

群組決策系統中 ——以多梯次密度叢集技術進行共識趨勢分析之研究

廖文豪

國立台北商業技術學院資管系

摘要

本研究提出利用多梯次密度叢集技術建構線上共識趨勢分析系統，並藉以改善線上公民決策模式下所面臨之低共識程度問題，該系統可於每一決策梯次結束後，產生以前梯次之趨勢資訊，如各主流群體之分布圖、各偏好方案支持度排行、樞紐分析圖表及共識程度指標等，使決策者在充分掌握共識趨勢資訊下，並可藉此逐梯次調整個人之偏好以求得全體高共識高滿意度之決策方案。為驗證共識趨勢分析系統對共識程度提昇之效益，本研究以旅遊決策範例進行線上公民決策實驗並以因徑分析法探討其因果關係，結果顯示共識程度指標如一致性程度、期望落差及決策滿意度皆隨者梯次增加而產生直接而顯著之變化，也就是說，多梯次共識趨勢分析系統對共識程度之提昇確有顯著影響。

關鍵字：密度叢集法、共識趨勢分析、決策滿意度、決策一致性程度

Employing Multistage Clustering Techniques to Analyze the Consensus Trend in Group Decision System

Wen Hao Liao

Department of Information Management,
National Taipei College of Business

Abstract

In order to improve consensus degree in the models of civil decision-making, this study proposes to employ multistage clustering techniques to construct a consensus-trend analytical system. This system can create consensus trend information that is needed in the consensus reaching process after each decision stage. After perceiving consensus trend of major groups, decision makers may adjust their preferences to reach a conclusion of high-consensus and high-satisfaction stage by stage. For verifying the benefits of the system, this study implements an on-line decision making experiment and then explores the causal relationship by path analysis. The results indicate that the consensus indices, such as coincidence degree, expectation gap and satisfaction degree, bring about significant and direct variations while we increasing the number of decision stage. This confirms that the consensus-trend analytical system is helpful for improving the problem of low consensus.

Keywords: density-based clustering techniques, consensus trend analysis, satisfaction degree, coincidence degree

壹、概論

本研究之貢獻分成兩部分，第一部份為應用叢集技術建構「多梯次共識趨勢分析系統」，第二部份為以實驗法驗證在線上公民決策模式之應用下，「多梯次共識趨勢分析系統」所產生之效益。將資訊科技及通訊科技應用在公共行政的過程，稱之為電子化民主（electronic democracy）（Calabrese and Borchert 1996），為改善地方政府與公眾間之聯繫，中央及各地方政府紛紛利用資訊網路科技建置電子化政府（行政院研考會，民 87，葉俊榮，民 86，項靖，民 88），以鼓勵公眾參予公共事務。Macpherson（1997）認為透過網際網路不獨可以使民眾取得公共議題資訊，並可對這些議題進行線上討論、辯論及投票等，以實踐公民參與的理想。

在民主化社會中，許多公共政策之制定常需要廣泛徵詢民意，而後再依主流民意之趨向來下決策，此種訴諸民意的決策模式即稱為公民決策模式，其主要特性為：(1)決策者群龐大，因此共識不易凝聚，(2)公民並無「必須參與決策」之強制義務：公民可以任意選擇「參與」或「不參與」決策，因此若要使其廣泛參與，必須使「參與決策所花費之時間成本」減到最低，如提供線上公民決策系統，(3)決策問題具主觀性，不易客觀評量：此類決策問題如某公司舉辦旅遊時，旅遊地點及時間之甄選，或某政黨以民意取向選擇該黨之市長後選人等，(4)決策結果將影響決策者之權益：由於決策者群之自身利益受決策結果影響，因此達成共識之決策結果應以多數決策者滿意為目標。

就解決低共識程度之問題而言，在決策群龐大、可選擇方案多之線上公民決策環境下，若採用單一梯次決策法，通常都不易達成高共識程度，而若決策過程又無適當的共識導引方法，則共識達成之時間將會拖延甚至無法在時限內達成共識。關於共識導引方式，過去之文獻皆著重於主席（session manager, coordinator）之協調功能探討，然而在網路環境下之線上決策會議裡，主席不易發揮協調功能，因此本研究提出以多梯次共識趨勢分析系統以提昇共識程度（consensus degree）。

