的观测值与不同时间的观察值均数的离均差平方和,反映重复测量因子的效应及重复测量因素与处理因素的交互效应。

3.2.3.6 自由度分解 decomposition of degrees of freedom

将总自由度按照方差分析中总变异的分解方法分解为若干部分的过程。

3.2.4 多重比较 multiple comparison

在方差分析拒绝零假设时,为明确哪些组之间存在差异而进行多个样本均数间两两比较的统计方法。

3.2.4.1 事后两两比较 post hoc comparison

一种多重比较方法,在整体检验出统计学意义后,进一步针对各组均值进行逐对对比的统计分析方法,常用方法包括图基(Tukey)、邦费罗尼(Bonferroni)、沙菲(Scheffé)检验等。

3.2.4.2 事前两两比较 planned comparison

按照资料收集前的研究设计,确定某一对或几对在专业上有特殊意义的均数之差是否具有统计学意义的统计方法。

3.2.4.3 SNK 法 Student-Newman-Keuls test

一种用于方差分析后多重比较的统计方法,适用于均衡设计的多组均值差异检验。其通过分层排序和逐步比较各组均值,控制I类错误膨胀。相较于 Tukey 法更灵敏,但较 LSD 法保守,属于中等严格度的事后检验。需注意:当组数较多时可能增加假阳性风险,适用于组间差异探索性分析。

3.2.4.4 邦费罗尼法 Bonferroni test

又称“Bonferroni 法”。用于多样本均数间两两比较的统计方法,其调整每次两两比较的检验水准以控制犯 I 类错误的累计概率。

3.2.4.5 邓肯法 Duncan test

又称“Duncan 法”。用于多个样本均数依次两两比较的统计方法,其根据各对比组内涵盖的均值的个数调整各自的检验水准限制实验误差并通过计算 q 统计量进行检验。

3.2.4.6 图基法 Turkey's test

又称“Turkey 法”。用于多个样本均数间的两两比较的统计方法,其要求各组样本含量相等,该法能将整个试验的误差控制在检验水准上。

3.2.4.7 邓尼特 t 检验 Dunnett- t test

简称“Dunnett- t 检验”。用于多个实验组均数逐一与对照组均数进行比较的多重比较方法,其使用 t 检验进行两两均数间的比较。

3.2.4.8 LSD- t 检验 least significant difference- t test

用于多组中某一对或几对在专业上有特殊意义的均数比较的统计方法,其使用 t 检验进行两两均数间的比较。

3.3 卡方检验

3.3 卡方检验 chi-square test

又称“ χ^2 检验”。利用两组或多组分类变量的样本信息,考察样本频数分布与假设成立条件下的理论频数分布之间的差异,对总体率的分布或总体频数分布进行推断的假设检验方法。也可用于两个分类变量之间的关联性分析。

3.3.1 拟合优度检验 goodness of fit test

用于评估观测数据与理论分布或模型的匹配程度,通过比较实际频数与期望频数,判断模型的适合性。常用方法包括卡方检验和科尔莫戈罗夫-斯米尔诺夫(Kolmogorov-Smirnov, K-S)检验,适用于正态性检验和模型拟合评价。帮助确定模型的有效性和预测

能力,广泛应用于统计建模、数据分析和科学研究中,确保所选模型合理反映数据特征。

3.3.1.1 实际频数 actual frequency

又称“观察频数(observed frequency)”。实际观察到的某一特定数值或类别的出现次数。

3.3.1.2 理论频数 theoretical frequency

又称“期望频数(expected frequency)”。在无效假设成立的条件下,由合并的样本率估算出的各组样本应分配的平均频数。

3.3.2 皮尔逊卡方检验 Pearson's chi-squared test

简称“Pearson 卡方检验”。针对两组或多组分类变量的总体率或总体频数分布进行推断的方法。

3.3.2.1 耶茨连续性校正 Yates's continuity correction

简称“Yates's 连续性校正”。在小样本情况下,基于频数计算出的离散卡方值近似到连续卡方统计量的过程中的误差增大,为改善卡方统计量分布连续性所提出的连续性校正方法。

3.3.3 麦克尼马尔卡方检验 McNemar's chi-squared test

简称“McNemar 卡方检验”。比较两个相关样本的分类变量频率的假设检验方法。

3.3.3.1 卡帕检验 Kappa test

简称“Kappa 检验”。比较两种不同方法得到的分类

结果是否一致的假设检验方法。通常采用卡帕系数来量化一致性。

3.3.4 趋势卡方检验 chi-square test for trend

分析多个率是否与等级变量间存在线性变化趋势的方法。

3.3.5 科克伦-曼特尔-亨塞尔卡方检验

Cochran-Mantel-Haenszel chi-squared test

又称“Cochran-Mantel-Haenszel 卡方检验”。用于评估在多个分层条件下,两个分类变量之间是否存在关联的统计方法。通过控制混杂因素,计算调整后的卡方统计量,检验变量间的关系是否有统计学意义。广泛应用于流行病学和医学研究,确保结果不受混淆变量的影响,提供更准确的关联分析。

3.3.6 费希尔确切概率法 Fisher's exact probability test

又称“Fisher 确切概率法”。在四格表周边合计不变的条件下,利用超几何分布直接计算样本事件及比样本事件更极端情形的概率的方法。作为卡方检验的补充,适用于样本量小于 40 或理论频数小于 1 的情况。

3.3.7 似然比卡方检验 likelihood ratio chi-squared test

基于似然函数值比较两个模型拟合优度的假设检验方法。两个模型的似然函数值之比经对数转换后服从卡方分布。

3.4 秩和检验

3.4 秩和检验 rank sum test

用于比较两组或多组样本中数据分布的非参数统计方法,通过对数据排序后计算秩和,评估组间差异。适合于样本不满足正态分布或存在离群值的情况,常用方法包括曼-惠特尼(Mann-Whitney U)检验和克鲁斯卡尔-沃利斯(Kruskal-Wallis)检验。广泛应用于生物统计学、心理学和社会科学中,提供对中位数差异的稳健分析。

3.4.1 基本概念

3.4.1.1 秩 rank

变量的一组数据中各个观测值按照大小顺序排列后,每个观测值所在的位置或顺位。

3.4.1.2 相持 tie

两个或多个样本具有相同的观测值,在按大小顺序排列后得到相同秩的现象。

3.4.2 威尔科克森符号秩和检验 Wilcoxon signed-rank test

又称“Wilcoxon 符号秩和检验”。比较两组相关样本或配对样本的中位数大小的假设检验方法。将样本的差值按绝对值大小排列,并赋予秩次,再求秩和,通

过比较正负差值两组的秩和,推断两个总体分布是否存在统计学差异。

3.4.3 威尔科克森秩和检验 Wilcoxon rank sum test

又称“Wilcoxon 秩和检验”。比较两组独立样本的中位数大小的假设检验方法。将两组样本合并,按照从小到大的顺序排列,并赋予秩次,再求秩和,通过比较两组的秩和,推断两个总体分布是否存在统计学差异。

3.4.4 克鲁斯卡尔-沃利斯 H 检验 Kruskal-Wallis H test

又称“Kruskal-Wallis H 检验”。用于比较三组或更多独立样本的非参数统计方法,通过对数据进行排序并计算秩和,计算检验统计量 H,检验组间中位数是否存在显著差异。适用于样本不满足正态分布或方差齐性的情况,是单因素方差分析的非参数替代方法。

3.4.4.1 奈梅尼检验 Nemenyi test

简称“Nemenyi 检验”。基于秩和的非参数多重比较方法。常用于克鲁斯卡尔-沃利斯 H 检验(Kruskal-Wallis H)检验检验拒绝零假设后。

3.4.5 弗里德曼秩和检验 Friedman rank sum test

针对多个相关样本的非参数检验方法，在数据不满足参数检验条件时使用。对每个区组内的数据进行编秩，计算各处理组的秩和，构建检验统计量，推断多个相关样本所来自的总体分布是否存在统计学差异。

3.4.5.1 q 检验 q test

一种多重比较方法，用于在方差分析发现显著差异后，比较各组均值间的具体差异。基于学生化极差分布，控制整体误差率，识别哪些组别间存在统计学差异，常用于图基（Tukey）检验中。

3.5. 相关分析

3.5. 相关分析 correlation analysis

用于评估两个或多个变量之间线性关系的统计方法，通过计算相关系数量化变量间的关联强度和方向。常用的相关系数包括皮尔逊相关系数和斯皮尔曼等级相关系数。适用于探索变量间的协同变化模式，帮助识别潜在的因果关系或预测模型构建。广泛应用于经济学、社会科学和自然科学研究中，提供对变量间关系的定量分析。

3.5.1 基本概念

3.5.1.1 线性相关 linear correlation

又称“简单相关(simple correlation)”。两个或多个随机变量之间呈直线趋势的关系。

3.5.1.1.1 线性相关系数 linear correlation coefficient

又称“Pearson 积矩相关系数(Pearson product moment coefficient)”。描述具有线性关系的两个随机变量间相关方向和密切程度的统计指标。

3.5.1.1.2 斯皮尔曼秩相关系数 Spearman rank correlation coefficient

描述两变量间相关关系的密切程度和相关方向的非参数统计指标，常用于两变量不服从正态分布或者分布类型未知，或为有序分类变量的情形。

3.5.1.1.3 ϕ 系数 phi coefficient

描述两个分类变量之间关联性的统计指标。常用于四格表数据。

3.5.1.1.4 克拉默 V 系数 Cramer's V coefficient

又称“Cramer's V 系数”。描述两个分类变量之间关联性的统计指标。常用于分析多行多列数据时。

3.5.1.1.5 皮尔逊列联系数 Pearson contingency coefficient

描述列联表中两个分类变量之间关联性的统计指标。

3.5.1.1.6 伽马系数 Gamma coefficient

描述两个有序分类变量之间关联性的统计指标。

3.5.1.2 非线性相关 nonlinear correlation

用于描述两个变量之间存在但非直线型的关系。通过非线性模型或变换方法，捕捉变量间的复杂依赖结构。适用于数据呈现曲线、周期性或其他非线性模式的情况。广泛应用于经济学、生物学等领域，提供更准确的变量关系描述，超越简单的线性分析。

3.5.1.3 相关矩阵 correlation matrix

由多个变量相关系数构成的矩阵，该矩阵对角线元素为 1 且对称。

3.5.1.4 相关系数 correlation coefficient

描述两个连续型随机变量间联系强度的统计量。

3.5.1.5 正相关 positive correlation

当一个变量增加时，另一个变量增加，两者呈同向变化的趋势。

3.5.1.6 负相关 negative correlation

当一个变量增加时，另一个变量减少，两者呈反向变化的趋势。

3.5.1.7 零相关 zero correlation

一个变量不随另一变量的改变而改变的趋势。

3.5.1.8 相互独立 mutual independence

一个变量的取值不受其他变量取值的影响。

3.5.1.9 双变量正态分布 bivariate normal distribution

由两个随机变量构成的联合概率分布，两个变量各自服从正态分布，且它们的任意线性组合也服从正态分布，变量间的关系通过协方差或相关系数描述，其概率密度函数呈现出特定的钟形曲面形态，体现两变量的联合变化特征。

3.5.2 复相关 multiple correlation

研究某个变量与多个变量之间的线性相关性及其程度。

3.5.2.1 复相关系数 multiple correlation coefficient

描述某个变量与多个变量间线性关系密切程度的指标。

3.5.3 偏相关 partial correlation

研究校正了其余变量的影响后两变量相关关系的方法。

3.5.3.1 偏相关系数 partial correlation coefficient

描述校正了其余变量的影响后两变量相关关系的指标。

3.5.3.2 部分相关系数 part correlation coefficient

描述某个变量校正其余变量的影响后与另外一个变量的相关关系的指标。

3.5.4 典则相关分析 canonical correlation analysis

用于评估两组多变量之间线性关系的统计方法。通过

提取每组变量的线性组合，最大化这些组合之间的相关性。通过计算典则相关系数衡量相关性强弱，揭示两组变量的内在联系。

3.5.4.1 典则变量 canonical variable

在典则相关分析中，通过线性组合原始变量得到的新变量，用于最大化两组变量间的相关性。每对典则变量反映两组数据间的最强关联模式，其相关系数称为典则相关系数，常用于揭示多变量数据集间的潜在关系。

3.5.4.2 典则系数 canonical coefficient

用于构建典则变量的线性组合系数，反映原始变量对典则变量的贡献大小。通过优化这些系数，最大化两组变量间的典则相关性。每个变量组产生一组系数，帮助解释典则变量的构成和变量间的关系。

3.5.4.3 典则载荷 canonical loading

描述原始变量与典型变量间的强度和方向，反映典则变量对原始变量的贡献程度。

3.5.4.4 冗余度分析 redundancy analysis, RDA

求解某一组典则变量所解释原始变量的变异被另一组典则变量重复解释的百分比的方法。

3.6 对应分析

3.6 对应分析 correspondence analysis, CA

用于分析分类数据的多维统计方法，通过将数据表中的行和列同时映射到低维空间，揭示变量间的相似性和关联模式。生成二维或三维图示，帮助可视化分类变量之间的关系和结构。适用于频数表或计数数据的分析，提供对复杂数据集的直观解释和洞察。

3.6.1 简单对应分析 simple correspondence analysis

分析二维列联表数据关联性的方法。以点的形式在较低维的空间中表示联列表的行与列中各元素的比例结构，通过空间距离反映两个分类变量间的关系。

3.6.2 多重对应分析 multiple correspondence analysis, MCA

分析三维及以上列联表数据关联性的方法。以点的形式在较低维的空间中表示联列表的行与列中各元素的比例结构，通过空间距离反映多个分类变量间的关系。

3.6.3 对应分析图 correspondence analysis map

可视化二维或多维列联表数据关联性的图。以点的形式在较低维的空间中表示联列表的行与列中各元素的比例结构，通过空间距离反映两个或多个分类变量间的关系。

3.6.4 最优对应表 optimal correspondence table

通过对应分析生成的重新排列的二维表，最大化行列

变量之间的关联结构。优化显示变量之间的相似性和差异性，便于识别潜在模式和关系。行和列经过转换后，反映数据集中最强的关联特征。

3.6.5 去趋势对应分析 detrended correspondence analysis, DCA

一种改进的对应分析方法，通过去除数据中的线性趋势，更清晰地揭示分类变量间的非线性关系。适用于处理具有明显趋势的列联表数据，增强低维空间中类别关联结构的解释能力。

3.6.6 正交对应分析 canonical correspondence analysis, CCA

一种改进的对应分析方法，通过正交变换消除维度间的相关性，使低维空间中的坐标轴相互独立。适用于复杂列联表数据的降维分析，增强可视化效果，更清晰地展示分类变量间的关联结构。

3.6.7 判别对应分析 discriminant correspondence analysis

融合对应分析和判别分析的多元统计方法，针对分类数据，在对应分析基础上，依据已知分类信息构建判别函数，使不同类别在低维空间中尽量分离，通过计算判别得分等展示样本归属及类别间关系，输出分类结果和分析图。

3.7 线性回归

3.7 线性回归 linear regression analysis

一种统计建模方法，通过拟合因变量与一个或多个自变量间的线性关系，预测因变量的取值，量化自变量对因变量的影响。

3.7.1 基本概念

3.7.1.1 回归 regression

研究一个变量与另外一个或多个变量间数量依存关系的方法。

3.7.1.2 因变量 dependent variable

又称“应变量，反应变量”。在回归分析中，被预测或被解释的变量。

3.7.1.3 自变量 independent variable

在回归分析中，用于预测或解释观测数据变化的变量。

3.7.1.4 线性回归方程 linear regression equation

利用回归分析所得到的，用于反映自变量与因变量间数量依存关系的方程式。

3.7.1.5 线性回归模型 linear regression model

一种用于定量描述自变量和因变量之间线性依存关系的统计模型。

3.7.1.6 截距 intercept

当所有自变量取值为零时，因变量的平均水平。

3.7.1.7 回归系数 regression coefficient

回归分析中量化自变量对因变量影响作用大小的指标。自变量每增加或减少一个单位，因变量平均改变的单位数。

3.7.1.8 残差 residual

因变量的观测值与基于回归方程得到的因变量预测值之间的差异。

3.7.1.9 最小二乘法 least square method

用于估计模型参数的统计方法，通过最小化观测值与预测值之间差值的平方和来确定最佳拟合线。适用于线性回归和其他模型，确保预测误差最小。

3.7.1.10 离均差平方和 sum of squares of deviations from mean

每个观测值与均值之差的平方和，反映了没有考虑样本量的观测值变异大小。

3.7.1.11 残差平方和 residual sum of squares

因变量的观测值与基于回归方程得到的因变量预测值之差的平方和。代表在总变异中无法用因变量与自变量的线性关系所解释的部分。

3.7.1.12 回归平方和 regression sum of squares

在回归分析中，模型预测值与均值之间差异的平方和，反映在总变异中可以用因变量与自变量的线性关系解释的部分。

3.7.1.13 回归方程的显著性检验 significance test for linear regression

检验自变量与因变量之间线性关系是否有统计学意义的假设检验，通常使用 F 检验或 t 检验，检验结果决定模型的有效性。

3.7.2 简单线性回归 simple linear regression

研究两个变量之间线性关系的统计方法，其中一个为自变量，另一个为因变量。通过拟合直线，描述因变量如何随自变量的变化而变化，常用于预测和趋势分析。

3.7.3 多重线性回归 multiple linear regression

又称“多元线性回归”。研究一个因变量与多个自变量之间线性关系的统计学方法。通过建立线性模型，评估各自变量对因变量的独立影响，常用于复杂数据

的预测和解释。

3.7.3.1 偏回归系数 partial regression coefficient

在多重回归分析中，某个自变量对因变量的影响。表示当方程中其他自变量保持不变时，该自变量每变化一个单位，因变量平均变化的数值。

3.7.3.2 标准化回归系数 standardized regression coefficient

消除了因变量与自变量所取量纲的影响之后的回归系数。其绝对值的大小直接反映了自变量对因变量的影响程度。

3.7.3.3 多重共线性 multicollinearity

在进行回归分析时，两个或多个自变量间存在高度相关，导致模型参数估计的稳定性降低与统计推断失真的现象。

3.7.3.4 变量筛选 variable selection

在数据分析或建模过程中，从一组可能的变量中选择最相关或最有意义的变量，以用于建立模型或进行统计分析。

3.7.3.4.1 最优子集回归 optimum subsets regression

通过遍历所有可能的自变量组合来筛选最优预测模型的方法。

3.7.3.4.2 逐步回归 stepwise regression

在建立多重回归方程时，逐个引入自变量并根据某种准则进行统计检验，保留效应显著的自变量，剔除效应不显著的自变量，直到不再引入和剔除自变量，得到最优的回归方程。

3.7.3.4.2.1 前进法 forward

从建立不包含或包含少量自变量回归方程开始，按照某种规则（如 P 值最小且有统计学意义）每次引入一个自变量，自变量由少到多，直至无可引入的自变量为止的一种变量筛选方法。

3.7.3.4.2.2 逐步前进法 forward stepwise

每选入一个自变量，对已在模型中的自变量进行检验，对达到剔除标准的变量逐一剔除的一种变量筛选方法。

3.7.3.4.2.3 后退法 backward

从建立一个包含全部自变量的回归方程开始，按照某种规则（如 P 值最小且有统计学意义）每次剔除一个自变量，自变量由多到少，直至无可剔除的自变量为止的一种变量筛选方法。

3.7.3.4.2.4 逐步后退法 backward stepwise

每剔除一个自变量，对方程外的变量进行检验，对符合入选标准的变量重新考虑将其纳入的一种变量筛选方法。

3.7.3.4.3 压缩估计 shrinkage estimation

在回归分析等中，对参数估计值进行向某个中心值

(如零)收缩的处理,通过引入正则化项等方法,在降低参数估计方差的同时,适当增加偏差,以提高模型的预测精度和稳定性,获取更可靠的估计结果。

3.7.3.4.3.1 最小绝对收缩选择算法 least absolute shrinkage and selection operator, LASSO

在拟合线性模型时,引入回归系数的L1范数惩罚项的有偏最小二乘回归。该算法倾向于部分自变量的系数压缩为零,实现特征的稀疏性,以达到高维数据中变量选择的目的。

3.7.3.4.3.2 自适应最小绝对收缩选择算法 adaptive least absolute shrinkage and selection operator

简称“adaptive LASSO”。一种同时进行特征选择和正则化的方法,旨在增强统计模型的预测准确性和可解释性,是对最小绝对收缩选择算法的推广。

3.7.3.4.3.3 平滑剪切绝对偏差算法 smoothly clipped absolute deviation, SCAD

一种用于变量选择和参数估计的正则化方法,旨在对模型参数进行平滑修剪,以处理高维数据集和变量选择问题

3.7.3.5 哑变量 dummy variable

又称“虚拟变量”。通常取值为0或1,一种二分类变量,为了使无序多分类变量可以在模型中使用,而进行的k个分类,转换成的k-1个0、1变量。

3.7.4 岭回归 ridge regression

通过在最小二乘法的目标函数中引入一个偏移常数,从而降低多重共线性对回归系数估计影响的统计方法。

3.7.4.1 基本概念

3.7.4.1.1 岭回归估计量 ridge regression estimator

通过在最小二乘法的目标函数中引入一个偏移常数所获得的参数估计量。

3.7.4.1.2 正则化参数 regularization parameter

损失函数中正则化项的系数,决定正则化项对整体目标函数的影响程度。用于抑制模型复杂度、防止过拟合,同时提高模型的泛化能力。

3.7.4.1.3 调节参数 tuning parameter

又称“岭参数(ridge parameter)”。岭回归中用于控制模型复杂度的关键超参数。其通过向损失函数添加正则化项来压缩回归系数,缓解多重共线性问题。

3.7.4.1.4 广义交叉验证 generalized cross-validation

一种用于模型选择和参数调优的统计方法,是交叉验证的一种变体,用于在给定模型和参数集合的情况下估计模型的预测误差。

3.7.4.1.5 岭迹 ridge trace

不同调节参数下,各个特征系数的变化轨迹,通常以图形方式展示。

3.7.4.1.6 方差膨胀因子 variance inflation factor, VIF
评估回归模型中多重共线性的指标,通过计算某个自变量可被其他自变量解释的程度,量化多重共线性对回归系数估计的影响。

3.7.4.2 其他类型估计

3.7.4.2.1 自适应广义岭回归估计量 adaptive generalized ridge regression estimator, AGRRE

将广义岭回归和基于数据特征的自适应调整相结合所得到的参数估计量。它根据数据特点动态调整参数估计,增强模型适应性,以提高预测准确性。

3.7.4.2.2 限制性岭回归估计量 restricted ridge regression estimator, RRRE

在岭回归中引入特定限制条件,通过对参数空间进行约束所获得的参数估计量,以提高模型稳定性与泛化能力,避免过拟合。

3.7.4.2.3 预试验岭回归估计量 preliminary test ridge regression estimator, PTRRE

在进行岭回归分析前,预先检验数据的统计特性,并根据检验结果对模型参数进行调整,最终获得的参数估计量。

3.7.4.2.4 施泰因型岭回归估计量 Stein-type ridge regression estimator, SRRE

又称“Stein型岭回归估计量”。以施泰因(Stein)估计为基础,利用岭回归技术对参数估计进行调整所得到的参数估计量。

3.7.4.2.5 正则性施泰因型岭回归估计量 positive-rule Stein-type ridge estimator, PRSRRE

又称“正则性Stein型岭回归估计量”。结合了施泰因(Stein)估计和岭回归,引入正则化概念的参数估计量,确保参数估计满足正定性规则,提高模型稳定性和预测性能。

3.7.5 稳健估计 robust estimate

一类能够在存在异常值或离群值的情况下获得更可靠的参数估计的统计方法。

3.7.5.1 基本概念

3.7.5.1.1 稳健性 robustness

一个分析方法在遇到异常情况或需求改变时,能够保持稳定和可靠的结果的能力;在统计学中,常用于衡量模型的抗干扰能力和对异常值的容忍程度。

3.7.5.1.2 异常值 outlier

又称“离群值”。与大多数观测数据有显著差异的数据点。

3.7.5.1.3 高杠杆点 high leverage points

回归分析中自变量(X)取值极端、远离数据中心的观测点,其杠杆值显著高于平均水平。这类点可能对回归系数的估计产生较大影响,但若因变量(Y)符

合模型预测趋势，则未必导致模型偏差。高杠杆点需结合残差分析判断其实际影响，常用杠杆值阈值进行识别。

3.7.5.1.4 强影响点 influential point

回归分析中对模型参数估计或预测结果具有显著影响的观测点，其存在可能大幅改变回归线的斜率、截距或拟合优度。这类点通常兼具高杠杆值（自变量 X 极端）和异常残差（因变量 Y 偏离预测）特征，可通过库克距离、DFFITS 或 DFBETAS 等指标量化其影响程度。

3.7.5.1.4.1 库克距离 Cook's distance

又称“Cook 距离”。一种用于衡量数据集中的个别数据点对回归模型拟合结果影响程度的统计量，该值越大，表示该数据点对回归模型的拟合结果产生的影响越大。

3.7.5.1.4.2 DFFITS 统计量 difference in fits statistic

用于衡量回归分析中单个观测点对模型拟合值影响程度的指标，反映删除该观测后预测值的变化。

DFFITS 绝对值越大，表明该点对模型的影响越显著，需结合库克距离等指标进一步诊断是否为强影响点。

3.7.5.1.4.3 安德鲁-普雷吉邦统计量 Andrew-Pregibon statistic

又称“Andrew-Pregibon 统计量”。用于识别对回归模型结果产生较大影响的观测值，通过检测剔除某个观测值后对模型参数估计的影响来评估该观测值对模型结果的影响程度。

3.7.5.1.4.4 似然距离 likelihood distance

用于衡量两个统计模型之间差异的指标，通过比较各自的似然函数值来量化模型与数据的拟合程度差异。常用于模型选择和假设检验，帮助判断哪个模型更适合描述数据。适合在复杂模型或大数据集的背景下应用，广泛用于统计学和机器学习中，提供对模型优劣的定量评估。

3.7.5.1.4.5 均值漂移模型 mean-shift outlier model

一种基于密度的聚类算法，通过找到数据集中局部密度最大值，从而识别数据中的聚类中心。

3.7.5.1.4.6 方差加权模型 variance-weighted model

通过对观测数据进行方差加权，以提高模型预测精度和稳健性的统计方法。分配较高权重给方差较小的观测值，减少噪声对模型的影响。常用于加权最小二乘法和组合预测，适用于异方差性或数据质量不一致的情境，增强模型对不确定性和异常值的抵抗力。

3.7.5.1.5 崩溃点 breakdown point

衡量统计估计方法对异常值抵抗能力的指标，定义为导致估计失效所需的异常值比例。崩溃点越高，方法的稳健性越强，能够抵御更多的异常值影响。用于评

估估计量在数据异常或污染情况下的可靠性。广泛应用于稳健统计分析，帮助选择适合处理噪声和异常数据的模型。

3.7.5.1.6 影响函数 influence function

用于衡量单个观测值对统计估计量影响的工具，描述估计量对数据微小扰动的响应。通过影响函数分析，评估估计方法的稳健性和敏感性，识别异常值或高杠杆点对结果的影响。在稳健统计中广泛应用，支持改进估计方法的设计，确保模型在异常数据或误差下的可靠性。

3.7.5.1.6.1 过失误差敏感度 gross error sensitivity

评估统计模型对误差或异常值敏感程度的指标，反映模型在数据偏离假设时性能下降的程度。高过失误差敏感度意味着模型对误差的抵抗力较弱，可能导致结果不稳定。用于比较不同模型的稳健性和可靠性，帮助选择适合处理不确定性和异常数据的分析方法。在统计学和机器学习中广泛应用，支持模型的优化和改进。

3.7.5.2 估计方法

3.7.5.2.1 R 估计 R-estimation

基于秩统计量的估计方法，通过利用数据的秩信息而非具体数值，提供对参数的稳健估计。适用于处理异常值或非正态分布数据，减少对分布假设的依赖。广泛应用于稳健统计和非参数统计中，增强模型对数据异常和分布偏离的抵抗力，确保估计结果的可靠性和解释性。

3.7.5.2.2 L 估计 L-estimation

基于线性组合顺序统计量的估计方法，通过对样本数据排序后进行加权平均，提供稳健的参数估计。适合处理非正态分布或含异常值的数据，减少对极端值的敏感性。广泛应用于稳健统计和非参数统计，支持在不确定性和数据偏离假设条件下的可靠分析，提升估计的鲁棒性和解释性。

3.7.5.2.3 M 估计 M-estimation

通过最小化损失函数对参数进行稳健估计的方法，广泛用于处理含有异常值或非正态分布的数据。利用加权残差的和，减少极端值对估计结果的影响，增强模型的鲁棒性。适用于回归分析和其他统计模型，提供对传统最小二乘法的改进，确保结果在数据异常情况下的可靠性。

3.7.5.2.3.1 休伯估计 Huber estimation

又称“Huber 估计”。结合最小二乘法和绝对偏差的稳健估计方法，通过调整损失函数来减少异常值影响。在小偏差时采用平方误差，在大偏差时采用绝对误差，提供平衡的估计精度。适用于存在离群点的数据集，提升模型的鲁棒性和可靠性。广泛应用于统计学、经

济学等领域，确保参数估计的稳定性。

3.7.5.2.3.2 双权数估计 bisquare weight estimation

根据未加权拟合的残差进行稳健地加权，删除极端离群点，降低轻度离群点权重，使得在异常值的情况下依然能够产生可靠的回归参数估计。

3.7.5.2.3.3 安德鲁估计 Andrew estimation

又称“Andrew 估计”。用于处理异方差问题的统计方法，通过引入权重函数来对样本观测值进行加权，提供对回归系数的一致和有效的估计。

3.7.5.2.3.4 图基估计 Tukey estimation

又称“Tukey 估计”。一种使用中位数和中位绝对偏差来稳健地估计参数的统计方法，对于受到异常值影响的数据具有较强的抵抗力。

3.7.5.2.4 广义 M 估计 generalized M-estimation

扩展了 M 估计的统计方法，通过引入更灵活的权重函数和损失函数，进一步提高对异常值的鲁棒性。适用于复杂数据结构和多种分布类型，提供更广泛的稳健估计。广泛应用于经济学、金融学等领域，增强模型在非标准数据条件下的稳定性和准确性。

3.7.5.2.5 高崩溃点回归 high breakdown point regression

一类稳健的回归估计方法，具有可承受数据中较大数量的异常值性质。

3.7.5.2.5.1 最小平方中位数回归 least median of squares regression

将最小二乘估计的目标函数改为使各残差平方的中位数最小，得到的一种稳健估计方法。

3.7.5.2.5.2 最小截尾二乘回归 least trimmed squares regression

通过截尾部分极端残差来提高回归模型稳健性的方法，将最小二乘法与截尾技术结合，减少异常值对估计结果的影响。适用于数据中存在显著异常值或非正态分布的情境，增强模型对异常数据的抵抗力。广泛应用于稳健统计分析中，提供可靠的参数估计和模型拟合。

3.7.5.2.5.3 S 估计 S-estimation

一种稳健回归方法，通过最小化残差的尺度函数来估计参数，具有较高的崩溃点，能够有效抵抗数据中大量异常值的干扰，适用于高污染数据集。

3.7.5.2.5.4 广义 S 估计 generalized S-estimation

扩展了 S 估计的统计方法，通过引入更灵活的权重函数和损失函数，进一步提高对异常值的鲁棒性。适用于复杂数据结构和多种分布类型，提供更高的崩溃点和更广泛的稳健估计，增强模型在非标准数据条件下的稳定性和准确性。

3.7.5.2.5.5 矩估计 method of moments estimation

结合了 S-估计的稳健性和 M-估计的效能所得到的改进的稳健估计方法。

3.7.5.2.5.6 τ 估计 τ -estimation

在 S-估计的基础上拓展残差尺度估计量的类型，使其不仅具有对离群值的稳健性、同时具有较高统计效能的一种稳健估计方法。

3.7.6 模型评价

3.7.6.1 模型假设 model assumption

对数据的某些特征或关系进行的假设，例如线性回归模型中的线性关系假设。

3.7.6.1.1 独立性 independence

指两个或多个随机变量之间相互独立、取值不会互相影响的这种关系。

3.7.6.1.1.1 序列相关 serial correlation

时间序列中不同时间点之间的观测值的相关性。

3.7.6.1.1.2 空间自相关 spatial autocorrelation

描述地理空间中邻近位置观测值之间相似性或依赖性的统计特性，常用 Moran's I 等指标度量，反映空间数据的聚集或分散模式，广泛应用于地理学、生态学等领域。

3.7.6.1.2 同方差性 homoscedasticity

不同观测值所对应的随机误差具有相同方差的特征，表明数据的离散程度在不同水平上是一致的。

3.7.6.1.2.1 BP 检验 Breusch-Pagan test

通过对模型残差的平方与自变量进行回归来检验线性回归模型中异方差是否存在的方法，相比怀特检验在残差平方回归模型中不包括原解释变量的平方值和交互项。

3.7.6.1.2.2 格莱泽检验 Glejser's test

又称“Glejser 检验”。一种用于检验回归模型中异方差性 (heteroscedasticity) 的统计方法，使用残差与自变量做回归，以判断误差项的方差是否与自变量存在相关性。

3.7.6.1.2.3 怀特检验 White's test

又称“White 检验”。通过对模型残差的平方与自变量进行回归来检验异方差是否存在及检测自变量是否能够显著地解释残差方差变化的统计检验方法。

3.7.6.1.2.4 异方差性 heteroscedasticity

在经典线性回归模型中，假定随机扰动项具有同方差性，即给定解释变量，随机扰动项的方差等于一个常数。给定解释变量条件下，随机扰动项的方差，也就是因变量的方差，不再保持不变，便称为具有异方差性。

3.7.6.1.3 正态性 normality

描述数据分布是否符合正态分布特征的统计属性。通过偏度、峰度和正态概率图等方法进行检验。适用于

多种统计分析，确保模型假设的有效性。提供数据分布形态的评估依据，确保后续分析的可靠性。

3.7.6.1.4 线性 linearity

两个变量间的关系呈简单直线关系的性质。

3.7.6.1.4.1 非线性 nonlinearity

两个变量间的关系不遵循直线关系。

3.7.6.1.4.2 线性化 linearization

将一个非线性系统、函数或方程在某一点附近近似为一个线性系统、函数或方程的过程。

3.7.6.1.4.3 非线性最小二乘 non-linear least squares

一种数学优化方法，用于估计非线性模型中的参数，以使模型的预测值与实际观测数据之间的残差平方和最小化。

3.7.6.2 残差分析 residual analysis

利用残差提供的信息来考察模型假设合理性与数据可靠性的方法。

3.7.6.3 标准化残差 standardized residual

残差除以其标准差所得到的值。

3.7.6.4 残差图 residual plot

以自变量的观测值或预测值作为横坐标，将对应的残差值作为纵坐标绘制的散点图。

3.7.6.4.1 标准化残差图 standardized residual plot

用于评估回归模型拟合效果的诊断工具，通过绘制标准化残差与预测值或自变量的散点图，检测异常值和模型假设的偏离。标准化残差消除单位影响，便于识别非线性模式、异方差性和异常数据点。广泛应用于回归分析和模型验证，帮助改进模型拟合和提升预测准确性。

3.7.6.4.2 偏残差图 partial residual plot

用于检测和可视化当回归模型存在其他自变量时单个自变量与因变量之间关系的图表工具，以回归模型的拟合值或单个自变量观测值为横坐标，回归模型的标准化残差为纵坐标。

3.7.6.5 决定系数 coefficient of determination

衡量一个回归模型对观测数据的拟合程度的统计量，可以理解为因变量的变异中被模型解释的比例。

3.7.6.5.1 校正决定系数 adjusted coefficient of determination

用于评估回归模型解释变量对因变量解释能力的指标，考虑到随着模型自变量数目的增加，决定系数也随之逐步增加的特点，在决定系数公式中引入一个惩罚项（自由度），对决定系数进行调整后所得的值。该值小于等于决定系数。

3.7.6.6 信息准则 information criterion

用于评估统计模型优劣的标准，平衡模型复杂度与拟合精度。通过量化模型对数据的解释能力与参数数量，

帮助选择最优模型，避免过拟合或欠拟合。常见形式包括赤池信息量准则、贝叶斯信息量准则等。

3.7.6.6.1 赤池信息量准则 Akaike information criterion, AIC

一种评估统计模型拟合优度的标准，由日本统计学家赤池弘次提出。通过权衡模型复杂度与拟合精度，选择最优模型。公式为-2 倍对数似然值加上 2 倍参数个数，惩罚过多参数以避免过拟合。值越小模型越好。广泛应用于时间序列分析、回归分析等领域。

3.7.6.6.2 贝叶斯信息量准则 Bayesian information criterion, BIC

一种基于贝叶斯理论的模型选择标准，由吉迪思·施瓦茨(Gideon Schwarz)提出。通过惩罚复杂模型，平衡拟合优度与参数数量，选择最优模型。公式为-2 倍对数似然值加上参数个数乘以样本量的对数，值越小模型越好。适用于大样本情况下的模型比较。

3.7.6.7 交叉验证 cross validation

一种评估模型泛化能力的统计方法。将数据集分为训练集和验证集，多次重复训练与验证过程，计算平均性能指标。常见形式包括 k 折、留一法等。有效防止过拟合，广泛应用于机器学习、数据挖掘等领域，为模型选择与参数调优提供可靠依据。

3.7.6.7.1 训练集 training set

用于构建和调整统计模型的数据子集。通过输入特征与对应标签，使模型和算法学习数据内在规律，优化模型参数。通常占总数据 60-80%，与验证集、测试集共同构成完整数据集。质量直接影响模型性能，需保证代表性、无偏性。

3.7.6.7.2 验证集 validation set

用于评估和选择统计模型和机器学习模型性能的数据子集。通过测试模型在未见数据上的表现，调整超参数，防止过拟合。通常占总数据 10-20%，独立于训练集和测试集。提供模型泛化能力的初步估计，为最终模型选择提供依据。

3.7.6.7.3 测试集 testing set

用于最终评估统计模型和机器学习模型性能的独立数据子集。通过模拟真实应用场景，测试模型在完全未见数据上的泛化能力。通常占总数据 10-20%，仅在模型训练和调优完成后使用。提供模型性能的客观评价，反映实际应用效果。

3.7.6.7.4 K 折交叉验证 K-fold cross validation

一种评估模型性能的交叉验证方法。将数据集均分为 K 个子集，依次以其中一个子集为验证集，其余为训练集，重复 K 次训练与验证。计算 K 次结果的平均值作为模型性能指标。有效利用有限数据，提供稳定可靠的模型评估结果。

3.7.6.7.5 留一法交叉验证 leave-one-out cross validation

一种模型验证的方法，每次从包含 n 个样本的数据集中留出 1 个样本作为测试集，其余 $n-1$ 个样本作为训练集，重复 n 次后以平均误差评估模型性能。其优势在于充分利用数据且无随机性，但因需训练 n 次模型，计算成本较高，适用于小样本场景。

3.7.6.7.6 留出法 holdout method

一种简单的模型评估方法。将数据集随机分为互斥的两部分，大部分用于训练模型，小部分用于测试模型性能。通常训练集占 70-80%，测试集占 20-30%。实现简单，计算成本低，但评估结果可能受数据划分影响较大。

3.7.6.7.7 内部验证 internal validation

利用训练数据本身评估模型性能的方法。通过重采样技术如交叉验证、自助法，在训练集内部分割出验证集。提供模型泛化能力的初步估计，用于模型选择与参数调优。计算成本较高，但充分利用有限数据，评估结果相对稳定。

3.7.6.7.7.1 过拟合 overfitting

模型在训练数据上表现优异，但在新数据上性能显著下降的现象。由于过度复杂，模型捕捉了训练数据中的噪声和细节，导致泛化能力差。常见于参数过多、训练时间过长的情况。可通过正则化、早停等方法缓解。

3.7.6.7.7.2 欠拟合 underfitting

模型在训练数据和新数据上均表现不佳的现象。由于过于简单，模型未能充分学习数据中的潜在规律。常见于参数过少、训练时间不足的情况。可通过增加模型复杂度、延长训练时间等方法改善。

3.7.6.7.8 外部验证 external validation

利用独立于训练数据的外部数据集评估模型性能的方法。通过测试模型在完全未见数据上的表现，客观反映其泛化能力和实际应用效果。常用于模型最终评估和比较，结果更具说服力，但需要额外收集数据。

3.7.6.7.8.1 泛化性 generalizability

机器学习模型在面对未见过的数据时的性能表现。

3.8 广义线性模型

3.8 广义线性模型 generalized linear regression

一般线性模型的扩展，通过分布函数选择因变量为非正态分布；通过连接函数建立因变量的数学期望值与自变量之间的回归关系。当因变量的分布为正态分布，连接函数为恒等(Identity link)时，可简化为一般线性模型。

3.8.1 基本概念

3.8.1.1 指数族分布 exponential family of distributions

一类重要的概率分布集合，概率密度函数可表示为指数形式，具有统一的数学结构，通过自然参数、充分统计量和规范化函数描述，包括正态分布、泊松分布、二项分布等。广泛应用于广义线性模型、统计力学等领域。具有共轭先验、充分统计量等良好性质，便于理论分析与计算。

3.8.1.2 连接函数 link function

将线性预测器与响应变量的期望值关联起来的函数。通过变换响应变量，使其与线性组合的自变量相关联，常见形式包括逻辑函数、对数函数和概率单位函数。适用于处理不同类型的响应变量，如二项分布和泊松分布，增强模型的灵活性和适应性。

3.8.1.3 最大似然估计 maximum likelihood estimation, MLE

又称“极大似然估计”。一种参数估计方法，通过极大化似然函数寻找最可能产生观测数据的参数值。利

用样本信息，使观测数据出现的概率最大。具有一致性、渐近正态性等优良性质。广泛应用于统计学、机器学习等领域，为模型参数估计提供理论基础。

3.8.2 逻辑斯谛回归 logistic regression

一种用于二分类或多分类问题的广义线性模型。通过对数单位函数(logit 函数)将线性预测结果映射为概率值，描述解释变量与类别概率之间的非线性关系。采用极大似然估计等方法求解参数，输出结果具有概率解释。广泛应用于疾病诊断等分类预测领域。

3.8.2.1 对数单位变换 logit transformation

又称“logit 变换”。一种常用的数据转换方法，将一个取值范围在 0 到 1 之间的概率值，通过逻辑函数转换为取值范围为负无穷大到正无穷大的实数值。

3.8.2.2 比值比 odds ratio

又称“优势比”。病例组暴露人数与非暴露人数的比值除以对照组暴露人数与非暴露人数的比值，是反映疾病与暴露之间关联强度的指标。

3.8.2.3 逻辑斯谛回归模型的假设检验

3.8.2.3.1 拟合优度检验 goodness of fit test

用于评估观测数据与理论分布或模型的匹配程度，通过比较实际频数与期望频数，判断模型的适合性。常用方法包括卡方检验和科尔莫戈罗夫-斯米尔诺夫(Kolmogorov-Smirnov, K-S)检验，适用于正态性检验和模型拟合评价。帮助确定模型的有效性和预测

