

南 投 縣 政 府 110 年 度 研 究 報 告

研析實價登錄資料庫價格結構及分佈 對區域地價影響因子評估

研究人

服務單位：草屯地政事務所

研究人員：李金輝

中 華 民 國 109 年 11 月 10 日

南 投 縣 政 府 110 年 度 研 究 報 告 摘 要 表

研 究 報 告 名 稱	研析實價登錄資料庫價格結構及分佈對區域地價影響因子評估
研 究 單 位 及 人 員	草屯地政事務所：李金輝
研 究 起 迄 年 月	109 年 7 月至 109 年 11 月
研 究 緣 起 與 目 的	本研究期望應用實價登錄資料庫加值應用，進而延伸到更複雜的模型。雖然估計的模型不斷改變，但對實務界而言，如何建立一個最能準確估計不動產大量估價的模型，才是最重要的課題。
研 究 方 法 與 過 程	1、根據研究動機與目的，界定研究內容與範圍，以及相關資料的選擇。 2、資料庫方面，是以實價登錄房地交易資料為統一資料來源。資料內容方面為草屯鎮 108 年 9 月 2 日～109 年 9 月 1 日實價登錄資料庫為範圍，並在過濾時考慮是否正確且經過時間及指數修正。 3、本研究主要是採用 R 語言為分析工具，其中使用的決策數函式為 Variance。 4、依實證分析結果所產生之變異數、平均數、中位數三個數值做判釋。 5、建構各資料庫參考模型並對分析判釋結果做成結論及建議。
研 究 發 現 與 建 議	經本研究分析結果，我們分別產製了土地及房地各個交易模型。有些中位數與平均數接近的模型中，其變異數較大，有些則較趨近。這充分的反映出本研究的方式，確實可以解析實價登錄資料庫中，價格結構之關聯性及合理性。 在未來若執行電腦大量估價時，資料庫數據也可

	先經由以此模型做初步解析，來決定此資料庫權值的釐定及參酌。
選擇獎勵	☒ 行政獎勵 ☐ 獎勵金

研析實價登錄資料庫價格結構及分佈對區域地價影響 因子評估

Research and analysis Actual land transaction price database price structure and distribution on the evaluation of regional land prices

李金輝
Chin Hui-Lee

梁崇智
Liang Chung-Chih

周永信
Yun Hsin-Chou

摘要

在大量估價資料庫的建置過程中，資料的來源及品質對於估價結果具有決定性影響。因此估價結果的準確度，除依賴估價模型的精進外，資料篩選更是影響估價結果的關鍵因素。我國於實價登錄制度實施後，不論是在不動產資料的量化與品質，相較過去已有大幅提升，這對於土地估價之實務運用可行性將大幅提高。

本研究應用變異數分析，來檢定實價登錄資料庫內的價格分佈，與各不同區域內地價變化情形，來建立一種特殊形式的統計假設檢定。在一定時間實範圍內的地區，經過指數修正後之價格，以我們目前實價登錄資料，來比較那些區域內價格的因素，會對未來估價方式產生影響。並且方便由估價人員應用該參考數據，作為掌握地價高低層次及衡量周邊土地價格之參考。

經本研究結果顯示，資料庫中各筆資料中的各個變數，是反映會影響價格的因素之一，而其中有一些地價區段其變異數與中位數及平均值間差異頗大。檢視其內容發現，當資料量越大越離散時，變異數與中位數及平均值間差異就越大越明顯。

關鍵字：變異數、中位數、平均值

Research and analysis Actual land transaction price database price structure and distribution on the evaluation of regional land prices

Abstract

In the process of constructing the large-scale valuation database, the source and quality of the data have a decisive influence on the valuation results. Therefore, in addition to relying on the refinement of the valuation model, the accuracy of the valuation results is also a key factor affecting the valuation results. After the implementation of the actual price registration system in China, both the quantification and quality of real estate materials have been greatly improved compared with the past, which will greatly improve the feasibility of land price estimation and practical application.

This study uses variance to analyze, to verify the price distribution in the actual price database, and the changes in land prices in different regions, to establish a special form test of statistical hypothesis. In the region within a certain time range, the price after index correction is compared with our current actual price data to compare the price factors in those regions, which will affect the future valuation method. And it is convenient for appraisers to apply the reference data as a reference for grasping the level of land price and measuring the price of surrounding land.

The results of this study show that each variable in each piece of data in the database reflects one of the factors that affect the price, and some of the land price segments are between Variance and Median and mean value The difference is quite big. Examining its content, it is found that when the amount of data is larger and more discrete, the difference between Variance and Median and mean value is larger and more obvious.

Keywords: Variance 、 Median 、 mean value

壹、研究源起與目的

一、研究源起

台灣在全球化及國際經貿自由化的影響下，不動產實價登錄制度的推行，確實有助把不動產實際的交易價格公開化及透明化，同時也可避免有不完全競爭市場出現。同時由於此制度的推行，使得政府對於不動產估價的方法與模型產生很大的改變。由傳統的人工查價方法，進而可應用此資料庫加值應用進而延伸到更複雜的模型。雖然估計的模型不斷改變，但對實務界而言，如何建立一個最能準確估計不動產大量估價的模型，才是最重要的課題。目前台灣不動產市場及實價登錄逐漸成熟，未來在不動產交易更為普遍透明的情況下，如何建立一個能夠迅速且正確評估不動產的大量估價模型，更是非常重要的。

二、研究目的

目前我國政府地政機關的地價制度，主要係依地價調查估計規則辦理，查估公告土地現值及公告地價。但由於人力、經費、時間有限，逐筆查估不符合經濟效益，故採區段地價之估價方式。因為採區段地價方式，大多以區段地價來代表宗地地價，如此是否未能考慮土地個別因素影響，是一個值得關注的課題。另政府建立不動產估價師制度，民間估價體系出現嶄新契機，且政府亦不斷研究改進估價方法，公私部門若各依不同的估價法規，運用不同的估價基準執行業務，因此建立一套估價機制，作為政府與民間共同使用之估價系統，是一項可以研究的課題。