共識程度之計算方式有三種，(1)以支持度為導向之計算法：以最獲偏好方案所得之支持率為計算基礎（Zadeh 1983; Yager 1988），(2)以一致性程度（coincidence degree）為基礎之計算法：將決策者間意見的一致性程度視為共識程度（Kacprzyk 1988; Herrera 1996），(3)以決策滿意度為導向之計算法：以決策者之期望落差為計算基礎。由於此三種方法所得出之共識程度，因其計算基礎之不同而將產生相當之差異，有關共識程度指標之定義與衡量，請參見第參節。

為探討「多梯次共識趨勢分析系統」對於提昇共識程度指標是否產生顯著而直接之影響，因此，本研究建構一雛形系統供決策者操作，就系統功能而言，除了須具有決策輸入所需介面外，還須具有叢集分析、偏好差距計算、決策項之屬性量化、共識指標計算、共識趨勢分析結果呈現等功能。就系統建構之技術而言，主要是利用叢集技術將支持相近方案之決策者歸屬為一叢集進而形成所謂之「叢集分布圖」，使決策者得以藉此叢集分布圖掌握各「主流群體」之偏好趨向。

為證實使用共識趨勢分析系統對決策一致性程度、期望落差及決策滿意度等共識指標確實產生具體影響，本文採實驗研究法加以驗證，其結果以兩階段加以檢定，第一階段為多變量變異分析（MANOVA）檢定，第二階段為因果關係檢定（causal analysis），結果證實：(1)決策一致性程度隨者梯次增加而逐漸升高，兩者產生直接而顯著之正相關，(2)決策者之期望落差隨者梯次增加而逐漸降低，兩者產生直接而顯著之負相關，(3)決策滿意度隨者梯次增加而逐漸增加，兩者呈現間接之正相關。因此，多梯次共識趨勢分析系統對於共識程度之提昇，經研究證實具有顯著效益。

貳、多梯次共識趨勢分析系統

達成決策共識並產生決議為決策者在決策過程中被賦予之重要任務，就共識達成支援系統（consensus reaching support system）之建構而言，以往文獻主要在探討達成共識之程序、促進群體討論之工具、共識程度或一致性程度等指標之計算，然而有關共識趨勢分析系統之建構與影響則鮮有探討，本節將介紹系統建構有關之共識達成流程與模型、叢集化處理之技術、系統功能及資料模型，而有關此系統之影響評估則留待第參節以後之實驗研究再予以探討。

一、共識達成之流程與模型

Herrera (1996) 將群組決策程序（group decision making process）分成共識凝聚程序（consensus process）與篩選程序（choice process），其中，共識凝聚程序為一迴圈結構，決策群在此迴圈中不斷進行共識收斂工作，直到共識程度滿意為止。但此模型並未具體提出如何進行共識收斂，也未提出共識趨勢分析方法，因此本研究將此模型予以更進一步詳細化及具體化。

本研究提出之共識達成架構如圖 1，輸入端為可選擇方案，輸出端為達成共識後之結果方案，決策者藉著此模型將輸入轉換成輸出需經過多梯次之共識收斂程序作業，每一梯次之作業分成兩階段，此兩階段之運作程序如下：