能力,广泛应用于统计建模、数据分析和科学研究中,确保所选模型合理反映数据特征。

3.8.2.3.2 似然比检验 likelihood ratio test

运用两个对立的似然函数之比检测某个假设是否有效的检验方法。

3.8.2.3.2.1 对数似然比检验 log-likelihood ratio test

运用两个对立的对数似然函数之比检测某个假设是否成立的检验方法。

3.8.2.3.3 计分检验 score test

又称“得分检验”。一种非参数统计方法,即将秩次转换为正态计分,并以此为基础计算检验统计量。

3.8.2.3.4 沃尔德检验 Wald test

又称“Wald 检验”。以测量无约束估计量与约束估计量之间距离为原理的统计检验方法,同时适用于线性与非线性的检验。

3.8.2.3.4.1 沃尔德统计量 Wald statistic

又称“Wald 统计量”。检验回归参数线性约束与非线性约束是否成立的统计量。

3.8.2.4 条件逻辑斯谛回归 conditional logistic regression

又称“条件 logistic 回归”。基于匹配或分层研究设计的逻辑斯谛(logistic)回归模型,考虑了分层和匹配的情况,同时结合条件概率,分析配对或分层数据的最灵活和通用的方法。

3.8.2.5 非条件逻辑斯谛回归 unconditional logistic regression

又称“非条件 logistic 回归”。用于分析二分类结果的统计模型,直接建模因变量与自变量之间的关系。不考虑分层或匹配因素,适用于大规模数据集。通过逻辑函数将线性预测转换为概率,广泛应用于流行病学、社会科学等领域。提供简单直观的分类结果,确保模型解释性和预测准确性。

3.8.2.6 二分类逻辑斯谛回归 binary logistic regression

又称“二分类 logistic 回归”。一种用于二分类问题的广义线性模型。通过对数单位函数(logit 函数)将线性预测结果映射为概率值,描述解释变量与类别概率之间的非线性关系。采用极大似然估计求解参数,输出结果具有概率解释。

3.8.2.7 多分类逻辑斯谛回归 multinomial logistic regression

又称“多分类 logistic 回归”。因变量为三个或以上分类,且链接函数服从多广义伯努利分布的回归模型。

3.8.2.7.1 有序多分类逻辑斯谛模型 ordinal multinomial logistic model

又称“有序多分类 logistic 模型”。以有序多分类变量

为因变量的回归模型。

3.8.2.7.1.1 累积逻辑斯谛回归模型 cumulative logistic regression model

又称“累积 logistic 回归模型”。因变量服从多项逻辑分布,且具有符合自然规律的多个有序类别的一类回归模型,对不同分类类别拟合的模型系数有平行斜率的假定。

3.8.2.7.2 无序多分类逻辑斯谛模型 multinomial logistic model

又称“无序多分类 logistic 模型”。用于处理没有固有顺序的多类别响应变量的回归模型,通过扩展二分类逻辑斯谛回归实现。使用一对多或多对多策略,将多分类问题分解为多个二分类问题,估计每个类别相对于基准类别的概率。

3.8.3 概率单位模型 probit model

一种同时分析众多变量对二分类因变量取值概率影响的广义线性模型,它把取值(分布在实数)范围为实数的因变量通过累积概率函数转换成分布在(0,1)区间的概率值,链接函数包括标准正态分布的累积概率函数的反函数,以及逻辑累积概率函数。

3.8.3.1 概率单位变换 probit transformation

采用标准正态分布的累积概率函数的反函数对概率做单位变换

3.8.4 泊松回归 Poisson regression

又称“Poisson 回归”。因变量服从泊松离散分布的一种回归模型,可用来分析计数资料的离散分布规律。

3.8.4.1 对数变换 logarithm transformation

对原始数据值取对数,是用于将偏态数据转换成一个新尺度的常用数据变换方法。

3.8.4.2 发病率比值 incident rate ratio

不同组别发生某疾病概率的比值。

3.8.4.3 尺度皮尔逊卡方 scaled Pearson chi-square

由数学家卡尔·皮尔逊提出,基于标准误校正的卡方统计量,反映模型的拟合离散度。

3.8.4.4 尺度偏移量 scaled deviance

基于标准误校正的残差偏移量,反映模型的拟合离散度。

3.8.4.5 离散程度 dispersion

衡量数据集中各观测值偏离中心趋势的程度,常用指标包括方差、标准差和极差。反映数据的分散性和变异性,帮助理解数据的分布特征和波动情况。

3.8.5 负二项回归 negative binomial regression

一种处理计数数据的广义线性模型。假设响应变量服从负二项分布,通过对数连接函数建立协变量与期望计数的关系。适用于数据过度离散的情况,比泊松回归更灵活。广泛应用于流行病学、生态学等领域,分

析事件发生次数的影响因素。

3.8.5.1 过离散现象 over-dispersion

随机变量的方差远远大于其均数的现象。

3.8.5.2 负二项分布 negative binomial distribution

描述在独立重复试验中，直到出现固定次数的成功前，失败次数的概率分布。适用于建模过度离散的计数数据，具有比泊松分布更大的方差。通过两个参数控制成功概率和成功次数，广泛应用于生物统计、保险和社会科学中，帮助分析事件发生的频率和变异性。

3.8.6 对数线性模型 log-linear model

可以解决分类变量（因素）之间是否相关，以及分析分类变量对频数的独立影响的一类广义线性模型

3.8.6.1 饱和模型 saturated model

独立参数的个数等于分类变量所有类别的组合数，即包含了所有变量的主效应，低阶交互效应和高阶交互效应的模型，其理论频数完全拟合了实际频数

3.8.6.1.1 一阶交互效应 first order interaction effect

在多元回归模型中，当两个自变量相互作用时，对因变量产生的联合影响。

3.8.6.1.2 二阶交互效应 second order interaction effect

在多元回归模型中，当三个自变量相互作用时，对因变量产生的联合影响。

3.8.6.1.3 两两关联模型 pairwise association model

给变量间的关系设定一个特殊结构，减少交互作用项的数量，在一定程度上描述变量间关系的统计方法。

3.8.6.2 不饱和模型 unsaturated model

独立参数的个数小于分类变量所有类别的组合数的一类对数线性模型模型。

3.8.6.2.1 无二阶交互效应模型 no second-order interaction model

当模型考虑纳入三个自变量，但并未考虑这三个自变量因素间的交互效应的一类不饱和模型。

3.8.6.2.2 条件独立模型 conditional independent model

利用条件独立性假设来对联合概率分布进行因式分解，从而简化概率分布的推理和计算过程的一类模型。

3.8.6.2.3 联合独立模型 jointly independent model

一种统计建模方法，假设多个随机变量在联合分布中相互独立，常用于简化复杂系统的概率计算。通过分解联合概率为边缘概率的乘积，降低模型复杂度，适用于高维数据分析与机器学习任务。

3.8.6.2.4 完全独立模型 completely independent model

在三维或多维分类变量的分析中，所有变量之间完全独立，在概率上可以累乘，该类模型包含全部一阶作用项，但不包含二阶作用项的模型。

3.9 生存分析

3.9 生存分析 survival analysis

用于研究事件发生时间及其影响因素的一类统计方法。主要处理时间至事件数据，考虑删失情况，常用模型包括考克斯（Cox）比例风险模型和卡普兰-迈耶（Kaplan-Meier）估计。广泛应用于医学、工程和社会科学领域。

3.9.1 基本概念

3.9.1.1 终点事件 terminal event

又称“失效事件(failure event)”。在研究中，某个感兴趣的研究事件的结局，例如某种疾病的发生，某种处理（治疗）的反应，疾病的复发或死亡等。

3.9.1.2 起始事件 initial event

在研究中，某个感兴趣的研究事件的开始，例如某健康事件的发病时间，第一次确诊时间，或接受正规治疗的时间等。

3.9.1.3 生存时间 survival time

又称“失效时间(failure time)”。广泛地定义为从规定的观察起点到某一特定终点事件出现的时间长度，狭义定义指研究对象在死亡前所经历的时间。

3.9.1.3.1 总生存时间 overall survival

又称“总生存期”。在纳入生存分析的观察对象中，每个对象从观察起点到终点事件发生（可以是任意原因死亡）所经历的所有时间。

3.9.1.3.2 无进展生存时间 progression-free survival

又称“无进展生存期”。从观察起点（如随机化分组）到研究事件（如疾病）进展或终点事件发生(如因病死亡)所经历的时间，包含了疾病恶化期的概念，可用于评估一些治疗的临床效益。

3.9.1.3.3 无病生存时间 disease-free survival

又称“无病生存期”。研究对象从观察起点（如随机化分组）到某终点事件（如疾病）复发或者任何原因导致的终点事件（如死亡）发生所经历的时间，常用于根治性手术治疗或放疗后的辅助治疗。

3.9.1.3.4 疾病进展时间 time-to-progression

研究对象从观察起点(如随机化分组)至终点事件(如疾病)进展恶化所经历的时间，该时间只包含了该研究事件的恶化期。

3.9.1.4 中位生存时间 median survival time

又称“半数生存期”、“中位生存期”。表示恰有一半观察对象尚存活的时间，取值越长，表示疾病的预后

越好，反之亦然。

3.9.1.5 生存率 survival rate

又称“生存函数(survival function)”。表示观察对象经历若干个单位时间段（单位时间常用1月、1年等）后仍存活的累积可能性。

3.9.1.6 风险率 hazard rate

又称“风险函数(hazard function)”。已经存活了某段时间的观察对象在当前时间点的瞬时死亡概率密度。

3.9.1.7 生存概率 probability of survival

某单位时间段开始时存活的观察对象，到该时段结束时仍存活的可能性。

3.9.1.8 死亡概率 probability of death

某单位时间段开始存活的观察对象，在该时间段内死亡的可能性。

3.9.1.9 累积生存概率 cumulative survival probability

观察对象自观察之日起经历了多个单位时间段后仍然存活的概率。

3.9.1.10 累积风险函数 cumulative hazard function

时间从起点到某一时刻，研究对象发生目标事件（如死亡、失效）的累积概率。通过积分风险函数得到，反映事件随时间的累积风险，常用于生存分析中评估长期风险趋势。

3.9.1.11 生存曲线 survival curve

以观察（随访）时间为横轴，以生存率为纵轴，将生存率连接在一起的曲线图。

3.9.1.12 风险曲线 hazard curve

以观察（随访）时间为横轴，以累积风险函数为纵轴，将累积风险函数连接在一起的曲线图。

3.9.1.13 删失数据 censored data

由于研究终止、失访或事件未发生，在观察期内未能获得研究对象事件发生时间的数据。分为右删失、左删失和区间删失，需采用特殊统计方法处理，如卡普兰-迈耶（Kaplan-Meier）估计。

3.9.1.13.1 左删失 left censoring

在生存分析中，事件发生时间早于观察期开始时间，导致事件时间仅知为某一时间点之前。常见于研究开始时已发生事件的个体，影响数据完整性和分析精度。处理左删失数据需要特殊统计方法。

3.9.1.13.2 区间删失 interval censoring

在生存分析中，事件发生时间仅能确定在两个观察时间点之间，而非确切时间。常见于定期随访研究中，导致部分数据不完整。需要特殊统计方法进行处理，如特思布尔（Turnbull）算法，以确保对事件时间的准确估计。

3.9.1.13.3 右删失 right censoring

在生存分析中，研究对象在观察期结束时未经历感兴

趣事件，确切事件时间未知，仅知其晚于某一时间点。常见于长期随访研究，影响数据的完整性和结果的准确性。处理右删失数据的常用方法包括卡普兰-迈耶（Kaplan-Meier）估计和考克斯（Cox）比例风险模型。

3.9.1.14 完全数据 complete data

从观察起点到发生终点事件所经历的时间完整之数据类型。

3.9.1.15 不完全数据 incomplete data

生存分析中，由于删失（右删失、左删失或区间删失）导致部分研究对象的事件发生时间未知的数据。

3.9.1.16 期初例数 number at risk

在某随访时间或时间段终点事件发生之前尚存活的人数。

3.9.1.17 有效期初例数 effective number at risk

某单位时间段开始尚存活观察对象数量扣去同一单位时间段删失人数的平均，所余下的存活对象数量。

3.9.2 生存率估计

3.9.2.1 寿命表法 life table method

一种基于时间区间的生存率估计方法，将观察期分为若干区间，计算每个区间的存活概率、死亡风险和预期寿命。适用于大样本数据，常用于人口学、保险学和流行病学研究。

3.9.2.2 乘积极限法 product-limited method

又称“卡普兰-迈耶，Kaplan-Meier 法(Kaplan-Meier method)”。通过逐步计算每个事件发生时间点的生存概率，估计生存函数的非参数统计方法。适用于右删失数据，生成阶梯形生存曲线，广泛应用于医学研究和临床试验。

3.9.3 生存率比较

3.9.3.1 时序检验 log-rank test

又称“log-rank 检验”。属于单因素的非参数检验，用于比较两组或多组生存曲线或生存时间是否相同，对远期差异敏感。

3.9.3.2 布雷斯洛检验 Breslow test

又称“Breslow 检验”。属于单因素方法，用于不同组生存曲线的比较，对近期差异敏感。

3.9.4 考克斯比例风险回归 Cox's proportional hazards regression model

又称“Cox 比例风险回归”、“Cox 回归模型”。以生存结局和生存时间为因变量，可同时分析多个因素对生存期和生存结局的影响，可分析截尾数据，不要求数据服从特定的生存分布。

3.9.4.1 比例风险 proportional hazards

任意两个个体风险函数之比，即风险比为常数，不随时间变化。

3.9.4.2 比例风险假定 proportional hazards assumption

简称“PH 假定”。比例风险回归模型的主要前提条件，即协变量对生存率的影响不随时间的改变而改变。只有满足该假定前提下，模型的分析预测才是可靠有效的。

3.9.4.3 基准风险函数 baseline hazard function

模型的所有协变量取值为零时的风险函数。

3.9.4.4 风险比 hazard ratio, HR

又称“风险比率”。比例风险回归模型的重要参数，任意两个个体风险函数之比。

3.9.4.5 半参数模型 semi-parametric model

对于模型构造中一部分具有参数形式，可进行有限维度的分析研究，而余下部分的模型内容不作任何设定的一类模型之总称。

3.9.4.6 偏似然函数 partial likelihood function

一种用于考克斯（Cox）比例风险模型的似然函数，仅依赖于风险因素的排序信息，忽略基线风险函数。通过最大化偏似然估计回归系数，广泛应用于生存分析中的多因素风险建模。

3.9.4.7 预后指数 prognosis index, PI

比例风险回归模型中，各自变量与其回归系数的线性组合，其取值越大，风险函数越大，预后越差。

3.9.5 竞争风险模型 competing risk model

分析多个相互排斥的终点事件对生存时间影响的统计方法，考虑每个事件的发生会影响其他事件的风险。通过识别和量化不同风险来源，常用方法包括 Fine-Gray 模型和累积发生函数。

3.9.5.1 兴趣事件风险 interest event risk

在生存分析的多种潜在结局中，被主要关注的某个生存结局。

3.9.5.2 竞争风险 competing risk

生存分析中，研究对象出现感兴趣的终点事件的同时可能会发生其他多种潜在结局，这些结局事件将阻止感兴趣的终点事件的出现或使其发生的概率降低。

3.9.5.3 累积发生函数 cumulative incidence function

又称“累积发生率”。当存在竞争风险时估计粗发生率的方法，假设竞争事件每次发生有且仅有一种，即各事件累积发生率之和等于复合事件累积发生率。

3.9.5.4 原因别风险函数 cause-specific hazard function

假设竞争风险事件互为独立，将非感兴趣的终点事件视为删失，可依次使用比例风险回归模型分别估计竞争事件的风险函数和危险因素的风险比。

3.9.5.5 累积原因别风险函数 cumulative cause-specific hazard function

考虑特定原因或者病因的累计发生率的一类函数模型。

3.9.5.6 部分分布风险函数 sub-distribution hazard function

在竞争风险模型中，描述特定事件类型在给定时间下发生的瞬时风险，考虑其他事件类型作为竞争风险。通过累积发生函数和子分布危害方法进行估计，提供对特定事件的独立效应分析。

3.9.5.7 法恩-格雷检验 Fine-Gray test

又称“Fine-Gray 检验”。用于竞争风险模型的单因素分析的统计检验方法。

3.9.6 加速失效模型 accelerated failure-time model

将生存时间的对数作为反应变量，研究多个协变量与对数生存时间之间关系的线性回归模型。

3.9.6.1 指数分布 exponential distribution

在任何时间点上的风险函数为一常数，风险函数的大小不受生存时间长短的影响，且生存时间分布服从指数函数。

3.9.6.2 对数正态分布 log-normal distribution

时间变量的对数分布近似中间高、两边低且左右对称，用以描述事件的发生率。

3.9.6.3 对数逻辑斯谛分布 log-logistic distribution

又称“log-logistic 分布”。时间变量遵从对数逻辑斯谛的连续概率分布，用以描述事件的发生率。

3.9.6.4 秩估计方法 rank based estimation procedure

基于秩次排序的非参数估计方法。

3.9.7 寿命表 life table

根据某一特定人群的年龄组死亡率编制出来的一种统计表，说明在特定人群年龄组死亡率的条件下，一群人的生命过程或死亡过程。

3.9.7.1 寿命表类型

3.9.7.1.1 现时寿命表 current life table, period life table

基于横断面资料，以某年（或某时期）所有年龄组死亡率为已知数据，假定同时出生的一代人（一般为 100000），按这年龄组死亡率先后死去，直至死完为止，算出一代人在不同年龄组的“期望寿命”等，由此编制的表格。

3.9.7.1.2 队列寿命表 cohort life table

又称“定群寿命表(generation life table)”。数据由纵向观察而得，反映某一特殊人群（队列）的死亡经历。

3.9.7.1.3 完全寿命表 complete life table

以每 1 岁作为年龄分组的寿命表。

3.9.7.1.4 简略寿命表 abridged life table

从实足 5 岁开始以 5 岁为一个年龄组的寿命表。

3.9.7.1.5 去死因寿命表 cause-eliminated life table

通过假设特定死亡原因被消除，重新计算群体生存概率和预期寿命的统计工具。用于评估特定疾病或风险因素对整体寿命的影响，帮助理解如果消除某一死因，群体寿命将如何改变。

3.9.7.1.6 模型寿命表 model life table

基于理论模型和经验数据构建的寿命表，用于估计不同人口群体的生存概率和预期寿命。通过标准化方法生成，适用于缺乏详细人口数据的地区，提供对人口动态的预测和分析工具。

3.9.7.2 寿命表指标

3.9.7.2.1 实足年龄 exact age

从出生到计算时为止，共经历的周年数或生日数。

3.9.7.2.2 平均存活年数 average years lived

在某年龄组中，每位死亡者平均可存活年数。

3.9.7.2.3 尚存人数 number of survivors

假想的同时出生的一代人中，某年龄段尚存者的平均人数，且一般假定“0~”岁组的人数为100000。

3.9.7.2.4 死亡人数 number of deaths

“理论”上死去的人数，即假想的同时出生的一代人中，某年龄段尚存者按死亡概率计算的死亡人数。

3.9.7.2.5 生存人年数 person-years of life

假想的同时出生的一代人中，某年龄段尚存者在该年龄段生存的人年数。

3.9.7.2.6 生存总人年数 total person-years of life

假想的同时出生的一代人中，某年龄段尚存者在今后所有年龄段生存人年数的总和。

3.9.7.2.7 期望寿命 life expectancy

同时出生的一代人活到X岁时，尚能生存的平均年数，一般使用现时寿命表计算。其中出生时的期望寿命为同时出生的一代人未来的平均存活年数，是比较不同地区、不同时期死亡水平的常用指标。

3.9.7.3 其他指标

3.9.7.3.1 潜在减寿年数 potential years of life lost, PYLL

发生在预先确定的年龄终点前的每一例死者所损失的寿命年数。

3.9.7.3.2 质量调整寿命年 quality adjusted life years, QALY

又称“质量调整生命年”。将生存时间按生存质量高低分段，生存质量高的权重大，权重值取0~1，得到的一种健康状况和寿命质量的正向综合测量指标，1单位取值反映1个健康生存年。

3.9.7.3.3 寿命损失年 years of life lost, YLL

因早死所致的寿命损失人年总和，体现某疾病所致生存数量的减少。

3.9.7.3.4 伤残调整寿命年 disability-adjusted life year, DALY

又称“伤残调整生命年”。从发病到死亡所损失的全部健康寿命年，包括因早死所致的寿命损失人年总和，以及疾病所致伤残引起的健康寿命损失人年总和。

3.9.7.3.5 疾病负担 burden of disease, BOD

疾病、伤残和过早死亡对个人、家庭及整个社会经济和卫生健康造成的负担。

3.9.7.3.6 健康期望寿命 healthy life expectancy, HLE

在期望寿命估算基础上，进一步考虑疾病致残状态和生命质量的情况，反映一个人在完全健康状态下生存的平均年数。

3.9.7.3.7 无伤残期望寿命 disability-free life expectancy, DFLE

扣除各年龄组疾病所致伤残引起的健康寿命损失后，推导出来的健康余寿，单位是年。

3.9.7.3.8 有活力期望寿命 active life expectancy, ALE

基于日常生活能力量表评分测算的人群期望寿命。

3.10 判别分析

3.10 判别分析 discriminant analysis

按照一定的准则，利用样本信息构建函数以判断样本所属类别的分类方法。

3.10.1 基本概念

3.10.1.1 监督分类 supervised classification

在构建分类模型时，利用数据带有的标签（即已知的数据点所属类别的标识）进行模型的训练，训练完成的模型用以对新的、标签未知的数据进行分类的方法。

3.10.1.2 非监督分类 unsupervised classification

不需要利用已知的数据标签，通过挖掘数据中的内在结构对数据进行分类的方法。

3.10.2 距离判别法 distance discriminant analysis

根据样本在空间中与各个总体相距远近来判断其属于某个总体的方法。

3.10.2.1 距离最近准则 minimum distance criterion

又称“最邻近方法(nearest neighbor method)”。样本和某个总体在空间中相距最近，则判断它属于该总体的原则。

3.10.3 费希尔判别法 Fisher discriminant analysis

又称“Fisher 判别法”。在空间中找到一个方向，使得多组高维数据在该方向的投影彼此之间的区分度最大的方法。

3.10.3.1 费希尔准则 Fisher criterion

又称“Fisher 准则”。选择一个投影方向，使得空间中样本点投影到该方向后组内差异达到最小，而组间的差异最大的原则。

3.10.3.2 费希尔判别函数 Fisher discriminant function

又称“Fisher 判别函数”。用于线性判别分析的统计方法，通过寻找最佳投影方向，使得不同类别的样本在投影后具有最大化的类间方差与类内方差比。实现样本分类的同时，降低数据维度。广泛应用于模式识别、机器学习和数据挖掘中，支持高维数据的分类任务，有助于提高分类器的性能和准确性。

3.10.4 贝叶斯判别法 Bayes discriminant analysis

基于贝叶斯定理的分类方法，通过计算后验概率来确定样本所属类别。结合先验概率和似然函数，选择最大后验概率的类别作为分类结果。适用于处理不确定性和复杂数据分布，广泛应用于模式识别、机器学习和统计推断中，提供对分类问题的概率解释和决策支持。

3.10.4.1 最大后验概率准则 maximum posterior probability criterion

基于贝叶斯定理的决策原则，通过选择具有最高后验概率的假设或参数值进行推断。结合先验信息和观测数据，优化决策过程，尤其在分类和参数估计中。广泛应用于机器学习、统计推断和信号处理，支持在不确定性条件下做出最优决策，提供对模型和数据的概率解释。

3.10.4.2 最小平均损失准则 minimum average loss criterion

一种决策准则，通过最小化分类或决策过程中的平均损失函数，确定最优分类规则。结合误分类成本和先验概率，适用于不平衡数据或误分类代价不同的场景。

3.10.4.3 广义平方距离 generalized square distance

为传统欧式距离的扩展，在距离计算中考虑权重或者将数据进行变换，以反应不同维度或特征的重要性。

3.10.5 极大似然法 maximum likelihood method

在参数估计时，找到使得当前观测数据出现的可能性最大的一组模型参数值的方法。

3.10.5.1 极大似然准则 maximum likelihood criterion

一种参数估计的原则，其目标是选择参数值，使得样本数据在给定模型下的似然函数达到最大化，即这组参数使观测数据在模型中的生成可能性最大。

3.10.6 贝叶斯公式法 Bayes formula method

基于贝叶斯条件概率公式为基础导出的判别法，根据

计算得出的条件概率大小判断样本属于某一总体。

3.10.7 逐步判别分析 stepwise discriminant analysis, SDA

一种基于样本变量信息，通过逐步添加或删除变量，利用统计准则筛选变量，最终选择一个变量子集的方法，该子集能够以最优的方式区分样本所属的不同类别。

3.10.8 典则判别分析 canonical discriminant analysis, CDA

一种基于样本的变量信息构建新的变量的线性组合（典则变量）的方法，这些变量能用于最佳区别样本所属的不同类别。

3.10.9 多水平判别分析 multilevel discriminant analysis

多水平分类的条件下，根据某一研究对象的自变量信息算出属于各类别的概率值，据以判别其类别归属问题的一种多变量统计分析方法。

3.10.10 判别效果评价 evaluation of discriminant validity

用以衡量判别模型的性能和分类准确度的指标或方法，常见的指标或方法有：混淆矩阵、准确度、精确度、召回率、F1 分数、受试者操作特征曲线（ROC）曲线、曲线下面积（AUC）等。

3.10.10.1 错判概率 misclassification probability

分类模型将样本归属总体判断错误的概率。

3.10.10.1.1 混淆矩阵 confusion matrix

以矩阵的形式显示模型的实际分类结果与预测分类结果之间的关系，从而帮助分析模型的性能指标。

3.10.10.2 回代检验 retrospective validation

模型建立后，将用于训练模型的全部原始数据重新代入模型，通过比较模型的预测值与真实值的差异来评估模型对训练数据的拟合效果的一种方法。

3.10.10.3 交叉验证 cross validation

一种模型验证的方法，将数据分为训练集和测试集（可分别为多个），用训练集来训练模型，并使用独立的测试集来评估模型的性能。

3.10.10.4 刀切法 jackknife

又称“弃一法”、“折刀法”。一种模型验证方法，从数据集中依次去除一个样本，使用剩余的样本构建模型，并用被排除的样本测试模型的性能。重复上述过程以至每个样本都被测试过一次，最终通过整合所有测试结果来评估模型的整体效果。

3.11 聚类分析

3.11 聚类分析 cluster analysis

又称“集群分析”。一种对数据集中各样本分类的方法，使得同一类别内部的差异性尽可能小，而不同类别间的差异性尽可能大。

3.11.1 样品聚类 sample clustering

又称“Q型聚类(Q-type clustering)”。一种基于样本间的数学距离将样本划分为不同的类的方法。

3.11.1.1 样品间距离 sample distance

用于衡量样本之间相似性的指标。

3.11.1.1.1 欧几里得距离 Euclidean distance

又称“欧氏距离”。用于计算两个点在空间中直线距离的方法，基于几何学中的欧几里得空间定义，为各维度坐标差值的平方之和取平方根。

3.11.1.1.2 曼哈顿距离 Manhattan distance

又称“城市街区距离”。通过计算样本点在多维空间中各坐标轴方向差值的绝对值之和来衡量两点之间的距离

3.11.1.1.3 闵可夫斯基距离 Minkowski distance

又称“闵氏距离”。一种测量多维空间中两个点之间距离的方法，它通过一个可调参数来定义，能够根据该参数的不同，泛化为曼哈顿距离、欧几里得距离等多种距离度量。

3.11.1.1.4 马哈拉诺比斯距离 Mahalanobis distance

又称“马氏距离”。用于计算两个样本点在多维空间中距离的方法，不仅考虑维度间的差异，还考虑维度间的相关性。

3.11.1.2 类间距离计算准则 calculation criterion of inter-class distance

评估不同聚类之间的相似性或差异性的度量准则。

3.11.1.2.1 最短距离法 nearest neighbor

以两类中距离最近的两个样本或变量之间的距离来进行聚类的方法。

3.11.1.2.2 最长距离法 furthest neighbor

以两类中距离最远的两个样本或变量之间的距离来进行聚类的方法。

3.11.1.2.3 中间距离法 intermediate neighbor

以两类中所有样本点之间距离的平均值来进行聚类的方法。

3.11.1.2.4 类平均法 average linkage

以两个聚类各自的中心点（同一类中所有样本点的均值）之间的距离来进行聚类的方法。

3.11.1.2.5 重心距离法 centroid neighbor

以两个聚类的各自的重心之间的距离来进行聚类的方法。

3.11.1.2.6 离差平方和法 Ward's method

简称“Wald法”。在每个聚类中，将所有的样本点到聚类中心的距离的平方相加，再把所有聚类的离差平方和相加，得到最后的评估指标，用于评估聚类质量，其值越小，表示聚类效果越好。

3.11.1.3 系统聚类 hierarchical clustering

又称“层次聚类”、“分层聚类”。将相似的样本或变量进行归类的方法，先将每个被聚对象各自视为一类，然后计算各类之间的距离，将类间距离最近的两类合并成一新类；接着计算新类与其他类间的距离，再将其中最近的两类合并。重复上述过程，逐步合并至所有的被聚对象都合并为一类。

3.11.1.3.1 系统聚类图 hierarchy diagram

一种以树状形式，呈现所有样本之间的相似性和聚类结构的图，图的纵轴表示聚类的相似性度量，横轴表示样本。

3.11.1.3.2 树状图 tree diagram

表示集群（包括单个样本）间内在联系与差异的一种结构图，其中“分枝”表示较小集群，“根”表示较大集群。用于指导在聚类过程中相似性水平的选取。

3.11.1.4 k均值聚类 k-means clustering

一种聚类的方法，首先指定需要划分类的个数，然后按照某种原则选择原始数据中根据预先指定分类个数的样本作为初始凝聚点；基于样本间距离，对除初始凝聚点外的所有样本进行逐个归类，将每个样本归入离初始凝聚点最近的那个类中，该类新的凝聚点更新为该类的均值。重复上述过程，直至所有样本归类完毕不需要再调整归类时为止。

3.11.1.4.1 初始聚类中心 initial cluster center

在多均值聚类开始迭代前，为每个类选择的起始点。

3.11.1.4.2 最终聚类中心 final cluster center

在多均值聚类迭代结束时，计算得到的代表每个类的中心点。

3.11.1.4.3 聚类数 number of clusters

在多均值聚类中，将数据集分为不同类的预定数量，预定的数量通常由分析人员事先决定。

3.11.1.4.4 快速聚类 quick cluster

采用近似计算、降维、采样、初始化优化等策略使得在更短时间内完成聚类分析的一类聚类算法。

3.11.2 变量聚类 variable clustering

又称“R型聚类(R-type clustering)”。基于变量之间的相似性或相关性将变量分组在一起的方法，每个聚类可以用一个单独的成分或变量来表示。

3.11.2.1 变量间相似度 variable similarity

在聚类分析中，对数据集内变量之间相似或相关性的评估。用于确定聚类过程中不同变量的相关性和重要性。

3.11.2.1.1 皮尔逊相关系数 Pearson correlation coefficient

又称“Pearson 相关系数”。用于度量两个连续型变量之间线性关联程度和方向的指标。数值介于-1 和 1 之间，其中 1 表示完全正相关，-1 表示完全负相关，0 表示不相关。

3.11.2.1.2 余弦相似度 cosine similarity

用于衡量两个向量在多维空间中相似程度的指标，定义为这两个向量的夹角余弦值，通过向量的点积除以它们的模长的乘积计算。

3.11.2.1.3 列联系数 coefficient of contingency

用于衡量两个分类变量之间关联程度的指标。它基于列联表计算，表示变量之间的依赖关系强度。其值介于 0 和 1 之间，其中 0 表示两个变量完全独立，1 表示完全依赖。

3.11.2.1.4 斯皮尔曼秩相关系数 Spearman rank correlation coefficient

描述两变量间相关关系的密切程度和相关方向的非参数统计指标。常用于两变量不服从整体分布或者分布型未知，或为有序分类变量数据。

3.11.2.1.5 肯德尔秩相关系数 Kendall rank correlation coefficient

简称“Kendall 秩相关系数”。衡量两个变量之间数据点的排序或排名的相似性。量化了协调对的数量（在两个变量中具有相同顺序的数据点对）和不协调对的数量（在两个变量中具有不同顺序的数据点对）。

3.11.2.2 类间相似度计算准则 calculation criterion of inter-class similarity

评估不同簇或类之间的相似性或差异的标准或指标，

帮助确定聚类质量的算法，特别是在没有真实类别信息的情况下。

3.11.2.2.1 最大相似系数法 maximum similarity coefficient method

以分属两类的两个对象两两距离（相似系数）的最大值，做两类间距离（相似系数）的一种系统聚类法。

3.11.2.2.2 最小相似系数法 minimum similarity coefficient method

以分属两类的两个对象两两距离（相似系数）的最小值，做两类间距离（相似系数）的一种系统聚类法。

3.11.3 模糊聚类 fuzzy clustering

允许数据点同时属于多个不同的簇，而不是严格地属于单个簇的聚类算法。每个数据点具有对于每个簇的隶属度值，表示它属于每个簇的程度，适用于存在交叉特征或模糊边界的数据集样本。

3.11.4 有序样品聚类 ordinal clustering methods

一种专门用于处理有序数据的聚类方法，对于按照一定顺序排列的数据，要求在聚类过程中保持原始顺序不被打乱。通过分析数据点之间的顺序或序列关系，将序列划分为若干连续的区段，使每段内部样品之间的差异最小，而不同区段之间的差异最大。

3.11.5 动态样品聚类 dynamical clustering methods

一种针对随时间变化或按时间顺序生成的数据集进行聚类分析的方法，能够根据数据的动态特性实时更新和调整，以适应新数据点或变化模式。常用的方法包括窗口式聚类、递增式聚类和在线聚类等。

3.11.6 两步聚类 two-step clustering

包含两个步骤的聚类算法。第一步，使用预处理方法将数据点分为多个较小的子集，以降低计算复杂度；第二步，在每个子类上应用聚类算法（通常使用系统聚类方法），将预处理得到的子类合并成最终的聚类结果，以进一步细化聚类结构。

3.12 主成分分析

3.12 主成分分析 principal component analysis, PCA

又称“主分量分析”。一种基于多个定量变量之间相互关系，利用降维思想，并通过线性变换提取少数几个关键综合变量的多元统计分析方法。

3.12.1 降维 dimensionality reduction

将高维度的数据保留下一些最重要的特征，去除噪声与次要特征，以提高数据处理速度的高维度特征数据预处理方法。

3.12.2 特征值 eigenvalue

又称“固有值”、“本征值”。主成分的方差，常用于奇异值分解。

3.12.3 特征向量 eigenvector

又称“固有向量”、“本征向量”。表示协方差矩阵的特征值的向量，通过计算数据矩阵的协方差矩阵并选择使矩阵方差最大的若干个特征，组成特征向量矩阵。可将数据矩阵转换到新的低维空间中，实现数据特征的降维。

3.12.4 主成分 principal component

原始变量的线性组合，旨在以较少的综合变量概括原始数据蕴含的绝大部分信息，从而达到降维的目的。

3.12.5 主成分得分 principal component scores

基于原来变量线性组合及其系数构造的综合变量取

值。

3.12.6 方差贡献率 variancecontributionrate

又称“方差解释率”。协方差矩阵中某特征值在总方差中所占比例。

3.12.7 累积方差贡献率

cumulativevariancecontributionrate

一种衡量所选主成分解释原始数据方差的累积比例的指标。可帮助判定保留主成分的数量，实现最大限度保留数据信息的同时降低数据维度。

3.12.8 成分矩阵 components matrix

又称“载荷矩阵”。由主成分法得到的多个定量变量的负荷矩阵，利用各个原始变量的因子表达式的系数，

表达提取的公因子对原始变量的影响程度。

3.12.9 碎石图 scree plot

一种可视化工具，图形以主成分或因子的方差为纵坐标，以主成分或因子的方差排序序号为横坐标。可直观评估需要保留的主成分或因子数量，通常通过观察特征值下降趋势的“拐点”来确定。

3.12.10 主成分回归 principle component regression, PCR

一种主成分分析的回归分析方法，通过将解释变量的主成分作为回归自变量，构造可用于估计因变量的回归模型，排除部分低方差主成分的影响。

3.13 荟萃分析

3.13 荟萃分析 meta-analysis

又称“元分析”。在系统综述中，针对某一具体问题，系统的收集相关研究结果并进行综合评价和定量分析的方法。

3.13.1 基本概念

3.13.1.1 效应量 effect size

反映各个研究的处理因素或水平与反应变量之间关联大小的无量纲统计量。

3.13.1.1.1 合并效应量 combined effect size

将多个独立研究的结果合并成某个单一的效应大小或效应尺度，即用合并统计量反应多个独立研究的综合效应。

3.13.1.1.1.1 随机效应模型 random effect model, REM

当多个研究的效应统计量存在异质性，且效应统计量之间的差异由非抽样误差来解释时，用于合并效应量的模型。

3.13.1.1.1.2 固定效应模型 fixed effect model, FEM

当多个研究的效应统计量具有同质性，且效应统计量之间的差异由抽样误差来解释时，用于合并效应量的模型。

3.13.1.1.1.3 贝叶斯统计过程 Bayesian statistical procedure

综合未知参数的先验信息与样本信息，依据贝叶斯定理，求出后验分布，根据后验分布推断未知参数的统计方法。经典随机效应模型聚焦于估计总的合并效应，而贝叶斯方法聚焦于试验特有效应。

3.13.1.2 异质性 heterogeneity

样本、数据或研究间存在的变异性，表现为特征、效应或结果的差异。需在数据解读和结果整合时加以识别和处理。

3.13.1.2.1 临床异质性 clinical heterogeneity

在临床研究或实践中，不同患者、干预措施或研究设计等之间存在的变异性

3.13.1.2.2 方法学异质性 methodological heterogeneity

由于各个研究的方法不同，如研究设计、结局测量工具和偏倚风险等差别引起的变异。

3.13.1.2.3 统计学异质性 statistical heterogeneity

由于临床异质性或方法异质性或两者共同作用导致研究效应的变异，超过了随机误差所致的变异。

3.13.1.3 科克伦协作网 Cochrane collaboration

又称“Cochrane 协作网”。一个以已故英国医生和流行病学家阿奇博尔德·考科蓝命名的国际化的公益性学术组织。该组织按疾病、疗法或临床问题收集全世界范围内众多高质量研究，提供所有最新证据，促进循证决策。

3.13.2 常规荟萃分析 conventional meta-analysis

对具有对照组的多个研究结果进行荟萃分析的方法。

3.13.2.1 检索策略 search strategy

在分析检索内容实质的基础上，选择检索系统、检索途径，确定检索词及其相互逻辑关系等的信息检索方案，其实质是对检索过程的科学规划。

3.13.2.1.1 医学主题词 medical subject headings

一种能概括地表示信息主题的语词，目的是为医学期刊文章和书籍建立索引。

3.13.2.1.2 自由词 free-text terms

检索过程中不受词表控制，用于主题标引和检索的自然语词。

3.13.2.2 文献质量评价 quality assessment of references

按照不同的研究类型，对文献所纳入研究的真实性、可靠性、实用性和偏倚风险进行质量评价。

3.13.2.2.1 **内部真实性 internal validity**
又称“内部效度”。研究设计中因果关系的准确性，反映研究排除混杂因素干扰的能力，从而确保结果的可靠性

3.13.2.2.2 **外部真实性 external validity**
又称“外部效度”。研究结论推广至其他情境或人群的能力，反映研究发现的广泛适用性。

3.13.2.3 **异质性检验 heterogeneity test**
用于检验多个具有相同目的的研究是否具有异质性的方法。

3.13.2.3.1 **Q 检验 Q test**
用于评估多个独立研究是否存在异质性的一种方法，通过卡方检验来确定纳入研究的异质性情况，其中检验统计量服从卡方分布，自由度为研究数量减1。

3.13.2.3.2 **I² 检验 I-square test**
用于评估多个独立研究结果异质性的统计方法。检验统计量为一个百分比值，表示异质性部分在效应量总的变异中所占的比重，其取值范围定义在0%~100%之间，其值越大异质性越大。

3.13.2.3.3 **亚组分析 subgroup analysis**
将研究样本按特定特征分组后分别进行统计分析，以探索不同群体间的效应差异。

3.13.2.3.4 **Meta 回归 meta regression**
又称“元回归”。通过线性回归方法，检验干预措施、研究对象特征等因素对合并效应量的影响，旨在探究研究间异质性的来源和协变量对效应大小的影响。

3.13.2.3.5 **敏感性分析 sensitivity analysis**
评估研究结果对关键假设或参数变化稳健性的分析方法，以判断结论在不同情境下的可靠性和合理性。

3.13.2.4 **合并效应量估计 combined effect size estimation**
根据研究设计类型、变量类型和异质性大小，计算效应合并值的点估计及区间估计的方法。

3.13.2.4.1 **森林图 forest plot**
在横坐标刻度为1或0的平面直角坐标系中以一条垂直的无效线为中心，用平行于横轴的多条线段的位置和长短代表每个被纳入研究的效应量大小和可信区间，用一个菱形的位置和大小描述合并效应量的大小和可信区间的一种图形。

3.13.2.5 **一致性检验 consistency test**
对多个独立研究结果统计量的一致性进行检验的方法，如果在统计学上不拒绝无效假设，可将多个研究的统计量进行加权合并。

3.13.3 **累积荟萃分析 cumulative meta-analysis, CMA**
一种荟萃分析的方法，把有共同研究目的的研究结果看成一个连续的统一体，每当有新的研究加入后，按

一定的顺序排列累积的结果，并用森林图表示，可以反映研究结果的动态变化趋势，用于评估各研究对综合结果的影响。

3.13.4 **网状荟萃分析 network meta-analysis, NMA**
一种扩展了传统荟萃分析的方法，通过直接比较将三种及以上的干预措施联系起来形成一个干预网络，在干预网络中再进行间接比较，进而将直接和间接估计值结合在一次分析中。