台灣在全球化及國際經貿自由化的影響下，不動產實價登錄制度確實有助於把不動產實際的交易價格公開及透明化，避免有不對稱競爭市場出現。不動產實價登錄主要目的是為了讓房市交易價格符合公開、透明的功能，並可配合房地合一課稅。因實價登錄制度提供了房地產實際的交易資訊，有助於平衡購屋者與廣大房市中之資訊不對稱的情況。而在實價登錄制度完備後，地方政府能有明確的參考依據。

因此，本研究想藉由變異數來分析實價登錄資料庫的價格結構，來檢視

若以此價格結構，對於不同地價區段所產生的變數及影響情形。另外，本研究也將對此影響因子及分佈情形做分析，並對於未來若實施電腦大量估價提出建議。

貳、相關理論及文獻探討

一、相關理論

當一個資料量收集完成後，可能會非常零亂及複雜，同時也很難看出該筆資料的特性。要如何將這些資料整理並歸整化成可讀性及可用性，常常會藉由圖示來表示資料的分布情形，有時也會計算其平均數 (mean)、中位數 (median) 等來看該筆資料的中心位置。同時還會經由計算變異數 (variance) 來看該筆資料的離散程度。如此一來，資料收集者可以簡單描述該資料的特性，讓有興趣者可以一目了然，取得所需的資訊。

往往在我們所關注的群體資料量會很大，而且在這龐大的資料量中，平均數實際上常常無法知道。因此為了減少調查時間與增加效率，常常會藉由抽樣來取得樣本資料，希望能藉由樣本資料，獲得樣本平均數與樣本變異數。利用樣本平均數來估計族群平均數，與利用樣本變異數來估計母體變異數，進而了解整個資料群體的狀況。

(一)變異數檢定種類：

變異數分析 (Variance) 是一種特殊形式的統計假設檢定，廣泛應用於實驗數據的分析中。在變異數分析的應用中，原假設是假設所有數據組都是整體測試物件的完全隨機抽樣。

變異數分析大致分為三種型態：

1、固定效應模式 (Fixed-effects models)

用於變異數分析模型中所考慮的因子為固定的情況，也就是目標因子是來自於特定的範圍。例如要比較 6 種不同的商品售量差異，目標因子就是 6 種不同的商品，其所反應變數為銷售量。該命題即限定了特定範圍，因此模

型的推論結果也將全部著眼在 6 種不同的商品售量差異上，故此種狀況下的因子便稱為固定效應。

2、隨機效應模式 (Random-effects models)

不同於固定效應模式中的因子特定性，在隨機效應中所考量的因子是來自於所有可能的母群體中的一組樣本，因子變異數分析所推論的並非著眼在所選定的因子上，而是推論到因子背後的母群體。例如藉由一間擁有全部商品的百貨公司，從所有商品中隨機挑選 5 種樣本，用於比較其銷售量的差異，最後推論到這間百貨公司的銷售狀況。因此在隨機效應模型下，研究者所關心的並非侷限在所選定的因子上，而是希望藉由這些因子推論背後的母群體特徵。

3、混合效應模式 (Mixed-effects models)

此種混合效應絕對不會出現在單因子變異數分析中，當雙因子或多因子變異數分析同時存在固定效應與隨機效應時，此種模型便是典型的混合型模式。

(二)其他相關變異模型介紹

1、單因子變異數分析：

進行變異數分析需滿足以下基本假設(1)各組樣本需取自於常態母體。
(2)各組母體變異數需假設相等。(3)各組樣本彼此獨立。

總變異(total sum of squares)

$$SST = \sum_{i=1}^k \sum_{j=1}^{n_i} (x_{ij} - \bar{x})^2$$

組內變異(error sum of squares)

$$SSE = \sum_{i=1}^k \sum_{j=1}^{n_i} (x_{ij} - \bar{x}_i)^2$$

組間變異(sum of squares between)

$$SSB = \sum_{i=1}^k n_i (\bar{x}_i - \bar{x})^2$$

2、多個平均數之多重比較：在 k 個母體平均數之多重比較中，兩母體平均數差 $(1-\alpha)*100\%$ 聯合信賴區間為：

$$\left[\bar{x}_i - \bar{x}_j - t_{\alpha/2a} (n-k) \sqrt{MSE \left(\frac{1}{n_i} + \frac{1}{n_j} \right)}, \bar{x}_i - \bar{x}_j + t_{\alpha/2a} (n-k) \sqrt{MSE \left(\frac{1}{n_i} + \frac{1}{n_j} \right)} \right]$$

3、二因子變異數分析

二因子變異數分析中，因子是指自變數，二因子(若以 A、B 表示)代表有二個自變數。二因子變異數分析的目的是要檢定主要效果(main effects)和交互效果(interaction effects)。

(1)主要效果：個別的自變數對依變數所造成的影響是指在二因子變異數分析中，若交互效果不顯著，就進行主要效果的比較，這項比較是針對個別的自變數其邊緣平均數進行比較。

(2)交互效果：所有自變數共同對依變項所造成的影響。也就是 A 因子對依變項的效果，會受到 B 因子的影響。同樣的 B 因子對依變項的效果，也會受到 A 因子的影響。所以二因子變異數分析，首先要檢視二因子的交互效果，若交互效果顯著則進行單純主要效果分析，否則進行主要效果分析。

二、文獻探討

Moran,P.和 C. Smith 最早在 1918 年提出變異數概念，其部分章節節錄如下：

“Several attempts have already been made to interpret the well-established results of biometry in accordance with the Mendelian scheme of inheritance. It is here attempted to ascertain the biometrical properties of a population of a more general type than has hitherto been examined, inheritance in which follows this scheme. It is hoped that in this way it will be possible to make a more exact analysis of the causes of human variability. The great body of available statistics show us that the deviations of a human measurement from its mean follow very closely the Normal Law of Errors, and, therefore, that the variability may be uniformly measured by the standard deviation corresponding to the square root of the mean square error. When there are two independent causes of variability capable of producing in an otherwise uniform population distributions with standard deviations σ_1 and σ_2 , it is found that the distribution, when both causes act together, has a standard deviation $\sqrt{(\sigma_1^2 + \sigma_2^2)}$.”