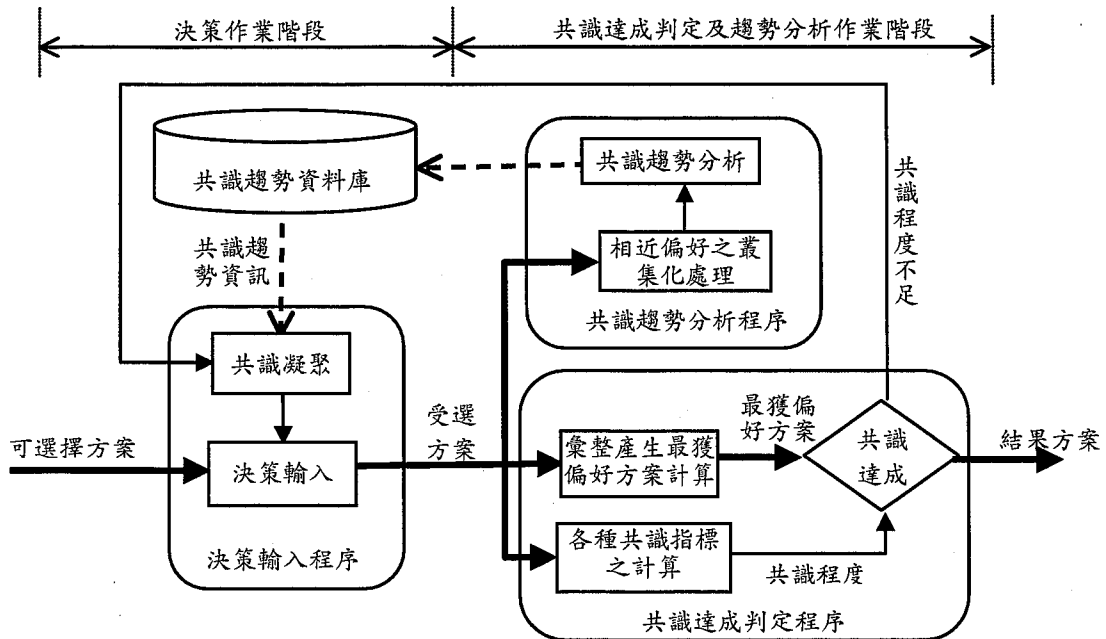


圖 1：共識達成之程序架構圖

(一) 決策作業階段：

決策者在決策作業階段之主要工作是取得共識趨勢資訊，之後再選取個人偏好之方案，被選取之方案集合稱之為受選方案。「共識凝聚」模組主要是協助決策者取得上一梯次之共識趨勢資訊，如各決策群之叢集分布圖、各方案支持度排行、平均期望落差及一致性程度等，之後，決策者在了解共識趨勢後，在所謂「多數影響」之驅策下，將可能調整個人偏好以尋求有利於達成整體共識又能兼顧個人偏好之另一方案。

依據 Asch 等之研究 (Asch 1951; Baker 1994)，決策者改變決策時有傾向主流意見之認同趨勢，此現象稱之為多數影響 (majority influence) 或多數力量 (strength in number)，依據 Asch 之研究結果顯示群組愈大此種歸順多數之效應愈強。雖然部分文獻亦提出少數影響 (minority influence) 之論述與實證 (Moscovici & Personnaz 1980)，然而其發生之原因主要在於少數處於高位者之獨裁或武斷 (assertive) 行為，亦有部份屬於少數者之強迫性爭辯 (compelling argument) 行為。就決策模式而言，在面對面或記名投票之決策會議中，少數影響之力量較明顯，然而在線上公民決策模式中，則是多數影響較為明顯。因此，雖然在公民決策之環境下欲找出高共識且高滿意度之決策結果極為困難，但只要提供決策者充分之共識趨勢資訊，再藉由群眾之多數影響力，則可驅使決策者在改變決策時願意選擇高共識高滿意度之方案。

(二) 共識達成判定及趨勢分析作業階段：

受選方案經彙整後可得出最獲偏好方案，若共識程度達到共識門檻，則最獲偏好方案即成為結果方案；若未達共識則決策者必須再進行另一梯次之決策程序。共識達成判定及共識趨勢分析兩程序之運作皆在此作業階段進行，說明如下：

共識達成判定程序

此程序中之「共識指標計算」模組，其功能是計算全體決策者對結果方案之平均期望落差、一致性程度及決策滿意度等共識程度指標，若決策全體對於最獲偏好方案之共識程度不足，則決策群組必須繼續進行下一梯次之共識收斂程序，若決策全體對於最獲偏好方案之已達共識，則共識收斂程序結束並宣告最獲偏好方案即為結果方案。

共識趨勢分析程序

此程序中之「叢集化處理」模組，其功能是將決策者予以分群，所用之方法是透過叢集技術將支持相近方案之決策者予以群聚並歸屬於同一個叢集，其叢集結果以偏好密度叢集分佈圖呈現。「共識趨勢分析」模組之功能是將叢集處理後之結果進行趨勢分析，而後再將分析後之結果存入共識趨勢資料庫內以供決策者進行下一梯次共識收斂之參考。