3.13.4.1 **累积顺位曲线下面积 surface under the cumulative ranking curve, SUCRA**

网状荟萃分析中，用于展示不同干预措施效果排序的图示方法。根据排序概率数据的累积概率数据制作出的各干预措施的累积排序概率图，计算出的曲线下面积。

3.13.4.2 **网状结构图 network plot**

网状荟萃分析中证据网络的图示方法，由代表网络中干预措施的节点和显示干预措施对之间现有直接比较的线条组成。

3.13.4.3 **直接比较效应 direct comparison effect**

干预措施效果估计来自于头对头比较。即通过随机对照试验对感兴趣的两种或多种干预措施进行直接比较。

3.13.4.4 **间接比较效应 indirect comparison effect**

通过整合两个干预措施与共同对照组的直接比较结果，推导出这两种干预措施之间的相对效果的方法

3.13.5 **诊断性荟萃分析 meta-analysis of diagnostic test accuracy studies**

一种针对诊断试验的荟萃分析。通过整合多项诊断试验研究，综合量化评估特定诊断工具或测试的准确性、灵敏度和特异度等指标，以提供更可靠的证据。

3.13.6 **个体数据荟萃分析 individual patient data meta-analysis**

不直接利用已经发表的研究结果总结数据进行荟萃分析，而通过向研究者获取原始数据，进行荟萃分析的一种方法。

3.13.7 **偏倚 bias**

实际测量值或估计值与真实值的非随机偏差。可由系统误差或过失误差造成。

3.13.7.1 **抽样偏倚 sampling bias**

检索文献时产生的偏倚，包括发表偏倚、查找偏倚、索引偏倚、引文偏倚和语种偏倚等。

3.13.7.1.1 **发表偏倚 publication bias**

荟萃分析中关于某主题的已发表文献对该主题下所有研究全貌的系统性偏离。

3.13.7.1.1.1 **漏斗图 funnel plot**

以效应尺度为横轴，效应估计精度（如估计方差的倒

数和样本量)为纵轴的散点图。其假设是随着样本量增加,效应估计的精度提高,变异幅度减小,最终形成倒置的对称漏斗形状。常用于识别发表偏倚:无偏倚时呈对称形状,存在偏倚时则不对称。

3.13.7.1.1.2 失安全系数 fail-safe number

当荟萃分析得到有统计学意义的“阳性”结果时,计算需要多少“阴性”结果才能使结论逆转。数值越大,说明结果越稳定,结论被推翻的可能性越小

3.13.7.1.1.3 剪补法 trim and fill method

一种用于估计缺失的研究个数并对发表偏倚进行矫正的非参数统计方法。基本思想是去掉引起漏斗图不对称的小样本研究,用修剪过的漏斗图估计漏斗的“中心”,接着在重估的中心周围替换去除的研究和其对应的缺失研究。

3.13.7.1.2 查找偏倚 search bias

因检索用词不准确或检索策略失误出现漏检或误检相关文献时产生的偏倚。

3.13.7.1.3 索引偏倚 index bias

因数据库中数据标引不准确而导致相关文献未被检出而产生的偏倚。

3.13.7.1.4 引文偏倚 citation bias

系统综述时,通过参考文献选择所引起的偏倚,如结果为阳性的研究可能比结果为阴性者更多地被作为参考文献加以引用。

3.13.7.1.5 语言偏倚 language bias

又称语种偏倚,在文献检索时因限定语种而产生的偏倚。

3.13.7.2 选择偏倚 selection bias

由于研究对象的选择不当,导致入选者与未入选者特征上存在差异,如缺乏代表性、暴露组与对照组可比性差等,而导致的研究结果偏离真实,从而产生的系统误差。

3.13.7.2.1 纳入标准偏倚 inclusion criteria bias

因文献的选择标准不准确产生的偏倚。

3.13.7.2.2 选择者偏倚 selector bias

根据纳入和剔除标准筛选文献时受到筛选者主观意愿的影响而产生的偏倚。

3.13.7.3 研究内偏倚 within study bias

从纳入的文献中提取数据信息时产生的偏倚。包括提取者偏倚,研究质量评分偏倚,报告偏倚。

3.13.7.3.1 提取者偏倚 extractor bias

荟萃分析者从纳入的研究中提取的数据信息不准确而产生的偏倚。

3.13.7.3.2 研究质量评分偏倚 bias in scoring study quality

对纳入研究的质量的评价不恰当或不够充分、不够全面而产生的偏倚。

3.13.7.3.3 报告偏倚 reporting bias

又称“调查对象偏倚”。在调查过程中研究对象对某些信息的故意夸大或缩小所导致的系统误差。

3.14 时间序列分析

3.14 时间序列分析 time series analysis

将按照时间顺序采集到的观测值序列作为基本信息建立统计学模型,刻画序列历史值演变的规律,实现对将来值的预测。

3.14.1.1 自相关函数 autocorrelation function

度量当前时刻序列值与滞后若干期序列值相关程度及方向的函数。

3.14.1.2 偏自相关函数 partial autocorrelation function

给定一个时段内序列值的条件下,度量该时段之前紧邻时刻与之后紧邻时刻序列值相关程度及方向的函数。

3.14.1.3 平稳时间序列 stationary time series

统计特性不随时间的推移而变化的时间序列,表现为将截取的任一段时间序列平移任意距离后,序列仍表现出良好的相似度。

3.14.1.3.1 单位根检验 unit root test

根据平稳序列所有特征根均在单位圆内、非平稳序列则在单位圆上或单位圆外存在特征根这一性质进行

检验统计量的构造,判断时间序列平稳与否的检验方法。

3.14.1.4 严平稳 strictly stationary

时间序列的统计特性比较严格地在时间上保持不变,表现为时间序列沿时间轴平移任意间隔,均能呈现与当前时间序列相同的所有统计特性。

3.14.1.5 宽平稳 weak stationary

时间序列的统计特性以比较宽松的条件在时间上保持不变,表现为时间序列沿时间轴平移任意间隔,均能呈现与当前时间序列相同的低阶矩统计特性。

3.14.1.6 差分 differencing

使非平稳时间序列变得平稳的数据转换技术,依次对原始序列中下一时刻观测值减去当前值,常用于消除时间序列中以均值漂移为表现形式的不平稳性。

3.14.1.7 季节差分 seasonal differencing

使非平稳时间序列变得平稳的数据转换技术,依次对具有一定周期长度的季节性时间序列中下一个周期的观测值减去对应的当前值,得到的新序列被视作消

除了这个周期长度的季节性。

3.14.1.8 p 阶差分 p-order difference

对原始时间序列重复 p 次差分运算。

3.14.2 指数平滑法 exponential smoothing

一种时间序列预测方法，通过对历史数据赋予指数递减的权重，使得较近的观测值对预测结果影响更大，从而更好地捕捉数据的近期趋势和动态变化。

3.14.2.1 简单移动平均 simple moving average

为早期采集到的时间序列观测值赋以相等的权重系数，形成对下一时刻的预测值。

3.14.2.2 简单指数平滑 simple exponential smoothing

特定的指数平滑法，通过对当前观测值和前一个平滑值进行加权平均来生成预测值，使用一个平滑常数 (α) 来确定当前值对预测的影响程度。

3.14.2.3 霍尔特两参数指数平滑 Holt two-parameter exponential smoothing

简称“Holt 两参数指数平滑”。对于呈线性趋势的时间序列，构建自变量为时间间隔，因变量为预测值的线性回归方程，实现提前期长度为这一时间间隔的预测。

3.14.2.4 霍尔特-温特斯三参数指数平滑

Holt-Winters three-parameter exponential smoothing

简称“Holt-Winters 三参数指数平滑”。对于兼具线性趋势、周期特性的时间序列，在自变量为时间间隔、因变量为预测值的线性回归方程基础上，叠加季节效应后实现提前期长度为这一时间间隔的预测。

3.14.2.4.1 霍尔特-温特斯加法模型 Holt-Winters' additive model

简称“Holt-Winters 加法模型”。对于兼具线性趋势、周期特性的时间序列，在自变量为时间间隔、因变量为预测值的线性回归方程基础上，以加法形式叠加季节效应后实现提前期长度为这一时间间隔的预测。

3.14.2.4.2 霍尔特-温特斯乘法模型 Holt-Winters' multiplicative model

简称“Holt-Winters 乘法模型”。对于兼具线性趋势、周期特性的时间序列，在自变量为时间间隔、因变量为预测值的线性回归方程基础上，以乘法形式叠加季节效应后实现提前期长度为这一时间间隔的预测。

3.14.3 自回归模型 autoregressive model

基于时间序列既往若干期的观测值，构建出的用于预测当前值的模型。

3.14.3.1 自回归移动平均模型 autoregressive moving average model

简称“ARMA 模型”。基于时间序列既往部分观测值及回代性预测所得部分残差值，构建出的用于预测当前值的模型。

3.14.3.1.1 沃尔德分解定理 Wold decomposition theorem

任何协方差平稳过程可以分解为两个不相关的平稳序列之和，其中，一个平稳序列是确定性分量，另一个平稳序列是随机性分量。

3.14.3.1.2 延迟算子 delay operator

表达出将时间序列当前值映射回前一时刻历史值这个转换关系的数学符号。

3.14.3.1.3 移动平均模型 moving average model

简称“MA 模型”。基于时间序列中对既往值作回代性预测所得残差值，构建出的用于预测当前值的模型。

3.14.3.2 克拉默分解定理 Cramer's decomposition theorem

全称“克拉默分解定理”。对于任何时间序列，可以分解为完全由历史信息确定的多项式确定性趋势部分和白噪声序列对应的随机波动部分。

3.14.3.3 自回归求和移动平均模型 autoregressive integrated moving average model

简称“ARIMA 模型”。先借助普通差分实现时间序列平稳化，然后对该序列既往部分观测值及回代性预测所得部分残差，构建出的用于预测当前值的模型。

3.14.3.4 季节性自回归求和移动平均模型 seasonal autoregressive integrated moving average model

先借助季节差分实现时间序列平稳化，然后对该序列既往部分观测值及回代性预测所得部分残差，构建出的用于预测当前值的模型。

3.14.3.5 带有外生变量的求和自回归移动平均模型

autoregressive integrated moving average with exogenous input model

简称“ARIMAX 模型”。同时考虑了外生变量影响的求和自回归移动平均模型。

3.14.3.5.1 协整 cointegration

两条或多条非平稳时间序列的某种线性组合是平稳的，表明这些时间序列在长期内具有稳定关系

3.14.3.5.1.1 协整检验 cointegration test

用于判识多个非平稳时间序列一定形式的线性组合是否呈现平稳特性的统计学检验方法。

3.14.4 自回归条件异方差模型 autoregressive conditional heteroskedasticity model

简称“ARCH 模型”。基于时间序列既往若干期的观测值，同时假设当前时刻的随机误差是此前若干期彼此独立且方差不等的残差项组合形成，进而构建出的用于预测当前值的模型。

3.14.4.1 波动集群 volatility clustering

在消除时间序列的确定性非平稳因素后，残差序列在大部分时间段表现为平稳，但在某些时段出现持续的

高波动或低波动现象。

3.14.4.2 广义自回归条件异方差模型 generalized autoregressive conditional heteroskedastic model

简称“GARCH 模型”。基于时间序列既往若干期的观测值，同时假设当前时刻的随机误差是此前若干期方差不相等且不苟求彼此独立的残差项组合形成，进而构建出的用于预测当前值的模型。

3.14.4.3 自回归-广义自回归条件异方差模型 a combination of autoregressive model and generalized autoregressive conditionally heteroscedastic model

简称“AR-GARCH 模型”。自回归模型指当前值可以表征为历史值的线性组合，广义自回归条件异方差模型则可以归纳时间序列中随机波动的强度因时而变的规律。两者结合形成能够同时描述时间序列中的动态均值和动态变异强度的复合模型。

3.14.5 转折点回归模型 joinpoint regression model

一种用于描述时间序列变化趋势的统计方法。其基本思想是根据时间序列特征，通过若干连接点将研究时间分割成不同区间，并对每个区间进行趋势拟合和优化，进而更详细地评价全局时间范围内不同区间序列的变化特征。

3.14.5.1 简单线性转折点回归模型 simple linear joinpoint regression model

一种回归模型，将时间序列分割为若干区段，并在每个区段内使用简单线性回归模型来描述被观测变量随时间的变化规律。该模型能够捕捉数据中潜在的转折点和趋势变化。

3.14.5.2 对数线性转折点回归模型 log-linear joinpoint regression model

一种回归模型，将时间序列分割为若干区段，并在每个区段内使用对数线性回归模型来描述被观测变量随时间的变化规律。该模型能够有效捕捉数据中的潜在转折点和趋势变化，并适用于处理呈指数增长或衰减的数据。

3.14.5.3 年度变化百分比 annual percent change

采样间隔以年为单位的时间序列中，下一年的观测值与当前年份的观测值之差除以当前年份的观测值算出的变化率。

3.14.5.4 年均变化百分比 average annual percent change

采样间隔以年为单位的时间序列中，时段内末期值相对于初期值的变化率再除以时段长度所得的数值。

3.14.6 中断时间序列分析 interrupted time series analysis

通过比较干预前后动态观测的数据，考察时间序列中可能出现的瞬时变化和趋势变化，以识别干预措施对目标变量的影响，评估干预效果的一种统计方法。

3.14.6.1 瞬时效应 immediate effect

基于干预实施后对时间序列模型参数产生的影响，算出的干预后立即发生的影响大小。

3.14.6.2 长期效应 sustained effect

基于干预实施后对时间序列模型参数产生的影响，算出的干预后持久存在的影响大小。

3.15 空间分析

3.15 空间分析 spatial analysis

研究地理空间中数据分布、变化和关联性的统计方法。

3.15.1 基本概念

3.15.1.1 地理信息系统 geographical information system, GIS

对整个或部分地球表层地理分布相关数据进行采集、储存、管理、运算、分析、显示和描述的技术系统。

3.15.1.2 空间数据 spatial data

用来表示地理空间系统中诸要素的数量、质量、分布、相互联系、变化规律等特征的数据，由位置、属性和时间三部分组成。

3.15.1.2.1 栅格数据 raster data

在二维表面上以规则矩形单元表示空间要素的位置和属性的空间数据类型。

3.15.1.2.2 矢量数据 vector data

以几何学中的点、线、面及其组合体来表示地理特征

及其属性的空间数据类型。

3.15.1.2.3 面板数据 panel data

对同一组个体在多个时间点上重复观测的数据类型，结合了横断面数据和时间序列数据的特点。

3.15.1.2.4 拓扑数据 topological data

描述空间结构中各个元素间的空间关系和连接方式的空间数据类型。

3.15.1.3 空间权重矩阵 spatial weight matrix

通过赋予不同的权重来表达空间邻近关系的工具，在空间统计分析运算中作为重要基础。

3.15.1.3.1 邻接矩阵 adjacent matrix

根据空间元素间相邻关系设定权重的一种矩阵，当空间元素间满足一定的相邻关系时赋予权重 1，否则权重为 0。

3.15.1.3.1.1 车式邻接 Rook adjacent matrix

又称“Rook 邻接”。以共边表示邻接关系，当两空间

元素间有共同的边时赋予权重 1，否则为 0 的一种邻接矩阵。

3.15.1.3.1.2 后式邻接 Queen adjacent matrix

又称“Queen 邻接”。以共边或共点表示邻接关系，当两空间元素间有共同的边或共同的顶点时赋予权重 1，否则为 0 的一种邻接矩阵。

3.15.1.3.1.3 象式邻接 Bishop adjacent matrix

又称“Bishop 邻接”。以共点表示邻接关系，当两空间元素间有共同的顶点时赋予权重 1，否则为 0 的一种邻接矩阵。

3.15.1.3.2 距离矩阵 distance matrix

根据空间元素间的距离关系设定权重的一种矩阵，当空间元素间的距离在一定范围内时给予权重，否则权重为 0。

3.15.1.3.2.1 加权距离 weighted distance

基于不同空间特征之间的相关性或重要性赋予其适当的权重，并以此权重计算距离的度量方式。

3.15.1.3.2.2 经济距离 economic distance

不同地区在经济特征上的差异程度，常用其人均国内生产总值差值的倒数进行衡量。

3.15.1.3.3 k 邻近空间权重矩阵 k-nearest neighbors matrix

表示研究区域内各空间单元间的相互关系的矩阵。计算每个空间单元与其他单元的距离，并按距离由小到大排序，将距离该单元最近的 k 个单元赋予权重 1，其余单元的权重为 0。

3.15.1.3.3.1 距离阈值法 distancethreshold

计算空间元素间的距离，当距离小于预先定义的距离阈值时赋予权重 1，否则为 0 的赋值方法。

3.15.1.3.3.2 高斯函数法 Gaussianfunction

计算空间元素间的距离，通过高斯函数对不同距离空间元素进行权重分配的赋值方法。

3.15.1.3.3.1.3 双重平方函数法 bi-square function

计算空间元素间的距离，并设置一定带宽，若距离在带宽范围内时，则基于二次函数对不同距离空间元素进行权重分配，否则为 0 的赋值方法。

3.15.2 空间模式 spatial pattern

空间数据的空间分布状态或形式，包括完全空间随机、空间聚集和空间规则分布等。

3.15.2.1 空间自相关 spatial autocorrelation

描述地理空间中邻近位置观测值之间相似性或依赖性的统计特性，常用 Moran's I 等指标度量，反映空间数据的聚集或分散模式，广泛应用于地理学、生态学等领域。

3.15.2.1.1 全局空间自相关 global spatial autocorrelation

在空间平稳性的基础上，衡量整个区域内某随机变量属性值空间聚集程度的一种度量指标。

3.15.2.1.1.1 全局莫兰指数 global Moran's index

又称“全局 Moran's 指数”。一种度量全局空间自相关性的指标。通过空间权重矩阵和相邻空间单元在某随机变量上的协方差关系构建，可揭示该随机变量在空间上的聚集或离散模式。

3.15.2.1.1.2 全局吉尔里系数 global Geary's coefficient

又称“全局 Geary's 系数”。一种度量全局空间自相关性的指标。通过空间权重矩阵和相邻空间单元在某随机变量上的差异平方构建，可揭示该随机变量在空间上的聚集或离散模式。

3.15.2.1.1.3 全局格蒂斯-奥德 G global Getis-Ord G

又称“全局 Getis-Ord G”。一种度量全局空间相关性的指标，通过分析所有空间单元间的空间关系以及每个空间单元所对应某随机变量属性值与其相邻单元属性值间数值的乘积大小，估算该随机变量的全局空间相关性，可区分出空间上的高值聚集和低值聚集。

3.15.2.1.2 局部空间自相关 local spatial autocorrelation

衡量特定空间位置与其邻近区域某随机变量属性值的空间关联程度，识别局部区域内的空间关联模式与不同区域间的空间异质性的一种度量指标。

3.15.2.1.2.1 局部莫兰指数 local Moran's index

又称“局部 Moran's 指数”。通过分析空间单元间的空间关系以及特定空间单元所对应某随机变量属性值与其相邻单元属性值间的协方差关系，估算各分区空间单元间的局部空间相关性的度量指标，可揭示某属性值空间正相关和空间负相关等局部空间分布模式。

3.15.2.1.2.2 局部吉尔里系数 local Geary's coefficient

又称“局部 Geary's 系数”。通过分析空间单元间的空间关系以及特定空间单元所对应某随机变量属性值与其相邻单元属性值间的差异平方和大小，估算各分区空间单元间的局部空间相关性的度量指标，可揭示某属性值空间正相关和空间负相关等局部空间分布模式。

3.15.2.1.2.3 局部格蒂斯-奥德 G local Getis-Ord G

又称“局部 Getis-Ord G”。一种衡量局部空间相关性的指标，通过分析空间单元间的空间关系矩阵以及特定空间单元相邻单元属性值与平均属性值间的比值，估算各分区空间单元间的局部空间相关性，可揭示某属性局部空间分布模式。

3.15.2.2 空间异质性 spatial heterogeneity

由于地理空间上的隔离，系统或系统属性在空间分布

上呈现不均匀性、复杂性和变异性的一种特性。

3.15.2.2.1 空间局部异质性 spatial local heterogeneity
特定区域内的观测属性值与整体空间趋势存在差异的现象。

3.15.2.2.2 空间分层异质性 spatial stratified heterogeneity

每个空间层内观测属性值同质而空间层间异质的现象，即不同空间层内方差小于空间层间方差。

3.15.2.2.3 变异分析 variance analysis

地统计学中以区域化变量理论为基础，通过距离函数估算其空间相关性大小，从而量化空间数据的变异结构的一种统计方法。

3.15.2.2.3.1 协方差函数 covariance function

某单个或多个自变量构成的随机过程或随机场在两个时间点或空间点处的随机取值的二阶混合中心矩，用以表示两随机取值间的相关性的一种函数。

3.15.2.2.3.2 半变异函数 semivariable function

某区域内数据点的变异值或变异性与数据点间距离的函数，用于描述数据点间的变异程度随着它们间距离增加而变化的情况，为地统计学的特有函数之一。

3.15.2.3 空间聚集性 spatial clustering

空间对象集中在一个或少数几个区域，其他大面积区域中没有或仅有少数空间对象的空间分布模式。

3.15.2.3.1 托布勒地理学第一定律 Tobler's first law of geography

又称“地理学第一定律”。任何事物都与其他事物相联系，但邻近的事物比较远事物联系更为紧密，形成了距离衰减概念的理论基础。

3.15.2.3.2 全局空间聚焦性 global spatial clustering

用以确定整个研究区域是否存在聚集性、而不考虑局部单个聚集簇的统计学方法。

3.15.2.3.2.1 样方分析法 quadrat analysis

将研究区域划分为面积相等的子区域（样方），统计每个子区域的空间点分布密度，并将其与理论标准分布（随机分布模式）比较，从而判断空间点分布模式的方法。

3.15.2.3.2.2 最邻近距离法 nearest-neighbor distance

通过将邻近点之间的距离和某种理论模式中邻近点之间的距离比较，从而推断研究区域空间点模式的分布特征的方法。

3.15.2.3.2.3 里普利 K 函数 Ripley's K function

通过计算多距离半径范围内空间点对象的密度，并将其与完全空间随机点分布比较，从而得出空间点要素的空间分布模式的一种多距离空间聚类分析方法。

3.15.2.3.2.4 库济克-爱德华法 Cuzick-Edwards method

又称“Cuzick-Edwards 法”。一种用于空间聚集性检验的统计方法。通过计算每个病例中心的 k 个最邻近点的病例数目构建全局统计量，可用于评估病例在空间上的聚集情况。

3.15.2.3.2.5 最大超额事件检验法 maximized excess events test, MEET

通过距离的指数权重函数计算空间单元间加权超额事件数，获得一定范围的空间聚集性尺度下、零假设成立的最小概率估计值（即最小检验统计量）的一种全局聚集性检验方法。

3.15.2.3.3 局部空间聚集性 local spatial clustering

在没有任何先验假设的情况下对局部聚集性进行定位，研究聚集簇的分布位置、并确定其是否具有统计学意义的方法。

3.15.2.3.3.1 空间扫描 spatial scanning

基于一定的扫描窗口探测某现象或某事件在一定区域范围内是否存在聚集现象的方法，可对集群进行精确定位与规模判定。

3.15.2.3.3.2 库济克-爱德华法 Besag-Newell method

又称“Besag-Newell 法”。一种用于局部空间聚集性检验的统计方法。以每个病例的空间位置为起点逐步扩展研究区域，计算达到预设聚集簇大小 K 时所需的研究单元数量，从而判断病例是否在局部空间上存在聚集。

3.15.2.3.3.3 聚集评价排列组合法 cluster evaluation permutation procedure

一种用于局部空间聚集性检验的统计方法。通过一系列重叠的圆对研究区域进行扫描，在保证扫描窗口内的风险人口数固定的前提下，计算窗口内病例数，构建最大值统计量从而进行局部聚集簇探测。

3.15.2.3.3.4 任意形状空间扫描统计量法 arbitrarily shaped spatial scan statistic

运用一系列自由形状的动态扫描窗口，探测研究区域疾病的空间局部聚集性、并确定最可能的和其他可能的聚集簇的方法。

3.15.2.3.3.5 空间热点 spatial hotspot

观察属性值大的空间位置及相互邻近的正空间自相关区域。

3.15.2.3.4 焦点空间聚集性 focused clustering

用以检验在一个事先确定的点、线或其他可疑源附近，是否具有局部聚集簇存在的方法。

3.15.2.3.4.1 斯通检验法 Stone analysis

又称“Stone 检验法”。基于简单有序性限制，使用各向同性回归估计方法检验污染源周围存在超额发病风险变化模式的局部聚集性检验方法。

3.15.2.3.4.2 劳森-沃勒得分检验法 Lawson-Waller score Test

又称“Lawson-Waller 得分检验法”。基于泊松分布，以暴露程度为权重，比较每一个区域的实际观察病例数与理论泊松分布下的期望病例数的一种最高功效检验方法。

3.15.2.3.4.3 迪格尔检验 Diggle test

又称“Diggle 检验”。通过构建最近邻累积分布函数，并将最近邻距离的观测分布与空间完全随机性零假设下预期的理论分布进行比较的检验方法。

3.15.3 空间抽样 spatial sampling

在地理空间或空间数据集中，按照特定规则抽取样本点或区域，以通过有限的样本推断空间区域总体的特征，同时考虑空间自相关性、异质性和分布模式。

3.15.4 空间插值 spatial interpolation

依据已知有限的样本数据点，对区域内未知点进行预测估值的过程。

3.15.4.1 反距离加权法 inverse distance weighting

一种空间插值方法，以待插值点与若干邻近样本之间距离的倒数为权重，采用加权平均的方式对未知点进行估计。距离越近的样本点对估计值的影响越大，距离越远的点影响越小。

3.15.4.2 克里金插值法 Kriging interpolation

基于变异函数理论和结构分析，最大限度利用空间取样所提供的空间位置关系、空间分布结构特征等信息，在有限区域内对区域化变量进行无偏最优估计的一种方法。

3.15.4.2.1 普通克里金法 ordinary Kriging

在区域化变量二阶平稳的假设下，利用拟合的半变异函数估算权重，通过已知观测值的简单线性加权对未知样本点属性值进行线性无偏最优估计的一种线性空间插值统计方法。

3.15.4.2.2 稳健克里金法 robust Kriging

通过引入稳健的权重函数、协方差函数和半变异函数，在提高对异常值的容忍度同时对未知样本点属性值进行预测的一种空间插值统计方法。

3.15.4.2.3 泛克里金法 universal Kriging

在区域化变量非平稳假设下，利用拟合的半变异函数和非平稳区域化变量的协方差函数对未知样本点属性值进行无偏最优估计的一种线性空间插值统计方法。

3.15.4.2.4 中位数平滑克里金法 median-polish Kriging

通过多次迭代提取空间网格各维度的中位数平滑估计值，构建更加稳健的变异函数，应用该平滑估计值对未知样本点属性值进行预测的空间插值统计方法。

3.15.4.2.5 协同克里金法 co-Kriging

假设多个区域变量属性值间存在交互关系，通过建立

交叉协方差函数和交叉变异函数，用易测变量或样品多的变量对难测变量或样品少的变量进行局部估计的空间插值统计方法。

3.15.4.2.6 析取克里金法 disjunctive Kriging

在已知任意区域化变量的二维概率分布的假设下，对待估点的值或待估点值超过给定阈值的概率进行估计的一种非线性空间插值统计方法。

3.15.4.2.7 回归克里金法 regression Kriging

通过构建目标属性值和辅助环境变量间的回归关系，基于拟合回归曲线与待插区域的辅助环境变量值对该区域目标属性值进行预测的空间插值统计方法。

3.15.4.2.8 神经网络克里金法 neural Kriging

基于已知空间样本点，通过人工神经网络训练获得兴趣变量与辅助环境变量间（如地表参数、遥感图像、地质、土壤和土地利用等）的复杂关系，并基于该训练网络对未知样本点属性值进行预测的一种空间插值统计方法。

3.15.4.2.9 贝叶斯克里金法 Bayesian Kriging

通过先验和后验分布对基础半变异函数进行最大似然估计，并基于该半变异函数进行预测的空间插值统计方法。

3.15.4.3 样条函数插值法 spline interpolation

以最小曲率充分逼近各观察点，在保证整个薄板表面的曲率最小的原则下对未知空间点进行插值的一种多用途插值方法。

3.15.4.3.1 三次样条插值法 cubic spline interpolation

基于一定节点将样本分段，在每个子区间构建不超过三次的多项式方程，从而拟合一个可充分逼近所有已知样本点的三次样条函数，以该函数对未知点进行预测的插值方法。

3.15.4.3.2 B 样条插值法 B-spline interpolation

通过构建基于 B 样条曲线的分段多项式函数，使之充分逼近已知样本点数据，从而对未知空间点进行插值预测的方法。

3.15.4.4 自然邻域插值法 natural neighbor interpolation

使用德劳内三角形，对所有样本点创建泰森多边形，利用距样本点最近的待插样本子集，按比例进行加权插值的一种空间局部插值法。

3.15.4.5 核密度估计法 kernel density estimation

在没有任何先验密度假设的前提下，以核函数为核心，基于观测样本点构建概率密度函数，从而获得单变量或多变量属性值概率密度的平滑估计值的一种非参数方法。

3.15.4.5.1 均匀核函数 uniform kernel function

假设每个观测点周围区域的权重相同的一类分布函

数，其函数值在区间 $[-1, 1]$ 内为常数，呈现矩形。

3.15.4.5.2 三角核函数 triangle kernel function

假设每个观测点周围区域的权重呈现线性递减趋势的一类函数，其函数值在区间 $[-1, 1]$ 内线性递减，呈现三角形。

3.15.4.5.3 伽马核函数 Gamma kernel function

通过指数函数对所估算样本点间的欧式距离平方进行非线性映射，在指数部分引入负号使函数值随着离中心点距离的增加而减小，并乘以调节参数调控函数值衰减速度的一类函数。

3.15.4.5.4 高斯核函数 Gaussian kernel function

通过指数函数对所估算样本点间的欧式距离平方进行非线性映射，在指数部分引入负号使函数值随着离中心点距离的增加而减小，并除以带宽参数调控函数作用宽度和平滑程度，从而形成沿径向对称的标量函数。

3.15.4.5.5 叶帕涅奇尼科夫核函数核函数

Epanechnikov kernel function

又称用“Epanechnikov 核函数”。通过对样本点间的标化距离平方进行简单的数学变化，并设置相应的边界带宽使距离大于1的样本点对应函数值为0的一类钟形分布函数。

3.15.4.6 多项式插值法 polynomial interpolation

通过对样本点建立多项式回归，从而拟合光滑面，利用趋势分析来对空间对象的属性进行预测插值的一种不精确插值方法。

3.15.4.6.1 拉格朗日插值法 Lagrange's interpolation

又称“Lagrange 插值法”。基于多个已有观测值构建拉格朗日多项式，获得可通过所有已有观测点的 $n-1$ 次插值多项式，从而对未知点进行预测的一种插值方法。

3.15.4.6.2 牛顿插值法 Newton interpolation

又称“Newton 插值法”。基于 n 个已有观察值，利用 n 阶差商的思想构建牛顿多项式，并以此对未知样本点预测的一种插值方法。

3.15.4.7 趋势面插值法 trend surface interpolation

根据样本点的属性值和空间坐标的关系，对整个研究区域的样本点采用全局多项式、最小二乘法进行拟合，生成反应要素整体渐变趋势的平滑表面，从而对未知点进行插值的一种全局多项式插值法。

3.15.4.7.1 线性趋势插值法 lineartrend interpolation

通过多项式回归将最小二乘表面与各输入点进行拟合，基于趋势面对未知点进行预测的插值方法。

3.15.4.7.2 逻辑型趋势插值法 logical trend

interpolation

当预测空间中的变量为某种现象存在与否的逻辑型

变量时，通过极大似然估计拟合非线性概率表面模型，从而对未知点进行预测的插值方法。

3.15.5 空间回归 spatial regression

考虑空间相邻因素和空间相依性影响的回归模型，根据模型形式可分为空间自回归模型、空间误差模型等。

3.15.5.1 空间自回归模型 spatialautoregressionmodel

在普通经典线性回归模型的基础上，融入空间依赖性或空间自相关性指标，从而形成能处理空间自相关性的空间线性回归模型。

3.15.5.1.1 空间滞后模型 spatial lag model

考虑邻近区域因变量对研究区域因变量的影响，利用空间矩阵来表示不规则空间数据的滞后性，将因变量空间滞后项纳入模型的一种空间回归模型。

3.15.5.1.2 空间误差模型 spatial error model

考虑到未包含于模型中的外生变量信息中存在的空间依赖性，将空间自相关误差项纳入模型，以修正因空间自相关的存在而产生的误差，从而形成的空间回归模型。

3.15.5.1.3 空间杜宾模型 spatial Dubin model

又称“空间 Dubin 模型”。同时考虑了因变量和自变量的空间交互效应，将因变量和自变量的空间滞后项同时纳入模型，从而形成的空间回归模型。

3.15.5.2 地理加权回归模型 geographically weighted regression model

基于局部回归原理，将数据的空间位置嵌入回归参数中，用以研究具有空间（或区域）分布特征的两个或多个变量间的数量依存关系的方法。

3.15.5.3 时空地理加权回归模型 geographically and temporally weighted regression model

基于局部回归原理，将数据的时空特性嵌入回归参数中，用以研究具有时空分布特征的两个或多个变量间的数量依存关系的方法。

3.15.5.4 空间面板模型 spatial panel model

基于具有空间相依性的面板数据，充分考虑其在截面维度和时间维度上的信息，在经典面板数据模型的基础上纳入表示空间单元或时间单元异质性的模型参数的一类统计模型。

3.15.5.4.1 静态空间面板模型 static space panel model

针对具有空间相依性的面板数据，假设空间相依性不随时间变化，仅考虑因变量空间滞后项、自变量空间滞后项和空间自相关误差项等空间维度的滞后项，从而构建的空间计量模型。

3.15.5.4.2 动态空间面板模型 dynamic spatial panel model

针对具有空间相依性的面板数据，同时考量了空间滞后、时间滞后以及时空滞后的复杂空间计量模型。

3.15.5.5 贝叶斯空间模型 Bayesian spatial model

基于贝叶斯层次模型对空间数据进行建模的统计学方法。

3.15.5.5.1 贝萨戈-约克-莫利模型 Besag-York-Mollié model

简称“BYM 模型”。将空间结构化效应与非结构化效应均纳入模型的一种贝叶斯空间统计学方法，主要

用于空间面板数据。

3.15.5.5.2 共同成分模型 shared component model

将每种疾病的潜在风险划分为共同成分和特定疾病成分，通过空间聚类模型对各成分同时建模，并进行贝叶斯估计获得模型参数的一类空间统计建模方法，常用于对两种或两种以上疾病进行联合空间分析。

3.16 结构方程模型

3.16 结构方程模型 structural equation model, SEM

一种融合因子分析与路径分析的多元统计方法，用于检验观测变量与潜在变量之间、以及潜在变量彼此之间的复杂关系。其核心特点是能够同时处理测量工具与构念的关系和构念间的因果关系。

3.16.1 测量模型 measurement model

用于描述潜在变量与观测变量之间关系的数学模型。它是结构方程模型的两大核心组成部分之一，主要用于验证潜在变量是否能够通过观测变量有效测量。

3.16.1.1 显变量 manifest variable

可以直接观测或度量的变量。

3.16.1.2 潜变量 latent variable

又称“潜在变量、隐变量(hidden variable)”。不能直接测量，而需要通过测量多个相关显变量去推测的变量。

3.16.1.3 内生变量 endogenous variable

指模型中被其他变量影响的需要解释的变量，通常是研究者感兴趣的输出变量。包括内生显变量和内生潜变量。

3.16.1.4 外生变量 exogenous variable

指模型中不受其他变量影响的独立变量，通常是指模型的输入变量。包括外生显变量和外生潜变量。

3.16.1.5 因子分析 factor analysis, FA

一种多变量统计方法，旨在从一组观测变量中提取出少数几个潜在的、不可直接测量的因子，以揭示观测变量之间的内在结构关系。

3.16.1.5.1 探索性因子分析 exploratory factor analysis, EFA

简称“因子分析”。一种多变量统计方法，旨在从一组观测变量中探索并提取出潜在的、不可直接测量的因子，以揭示观测变量之间的内在结构关系。通常在研究者对数据结构缺乏明确理论假设时使用。

3.16.1.5.1.1 公共因子 common factor

简称“因子”，又称“共性因子”。从一组观测变量中提取出的、能够解释这些变量之间共同变异的潜在变量。反映观测变量背后的共同维度或潜在结构，是因

子分析的核心概念之一。

3.16.1.5.1.2 特殊因子 specific factor

又称“误差因子”、“个性因子”。与某个特定观测变量相关的、无法被公共因子解释的变异部分。反映观测变量中独特的、独立于公共因子的部分，通常包括测量误差和变量特有的变异。

3.16.1.5.1.3 因子载荷 factor loading

又称“因子负荷”。某个观测变量与对应的因子间的相关系数，取值范围一般在-1 到+1 之间。其绝对值越大，表示该观测变量与对应的因子高度相关。

3.16.1.5.1.4 碎石图 scree plot

以特征值由大到小排序的因子序号为横轴、以每个因子对应的特征值为纵轴绘制的折线图，可反映各因子对数据总方差的解释程度。通常以特征值大于等于 1 或折线下降坡度变缓的拐点为标准，来确定提取因子的数量。

3.16.1.5.1.5 正交因子模型 orthogonal factor model

假设潜在因子之间相互正交（相互独立）的一种探索性因子分析方法，旨在简化因子结构以提高模型的解释度。

3.16.1.5.1.5.1 共性方差 communality

又称“共同度”。观测变量中与潜在因子相关的部分方差，即观测变量中能被潜在因子解释的部分。

3.16.1.5.1.5.2 个性方差 specific variance

又称“特殊方差(unique variance)”。观测变量中与潜在因子无关的部分方差，即观测变量中不能被潜在因子解释的方差。

3.16.1.5.1.5.3 因子得分 factor score

指每个个体或者观测单位在潜在因子上的得分。

3.16.1.5.1.5.3.1 汤姆孙法 Thomson method

又称“Thomson 法”。通过计算标准化观测变量与特征向量的线性组合来估计因子得分的方法。

3.16.1.5.1.5.3.2 巴特利特法 Bartlett method

又称“Bartlett 法”。通过加权最小二乘法得到的因子载荷矩阵和观测变量的值来估计因子得分的方法。

3.16.1.5.1.6 因子模型估计 factor model estimation

通过对观测数据建模，估计潜在因子与观测变量之间的关系的方法。以揭示数据的潜在结构和关联，估计方法包括主成分法、极大似然法、主因子法和迭代主因子法等。

3.16.1.5.1.6.1 主成分法 principal components factor method

通过对观测数据的协方差矩阵进行特征值分解，提取主成分并计算主成分载荷和得分，以解释数据方差的因子模型估计方法。

3.16.1.5.1.6.2 主因子法 principal factor method

通过对观测数据的协方差矩阵进行因子分解，提取与潜在因子相关的方差，以揭示数据中的共性结构的因子模型估计方法。

3.16.1.5.1.6.3 迭代主因子法 iterated principal factor method

通过迭代估计观测数据的因子载荷和个性方差，反复调整模型参数以更准确地拟合数据的因子模型估计方法。

3.16.1.5.1.7 因子旋转 factor rotation

因子分析中，通过调整因子载荷矩阵来提高潜在因子可解释性的数学变换技术，通过对原始潜在因子作线性变换转化为一组新的潜在因子，使得新的潜在因子与可测变量之间的因子载荷向 0 或 1 两极分化，旋转后原可测变量隐含的潜在因子的个数及实际意义不会改变，有正交旋转和斜交旋转。

3.16.1.5.1.7.1 方差最大正交旋转 varimax orthogonal rotation

因子分析中常用的旋转方法，通过调整因子载荷矩阵，使每个公共因子的变量相对载荷的方差之和最大化，且保持原共性因子的正交性及公共方差总和不变，以减少每个潜在因子上具有最大载荷的可测变量数，简化对潜在因子的解释。

3.16.1.5.1.7.2 斜交旋转 oblique rotation

通过调整因子载荷矩阵，允许经旋转后的因子载荷矩阵中因子间存在一定程度的相关性，使新的因子轴穿过因子图上的聚集点，这些点在新因子轴上呈较大负荷，而在其它因子轴上的负荷几乎等于零，从而提高新因子的可解释性。

3.16.1.5.2 验证性因子分析 confirmatory factor analysis, CFA

又称“证实性因子分析、确定性因子分析”。在观测变量间与潜在变量之间内在结构已知的情况下，确定观测变量在潜在变量上负荷的大小、验证数据与已知结构的吻合程度的分析方法。

3.16.2 结构模型 structural model

描述潜在变量之间关系的模型部分。它是结构方程模

型的两大核心组成部分之一，主要用于分析和验证潜在变量之间的因果关系或预测关系。

3.16.2.1 路径图 path diagram

反映变量之间相互关系的路线图。应展示观察变量对潜变量的回归系数，潜变量之间的回归系数，与观测变量相关的测量误差，潜变量预测值的残差 4 个方面的内容。单箭头表示因果关系，双箭头表示相关关系。

3.16.2.2 载荷 load

又称“负荷”。观测变量与潜在变量或因子之间的相关性强度。反映观测变量对潜在变量或因子的贡献程度。

3.16.3 模型设定 model specification

构建结构方程模型时对模型的基本结构和参数的设定，包括纯粹验证模型、选择最优模型和导出模型三种形式。

3.16.3.1 纯粹验证模型 strictly confirmatory model

在构建结构方程模型时，采用严格的、事先设定的理论结构和参数设定来拟合观测数据的模型。

3.16.3.2 备选模型 alternative model, AM

在构建结构方程模型时，用于与主要理论模型比较的替代性模型，以评估不同理论模型对观测数据的拟合程度。

3.16.3.3 导出模型 generative model

在构建结构方程模型时，对原始模型作限制、简化、扩展或调整参数而产生的次级模型。

3.16.3.4 自由参数 free parameter

在结构方程模型中，未被理论设定的参数，其取值通过数据进行估计。

3.16.3.5 固定参数 fixed parameter

在结构方程模型中，由研究者根据先验知识或理论知识而事先设定的参数，其取值不由模型拟合过程估计。

3.16.3.6 约束参数 constrained parameter

在结构方程模型中，根据研究者的理论假设，通过设定等式或限制条件以规定其取值的参数。

3.16.4 模型识别

3.16.4.1 矩结构 moment structure

样本矩阵（样本数据的协方差或相关矩阵）与模型参数和总体矩阵（总体协方差或相关矩阵）之间的关系结构。

3.16.4.2 测量尺度 measurement scale

用于度量变量的一种尺度，在结构方程模型中，是模型识别依据的必要条件之一。可通过将潜变量的方差设为 1，即将潜变量标准化，或将潜变量的观测标识中任何一个因子负载设为 1 来设定。

