

S T A T I S T I C S

Theory and Application Vol. 2

統計學

理論與應用 下

程大器 編著

智勝
BEST-WISE

統計學

理論與應用

下

程大器 編著

Theory and Application Vol. 2

智勝文化

再版序

統計學是一門理論與實務應用並重之學科，其理論基礎包含了代數、微積分及或然率等基礎數學，而其應用更是廣泛，可說上至天文、下至地理，一切人文社會與自然生物等科學皆可應用。古人常言：「工欲善其事，必先利其器」，是故進行任何研究工作，必須先研究各種研究方法，如此不僅事半功倍，更對科學領域的開拓有長足進步之影響。

一般人皆認為統計學是一門深奧難懂的學科，因此學習統計的態度總是以背誦公式為主要學習方法，而對於其理論基礎卻不予理會，因此在實務應用中，常會有不知要使用哪種類型的統計方法之挫折感。事實上，學習統計應該理論與應用並重，才能相得益彰，但對於大多數的學習者而言，由於基礎數學的程度較差，在學習統計理論部分時，常會感到窒礙難行，因此常會放棄統計理論部分，而只學習統計應用方法，如此惡性循環下去，更造成了統計方法之濫用，進一步導致錯誤之決策。有鑑於此，筆者為了讓學習者對統計方法之學習有系統性之學習架構，因此才有編寫本書之動機。

本書共分為上、下兩冊，上冊內容主要部分為敘述統計與機率論。敘述統計即大量觀察統計資料之統計分析，包括資料之蒐集、整理與分析（見第一及第二章），而本書上冊主要精華在於機率論部分，由於一般商用統計學書籍皆有共同缺失，即機率論之部分描述太少，使讀者在學習推論統計學時常會造成困擾，不知其推論方法之理論基礎，因此為了使讀者能對上述之缺失能予以改進，故本書對於機率理論部分有詳細的描述（見第三及第八章），且編排之方式由淺入深，而對於涉及基礎數學部分更是詳細交代其來源及觀念，對於每個

理論之證明皆盡量以較淺顯之方式交代，使讀者能消除學習統計學理論部分之恐懼。

在一般大學統計教學中，由於教學時間有限，因此常常以些許之時間來討論機率部分，而大部分時間以探討推論及應用統計為主。但學生如以此方式學習統計，比較不容易深入瞭解統計方法之奧妙，況且在一般就業及研究所考試時，也無法應付大量之機率理論部分之考題，因此學習統計學若想要融會貫通，對於機率部分之研習一定要詳細，如此在探討推論及應用統計時，才能瞭解統計方法之架構與應用。在此必須強調一點，即機率與微積分等基礎數學對統計均是工具，而學習統計主要的是分析資料及實務應用，因此所研究或討論的內容均以統計為目標，切勿將機率與統計混淆，而捨本逐末。

本書下冊內容就是以統計方法為基礎，讓讀者能有效地將統計方法應用在日常生活中，而其內容主要區分為推論及應用兩大部分。在推論統計中，主要是將上冊所學之知識用來發展一些推論方法，內容以估計與檢定為主，此內容也是統計學之主要精華，若讀者能瞭解其架構與觀念，則對於學習應用部分之變異數分析及迴歸分析，將會有莫大之助益。變異數分析及迴歸分析屬於統計應用部分，由於其應用廣泛，而且成效卓著，因此在實務上是最常用之分析工具。而在下冊內容中，最後將探討無母數統計方法，此部分是統計在近幾年來發展最快速之領域，也是一般升學考試較容易疏忽之部分，且在實務應用中也是常見之分析方法，讀者對於各種方法要釐清觀念，以方便往後應用。

本書因為考慮統計學之整體性與應用性，故其內容較多且廣泛，因此對於初學者在學習本書內容所花費之時間將會較久，但若讀者對於一些機率論的基礎理論證明、二維隨機變數之機率分配以及變數變換單元無興趣，可直接略過，如此對於數學基礎不好之同學在閱讀本