二、叢集化處理

(一) 叢集技術之分類

現存幾種常用之叢集技術如下：

1. 區塊分割法 (partition methods)：其方法是將資料庫內所有個體 (objects or data tuples) 分成 k 個叢集 (或群組)，每一個分割區塊表示一個叢集，若同一叢集內之個體彼此靠近而相異叢集間之個體彼此遠離則視為良好分割。區塊分割法主要分成兩種計算方法，其一為 MacQueen (1967) 所提出之 k -means 演算法，另一為 Kaufman 及 Rousseeuw (1990) 所提出之 k -medoids 演算法。
2. 階層法 (hierarchical methods)：其方法是將資料庫內所有個體進行階層式分解，最後產生一棵叢集樹，其中每一個節點即為一個叢集。階層叢集法主要分成結塊法 (agglomerative) 及分裂法 (divisive) 兩種演算法，結塊法是以由下而上 (bottom-up) 之不斷合併程序來進行叢集作業，其中較知名者如 AGNES (Agglomerative NESTing) 演算法；分裂法則是以由上而下 (top-down) 之不斷分裂程序來進行叢集作業，其中較知名者如 DIANA (Divisive ANALysis) 演算法 (Kaufman & Rousseeuw 1990)。
3. 密度基礎法 (density-based methods)：密度基礎法定義叢集為一高密度連結點所成之集合。Ester (1996) 等人所提出之 DBSCAN (Density Based Spatial Clustering of Applications with Noise) 演算法，只要某資料點之鄰近區域的密度超過門檻值 (density threshold)，則將此鄰近區域納入此叢集，最後，不歸屬於任何叢集之資料點視為雜訊；由於 DBSCAN 需輸入鄰近半徑及最少資料點兩參數，然而輸入參數之微小變動可能導致叢集結構之巨大變動，應如何選擇適當之參數才能得出較佳之叢集結構往往需要依賴經驗，因此 Ankerst (1999) 等人提出 OPTICS (Ordering Points To Identify the Clustering Structure) 演算法以產生叢集結構與輸入參數之對應關係圖，使得分析者更容易全面性且精確地控制輸入參數。另外

Hinneburg 及 Keim (1998) 提出 DENCLUE (DENSity-based CLUstering) 演算法以密度分佈函數為基礎以產生叢集結構。

(二) 叢集技術之選擇與考慮

就共識趨勢分析系統之需求而言，叢集法之選擇應考慮表 1 所列之 4 項因素，因此，本研究採用能滿足表 1 所列四種需求之密度叢集法，說明如下：

表 1：叢集法之選擇因素列表

比較項目	區塊分割法	階層法	密度基礎法
高資料量之處理能力	Yes	No	Yes
雜訊過濾能力	No	No	Yes
最適叢集數之自動計算能力	No	No	Yes
不規則叢集形狀之發現能力	No	No	Yes

資料整理：本研究

1. 高資料量之處理能力：線上公民決策型態所具有的特性之一是參與決策者多且可選擇量大，因此所產生之決策資料數量多，故必須選擇能處理龐大資料量之叢集方法。而階層法處理高資料量時，將特別耗時 (Kaufman & Rousseeuw 1990)。
2. 雜訊過濾能力：在共識凝聚過程中，當極少數人所堅持之意見並不被多數人認同時，若這類決策者不願改變原先之意見反而不斷發出干擾信息，則會使決策全體難以達成最終共識。因此，若能將低於雜訊門檻值的個體予以過濾，則有助於整體共識之達成。密度叢集法可以過濾雜訊，區塊分割法及階層法則否 (Hinneburg & Keim 1998)。
3. 最適叢集數之自我發覺能力：決策系統必須有能力將相似意見之決策者予以自動分群，以便使所有決策者清楚知道目前之群組數目與分布、各群組之屬性特徵及本身隸屬於那一個群組，繼而判斷在下一梯次決策時自己應採何種決策行為。區塊分割法並無法自行找出最適之叢集數目，因此必須由使用者輸入叢集數 (Kaufman & Rousseeuw 1990)；階層法無論是使用分裂法或結塊法都無法找出最適之叢集數。密度叢集法可以找出具有連續高密度之連結區塊，此高密度之連結區塊即為一個叢集 (Hinneburg & Keim 1998)，各叢集間則以密度斷層作為分界，密度斷層指的是密度低於門檻值的連結區，因此用密度法所產生之叢集數目為符合實際之數目。
4. 不規則叢集形狀之發現能力：區塊分割法所產生之叢集是以 mean 或 medoid 為中心畫出一橢圓形為範圍 (MacQueen 1967; Hinneburg & Keim 1998)，不同叢集間容易出現重疊現象，以致在重疊區之資料點，其歸屬品質變得極差，因此，無法依照實際群聚分布而產生分隔方案數。然而密度叢集法可依實際分布狀況產生不規則形狀之叢集，叢集間也不會產生重疊現象 (Hinneburg & Keim 1998)。