3.16.5 模型估计

3.16.5.1 最大似然估计 maximum likelihood estimation, MLE

又称“极大似然估计”。基于观测数据，通过使得模型似然函数最大而获得一组模型参数，模型拟合结果最接近实际数据分布。

3.16.5.2 未加权最小二乘 unweighted least squares, ULS

一种使实际观测值与估计值之差的平方和达到最小的参数估计方法，其中所有观测值的权重相同。

3.16.5.3 加权最小二乘 weighted least squares, WLS

一种使实际观测值与估计值之差的平方和达到最小的参数估计方法，其中不同的观测值被赋予不同的权重，方差较小的观测数据会被赋予更高的权重，而方差较大的观测数据则被赋予较低的权重，通常在数据的方差存在异质性时使用。

3.16.5.4 广义最小二乘 generalized least squares, GLS

一种通过对数据的协方差矩阵进行变换来估计模型参数的最小二乘估计方法，主要思想是通过权重矩阵来对观测数据进行加权，以考虑观测数据之间的相关性和方差异质性。

3.16.5.5 对角加权最小二乘 diagonal weighted least squares, DWLS

一种简化的加权最小二乘法。其中权重矩阵仅在对角线上有非零值，每个观测值根据其自身的方差被独立加权，而不考虑观测值之间的协方差。适用于处理有序分类数据。

3.16.6 模型评价

3.16.6.1 绝对拟合指数 absolute fit index

衡量模型对数据拟合情况的指数。比较模型预测数据与观测数据之间的差异，不涉及与其他模型的比较。

3.16.6.1.1 拟合优度指数 goodness of fit index, GFI

模型能够解释的协方差比例指数。取值在 0 到 1 之间，取值越大越好，一般取值大于 0.9 时，认为拟合程度好。

3.16.6.1.2 调整的拟合优度指数 adjusted goodness of fit index, AGFI

在拟合优度指数的基础上，额外考虑模型的自由度和样本量大小的指数。

3.16.6.1.3 均方根残差 root mean square residual, RMR

指模型预测值与实际观测值之间差异的平方根均值，较低的值表明模型预测值与实际观测值之间的差异较小，即模型拟合度较好。

3.16.6.1.4 近似误差均方根 root mean square error of approximation, RMSEA

评估模型与实际观测值之间的近似误差，考虑了模型的自由度和样本量大小。

3.16.6.2 相对拟合指数 comparative fit index

通过与基准模型的拟合优度比较来评估模型相对拟合情况的指数。

3.16.6.2.1 规范拟合指数 normed fit index, NFI

相对于基准模型的卡方值降低的比例指数。取值在 0 到 1 之间，越接近 1 表明模型拟合程度越好。

3.16.6.2.2 不规范拟合指数 non-normed fit index, NNFI

在规范拟合指数的基础上校正了自由度对其影响，从而避免模型复杂度影响的指数。

3.16.6.2.3 增值拟合指数 incremental fit index, IFI

一个反映假设模型相对于基准模型（一般是零模型）的改善程度的拟合度指标，通过比较两个模型的卡方值和自由度差异来评估，但不对样本量进行调整。

3.16.6.3 信息准则拟合指数 information criteria fit index

又称“信息标准指数(information criteria index)”。用于比较同一数据集上的两个或多个模型拟合优度的统计指标。包括赤池信息准则、一致性赤池信息准则和期望交叉验证指数。

3.16.6.3.1 一致性赤池信息准则 consistent Akaike information criterion, CAIC

赤池信息准则的一种修正形式，考虑了样本量大小对模型选择的影响。

3.16.6.3.2 期望交叉验证指数 expected cross validation index, ECVI

评估模型在交叉验证情境下的预期性能，度量了由样本得到的拟合协方差矩阵和其他具有相同样本数的样本得到的预期协方差矩阵的差异。度量了由样本得到的拟合协方差矩阵和其他具有相同样本数的样本得到的预期协方差矩阵的差异，可以取任意值，取值越小表示模型重复的可能性越大。

3.16.6.4 节俭拟合指数 parsimony fit index, PFI

一种在评估模型拟合优度时同时考虑模型参数的数量的拟合指数。指在两个或多个模型在数据拟合度上相似或相同的情况下，参数数量较少的模型被认为更优。

3.16.7 模型修正 model modification

通过适当地改变模型中某些变量之间的关系，从而使模型拟合效果更好。如改变测量模型和结构参数，设定某些误差项相关，或者限制某些结构参数。

3.16.7.1 修正指数 modification index, MI

当单个固定参数或约束参数被释放为自由参数时，反映新拟合的模型所引起的卡方值减小的量。

3.17 多水平模型

3.17 多水平模型 multilevel statistical model

又称“混合效应模型 (mixed effect model)”、“混合模型”。一组处理允许同一层级内部的观测数据间存在相关关系的多层数据的统计模型。

3.17.1 基本概念

3.17.1.1 多水平数据 multilevel data

反应变量的分布在个体间不具备独立性、存在地理距离内或特定空间/时间范围内聚集性的数据。

3.17.1.2 固定效应 fixed effect

处理水平是经过人为有目的地选择、或可以人工控制、可以重复的因素所对应的效应，用于估计这些水平的效应或它们之间的差异。

3.17.1.3 随机效应 random effect

处理水平是从所有可能水平中随机抽取出来的因素所对应的效应，用于了解该因素不同水平的总体变异情况。

3.17.1.4 组群 cluster

在研究暴露因素与某结局的关系时，处于暴露状态及未处于暴露状态这两个不同水平下的被研究对象按一定依据形成的若干组、群或队列。

3.17.1.5 交叉分类 cross classification

将数据按照两个或多个因素进行同时分类或划分的方法。

3.17.1.6 设计矩阵 design matrix

描述各个结局变量与解释变量所对应所有可能效应的矩阵。其中的元素取值为0或1的指示变量，1表示所描述的效应存在，0表示所描述的效应不存在。

3.17.1.7 水平n变异 level n variation

多层次数据结构的资料中，第n个层级单位所呈现的变异。

3.17.1.8 多群组 multiple membership

多层次数据结构的资料中，低层级单位同时归属于同一高层级的多个单位的现象。

3.17.1.9 组内相关系数 intraclass correlation coefficient, ICC

在一组观察值中，两两之间间的相关系数。等于组间方差与总方差(组间方差与组内方差之和)之比。

3.17.1.10 方差分割系数 variance partition coefficient

用于衡量多层次模型或聚类数据中，某一层级的方差占总方差的比例，反映该层次对总体变异的贡献程度。

3.17.1.11 两水平模型 two-level model

一种用于分析具有多层次结构数据的统计模型，适用

于描述数据中存在两个层次或级别的情况，其中个体或观察值被嵌套在高一级的层次单位内。

3.17.1.12 高水平解释变量 higher level explanatory variable

多层次数据结构的资料中，高层级单位所对应的解释变量。

3.17.1.13 成分效应 compositional effect

多水平模型中，高水平解释变量对反应变量的效应。

3.17.1.14 非嵌套模型 non-nested model

两个或多个统计模型之间不存在包含关系，即它们是互相独立的模型，各个模型的模型形式与参数估计结果不会互相影响。

3.17.1.15 单位内相关 intra-unit correlation

同一个单位或群体内个体之间的相关关系，衡量了同一单位内个体之间的相关程度。

3.17.2 边际模型 marginal model

用于描述因变量与自变量之间的关系，同时边际化（亦即控制）其他相关因素的一种统计学方法，常用于多水平模型中自变量对因变量边际效应的估计。

3.17.3 模型估计

3.17.3.1 加权最大似然估计 weighted maximum likelihood estimation

一种参数估计方法，基于似然函数最大化的原则实施参数估计时，为每个被观测个体的似然函数赋以一定权重。

3.17.3.2 加权马尔可夫链蒙特卡罗估计 weighted Markov chain Monte Carlo estimation

一种基于马尔可夫链蒙特卡罗(MCMC)的估计方法，通过在对复杂或未知概率分布抽样中为每个样本分配权重，以更准确地估计目标分布或目标参数。

3.17.3.3 自举法 bootstrap

一种通过有放回地随机重采样从原始样本中生成大量新样本的统计方法，用于估计统计量的分布和置信区间等，尤其适用于小样本或分布未知的情况。

3.17.3.3.1 非参数自举法 nonparametric bootstrap

纯粹依赖于从给定样本中进行重采样、不对样本来源的总体作分布限定的一种自举法。

3.17.3.3.2 参数自举法 parametric bootstrap

从基于数据拟合得到的分布函数中进行重采样的一种自举法。

3.17.3.3.3 残差自举法 residual bootstrap

通过有放回的重抽样依次创建多个残差数据集，据以

对描述残差分布的参数进行估计的一种自举法。

3.17.4 方差成分模型 variance component model

又称“方差分量模型”。用于分析数据集中方差来源的统计模型，模型将数据中的总方差分解为不同的分量，通过估计和对比这些方差分量，可以了解不同水平或群体对数据总体变异的相对贡献。

3.17.4.1 中心化处理 grand mean centering

将数据中的每个值减去该数据的中心值（如算术均数、中位数），以得到中心化后的数据的一种处理方法。

3.17.5 随机截距模型 random intercept model

混合效应模型的一种，在模型中引入随机截距，允许不同组别的截距存在差异，但假设自变量的效应在各组间保持一致，用于分析具有层次结构的数据。

3.17.5.1 协方差矩阵 covariance matrix

一个方阵，主对角线上的元素为各个变量的方差，主对角线外的元素为随机变量两两之间的协方差。用于描述多个随机变量之间的相关性。

3.17.5.1.1 无约束因子协方差矩阵 unconstrained factor covariance matrix

在因子分析中，放弃潜在因子之间正交性的约束而得到的因子协方差矩阵。该矩阵中的元素乃潜在因子之间的协方差，反映因子之间的相关性。

3.17.5.1.2 反应变量协方差矩阵 covariance matrix of response variable

用于描述多个反应变量之间的协方差关系的矩阵。它可以通过基于观测数据进行估计，显示反应变量之间的关联关系。

3.17.5.1.3 随机截距协方差矩阵 covariance matrix of random intercept

描述多层次数据模型中的随机截距之间关系的协方差矩阵，揭示不同组之间截距的相关性和变异程度。

3.17.5.2 协方差结构 covariance structure

在多元统计分析中，变量之间关联关系的理论模式，可包括线性和非线性形式，也蕴含变量间的相关性和独立性等信息。

3.17.6 随机斜率模型 random slope model

又称“随机系数模型”。针对层次结构数据，各自变量的效应在不同组间随机变化，揭示各组间自变量与因变量关系的异质性。

3.17.6.1 随机斜率协方差矩阵 covariance matrix of random coefficient

描述多层次数据模型中的随机斜率之间关系的协方差矩阵，揭示不同组之间斜率的相关性和变异程度。

3.17.7 非线性多水平模型 nonlinear multilevel model

因变量与部分或所有自变量之间呈非线性关系且相应参数未知的一类多水平模型

3.17.7.1 非线性模型 nonlinear model

因变量与部分或所有自变量之间呈非线性关系且相应参数未知的模型，可通过逐次逼近的方法进行模型拟合，常使用的函数有对数函数、三角函数、指数函数、幂函数等。

3.17.8 多水平广义线性模型 multilevel generalized linear model

反应变量为来自除高斯分布以外的不同分布，在多水平框架内处理这类资料的统计模型。

3.17.8.1 多水平泊松模型 multilevel Poisson model

用于具有层级结构的计数数据的建模方法。模型结合了泊松回归与多层次结构，允许在不同层次上考虑随机效应。

3.17.8.2 多水平逻辑斯谛模型 multilevel logistic model

用于具有多层次结构的二分类或多分类计数数据分析建模的方法。模型的连接函数为逻辑斯蒂函数，并考虑到不同层次之间的变异允许解释变量所起的作用随着层次的变化而变化。

3.17.9 多水平考克斯模型 multilevel Cox model

又称“多水平 Cox 模型”。考克斯比例风险模型的扩展，在模型中引入随机效应，以捕捉数据中不同层次的变异，同时评估协变量对事件发生风险的影响，用于分析具有层次结构的生存数据。

3.17.9.1 对数持续时间模型 log duration model

生存分析中对数尺度上对时间变量进行建模的一类参数模型，资料呈现或不呈现层级结构时均适用。

3.17.9.1.1 时变协变量 time-varying covariate

在随访期间随时间变化而变化，且在预测模型中角色为预测因子的变量。

3.17.10 多水平荟萃分析 multilevel meta-analysis

将具有协变量的多水平模型应用于荟萃分析的统计方法。

3.17.10.1 异质方差 heterogeneous variance

在被纳入的各个研究中，不同研究的效应大小之间的变异程度

3.17.10.2 层次结构 hierarchical structure

数据中存在的不同等级或层次的组织方式，这些层次通常反映了数据的自然分组或嵌套关系。同一层次内的数据可能存在一定相似性。

3.17.10.3 效应尺度 effect scale

用来衡量处理效应的大小或因变量与自变量关联强度的指标。

3.17.11 多水平主成分分析 multilevel principal component analysis

用于研究大量维度/特征的大型数据集的主成分分析

方法，在保留最大信息量的同时提高数据的可解释性，实现多维数据的可视化，提升高维函数数据的信息挖掘效果。

3.17.12 多水平因子分析 multilevel factor analysis

用于研究不独立的层次数据或重复测量数据潜在结构的因子分析方法。

3.17.13 多水平时间序列分析 multilevel time series analysis

一种用于分析具有层次结构的单一变量的时间序列数据的分析方法。

3.17.13.1 多元时间序列分析 multivariate time series analysis

一种用于分析包含同一层级的多个变量的时间序列数据的分析方法。

3.17.14 多元多水平模型 multivariate multilevel models

处理允许同一水平内部的向量形式的观测数据间存在相关关系的多层数据的统计模型。

3.17.14.1 旋转设计 rotation design

具备在因子空间内与中心点等距的同一球面上各点的预测值方差都相等的统计性质的回归设计。

3.17.15 多水平结构方程模型 multilevel structural

equation model

结构方程模型的一种拓展，适用于在具有层次结构的数据中，基于潜变量的协方差矩阵来分析变量之间关系。

3.17.15.1 恰好识别 just-identified

模型参数的数量严格等于数据可提供的独立信息量，具有唯一解和不可检验的特性。

3.17.15.2 识别不足 under-identified

模型参数数量超过数据可提供的独立信息量，具有无唯一解和需修正的特性。

3.17.15.3 过度识别 over-identified

模型参数数量少于数据可提供的独立信息量，具有可检验性和最优解的特性。

3.17.15.4 调整拟合优度指数 adjusted goodness fit index

一个用于评估结构方程模型拟合优度的统计指标。它是拟合优度指数经过调整后的值，用于惩罚模型中自由参数数量过多的情况。取值范围在 0 到 1 之间，值越接近 1 表示模型的拟合优度越好。

3.17.15.5 迭代 iteration

参数估计过程中逐步逼近最优解的数值计算步骤。

3.18 广义估计方程

3.18 广义估计方程 generalized estimating equation

在广义线性模型框架下通过考虑观测数据之间的相关性结构对非独立数据进行分析的统计方法，适用于分析纵向数据或重复测量数据。

3.18.1 基本概念

3.18.1.1 纵向数据 longitudinal data

在一段时间内或随特定事件发生时，对某些感兴趣的个体（同一研究对象）进行多次观察或重复测量所获得的数据。

3.18.1.2 作业相关矩阵 working correlation matrix

表示某一个体或区组内相关个体中因变量的各次重复测量值两两之间相关性矩阵。

3.18.1.2.1 等相关 exchangeable correlation

又称“可交换的相关”。任意两次观测之间的相关系数相等，且不随时间推移而改变。

3.18.1.2.2 相邻相关 stationary correlation

又称“稳态相关”。只有相邻的两次观测值间有相关。例如，在重复测量资料中，相关矩阵服从某一稳定结构，等相关、自相关等均出现稳定相关的特殊情况。

3.18.1.2.3 自相关 autoregressive correlation

同一变量在不同时间点上的观测之间的相关关系。相

关与时间间隔有关，间隔越长，相关系数随之衰减。

3.18.1.2.4 非结构化相关 unstructured correlation

在处理重复测量或多次测量数据时的一种相关性假设，认为不同测量之间的相关性缺乏可预测模式。该假设下能够捕捉任意相关模式，但模型复杂度高，可能导致不收敛。

3.18.1.2.5 固定相关 fixed correlation

不同时间点或观测之间存在稳定的、固定的相关性结构；亦可解释为根据专业知识提前指定的相关结构。

3.18.1.2.6 独立结构 independent

任意两次观测值之间无相关性，即相关性矩阵为单位矩阵。

3.18.1.4 离散参数 dispersion parameter

又称“扩散参数”。指描述观测数据相关性结构的参数，影响参数估计的标准误。

3.18.2 模型估计

3.18.2.1 分布簇 distribution family

又称“分布族”。一类具有相似性质的概率分布族群。

3.18.2.1.1 逆正态分布 inverse normal distribution

又称“逆高斯分布”(inverse Gaussian distribution)。一种单峰右偏的分布，可用于描述布朗运动现象中达

到单位距离所需时间的分布。

3.18.2.1.2 恒等连接函数 identity link function

将线性预测值映射到响应变量的函数。

3.18.2.1.3 对数连接函数 log link function

将线性预测值映射到响应变量的对数值的函数。

3.18.2.1.4 对数单位连接函数 logit link function

又称“logit 连接函数”。将线性预测值映射到因变量的 logit() 函数值。

3.18.2.1.5 概率单位连接函数 probit link function

又称“probit 连接函数”。将线性预测值映射到响应变量的 probit() 函数值。

3.18.2.1.6 重对数连接函数 log-log link function

将线性预测值映射到响应变量的 log(-log()) 函数值

3.18.2.1.7 幂连接函数 power link function

将线性预测值映射到响应变量的幂函数值

3.18.2.1.8 优势幂连接函数 odds power link function

将线性预测值映射到响应变量的优势(odds)的幂函数值。

3.18.2.1.9 负二项连接函数 negative binomial link function

将线性预测值映射到响应变量的负二项函数值

3.18.2.1.10 互反连接函数 reciprocal link function

将线性预测值映射到响应变量的倒数。

3.18.2.2 方差成分估计 variance component estimation

将观测数据总方差分解为不同来源或组分。

3.18.2.3.1 全信息极大似然 full information maximum likelihood

利用全部数据信息估计模型参数，尤其适用于存在缺失数据的情形。

3.18.2.3.2 有限信息极大似然 limited information maximum likelihood

利用模型中的部分信息估计模型参数。

3.18.2.3.3 有限信息极大似然 limited information maximum quasi-likelihood

利用模型中的部分信息，通过对似然方程的修正，估计模型参数。

3.18.2.4 效应估计 effect estimation

估计由某个自变量引起因变量的变化程度。

3.18.2.4.1 群体平均效应 population average effect

总体中所有个体的平均效应，而不特指某一个体的效应。

3.18.2.4.2 特定个体效应 subject specific effect

某一个体与总体或其他所有个体相比所表现出的独特效应。

3.18.3 模型评价

3.18.3.1 准似然独立准则 quasi-likelihood under the independence model criterion

又称“拟似然独立准则”。基于准似然估计的模型选择准则，通过同时考虑模型的拟合优度和复杂度来评价模型。

3.19 遗传统计

3.19 遗传统计 genetic statistics

用于研究生物群体性状遗传规律的数理统计和数学方法

3.19.1 群体遗传研究 population study

以群体为单位对遗传因素及性状表型之间的规律进行的一类研究，一般情况下假设所研究的个体不存在亲缘关系。

3.19.1.1 群体结构 population structure

群体内部各亚种群间由于个体间的非随机交配而产生的等位基因频率的系统差异，可反映祖先的遗传信息。

3.19.1.1.1 哈迪-温伯格平衡 Hardy-Weinberg equilibrium

在一个没有突变、选择和迁移的遗传漂变的无限大的随机交配群体中，一对等位基因在常染色体上遗传时，无论群体起始基因频率如何，只要经过一代的随机交配，群体的基因型频率和基因频率即达到平衡状态。

3.19.1.2 病例对照关联分析 case-control association

study

一种关联分析方法，选择一组患有某疾病的患者与一组未患此病的个体，调查他们发病前对某因素的暴露情况，比较两组中暴露率和暴露水平的差异，以研究该疾病与该因素关系。

3.19.1.2.1 外显率 penetrance

一个群体中一定基因型的个体在特定环境中显示预期表型的百分比。

3.19.1.3 全基因组关联研究 genome-wide association study, GWAS

通过对大规模的群体 DNA 样本进行全基因组高密度遗传标记分型，在全基因组层面上识别与复杂性状或者疾病相关的遗传位点的研究方法。常用遗传标记为单核苷酸多态性，即 SNP。

3.19.1.3.1 广义遗传度 broad-sense heritability

在特定群体中，某一性状由遗传解释的变异在全部表型变异中所占的比例。

3.19.1.3.2 狭义遗传度 narrow-sense heritability

在特定群体中，某一性状由遗传加性效应解释的变异在全部表型变异中所占的比例。

3.19.1.3.3 基因组膨胀因子 genomic inflation factor
实际观察的检验统计分布中位数与预期中位数之比。预期的基因组膨胀因子为1，当基因组膨胀因子大于1时，提示可能需要校正群体分层的影响。

3.19.1.3.4 汇总统计量 summary statistics
描述关联分析结果和样本特征的数值。一般包括全基因组关联分析的效应值、效应值的标准误、*P*值、样本量等。

3.19.1.3.5 遗传相关 genetic correlation
群体中不同表型，比如身高、血压等，在剔除环境影响之后，由于遗传因素导致的相关性。

3.19.1.3.6 多效性位点 pleiotropic locus
对多种性状产生影响的遗传位点。

3.19.1.3.7 精细化定位 fine-mapping
在基因组候选位点进一步缩小因果变异潜在区域的定位方法。

3.19.1.3.8 多基因遗传风险评分 polygenic risk score, PRS
由多个遗传位点的基因型和及其相应权重计算得到的分值。可以用于评价遗传因素导致的风险。

3.19.1.4 罕见变异分析 rare variant analysis
针对罕见遗传变异（通常指次等位基因频率小于0.01的遗传变异）进行的一类分析。

3.19.1.4.1 常见变异 common variant
在一个特定群体中，次等位基因频率比较大的遗传变异，通常情况为次等位基因频率大于0.01的遗传变异

3.19.1.4.2 罕见变异 rare variant
在一个特定群体中，次等位基因频率比较小的遗传变异，通常情况为次等位基因频率小于0.01的遗传变异

3.19.1.4.3 全基因组测序 whole genome sequencing, WGS
利用高通量测序平台对一种生物的基因组中的全部基因进行测序，测定其DNA碱基序列的方法

3.19.1.4.4 全外显子测序 whole exome sequencing, WES
利用序列捕获技术将全基因组中所有外显子区域DNA序列捕获，富集后进行高通量测序的方法。

3.19.1.4.5 基因型填充 genotype imputation
以单体型图谱为参考，对个体缺失基因型进行推断的一种方法。

3.19.1.4.6 负荷检验 burden test
假设基因组某个区域内的一组罕见变异对疾病的效应具有方向一致性，从而将其合并成一个新的变量，然后通过分析新变量与疾病之间的关系研究罕见变

异与疾病之间的关联性的方法

3.19.1.4.7 方差成分检验 variance component test
假设基因组某个区域内的一组罕见变异与疾病的效应服从正态分布，从而通过检验随机效应的方差成分研究罕见变异与疾病之间关联性的方法

3.19.1.5 全转录组关联研究 transcriptome-wide association study, TWAS
通过对大规模的群体全部基因表达量进行测量或者预测，在全转录组层面上识别与复杂性状或者疾病相关的基因的研究方法

3.19.1.6 全表观基因组关联研究 epigenome-wide association study, EWAS
通过对大规模的群体DNA样本进行全基因组高密度表观遗传标记测量，在全基因组层面上识别与复杂性状或者疾病相关的表观遗传修饰位点的研究方法。常用表观遗传标记为DNA甲基化。

3.19.2 家系研究 family-based study
以家庭为单位研究遗传因素及性状表型之间规律的一种研究。一般情况下研究个体存在亲缘关系。

3.19.2.1 亲缘系数 coefficient of relationship, kinship coefficient
两个个体的基因组中相同且同源基因的比例。可以用于度量两个个体间亲缘关系的一个指标。

3.19.2.1.1 血缘一致性 identity by descent, IBD
一种一致性现象，指两个个体拥有的相同的等位基因或者片段来源于同一祖先。

3.19.2.1.2 状态一致性 identity by state, IBS
一种一致性现象，指两个个体拥有相同的等位基因或者基因片段。

3.19.2.1.3 遗传关系矩阵 genetic relationship matrix, GRM
由个体间的遗传相关系数构成的矩阵。个体之间的遗传关系一般可以通过基因型或者系谱进行推断。

3.19.2.2 遗传连锁分析 genetic linkage analysis
一种基于遗传连锁原理，通过分析家庭成员之间的遗传关系，确定与某一性状或疾病相关的基因位置的遗传学研究方法

3.19.2.2.1 连锁不平衡 linkage disequilibrium
不同基因座上的等位基因同时遗传的频率明显高于随机组合的预期值的现象。

3.19.2.2.2 连锁分析 linkage analysis, LD
一种基于遗传连锁原理，通过分析家庭成员之间的遗传关系，确定与某一性状或疾病相关的基因位置的遗传学研究方法

3.19.2.2.3 对数优势记分法 log odds score, LOD
利用两个基因座不连锁的样本观察值的最大似然函

数与两个基因座彼此连锁的样本观察值的最大似然函数之比检验两个基因座间是否彼此连锁的方法。

3.19.2.2.4 受累同胞对 affected sib-pair, ASP

通过比较单基因遗传病中同胞对的观测值和期望值的分布,分析标记位点与疾病连锁关系的方法。

3.19.2.3 家系关联分析 family-based association study, FBAT

利用具有亲缘关系的研究群体,研究基因与性状之间关联性的方法。

3.19.2.3.1 传递不平衡检验 transmission

disequilibrium test, TDT

一种基于家系的基因关联分析方法,通过观察基因从亲代到子代的传递来分析基因与性状的遗传连锁关系。

3.19.2.3.2 家系不平衡检验 pedigree disequilibrium test, PDT

一种基于家系的基因关联分析方法,一般基于相关核心家庭和扩展谱系同胞的数据来分析基因与性状的遗传连锁关系。

3.20 贝叶斯统计

3.20 贝叶斯统计 Bayesian statistics

英国科学家托马斯·贝叶斯(Thomas Bayes)提出的一种归纳推理统计推断方法。当存在不确定性时应如何进行推理,用概率语言描述不确定性,根据已有的证据去推断事件发生的概率可以实现用先验的知识去计算后验的概率。

3.20.1 贝叶斯定理 Bayes' theorem

又称“贝叶斯公式(Bayes'formula)”、“贝叶斯法则(Bayes' rule)”。以托马斯·贝叶斯(Thomas Bayes)命名的基于可能与事件相关的条件的先验知识,描述事件发生概率的一则定理。

3.20.1.1 先验概率 prior probability

是在观测到任何新证据之前,根据先前的经验、专业知识、或其他信息得到的关于某个事件是否发生的主观概率。

3.20.1.2 后验概率 posterior probability

通过调查或其他方式获取新的附加信息,利用贝叶斯公式对先验概率进行修正,而后得到的概率。

3.20.1.3 条件概率 conditional probability

一个事件在另一个前置事件已经发生条件下的发生概率,由前置事件发生概率和该事件后续发生概率相乘得到。

3.20.1.4 全概率公式 total probability formula

描述了在考虑多个不相交事件的情境下,一个特定事件发生的总概率,通过将所有相关事件的条件概率加权求和得到。是概率论的重要公式之一。

3.20.1.4.1 完备事件组 collectively exhaustive events

一组相互独立且互斥的事件,这组事件在一起涵盖了所有可能的结果,确保其中的任意一个事件发生的总概率为1。即,总有一个事件在该组中发生,而不会有两个或更多的事件同时发生。

3.20.1.5 归一化常数 normalising constant

一个使非正规化的概率或概率密度函数的总和或积

分等于1的常数。它确保概率分布函数或概率密度函数满足其相应的总和或积分的要求。

3.20.1.6 离散型贝叶斯公式 Bayes' theorem for discrete random variables

概率论和统计学中计算离散型随机变量发生概率的一种方法,用于计算某一变量在已知另一变量的情况下发生的概率,该离散型变量常为某事件发生概率

3.20.1.7 连续型贝叶斯公式 Bayes' theorem for continuous random variables

概率论和统计学中计算连续型随机变量概率密度的一种方法,用于计算某参数在已知样本的情况下的概率密度即后验密度,由参数先验概率和样本似然函数相乘得到

3.20.2 贝叶斯推断 Bayesian inference

一种统计推断方法。使用当前所观察到的样本信息和先验知识,更新或重新推断后验概率。

3.20.2.1 先验信息 prior information

在决策分析中,尚未通过试验收集状态信息时所具有的信息。

3.20.2.2 先验分布 prior distribution

在观察数据之前对一个或多个随机变量的不确定性或变化性的描述。它表示我们在观测到新数据之前关于某一参数或变量的知识或信念。可以基于之前的研究、专家知识或其他信息来源。

3.20.2.2.1 无信息先验 non-informative prior

一种尽量不对参数或变量的可能值做出任何假设或预设的先验分布。这种先验的目标是尽量避免对分析结果施加太多的主观偏见,从而让数据本身在贝叶斯分析中产生更大的影响。

3.20.2.2.2 有信息先验 informative prior

一种反映明确或特定知识的先验分布。它基于先前研究、经验、专家意见或其他信息源。与非信息先验相比,为参数或变量设定了明确的信念或约束,从而在

分析中起到引导作用。

3.20.2.2.3 部分信息先验 weakly informative prior

又称“弱信息先验”。一种基于已有知识和部分信息的先验分布。它为参数提供了某种基本的、广泛的指导或限制，但并没有强烈地偏向于任何特定值。这种先验有助于提供模型的稳定性，同时仍然允许数据对参数估计产生主导作用。

3.20.2.2.3.1 单位信息先验 unit information prior

一种在模型参数估计中赋予先验知识与单一数据样本相同的权重的先验分布，包含了与单一数据点等量的信息。它旨在平衡先验和似然函数的贡献，使得先验的影响等同于一个观测值。

3.20.2.2.4 共轭先验 conjugate prior

一种当先验分布与似然函数结合后，所得到的后验分布与先验分布属于同一家族或形式的先验分布。因后验分布的形式与先验分布相同，可避免在更新估计时的复杂计算。

3.20.2.2.5 非共轭先验 non-conjugate prior

与似然函数结合后不会产生同类的后验分布的先验分布。可能需要更复杂的方法（如 MCMC）来得到后验分布。简言之，它不保证与特定的似然函数结合时产生易于处理的后验分布。

3.20.2.2.6 半共轭先验 semiconjugate prior

一种先验分布，它与似然函数结合后，得到的后验分布与原先验形式相似，但可能涉及更多的参数。只部分地保留了形式，从而为建模提供了更大的灵活性，同时仍然简化了计算。

3.20.2.2.7 层次先验 hierarchical prior

参数本身有其先验分布的情况下，该先验分布的参数又有另外的先验分布的结构，允许在多个层次上共享信息。适用于处理复杂的数据结构或组间变异，可以考虑到不同的信息源，并更好地估计模型参数。

3.20.2.2.7.1 层次贝叶斯分析 hierarchical Bayesian

analysis

一种基于贝叶斯框架，通过结构化地考虑多层参数关系，使得一个层级的参数估计能够受到另一个层级数据的影响，允许共享信息的分析方法，通常分为“全局”或“总体”参数和“局部”或“子群”参数。适用于存在群体差异或结构化数据。

3.20.2.2.7.2 超先验分布 hyperprior distribution

在贝叶斯统计中，用于描述先验分布参数的先验分布。当为模型参数设定先验分布时，如该先验分布本身还有未知参数，可以为这些未知参数设置的先验分布。

3.20.2.3 后验分布 posterior distribution

根据样本信息和参数的先验分布，通过抽样并根据条件概率分布的计算原理更新已知条件下参数的条件

分布。

3.20.2.3.1 贝叶斯可信区间 Bayesian credible interval

简称“可信区间”，又称“贝叶斯置信区间”。利用参数后验分布获得的参数区间。90%的可信区间在样本给定后，可通过后验分布的分位数求得，总体参数落入可信区间的概率为 0.9。贝叶斯学派为了与频率学派置信区间相区别而提出的一个概念。

3.20.2.3.2 最高后验密度 highest posterior density, HPD

后验分布中具有最大概率密度的点或值，代表了对未知参数的最可能估计。

3.20.2.4 边际似然 marginal likelihood

又称“边缘似然”。由先验加权的似然在整个参数空间的平均值，代表了模型对数据集的平均拟合度，是贝叶斯框架中模型比较的常用方法。

3.20.2.5 后验近似 posterior approximation

结合先验信息，根据观测结果得出的所有未知变量的后验概率的近似值。

3.20.2.5.1 马尔可夫链蒙特卡罗 Markov chain Monte Carlo, MCMC

构造一个随机模型来描述一系列发生概率只受到前一事件状态影响的可能事件，使其平稳分布为待估参数的后验分布并产生样本，基于该样本通过重复的模拟和抽样进行参数估计的统计方法

3.20.2.5.1.1 马尔可夫链蒙特卡罗诊断 Markov chain Monte Carlo diagnostics

用于评估马尔可夫链蒙特卡罗方法的收敛性和性能的一系列统计工具和指标。

3.20.2.5.1.2 老化 burn-in

在使用马尔可夫链蒙特卡罗(MCMC)等采样方法时，初始阶段生成的样本被丢弃的过程，以避免初始偏差对最终结果的影响。

3.20.2.5.1.3 收敛 convergence

某个统计量、模型参数或目标函数值，随着数据量、迭代次数或训练时间的增加，逐渐趋近于某个稳定值或理论值的过程。

3.20.2.5.2 吉布斯抽样 Gibbs sampling

简称“Gibbs 抽样”。马尔可夫蒙特卡罗迭代的一种算法，用于构建一系列相互依赖的参数值，这些参数值的分布逐渐收敛到目标联合后验分布。

3.20.2.5.2.1 吉布斯采样器 Gibbs sampler

简称“Gibbs 采样器”。统计学中用于马尔可夫链蒙特卡罗(MCMC)的一种算法，用于在难以直接采样时从某一多变量概率分布中近似抽取样本序列。

3.20.2.5.3 梅特罗波利斯-黑斯廷斯算法

Metropolis-Hastings algorithm

又称“Metropolis-Hastings 算法”。一种用于从分布函数未知的概率分布中抽样的马尔科夫链蒙特卡洛方法。通过构建一个马尔科夫链，使得该链的平稳分布与所需的概率分布相对应，然后利用该马尔科夫链生成的样本进行统计推断和模型估计。

3.20.2.5.4 积分嵌套拉普拉斯近似 integrated nested Laplace approximation, INLA

用于快速获得后验边缘分布精确近似的贝叶斯推断方法。

3.20.2.6 贝叶斯估计 Bayes estimation

结合关于待估计参数的先验知识和数据的似然函数来估算参数的贝叶斯方法。

3.20.2.7 贝叶斯因子 Bayes factor

一种用于比较两个或多个统计模型相对拟合优度的指标。它是将数据观测到的概率与模型参数的先验概率相乘，再对所有可能的参数值进行积分求和所得。

3.20.2.8 贝叶斯方法 Bayesian approach

应用所观察到的现象对有关概率分布的主观判断(即先验概率)进行修正的标准方法。

3.20.3 贝叶斯分类器 Bayesian classifier

基于贝叶斯公式，通过计算最大后验概率的方法来预测待分类数据类别的分类方法。

3.20.3.1 贝叶斯判定准则 Bayes decision rule

一种决策规则，基于贝叶斯定理和概率分布，选择使后验期望损失最小的决策。

3.20.3.2 属性条件独立 attribute conditional independence

在给定其他属性的条件下，两个属性之间互相独立。可用于属性间依赖关系的推断，有助于理解数据关联。

3.20.3.3 朴素贝叶斯分类器 naive Bayes classifier

假定给定类别标记时特征相互之间独立的贝叶斯分类器。

3.20.3.4 半朴素贝叶斯分类器 semi-naive Bayes classifier

放宽各个属性间属性独立性假设，适当考虑一部分属性间的相互依赖关系，放松后的朴素贝叶斯分类器。

3.20.3.5 贝叶斯网络 Bayesian network

一种用有向无环图描述随机变量及其条件依赖关系的概率图模型。

3.21 缺失数据

3.21 缺失数据 missing data

又称“缺失值(missing values)”。一组数据中由于某种原因而缺失掉的数据。通常分为完全随机缺失、随机缺失和非随机缺失三种类型。

3.21.1 缺失数据模式

3.21.1.1 单调缺失模式 monotone missing data pattern

在包含多个观测变量的数据集中，缺失值的出现遵循一种特定的顺序特征，如果个体在某个变量上的值缺失，则该个体在所有后续变量上的值也必将缺失。

3.21.1.2 任意缺失模式 arbitrary missing data pattern

对含有多个观测变量的数据各种行交换或/和列交换均不能使缺失数据呈现层级缺失的模式。即数据缺失具有偶然性，呈现没有规律可循的缺失模式。

3.21.2 缺失数据机制

3.21.2.1 完全随机缺失 missing completely at random, MCAR

数据的缺失与已观测到和未观测到的数据均无关，所有数据缺失概率都是相同的。

3.21.2.2 随机缺失 missing at random, MAR

在一定条件下，缺失值与已知其他数据或变量有关，而与自身取值大小无关。

3.21.2.3 非随机缺失 missing not at random, MNAR

缺失值不仅与其他变量的取值有关，也与自身变量的

取值有关，此情况不可忽略缺失数据带来的影响。

3.21.3 缺失数据处理

3.21.3.1 删除法 deletion

直接对未观测的数据值进行排除的分析方法。

3.21.3.1.1 列表删除 listwise deletion

又称“个案剔除法(case deletion method)”、“完全案例分析法”。一种删除数据集中任何包含缺失值的个案，仅保留所有变量都完整的个案的分析方法。

3.21.3.1.2 成对删除 pairwise deletion

又称“有效案例分析法”。一种根据要做的检验所涉及的变量进行个案剔除的分析方法。比如做两个变量的相关时，仅删除在这两个变量上有缺失的个案，可最大限度地保留了数据集中的可用信息。

3.21.3.2 加权校正法 weighting adjustment

当出现缺失单元时，通过估算单元的应答率，估计缺失产生的影响，并计算出相应权重，把缺失单元的权数分解到非缺失单元上，通过增大样本中观测数据的权数，以减小由于缺失值可能对估计量带来偏差的方法。

3.21.3.3 期望值最大化法 expectation maximization, EM

一种在不完全数据情况下计算极大似然估计或者后验分布的迭代算法。数据集满足大样本、多元正态分

布条件，呈任意缺失模式时，通过缺失值和模型参数间的迭代关系，在假定模型参数基础上进行缺失估计，再利用缺失估计值修正模型参数，迭代至模型收敛。

3.21.3.4 填补法 imputation

又称“插补法”。基于某种已知数据分布特征对（未观测值）缺失值进行弥补的方法。

3.21.3.4.1 单一填补 single imputation, SI

又称“单一插补”。一种处理缺失数据的方法。它为数据集中的每个缺失值仅生成一个填补值，形成完整数据集后进行分析。

3.21.3.4.1.1 均值填补 mean imputation

又称“均值插补”。一种直接将缺失值填补为平均值的方法。它的假设是缺失的数据在期望上与总体相等，常被应用于随机缺失和完全随机缺失的情况中。

3.21.3.4.1.2 回归填补 regression imputation

又称“回归插补”。一种利用已知变量信息，依据变量间的关系建立回归模型并通过回归模型的预测值作为填补值的方法。

3.21.3.4.1.3 热卡填补 hot-deck imputation, HI

又称“热卡插补”。一种基于相似性匹配的单一填补方法。其核心思想是从同一数据集中选择与缺失个案相似的完整个案的值进行填补。

3.21.3.4.1.3.1 随机热卡填补 random hot-deck imputation

又称“随机热卡插补”。一种基于相似性填补缺失值的插补方法，对于一个包含缺失值的数据集，在完整数据中通过特定规则找到一组与缺失值最相似的观测值，然后从这组观测值中随机选择一个来进行填充。

3.21.3.4.1.3.2 最近邻填补 nearest neighbor imputation

又称“最近邻插补”。利用辅助变量定义测量间距离度量，以距离含缺失值的数据集观测距离最近的某一变量值作为插补值的方法。

3.21.3.4.1.3.2.1 最近距离填补 nearest distance imputation

又称“最近距离插补”。利用辅助变量，定义一个已观测样本之间的距离函数，在与缺失值临近的非缺失样本中，选择满足所设定距离条件的辅助变量，将这个辅助变量中的非缺失样本所对应的目标变量上的非缺失值作为填补值的一种方法。

3.21.3.4.1.3.2.2 k-最近邻填补 k-nearest neighbor imputation, KNN

如果数据的某一变量上有缺失，找出离缺失值最近的若干个数据，用这若干数据的平均值来填充缺失值，每个样本的缺失值使用数据集中找到的这若干数据邻域的平均值进行插补。

3.21.3.4.1.3.2.3 象限近邻填补

quadrant-encapsidated-nearest-neighbor-based imputation

又称“象限近邻插补”。一种通过找出缺失值周围一定范围内的所有最近数据点对其加权来弥补缺失值的方法。

3.21.3.4.1.3.2.4 序贯热卡填补 sequential hot-deck imputation

又称“序贯热卡插补”。对样本分层后再每层中按照某些变量对样本进行排序，对于有数据缺失的单元，用同一层中其前后相邻的若干个数据，找到使某距离函数的值达到最小的单元，该单元所对应的目标变量上的非缺失数值的方法。

3.21.3.4.1.4 冷卡填补 cold-deck imputation

又称“冷卡插补”。一种基于与缺失值样本相似的其他研究中样本的观测值来填补缺失数据的方法。

3.21.3.4.1.5 观测值结转 observation carried forward

将先前观察到的数值、指标或特征延续到未来的时间段或不同的条件下进行分析、预测或比较的方法。

3.21.3.4.1.5.1 末次观测值结转 last observation carried forward, LOCF

又称“末次访视观测值向前结转”、“末次观测结转”。一种基于最后观测值填补缺失数据的方法。当研究的某个参与者或对象在某个时间点缺失了数据时，将该参与者最后一次观察的数据应用于缺失的时间点，作为缺失值的替代。常在队列研究或临床试验中使用。