書時，也能快速且有效率地達到學習的效果。

筆者以多年教學經驗編寫本書，希望能對莘莘學子學習統計有所幫助，因此內容盡量翔實，但筆者學識與經歷均屬有限，故對於遺漏或謬誤之處盼請見諒，尚請各方先進及讀者不吝賜教。

最後，感謝周光凱老師給予書籍內容之指正及編排順序上之建議，以及感謝智勝文化事業所有同仁之支持與協助，使本書不論在編排及版面設計上，皆能引人入勝。

程大器

2005 年 9 月

目錄

Contents

再版序

第九章 估計 1

9.1 緒論 2

9.2 點估計式之評判標準 4

9.3 尋找點估計式之方法 33

習題 55

第十章 區間估計 59

10.1 緒論 60

10.2 母體平均數之信賴區間 66

10.3 兩母體平均數 $\mu_1 - \mu_2$ 之區間估計 82

10.4 單一母體比例之信賴區間 102

10.5 二獨立母體比例差之信賴區間 110

10.6 母體變異數之信賴區間 112

習題 121

第十一章 假設檢定 (一) 125

11.1 緒論 126

11.2 假設檢定之基本觀念 129

11.3 假設檢定之步驟 151

11.4 母體平均數之檢定——以常態分配處理 158

11.5 母體平均數之檢定——以 T 分配處理 171

11.6 樣本數 n 之規劃——在 α 、 β 控制下 175

11.7 兩母體平均數差 $\mu_1 - \mu_2$ 之檢定 183

習題 198

第十二章 假設檢定 (二) 205

12.1 母體比例之檢定 206

12.2 樣本數 n 之規劃——在 α 、 β 控制下 219

12.3 兩母體比例差 $p_1 - p_2$ 之檢定 224

12.4 母體變異數之檢定 233

習題 241

第十三章 卡方檢定 245

13.1 適合度檢定 246

13.2 列聯表——獨立性檢定 268

13.3 列聯表——齊一性檢定 276

習題 288

第十四章 變異數分析 293

14.1 變異數分析基本觀念 294

14.2 單因子變異數分析 299

14.3 多重比較法 325

14.4 k 個常態母體變異數相等之檢定 331

14.5 一因子隨機集區設計 335

14.6 二因子變異數分析 351

習題 373

第十五章 簡單線性迴歸及相關分析 383

15.1 迴歸分析之觀念 384

15.2 簡單直線迴歸模式 388

15.3 迴歸參數之點估計 392

15.4 迴歸係數的統計推論 409

15.5 平均值與新觀測值的預測 421

15.6 迴歸分析之變異數分析法 431

15.7 相關分析 441

15.8 殘差分析 454

習題 468

第十六章 複迴歸分析與複相關 473

16.1 二元迴歸模式 474

16.2 k 個自變數迴歸模式 497

- 16.3 特殊迴歸模式 506
- 16.4 複相關係數與偏相關係數分析 516
- 16.5 虛擬變數 527
- 16.6 線性重合 535
- 習題 541

第十七章 無母數統計 547

- 17.1 概論 548
- 17.2 符號檢定 550
- 17.3 Wilcoxon 符號等級檢定 559
- 17.4 Wilcoxon 等級和檢定 567
- 17.5 Kruskal-Wallis 檢定 575
- 17.6 Friedman 檢定 578
- 17.7 Spearman 等級相關係數 582
- 17.8 連數檢定——隨機性檢定 588
- 習題 593

附錄

- 附錄一 二項機率分配機率表 598
- 附錄二 波瓦松機率分配機率表 609
- 附錄三 標準常態分配機率值 616
- 附錄四 卡方分配右尾機率值之臨界點 619
- 附錄五 T 分配臨界值 620