(三) 共識問題領域之限制與名詞定義

整個共識問題領域包含可選擇方案、決策者、偏好值三個維度，而此三維度本身又

有複雜之組成結構，因此欲分析此複雜之共識問題，研究者必須將此問題領域作適當之限制並對相關名詞予以清楚定義。本研究針對共識問題領域所作之限制及名詞定義說明如下：

1. 共識問題領域之限制

- (1) 決策者之限制部分：在線上公民決策環境下，決策者不分職位等級，其決策權重值皆均等。就決策者之群組架構而言，決策者是依照其偏好選擇予以分群，並不依決策者之職位、所屬單位、專長等屬性建立群組結構。
- (2) 可選擇方案之限制部分：為了進行叢集處理，所有之方案必須以空間 (spatial) 座標表示，因此必須分析各方案內涵並選取重要屬性中具有數值或次序型態者為座標軸，並以此座標軸建立可選擇方案之決策空間。

2. 名詞定義

- 決策空間：若將所有的可選擇方案以座標空間表示，則稱此空間為決策空間，例如可選擇方案具 A、B 兩個屬性，則其所顯示之決策空間為二維平面空間如圖 2。
- 決策點：決策空間上之座標點稱為決策點，每一個決策點亦代表 1 個可選擇方案，例如圖 2 中共有 100 個決策點。
- 支持者：位於決策點之所有資料點即為支持該決策點之決策者。例如圖 2 中各決策點上之黑色圖點即為其支持者，當決策者在決策空間上進行偏好選擇後，其偏好選擇結果即以此黑色圖點呈現於決策空間上。
- 支持度：決策點上之資料點數目即為該決策點之支持度，此支持度可視為支持該方案之人數。例如圖 2 中，決策點 (a7, b8) 之支持度為 4。
- 決策叢集：決策叢集為一連續高密度之不規則區塊，每一個叢集代表 1 個決策群，如圖 3 所呈現之兩大決策叢集。
- 群聚量：某一叢集所獲得之支持度總和稱之為群聚量。
- 密度：一個決策點之密度為其對所有決策者之影響力總和，請參見隨後定義之密度及影響力函數。
- 引力中心 (attractor)：在某一叢集內，若某一決策點具有局部最大密度，則稱此決策點為此叢集之引力中心。

(四) 密度之計算

Hinneburg 及 Keim (1998) 引用 Gaussian function 建立影響力函數並提出 DENCLUE 密度叢集演算法，為將其應用於共識問題領域之決策空間中以便描述某決策點對決策者之影響力，本研究將決策點 x 對位於決策點 y 之決策者的影響力函數 (influence function) 定義如下：

$$f_{Gauss}(x, y) = e^{-\frac{d(x, y)^2}{2\sigma^2}} \dots \dots \dots (1)$$

d(x, y) 為 x 與 y 兩決策點間之距離，y 若距離 x 愈近則表示其受決策點 x 之影響力愈

大，影響力意指“匯聚”或“吸引”之能力，某決策點具高影響力，表示該決策點能吸引許多決策者向其靠攏，因此造成資料點之群聚現象，此現象稱之為叢集效應。

密度函數為決策點對所有決策者之影響力總和，定義如下：

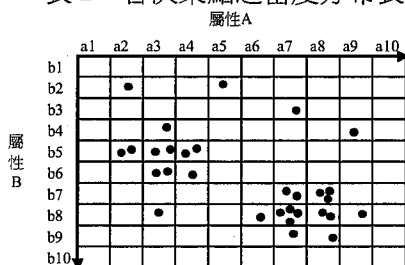
$$f_{Gauss}^D(x) = \sum_{i=1}^n e^{-\frac{d(x,x_i)^2}{2\sigma^2}} \dots\dots\dots (2)$$