3.21.3.4.1.5.2 基线观测值结转 baseline observation carried forward, BOCF

将基线观察应答视作研究终点，是一种非随机缺失数据处理方法。

3.21.3.4.1.5.3 最差观测值结转 worst observation carried forward, WOCF

又称“最差观测值向前结转”、“最差观测结转”。采用该样本既往观测数据中，具有某种特定意义上最差涵义（如临床上测血压最大值等）的观测值替代缺失值的方法。

3.21.3.4.2 多重填补 multiple imputation, MI

又称“多重插补”。一种通过生成多个数据集来提高缺失值估计准确性的方法。对每个缺失值用若干个可能值进行填补，再对填补后的若干个“完整”数据集分别进行分析，综合所有分析的结果，从而得到该缺失值的估计值。

3.21.3.4.2.1 联合模型策略 joint modeling, JM

在假定模型参数服从多元概率密度分布的情况下，从贝叶斯联合后验分布函数中抽取来做填补值的方法。

3.21.3.4.2.2 完全条件定义法策略 full conditional specification

又称“链式方程多变量插补”、“逐步回归多变量插补”。对多变量任意缺失，且数据不满足多元正态分布或其

他联合分布时，利用条件概率分布原理，基于单变量回归模型对每个变量依次弥补的方法。

3.21.3.4.2.3 马尔可夫链蒙特卡洛法 Markov Chain Monte Carlo method, MCMC

如果缺失值为连续型，通过填补和后验抽样两步循环交替进行缺失值填补的方法。

3.21.3.4.2.4 回归分析法 regression analysis method

一种以无缺失值且与有缺失值的变量相关的变量作为自变量，建立适当的线性回归模型，对缺失值进行填补的方法。

3.21.3.4.2.5 判别函数法 discriminant function method

一种以没有缺失值且与缺失值相关的变量作为自变量，建立适当的判别函数模型，根据得到的模型对缺失值进行填补的方法。

3.21.3.4.2.6 预测均值匹配法 predictive mean match, PMM

一种通过建立有缺失值的变量与完整变量之间的回归模型，得到回归系数的参数估计后，从回归系数估计值的后验分布中随机抽取多个新参数计算预测值再对缺失值进行填补的方法。

3.22 量表评价

3.22 量表评价 scale evaluation

由反映某一抽象概念的测量变量构成的问卷作为标准化工具的测量效果的评价。

3.22.1 基本概念

3.22.1.1 条目 item

为了解被测者的状况，从不同的侧面提出的与测评特征有关的问题或进行的测量。

3.22.1.2 条目池 item pool

由量表编制的议题小组成员分别独立地根据个人理解和经验拟订的与测评特征概念有关的建议条目的汇总。

3.22.1.3 离散趋势法 tendency of dispersion

采用标准差或变异系数表示条目的变异程度，反映条目对被测对象间差异的区别能力，是从敏感性角度筛选条目的统计方法。

3.22.1.4 克龙巴赫 α 系数 Cronbach's alpha coefficient

又称“Cronbach's α 系数”。组成量表的条目得分的方差在量表总分方差中所占的比例，是从内部一致性的角度评定条目适用性的统计方法。

3.22.2 反应度 responsiveness

为了解被测者的状况，从不同的侧面提出的与测评特征有关的问题或进行的测量。

3.22.2.1 内部反应度 internal responsiveness

量表在特定时间段内测量出研究对象自身变化的能力。它关注的是量表在同一群体内对变化的敏感度，通常用于评估干预或治疗前后的变化。

3.22.2.2 外部反应度 external responsiveness

量表测量不同时间目标特征的变化与其相应健康状况参考指标变化之间关系的能力，常用相关系数衡量。

3.22.2.3 标准化反应均数 standardized response mean, SRM

干预前后量表得分差值的均值与干预前标准差的比值，以反映量表测量干预作用下目标特征变化的能力。

3.22.3 信度 reliability

又称“可靠性”。量表测量过程中由随机误差造成的测量值的变异程度，评价量表的稳定性、精确性和一致性。

3.22.3.1 内部一致性信度 internal consistency reliability

量表中题目之间的一致性程度，反映的是量表各题目是否在测量同一特质。它是评估量表信度的重要指标之一。

3.22.3.2 重测信度 test-retest reliability

同一量表对同一群体在不同时间点进行重复测量时，所得结果的一致性程度。它用于评估量表的稳定性和可靠性，反映量表在不同时间点测量同一特质时的稳定性。

3.22.3.3 复本信度 alternate-form reliability

使用两个平行或等价的量表版本对同一群体进行测量时，所得结果的一致性程度。它用于评估量表在不同版本之间的一致性，反映量表在测量同一特质时的稳定性。

3.22.3.4 分半信度 split-half reliability

将相同量表分为对等的两部分，如按照序号的奇偶数分成两半，计算两部分得分的相关系数，以反映内部一致性。

3.22.4 效度 validity

量表能够准确测量其所声称要测量的目标特质的程度。效度是评估量表质量的核心指标之一，反映量表是否真正测量到了研究者想要测量的内容。

3.22.4.1 内容效度 content validity

又称“逻辑效度”。测量工具效度的一种类型，用于评估量表或问卷的题目是否全面、准确地覆盖了目标构念的全部理论维度。其核心关注的是测量工具的内容代表性和逻辑完整性。

3.22.4.2 结构效度 construct validity

- 又称“构想效度”。测量工具效度的一种类型，用于评估量表或问卷是否真实、准确地测量了目标理论构念，并反映其内在理论结构。
- 3.22.4.3 表面效度 face validity
量表使用者或测量对象从表面上直观感觉测得内容与拟测量内容之间的相符程度。
- 3.22.4.4 效标效度 criterion-related validity
又称“效标关联效度”。检验新量表与一个公认有效的标准量表分数的相关系数。
- 3.22.4.5 同时效度 concurrent validity
新量表与同时期的标准量表分数之间的相关程度。
- 3.22.4.6 预测效度 predictive validity
量表分数用以预测某种行为或特质发展状况时呈现出的效标效度。
- 3.22.4.7 交叉效度 cross-validity
评估测量工具在独立样本或不同情境下稳定性和泛化能力的效度类型，旨在验证其统计结果的可重复性。
- 3.22.4.8 判别效度 discriminant validity
又称“区分效度”。多个具有不同构想的量表测量分数之间的相关程度。
- 3.22.4.9 效度系数 validity coefficient
量表分数与标准变量或结局变量间的相关程度，为两变量方差的比。
- 3.22.4.9.1 非标准化效度系数 unstandardized validity coefficient
基于原始量纲的量表分数与量表目标特征变量（潜变量）取值的相关系数。
- 3.22.4.9.2 标化效度系数 standardized validity coefficient
标准化后的量表分数与量表目标特征变量（潜变量）取值的相关系数。
- 3.22.4.9.3 唯一效度方差 unique validity variance
量表目标特征变量（潜变量）可由自变量解释的方差比例。
- 3.22.5 项目反应理论 item response theory, IRT
又称“条目反应理论”。以潜在特质理论为基础，研究被测者在测验项目上的反应行为（对某条目的选择概率）与其潜在特质估计值之间关系的测量理论。
- 3.22.5.1 难度系数 coefficient of difficulty
又称“位置参数(location parameter)”。条目特征曲线上上升最快点（拐点）对应的能力大小，具备此能力的受试者答对条目的概率与无任何潜质者答对的概率相等。
- 3.22.5.2 区分度 discrimination
又称“尺度参数(scale parameter)”。条目特征曲线的拐点处的斜率，反映不同能力被试者的区分程度。
- 3.22.5.3 猜测系数 coefficient of guessing
又称“渐近参数(asymptotic parameter)”。完全靠猜测亦能答对问题的概率。
- 3.22.5.4 项目特征函数 item characteristic function
在项目反应条目理论下，体现被测者的待测能力与对应每个条目的反应概率关系的函数表达。
- 3.22.5.5 项目特征曲线 item characteristic curve, ICC
反映被试者的反应概率与其能力之间的关系曲线，曲线的不同部位反映着条目难度、区分度和猜测度的信息。
- 3.22.5.6 项目信息函数 item information function
以方差的倒数作为调查精度的函数。表示条目在评价被试者特质水平时贡献的信息量大小，以评估条目和量表的测量误差。
- 3.22.5.7 项目功能差异 differential item functioning, DIF
潜在特质相同的不同群体在某一条目得分的分布不尽相同的情况。
- 3.22.5.8 项目反应模型 item response model
一种基于概率的用于描述被试者的潜在特质与其对测试项目的反应之间关系的统计模型。
- 3.22.5.8.1 拉施模型 Rasch model
又称“Rasch 模型”。一种单参数逻辑斯谛模型，用于描述被试者的潜在特质与其对项目的反应之间的关系。
- 3.22.5.8.2 局部独立性 local independence
具有相同能力的被试者对某一条目的作答不受其他条目的影响，即扣除被试能力的影响后各条目的作答反应间相互独立。
- 3.22.5.8.3 计算机自适应测试 computerized adaptive testing, CAT
一种利用计算机技术和项目反应理论的测试方法。以项目反应理论为基础，由计算机根据被试者的答题情况构建出使信息函数达到最大值的最佳条目组合，据其对条目组合的答题情况估计被试者能力。
- 3.22.6 概化理论 generalizability theory, GT
根据研究的要求及资料特点，考虑多个误差来源，应用方差分析估计各个误差来源对测量结果的影响，进而优化测量过程，实现测量结果可靠性的最大化的测量学理论。
- 3.22.6.1 测量目标 measurement objective
研究者拟测量的潜在特质。不仅限于被试的某种潜在特质，也可以是条目或评分者的某种特质。
- 3.22.6.2 测量侧面 facet of measurement
除了测量目标以外，影响测量目标观察值的各种条件因素，包括测量工具、测量环境、测量过程、评分专

家，以及观察的场合、情景、时间等。

3.22.6.3 全域分数 universe score

测试或量表中所有项目得分的总和。用于反映被试者在测试或量表所测量的目标领域中的整体表现。

3.22.6.4 观测全域 universe of admissible observations

实际测量活动中所有测量侧面条件全域的集合。

3.22.6.5 概化全域 universe of generalizability

概化理论研究中涉及到的测量面条件全域的集合。

3.22.6.6 决策研究 decision study

概化理论研究中，考察不同测量条件下概化系数的变化，从而探索控制误差的方法，做出最佳决策的过程。

3.22.6.7 概化研究 generalizability study

概化理论研究中，研究者针对所有侧面和侧面目标以及它们之间的交互关系，应用方差分析将总分数变异的方差分解为不同来源的方差，以估计各来源的方差分量（成分相对大小）的过程。

3.22.6.8 概化系数 generalizability coefficient

用相对误差估计出来的信度系数，全域分变异与观测分期望值之比，即测量目标的有效变异与总变异的比值，反映测量的精度。

3.22.6.9 可靠性指数 index of dependability

用绝对误差估计出来的信度系数，被试者观测分均值与全域分总均值之差的平方和，即一个测验或测量的

被测者得分到施测者同等程度接受的所有可能条件下被测者均分的概化的精确性。

3.22.6.10 测量模式 measurement design

测量面的条件样本从观测全域中的抽取模式。如果测量面的条件样本是从观测全域中随机抽取的，则为随机测量模式；如果测量的所有面的条件样本都是固定不变的，则称为固定测量模式；如果一次测量中兼有这两种面，则称为混合测量模式。

3.22.6.10.1 随机单面交叉设计 random single facet crossover design

仅有一个测量侧面，且测量侧面和测量目标中的每个水平都要一一结合，侧面和目标都是随机取样的，总体和全域都是无限的测量设计。

3.22.6.10.2 随机双面交叉设计 random double facets crossover design

观测全域由两个测量侧面的全域组成，且每个侧面的水平及测量目标的水平之间都要一一结合，侧面和目标都是随机取样的，总体和全域都是无限的测量设计。

3.22.6.10.3 多元概化理论 multivariate generalizability theory, MGT

用于由多个相关变量组成的多维度的测量目标研究的概化理论。

3.23 因果推断

3.23 因果推断 causal inference

通过统计学理论与方法，在反事实框架下量化干预变量对结果的干预效应，依赖严格假设（如可忽略性、正性）及数据生成机制的透明性，以区分因果与关联，为科学决策提供依据。

3.23.1 鲁宾因果模型 Rubin causal model, RCM

又称“潜在结果模型(potential outcome model)”、“Rubin 因果模型”。由 保罗·霍兰(Paul W. Holland)提出，以唐纳德·鲁宾(Donald Rubin)的名字命名，以反事实框架为核心，对每个研究对象在一个或多个处理因素作用下所可能产生的预期结果进行建模，从而进行因果效应估计的一类统计模型。

3.23.2 因果图模型 causal diagram

基于直观、透明化的图形完整地展现因果关联系统中各变量间复杂的相互关系的理论模型，是因果关系的识别和因果效应的估计的基础。

3.23.2.1 因果贝叶斯网络 causal Bayesian network, CBN

结合概率论和图论的知识，利用条件概率和图形来描述变量间的复杂相互关系的网络结构，其包含两部分，

其一为表示随机变量间条件独立关系的有向无环图，其二为图节点间变量的条件概率分布。

3.23.2.2 有向无环图 directed acyclic graph, DAG

以每一随机变量作为节点，使用单向箭头的有向边描述变量间的因果关联，且从任意顶点出发无法经过若干路径回到该点的一类没有回环的有向图模型。

3.23.2.2.1 前门准则 frontdoor criterion

给定有向无环图，某变量集合相对于一个由 X 指向 Y 的变量对(X,Y)的一套前门路径的准则，包含该变量集合解释了所有由 X 指向 Y 的有向路径，X 到该变量集合间没有后门路径，且该变量集合到 Y 的后门路径均被 X 所阻断。

3.23.2.2.2 后门准则 backdoor criterion

给定有向无环图，某变量集合相对于一个由 X 指向 Y 的变量对(X,Y)的一套后门路径的准则，包含该变量集合中没有 X 直接指向的节点，且阻断了所有 X 和 Y 之间每条含有指向 X 的路径。

3.23.2.2.3 d 分离 d-separation

又称“有向分离”。概率图模型的一个概念，用于确定图中随机变量之间的条件独立关系。基于有序无环

图，若在给定某变量集合的条件下，变量集合 X 与变量集合 Y 间条件独立，即变量集合 X 和变量集合 Y 被该给定集合分离。

3.23.3 反事实框架 counterfactual framework

通过比较实际观察到的结果与在相同条件下可能的替代情境中的预期结果，来推断因果效应的一种分析框架。

3.23.3.1 反事实 counterfactual

每一研究对象真实仅能观察到一种状态，即真实看到、经历的事实状态。当假设处理变量取值为“现实生活中没有发生”的状态下，研究响应变量在该反事实状态下的取值的一种分析思想。

3.23.3.2 潜在结果 potential outcome

每一研究对象在处理变量多种可能取值下感兴趣结局变量的真实取值。如对于二分类结局，其潜在结果分别为“处理状态”和“对照状态”下，响应变量的真实取值。

3.23.4 因果效应 causal effect

某一暴露或处理对于特定结果产生的干预作用，而不是由于混杂偏倚、选择偏倚或者测量误差等因素或偶然导致的关联。

3.23.4.1 一致性假设 consistency assumption

确保因果效应的定义和估计具有明确的意义假设，即暴露于某种状态的个体，其响应变量的实际取值与暴露于该状态下的真实响应变量取值一致。

3.23.4.2 个体处理稳定性假设 stable unit treatment value assumption, SUTVA

要求每一研究对象的响应变量潜在取值不受所有其他研究对象暴露取值变化的影响，且对于每个研究对象，其所接受的每种暴露水平不存在不同的形式或版本的假设。

3.23.4.3 条件独立假设 conditional independence assumption

又称“条件可交换性假设(exchangeability assumption)”。给定协变量后，响应变量的潜在取值与处理变量的取值无关的假设，即暴露于某种状态的个体，若暴露于另一种状态，其响应变量取值与已知暴露于该状态的个体一致。

3.23.4.4 共同支撑条件 common support condition

基于已知协变量取值，每个研究对象具有非零的条件概率分入某处理组，从而确保每个处理水平都有足够的观测样本以支撑因果效应的推断。

3.23.4.5 平均因果效应 average treatment effect, ATE

衡量一个处理或暴露对总体结果的平均影响。假定存在明确定义的某“处理状态”和与之相对立的“对照状态”。估算某研究群体分别处于“处理状态”和“对

照状态”下响应变量期望取值的差异。

3.23.4.6 个体因果效应 individual causal effect

评估特定个体的处理或暴露对其结果的影响。假定存在明确定义的某“处理状态”和与之相对立的“对照状态”。估算每个研究对象分别处于“处理状态”和“对照状态”下响应变量取值差异。

3.23.5 倾向性评分 propensity score, PS

在给定一组观测到的协变量条件下，个体接受某种干预或暴露的条件概率。它可用于平衡观察性研究中混杂变量，其核心目的是模拟随机对照试验的条件，减少选择偏差。

3.23.5.1 倾向评分分层 propensity score stratification

基于已知协变量的取值估算每个研究对象分入某处理组的条件概率，并基于该条件概率对个体进行分层从而控制混杂因素的统计方法。

3.23.5.2 倾向评分标准化 propensity score standardization

基于已知协变量的取值估算每个研究对象分入某处理组的条件概率，并基于该条件概率对个体进行标准化从而控制混杂因素的统计方法。

3.23.5.3 倾向评分匹配 propensity score matching

基于已知协变量的取值估算每个研究对象分入某处理组的条件概率，并将该条件概率相似的个体进行匹配，从而提高对比组间的均衡性，从而控制混杂的统计方法。

3.23.5.3.1 最邻近匹配 nearest neighbor matching

一种用于因果推断的统计方法，常用于估计处理或暴露的因果效应。通过找到处理组和对照组中特征相似的个体，在给定协变量条件下，匹配到与某处理组概率值最相近的对照组个体作为配比对象的一种匹配方法。

3.23.5.3.2 半径匹配 radius matching

因果推断中一种改进的匹配方法，可用于估计处理效应。基于预先设定的卡尺或半径，从对照组中寻找与该处理组个体，在给定协变量下分入某处理组概率值差异在一定区间内的个体作为配比对象的匹配方法。

3.23.5.3.3 核匹配 kernel matching

基于核函数，以在给定协变量下分入某处理组概率越小权重越大为原则，计算所有对照个体的加权平均值并与该处理个体进行匹配的策略。

3.23.5.4 逆概率加权 inverse probability weighting

一种因果推断方法，用于估计处理效应。基于已知协变量的取值估算每个研究对象分入某处理组的条件概率的倒数作为权重，给每个研究对象分配权重来重新获得一个混杂在组间均衡分布的伪人群。

3.23.6 中介分析 mediation analysis

一种分析因果效应当中的中介现象的分析方法。研究某原因变量通过某中介变量介导而影响结局变量的“中介”现象，并基于一系列统计手段量化该“中介”效应，估算直接和间接因果效应。

3.23.6.1 通径分析 path analysis

基于一组多元线性回归方程来分析变量之间因果关系的一种统计建模方法。

3.23.6.2 因果中介分析 causal mediation analysis

评估和量化原因变量在影响结果变量时的中介机制路径的统计推断过程。

3.23.6.3 自然直接效应 natural direct effect, NDE

因果推断中用于分解因果效应的一种概念，其假定存在明确定义的某“处理状态”和与之相对立的“对照状态”，将中介变量控制在某个暴露状态下的随机取值，估算某研究群体的暴露变量分别处于“处理状态”和“对照状态”下时，响应变量期望取值的差异。

3.23.6.4 自然间接效应 natural indirect effect, NIE

因果推断中用于分解因果效应的一种概念，其假定存在明确定义的某“处理状态”和与之相对立的“对照状态”，将暴露变量控制在某个状态下，估算某研究群体的中介变量分别处于“处理状态”和“对照状态”下的随机取值时，响应变量期望取值的差异。

3.23.7 工具变量 instrumental variable

某变量与模型中随机解释变量高度相关，但却与未知混杂变量独立的一类变量，即该变量仅通过该解释变量来影响研究结局，是估算该解释变量与研究结局间因果效应的有效工具。

3.23.7.1 内生性 endogeneity

统计模型中的一个或多个解释变量与误差项存在相关关系的现象。

3.23.7.2 内生解释变量 endogenous explanatory variable

统计模型中，受模型内部其他变量影响的一类变量，也体现于模型中试图解释的一类变量以及模型中与误差项相关的解释变量。

3.23.7.3 外生解释变量 exogenous explanatory variable

统计模型中，由模型外部因素决定的变量，也体现于模型中试图解释兴趣变量差异的一类变量或模型中与误差项无关的解释变量。

3.23.7.4 过度识别检验 overidentified test

为了验证工具变量的个数是否超过了模型内生变量的个数，即存在多余的工具变量模型过度识别，对所有工具变量进行外生性检验的统计检验方法。

3.23.8 孟德尔随机化 Mendelian randomization

一种利用遗传变异作为工具变量来推断因果关系的

方法。基于“亲代等位基因随机分配给子代”的孟德尔遗传定律，选择一个或多个合适的“基因变异”作为工具变量，通过测量遗传变异与暴露因素、遗传变异与疾病结局之间的关联，进而推断暴露与疾病结局间因果关联。

3.23.8.1 一阶段孟德尔随机化 one stage Mendelian randomization

一种基于孟德尔随机化方法的因果推断策略，旨在使用遗传变异作为工具变量来推断暴露因素对结果变量的因果效应，即通过估计遗传变异变量分别与暴露变量和结局变量间的关联来估计暴露与结局间的因果效应。

3.23.8.2 独立样本孟德尔随机化 one-sample Mendelian randomization

一种基于孟德尔随机化方法的因果推断策略，利用单一研究样本，通过使用二阶段最小二乘回归等模型，定量估计暴露与结局间的因果效应。

3.23.8.3 两样本孟德尔随机化 two-sample Mendelian randomization

一种基于孟德尔随机化方法的因果推断策略，基于来自相同人群的两个独立样本，分别建立遗传变异变量与暴露变量和结局变量的关联来估计暴露与结局间的因果效应，从而提高效应推断的把握度。

3.23.8.4 双向孟德尔随机化 bidirectional Mendelian randomization

基于孟德尔随机化框架的扩展方法，旨在利用遗传变异作为工具变量，同时评估暴露因素对结局以及结局对暴露因素的潜在双向因果关系，从而检验因果关系的方向性并排除反向因果干扰。其核心是通过两组独立的遗传关联数据分别估计双向因果效应，为复杂性状间的因果网络提供统计学证据。

3.23.8.5 两阶段孟德尔随机化 two-step Mendelian randomization

一种基于孟德尔随机化方法的因果推断策略，为探讨暴露因素是否通过中介变量而导致疾病结局的改变，针对暴露变量和中介变量分别选择合适的遗传变异，并分两阶段分别探讨暴露到中介和中介到结局间的因果效应，从而阐明暴露到结局的中间机制。

3.23.9 结构嵌套模型 structural nested model, SNM

一类存在随时间变化的暴露或混杂因素时，通过在多个时间点上构建暴露、混杂和结局的联合分布函数，从而模拟每个研究对象在不同暴露取值下响应变量取值，从而估计时间依赖因果效应的模型。

3.23.9.1 结构嵌套均值模型 structural nested mean model, SNMM

一类存在随时间变化的暴露或混杂因素时，通过在多

个时间点上构建暴露、混杂和结局的联合分布函数，估计给定历史处理和协变量下的每个研究对象在不同暴露取值下的响应变量均值函数，从而估计时间依赖因果效应的模型。

3.23.9.2 结构嵌套分布模型 structural nested

distribution model, SNDM

一类存在随时间变化的暴露或混杂因素时，通过在多个时间点上构建暴露、混杂和结局的联合分布函数，估计给定历史处理和协变量下的响应变量的分布函数，从而估计时间依赖因果效应的模型。

4 研究设计

4 研究设计 research design

科学研究中为实现研究目标而制定的系统性计划。它规定了数据的收集、分析和解释方法，以确保研究结

果的可靠性、有效性和可重复性。其核心是科学地回答研究问题，同时控制混杂因素和偏差。

4.1 实验研究设计

4.1 实验研究设计 design of experiment

一种通过主动操纵自变量、控制混杂变量，并随机分配被试到不同实验组别，以检验因果关系的科学研究方法。其核心目标是确立自变量对因变量的因果效应，并通过严格控制实验条件来最小化偏倚。

既接受研究的干预措施也接受被用于对比的措施的同一组对象组成的对照。

4.1.1 基本概念

4.1.1.1 实验单位 experiment unit

又称“实验对象(experiment subject)”。接受处理因素的基本单位。

4.1.1.4.4.1 单病例随机对照试验 single case randomized controlled trial, N-of-1 trial

基于单个病例进行随机安排干预措施、多周期二阶段交叉设计的对照试验，是一种以病人为中心的个体化临床试验研究手段。

4.1.1.2 处理因素 treatment, treatment factor

研究者根据研究目的施加于实验单位，在实验中需要设定并阐明其效应的因素。

4.1.1.4.5 外部对照 external control

实验以外的对照，可以是过去的研究结果，也可以是同期的另一个研究结果。

4.1.1.2.1 单因素设计 single factor design

实验中仅安排了1个处理因素的一类研究。

4.1.1.4.5.1 历史对照 historical control

又称“潜在对照(potential control)”。以过去的研究结果作为对照。

4.1.1.2.2 多因素设计 multifactor design

实验中安排了1个以上处理因素的一类研究。

4.1.1.4.6 安慰剂对照 placebo control

使用一种与实验药在外观、气味、口味、重量等完全相同，不能被实验对象所识别，且无药理作用的制剂做成的对照。常用于临床试验。

4.1.1.3 实验效应 experiment effect

将处理因素施加于实验单位而观察到的影响。

4.1.1.4.7 阳性对照 positive control, active control

预期会产生所研究效应的对照，通常接受已知的有效处理或标准处理。

4.1.1.4 对照 control

研究过程中为了控制研究对象的复杂性与非处理因素干扰而设置的一组个体，通过与这个组的对比可以更好地揭示处理因素的作用。

4.1.1.4.8 阴性对照 negative control

预期不会产生所研究效应的对照。

4.1.1.4.1 空白对照 blank control

又称“无治疗对照”。不接受任何处理的对照。

4.1.1.4.9 量效对照 dose-response control

又称“剂量-反应对照”。实验药的一系列设定剂量互为对照。

4.1.1.4.2 标准对照 standard control

接受现有标准处理或常规处理的对照。

4.1.1.5 随机化 randomization

实验分组时的操作措施，旨在使每个受试对象均有相同的概率或机会被分配到实验组和对照组。

4.1.1.4.3 互相对照 mutual control

各实验组间互为对照，不专门设立对照。

4.1.1.6 重复 replication

4.1.1.4.4 自身对照 self control

实验研究中在相同条件下进行多次观察,以提高实验的可靠性与科学性。

4.1.1.7 单独效应 simple effect

医学研究中其它影响因素的取值固定时,观察到的研究因素在不同取值下所对应效应的差异。

4.1.1.8 主效应 main effect

某一因素各水平上诸多实验效应平均值的差异。

4.1.1.9 交互作用 interaction

某一因素在另一因素不同水平下的诸单独效应呈现出的变化。

4.1.1.10 准备期 run in

又称“导入期”。实验对象在开始实验前为了进入自然状态而需要经过的一段不加任何处理的观察时间段。

4.1.1.11 洗脱期 washout period

又称“清洗期”。前一阶段处理结束后,为了保证后一阶段的处理效应不受前者影响,由课题组设置的一段不施加任何处理的时间段。旨在令前一阶段的处理效应消失,实验对象又回到自然状态。

4.1.1.12 滞后效应 carry-over effect

实验研究中,前一阶段的处理在后续处理阶段中仍然留存的效应。

4.1.1.13 阶段效应 period effect

又称“时期效应”。实验研究中,在没有施加处理的自然状态下也存在的由于时间流逝而产生的对实验结果的影响。

4.1.2 配对设计 paired design

将实验单位按一定条件配成对子,再将每个对子中的两个实验单位随机分配到不同处理组的实验设计。

4.1.2.1 同源配对 autsyndetic pairing

对照和实验的处理措施作用于同一受试者的配对设计。

4.1.2.2 异源配对 heterogenetic pairing

对照和实验的处理措施作用于不同受试者的配对设计。

4.1.3 完全随机设计 completely randomized design

又称“简单随机设计(simple randomized design)”。采用完全随机化分组方法将同质的实验单位分配到各实验组和对照组的实验设计。

4.1.4 区组设计 block design

按性质(如动物的性别、体重,患者的病情、性别、年龄等非处理因素)相同或相近原则将实验单位组成区组,再将各区组内的实验单位分配到各处理组别的实验设计。

4.1.4.1 随机区组设计 randomized block design

又称“配伍组设计”。按性质(如动物的性别、体重,

患者的病情、性别、年龄等非处理因素)相同或相近原则将实验单位组成区组,再将各区组内的实验单位随机分配到各处理组别的实验设计。各处理组别的实验单位数相等,各区组内处理因素的影响均衡。

4.1.4.2 平衡不完全区组设计 balanced incomplete block design

不苛求处理因素全部水平被施加于同一个区组内实验对象的区组设计。各个处理组别内的个体数相同,每个区组包含的个体数也相同。

4.1.5 拉丁方设计 Latin square design

按行、列形成的方阵及拉丁方字母安排处理及影响因素的实验设计。

4.1.5.1 拉丁方 Latin square

用g个拉丁字母排成使每行每列中每个字母都只出现一次的g行g列方阵。

4.1.5.2 希腊拉丁方 Graeco-Latin square

又称“正交拉丁方”。两个拉丁方阵相正交所得到的方阵。

4.1.6 析因设计 factorial design

将两个或两个以上处理因素的各水平进行全面交叉组合,对各种可能的组合都安排实验的设计方法。常用以评价两个或多个处理因素的效应及各因素间的交互作用。

4.1.6.1 完全析因设计 full factorial design

实验单位被采用完全随机化分组方法分配到全部可能组别的析因设计。

4.1.6.2 随机区组析因设计 randomized block factorial design

依照区组设计方案处于相同区组的个体,被通过随机分组的方式进入析因设计中各个处理因素诸水平全面交叉组合所形成组别的实验设计。

4.1.7 重复测量设计 repeated measurement design

同一实验单位在各种实验条件下进行重复观测的实验设计。

4.1.7.1 前后测量设计 premeasure-postmeasure design

同一实验单位在不同实验条件下进行前后两次观测的实验设计。

4.1.8 正交设计 orthogonal design

一种基于正交表的高效多因素实验设计方法,通过科学安排多因素多水平的实验组合,以最少的试验次数分析各因素的主效应及交互作用。

4.1.8.1 正交性 orthogonality

正交设计中安排实验所依据的专用表中,所有处理因素的水平与任意另一因素各水平间必定搭配1次且仅搭配1次的特性。

4.1.8.2 正交表 orthogonal table

一种数学表格工具，用于科学安排多因素多水平的实验组合，确保实验设计满足均衡性和正交性，从而高效分析各因素的主效应及交互作用。

4.1.9 裂区设计 split-plot design

将区组作为一级实验单位，区组中每个成员作为二级实验单位的设计方法。每个成员将同时接受两个级别因素的处理。

4.1.9.1 完全随机裂区设计 completely randomized split-plot design

一级实验单位随机等分成的 I 组与其二级实验单位随机等分成的 J 组分别接受 A、B 两种因素不同水平处理的实验设计。I、J 各自等于 A、B 两因素的水平数，任何处理组内的例数不小于 2。

4.1.9.2 随机区组裂区设计 randomized block split-plot design

完全随机裂区设计中的一级实验单位先按照随机区组设计的方式分为 r 个组后，每个组内的 I 个一级实验单位随机接受 A 因素 I 种不同处理的设计方案。构建的 $r \times I$ 个一级实验单位中，所有二级实验单位将随机接受 B 因素的 J 种处理。

4.1.9.2.1 全区组 main plot

一级实验单位被分配进入的区组。

4.1.9.2.2 裂区组 split-plot

全区组内部实施的分组。

4.1.9.3 裂区-裂区设计 split-split-plot design

二级实验单位被进一步分为三级实验单位或更细级别实验单位的裂区设计。

4.1.10 嵌套设计 nested design

多个实验因素按照主次级别的方式逐级嵌套的设计方案，主要因素的各水平下嵌套着次要因素的所有可能水平。

4.1.10.1 主因素 main factor

实验研究中对观测指标有主导影响的处理因素。

4.1.10.2 嵌套因素 nested factor

实验因素中对观测指标的影响隶属于上一级影响因素的因素。

4.1.11 平行组设计 parallel group design

将受试对象按事先指定的概率或指定算法算得的概率随机地分配到各试验组，各组同时进行、平行推进，是最常用的临床试验设计类型。

4.1.11.1 双臂试验 two-arm study

一个试验组与一个对照组比较的平行组设计。

4.1.11.2 多臂试验 multi-arm study

有 3 个及以上试验组的平行组设计。

4.1.12 交叉设计 cross-over design

按事先确定的试验顺序，在各个阶段对研究对象先后

实施不同处理，以比较各处理组差异的一种临床试验设计。

4.1.12.1 两阶段交叉设计 two-stage cross-over design

按事先确定的试验顺序，研究对象在两个阶段先后接受两种处理的交叉设计。

4.1.12.2 多阶段交叉设计 multiple-stage cross-over design

按事先确定的试验顺序，研究对象在 2 个以上阶段先后接受不同处理的交叉设计。

4.1.12.3 两因素交叉设计 two-factor cross-over design

研究对象按照随机顺序在不同阶段依次接受两种处理的交叉设计。

4.1.12.4 威廉设计 William design

当存在多个处理时，每种处理会出现在每个阶段和每个顺序，并且在每个阶段和每个顺序中出现次数都相同的一种交叉设计。

4.1.12.5 平衡不完全交叉设计 balanced incomplete cross-over design

当处理较多时，使研究对象尽可能少地接受其中几种处理，每种处理先于彼此出现的次数相等且同种处理在一个顺序中不会重复出现的一种交叉设计。

4.1.13 群随机设计 cluster randomized design

将一些群体随机分配到不同处理组，同一群体内受试对象接受相同处理的一种试验设计，可减少因个体之间相互干扰对处理效果造成的影响。

4.1.13.1 推断单位 inference unit

群随机设计中统计分析的基本试验单位，可以是群或者受试对象。

4.1.14 准实验设计 quasi-experimental design

对随机对照设计的严苛条件做出一定妥协而形成的研究方案。对混杂因素的控制好于非实验设计，而差于随机对照设计的一种实验设计。

4.1.14.1 不等同对照组设计 nonequivalent control group design

一种未贯彻随机分组原则而不能保证非处理因素在组间分布相同的组间比较设计。

4.1.14.1.1 前测 pre-test

实验开始前受试对象接受的对某种属性的测试。

4.1.14.1.2 后测 post-test

实验结束后受试对象接受的对某种属性的测试。

4.1.14.2 交叉滞后组相关设计 cross-lagged panel correlational design

基于面板数据或纵向数据，通过估计不同时间点一个变量与另一个变量的交叉滞后关系和同一个变量的自回归关系，评估实验效应的设计。

4.1.14.3 断点回归设计 regression-discontinuity design

通过设定随机分组变量的取值高于或低于固定截断值或阈值来划分实验组别，通过回归模型评估实验效应的一种设计。

4.1.14.4 间断时间序列设计 interrupted time series design

收集在实验前后多个时间点的测量数据，比较被试接

受实验处理前后的反应模型，以评估实验效应的设计。

4.1.14.5 双重差分法 difference in difference, DID

比较准实验中实验组干预前后的均值的差与对照组干预前后均值的差，以评估实验效应的一种因果推断方法。

4.2 临床试验

4.2 临床试验 clinical trial

一种在人类受试者中进行的系统性医学研究，旨在评估医疗干预措施的安全性、有效性及作用机制，为医学决策提供科学依据。

4.2.1 基本概念

4.2.1.1 意向性治疗 intention-to-treat, ITT

一种临床试验数据分析方法。要求将所有随机分组的受试者纳入最终分析，并按照最初分配的治疗组进行评估，无论其是否完成了治疗或遵守了研究方案。这种方法旨在保留随机化的优势，减少偏倚，并更真实地反映干预措施在实际临床环境中的效果。

4.2.1.1.1 意向性治疗原则 Intention-To-Treat principle

基于治疗意向（如计划的治疗方案）而非实际给予的治疗评估治疗效应的原则。

4.2.1.1.2 基线 baseline

在研究开始时，对受试者或研究对象的关键特征、指标或状态进行的初始测量值。

4.2.1.1.2.1 基线观测值结转 baseline observation carried forward, BOCF

将基线观察应答视作研究终点，是一种非随机缺失数据处理方法。

4.2.1.1.2.2 基线均衡 baseline balance

在实验或观察性研究中，不同组别的受试者在研究开始时，其关键基线特征的分布不存在统计学显著差异。

4.2.1.1.2.3 基线校正 baseline adjustment

通过统计学分析方法对研究对象的基线测量值进行调整，以控制其对结局变量的潜在影响，从而更准确地评估干预或暴露因素的独立效应。

4.2.1.1.3 协变量 covariate

研究中可能与因变量相关，且可能对研究结果产生干扰的变量，需要在统计分析中加以控制以消除潜在混杂。

4.2.1.1.3.1 协变量校正 covariate adjustment

通过统计方法控制协变量对研究结果的影响，从而更准确地估计自变量与因变量之间的独立关联。

4.2.1.1.4 亚组分析 subgroup analysis

针对研究对象的某一特征如性别、年龄段或疾病的亚

型等进行的分析，以探讨这些特征对总效应的影响。

4.2.1.1.4.1 探索性亚组分析 exploratory subgroup analysis

用于探索药物在不同亚组间疗效和/或安全性差异的分析，事先可不在试验方案中明确。

4.2.1.1.4.2 支持性亚组分析 supportive subgroup analysis

当全人群的主要终点同时具有统计学意义和临床意义时，进一步考察试验药物在各个亚组中疗效一致性的分析。

4.2.1.1.4.3 确证性亚组分析 confirmatory subgroup analysis

在确定治疗关键证据的临床试验中，按照临床试验方案和/或统计分析计划中预先规定的人群和多重性调整方法，考察试验药物在目标人群和/或全人群疗效的分析。

4.2.1.1.4.4 一致性 consistency

测量工具、分析方法或研究结果在不同条件下保持稳定、可靠或相互吻合的程度。

4.2.1.1.4.5 异质性 heterogeneity

描述多个研究效应统计量之间的差异与多样性，或内在真实性的变异。

4.2.1.1.4.6 可信度 credibility

研究结果、测量工具或科学主张被合理信任的程度，反映其真实性、可靠性和可接受性。

4.2.1.5 多重性 multiplicity

在统计分析中同时进行多次假设检验或模型比较时，导致I类错误概率增加的现象。

4.2.1.5.1 平行策略 parallel strategy

所包含的各个假设检验相互独立，每个假设检验的推断结果不依赖于其它假设检验的推断结果，是一种多重性问题的处理方法。

4.2.1.5.1.1 前瞻性 α 分配法 prospective alpha allocation scheme, PAAS

一种平行策略的多重性调整方法。令各个假设检验的名义检验水准的互余的乘积等于总的I类错误的方法，各名义检验水准可以相同，也可以不同。

4.2.1.5.2 序贯策略 sequential strategy

按一定顺序对原假设进行检验，直到满足相关条件而停止检验，是一种多重性问题的处理方法。

4.2.1.5.2.1 霍尔姆法 Holm method

一种基于邦费罗尼法的序贯策略的多重调整方法。将 P 值由小到大排序，并逐个与相应的名义检验水平进行比较。

4.2.1.5.2.2 霍赫贝格法 Hochberg method

一种基于西梅斯法的序贯策略的多重调整方法。将 P 值由大到小排序，并逐个与相应的名义检验水平进行比较。

4.2.1.5.2.3 固定顺序法 fixed order method

按预先定义的顺序进行假设检验的一类错误控制方法。每个假设检验的名义检验水准与总的检验水准 α 相同，只有在上一个假设检验拒绝原假设时才进行到下一个假设检验，直到某一个假设检验不拒绝原假设为止。

4.2.1.5.2.4 回退法 fallback method

一种序贯策略的多重性调整方法，事先根据固定顺序法对各假设检验排序，并确定每个假设检验的名义检验水准，然后依顺序进行假设检验，在检验中名义检验水准可以重新分配给下一个假设检验。

4.2.1.5.3 α 消耗函数 α spending function

一种根据试验进行进行 α 分割的函数。当某个临床研究分若干阶段进行整体决策时（如基于有效性所做的期中分析），每个阶段都要消耗一定的 α ，随着研究进展，研究所完成的比例（如 1/3、1/2、3/5 等）与累积的 I 类错误率呈现某种函数关系。

4.2.1.5.3.1 波科克型 α 消耗函数 Pocock-type α spending function

又称“Pocock 型 α 消耗函数”。一种均等分割 α 的函数。在已确定期中分析时间间隔和次数的前提下，以固定的名义检验水准判定期中分析结果，在每次期中分析和终末分析中采用相等的名义检验水准和界值。

4.2.1.5.3.2 奥布赖恩-弗莱明 α 消耗函数

O'Brien-Fleming α spending function

又称“O'Brien-Fleming α 消耗函数”。一种显著性水平逐渐放宽的名义检验水准计算函数，在已确定期中分析时间间隔和次数的前提下使用，其界值随着受试者信息的增加成比例地减小。

4.2.1.5.3.3 海比特尔-佩托 α 型消耗函数

Haybittle-Peto-type α spending function

又称“Haybittle-Peto α 型消耗函数”。用于确定较大的相同临界值进行期中分析，并调整最终分析的临界值以控制总的 I 类错误的函数。通常期中分析检验统计量的临界值取 3。

4.2.1.5.3.4 指数族 α 消耗函数 exponential family α spending function

名义检验水准的指数型计算函数，用于控制成组序贯试验的总 I 类错误。

4.2.1.5.3.5 γ 族 α 消耗函数 gamma family α spending function

一种特殊的 α 消耗函数，通过参数 γ 调节期中分析的界值和名义检验水准，其函数形式为截断的指数分布族。

4.2.1.5.4 β 消耗函数 β spending function

总的 II 类错误率和信息时间的连续函数，用于计算离散的期中分析时间点无效停止的界值。

4.2.1.6 有效性评价 efficacy evaluation

干预措施（如药物或手术）产生预期临床收益能力的评价。包括评估药物对特定疾病或症状的治疗效果，例如疗效的大小、持续时间等。

4.2.1.6.1 非劣效试验 non-inferiority trial

主要目的为显示试验药物的效应在临床上不劣于阳性对照药的试验。用于评估新疗法在效果上不比现有方法差。

4.2.1.6.1.1 非劣效界值 non-inferiority margin

试验药与阳性对照药相比在临床上可接受的最大疗效损失。通常基于阳性对照药相对于安慰剂的疗效差异和临床判断确定临床上可接受的最大损失比例而确定。

4.2.1.6.1.1.1 固定界值法 fixed margin method

在决策或分析过程中，预设一个固定值作为判断结果是否达到标准的方法。临床试验中，将研究药物与阳性对照药的疗效差异与固定界值进行比较，以评估研究药物的疗效。

4.2.1.6.1.1.2 综合法 synthesis method

一种确定非劣效界值并进行统计推断的方法。核心思想是将既往阳性对照药与安慰剂的优效试验和当前试验药与阳性对照药的非劣效试验的数据进行综合，构建检验统计量并进行统计分析。