附錄六 F 分配臨界值 621

附錄七 隨機亂數表 629

附錄八 **Hartley** 檢定的臨界值 630

附錄九 **Wilcoxon** 符號等級檢定 W^+ 的左尾和右尾臨界值 631

附錄十 **Wilcoxon** 等級檢定 W_1 的左尾和右尾臨界值 632

附錄十一 隨機性連串檢定的左尾和右尾臨界值 633

參考書目 635

索引 637



估計

緒論

9.1

點估計式之評判標準

9.2

尋找點估計式之方法

9.3



緒論

在抽樣理論中，曾提及由於研究的母體數量過大，限於人力、物力、財力之不足，故僅能抽取有限的樣本資料，並利用此樣本資料所提供之訊息，依據推論統計學之理論方法，來推論並瞭解母體資料所隱含的特性。但由於樣本僅為母體的一小部分，且此組樣本資料之代表性如何並無法得知，因此利用樣本推論母體未知特性時，往往可能會因為抽樣資料的偏差，而使得在估計時產生估計誤差的問題，所以如何尋找優良的估計式，以及運用好的統計估計方法去推論使得估計誤差變小，便成為估計過程中不可忽略的考慮要件，而這正是本章所要討論的主題。

在機率問題中曾討論過，大部分研究之母體現象皆可以利用機率模型來替代，因此在上冊第七章的機率理論裡，曾利用許多不同類型的模式，如常態分配、二項分配、波瓦松分配、指數分配等機率函數，來替代許多各種不同的自然及社會等現象。但這些機率模式中皆包含有一個或二個以上的未知常數（即參數），如 μ 、 σ^2 、 λ 、 p 等，因此在實務上若要利用這些機率模式去作應用及統計分析，便會造成很大的困擾，所以在本章中，我們將先討論如何利用樣本資料給予這些未知的參數估計 (estimation)。而此處所謂的估計，即是利用上冊第八章所討論的樣本統計量去估計母體中未知的參數，而其內容又可區分為點估計 (point estimation) 及區間估計 (interval estimation) 兩大類。其中，點估計之意義是指根據一組樣本資料來求出某個樣本統計量或稱某一個估計式 (estimator) 之一數值，即估計值 (estimate)，以估計未知參數值為何的一種方法，由於此方法只提供未知參數一個具體的數值，故稱之為點估計。

定義

9.1

設 $\hat{\theta}(X_1, X_2, \dots, X_n)$ 為一個統計量，若它被使用來估計某一未知參數 θ 時，則 $\hat{\theta}(X_1, X_2, \dots, X_n)$ 即稱為參數 θ 之估計式，常以符號 $\hat{\theta}$ 表示之。而當獲取一組實際樣本觀察值 $(x_1, x_2, \dots, x_n)$ 所計算出來之 $\hat{\theta}(x_1, x_2, \dots, x_n)$ 則稱為參數 θ 之估計值。

例如，以 $\bar{X}$ 估計 μ 或以 S^2 估計 σ^2 ，則 $\bar{X}$ 及 S^2 即為參數 μ 及 σ^2 的估計式，但若根據抽樣所獲取之真正樣本資料所計算出來之值 $\bar{x}$ 及 s^2 即為點估計值。而區間估計係根據樣本資料建立兩個估計式之值構成一個區間，再利用此區間來說明此區間包含未知參數之信心為何？在本章，我們將先探討點估計之問題，而區間估計之問題則在下一章才進一步陳述。