此函數可以解釋為“決策點 x 對決策空間 D 內所有 n 個決策者（或資料點）之影響力總和”，其中，決策空間 $D = \{x_1, \dots, x_n\}$ ， σ 為密度參數（density parameter）。

（五）將決策者之偏好轉換成密度叢集圖

圖 2 為 30 位決策者針對 100 個可選擇方案所作之偏好選擇，假設可選擇方案具有兩個屬性，分別為屬性 A 及屬性 B，由圖中可看出決策點（a7,b8）之支持度為 4，決策點（a3,b5）之支持度為 2。

表 2：各決策點之密度分布表



	a1	a2	a3	a4	a5	a6	a7	a8	a9	a10
b1	0.0	0.1	0.0	0.0	0.1	0.0	0.0	0.0	0.0	0.0
b2	0.1	1.0	0.1	0.1	1.0	0.2	0.1	0.0	0.0	0.0
b3	0.0	0.2	0.2	0.0	0.1	0.2	1.0	0.2	0.1	0.0
b4	0.0	0.4	1.3	0.4	0.0	0.0	0.1	0.2	1.0	0.1
b5	0.3	2.3	3.0	2.5	0.3	0.0	0.0	0.0	0.1	0.0
b6	0.0	0.6	2.5	1.6	0.2	0.0	0.3	0.4	0.1	0.0
b7	0.0	0.0	0.3	0.2	0.0	0.5	3.0	3.6	0.6	0.0
b8	0.0	0.0	0.1	0.0	0.1	1.6	4.9	3.3	1.3	0.1
b9	0.0	0.1	1.0	0.1	0.0	0.3	1.7	1.5	0.3	0.0
b10	0.0	0.0	0.1	0.0	0.0	0.0	0.2	0.2	0.0	0.0

圖 2：30 位決策者在決策空間之偏好選擇範例

表 2 及圖 3 分別為各決策點之密度分布表與分布圖，此兩圖表是圖 2 經過密度函數計算後所得之結果。若設定雜訊過濾門檻為 1.2 則可將其中 5 個無法形成叢集之資料點去除，則僅剩餘兩個大叢集，圖 4 為由圖 3 俯視所得之等高線圖，由圖中可清楚看出此兩個決策叢集，每一個叢集代表一決策群，左上角所呈現之叢集是以（a3,b5）為引力中心，其群聚量為 10；右下角所呈現之叢集是以（a7,b8）為引力中心，其群聚量為 15，此為最大叢集。因此藉者此密度叢集圖，決策者便可觀察決策群之分布概況並得出各叢集之引力中心及群聚量等資訊。

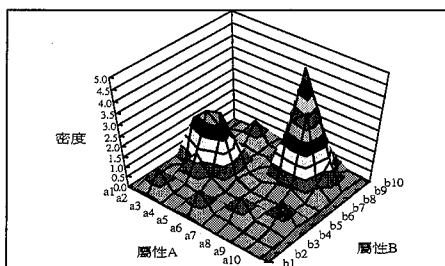


圖 3：偏好密度叢集圖

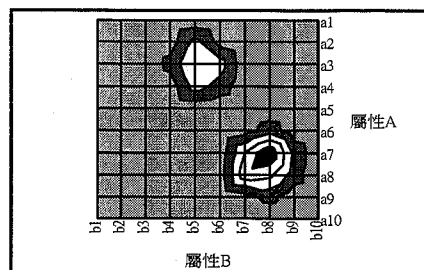


圖 4：各密度叢集之等高線圖

三、系統之功能與資料模型說明

Zadrozny (1993) 認為共識達成支援系統不僅要能促進群組討論，且當討論過程中專家意見改變時，還要能測量偏好之變化。就測量偏好變化之功能而言，本研究所建構之系統不僅具有呈現一致性程度、期望落差及決策滿意度等共識程度指標之功能，還可提供叢集分布圖給決策者，使其能掌握全體決策者於各梯次決策後之偏好變化。本實驗系統所使用之決策範例為受測者所熟悉之「旅遊決策」，相關之實驗設計與程序說明請參見第四節，系統主要功能如下：