4.2.1.6.2 优效性试验 superiority trial

主要目的为显示试验药物的效应优于对照药（阳性药或安慰剂）的试验。

4.2.1.6.3 等效性试验 equivalence trial

主要目的为确认两种或多种治疗效果的差别大小在临床上并无重要意义的试验，通常以真正的治疗效果差异落在临床上可接受的等效性界值上下限之间来表明等效性。

4.2.1.6.3.1 等效性界值 equivalence margin

等效性试验中两种药物疗效差值的最大允许值，包含一个上界值，一个下界值。

4.2.1.7 安全性评价 safety evaluation

临床试验中对受试者医学风险的评价，以评估药物的安全性，为药物的批准和使用提供科学依据。

4.2.1.7.1 安全性 safety

受试者因药物或治疗方法而面临的医学风险的程度。在临床试验中通常由实验室检查、生命体征、临床不良事件，以及其他特殊的安全性检查等来判定。

4.2.1.7.2 耐受性 tolerability

受试者对药物的反应和耐受能力，即在使用药物后出现的不良反应或副作用的程度和频率。

4.2.1.7.3 不良事件 adverse event

受试者接受试验用药品后出现的所有不良医学事件，可以表现为症状体征、疾病或者实验室检查异常，但不一定与试验用药品有因果关系。

4.2.1.7.4 不良反应 adverse reaction

用药后产生与用药目的不相符的，并给患者带来有害的和非预想的反应。

4.2.1.8 评价指标

4.2.1.8.1 主要指标 primary endpoint

全称“主要终点指标”。与试验主要研究目的有本质联系的，能确切反映药物有效性或安全性的观察指标。一般应根据试验目的选择在相关研究领域已有公认标准的指标。

4.2.1.8.2 次要指标 secondary endpoint

全称“次要终点指标”。与次要研究目的相关的效应指标，或与试验主要目的相关的支持性指标。

4.2.1.8.3 复合指标 composite indicator

全称“复合终点指标”。当试验难以确定单一指标时，按预先确定的计算方法，将多个指标组合构成单一复合变量。

4.2.1.8.4 全局评价指标 global assessment indicator

将客观指标和研究者对受试者疗效的总印象有机结合的综合指标，它通常是等级指标，其判断等级的依据和理由应在试验方案中明确。

4.2.1.8.5 替代指标 surrogate indicator

在直接测定临床效果不可能或不实际时，用于间接反映临床效果的指标。

4.2.1.8.6 探索性指标 exploratory endpoint

全称“探索性终点指标”。具有探索性分析目的的终点指标。

4.2.1.9 单中心临床试验 single-center clinical trial

只在一个研究单位进行的临床试验。

4.2.1.10 多中心临床试验 multicenter clinical trial

由一个单位的主要研究者总负责，多个单位的研究者参与，按同一个试验方案同时进行的临床试验。

4.2.1.10.1 中心效应 center effect

多中心临床试验中，不同研究中心之间在试验结果上

出现的系统性差异。

4.2.1.10.2 中央随机化 central randomization

多中心临床试验中，统一实现多个临床研究中心的受试者随机化。

4.2.1.11 0 期临床试验 phase 0 clinical trial

I 期试验前开展的新药探索性试验，旨在快速评估新药在人体中的药代动力学和药效学特性。

4.2.1.12 I 期临床试验 phase I clinical trial

药物或医疗干预在人体中进行的首次系统性测试阶段，主要目标是评估安全性、耐受性及初步药代动力学/药效学特征，为后续研究确定合适的剂量范围。

4.2.1.12.1 剂量探索 dose-finding

在早期药物开发中，探索剂量-毒性和剂量-有效性关系的过程。

4.2.1.12.1.1 3+3 设计 3+3 design

一种基于规则的 I 期临床试验设计，用于探索最大耐受剂量。每 3 名受试者序贯接受治疗和评估，根据前序受试者的安全性评估结果决定当前受试者的剂量分配。

4.2.1.12.1.2 改进的毒性概率区间设计 modified toxicity probability interval design, mTPI

一种基于区间模型的贝叶斯 I 期剂量探索方法，将毒性后验概率划分为三个区间：剂量不足、适当剂量、过剂量，通过计算剂量落在上述三个区间的后验概率确定药物的最大耐受剂量的方法。

4.2.1.12.1.3 贝叶斯最优区间设计 Bayesian optimal interval design, BOIN

一种基于贝叶斯理论，根据预设的目标毒性率设置毒性耐受区间的下界和上界，通过最小化剂量分配错误决策的概率来确定药物的最大耐受剂量的方法。

4.2.1.12.1.4 键盘设计 keyboard design

一种基于贝叶斯后验分布区间划分的剂量探索方法。将毒性后验概率划分为一组等宽的剂量区间（键），并根据具有最高后验概率的键计算剂量增减边界。试验结束后，根据毒性率的保序估计选择最大耐受剂量。

4.2.1.12.2 药代动力学分析 pharmacokinetics analysis

全称“药物代谢动力学分析”。对药物在人体内的吸收、分布和消除过程的定量分析。

4.2.1.12.3 扩展队列 expansion cohort

在临床研究早期的初始剂量递增阶段获得一定研究数据之后（剂量递增研究尚在进行中），或紧随递增研究之后进行的三个或以上具有特定队列研究目的的额外患者队列。

4.2.1.12.4 剂量限制毒性 dose limit toxicity, DLT

限制药物剂量进一步增加的不可接受的毒性反应阈值。

4.2.1.12.5 最大耐受剂量 maximum tolerated dose, MTD

临床试验中，药物不引起受试对象出现某种严重不良反应的最大剂量。

4.2.1.12.6 II期推荐剂量 recommended phase II dose, RP2D

根据I期试验结果以及受试者反应确定的适用于II期试验的药物临床用量。

4.2.1.12.7 最优生物学剂量 optimal biological dose, OBD

根据试验目标确定的临床最优剂量。

4.2.1.12.8 剂量毒性曲线 dose-toxicity curve

以剂量为自变量，以毒性为因变量的表现剂量与毒性之间关系的曲线图。

4.2.1.13 II期临床试验 phase II clinical trial

基于I期试验结果，初步探索药物治疗有效性并继续监测药物安全性与病人不良反应的试验。

4.2.1.13.1 概念验证 proof of concept, PoC

在早期临床研究阶段进行的验证候选药物的药理效应是否可以转化成临床获益的试验。

4.2.1.13.2 关键试验 pivotal trial

为药品审批提供决定性证据的临床试验。证明一种治疗的安全性和有效性，并估计常见不良反应的发生率而开展的试验。通常是临床试验过程中的III期试验。

4.2.1.14 III期临床试验 phase III clinical trial

临床开发的确证性研究阶段，其目的是进一步验证药物对目标适应症患者的治疗作用和安全性，评价利益与风险关系，是药物注册申请和审查的主要依据。

4.2.1.15 IV期临床试验 phase IV clinical trial

新药上市后的进一步研究。其目的是考察在广泛使用条件下的药物的疗效和不良反应，评价在普通或者特殊人群中使用的利益与风险关系以及改进给药剂量等。

4.2.1.16 统计分析集 analysis set

用于统计分析的数据集。事先需要明确定义，并在盲态审核时确认每位受试者所属的分析集。

4.2.1.16.1 全分析集 full analysis set, FAS

由尽可能接近符合意向性治疗原则的理想的受试者组成的分析集，从所有随机化的受试者中以最少的和合理的方法剔除受试者后得到。

4.2.1.16.2 符合方案集 per protocol set, PPS

由充分依从于试验方案的受试者组成的分析集，以确保这些数据可能会展现出治疗的效果。

4.2.1.16.3 安全集 safety analysis set, SS

由所有随机化后至少接受一次治疗且有安全性评价的受试者组成的分析集，用于评价安全性与耐受性。

4.2.1.16.4 意向性治疗分析集 intention-to-treat analysis set, ITTS

又称“ITT分析集”。由所有经过随机分组的受试者组成的分析集。

4.2.1.17 统计分析计划 statistical analysis plan, SAP

基于临床试验方案，确定的涉及统计分析技术和操作细节的独立文件，包括对主要和次要评价指标及其他数据进行统计分析的详细过程。

4.2.1.18 统计分析报告 statistical analysis report, SAR

基于统计分析计划，应用统计分析软件编写分析程序输出的统计分析表格和统计分析图形汇总形成的试验结果报告文档。

4.2.1.19 偏倚 bias

由系统误差或过失误差造成的实际测量值或估计值与真实值的非随机偏差。

4.2.1.19.1 非随机对照试验的偏倚

4.2.1.19.1.1 分配偏倚 allocation bias

因非随机化分配导致试验组间在重要的特征上不同所产生的偏倚。

4.2.1.19.1.2 意愿偏倚 intention bias

受试者根据个体情况和环境产生适合自己的入组选择所导致的偏倚。

4.2.1.19.1.3 排除偏倚 exclusion bias

因排除标准导致纳入研究的人群与总体人群存在差异所导致的偏倚。

4.2.1.19.1.4 参与偏倚 participation bias

参与试验的人群与总体人群之间的差异造成的偏倚。

4.2.1.19.2 随机对照试验的偏倚

4.2.1.19.2.1 选择偏倚 selection bias

由于研究对象的选择不当，导致入选者与未入选者特征上存在差异，如缺乏代表性、暴露组与对照组可比性差等，而导致的研究结果偏离真实，从而产生的系统误差。

4.2.1.19.2.2 表现偏倚 performance bias

在临床试验或观察性研究中，由于干预组和对照组在实施干预以外的环节存在差异，导致结局评估失真的一种系统性误差。

4.2.1.19.2.3 测量偏倚 detection bias

组间的结果检测存在系统差异导致的偏倚。

4.2.1.19.2.4 失访偏倚 loss to follow-up bias

因组间患者失访情况不同所导致的偏倚。

4.2.1.19.2.5 报告偏倚 reporting bias

又称“调查对象偏倚 (survey respondent bias)”。在调查过程中研究对象对某些信息的故意夸大或缩小所导致的系统误差。

4.2.1.19.3 混杂 confounding

暴露因素与结局之间的关系被其他变量扭曲的现象

4.2.1.19.3.1 已观测混杂 measured confounding

研究中暴露因素与结局之间的关系被其他观测到的变量扭曲的现象

4.2.1.19.3.2 未观测混杂 unmeasured confounding

研究中暴露因素与结局之间的关系被其他未观测到的变量扭曲的现象

4.2.1.20 估计目标 estimand

在临床试验中对治疗效应的精确描述，反映了针对临床试验目的提出的临床问题。包括治疗、目标人群、变量、伴发事件的处理和总结方式五个属性。

4.2.1.20.1 伴发事件 intercurrent event

治疗开始后发生的、可影响与临床问题相关的观测结果的解读或存在的事件。

4.2.1.20.2 伴随治疗 concomitant therapy

受试者在临床研究期间接受的试验治疗以外的所有治疗。

4.2.1.20.3 疗法策略 treatment policy strategy

假设伴发事件的发生与定义治疗效应无关的疗效估计策略，即无论是否发生伴发事件，均使用相关变量的值。

4.2.1.20.4 假想策略 hypothetical strategy

设想一种没有发生伴发事件的情景，从而进行疗效估计的策略。采用在该假设场景下的变量值作为体现临床问题的变量值。

4.2.1.20.5 复合变量策略 composite variable strategy

将伴发事件纳入关注变量的定义中的疗效估计策略。当伴发事件本身可提供患者结局的信息时，采用此信息进行疗效估计。

4.2.1.20.6 在治策略 while on treatment strategy

只关注伴发事件发生之前的治疗效应来进行疗效估计的策略，伴发事件发生之后的结局被认为与试验治疗无关。

4.2.1.20.7 主层策略 principal stratum strategy

将目标人群分层的疗效估计策略。根据研究目的将目标人群看作是会发生或不会发生伴发事件的“主层”，临床问题仅在该主层中与治疗效应相关。

4.2.1.20.8 估计值 estimate

由估计方法计算得出的数值。

4.2.1.20.9 主分层 principal stratification

一种根据所有治疗中伴发事件的潜在发生情况，将受试者划分为不同的潜在子群体，来估计因果效应的方法。主要用于处理随机化试验中因非依从或死亡截断等问题导致的随机化被破坏的情况。

4.2.1.21 独立数据监查委员会 independent data monitoring committee, IDMC

一个独立的具有相关专业知识和经验的专家组，负责定期审阅来自一项或多项正在开展的临床试验的累积数据，从而保护受试者的安全性、保证试验的可靠性以及试验结果的有效性。

4.2.2 随机对照试验 randomized controlled trial, RCT

把研究对象随机分配到不同的比较组，每组施加不同的干预措施，估计比较组间重要临床结局的差别，以定量估计不同措施的作用或效果的临床试验。

4.2.2.1 随机化 randomization

试验分组时的操作措施，旨在使每个受试对象均有相同的概率或机会被分配到试验的各组。

4.2.2.1.1 简单随机化 simple randomization, SR

又称“完全随机化(complete randomization)”。除了对受试者的数量及组间分配比例有所要求外，对随机化序列不附加任何限制的随机化过程。

4.2.2.1.2 限制性随机化 restricted randomization

包含限制条件的随机化方法。通过设定区组、规定分配组间的最大允许不平衡程度等限制条件，将受试对象进行随机分配，以达到保护事先确定的分配比例的目的的随机化过程。

4.2.2.1.2.1 区组随机化 block randomization

根据区组进行随机化的方法。将受试对象设定区组并在区组内按事先确定的分配比例进行随机分配，区组中对象的数目（区组长度的）需事先确定的随机化过程。

4.2.2.1.2.1.1 固定区组设计 permuted block design, PBD

一种将试验对象划分为若干区组，在每个区组内，处理的分配顺序是随机置换的，确保每个区组内的处理分配是平衡的试验设计方法。

4.2.2.1.2.1.2 可变区组设计 variable block design, VBD

一种将试验对象划分为若干区组，但每个区组的大小可以根据试验需求或试验对象特征动态调整的试验设计方法。

4.2.2.1.2.1.3 区组瓮设计 block urn design, BUD

一种结合了区组设计和瓮模型的试验设计方法，在每个区组内，处理的分配基于瓮模型进行动态调整。

4.2.2.1.3 动态随机化 dynamic randomization

在试验中根据已收集的数据动态调整处理组的分配，以确保组间在某些关键特征上的平衡性。

4.2.2.1.3.1 分层随机化 stratified randomization

根据受试对象的某些重要因子将研究对象分层，然后在每一层内对研究对象进行随机分组的随机化方法。

4.2.2.1.3.2 分层区组随机化 stratified block randomization

根据受试者的某些重要因子将受试对象分层后，逐层

进行区组随机化的试验设计方法。

4.2.2.1.3.3 大棒法 big stick method

一种预设组间不平衡程度的偏币设计，试验中根据两组间人数的差异分配受试者入组。

4.2.2.1.4 适应性随机化 adaptive randomization

一种动态调整受试者分组概率的随机化方法，其核心特点是根据试验中已积累的数据实时优化分配策略，以提升试验的科学性或伦理合理性。

4.2.2.1.4.1 偏币法 biased coin design, BCD

一种通过调整分配概率使双臂人数趋于均衡的动态随机化方法。

4.2.2.1.4.2 瓮法 urn method

一种基于瓮模型的动态随机化分配技术，主要用于平衡试验组和对照组之间的基线特征，尤其适用于小样本研究或适应性临床试验。

4.2.2.1.4.3 最小化法 minimization method

逐一分配病人以使得不同组之间协变量的不平衡程度达到最小的随机化方法。

4.2.2.1.4.4 协变量适应性随机化 covariate adaptive randomization

一种考虑组间协变量均衡性的适应性随机化方法，根据每个协变量组间人数差异，同时考虑协变量重要性权重进行随机化。

4.2.2.1.4.5 反应-适应性随机化 response adaptive randomization

一种根据试验过程中已观察到的试验结果动态调整后续受试者的处理分配概率的动态随机化方法。

4.2.2.1.4.5.1 广义弗里德曼瓮模型 generalized Friedman's urn model, GFU

又称“广义Friedman's 瓮模型”。针对多臂试验的一种适应性随机化方法，根据当前受试者分配的处理，增加后续入组的受试者分配到其他处理的概率。

4.2.2.1.4.5.2 随机化胜者优先法 randomized play-the-winner

利用瓮模型动态调整治疗分配概率，使更优治疗的分配概率随成功次数增加而提高。

4.2.2.1.4.5.3 三重瓮模型 ternary urn model

一种基于瓮随机化的扩展方法，通过模拟从多个瓮中依次抽取球的过程，实现动态随机化。

4.2.2.2 盲法 blinding

在临床试验或实验研究中，有意不让某些相关方知道具体的分组情况或干预措施，以消除主观期望或偏见对结果的影响。

4.2.2.2.1 双盲 double blind

一种盲法，临床试验中受试者和研究者均不知道分组情况。

4.2.2.2.2 单盲 single blind

一种盲法，临床试验中受试者不知道分组情况，而研究者知道。

4.2.2.2.3 非盲 open label

又称“开放(non-masked)”。临床试验中受试者和研究者均知道治疗分配情况。

4.2.2.2.4 破盲 breaking of blindness

揭盲前任何非规定情况所致的盲底泄露的情况。

4.2.2.2.5 揭盲 unblinding

在试验结束后，正式的统计分析前揭晓试验中受试者的治疗分组信息的过程。

4.2.2.2.5.1 紧急揭盲 emergency unblinding

因紧急情况而揭盲的操作。按照临床试验方案规定，基于受试者安全考虑和其他特殊原因，通过预先制定的标准操作规程，在紧急情况下获得单个或部分受试者的治疗分组信息。

4.2.2.2.5.2 安全性揭盲 safety unblinding

因安全性而提前揭盲的操作。按照临床试验方案规定，基于受试者安全考虑，通过预先制定的标准操作规程，获得单个或部分受试者的治疗分组信息。

4.2.2.2.6 模拟技术 dummy

在临床试验中，制备与研究药物感观相似的但不含试验药物有效成分的模拟剂的技术。以保证受试者、研究者不能通过对研究药物的感观获知所使用的具体药物情况，从而保持盲态，减少偏倚，提高实验结果的可靠性。

4.2.2.2.6.1 单模拟技术 single-dummy

在临床试验中，仅对一种治疗方法使用模拟剂，而另一种治疗方法则直接使用真实药物的技术。

4.2.2.2.6.2 双模拟技术 double-dummy

在临床试验中，同时为试验组和对照组提供模拟剂的技术。

4.2.2.2.7 遮蔽 concealment

通过隐藏治疗分配信息，防止受试者、研究者或评估者知晓具体的治疗组别，以减少主观偏倚的方法。

4.2.2.3 对照 control

在实验中设置对照组，用于与试验组进行比较，以评估干预措施的效果。

4.2.2.4 重复 replication

在实验中通过多次重复实验或增加样本量，以提高结果的可靠性和稳定性

4.2.3 非随机对照试验 non-randomized controlled trial

处理组与对照组之间无法实施或因各种原因没有实施随机化分组的前瞻性试验研究。

4.2.3.1 单组目标值试验 single arm objective performance criteria clinical trial

试验设计时设定主要评价指标有临床意义的有效性或安全性最低标准，获得试验组的临床研究数据后，计算相应的统计量并与目标值进行统计学比较，据此来评价被试产品有效性/安全性的一类设计方法。

4.2.3.1.1 目标值 objective performance

临床试验中基于历史数据或科学文献确定的、用于评估新产品或技术性能的定量标准。

4.2.3.1.1.1 客观性能标准 objective performance criteria, OPC

基于现有技术或产品的性能数据，用于评估新产品性能参考标准

4.2.3.1.1.2 性能目标 performance goal, PG

基于临床需求或专家共识设定的、用于评估新产品性能的目标值

4.2.3.1.2 威尔逊计分区间法 Wilson score interval method

威尔逊于 1927 年提出的一种计算总体率置信区间的方法。

4.2.3.1.3 克洛珀-皮尔逊法 Clopper-Pearson method

一种基于二项分布的原理直接计算总体率置信区间的方法。

4.2.3.2 单组多阶段试验 single-arm multi-stage design

一种在单臂试验中可进行多次决策的设计方法，在试验进程中一旦发现试验药物疗效没有达到设定的有效标准，即可早期终止试验。常用于 II 期探索性试验。

4.2.3.2.1 单阶段设计 single-stage design

等待所有受试者完成研究，再根据事先预设的统计分析方法或规则评估药物有效性的设计。

4.2.3.2.2 二阶段设计 two-stage design

一种单组多阶段设计，试验分两个阶段，根据第一阶段试验结果决定是否进入试验第二阶段，试验结束综合两阶段信息决定是否终止研究。

4.2.3.2.2.1 西蒙二阶段设计 Simon's two-stage design

西蒙于 1989 年提出的一种二阶段设计方法，可采用最优化设计和最小最大设计两种策略。

4.2.3.2.2.2 吉亨二阶段设计 Gehan's two-stage design

吉亨于 1961 年提出的一种两阶段设计方法，主要关注第一阶段。根据有关参数确定第一阶段的样本量，如果第一阶段受试者均无效，则终止试验，如果有至少 1 位受试者有效即可进入第二阶段。第二阶段所需的样本量根据第一阶段的有效人数进行设计。

4.2.3.2.2.3 弗莱明二阶段设计 Fleming's two-stage design

弗莱明于 1989 年提出的一种二阶段设计方法，既允许在各阶段因无效而终止试验，也允许在各阶段因有效而终止试验。也可采用三阶段的设计策略。

4.2.3.2.3 三阶段设计 three-stage design

一种单组多阶段设计，试验分三个阶段，根据前一阶段试验结果决定是否进入下一阶段，试验结束综合三个阶段信息决定是否终止研究。

4.2.3.2.3.1 恩赛因三阶段设计 Ensign's three-stage design

恩赛于 1993 年提出的一种三阶段设计方法，综合了吉亨(Gehan)第一阶段设计和西蒙(Simon)最优设计的三阶段设计方法。

4.2.3.2.3.2 弗莱明三阶段设计 Fleming's three-stage design

全称“弗莱明三阶段设计”。弗莱明于 1989 年提出的一种三阶段设计方法，既允许在各阶段因无效而终止试验，也允许在各阶段因有效而终止试验。也可采用二阶段的设计策略。

4.2.4 成组序贯设计 group sequential design, GSD

预先计划在试验过程中进行一次或多次期中分析，依据每一次期中分析的结果做出后续试验决策的试验设计。

4.2.4.1 期中分析 interim analysis

在临床试验过程中，按照事先制订的分析计划，根据试验已累积的数据进行的分析评价，用于评估数据以指导试验决策。

4.2.4.1.1 日历时间 calendar time

通常意义上的时间，根据试验计划完成需要的时间，选择在一定的日历日期进行期中分析。

4.2.4.1.2 信息时间 information time

试验过程中积累的信息占计划总信息的百分比，是用于量化累积信息量的指标。

4.2.4.1.3 名义检验水准 nominal significance level

在统计分析中预先设定的单次检验的显著性水平。

4.2.4.1.4 I 总 I 类错误 family-wise type I error

在多次检验或多阶段试验中，至少一个真的原假设被拒绝的概率。

4.2.4.1.5 随机缩减设计 stochastic curtailment method

一种在临床试验中基于中期分析结果提前终止试验的设计。基于试验累积数据通过概率模型预测试验的最终结论，如疗效已明确或无效，则提前停止试验，避免资源浪费或伦理风险。

4.2.4.2 早期停止 early stopping

期中数据满足预先规定的停止准则时，提前停止试验的策略。

4.2.4.2.1 有效停止 stopping for efficacy

期中分析结果显示试验药明显优于对照时提前停止试验的策略。

4.2.4.2.2 无效停止 stopping for futility

期中分析结果显示试验药不优于对照时提前停止试验的策略。

4.2.4.2.3 安全停止 stopping for safety

在试验期间，出于安全性考虑停止试验的策略

4.2.4.3 重复置信区间 repeated confidence interval, RCI

在试验的各个期中分析时间点建立一系列的置信区间，并保证这些置信区间同时包含试验总体参数目标值的概率可以达到置信水平。

4.2.5 适应性设计 adaptive design

又称“自适应设计”。按照预先设定的计划，在期中分析时使用试验期间累积的数据对试验做出相应修改的临床试验设计。

4.2.5.1 亚组效应 subgroup effect

不同人群之间治疗效果的差异，体现了分组因素对研究药物治疗效果的影响。

4.2.5.2 样本量再估计 sample size re-estimation, SSR

又称“样本量重新估计”。在试验进行过程中，根据预先设定的期中分析计划，利用累积的试验数据重新计算样本量，以保证最终的统计检验能达到预先设定的目标或修改后的目标，并同时能够控制 I 类错误率。

4.2.5.2.1 盲态样本量再估计 blinded sample size re-estimation

在试验进行过程中，基于治疗分配信息未知的期中数据重新估计样本量的过程。

4.2.5.2.2 非盲态样本量再估计 unblinded sample size re-estimation

在试验进行过程中，基于治疗分配信息已知的期中数据重新估计样本量的过程。

4.2.5.3 希望区间 promising zone

又称“前景区”。在临床试验期中分析中，预先设定的允许调整试验设计的区间，若实际的期中分析统计量落入该区间，则需试验的增加样本量。

4.2.5.4 条件功效 conditional power, CP

在给定当前观察到的治疗效果和试验进展程度的前提下，如果假设当前趋势持续不变，到研究结束时能够检测出具有临床意义的治疗差异的可能性。

4.2.5.5 预测功效 predictive power, PP

在考虑当前观察到的期中分析数据和未来数据的不确定性条件下，试验最终达到显著性结果的概率。

4.2.5.6 无缝设计 seamless design

一种将多期独立的试验合并为一个试验进行或整合多期试验数据进行分析的设计方法。主要包括操作无缝设计和推断无缝设计两种类型。

4.2.5.6.1 适应性无缝剂量选择设计 adaptive seamless dose-finding design

将两个试验无缝连接，在前期试验结束时做剂量选择，并将所选剂量用于后期试验，最终分析所有受试者数据的试验设计方法。

4.2.5.6.2 II/III期无缝设计 seamless phase II/III design

一种在同一研究中包含 II 期和 III 期试验，允许在期中分析时进行治疗或剂量选择并与对照组进行有效性比较的试验设计。

4.2.5.6.3 I/II期无缝设计 seamless phase I/II design

一种组合了 I 期和 II 期试验的设计。在 I 期试验中确定最大耐受剂量水平，在 II 期试验中适应性地将患者随机分配到合适的剂量组中以确定临床给药剂量。

4.2.5.6.4 I/II/III期无缝设计 seamless phase I/II/III design

一种创新性的研究策略，通过将传统上分阶段进行的 I 期安全性评估、II 期剂量探索和疗效初步验证以及 III 期确证性研究整合到一个连续统一的试验框架中，实现研究阶段间的动态过渡与数据共享。

4.2.5.7 两阶段适应性设计 two-stage adaptive design

将一个试验分为两个阶段，适应性调整前是第一阶段，适应性调整后是第二阶段。在第一阶段结束时进行期中分析，依据预先设定的修改计划，对第 2 阶段的试验进行适应性修改。

4.2.5.8 多重适应性设计 multiple adaptive design

一个试验中采用两种或两种以上适应性调整方法的试验设计。

4.2.5.9 鲍尔-克内法 Bauer-Köhne method

一种根据费希尔结合检验原理将多次独立检验的 P 值结合起来进行显著性检验的方法，由鲍尔和克内在 1994 年提出，并应用于适应性设计临床试验中。

4.2.5.10 逆正态合并 P 值法 inverse normal method of combining independent p values

一种将适应性设计中各阶段统计检验的 P 值变换为其标准正态分布下的分位数后进行线性合并的分析方法。

4.2.5.11 P 值累加法 method based on sum of P -values, PSM

一种对各个阶段统计检验的 P 值分别取权重相加作为检验统计量的方法。

4.2.6 富集设计 enrichment design

一种选择对试验药物最有可能获益的受试者的设计方法。

4.2.6.1 同质化富集 reducing heterogeneity strategy

通过减少受试者间的异质性以提高临床试验的检验效能的一种研究策略。

4.2.6.2 预后型富集 prognostic enrichment

通过对预后型标志物的识别，入选更有可能观察到终

- 点事件或疾病进展的高风险人群，以增加检验效能的一种策略。
- 4.2.6.3 预测型富集 predictive enrichment**
根据受试者的病理生理、应答史或与药物作用机制有关的疾病特征选择对试验药物最可能有应答的受试者，以提高试验效率的一种研究策略。
- 4.2.6.4 复合型富集 mixed enrichment**
同时使用多个标志物以减少受试者异质性的富集策略。
- 4.2.6.5 适应性富集 adaptive enrichment**
按照预先制定的计划，根据临床试验期中分析结果，在保证试验的合理性和完整性的前提下，对目标人群进行调整的一种研究策略。
- 4.2.7 桥接试验 bridging study**
为了将某个药物或医疗产品在特定人群或地区已获得的疗效和安全性数据外推至新的目标人群或地区而开展的补充性临床研究。
- 4.2.7.1 重现概率 reproducibility probability**
在新地区中重复原始试验结果的概率。
- 4.2.7.1.1 估计功效法 estimated power approach**
一种在研究已经完成后估计研究的统计功效的方法。
- 4.2.7.1.2 置信限法 confidence bound approach**
一种基于统计学区间估计的定量分析方法，通过构建目标人群与原研人群关键疗效参数差异的置信区间，来科学判断原研地区临床试验结果在新地区重现的可能性。
- 4.2.7.1.3 贝叶斯重现概率 Bayesian reproducibility probability**
一种基于贝叶斯统计框架的动态分析方法，通过将原研地区已有的临床试验数据作为先验信息，与目标地区桥接试验获得的当前数据进行概率化整合，从而计算出目标人群能够重现原研人群疗效结果的量化可能性。
- 4.2.7.2 可推广概率 generalizability probability**
一个基于概率论框架的量化指标，用于评估原研地区获得的药物疗效和安全性数据在目标地区临床实践中能够保持其治疗效果的统计可能性。
- 4.2.7.3 种族敏感性分析 racial sensitivity analysis**
判断新地区患者人群在治疗反应方面与原试验开展地区人群是否存在差异的分析，内容包括相关体外、人体药物代谢动力学、药物效应动力学、有效性和安全性评价。
- 4.2.8 生物等效性试验 bioequivalence trial**
通过比较两种药物或同一种药物的两套制剂在人体内的药代动力学数据，来证明它们是否具有几乎相同的生物利用度和人体药物代谢动力学、药物效应动力学参数的临床试验。
- 4.2.8.1 平均生物等效性 average bioequivalence**
人群使用试验药物与使用对照药物所得效应，其平均值相同。
- 4.2.8.1.1 几何均数比 geometric mean ratio, GMR**
两个几何均数的比值。
- 4.2.8.1.2 双单侧检验 two one-sided test, TOST**
统计学中用于验证等效性或生物等效性的假设检验方法，核心思想是通过两次单侧检验，判断目标参数是否落在预先定义的等效区间内。
- 4.2.8.2 群体生物等效性 population bioequivalence**
药物在整个人群中的生物利用度和药代动力学特征的等效性
- 4.2.8.3 个体生物等效性 individual bioequivalence**
对每个个体而言，使用试验药物与使用参比药物所得效应值接近，具有药物可转换性，即正在使用参比药物的患者，如果要转换为试验药物，也能保证具有相同的安全性和有效性。
- 4.2.8.3.1 可转换性 switchability**
患者在服用某药物一段时间后，如果改用另一个与之具有个体生物等效性的药物，可以得到同样效果的特性
- 4.2.8.4 对数变换 logarithm transformation**
对原始数据值取对数，是用于将偏态数据转换成一个新尺度的常用数据变换方法。
- 4.2.8.5 洗脱期 washout period**
又称“清洗期”。为了保证后一阶段的处理效应不受前一阶段处理的影响，需要经过一段不加任何处理的观察时间段以确保前一阶段的处理效应已经消失，实验对象又回到自然状态。
- 4.2.9 诊断试验 diagnostic test**
评价一种诊断方法区分疾病或不同健康状态能力的试验方法。诊断方法包括各种实验室检查、影像学检查、量表评测及其他医学检测等。
- 4.2.9.1 金标准 gold standard**
在现有条件下，公认的、可靠的、权威的诊断方法。
- 4.2.9.2 准确性 accuracy**
观察值与真实值的接近程度。
- 4.2.9.2.1 灵敏度 sensitivity**
实际有病，按诊断试验的标准被正确地判为有病的百分比。反映诊断试验确定患者的能力。
- 4.2.9.2.2 特异度 specificity**
实际无病，按诊断试验的标准被正确地判为无病的百分比。反映诊断试验确定非患者的能力。
- 4.2.9.2.3 约登指数 Youden index**
灵敏度和特异度之和减 1，反映了诊断试验发现真正

的患者和非患者的总能力。

4.2.9.2.4 假阴性率 false negative rate

又称“漏诊率”。金标准确诊的病例中，被错判为阴性者的概率。

4.2.9.2.5 假阳性率 false positive rate

又称“误诊率”。金标准确诊的非病例中，被错判为阳性者的概率。

4.2.9.2.6 阳性符合率 positive percent agreement

在没有金标准时，采用参比方法判断为阳性的个体中，新方法判断为阳性的比例。

4.2.9.2.7 阴性符合率 negative percent agreement

在没有金标准时，采用参比方法判断为阴性的个体中，新方法判断为阴性的比例。

4.2.9.2.8 布兰德-奥尔特曼分析 Bland-Altman analysis

由布兰德-奥尔特曼提出的一种评价定量数据两种测量结果一致性的分析方法，同时控制系统误差和随机误差。

4.2.9.2.8.1 一致性限值 limits of agreement

两种测量结果差值的 95% 参考值范围。

4.2.9.2.8.2 临床可接受范围 clinically acceptable range

由临床医生决定的有临床意义的一致性界限。

4.2.9.2.9 允许总误差/错误结果限区域法 allowable total error and limits for erroneous results zones

一种通过比较允许总误差范围与错误结果限值范围，进行两种定量检测诊断器械一致性评价的方法。

4.2.9.2.9.1 允许总误差 total allowable error

两种检测方法的差异临床允许的总误差大小，表现为评价一致性时允许的误差限值范围。

4.2.9.2.9.2 错误结果限值 limits for erroneous results

两种检测方法的差异视为错误结果（不一致）的限值。

4.2.9.3 精密度 precision

在相同条件下重复测量同一受试者时，诊断试验结果的一致性。

4.2.9.3.1 总符合率 agreement

分类结果与实际结果完全一致的样本占总样本的比例，反映整体分类的准确性。

4.2.9.3.2 卡帕值 Kappa value

又称“Kappa 值”。校正了随机一致性的总符合率，适用于评价观察者间一致性。

4.2.9.4 临床效用 clinical utility

临床实践中的实用性或价值。

4.2.9.4.1 阳性预测值 positive predictive value

试验结果被预测为阳性的人数中，真阳性人数所占比例。

4.2.9.4.2 阴性预测值 negative predictive value

试验结果被预测为阴性的人数中，真阴性人数所占比例。

4.2.9.4.3 阳性似然比 positive likelihood ratio

试验结果真阳性率与假阳性率之比，表示诊断试验正确判断阳性可能性是错误判定阳性可能性的倍数。

4.2.9.4.4 阴性似然比 negative likelihood ratio

试验结果假阴性率与真阴性率之比，表示诊断试验错误判断阴性可能性是正确判断阴性可能性的倍数。

4.2.9.5 受试者工作特征曲线 receiver operator characteristic curve

又称“ROC 曲线”。以 1-特异度为横坐标，灵敏度为纵坐标，依照连续变化的诊断阈值，由不同的灵敏度和特异度画出的曲线。

4.2.9.5.1 曲线下面积 area under curve, AUC

诊断试验中 ROC 曲线下的面积，用来综合评价诊断的准确性，可以理解为所有特异度下的平均灵敏度。

4.2.9.5.2 最佳分界值 optimal cut-off value

诊断试验中区分患病和未患病个体的一个临界值。当诊断试验的结果大于或等于这个阈值时，判断为患病；当结果小于这个阈值时，判断为未患病。

4.2.9.5.3 时间依赖性曲线下面积 time-dependent area under curve

时间依赖设计下的曲线下面积，等于随机选择的一对患病和非患病个体的诊断测试结果正确排序的概率。

4.2.9.6 校准度 calibration

对事件的预测概率与事件的实际频率是否一致的一种度量。

4.2.9.6.1 校准度曲线 calibration curve

事件的预测概率与事件的实际频率之间的关系(或特性)曲线。

4.2.9.6.2 霍斯默-莱梅肖检验 Hosmer-Lemeshow test

又称“Hosmer-Lemeshow 检验”。由霍斯默(Hosmer)和莱梅肖(Lemeshow)提出的一种评价二分类结局变量模型拟合优度的假设检验方法。它将研究对象按照模型预测风险从小到大分为若干组，然后计算每组的平均预测事件发生率与实际观察事件发生率的差异，通过构建卡方统计量来量化这些差异的总体偏离程度。