點估計是利用樣本統計量來估計母體中的未知參數，但統計量的形式很多，因此在進行點估計時，必須注意下列兩個問題，即：

(1) 對母體中的未知參數作估計時，應取最佳估計式，但是所謂的最佳，應具備哪些性質？又該如何去尋找此種最佳估計式呢？

(2) 假若已知最佳的估計式，而此種估計式仍不能確切表示母體中的未知參數，只能在某種可靠性內求得此種估計式推估母數之範圍時，此種範圍要如何尋找呢？

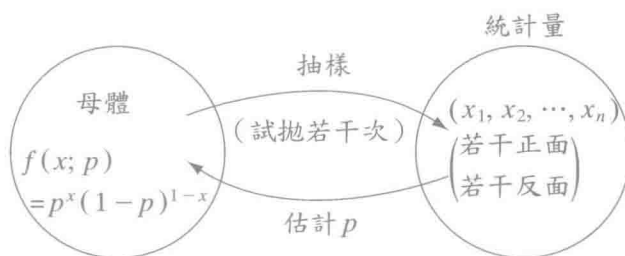
上述第(1)個問題便是本章所要研究的主題。茲舉一例說明點估計之觀念。例如，投擲一枚錢幣出現正面之次數，其母體分配為一個點二項分配，即：

$$f(x; p) = p^x(1-p)^{1-x}, \quad x = 0, 1, \quad 0 < p < 1$$

p 為錢幣出現正面的機率， p 的變化範圍在 0 與 1 之間，由不同 p 決

定此錢幣出現正面之情況。在不採用科學儀器以檢查此錢幣之正、反兩面構成如何，而完全要由統計推定去瞭解此錢幣出現正面或反面的機率時，唯一的方法就是將錢幣試拋若干次，亦即抽樣，再根據此抽樣結果給予 p 估計。此時就必須思考，究竟應根據何種估計式來推論母體中的 p 較適合呢？如果找到某種估計式時，以此種估計式解釋母體未知母數之可靠性又為如何呢？利用抽樣估計參數之概念可以圖 9.1 表示之。

Figure 9.1 估計之觀念



9.2

點估計式之評判標準

在估計一個未知參數 θ 時，我們可利用很多不同形式的樣本統計量 $\hat{\theta}$ 作為其估計式，例如要估計母體中未知母體平均數 μ 時，我們可採用樣本平均數 $\bar{X}$ 與樣本中位數 Me 或其他各種估計式來估計。但若要問這些估計式中到底何者較為優良時，就必須提供一個評判標準來判別。因此評判點估計式之優劣，即為點估計問題首先要討論之內容。一般常見判斷點估計式優劣之準則有下述幾類，即：

- (1) 不偏性 (unbiasedness)。
- (2) 有效性 (efficiency)。
- (3) 一致性 (consistency)。
- (4) 充分性 (sufficiency)。

茲將各種判斷準則之觀念及內容分述如下。

9.2.1 不偏性

在定義 9.1 中曾提過一個參數 θ 之估計式 $\hat{\theta}$ ，本身是一個統計量，為隨機樣本 $(X_1, X_2, \dots, X_n)$ 之函數，因此估計式 $\hat{\theta}$ 之估計值在估計參數 θ 時，可能會因不同之樣本觀察值 $(x_1, x_2, \dots, x_n)$ 而不同。但若這些估計值之平均數能接近 θ 值時，則表示以 $\hat{\theta}$ 估計 θ 時可能較好，此種觀念即為不偏性之觀念。

定義

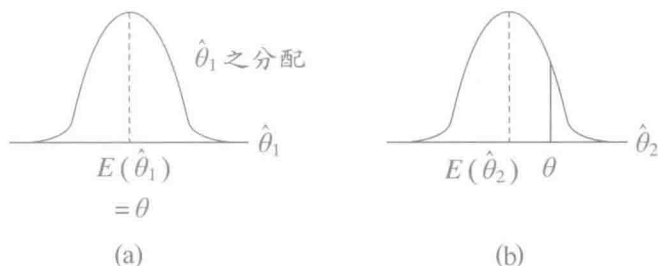
9.2

若統計量 $\hat{\theta}$ 為母體未知參數 θ 的一個估計式，且 $\hat{\theta}$ 的期望值等於未知參數 θ ，即 $E(\hat{\theta}) = \theta$ ，則稱 $\hat{\theta}$ 為母體未知參數 θ 的不偏估計式 (unbiased estimator)。但若 $E(\hat{\theta}) \neq \theta$ ，則稱 $\hat{\theta}$ 為未知參數 θ 之偏誤估計式 (biased estimator)，且其偏誤為 $E(\hat{\theta}) - \theta = Bias$ 。