在這篇節錄中，探討孟德爾遺傳計劃並期望對變異的原因進行精確分析，另外在變異性上可以通過對應於平均值平方根的標準偏差來統一度量平方誤差。當有兩個獨立的變異導致能夠影響原本均勻的總體分佈時，也就是標準偏差 σ_1 和 σ_2 ，而發現當兩個原因共同作用時，該分佈具有標準偏差 $(\sqrt{\sigma_1^2 + \sigma_2^2})$ 。

Although Fisher states that random mating is assumed at first, the theory is developed in terms more general than this and he is careful to state when the additional assumption is introduced. He also assumes for the present that each measured character is the result of summing a large number of small factors which are independent, i.e. that there is no linkage. It then follows from the standard properties of means, variances and covariances that the mean value of the character in the population is equal to the sum of the means of the individual small factors, the variance is similarly the sum of the individual variances, and the same is true of the covariances.

另外在以上以費舍爾隨機配對的理論及假設中，論證從均值、方差和協方差的標準屬性得出的結論，是總體中的平均值等於各個小因子的均值之和，方差是各個方差之和，並且協方差也是如此。

This assumes that the causes act additively and not, for example, multiplicatively.

"It is therefore desirable in analysing the causes of variability to deal with the square of the standard deviation as the measure of variability. We shall term this quantity the Variance of the normal population to which it refers, and we may now ascribe to the constituent causes fractions or percentages of the total variance

這章節主要是在敘述分析變異性的原因時，最好將標準差的平方作為變異性的度量。我們將這個量稱為其所指該群體的方差，在此我們可以將其構成原因歸因於總方差的分數或百分比。

另外王俊明(1995)所提出的論點，是討論變異數分析的性質，並比較不同種類的變異數分析之優缺點。其中並以例子說明各類型變異數分析的結果，所舉的例子包括有獨立樣本單因子變異數分析、重複量數單因子變異數分析、獨立樣本二因子變異數分析、混合設計二因子變異分析及重複量數二因子變異數分析等。此項結果可以提供研究者們在事研究工作時，能選擇比較適當的實驗設計。

范德鑫(1988)認為一些研究者在使用變異數方法時，習慣採用綜合的F考驗，如果發現F考驗達統計顯著(Statistical significance)則繼續進行事後比

較(Post hoc Comparison)，若無顯著，則告一段落。不可否認的，這種作法對於探究性的研究極為有用。

另外又敘明此種統計方法所以具有如此廣大的應用性，主要歸因於下列兩點：(1)它不但具有考驗處理樣本平均數差異功能，而且不受平均數個數之限制；(2)它能同時處理兩個以上的自變項的交互的效果(Howell, 1985)。

綜上，對變異數分析性質的描述，得知變異數分析的用途及其使用限制，同時得以了解各種變異數分析的差異。基本上，各種的實驗設計都有其優缺點。因此，研究者必須從研究目的加以衡量，來如何選擇最適合自己研究的實驗設計。

叁、研究範圍與方式

一、研究範圍

本研究區範圍選取是以草屯鎮 108 年 9 月 2 日~109 年 9 月 1 日實價登錄資料庫為範圍，該資料庫剔除了特殊交易態樣，同時並篩選出都市計畫內之住宅區，為本實驗研究樣區。

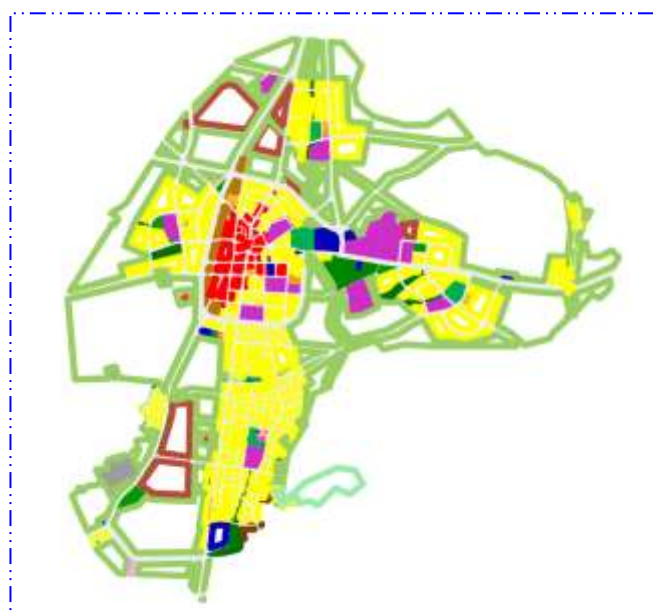


圖 1 本研究區試驗範圍圖

二、研究方法與過程

本研究之研究步驟及方法如下：

- 1、根據研究動機與目的，界定研究內容與範圍，以及相關資料的選擇。

2、資料庫方面，是以實價登錄房地交易資料為統一資料來源。資料內容方面為草屯鎮 108 年 9 月 2 日~109 年 9 月 1 日實價登錄資料庫為範圍，並在過濾時考慮是否正確且經過時間及指數修正。