4.2.10 主方案 master protocol

一种综合性的临床试验设计框架，用于指导复杂的临床试验，能够同时评估多种药物、多个疾病亚组或者多种治疗组合等情况。

4.2.10.1 篮式试验 basket trial

一种临床试验设计方法，基于生物标志物或其他共同的分子特征来选择患者，将具有相同分子靶点的不同

疾病亚型集中在一个试验中进行研究。

4.2.10.2 伞式试验 umbrella trial

一种临床试验设计方法，主要针对一种疾病，同时研究多种不同的治疗药物或治疗方法。

4.2.10.3 平台试验 platform trial

适应性多臂、多阶段试验设计的延伸，允许进行中期

评估和在试验期间添加新的干预措施来评估针对单一疾病人群的多种干预措施。

4.2.10.4 共同对照组 common control arm

在试验中，所有接受不同试验药物或治疗方法的各个试验组共同参照对比的一个组。

4.3 调查设计

4.3 调查设计 survey design

在开展调查研究时，为确保数据的科学性、可靠性和有效性，对调查全过程进行的系统性规划与安排，涵盖从研究目标设定到最终数据收集和分析的各个环节。

4.3.1 基本概念

4.3.1.1 调查对象 survey objects

又称“分析单位”。调查研究者进行调查和抽样的基本单位。

4.3.1.1.1 总体 population

研究对象的全体，具有同质性和变异性的个体组成的集合。

4.3.1.1.1.1 目标总体 target population

根据研究目的选择的，其研究结果能够适用和推论到总体的人群

4.3.1.1.1.2 抽样总体 sampled population

用于抽样的调查单位集合。

4.3.1.1.1.3 子总体 subpopulation

据研究需要确定的具有不同特征（性别、职业、地区）的调查单位集合。

4.3.1.2 调查单位 survey units

调查对象中互不重叠且有限的具体单位。

4.3.1.3 调查条目 survey items

调查问卷或量表的基本组成单元，通常是一个具体的问题或陈述，用于收集受访者的信息、态度、行为或意见。

4.3.1.4 问卷 questionnaire

又称“调查表（survey form）”。根据研究目的确定的由一系列条目组成的一类调查工具，常包含备选答案、必要说明、编号、项目和选项编码、调查证明记载等。

4.3.1.4.1 结构式问卷 structured questionnaire

又称“封闭式问卷”。条目统一制作且每个条目均事先列出若干备选答案的问卷，由被调查者根据个人情况选择作答。

4.3.1.4.1.1 开放式问题 open-ended question

问卷中未列出备选答案的条目，由被调查者自由作答。

4.3.1.4.2 非结构式问卷 unstructured questionnaire

又称“开放式问卷”。条目统一制作但事先并未给出任何备选答案的问卷，由被调查者根据个人情况自由作答。

4.3.1.4.2.1 封闭式问题 close-ended question

设有备选答案的条目，由被调查者从中选择一个或若干个作答。

4.3.1.4.3 半结构式问卷 semi-structured questionnaire

同时包含开放式问题和封闭式问题的问卷。

4.3.1.4.4 敏感性问题 sensitive question

涉及机密或个人隐私等，可能导致拒答率高或真实性低的问卷条目。

4.3.1.4.4.1 随机应答技术 randomized response technique, PRT

一种用于保护受访者隐私的调查方法，使用特定的随机化装置使被调查者以预定的概率来回答敏感性问题，从而保护其真实答案不被直接暴露。

4.3.1.4.4.1.1 沃纳模型 Warner model

一种随机应答模型，在调查中设置敏感问题的正反形式，通过随机化装置，让受访者回答两个互斥问题中的一个。

4.3.1.4.4.1.2 西蒙斯模型 Simmons model

一种随机应答模型，在调查中将敏感问题与一个无关的非敏感问题结合，受访者根据随机装置决定回答哪一类问题。

4.3.1.5 定量调查 quantitative survey

通过量化的问题对调查对象进行数据信息的搜集与分析。

4.3.1.5.1 电话调查 phone survey

一种通过电话与受访者进行交流的信息收集方法。

4.3.1.5.1.1 传统电话调查 traditional phone survey

调查员通过电话直接与受访者进行交流的电话调查方法。

4.3.1.5.1.2 计算机辅助电话调查 computer aided phone survey

一种利用计算机软件和技术辅助进行的电话调查方法。

4.3.1.5.2 面访调查 face to face survey

调查者根据事先制定的调查问卷或访谈提纲面对面地与调查者交谈并记录相关信息的信息收集方法。

4.3.1.5.2.1 入户调查 door to door survey

调查者被允许进入被调查者居所，依照事先制订的调查问卷或访谈提纲对被调查者进行面对面调查或访谈的信息收集方法。

4.3.1.5.2.2 拦截式调查 interception method survey

调查者在路口、商场和街道等人流量较大的公共场合，拦截被调查者进行面对面调研的信息收集方法。

4.3.1.5.2.3 神秘顾客法 mystery customer

调查者经严格培训，在指定时间和地点扮演成顾客，通过参与观察的方式对调查内容进行评估的调查方法。

4.3.1.5.3 网络调查 web survey

调查者利用互联网、计算机通信和数字交互式媒体等方式搜集相关数据和信息的方法。

4.3.1.5.4 电子邮件调查 e-mail survey

调查者以电子信件的形式，把调查问卷发送给被调查者，由被调查者填写、回传，完成调查工作的信息收集方法。

4.3.1.5.5 邮寄调查 mail survey

调查者通过邮寄将调查问卷送至被调查者，由被调查者填写并寄回调查机构的信息收集方法。

4.3.1.5.4.1 留置问卷调查 returned questionnaire survey

调查者事先将调查问卷交给被调查者，由被调查者按要求如实填写，调查者约定日期登门回收或由被调查者寄回的信息收集方法。

4.3.1.5.4.2 固定样本邮寄调查 fixed sampling survey

事先抽取一定区域内自愿参与邮寄调查的固定样本，调查者定期向固定样本中的受调查者发送调查问卷，受调查者按要求填写问卷并回寄调查机构的信息收集方法。

4.3.1.6 定性调查 qualitative survey

一种以深入理解社会现象、人类行为及其背后原因为核心的探索性研究方法，它通过非数值化的资料收集与分析，揭示研究对象的内在意义、行为动机、文化模式或社会过程。

4.3.1.6.1 深度访谈法 in-depth interview

通过调查者与被调查者的当面交流搜集所需资料，调查者可根据被调查者的回答，在访谈中随时提出新的问题以逐步深入主题的一种调查研究方法。

4.3.1.6.2 头脑风暴法 brain storm method

一种组织专家围绕特定研究目的和主题相互交流，自由发表意见，再对意见进行分类、筛选和优化的一种定性调查方法。

4.3.1.6.3 专家咨询法 expert consultation method

通过开会或函询等方式，向有关行业专家咨询相关问题，以充分利用专家既有经验获得有益信息的调查方法。

4.3.1.6.3.1 德尔菲法 Delphi method

一种通过多轮匿名问卷调查收集专家意见，经过归纳与修改，直至专家意见逐步趋于一致，以获得足够可靠的结论或方案的调查方法。

4.3.1.6.4 投影技术 projection technique

又称“摄影技法”。一种通过间接提问或任务设计，以获知受访者潜在态度、动机或行为的调查方法。

4.3.1.7 误差 error

研究对象的实际测量值与真实值之间的差异。

4.3.1.7.1 普查误差 census error

又称“净误差”。普查数据登记值与其真实值之间的差异，通常依据普查数据所反映的内容分为覆盖误差和内容误差。

4.3.1.7.1.1 覆盖误差 coverage error

因抽样总体与目标总体不一致所导致的误差，通常因抽样总体遗漏或包含不属于目标总体的个体所致。

4.3.1.7.1.2 内容误差 content error

因对普查项目的错误报告或错误记录所导致的误差。

4.3.1.7.2 抽样误差 sampling error

由于随机抽样的偶然因素导致样本统计量与总体参数之间的差异，或样本统计量之间的差异。

4.3.1.7.2.1 抽样框误差 sample frame error

由于抽样框与目标总体的不一致所导致的抽样调查估计量的偏差。

4.3.1.7.3 非抽样误差 non sampling error

除随机抽样所致差异外，其他各种原因引起的差异。

4.3.1.7.3.1 无应答误差 non-response error

因某些被调查者拒绝接受调查或失访所导致的测量值与真实值之间的差异。

4.3.1.7.3.2 测量误差 measure error

测量或检测数据与欲调查指标的真值之间的差别。

4.3.2 普查 census

又称“全面调查”。在特定时间对特定范围内的全部观察单位进行的调查或检查。

4.3.3 抽样调查 sampling survey

一种通过从目标总体中科学选取部分个体进行数据收集，进而推断整体特征的研究方法。

4.3.3.1 抽样单元 sampling unit

总体中被实施抽样的基本单元。

4.3.3.2 抽样框 sampling frame

又称“抽样框架”、“抽样结构”。全部抽样单元的完整列表，包括编号和对应的抽样单元名称。

4.3.3.3 抽样比 sampling fraction

抽取的样本单元数与总体单元数的比值。

4.3.3.4 样本量估计 estimation of sample size

为保证研究结论的足够准确和可靠，估计研究所需的最小样本例数。

4.3.3.4.1 相对误差 relative error

某观察指标的绝对误差与其真值之比。

4.3.3.5 抽样设计评价 sampling design evaluation

对不同抽样设计过程中所产生的抽样误差进行科学评估。

4.3.3.5.1 抽样效率 sampling efficiency

针对相同调查单位的同一目标量，两个抽样设计均方误差的比值。

4.3.3.5.2 设计效应 design effect

针对相同调查单位的同一目标量，一个特定的抽样设计估计量方差与相同样本量下简单随机抽样估计量方差的比值。

4.3.3.6 估计方法

4.3.3.6.1 简单估计 simple estimation

利用样本统计量(样本矩)对总体参数(总体矩)进行估计。

4.3.3.6.2 比率估计 ratio estimation

利用目标变量与辅助变量的比例关系进行估计，用以提高精度。

4.3.3.6.3 乘积估计 product estimation

利用目标变量与辅助变量的负相关关系进行估计，用以提高精度。

4.3.3.6.4 回归估计 regression estimation

利用目标变量与辅助变量的线性回归关系进行估计，用以提高精度。

4.3.3.7 概率抽样 probability sampling

又称“随机抽样(random sampling)”。按照概率论原理和统计学抽样方法，以一定的概率从总体中随机选择观察单位的抽样方法。

4.3.3.7.1 放回抽样 sampling with replacement

在逐个抽取单元的过程中，将每次抽中的单元先放回总体后再进行下次抽取的抽样方法。

4.3.3.7.1.1 多项抽样 multi-normal sampling

一种在同一研究框架内结合多种抽样技术的方法，根据研究目标、总体特征和数据收集的可行性，在不同阶段或不同子群体中有针对性地采用最适合的抽样技术。

4.3.3.7.2 不放回抽样 sampling without replacement

在逐个抽取单元的过程中，每次被抽中的单元不再放回总体，或将样本一次性同时抽取的抽样方法。

4.3.3.7.3 简单随机抽样 simple random sampling

从总体中随机抽取若干单位作为样本，每个单位被抽中的概率相等的抽样方法。

4.3.3.7.4 系统抽样 systematic sampling

又称“等距抽样”。一种基于固定间隔规则从有序总体中选取样本的概率抽样方法，其操作过程先将总体中的所有个体按某种顺序进行系统化排列，然后随机确定一个起始点，之后按照预先计算好的固定抽样间隔依次选取样本单元。

4.3.3.7.4.1 抽样间隔 sampling interval

系统抽样的时间间隔或空间间隔。

4.3.3.7.4.2 中心位置系统抽样 centrally located systematic sampling

又称“中心位置等距抽样”。一种改进型的系统抽样方法，通过在传统系统抽样框架中引入中心化选择机制来增强样本对总体核心特征的代表性，其核心操作流程是在确定抽样间隔后，将每个抽样区间划分为前、中、后三个位置段，并固定选取每个区间中点位置或其邻近范围内的单元作为样本。

4.3.3.7.4.3 直线系统抽样 linear systematic sampling

又称“直线等距抽样”。一种将系统抽样原理应用于线性空间或连续序列的标准化抽样技术，其核心特征在于沿着预设的直线路径或一维排列体系，以严格固定的空间或时间间隔选取观测单元。

4.3.3.7.4.4 循环系统抽样 cyclic systematic sampling

一种将传统系统抽样方法应用于周期性或环形结构总体的改进技术，它先识别并量化总体的周期性规律，然后根据周期长度重新设计抽样框架，使得抽样间隔与总体周期形成非整数倍关系，从而避免抽样点总是落在循环的相同相位上。

4.3.3.7.4.5 对称系统抽样 balanced systematic sampling

又称“对称等距抽样，平衡系统抽样”。一种通过精心设计抽样间隔和起始点配置来确保样本单元在总体空间或序列中呈现镜像对称分布的改进的系统抽样方法，它先将研究总体视为一个具有潜在对称性特征的系统，在确定基础抽样间隔后，不再采用单一随机起始点，而是设置互为对称的多个起始点，随后按照固定间隔双向展开抽样，形成在空间、时间或序列维度上具有镜像对应关系的样本网络。

4.3.3.7.5 分层抽样 stratified sampling

将总体中全部个体按某种特征分成若干“层”，再在各层内随机抽取一定数量的个体组成样本的抽样方法。

4.3.3.7.5.1 按比例分配 proportional allocation

根据总体中各层的抽样单元数，从每层抽取的单元个数与该层的单元总数成比例地抽样的方法。

4.3.3.7.5.2 最优分配 optimum allocation

确定分层随机抽样中每个分层的最佳抽样分数和样本量,根据奈曼分配法或精确最佳样本分配法使样本平均数的方差最小化的抽样方法。

4.3.3.7.5.3 多重分层 multiple stratification

研究中若存在与主要变量高度关联的辅助变量,可将其按重要性逐级分层,直至满足研究所需的抽样方法。

4.3.3.7.5.4 事后分层 post-stratification

在抽样完成后依据辅助信息对所抽出的样本进行分层,以利用层内的相似性提高估计精度的抽样方法。

4.3.3.7.6 整群抽样 cluster sampling

将总体按某种标志划分成若干个群,再随机抽取一些群,对抽中群内的所有个体进行研究的抽样方法。

4.3.3.7.6.1 群内相关系数 inter-class correlation coefficient

同一群内不同小单元的指标值对总体均值的离差乘积的期望值与总体中所有小单元指标值对总体均值离差平方的期望值之比,用以描述群单元内基本单元指标的相关性。

4.3.3.7.7 多阶段抽样 multi-stage sampling

又称“多阶抽样”。将抽样过程分阶段进行,各阶段抽样方法往往不同,常在大规模抽样调查中使用的抽样方法。

4.3.3.7.7.1 按容量比例概率抽样 proportionate to population size sampling

又称“概率比例规模抽样”(probability-proportional-to-size sampling)、“PPS 抽样”。每个群被抽取的可能性与其群规模的大小成比例的抽样方法

4.3.3.7.7.2 初级抽样单元 primary sampling unit, PSU

多阶段抽样过程中,首轮随机抽取的群体。

4.3.3.7.7.3 二级抽样单元 second-stage sampling unit, SSU

多阶段抽样过程中,从初级抽样单元中进一步抽取的单元。

4.3.3.7.7.4 三级抽样单元 third-stage sampling unit, TSU

多阶段抽样过程中,从二级抽样单元中进一步抽取的单元。

4.3.3.7.8 连续性调查 successive survey

对同一总体在不同时间进行多次调查,掌握调查对象的动态变化,从而把握事物发展、变化规律及因果关系的调查方法。

4.3.3.7.8.1 重复样本调查 repeated survey

在不同时间点对同一总体进行多次调查,但每次调查的样本是独立抽取的调查方法。

4.3.3.7.8.2 固定样本调查 panel survey

在不同时间点对同一组样本进行多次调查,追踪其变化的调查方法。

4.3.3.7.8.3 轮换样本调查 rotating panel survey

固定样本调查的改进形式,每次调查保留部分样本并替换另一部分样本的调查方法。

4.3.3.7.8.3.1 样本轮换 sample rotation

在多次调查中,每次保留部分样本并替换另一部分样本的设计方法

4.3.3.7.8.3.2 重叠样本 overlap sample

在多次调查中,在不同调查周期中重复出现的样本。

4.3.3.7.8.3.3 样本轮换率 sample rotation rate

从总体中抽取一些新的单元替换上期调查原样本中的部分单元,新单元数占本次调查样本总单元数的比例。

4.3.3.7.8.4 分裂样本调查 split panel survey

一种结合了固定样本调查和重复样本调查的混合调查方法,将样本分为两部分,一部分作为固定样本进行追踪,另一部分作为独立样本进行横截面分析。

4.3.3.8 非概率抽样 non-probability sampling

根据调查者的主观意愿、既有经验或调查实施方便程度等条件来抽取调查对象的抽样方法。

4.3.3.8.1 随意抽样 haphazard sampling

研究者根据个人的方便或直觉,在特定的时间和地点随机选择调查对象,而不考虑总体中个体的特征或分布的抽样方法。

4.3.3.8.1.1 偶遇抽样 accidental sampling

又称“便利抽样(convenience sampling)”。调查者根据实际情况使用对自己最为便利的方式选取样本,可以抽取偶然遇到的人或者选择那些距离最近的、最容易找到的人作为调查对象的抽样方法。

4.3.3.8.2 立意抽样 purposive sampling

又称“判断抽样(judgmental sampling)”。调查者依据研究目标和对情况的主观判断选取调查对象,“有目的”地去选择对总体具有代表性的样本的抽样方法。

4.3.3.8.3 重点调查 key survey

一种非全面调查方法,在所调查的总体中选择一部分重点单位进行的调查。

4.3.3.8.4 滚雪球抽样 snowball sampling

一种通过已有受访者推荐新受访者的非概率抽样方法,常用于难以接触或隐蔽的群体的抽样研究。

4.3.3.8.5 应答推动抽样 respond-driven sampling

一种改进的滚雪球抽样方法,在滚雪球抽样的基础上通过数学模型校正样本偏差,常用于难以接触或隐蔽的群体的抽样研究。

4.3.3.8.6 配额抽样 quota sampling

按照总体的某种特征进行分层，然后在每一层中按照事先规定的比例或数量选取样本的抽样方法。

4.3.3.8.7 空间抽样 spatial sampling

一种专门用于研究地理或三维空间分布现象的抽样方法，通过科学设计的空间布点策略，捕捉研究变量在特定区域内的分布特征和空间依赖关系。

4.3.3.8.8 志愿者抽样 volunteer sampling

研究者邀请和选择自愿参与调查的个体作为样本的抽样方法。

4.3.3.8.9 典型调查 typical survey

一种非全面调查方法，根据调查目的和要求，在对研究对象进行全面分析的基础上，选择部分有代表性的典型单位进行深入细致的调查方法。

4.3.3.9 其他抽样方法

4.3.3.9.1 二重抽样 double sampling

在抽样过程中分两次进行样本选择的抽样方法，研究者首先从总体中抽取一个较大的初步样本，然后在此基础上再从该样本中抽取一个较小的子样本，对其开展调查。

4.3.3.9.2 捕获再捕获抽样 capture-recapture sampling

从总体中无放回随机抽取第一个样本，对每个样本单元做标记后放回总体并充分混合，再从总体中无放回

随机抽取第二个样本，利用样本中标记单元的比例来估计总体单元总数的抽样方法。

4.3.3.9.2.1 逆抽样 inverse sampling

事先确定具有某特征的样本单元数，进行逐个随机抽样，直至抽中满足所考虑特征的单元数，得到样本量不定的抽样方法。

4.3.3.9.3 拉希里-水野抽样 Lahiri-Midzuno sampling

又称“Lahiri-Midzuno 抽样”。一种不等概率抽样方法，通过赋予每个抽样单元不同的入样概率，确保样本能够更好地反映总体特征，适用于分层或多阶段抽样设计。

4.3.3.9.4 泊松抽样 Poisson sampling

各单元的入样概率与单元大小严格成比例、不放回，以入样概率为成功概率进行一次伯努利试验，若试验成功，则相应单元入样，总体均单元进行该试验，得到样本量不定的抽样方法。

4.3.3.9.5 目录抽样 list sampling

将呈偏斜分布的调查总体按其相关标志的最优比例划分为两个部分，对内部差异较大的少部分进行全面调查，而对内部差异较小的大部分进行抽样调查的方法。

4.4 现实世界研究

4.4 现实世界研究 real world study

俗称“真实世界研究”。在现实世界环境下收集与研究对象健康有关的数据或基于这些数据衍生的汇总数据，通过分析获得药物使用等临床诊疗情况及潜在获益-风险的临床证据的研究过程。

4.4.1 现实世界数据 real world data

俗称“真实世界数据”。来源于日常所收集的各种与患者健康状况和/或诊疗及保健有关的数据。

4.4.1.1 医院信息系统数据 hospital information system data

分散存储于医疗卫生机构的电子病历、电子健康档案、实验室信息管理系统、医学影像存档与通讯系统、放射信息管理系统等不同信息系统中的数据。

4.4.1.1.1 电子健康档案 electronic health records, EHR

以电子化形式存储的患者健康信息，涵盖患者的医疗历史、诊断、治疗、用药、检验检查等内容。

4.4.1.1.2 电子病历 electronic medical record, EMR

由医疗机构中授权的临床专业人员创建、收集、管理和访问的个体患者的健康相关信息电子记录。

4.4.1.2 登记研究数据 registry data

通过有组织的系统，利用观察性研究的方法搜集临床和其它来源的数据。可用于评价特定疾病、特定健康状况和暴露人群的临床结局。

4.4.1.3 医保支付数据 administrative and healthcare claim data

医疗保险体系中，包含有关患者基本信息、医疗服务利用、处方、结算、医疗索赔和计划保健等结构化字段的数据。

4.4.1.4 组学数据 genomic data

通过高通量技术在分子水平上大规模测量的生物学数据，涵盖了基因组、转录组、蛋白质组、代谢组等多个层次。

4.4.1.5 社会经济数据 socioeconomic data

以基层行政区为基本单元，通过普查、抽样统计等方式逐级获取的各类人口及经济数据，包括全国人口、国内生产总值等。

4.4.1.6 药品安全性主动监测数据 active drug safety surveillance data

通过国家或区域药品安全性监测网络，从医疗机构、制药公司、医学文献、网络媒体、患者报告结局等渠道收集的数据

4.4.1.7 公共卫生监测数据 public health surveillance data

通过监测系统持续收集的与健康相关的信息，包括疾病发病率、死亡率、危险因素、健康行为、环境暴露等。

4.4.1.8 患者报告结局数据 patient-reported outcome, PRO

来自患者自身测量与评价疾病结局的数据，通常包括症状、生理、心理、医疗服务满意度等。

4.4.1.9 自然人群队列数据 natural population cohort data

对健康人群和/或患者人群通过长期前瞻性动态追踪观察，获取的各种数据。具有统一标准、信息化共享、时间跨度长和样本量较大的特点。

4.4.1.10 现实世界数据适用性 the applicability of real world data

现实世界数据是否适用于产生现实世界证据的特性，包括相关性和可靠性。

4.4.1.10.1 相关性 relevance

现实世界数据与所关注的临床问题的密切程度。重点关注关键变量的覆盖度、暴露或干预和临床结定义的准确性、目标人群的代表性和多源异构数据的融合性。

4.4.1.10.1.1 覆盖度 coverage

现实世界数据包含与临床结局相关的重要变量和信息完整性，如药物使用、患者人口学和临床特征、协变量、结局变量、随访时间、潜在安全性信息等。

4.4.1.10.1.2 代表性 representativeness

现实世界研究人群对现实世界环境下目标人群的符合性。

4.4.1.10.1.3 融合性 linkage

不同来源数据在个体水平进行数据链接、整合和同构

处理的特性。

4.4.1.10.2 可靠性 reliability

衡量现实世界数据准确性的一种度量，主要包括完整性、准确性和透明性。

4.4.1.10.2.1 完整性 integrity

数据信息的完全程度，通常用变量的缺失和变量值的缺失衡量。

4.4.1.10.2.2 准确性 accuracy

数据与其描述的客观特征是否一致，包括源数据是否准确、数据值域是否在合理范围、编码映射关系是否对应且唯一等。

4.4.1.10.2.3 透明性 transparency

现实世界数据的治理方案和治理过程的清晰程度，应确保数据中的关键暴露或干预变量、协变量和结局变量能够追溯至源数据，并反映数据的提取、清洗、转换和标准化过程。

4.4.2 现实世界证据 real world evidence

又称“真实世界证据”。通过对适用的真实世界数据进行恰当和充分的分析所获得的关于药物的使用情况和潜在获益-风险的证据。

4.4.3 比较效果研究 comparative effectiveness research

对比评价用于预防、诊断、治疗疾病和监测健康状况的不同干预措施或策略在真实世界中的患者获益和损害的研究。

4.4.3.1 效果 effectiveness

现实医疗环境中干预的疗效和远期不良反应。

4.4.3.2 实用临床试验 pragmatic clinical trial

尽可能接近真实世界临床实践的临床试验，是介于随机对照试验和观察性研究之间的一种研究类型。

5 卫生统计指标

5.1 居民健康状况指标

5.1 居民健康状况指标 indicators of health status

反映一定时期、一定地区居民寿命、死亡、发病、患病及残障等健康水平的统计指标。

5.1.1 期望寿命 life expectancy

同时出生的一代人活到X岁时，尚能生存的平均年数，一般使用现时寿命表计算。其中出生时的期望寿命为同时出生的一代人未来的平均存活年数，是比较不同

地区、不同时期死亡水平的常用指标。

5.1.2 粗死亡率 crude death rate

某年某地区平均每千人口中的死亡人数，反映当地居民总的死亡水平。是反映居民健康状况的指标之一。

5.1.3 年龄标化死亡率 age-standardized mortality rate

按照某一标准人口的年龄结构计算的死亡率，主要用于不同地区、不同时期的比较。是反映居民健康状况

的指标之一。

5.1.4 死因别死亡率 cause-specific mortality rate

某年某地区每 10 万人口中死于某类疾病或损伤的人数，是反映居民健康状况的指标之一。

5.1.5 死因构成 distribution of causes of death

某年某地区因某疾病或损伤死亡人数占当年该地区全部死亡人数的比例。是反映居民健康状况的指标之一。

5.1.6 5 岁以下儿童死亡率 under-five mortality rate

某年某地区 5 岁以下儿童死亡数与当年该地区活产数之比，是反映居民健康状况的指标之一。

5.1.7 婴儿死亡率 infant mortality rate

某年某地区未满 1 岁的婴儿死亡数与当年该地区活产数之比，是反映居民健康状况的指标之一。

5.1.8 新生儿死亡率 neonatal mortality rate

某年某地区未满 28 天的新生儿死亡数与当地活产数之比，是反映居民健康状况的指标之一。

5.1.9 围产儿死亡率 perinatal mortality rate

某年某地区从怀孕 28 周到出生后 7 天内的死胎死产数和新生儿死亡数与当地该地区围产儿总数之比，是反映居民健康状况的指标之一。

5.1.10 5 岁以下儿童死因构成 distribution of causes of death among children under five years old

某年某地区 5 岁以下儿童死亡总数中因某类疾病或损伤死亡所占比例。是反映居民健康状况的指标之一。

5.1.11 新生儿死因构成 distribution of causes of neonatal death

某年某地区新生儿死亡总数中因某类疾病或损伤死亡所占比例。是反映居民健康状况的指标之一。

5.1.12 孕产妇死亡率 maternal mortality ratio

某年某地区由于怀孕和分娩及并发症造成的孕产妇死亡数与同年该地区活产数之比，是反映居民健康状况的指标之一。

5.1.13 孕产妇死因构成 distribution of causes of maternal death

某年某地区孕产妇死亡总数中因某类疾病死亡所占比例。是反映居民健康状况的指标之一。

5.1.14 法定传染病报告死亡率 reported mortality rate of notifiable infectious diseases

某年某地区法定传染病报告死亡人数占同年该地区总人口数的比例，是反映居民健康状况的指标之一。

5.1.15 职业病病死率 fatality rate of occupational disease

某年某地区因各类职业病死亡人数与同年该地区各类职业病患者人数之比，是反映居民健康状况的指标之一。

5.1.16 法定传染病报告发病率 reported incidence rate of notifiable infectious disease

某时期某地区法定传染病报告的新发病例数与同期该地区人口数之比，是反映居民健康状况的指标之一。

5.1.17 新生儿破伤风发病率 incidence rate of neonatal tetanus

某年某地区新生儿中破伤风发病人数与活产人数之比，是反映居民健康状况的指标之一。

5.1.18 职业病发病率 incidence rate of occupational disease

某年某地区各类职业病新发病人数与所有接触能导致职业病的有害物质的人数之比，是反映居民健康状况的指标之一。

5.1.19 患病率 prevalence

某特定时间内总人口中某病新旧病例所占的比例。是反映居民健康状况的指标之一。

5.1.20 成人高血压患病率 prevalence of hypertension among adults

某时期某地区 18 岁及以上人口中高血压患者所占比例，是反映居民健康状况的指标之一。

5.1.21 成人糖尿病患病率 prevalence of diabetes among adults

某时期 18 岁及以上人口中糖尿病患者所占比例，是反映居民健康状况的指标之一。

5.1.22 宫颈癌患病率 prevalence of cervical carcinoma

某年某地区妇女中宫颈癌患者所占比例，是反映居民健康状况的指标之一。

5.1.23 乳腺癌患病率 prevalence of breast cancer

某年某地区妇女中乳腺癌患者所占比例，是反映居民健康状况的指标之一。

5.1.24 孕产期中重度贫血患病率 prevalence of maternal moderate or severe anemia

某年某地区孕产期发生中重度贫血的产妇数占产妇总数的比例，是反映居民健康状况的指标之一。

5.1.25 孕产妇艾滋病病毒感染率 HIV prevalence among pregnant women

某地区某时期内孕产妇中艾滋病病毒感染人数所占比例，是反映居民健康状况的指标之一。

5.1.26 孕产妇梅毒感染率 syphilis prevalence among pregnant women

某年某地区孕产妇中梅毒感染人数所占比例，是反映居民健康状况的指标之一。

5.1.27 孕产妇乙肝表面抗原阳性率 positive rate of HBsAg among pregnant women

某年某地区孕产妇中乙肝表面抗原阳性人数所占比例，是反映居民健康状况的指标之一。

5.1.28 人群血吸虫感染率 infection rate of schistosomiasis

某年某地区血吸虫流行人口中血吸虫感染人数的比例，是反映居民健康状况的指标之一。

5.1.29 血吸虫病病人数 number of schistosomiasis patients

某年某地区急性血吸虫病病人数、晚期血吸虫病病人数及慢性血吸虫病病人数之和，是反映居民健康状况

的指标之一。

5.1.30 地方病现症病人数 number of patients with endemic diseases

年底实有的某种地方病现症病人数，是反映居民健康状况的指标之一。

5.1.31 出生缺陷发生率 birth defect rate

某年某地区患有先天性缺陷的围产儿数占当地围产儿总数的比例。是反映居民健康状况的指标之一。

5.2 健康影响因素指标

5.2 健康影响因素指标 indicators of health-related risk factors

反映一定时期、一定地区影响居民健康状况的环境、行为/生活方式、遗传生物学、卫生服务等因素的统计指标。

5.2.1 农村卫生厕所普及率 proportion of rural households with sanitary toilets

符合国家农村户厕卫生规范的累计卫生厕所农户数占当地农村总户数的比例，是反映健康影响因素的指标之一。

5.2.2 农村无害化卫生厕所普及率 proportion of rural households with harmless sanitary toilets

符合国家农村户厕卫生规范中无害化卫生厕所要求建设的累计卫生厕所户数占农村总户数的比例，是反映健康影响因素的指标之一。

5.2.3 成人现在吸烟率 prevalence of current smoking among adults

18岁及以上人口中现在吸烟者所占比例，是反映健康影响因素的指标之一。

5.2.4 成人人均酒精消费量 per capita consumption of pure alcohol among adults

18岁及以上人口酒精人均消费量，是反映健康影响因素的指标之一。

5.2.5 成人蔬菜水果摄入不足率 prevalence of low consumption on fruit and vegetable among adults

18岁及以上人口中蔬菜水果摄入不足者所占比例，是反映健康影响因素的指标之一。

5.2.6 成人身体活动不足率 prevalence of insufficiently physically active among adults

18岁及以上人口中身体活动不足者所占比例，是反映健康影响因素的指标之一。

5.2.7 人均日食物摄入量 per capita daily food intake

每人每日平均所有食物的摄入量，是反映健康影响因素的指标之一。

5.2.8 人均日能量摄入量 daily calorie intake per

capita

由一定时期内调查人群能量摄入总量除以同期调查人群总人日数计算得出，反映人群每人每日平均能量摄入量，是反映健康影响因素的指标之一。

5.2.9 居民日均营养素摄入量 daily nutrients intake per capita

由一定时期内调查人群某种营养素摄入总量除以同期调查人群用餐总人日数计算得出，反映人群每人每日平均某种营养素的摄入量，是反映健康影响因素的指标之一。

5.2.10 膳食结构 dietary structure

定时期调查人群膳食中各类食物数量所占比重，表示膳食中各种食物间的组成关系，是反映健康影响因素的指标之一。

5.2.11 健康素养水平 health literacy level

由一定时期某地区具备基本健康素养人数除以同期该地区健康素养调查人数计算得出，反映具备基本健康素养人数占全人群中的比例，是反映健康影响因素的指标之一。

5.2.12 中医药健康文化素养水平 health literacy level of traditional Chinese medicine

由一定时期某地区具备基本中医药健康文化素养人数除以同期该地区健康素养调查人数计算得出，反映具备基本中医药健康文化素养人数占全人群中的比例，是反映健康影响因素的指标之一。

5.2.13 成人肥胖率 prevalence of obesity among adults

由某时期18岁及以上肥胖人数除以同期18岁及以上调查人数计算得出，反映18岁及以上人口中肥胖人数所占比例，是反映健康影响因素的指标之一。

5.2.14 5岁以下儿童肥胖率 prevalence of obesity among children under five years old

由报告期内某地区5岁以下儿童肥胖人数除以同期该地区5岁以下儿童身高(长)体重检查人数计算得出，反映5岁以下儿童肥胖人数占全部5岁以下儿童数的比例，是反映健康影响因素的指标之一。

5.2.15 成人血压升高流行率 prevalence of raised blood pressure among adults

由某时期 18 岁及以上血压升高者的人数除以同期 18 岁及以上调查人数计算得出，反映 18 岁及以上人口中血压升高者的人数占 18 岁及以上总人口的比例，是反映健康影响因素的指标之一。

5.2.16 成人总胆固醇升高流行率 prevalence of raised total cholesterol among adults

由某时期 18 岁及以上人口中检出的总胆固醇升高者的人数除以同期 18 岁及以上调查人数计算得出，反映 18 岁及以上人口中总胆固醇升高者所占比例，是反映健康影响因素的指标之一。

5.2.17 成人血糖升高流行率 prevalence of raised blood glucose among adults

由某时期 18 岁及以上血糖升高者的人数除以同期 18 岁及以上调查人数计算得出，反映 18 岁及以上人口中血糖升高者的人数占全 18 岁以上人群的比例，是反映健康影响因素的指标之一。

5.2.18 5 岁以下儿童低体重率 prevalence of underweight among children under five years old

由报告期内某地区 5 岁以下儿童低体重人数除以同期该地区 5 岁以下儿童体重检查人数计算得出，反映 5 岁以下儿童低体重人数占全部 5 岁以下儿童的比例，是反映健康影响因素的指标之一。

5.2.19 5 岁以下儿童生长迟缓率 prevalence of stunting among children under five years old

由报告期内某地区 5 岁以下儿童生长发育迟缓的人数除以同期该地区 5 岁以下儿童身高（长）体重检查人数计算得出，反映 5 岁以下儿童生长迟缓人数占有 5 岁以下儿童的比例，是反映健康影响因素的指标之一。

5.2.20 低出生体重率 incidence of low-birth-weight newborns

由报告期内某地区低出生体重儿数除以同期该地区活产数计算得出，反映出生体重低于 2500 克的活产数占全部活产数的比例，是反映健康影响因素的指标之一。

5.2.21 碘盐覆盖率 proportion of households using iodized salt

由抽样家庭中碘含量 ≥ 5 毫克每千克盐样份数除以同期抽样居民食用盐样份数计算得出，反映食用碘盐的家庭数占全部家庭的比例，是反映健康影响因素的指标之一。

5.2.22 碘盐合格率 qualified rate of iodized salt

由一定时期内食用合格碘盐的家庭数除以同期食用碘盐家庭数计算得出，反映食用合格碘盐的家庭数占食用碘盐家庭的比例，是反映健康影响因素的指标之一。

5.3 疾病控制指标

5.3 疾病控制指标 indicators of disease control

反映一定时期规划免疫、传染病、慢性病、严重精神障碍、寄生虫病、地方病防治等疾病控制的统计指标。

5.3.1 国家免疫规划疫苗接种率 routine immunization coverage

由某疫苗（某剂次）实际接种人数除以该疫苗（该剂次）应接种人数计算得出，反映某年适龄儿童中完成国家免疫规划疫苗常规免疫接种人群所占比例，是疾病控制的统计指标之一。

5.3.2 传染病报告及时率 timely reporting rate of infectious diseases

由某时期某地区在规定时限内报告的甲乙丙类传染病病例数除以同期该地区传染病直报系统报告的甲乙丙类传染病病例总数计算得出，反映某时期某地区通过传染病报告管理系统报告的法定甲乙丙类传染病病例数中，在规定时限内报告的病例数所占比例，是疾病控制的统计指标之一。

5.3.3 艾滋病病毒感染者和病人接受抗病毒治疗的比例 antiretroviral therapy coverage among people living

with HIV/AIDS

截至统计时间某地区报告存活的艾滋病病毒感染者和病人中正在接受抗病毒治疗人数所占的比例，是疾病控制的统计指标之一。

5.3.4 新涂阳肺结核病人治愈率 treatment success rate for new pulmonary smear-positive tuberculosis cases

某时期治愈的新涂阳肺结核患者占同期全部登记的新涂阳肺结核患者数的比例，是疾病控制的统计指标之一。

5.3.5 居民健康档案规范化电子建档率 percentage of population with standardized electronic health records

年末某地区城乡居民累计建立规范化电子健康档案人数与常住人口数之比，是疾病控制的统计指标之一。

5.3.6 高血压规范管理人数 number of people with hypertension under standardized management

年末某地区建立高血压患者管理档案并进行规范管理的人数，是疾病控制的统计指标之一。

5.3.7 糖尿病规范管理人数 number of people with diabetes under standardized management

年末某地区建立糖尿病患者管理档案并进行规范管理的人数，是疾病控制的统计指标之一。

5.3.8 成人高血压控制率 proportion of people with raised-blood pressure controlled among adults with hypertension

某时期18岁及以上高血压人群中血压达标者的比例，是疾病控制的统计指标之一。

5.3.9 成人糖尿病控制率 proportion of people with fasting blood glucose controlled among adults with diabetes

某时期18岁及以上糖尿病患者中空腹血糖达标者的比例，是疾病控制的统计指标之一。

5.3.10 65岁以上人口健康管理人数 number of populations over 65 under health management

某年某地区建立65岁以上人口健康档案并进行健康管理人数，是疾病控制的统计指标之一。

5.3.11 12岁人口患龋率 prevalence of caries among 12-year-olds

由某时期检出的12岁龋病患者数除以同期12岁受检人数计算得出，反映一定时期内12岁人口中龋病患者所占比例，是疾病控制的统计指标之一。

5.3.12 死因登记覆盖率 civil registration coverage of cause-of-death

由某年某地区医疗卫生机构签发的《居民死亡医学证明（推断）书》数除以同年该地区估计的总死亡人数计算得出，反映某年某地区医疗卫生机构签发的《居民死亡医学证明（推断）书》数占同年该地区死亡人数的比例，是疾病控制的统计指标之一。

5.3.13 严重精神障碍患者报告患病率 registration rate of severe psychiatric patients

由年末某地区所有登记在册的确诊严重精神障碍人数除以同年该地区常住人口数计算得出，反映年末某地区所有登记在册的确诊严重精神障碍人数与同年该地区常住人口数之比，是疾病控制的统计指标之一。

5.3.14 病区改水率 proportion of epidemic regions with improved drinking-water supply

某地区历年累计完成改水任务的病区村数与该地区该时期末病区村总数之比，是疾病控制的统计指标之一。

5.3.15 病区改炉改灶率 proportion of households installed improved stoves in epidemic regions

由某地区历年累计完成改炉改灶的病区户数除以该时期末该地区病区总户数计算得出，反映某地区历年累计完成改炉改灶的病区户数与该地区病区总户数之比，是疾病控制的统计指标之一。

5.4 妇幼保健指标

5.4 妇幼保健指标 indicators of maternal and child health care

反映一定时期、一定地区内婚前保健、孕产期保健、儿童保健、妇女保健、计划生育技术服务等统计指标。

5.4.1 婚前医学检查率 proportion of people receiving premarital medical examination

由报告期内某地区婚前医学检查人数除以同期该地区结婚登记人数计算得出，是妇幼保健领域的统计指标之一。

5.4.2 婚前医学检查检出疾病率 detection rate of diseases among people receiving premarital medical examination

由报告期内某地区检出疾病人数除以同期该地区婚前医学检查人数计算得出，反映一定时期内某地区婚检检出疾病人数与婚前医学检查人数之比。

5.4.3 妇女常见病筛查率 screening coverage of gynecological and breast diseases

由报告期内某地区实查人数除以同期该地区应查人数计算得出，反映一定时期内某地区妇女常见病实查人数与应查人数之比，是妇幼保健领域的统计指标之一。

5.4.4 计划生育技术服务例数 number of family planning service cases

报告期内某地区施行放置（取出）宫内节育器术，输精（卵）管绝育术及吻合术，人工流产（负压吸引术、钳刮术、药物流产），放置（取出）皮下埋植剂的人次数之和，是妇幼保健领域的统计指标之一。

5.4.5 活产数 number of live births

报告期内某地区妊娠满28周及以上（如孕周不清楚，可参考出生体重达1000克及以上），娩出后有心跳、呼吸、脐带搏动、随意肌收缩4项生命体征之一的新生儿数，是妇幼保健领域的统计指标之一。

5.4.6 产妇建卡率 health care recording rate of pregnant women

由报告期内某地区产妇建卡人数除以同期该地区产妇数计算得出，反映一定时期内某地区产妇建卡人数与产妇数之比，是妇幼保健领域的统计指标之一。

5.4.7 高危产妇占产妇总数的百分比 proportion of pregnant women with high risk factors

由报告期内某地区高危产妇人数除以同期该地区产

妇数计算得出，反映一定时期内某地区产妇总数中高
危产妇所占百分比，是妇幼保健领域的统计指标之一。

5.4.8 孕产妇系统管理率 maternal systematic management rate

由报告期内某地区产妇系统管理人数除以同期该地
区活产数计算得出，反映一定时期内某地区产妇系统
管理人数与当地活产数之比，是妇幼保健领域的统计
指标之一。

5.4.9 产前检查率 antenatal examination rate

由报告期内某地区产妇产前检查人数除以同期该地
区活产数计算得出，反映一定时期内某地区接受过 1
次及以上产前检查的产妇人数与当地活产数之比，是
妇幼保健领域的统计指标之一。

5.4.10 孕产妇艾滋病病毒检测率 HIV testing coverage among pregnant women

由报告期内某地区产妇接受艾滋病病毒检测人数除
以同期该地区产妇数计算得出，反映一定时期内某地
区产妇艾滋病病毒检测人数与产妇数之比，是妇幼保
健领域的统计指标之一。

5.4.11 孕产妇梅毒检测率 syphilis testing coverage among pregnant women

一定时期内某地区产妇梅毒检测人数与产妇数之比，
是妇幼保健领域的统计指标之一

5.4.12 孕产妇乙肝表面抗原检测率 HBsAg testing coverage among pregnant women

一定时期内某地区产妇乙肝表面抗原检测人数与产
妇数之比，是妇幼保健领域的统计指标之一

5.4.13 孕妇产前筛查率 prenatal screening coverage of birth defects

一定时期内某地区出生缺陷产前筛查孕产妇数与当
地产妇数之比，是妇幼保健领域的统计指标之一

5.4.14 住院分娩率 hospital delivery rate

一定时期内某地区住院分娩活产数与当地活产数之
比，是妇幼保健领域的统计指标之一

5.4.15 剖宫产率 caesarean delivery rate

一定时期内某地区采用剖宫产手术分娩的活产数与
当地活产数之比，是妇幼保健领域的统计指标之一

5.4.16 产后访视率 coverage of postnatal care visit within 28 days of childbirth

一定时期内某地区接受产后访视的产妇人数与活产
数之比，是妇幼保健领域的统计指标之一

5.4.17 7 岁以下儿童健康管理率 health management rate for children under seven years old

一定时期内某地区 7 岁以下儿童中接受健康管理服
务的人数所占比例，是妇幼保健领域的统计指标之一

5.4.18 3 岁以下儿童系统管理率 systematic management rate for children under three years old

一定时期内某地区 3 岁以下儿童中接受系统管理的
人数所占比例，是妇幼保健领域的统计指标之一

5.4.19 0-3 岁儿童中医药健康管理人数 number of children under three years old receiving health management of traditional Chinese medicine

年末按照国家本年度《国家基本公共卫生服务规范》
要求，为 0-3 岁儿童、65 岁以上老人建立 健康档
案并提供儿童中医药调养的人数，是妇幼保健领域的
统计指标之一

5.4.20 新生儿甲状腺功能减低症筛查率 screening coverage of neonatal congenital hypothyroidism

一定时期内某地区接受过甲状腺功能减低症筛查的
新生儿数与同期该地区活产数之比，是妇幼保健领域
的统计指标之一

5.4.21 新生儿苯丙酮尿症筛查率 screening coverage of neonatal phenylketonuria

一定时期内某地区接受过苯丙酮尿症筛查的新生儿
数与当地活产数之比，是妇幼保健领域的统计指标之
一

5.4.22 新生儿听力筛查率 screening coverage of newborn hearing

一定时期内某地区接受过听力筛查的新生儿人数与
活产数之比，是妇幼保健领域的统计指标之一

5.4.23 6 个月内婴儿纯母乳喂养率 exclusive breastfeeding under 6 months

一定时期内某地区母乳喂养人群中纯母乳喂养所占
比例，是妇幼保健领域的统计指标之一

5.5 卫生监督指标

5.5 卫生监督指标 indicators of health inspection and supervision

反映一定时期、一定地区内卫生监督许可、执法、处
罚等管理相关统计指标

5.5.1 有效卫生许可证份数 number of valid hygienic licenses

报告期末某地区经卫生行政部门批准的、在有效期内的
卫生许可证份数，是反映卫生监督情况的统计指标
之一

5.5.2 放射防护预评价审核总数 number of pre-evaluation reports audited for radiological protection

报告期内某地区开展建设项目放射防护预评价审核

总数，是反映卫生监督情况的统计指标之一

5.5.3 选址和设计卫生审查总数 number of sanitation review for siting and design of building projects