(一) 前端使用者

1. 決策者之基本資料輸入，如決策者姓名、年齡、性別等
2. 決策者之偏好權值輸入：決策者依其偏好來設定各決策項屬性之權值，如 $W_{quarter}$ 、 W_{month} 、 W_{week} 、 W_{city} 、 W_{region} 、 W_{type} 、 $W_{fee-rank}$ ，權值大小顯示決策者對該屬性之重視程度。
3. 決策輸入：於各決策梯次進行時，由決策者輸入其所選擇之方案。
4. 線上分析處理：提供決策者對以前梯次之決策結果進行相關之樞紐分析、資料方塊圖或彙總表分析功能。
5. 叢集分布圖分析：系統將以前梯次之決策結果以叢集分布圖呈現，使決策者藉以觀察並掌握主流群體之分布與動向。
6. 各方案所獲支持率之排行表呈現：使決策者得以了解，於各梯次中，獲偏好方案之支持率變化。
7. 各梯次之共識程度指標之呈現，如決策滿意度、一致性程度、平均期望落差等指標。

(二) 後端（管理者之線上管理）

1. 決策方案、決策者等基本資料檔建構與維護。
2. 決策項之屬性設定。
3. 決策梯次設定：如各梯次之共識程度門檻、起始及截止時間。
4. 決策限制條件之管理：針對各決策項屬性進行欄位限制（constraints）。如費用限制、特定時間之限制等。
5. 共識資訊之彙整：各梯次截止後，彙整決策資料以產生各種決策者所需之共識資訊。

另外，決策過程所需之旅遊資訊及相關應用軟體，決策者可透過網際網路任意取得。

(三) 資料模型

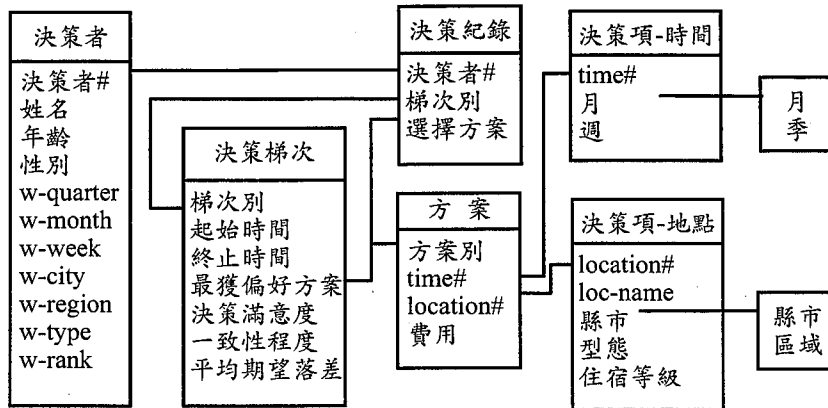


圖 5：系統之作業資料庫模型

季 (quarter)：旅遊季【1~4】，月 (month)：旅遊月份【1~12】

週 (week)：當月份之星期序號【1~5】

區域 (region)：旅遊地之區域【N：北部，M：中部，S：南部，E：東部】

型態 (type)：旅遊地之型態【A：山林湖泊，B：濱海，C：健康休閒中心】

住宿等級 (fee-rank)：依住宿費用由低至高進行排序，住宿等級愈高收費愈高。

(四) 共識趨勢線上分析處理系統之星狀結構圖

本研究所建構之線上分析模型為一星狀結構如圖 6，每一決策梯次結束後，作業資料庫之內容經過萃取 (extracting)、轉換 (transferring) 程序後，即可載入 (loading) 此資料庫供決策者進行各種線上分析。