由定義 9.2 中可知，要測定一個估計式是否為參數之不偏估計式，只要依上述定義求估計式的期望值即可（見圖 9.2）。在圖 9.2(a) 中，可以很明顯地發現 $\hat{\theta}_1$ 之抽樣分配之期望值 $E(\hat{\theta}_1) = \theta$ ，故 $\hat{\theta}_1$ 為 θ 之不偏估計式；但在圖 9.2(b) 中， $\hat{\theta}_2$ 之期望值並不等於 θ ，因此 $\hat{\theta}_2$ 即為具有負偏誤之估計式。例如，在上冊第八章 $\bar{X}$ 的抽樣分配中曾證明 $E(\bar{X})$

Figure

●圖 9.2 不偏及偏誤估計式



$= \mu$ ，因此，樣本平均數 $\bar{X}$ 為母體平均數 μ 的不偏估計式。又樣本變異數 S^2 是否為母體變異數 σ^2 之不偏估計式呢？以例 9.2.1 說明之。

Example

9.2.1

設 $(X_1, X_2, \dots, X_n)$ 為來自任意母體 $f(x)$ 的一組隨機樣本，若樣本變異數為

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

試證明 S^2 為 σ^2 的不偏估計式。

□ Solution

因

$$\begin{aligned} E(S^2) &= E\left[\frac{\sum_{i=1}^n (X_i - \bar{X})^2}{n-1}\right] = \frac{1}{n-1} E\left[\sum_{i=1}^n (X_i - \bar{X})^2\right] \\ &= \frac{1}{n-1} E\left[\sum_{i=1}^n X_i^2 - n\bar{X}^2\right] = \frac{1}{n-1} \left\{ \sum_{i=1}^n E(X_i^2) - nE(\bar{X}^2) \right\} \\ &= \frac{1}{n-1} \left\{ \sum_{i=1}^n [Var(X_i) + (E(X_i))^2] - n[Var(\bar{X}) + (E(\bar{X}))^2] \right\} \end{aligned}$$

$$\begin{aligned}
 &= \frac{1}{n-1} \left\{ \sum_{i=1}^n (\sigma^2 + \mu^2) - n \left(\frac{\sigma^2}{n} + \mu^2 \right) \right\} \\
 &= \frac{1}{n-1} \{ n\sigma^2 + n\mu^2 - \sigma^2 - n\mu^2 \} = \frac{1}{n-1} (n-1) \sigma^2 \\
 &= \sigma^2
 \end{aligned}$$

故樣本變異數 S^2 為 σ^2 之不偏估計式。

□

在例 9.2.1 中，若將樣本變異數 S^2 計算公式中之分母部分 $n-1$ 更改為 n ，即

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

此時可發現 $E(S_*^2) = \left(\frac{n-1}{n} \right) \sigma^2 \neq \sigma^2$ ，即 S_*^2 為偏誤估計式，因此在母體變異數 σ^2 之估計中，我們常以不偏估計式 $S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$ 來估計 σ^2 。但事實上，在樣本數 n 大時，可發現

$$\lim_{n \rightarrow \infty} E(S_*^2) = \lim_{n \rightarrow \infty} \left(\frac{n-1}{n} \right) \sigma^2 = \sigma^2$$

亦即當樣本數 n 大時， $S_*^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2$ 亦為不偏估計式。而此種不偏估計式又稱之為極限不偏 (asymptotic unbiased) 估計式。

Example

9.2.2

設 $(X_1, X_2, \dots, X_5)$ 為抽自平均數為 θ 之指數分配中的一組隨機樣本，試問下列估計式中何者為參數 θ 之不偏估計式？

(1) $\hat{\theta}_1 = X_3$