报告期内某地区开展建设项目选址和设计卫生审查项目总数，是反映卫生监督情况的统计指标之一

5.5.4 应监督户数 number of units to be supervised

报告期内某地区应实施经常性监督的被监督单位总数，是反映卫生监督情况的统计指标之一

5.5.5 实监督户数 number of units supervised

报告期内某地区应监督户数中实际实施经常性监督的被监督单位总数，是反映卫生监督情况的统计指标之一

5.5.6 实监督户次数 number of times supervised

报告期内某地区实际实施经常性卫生监督的户次数量，是反映卫生监督情况的统计指标之一

5.5.7 查处案件数 number of cases been investigated and treated

报告期内某地区实施卫生行政处罚、行政强制及其他措施的案件，包括行政处罚案件和未做处罚但予以行政处理的案件数。是反映卫生监督情况的统计指标之一。

5.5.8 责令限期改正户次数 number of times ordered to improve within a deadline

报告期内某地区卫生监督部门针对违法行为依法实施责令限期改正的户次数，是反映卫生监督情况的统计指标之一。

5.5.9 行政处罚户次数 number of times given administrative penalties

报告期内某地区卫生监督部门针对违法行为依法实

施的卫生行政处罚户次数，是反映卫生监督情况的统计指标之一

5.5.10 警告户次数 number of times warned

报告期内某地区卫生监督部门实施卫生行政处罚行为中给予警告的户次数，是反映卫生监督情况的统计指标之一。

5.5.11 罚款户次数 number of times fined

报告期内某地区卫生监督部门实施卫生行政处罚行为中给予罚款的户次数，是反映卫生监督情况的统计指标之一。

5.5.12 行政处罚金额 amount of administrative penalty

报告期内某地区卫生监督部门实施卫生行政处罚行为中给予违法行为单位或个人的罚款金额，是反映卫生监督情况的统计指标之一。

5.5.13 吊销卫生行政许可证份数 number of health administrative license revoked

报告期内某地区实施卫生行政处罚行为中给予吊销卫生行政许可证的份数，是反映卫生监督情况的统计指标之一

5.5.14 行政复议案件数 number of administrative reconsideration cases

报告期内某地区依据法律法规实施卫生行政处罚、行政强制及其他具体行政行为的案件中引起行政复议的案件数，是反映卫生监督情况的统计指标之一

5.5.15 行政诉讼案件数 number of administrative cases

报告期内某地区依据法律法规实施卫生行政处罚、行政强制及其他具体行政行为的案件中引起行政诉讼的案件数，是反映卫生监督情况的统计指标之一

5.6 医疗服务指标

5.6 医疗服务指标 indicators of health care services

反映一定时期、一定地区医疗服务利用、医疗服务效率、医疗服务质量与安全、病人医药费用等统计指标

5.6.1 总诊疗人次 total number of visits

报告期内某地区所有诊疗活动的总人次，包括医疗卫生机构的门诊、急诊、出诊、单项健康检查、健康咨询指导人次。是反映医疗服务情况的统计指标之一

5.6.2 门急诊人次 number of outpatient and emergency visits

报告期内某地区医疗卫生机构的门诊和急诊人次之和。是反映医疗服务情况的统计指标之一

5.6.3 预约诊疗人次 number of clinic appointments

报告期内某地区患者采用网上、电话、院内登记、双向转诊等方式成功预约诊疗人次之和，含中医。是反

映医疗服务情况的统计指标之一

5.6.4 健康检查人次 number of health examinations

报告期内某地区医疗机构全身体检人次和体检中心全身及单项健康检查人次之和。是反映医疗服务情况的统计指标之一

5.6.5 中医治未病服务人次 number of patients receiving preventive health care services of traditional Chinese medicine

年内某地区中医治未病科、中医治未病中心的门诊服务人次之和，是反映医疗服务情况的统计指标之一。

5.6.6 使用中药饮片的出院人数占比 proportion of discharged patients using Chinese herbal medicine decoction pieces

报告期内使用中药饮片的出院人数占同类机构出院

人数的比例。是反映医疗服务情况的统计指标之一

5.6.7 门诊中医非药物治疗诊疗人次占比 proportion of visits using non-pharmaceutical therapy

报告期内门诊中医非药物治疗诊疗人次数（以挂号人次计）占门诊人次数之比。是反映医疗服务情况的统计指标之一

5.6.8 居民年平均就诊次数 average number of annual visits per capita

年内某地区医疗卫生机构的门诊、急诊、出诊、单项健康检查、健康咨询指导人次与同年该地区常住人口数之比。是反映医疗服务情况的统计指标之一

5.6.9 两周就诊率 two-week medical consultation rate

调查前两周内居民因病或身体不适到医疗机构就诊的人次数与调查人口数之比。是反映医疗服务情况的统计指标之一

5.6.10 两周患病未就诊率 two-week morbidity without medical consultations

调查前两周内居民患病而未就诊的人次数与两周患病总人次数之比。是反映医疗服务情况的统计指标之一

5.6.11 非公医疗机构门诊量占门诊总量的比例 number of visits in non-public health institutions as a percentage of total visits

年内某地区非公医疗机构门诊量与医疗卫生机构门诊总量之比。是反映医疗服务情况的统计指标之一

5.6.12 入院人数 number of inpatients

某年某地区居民到医疗卫生机构住院的总人次数，包括已办理入院手续的住院人次数和未办理入院手续而实际住院的人次数，用以反映医疗服务利用情况

5.6.13 居民年住院率 annual hospitalization rate

某年某地区医疗卫生机构的入院人数占同年该地区常住人口数的比例，通常以百分率表示，反映一年内某地区每百名居民的住院次数，属于医疗服务利用指标之一。

5.6.14 每百门急诊入院人数 number of inpatients per 100 outpatients

某年某地区每百名门急诊患者中接受入院治疗的人次数，数值上等于某年某地区医疗卫生机构的入院人数与同年该地区医疗卫生机构的门急诊人次数之比的一百倍。其中门急诊人次数包括门诊和急诊人次数之和。

5.6.15 出院人数 number of discharged patients

某年某地区医疗卫生机构所有住院后出院的人次数，包括医嘱离院、医嘱转其他医疗机构、非医嘱离院、死亡及其他原因离院的人次数，不含家庭病床撤床人次数，用以反映医疗服务利用情况。

5.6.16 出院病人疾病构成 distribution of disease in discharged patients

某年某地区某类疾病出院人数占同年该地区出院人数的百分比，用以反映医疗服务利用情况。

5.6.17 民营医院住院量占医院住院量的比例 number of inpatients in non-public hospital as a percentage of hospital inpatients

某年某地区某地区民营医院出院人数占同年该地区医院出院人数的百分比，用以反映医疗服务利用情况。

5.6.18 血液采集总量 total number of blood units collected

某年某地区一般血站采集的血液总单位数，包括全血单位数，红细胞类、血小板类和血浆类血液成分的单位数，用以反映医疗服务利用情况。

5.6.19 临床用血总量 total number of blood units transfused

某年某地区各级各类医疗卫生机构临床用血总单位数，包括全血单位数，红细胞类、血小板类和血浆类血液成分的单位数，用以反映医疗服务利用情况。

5.6.20 报废血液总量 total number of blood units discarded

某年某地区因血液筛查结果不合格，血液过期，或血液采集、制备、储存和运输过程中由物理或化学因素造成的血液破损、变性、非标准剂量等，而导致的不能用于临床输注的血液总单位数，用以反映医疗服务利用情况。

5.6.21 血液检测不合格率 unqualified rate of blood tested

某年某地区血液筛查检测结果不合格数占同年该地区血液筛查检测总数的比例，通常以百分率表示，反映一年内某地区每百份被查血液样本中筛查结果不合格数，用以反映医疗服务利用情况。

5.6.22 献血人次数 number of blood donors

某年某地区采供血机构为献血者提供服务，并完成血液采集的总人次数，包括完成全血，机采血小板、红细胞类或血浆类血液成分捐献的献血者，用以反映医疗服务利用情况。

5.6.23 千人口献血人数 number of blood donors per 1000 population

某年某地区每千人口中的献血者人数，数值上等于某年某地区献血者人数与同年该地区常住人口数之比的一千倍。其中献血者人数是指完成1次（含）以上全血、机采血小板、红细胞类或血浆类血液成分捐献的献血者总数。

5.6.24 人均用血量 number of blood units transfused per capita

某年某地区人均临床用全血和红细胞类血液成分的量，数值上等于某年某地区临床用血总量与同年该地区常住人口数的比值，用以反映医疗服务利用情况。

5.6.25 住院病人手术人次数 number of inpatients receiving surgery

某年某地区住院病人中施行手术和操作的人次数。1次住院期间施行多次手术的，按实际手术次数统计；1次实施多个部位手术的按1次统计。

5.6.26 病床使用率 occupancy rate of hospital beds

报告期内某地区医疗卫生机构实际占用总床日数与实际开放总床日数之比。实际开放总床日数指报告期内医疗卫生机构各科每天24点开放病床数总和；实际占用总床日数指同期医疗卫生机构各科每天24点实际占用的病床数之和。

5.6.27 平均住院日 average length of stay in hospital

报告期内某地区医疗卫生机构出院者占用总床日数与同期该地区医疗卫生机构出院人数之比，反映报告期内某地区平均每个出院者占用的住院床日数。出院者占用总床日数指报告期内医疗卫生机构所有出院者的住院床日之和。

5.6.28 病床周转次数 turnover rate of hospital beds

某年某地区医疗卫生机构出院人数与同年该地区医疗卫生机构平均开放病床数之比。其中平均开放病床数指某年度医疗卫生机构实际开放总床日数与同年度工作日数之比。工作日数即日历日数扣除节假日数。

5.6.29 医师日均担负诊疗人次 number of visits per doctor per day

报告期内某地区医疗卫生机构平均每位执业（助理）医师每日担负的诊疗人次数。其中诊疗人次数包括医疗卫生机构的门诊、急诊、触诊、单向健康检查、健康咨询指导人次数。

5.6.30 医师日均担负住院床日 number of inpatients per doctor per day

报告期内某地区医疗卫生机构平均每位执业（助理）医师每日担负的实际占用床日数，反映医疗服务效率。

5.6.31 急救车出车率 ambulance dispatching rate

某年某地区急救机构出动急救车的次数占同年该地区急救呼叫次数的比例，通常以百分率表示，反映医疗服务效率。

5.6.32 院前急救服务网络平均反应时间 average reaction time of pre-hospital emergency service network

某年某地区急救机构的每次急救人员到达时间与患者发出急救呼救时间之差的总和与同年该地区急救机构的总急救次数之比，反映一年内某地区自患者发出急救呼救请求至急救人员到达呼救患者驻地的平均时间。

5.6.33 急诊病死率 mortality rate of emergency patients

某年某地区医疗卫生机构急诊死亡人数占同年该地区医疗卫生机构急诊人次数的比例，通常以百分率表示，反映医疗服务质量与安全。

5.6.34 住院病死率 mortality rate of inpatients

某年某地区医疗卫生机构住院死亡人数占同年该地区医疗卫生机构出院人数的比例，通常以百分率表示，反映医疗服务质量与安全。

5.6.35 入院与出院诊断符合率 coincidence rate of admitting and discharge diagnosis

某年某地区医疗卫生机构入院与出院诊断符合人数占出院人数的比例。其中入院与出院诊断符合人数指《住院病案首页》中主要诊断的“入院病情”为“有”或“临床未确定”的人数。

5.6.36 临床与病理诊断符合率 coincidence rate of clinical and pathological diagnosis

某年某地区医院病理诊断与出院诊断符合人数占同年该地区医院出院病理检查人数的比例，通常以百分率表示，反映医疗服务质量与安全。

5.6.37 某类疾病院内感染率 hospital infection rate

某年某地区某类疾病发生院内感染人次数占同年该地区该类疾病出院人次数的比例。院内感染包括在住院期间发生的感染和在医院内获得出院后发生的感染、医院工作人员在医院内获得的感染。

5.6.38 I类切口甲级率 healing rate of type I incision

某年某地区医院I类切口甲级愈合例数占同年该地区医院I类切口愈合例数的比例，用以反映医疗服务质量与安全。

5.6.39 I类切口感染率 infection rate of type I incision

某年某地区医院I类切口丙级愈合例数占同年该地区医院I类切口愈合例数的比例，反映医疗服务质量与安全。

5.6.40 医疗纠纷例数 number of medical dispute cases

某年由医疗卫生机构相关管理部门受理的医疗纠纷例数，包括门诊和住院发生的医疗纠纷例数，反映医疗服务质量与安全。

5.6.41 医疗事故例数 number of medical malpractice cases

某年由医疗事故鉴定机构依据《医疗事故处理条例》鉴定的事故例数，按鉴定日期统计，用以反映医疗服务质量与安全。

5.6.42 基本医疗保险收入占医疗收入比重 basic medical insurance revenue as a percentage of medical income

报告期内医疗卫生机构三项基本医疗保险收入占同

期医疗卫生机构医疗收入总额的比例。其中三项基本医疗保险收入包括城镇职工基本医疗保险收入、城镇（城乡）居民基本医疗保险收入、新农合补偿收入。

5.6.43 门诊病人次均医药费用 average medical expense per outpatient visit

报告期内医疗收入中的门诊收入与健康检查收入之差与同期总诊疗人次数的比值，反映报告期内门诊病人平均每次就诊医药费用。

5.6.44 门诊病人次均药费 average medication expense per visit

报告期内门诊药品收入与同期总诊疗人次之比，反映报告期内门诊病人平均每次就诊药费。

5.6.45 住院病人次均医药费用 average medical expense per inpatient

报告期内出院者住院医药费用与同期出院人次之比，反映报告期内出院者平均每次住院医药费用。

5.6.46 住院病人次均药费 average medication expense per inpatient

报告期内出院者住院药费与同期出院人次之比，反映报告期内出院者平均每次住院药费。

5.6.47 住院病人日均医药费用 average medical expense of inpatients per day

报告期内出院者医药费用总额与同期出院者住院天数之比，反映报告期内住院病人平均每日医药费用。

5.6.48 病种住院费用 average inpatients medical expense by disease

报告期内医院某病种住院医药费用与同期医院该病种出院人次之比，反映报告期内某种疾病出院者平均每次医药费用。

5.6.49 病人医药费用构成 composition of patient medical expense

报告期内门诊或住院收入中某项收入占同期门诊或住院收入的百分比。其中医药费用分类采用《医院会计制度》或《基层医疗卫生机构会计制度》。

5.6.50 病人医药费用增长率 growth rate of patient medical expense

报告期内门诊或住院病人医药费用与上期门诊或住院病人医药费用之差与上期上期门诊或住院病人医药费用的比值，通常以百分率表示。

5.7 药品与卫生材料供应保障指标

5.7 药品与卫生材料供应保障指标 indicators of drugs and medical material supply

反映一定时期、一定地区药品与卫生材料生产、流通、采购、配送、使用、质量与安全等的统计指标。

5.7.1 药品注册数 number of registered drugs

年末经国家食品药品监管部门审查并批准上市的药品品种总数，按批件数统计，是反映药品生产和流统的指标之一。

5.7.2 中药保护品种数 number of protected traditional Chinese medicine

包括初次保护品种数、延长保护品种数及同品种保护数，是反映药品生产和流统的指标之一。

5.7.3 《药品生产许可证》持证企业数 number of enterprises with drug production licenses

包括报告期内发给、年前发给并在有效期内的《药品生产许可证》的持证企业数，是反映药品生产和流统的指标之一。

5.7.4 《药品经营许可证》持证企业数 number of enterprises with drug business licenses

年末省级食品药品监管部门发给药品批发企业和地县级食品药品监督管理部门发给药品零售企业的《药品经营许可证》数之和，按报告期内发给、年前发给且在有效期内的证书统计，是反映药品生产和流统的

指标之一。

5.7.5 通过 GMP 认证的药品生产企业数 number of drug manufacturing enterprises with good manufacturing practice certifications

年末按照《药品生产质量管理规范认证管理办法》的规定，经现场检查和审核符合《药品生产质量管理规范》要求，国家食品药品监管部门发给《药品生产质量管理规范证书》的药品生产企业数，是反映药品生产和流统的指标之一。

5.7.6 药品生产企业主营业务收入 main business income of drug manufacturing enterprises

报告期内某地区药品生产企业销售商品、提供劳务等主营业务收入，是反映药品生产和流统的指标之一。

5.7.7 药品生产企业利润总额 total profits of drug manufacturing enterprises

报告期内某地区药品生产企业在生产经营过程中各种收入扣除各种耗费后的盈余，是反映药品生产和流统的指标之一。

5.7.8 国家基本药物目录品种数 number of drug varieties in national essential drug list

国家卫生健康委颁布的《国家基本药物目录》内的药品品种数，是反映药品采购和配送的指标之一。

5.7.9 省级增补药品品种数 number of supplemented

provincial essential drug varieties in essential drug lists

年末省级人民政府确定的《国家基本药物目录》以外的、供基层医疗卫生机构使用的药品品种数，是反映药品采购和配送的指标之一。

5.7.10 参与省级药品集中采购的药品生产企业数
number of drug manufacturing enterprises involved in
provincial centralized drug procurement

按照年末省级药品集中采购平台上注册的生产企业用户数统计的药品生产企业数，是反映药品采购和配送的指标之一。

5.7.11 参与省级药品集中采购的医疗卫生机构数
number of health care institutions involved in provincial
centralized drug procurement

按照省级药品集中采购平台统计的年末在省级药品集中采购平台注册并通过平台采购药品的各级各类医疗卫生机构数，是反映药品采购和配送的指标之一。

5.7.12 参与医疗卫生机构药品配送的企业数 number
of enterprises involved in medical institutions drug
delivery

年末取得食品药品监管部门《药品经营许可证》并将药品配送到医疗卫生机构的药品经营企业数，是反映药品采购和配送的指标之一。

5.7.13 药品3日配送到位率（按金额） ratio of
delivered amount to ordered amount of drugs in 3 days

一定时期内配送企业收到医疗机构采购订单后3天内送达的药品金额与同期配送企业收到医疗机构采购订单金订单金额之比，是反映药品采购和配送的指标之一。

5.7.14 医疗机构30天内回款率 ratio of payment to
warehoused value in 30 days

一定时期内医疗机构在30天内结算药款金额与同期入库金额之比，是反映药品采购和配送的指标之一。

5.7.15 药品费用 drug expenditure

某年某地区用于治疗 and 预防人类疾病的药品费用总额，包括辖区内医疗卫生机构门诊和住院药品收入、零售药店的药品销售额。是反映药品使用情况的指标之一。

5.7.16 人均药品费用 drug expenditure per capita

某年某地区药品费用与同年该地区平均人口数之比。是反映药品使用情况的指标之一。

5.7.17 药品费用占卫生总费用的比重 drug
expenditure as a percentage of total health expenditure

某年某地区药品费用与同年该地区卫生总费用之比。是反映药品使用情况的指标之一。

5.7.18 国家基本药物使用金额比例 utility rate of
national essential drugs

一定时期内医疗机构使用的药品中，国家基本药物所占的比例。由医疗机构国家基本药物采购金额与同期医疗机构药品采购总金额之比计算得出。是反映药品使用情况的指标之一。

5.7.19 药品电子监管系统覆盖率 coverage of
electronic drug supervision system

某年某地区药品生产经营企业和医疗卫生机构中纳入国家药品电子监管系统的企业或机构数所占比例。是反映药品质量管理的指标之一。

5.7.20 每百万人口药品不良反应报告例数 number of
reported cases of adverse drug reaction per 1000000
population

某年药品不良反应报告例数与同年该地区常住人口数之比，通常用百万分之一表示。是反映药品质量管理的指标之一。

5.7.21 查处药品案件数 number of drug cases
investigated and punished

某年某地区食品药品监管部门依据《药品监督行政处罚程序规定》查处药品生产、经营企业违法药品案件数之和。是反映药品质量管理的指标之一。

5.7.22 取缔药品无证经营户数 number of banned
enterprises without drug business licenses

年内某地区药品监督管理部门依据《药品监督行政处罚程序规定》，取缔无《药品生产经营许可证》或《药品生产经营许可证》的药品生产、经营企业活动的户数。是反映药品质量管理的指标之一。

5.7.23 卫生材料生产企业主营业务收入 main
business income of medical material manufacturing
enterprises

某年某地区卫生材料生产企业销售商品、提供劳务等主营业务收入。是反映卫生材料供应的指标之一。

5.7.24 卫生材料生产企业利润总额 profits of medical
material manufacturing enterprises

某年某地区卫生材料生产企业在生产经营过程中各种收入扣除各种耗费后的盈余。是反映卫生材料供应的指标之一。

5.7.25 参与省级集中采购的高值医用耗材生产企业数
number of high-value medical material manufacturing
enterprises involved in provincial centralized procurement

年末在省级药品集中采购平台上实际注册的高值医用耗材生产企业数。是反映卫生材料供应的指标之一。

5.7.26 参与高值医用耗材配送的企业数 number of
enterprises involved in high-value medical material
delivery

年末通过公开招标、邀请招标、询价采购、单独议价等方式确定的高值医用耗材配送企业数。是反映卫生

材料供应的指标之一。

5.8 卫生资源指标

5.8 卫生资源指标 indicators of health resources

反映一定时期、一定地区卫生人力、卫生经费、卫生设施等的统计指标。

5.8.1 卫生人员数 number of health personnel

报告期末在医疗卫生机构工作并由单位支付年底工资的在岗职工数之和，包括卫生技术人员、乡村医生和卫生员、其他技术人员、管理人员和工勤技能人员。是反映卫生人力的指标之一。

5.8.2 卫生技术人员数 number of health technical personnel

报告期末医疗卫生机构中执业医师、执业助理医师、注册护士、药师士、检验及影像技师士、卫生监督员和见习医（药、护、技）师（士）等卫生专业人员之和。是反映卫生人力的指标之一。

5.8.3 执业（助理）医师数 number of licensed physicians and physician assistants

报告期末医疗卫生机构中取得医师执业证书且实际从事医疗、妇幼保健、疾病防治等工作的执业医师和执业助理医师数之和。是反映卫生人力的指标之一。

5.8.4 全科医生数 number of general practitioners

报告期末医疗卫生机构中取得医师执业证书且执业范围为全科医学专业的执业（助理）医师数，以及基层医疗卫生机构取得全科医生转岗培训、骨干培训、岗位培训和住院医师规范化（全科医生）培训合格证的执业助理医师数之和。是反映卫生人力的指标之一。

5.8.5 注册护士数 number of registered nurses

报告期末医疗卫生机构中取得注册护士证书且实际从事护理工作的人员之和。是反映卫生人力的指标之一。

5.8.6 卫生监督员数 number of health supervisors

报告期末取得卫生监督员证且从事各类卫生监督执法的人员之和。是反映卫生人力的指标之一。

5.8.7 乡村医生数 number of village doctors

年末基层医疗卫生机构中从当地卫生行政部门取得乡村医生证书的人数之和。是反映卫生人力的指标之一。

5.8.8 管理人员数 number of administrative staffs

年末医疗卫生机构中担负领导职责或管理任务的工作人员之和。包括从事医疗、公共卫生、医学科研与教学等业务管理工作的人员以及主要从事党政、人事、

财务、统计、信息、安全保卫等行政管理工作的人员。是反映卫生人力的指标之一。

5.8.9 工勤技能人员数 number of logistics technical workers

年末医疗卫生机构中承担技能操作和维护、后勤保障、服务等职责并获得技能等级证书的工作人员之和。包括技术工和普通工，技术工包括护理、药剂、检验、收费、挂号等工作人员。是反映卫生人力的指标之一。

5.8.10 每千人口卫生技术人员数 number of health technical personnel per 1000 population

年末每千人口拥有的卫生技术人员数，用年末卫生技术人员数除以年末常住人口数计算得出，通常用千分之一表示。是反映卫生人力的指标之一。

5.8.11 每千人口执业（助理）医师数 number of licensed physicians and physician assistants per 1000 population

年末每千人口拥有的执业医师和执业助理医师数，用年末执业医师数及执业助理医师数除以年末常住人口数计算得出，通常用千分之一表示。是反映卫生人力的指标之一。

5.8.12 每万人口全科医生数 number of general practitioners per 10000 population

年末每万人口拥有的全科医生数，用年末全科医生数除以年末常住人口数计算得出，通常用万分之一表示。是反映卫生人力的指标之一。

5.8.13 每万人口口腔医师数 number of dentists per 10000 population

年末每万人口拥有的口腔医师数，用年末口腔医师数除以年末常住人口数计算得出，通常用万分之一表示。是反映卫生人力的指标之一。

5.8.14 每千人口注册护士数 number of registered nurses per 1000 population

年末每千人口拥有的注册护士数，用年末医疗卫生机构注册护士数除以年末常住人口数计算得出，通常用千分之一表示。是反映卫生人力的指标之一。

5.8.15 每万人口药师（士）数 number of pharmacists per 10000 population

年末每万人口拥有的药师（士）数，用年末医疗卫生机构药师（士）数除以年末常住人口数计算得出，通常用万分之一表示。是反映卫生人力的指标之一。

5.8.16 每千农村人口乡镇卫生院人员数 number of township health centre personnel per 1000 rural population

86

公开征求意见期限：

2025年7月8日——2025年10月8日

- 年末每千农村人口拥有的乡镇卫生院人员数。是反映卫生人力的指标之一。
- 5.8.17 每千农村人口村卫生室人员数** number of village clinic staffs per 1000 rural population
年末每千农村人口拥有的村卫生室人员数。是反映卫生人力的指标之一。
- 5.8.18 每万人口公共卫生人员数** number of public health workers per 10000 population
年末每万人口拥有的专业公共卫生机构人员数。是反映卫生人力的指标之一。
- 5.8.19 医护比** ratio of physicians to nurses
年末执业（助理）医师数与注册护士数的比值。是反映卫生人力的指标之一。
- 5.8.20 继续医学教育比例** percentage of physicians received continuing medical education
年末主治医师中接受过继续医学教育人数所占比例。是反映卫生人力的指标之一。
- 5.8.21 医学专业招生数** number of medical specialty recruitments
某学年初实际注册的医学专业新生数。是反映卫生人力的指标之一。
- 5.8.22 医学专业毕业人数** number of medical specialty graduates
某学年末取得毕业证书的医学专业毕业生数。是反映卫生人力的指标之一。
- 5.8.23 医学在校生数** number of students in medical schools
某学年初某地区具有学籍的医学专业学生数之和。是反映卫生人力的指标之一。
- 5.8.24 卫生总费用** total health expenditure
某年某地区用于医疗卫生保健服务的资金总量，包括政府卫生支出、社会卫生支出和个人现金卫生支出。是反映卫生经费的指标之一。
- 5.8.25 卫生总费用构成** constitution of total health expenditure
某年政府卫生支出、社会卫生支出、个人现金卫生支出占卫生总费用的比例。是反映卫生经费的指标之一。
- 5.8.26 个人现金卫生支出占卫生总费用的比重** out-of-pocket expenditure on health as a percentage of total health expenditure
某年某地区个人现金卫生支出占卫生总费用的比例。是反映卫生经费的指标之一。
- 5.8.27 政府医疗卫生支出增长率** growth rate of government health expenditure
某年某地区政府医疗卫生支出较上年增长百分比。是反映卫生经费的指标之一。
- 5.8.28 政府医疗卫生支出占财政支出的比重** general government health expenditure as a percentage of total government expenditure
某年财政支出中医疗卫生支出所占比例。是反映卫生经费的指标之一。
- 5.8.29 人均卫生费用** per capita health expenditure
某年某地区卫生总费用与该年该地区年平均人口数的比值。是反映卫生经费的指标之一。
- 5.8.30 卫生总费用占国内生产总值比例** total health expenditure as a percentage of gross domestic product
某年某地区卫生总费用占国内生产总值（GDP）的比例。是反映卫生经费的指标之一。
- 5.8.31 卫生消费弹性系数** health consumption elasticity coefficient
某年某地区卫生总费用增长速度与国内生产总值（GDP）增长速度之比。是反映卫生经费的指标之一。
- 5.8.32 人均基本公共卫生服务补助经费** per capita subsidy of basic public health service
某年某地区平均每人中央和地方财政拨付的基本公共卫生服务项目经费的补助。是反映卫生经费的指标之一。
- 5.8.33 财政补助收入占医疗卫生机构总支出的比例** financial subsidy as a percentage of total expenditure of health institutions
某年某地区医疗卫生机构财政补助收入占总支出的比例。是反映卫生经费的指标之一。
- 5.8.34 医疗卫生机构总收入** total income of health institutions
某年某地区医院、基层医疗卫生机构、专业公共卫生机构、其他医疗机构的总收入之和。是反映卫生经费的指标之一。
- 5.8.35 医疗卫生机构总收入构成** constitution of total income of health institutions
某年某地区医疗卫生机构总收入中不同来源收入所占比例。是反映卫生经费的指标之一。
- 5.8.36 医疗收入构成** constitution of medical revenue
某年某地区医疗卫生机构医疗收入中各项收入所占比例。是反映卫生经费的指标之一。
- 5.8.37 医疗卫生机构总支出** total expenditure of health institutions
某年某地区医院、基层医疗卫生机构、专业公共卫生机构、其他医疗机构的总支出之和。是反映卫生经费的指标之一。
- 5.8.38 医疗卫生机构总支出构成** constitution of total expenditure of health institutions
某年某地区医疗卫生机构总支出中各项支出所占比例。

- 例。是反映卫生经费的指标之一。
- 5.8.39 医疗业务成本构成** constitution of medical business cost
某年某地区医院医疗业务成本中各项支出所占比例。是反映卫生经费的指标之一。
- 5.8.40 医疗支出构成** constitution of total medical expenditure
某年某地区基层医疗卫生机构医疗支出中各项支出所占比例。是反映卫生经费的指标之一。
- 5.8.41 医疗卫生机构总资产** total assets of health institutions
年末某地区医院、基层医疗卫生机构、专业公共卫生机构和其他医疗卫生机构等事业单位的总资产之和。是反映卫生经费的指标之一。
- 5.8.42 医疗卫生机构负债** liabilities of health institutions
年末某地区医院、基层医疗卫生机构、专业公共卫生机构和其他医疗卫生机构等事业单位的负债之和。是反映卫生经费的指标之一。
- 5.8.43 医疗卫生机构净资产** net assets of health institutions
年末某地区医疗卫生机构资产减去负债后的余额。是反映卫生经费的指标之一。
- 5.8.44 业务收支结余率** ratio of business balance to revenue
某年某地区医院业务收支结余与医疗收入、财政基本支出补助收入、其他收入之和的比值。是反映卫生经费的指标之一。
- 5.8.45 百元收入药品及卫生材料消耗** drug and medical material consumption per 100 yuan revenue
某年某地区医院每百元收入中，药品及卫生材料费用所占比例。是反映卫生经费的指标之一。
- 5.8.46 平均每床固定资产** fixed assets per bed
年末某地区医院固定资产原值与实有床位数的比值。是反映卫生经费的指标之一。
- 5.8.47 总资产周转率** turnover rate of total assets
某年某地区医院医疗收入、其他收入之和与平均总资产的比值。是反映卫生经费的指标之一。
- 5.8.48 资产负债率** asset-liability ratio
年末某地区医院负债总额与资产总额的比值。是反映卫生经费的指标之一。
- 5.8.49 流动比率** current ratio
年末某地区医院流动资产与流动负债的比值。是反映卫生经费的指标之一。
- 5.8.50 人员经费支出比率** personnel expenses ratio
某年某地区医院人员经费与医疗业务成本、管理费用、其他支出之和的比值。是反映卫生经费的指标之一。
- 5.8.51 在岗职工年平均工资** annual average salary of on-post staffs
某年某地区医疗卫生机构在岗职工年平均工资收入。是反映卫生经费的指标之一。
- 5.8.52 绩效工资所占比例** performance-related salary as a percentage of total salary
某年某地区医疗卫生机构绩效工资占工资总额的比值。是反映卫生经费的指标之一。
- 5.8.53 管理费用率** administrative cost ratio
某年某地区医院管理费用与医疗业务成本、管理费用、其他支出之和的比值。是反映卫生经费的指标之一。
- 5.8.54 药品和卫生材料支出率** medication and health material expenditure ratio
某年某地区医疗卫生机构药品费、卫生材料费之和与医疗业务成本、管理费用、其他支出之和的比值。是反映卫生经费的指标之一。
- 5.8.55 药品收入占医疗收入比重** medication revenues as a percentage of medical revenue
某年某地区医疗卫生机构药品收入占医疗收入的比重，是反映卫生经费的指标之一。
- 5.8.56 总资产增长率** growth rate of total assets
某地区医院年末总资产、年初总资产之差与年初总资产的比值，是反映卫生经费的指标之一。
- 5.8.57 净资产增长率** growth rate of net assets
某地区医院年末净资产、年初净资产之差与年初净资产的比值，是反映卫生经费的指标之一。
- 5.8.58 固定资产净值率** rate of net value of fixed assets
年末某地区医院固定资产净值与固定资产原值的比值，是反映卫生经费的指标之一。
- 5.8.59 公共卫生支出占总支出比例** public health expenditure as a percentage of total expenditure
某年某地区基层医疗卫生机构公共卫生支出与总支出的比值，是反映卫生经费的指标之一。
- 5.8.60 编制床位数** number of authorized beds
年末由卫生行政部门核定的医疗卫生机构床位数。是反映卫生设施的指标之一。
- 5.8.61 床位数** number of beds
年末医疗卫生机构实有床位数，其中实有床位包括正规床、简易床、监护床、超过半年的加床、正在消毒和修理的床位、因扩建或大修而停用的床位。不包括产科新生儿床、接产室待产床、库存床、观察床、临时加床和病人家属陪侍床。是反映卫生设施的指标之一。
- 5.8.62 非公医疗机构床位占医疗机构床位总数的比例** number of private health institution beds as a percentage of

total number of beds

年末某地区非公医疗机构床位数占该地区医疗机构床位总数的比例。是反映卫生设施的指标之一。

5.8.63 民营医院床位数占医院床位数的比例 number of private hospital beds as a percentage of total number of hospital beds

年末民营医院床位数与医院床位总数的比值是反映卫生设施的指标之一。是反映卫生设施的指标之一。

5.8.64 每千人口医疗卫生机构床位数 number of health institution beds per 1000 population

年末每千人口拥有的医疗卫生机构床位数。是反映卫生设施的指标之一。

5.8.65 每千农村人口乡镇卫生院床位数 number of township hospital beds per 1000 rural population

年末某地区每千农村人口拥有的乡镇卫生院床位数。是反映卫生设施的指标之一。

5.8.66 医师与床位之比 ratio of physicians to beds

年末医疗卫生机构执业（助理）医师数与床位数的比值是反映卫生设施的指标之一。

5.8.67 护士与床位之比 ratio of nurses to beds

年末医疗卫生机构注册护士数与实有床位数的比值。是反映卫生设施的指标之一。

5.8.68 万元以上医疗设备台数 number of medical equipment with value of 10,000 yuan and above

年末医疗卫生机构单价超过万元的医用设备数。是反映卫生设施的指标之一。

5.8.69 设备配置率 equipment configuration rate

年末配置某种医用设备的医疗卫生机构数占同类机构总数的比例。是反映卫生设施的指标之一。

5.8.70 房屋建筑面积 housing construction area

年末医疗卫生机构购买、自建或主管部门划拨的且有产权证的房屋建筑面积，但不包括租房面积、供本单位使用无租金也无产权的房屋建筑面积。是反映卫生设施的指标之一。

5.8.71 业务用房面积 operating area

年末医疗卫生机构除职工住宅之外的所有房屋建筑面积，包括医疗服务（急诊、门诊、住院、医技）、公共卫生服务、医学教育与科研、后勤保障、行政管理和院内生活等设施用房，但不包括租房面积。是反映卫生设施的指标之一。

5.8.72 危房所占比例 dilapidated house area as a percentage of operating area

年末医疗卫生机构经过专业机构鉴定后认定的 D 级危房占业务用房面积的比例。是反映卫生设施的指标之一。

5.8.73 每床占用业务用房面积 occupied operating area per bed

年末医疗卫生机构平均每张床位所占用的业务用房面积。是反映卫生设施的指标之一。

5.8.74 房屋竣工面积 floor area of buildings completed

年内房屋建筑按照设计要求，已全部完成，达到了住人或使用条件，经验收鉴定合格，正式移交给使用单位（或建设单位）的各栋房屋建筑面积的总和。是反映卫生设施的指标之一。

5.8.75 当年实际完成投资来源构成 constitution of annual actual investment by source

医疗卫生机构当年实际完成投资额中各投资来源金额所占构成比，其中当年实际完成投资额按投资来源分为财政性投资、单位自有资金、银行贷款。是反映卫生设施的指标之一。

5.8.76 医疗卫生机构数 number of health institutions

报告期末从卫生行政部门取得《医疗机构执业许可证》、《计划生育技术服务许可证》或从民政、工商行政、机构编制管理部门取得法人单位登记证书，为社会提供医疗服务、公共卫生服务或从事医学科研和在职培训等工作的单位数。包括医院、基层医疗卫生机构、专业公共卫生机构、其他医疗卫生机构数。是反映卫生设施的指标之一。

5.8.77 非公医疗机构数 number of non-public health institutions

报告期内登记注册类型为私营、联营、股份合作（有限）、台港澳和中外投资等的医院、基层医疗卫生机构、妇幼保健机构、专科疾病防治机构、疗养院、急救中心（站）数之和。是反映卫生设施的指标之一。

5.8.78 医院数 number of hospitals

报告期末综合医院、中医医院、中西医结合医院、民族医院、各类专科医院和护理院（含医学院校附属医院）数之和，不包括专科疾病防治院、妇幼保健院和疗养院。是反映卫生设施的指标之一。

5.8.79 公立医院数 number of public hospitals

报告期末登记注册类型为国有和集体的医院数之和。是反映卫生设施的指标之一。

5.8.80 民营医院数 number of non-public hospitals

报告期末登记注册类型为国有和集体以外的医院，包括私营、联营、股份合作（有限）、港澳台合资合作、中外合资合作等医院数之和。是反映卫生设施的指标之一。

5.8.81 基层医疗卫生机构数 number of grass-roots health institutions

报告期末社区卫生服务中心、社区卫生服务站、街道卫生院、乡镇卫生院、村卫生室、门诊部、诊所、医务室数之和。是反映卫生设施的指标之一。

5.8.82 政府办基层医疗卫生机构数 number of government-run grass-roots health institutions

报告期末卫生行政部门、街道办事处、生产建设兵团、林业局、农垦局等机关举办等行政机关举办的社区卫生服务中心(站)、乡镇卫生院与街道卫生院数之和。但不包括公立医院举办的社区卫生服务中心和社区卫生服务站(属事业单位举办)。是反映卫生设施的指标之一。

5.8.83 专业公共卫生机构数 number of specialized public health institutions

报告期末疾控中心、专科疾病防治机构、妇幼保健机构、健康教育机构、急救中心(站)、采供血机构、卫生监督机构、计划生育技术服务机构数之和。是反映卫生设施的指标之一。

5.8.84 提供中医服务的基层医疗卫生机构所占比例 number of grass-roots health care institutions providing traditional Chinese medicine services as a percentage of grass-roots health institutions

年末提供中医药服务的社区卫生服务中心数(或社区卫生服务站数、乡镇卫生院数、村卫生室数)占年末同类机构总数的比例。是反映卫生设施的指标之一。

5.8.85 医疗卫生机构基础设施建设达标率 compliance rate of the construction of health institutions

报告期末由主管部门审核达到由上级主管部门按照国家发改委和原卫生部下发的《中央预算内专项资金项目一县医院、县中医院、中心乡镇卫生院、村卫生

室和社区卫生服务中心建设指导意见》审核达到业务用房面积和设备配置标准的某类医疗卫生机构数与同期该类医疗卫生机构总数的比例。是反映卫生设施的指标之一。

5.8.86 村卫生室覆盖率 coverage of village clinics

年末设立卫生室的行政村数占年末行政村总数的比例。是反映卫生设施的指标之一。

5.8.87 实行乡村一体化管理的村卫生室所占比例 number of village clinics implementing integrated management as a percentage of village clinics

年末实行乡村一体化管理的村卫生室数占年末村卫生室总数的比例,其中乡村一体化管理是指按照原卫生部办公厅《关于推进乡村卫生服务一体化管理的意见》的要求,对乡镇卫生院和村卫生室行政业务、药械、财务和绩效考核等方面予以规范的管理体制。是反映卫生设施的指标之一。

5.8.88 实行乡村一体化管理的乡镇卫生院所占比例 number of township hospitals implementing integrated management as a percentage of township hospital

年末实行乡村一体化管理的乡镇卫生院数占年末乡镇卫生院总数的比例,其中乡村一体化管理是指按照原卫生部办公厅《关于推进乡村卫生服务一体化管理的意见》的要求,对乡镇卫生院和村卫生室行政业务、药械、财务和绩效考核等方面予以规范的管理体制。是反映卫生设施的指标之一。

5.9 医疗保障统计指标

5.9 医疗保障统计指标 indicators of medical insurance

反映一定时期、一定地区内基本医保覆盖、基本医保筹资、基金使用与受益的统计指标。

5.9.1 新型农村合作医疗 new rural cooperative medical scheme

由政府组织、引导、支持,农民自愿参加,个人、集体和政府多方筹资,以大病统筹为主的农民医疗互助共济制度。通过中央财政补助、地方财政补助、集体扶持和农民个人缴费等渠道筹集资金,主要对农民住院及大病医疗费用给予补偿。是反映基本医保覆盖的指标之一。

5.9.2 新型农村合作医疗参合人数 number of enrollees of the new rural cooperative medical scheme

根据本地新型农村合作医疗的实施方案,到本年度新型农村合作医疗筹资截止时已缴纳参加新型农村合作医疗资金的人口数。是反映基本医保覆盖的指标之

一。

5.9.3 新型农村合作医疗参合率 rate of enrollment of the new rural cooperative medical scheme

本年度某地区新型农村合作医疗参合人数占上年末该地区农业人口数的比例,其中农业人口数系当地统计局数字,部分地区为应参合人数(即农业人口数一参加城镇居民基本医保或城镇职工基本医保的农民及学生数)。是反映基本医保覆盖的指标之一。

5.9.4 新型农村合作医疗当年筹资总额 fund raised by the new rural cooperative medical scheme for the current year

本年度某地区筹集的、实际进入新型农村合作医疗专用账户的基金数额,包括本年度中央及地方财政配套资金、农民个人缴纳资金(含民政部门及其他相关部门代缴的救助资金)、新型农村合作医疗基金本年度产生的全部利息收入及其他渠道实际筹集到的新型农村合作医疗基金额,不含上年结转资金。是反映基

本医保筹资的指标之一。

5.9.4 新型农村合作医疗人均筹资水平 the new rural cooperative medical scheme premium of per capita

本年度某地区平均每一参合农民的实际筹资水平。是反映基本医保筹资的指标之一。

5.9.5 新型农村合作医疗当年基金支出总额 the aggregate spending of the new rural cooperative medical scheme fund for the current year

本年度某地区实际从新型农村合作医疗基金账户中支出用于参合农民补偿的金额。是反映医保基金使用与受益的指标之一。

5.9.6 新型农村合作医疗当年基金使用率 utility rate

of the new rural cooperative medical scheme fund for the current year

本年度某地区新型农村合作医疗基金支出总额占实际筹资总额的比例。是反映医保基金使用与受益的指标之一。

5.9.7 新型农村合作医疗受益总人数 number of person times claimed benefits from the new rural cooperative medical scheme

本年度某地区参合农民从新型农村合作医疗中获得受益的总人数。是反映医保基金使用与受益的指标之一。