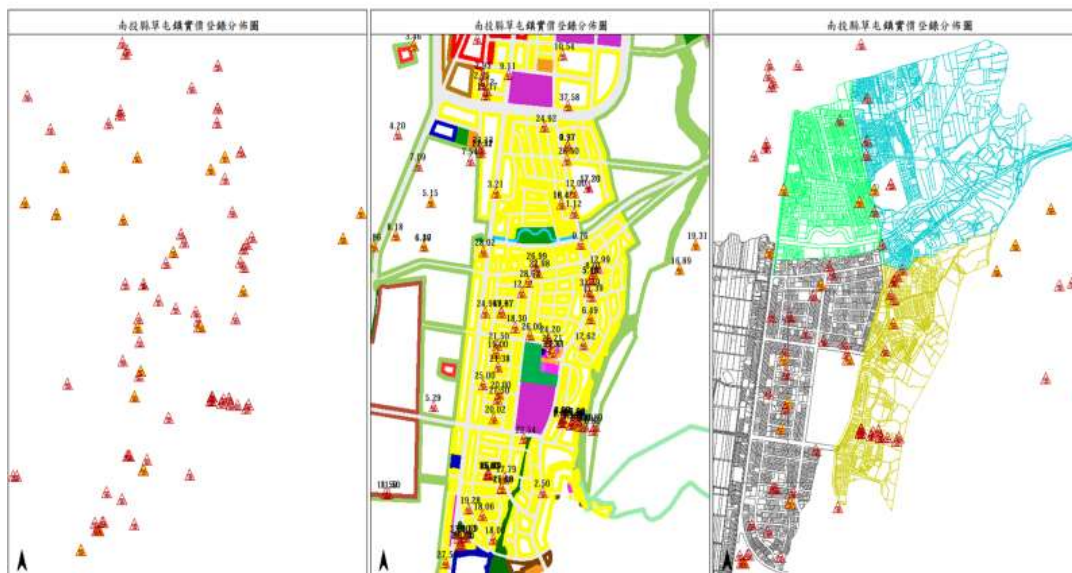


圖 2 草屯鎮都市內實價登錄分佈圖

3、本研究主要是採用 R 語言為分析工具，其中使用的決策數函式為 Variance。

4、依實證分析結果所產生之變異數、平均數、中位數三個數值做判釋。

$$\text{Variance} = \text{Var}(X) = E[(X - \mu)^2]$$

Median=可以通過把所有觀察值高低排序後找出正中間的一個作為中位數。

$$\text{mean value} = \bar{x} = \frac{x_1 + x_2 + \cdots + x_n}{n}$$

5、建構各資料庫參考模型並對分析判釋結果做成結論及建議。

肆、實證分析

本研究採用次程式語言去撰寫 R studio 為系統分析方式，因為 R 語言是結合統計分析與繪圖功能的原始碼軟體，R Studio 則是能讓我們編寫 R 程式時候使用體驗更好的整合開發環境(Integrated Development Environment, IDE)，同時本研究架構也依循前述特性去設計。

設 X 為服從某分布的隨機變數，如果 $E[X]$ 是隨機變數 X 的期望值（平均數 $\mu=E[X]$ ）隨機變數 X 的變異數為：

$$\text{Var}(X) = E[(X - \mu)^2]$$

中位數

可以通過把所有觀察值高低排序後找出正中間的一個作為中位數

平均數。

算術平均數： n 個數據相加後除以 n 。

算術平均數（或簡稱平均數）是一組樣本 $x_1, x_2, \dots, x_n$ 的和除以樣本的數量。

其通常記作：

$$\bar{x} = \frac{x_1 + x_2 + \dots + x_n}{n}$$

程式碼撰寫及設計分析結果：

土地交易資料分析

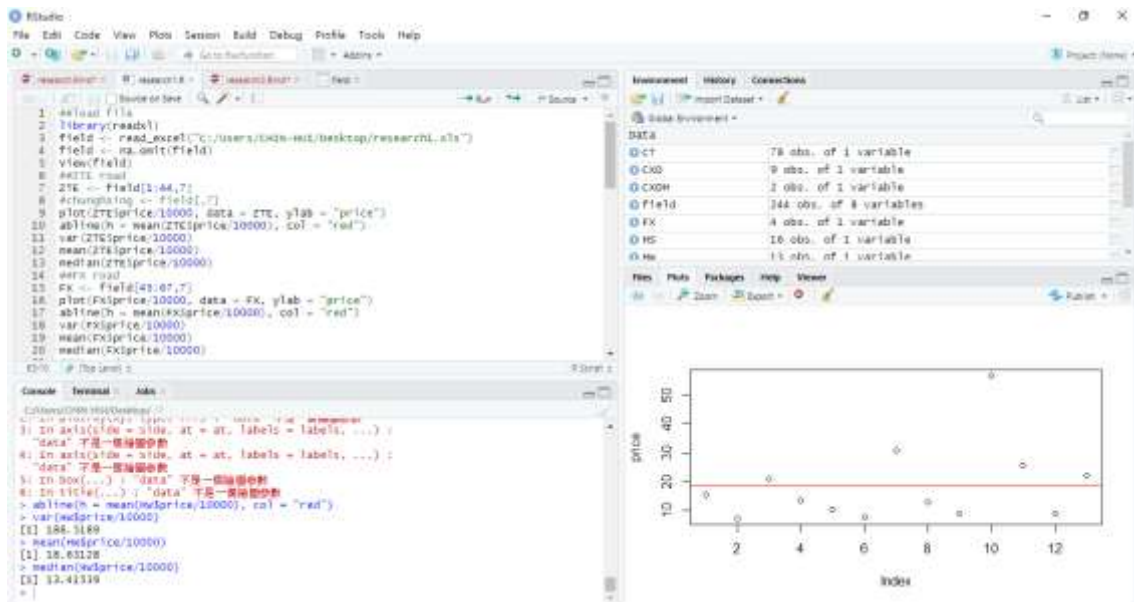


圖 3 程式碼分析圖

中興段(土地)分析結果數據：

變異數 [1] 42.43064

平均數 [1] 20.78452

中位數 [1] 19.51587

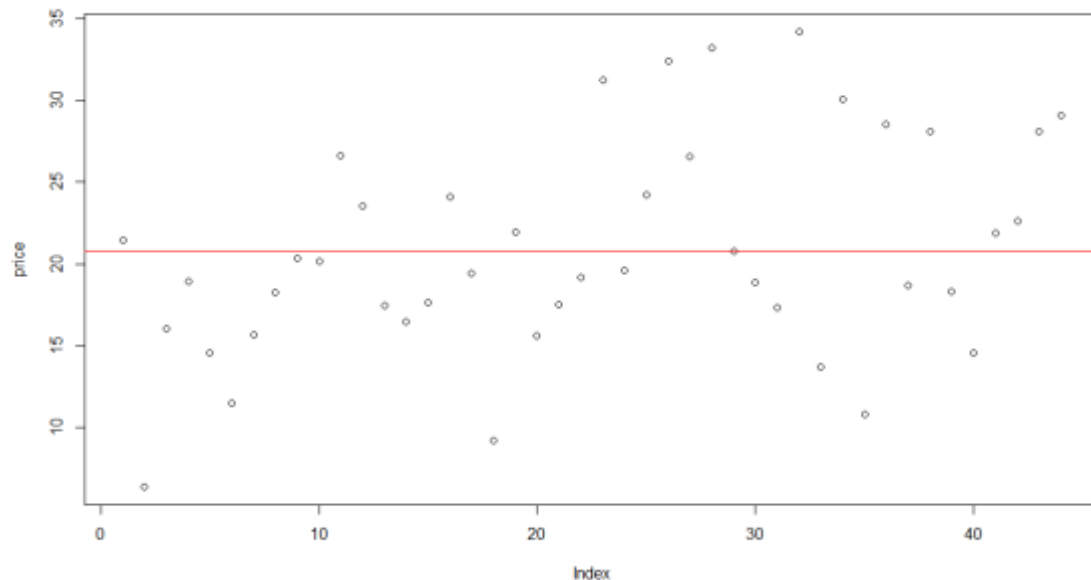


圖 4 中興段(土地)程式分析分佈圖

以上是中興段土地交易實價登錄資料庫分析結果，變異數值為42.43064，平均數及中位數分別為20.78452、19.51587，依平均數與中位數越接近時，代表分佈越平均概念。其分析圖表可看出中興段土地交易價格分佈看似尚屬平均，但其變異數為42.43064，故其價位離紅色均線上下分佈較離散。

由以上推論，中興段地價分佈其個別區域間差異較明顯，分析其原因，是否與南北狹長區塊有關，以至於造成不同區位間因不同之便利性而產生個別差異所致？若此推論成立，則該區未來若在做電腦大量估價時，其區域的劃分就格外顯得非常重要了。

府西段(土地)分析結果數據:

變異數 [1] 1.274952

平均數 [1] 8.761881

中位數 [1] 8.116967

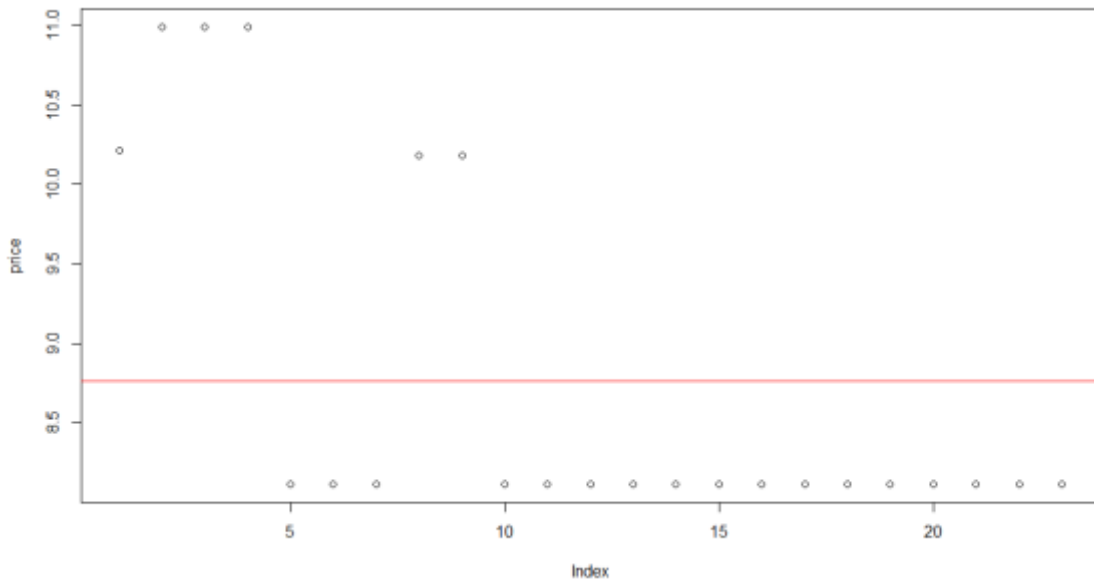


圖 5 府西段(土地)程式分析分佈圖

再來是府西段土地交易實價登錄資料庫分析結果，變異數值為 1.274952，平均數及中位數分別為 8.761881、8.116967，依平均數與中位數越接近時，代表分布越平均概念，其分析圖表可看出府西段土地交易價格分佈平均，另外其變異數為 1.274952，在價格一致性方面看來頗均衡。

以上初步看出，府西段地價分佈其個別區域間差異不明顯，但再分析其內容，其中有好幾筆價格分佈相近，以至於呈現出不論是變異數、平均數、中位數皆相近。深入觀察，這些相近的交易其買賣目的及主體是否相同？以致於造成其中一批交易價格其價格帶幾乎相同。若此推論結果，則樣本未來若在做電腦大量估價時，其資料來源是否加入其他數據值得參考。

建興段(土地)分析結果數據:

變異數 [1] 13.34592

平均數 [1] 21.68837

中位數 [1] 21.9788

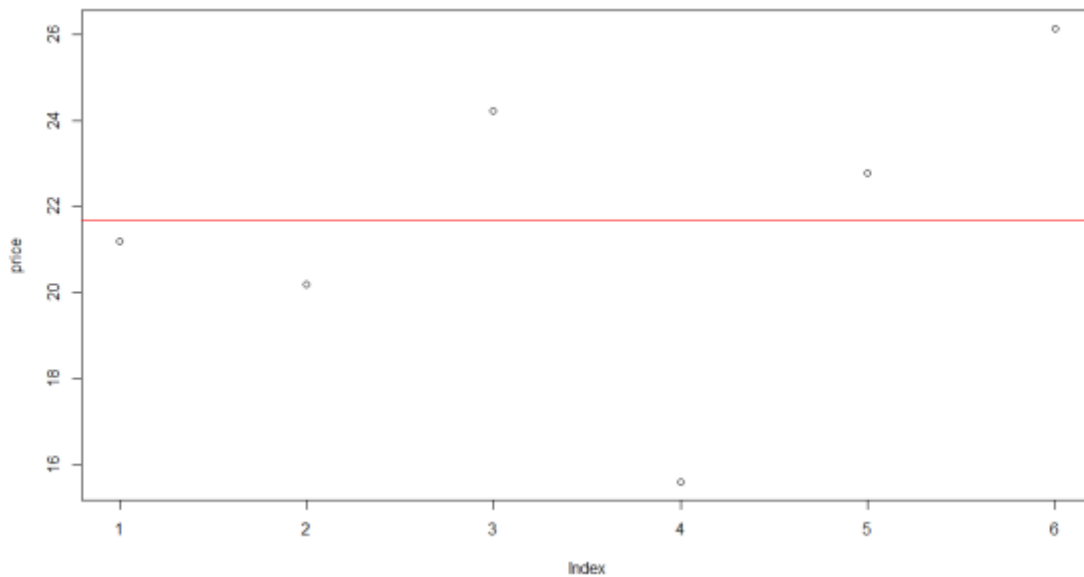


圖 6 建興段(土地)程式分析分佈圖

以上是建興段土地交易實價登錄資料庫分析結果，變異數值為 13.34592，平均數及中位數分別為 21.68837、21.9788，由以上分析圖表分析建興段土地交易價格分佈平均，另外其變異數為 13.34592，分析原因其成交量過低，可能間接造成變異較為離散。

經由以上評斷，建興段地價分佈其個別區域間差異不明顯，雖然成交量能不高，以至於呈現出變異數較為離散，但是平均數、中位數皆相近。經觀察雖然成交量能不足，但也有可能是該區本身所能提供空地量能有限所致。依此推論，該情形未來在做電腦大量估價時，因其腹地及量能不大，未來在估價上，分歧也會較小。

草屯段(土地)分析結果數據:

變異數 [1] 134.1124

平均數 [1] 22.28707

中位數 [1] 18.96229

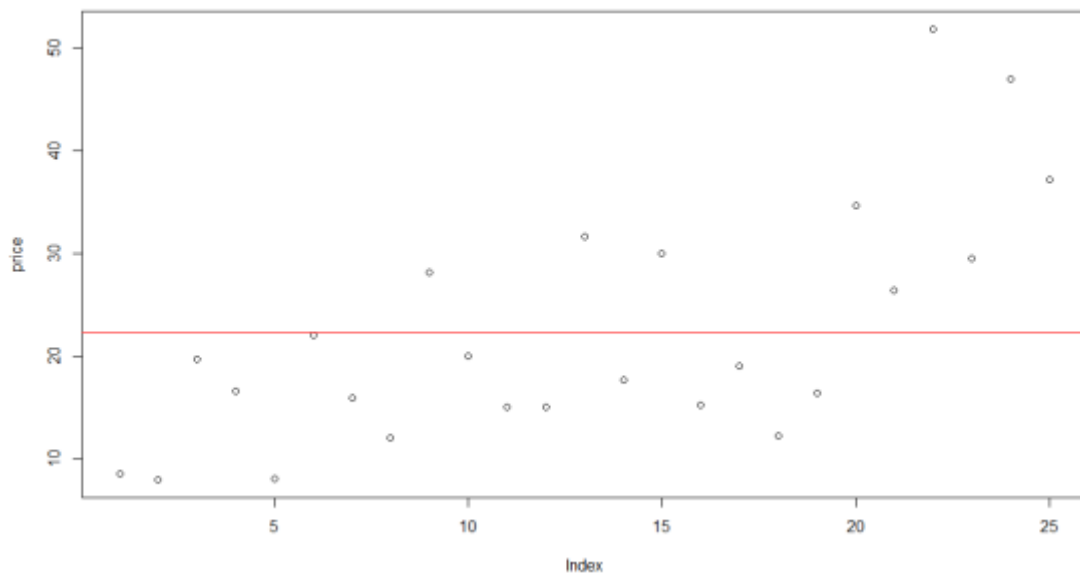


圖 7 草屯段(土地)程式分析分佈圖

以上是草屯段土地交易實價登錄資料庫分析結果，變異數值為 134.1124，平均數及中位數分別為 22.28707、18.96229，由以上圖表分析草屯段土地交易價格分佈較不平均，另外其變異數為 134.1124，觀察其內容，除了其成交量能較高，其價格分佈離紅色均線也較為離散。

經分析及評斷，草屯段本身夾雜住、商、工、農等分區，因此雖然實價登錄資料庫已過濾為同一住宅區。但由於有些區域因地域之鄰近性，呈現住商或其他型態，也因此相同分區及區段地價分佈差異會較大，雖然成交量能較高，也有可能間接造成變異數較大，但搭配平均數、中位數分析發現，該段的確在區域因素的個別差異上產生了分歧。

分析以上結果，草屯段未來在做電腦大量估價時，可能要搭配實價登錄資料庫的劃分，並同時注意各區域不同性質的律定及細分。也就是說，未來在本模型案例上，大量估價可能會偏向宗地之個別性，以往的區段地價，恐無法較完整呈現本區域之實際量價情形。

中興段(房地)分析結果數據:

變異數 [1] 42.79788

平均數 [1] 17.22229

中位數 [1] 16.65646

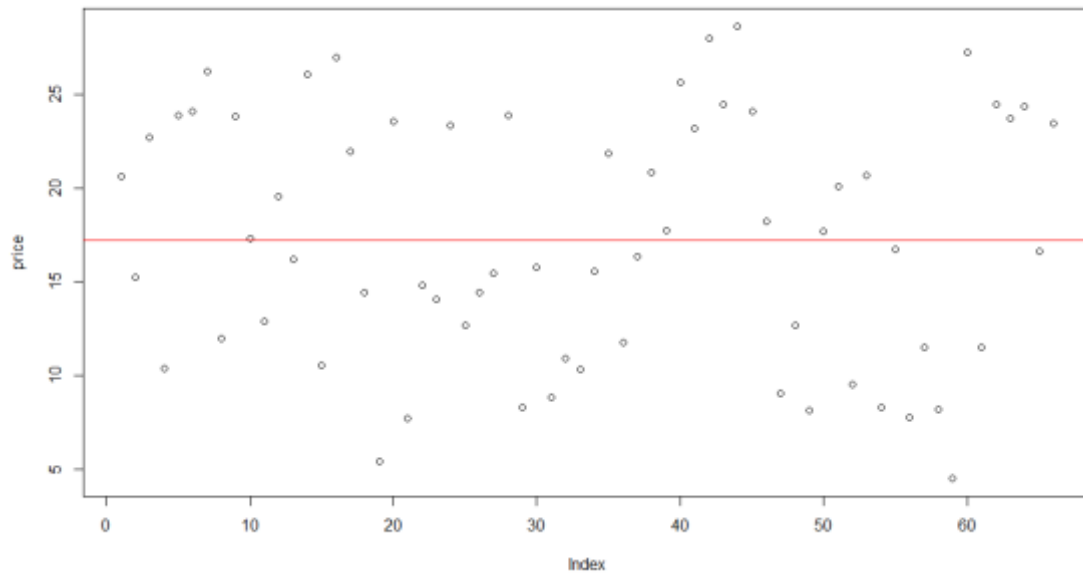


圖 8 中興段(房地)程式分析分佈圖

以上是中興段房地交易實價登錄資料庫分析結果，變異數值為42.79788，平均數及中位數分別為17.22229、16.65646，由以上分析圖表可看出，中興段房地交易價格分佈均衡情形與土地相近，變異數為42.79788，故其價位離散分佈情形也與土地分析結果類似，由以上推論，該段實價登錄資料庫模型，房地價格的一致性頗高。

府西段(房地)分析結果數據:

變異數 [1] 131.4659

平均數 [1] 18.92706

中位數 [1] 15.58165

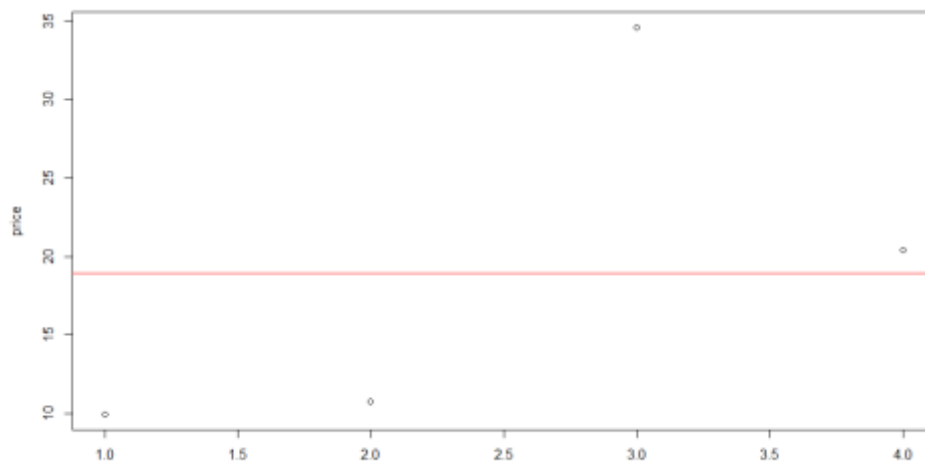


圖 9 府西段(房地)程式分析分佈圖

前揭府西段房地交易實價登錄資料庫分析結果，變異數值為 131.4659，平均數及中位數分別為 18.92706、15.58165，由分析圖表可看出府西段土地交易價格分佈平均數及中位數相差 3.34541，但變異數高達 131.4659，其價位分佈數據方面看來差異頗大。

經由以上內容分析，府西段房地交易實價登錄資料庫交易內容稀少，因此造成價格變異分佈較大，以至於呈現出變異數與平均數、中位數皆負相關。深入細究，該模型未來若在做電腦大量估價時，其資料來源是否能更多元值得參考。

建興段(房地)分析結果數據：

變異數 [1] 22.77364

平均數 [1] 24.27948

中位數 [1] 23.86232

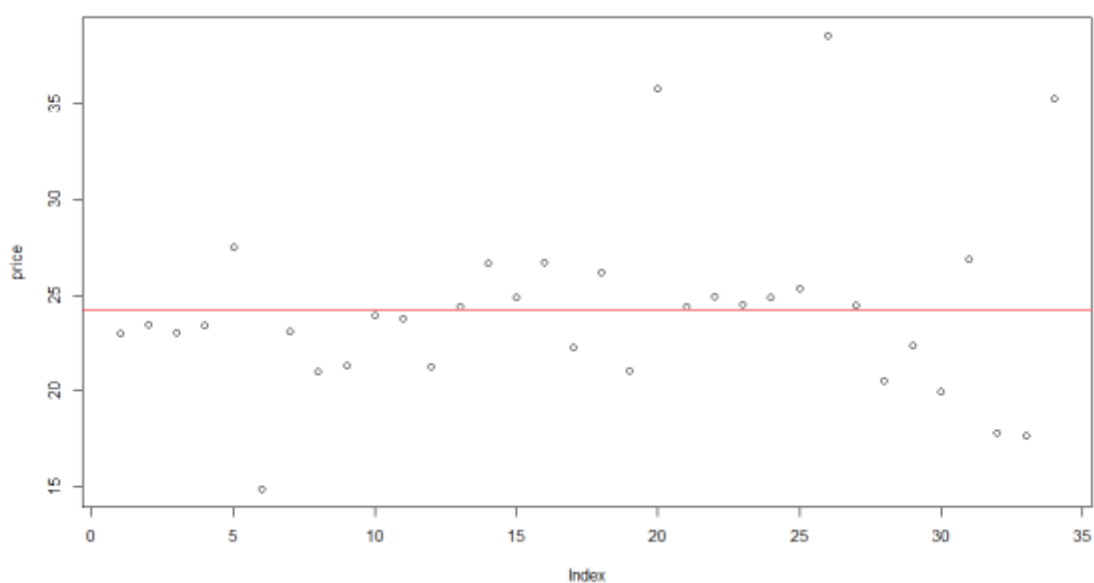


圖 10 建興段(房地)程式分析分佈圖

以上是建興段房地交易實價登錄資料庫分析結果，變異數值為 22.77364，平均數及中位數分別為 24.27948、23.86232，由以上分析圖表分析，建興段房地交易價格分佈與變異數皆平均，可證明該模型房地交易價格分歧不大，區域地價相對穩定，實價登錄資料庫權值相對可提高。

草屯段(房地)分析結果數據:

變異數 [1] 100.8608

平均數 [1] 17.71474

中位數 [1] 16.13761

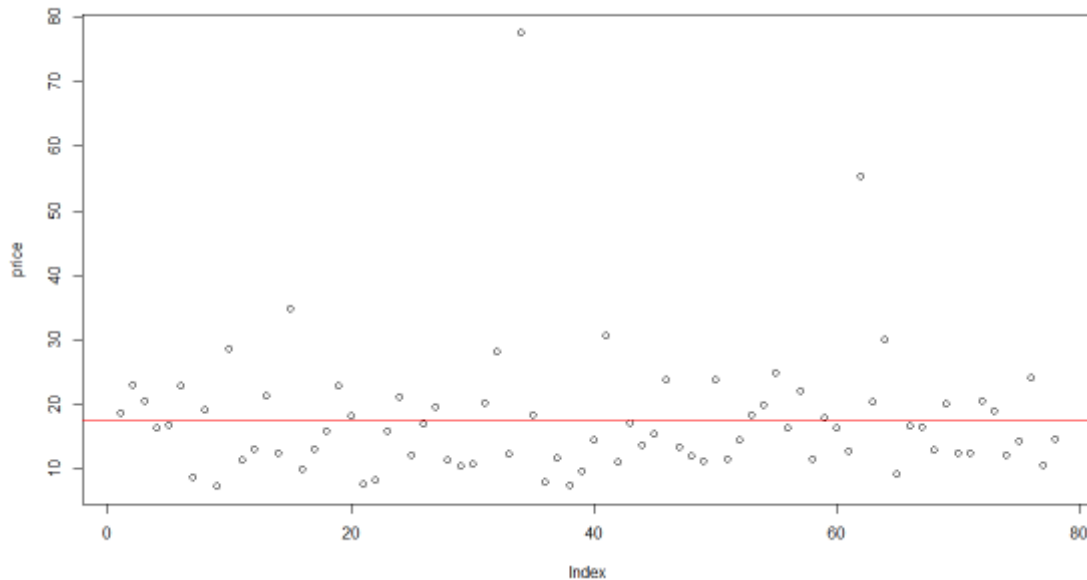


圖 11 草屯段(房地)程式分析分佈圖

草屯段房地交易實價登錄資料庫分析結果顯示，變異數值為 100.8608，平均數及中位數分別為 17.71474、16.13761，以上分析可看出草屯段房地交易價格分佈較土地平均，另外其變異數為 100.8608 也較土地為小。觀察其內容，除了其成交量能較高，其價格分佈似乎也較能達到收斂的效果。

經分析及評斷，草屯段房地部分同樣也夾雜住、商、工、農等分區，雖然實價登錄資料庫內容同樣為住宅區。但有些區域因周邊之相鄰效應，呈現類似住商混合型態。也因此在此較少數房地的交易點上，地價分佈差異會較大，但也由於成交量能較土地部分高，在搭配變異數及平均數、中位數分析上，能夠較為縮小價位的分歧。

以上結果分析，該模型未來在做電腦大量估價時，實價登錄資料庫不同使用類型的劃分，更能提高各項數據的相關性。因此建議未來在本模型案例的應用上，估價將偏向各地價區域之細分及獨立性，方能較完整呈現本區域合理價位。

小結:經由以上實證分析,若以實價登錄資料庫為未來電腦大量估價為背景,則各個不同區段間皆會產生各自不同的態樣。這也許是反映出不動產價格之多樣性,同時也呈現出區域及個別因素,確實會左右不動產實際的成交價格,這對於電腦大量估價上具參考意義。

伍、結論及建議

一、結論

經本研究分析結果，我們分別產製了土地及房地各個交易模型。從以上模型中我們經由變異數、中位數、平均數可以看出，在有些中位數與平均數接近的模型中，其變異數較大，有些則較趨近。這充分的反映出本研究的方式，確實可以解析實價登錄資料庫中，價格結構之關聯性及合理性。在未來若執行電腦大量估價時，資料庫數據也可先經由以此模型做初步解析，來決定此資料庫權值的釐定及參酌。

在以上實驗數據所呈現各樣本母體彼此間差異及分析可知，實價登錄資料庫確實能反映出各自不同區段間，不同的量價結構差異。換言之，若作業方式不洽當，其各自系統間的差異情形將無可避免，因此在進入大量估價前，資料庫內容結構的檢核及規整化，就顯得有其必要性。

本研究的目的是針對在一個跟區間有關的科學應用加以整理、分析，以解析成果，進而探討某些因子對資料庫量價結構的關聯性，並嘗試提供有相關此科學應用的人有更多的參考。

二、建議

透過以上分析，本研究有以下建議：

1. 依誤差傳播理論，實價登錄作業成果好壞，將直接影響未來估價的品質。
2. 在不同地價區段因其不動產資料結構及特性不同，往往是無法採單一模型去做統一解釋。
3. 因視各地區不動產交易資料的情況，來決定是否適用於那種估價模型，非全部區段皆適合電腦大量估價的方式。

陸、參考文獻

- 一、《台北都會區不同住宅類型價差之研究》台灣土地研究，2006，李泓見、張金鶚、花敬群。
- 二、不同變異數分析考驗的優缺點比較，1995，王俊明:15-30。
- 三、談談平均數, 眾數與中位數. 中學課程輔導(初二版) 12，2004，蔡連信:18。
- 四、統計學導論，2007，華泰圖書出版社，方世榮。
- 五、統計學的世界，2002，台北：天下遠見，莫爾（鄭惟厚譯）。
- 六、臺灣大量估價問題分析其及改進方法之研究，2004 年，海峽兩岸土地學術研討會論文，張梅英、施昱年。
- 七、距離最近的位置—談統計中位數的應用，1983，財經問題研究，馬家善：45-50。
- 八、算術平均數概念的四個理解水平及測試結果，2006，數學教育學報 15.3，梁紹君: 35-37.
- 九、平均數事前比較的重要性及其統計方法，1988，師大學報，范德鑫。
- 十、Moran, P., and C. Smith. "The correlation between relatives on the supposition of mendelian inheritance." *Transactions of the Royal Society of Edinburgh* 52 (1918): 899-438.