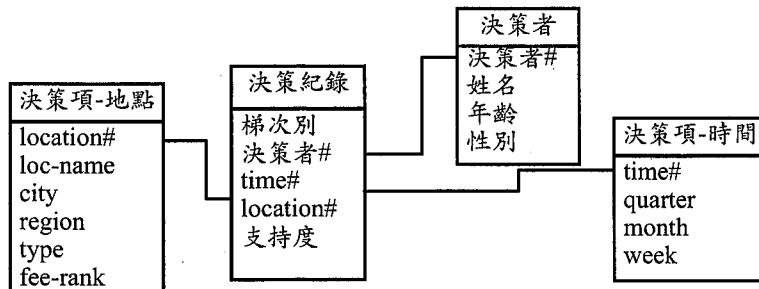


圖 6：星狀結構圖

參、實驗研究之模型與假說

本實驗研究旨在探討「多梯次共識趨勢分析處理系統」對於提昇共識程度是否產生顯著貢獻。當決策者使用「多梯次共識趨勢分析處理系統」取得共識趨勢資訊後，由於個人期望與各主流群體之期望，彼此之間會產生拉鋸現象，若採用傳統「單一梯次」之決策方式，則要達到高共識程度將極為困難。為改善此問題，本研究以「多梯次」決策方式來進行實驗，每一梯次結束後，系統會將決策結果經分析處理過程後轉換成共識趨勢資訊，供決策者進行下一梯次之參考。當決策者充分掌握得全體之共識趨勢資訊後，在「多數影響」效應影響下，決策結果將可能會產生幾種效果，(1)整體期望落差將會逐梯次降低，(2)整體一致性程度將會逐梯次升高，(3)由於整體期望落差降低以及一致性程度之提昇兩因素影響，將會促使決策滿意度逐梯次升高。因此，本研究提出如圖 7 之因果模型 (causal model)，其中，H1~H6 分別代表假說一至假說六。

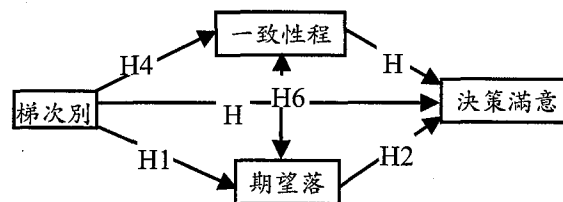


圖 7：本研究之因果模型

一、決策滿意度及期望落差之定義、衡量與假說

滿意度之構念出現在在各種領域中，尤其是行銷學中的顧客行為之研究文獻 (Tsiros 1998; Voss Parasuraman & Grewal 1998; Tears 1993.)，顧客滿意度為期望效能與認知效能之函數，即 $\text{satisfaction} = f(\text{expectation}, \text{perceived performance})$ ，顧客對產品之期望效能與所認知之實際效能間之差距稱為期望落差，研究文獻之實證結果顯示顧客滿意度與各類期望落差之間呈現明顯的負相關。在決策領域中，決策滿意程度是比較過程後所產生的結果 (Yi, 1990)，因此決策滿意度為期望落差之函數，期望落差指的是決策者所選方案與決策結果間之偏差量，此偏差量可以解釋為“個別”決策者之期望與“大多數”決策者所選擇結果之落差。

$$G_i = d(X_i, X_{\text{most}})$$

其中， G_i 為決策者 i 之期望落差， X_i 表示決策者 i 所選擇之方案， X_{most} 為最獲偏好方案，亦即最多人支持之方案， $d(X_i, X_j)$ 為兩方案之偏好差距函數。若決策方案包含 K 個屬性， $A_1, A_2, \dots, A_K$ ，則 X_i 及 X_{most} 可表示如下：

$$X_i = \langle A_{1i}, A_{2i}, \dots, A_{Ki} \rangle$$

$$X_{\text{most}} = \langle A_{1\phi}, A_{2\phi}, \dots, A_{K\phi} \rangle$$

二、偏好差距計算與相關屬性之量化

由於決策者對於決策項下各種屬性之重視程度不一，因此必須針對各屬性設定其偏好權值 (preference-weighting)，才能較精確地計算出該決策者對決策結果所產生之期望落差值。 ω_{A_1} 、 ω_{A_2} 、 $\dots$ 、 ω_{A_K} 為各屬性之偏好權值，決策者可針對此個人偏好設定權值，如某決策者比較重視 A_1 及 A_2 屬性者，則可將 ω_{A_1} 、 ω_{A_2} 調高，並將其他屬性權值調低。為計算距離值，本研究之權值處理採平方和法，故權值之平方和必須為 1。

$$\sqrt{\omega_{A_1}^2 + \omega_{A_2}^2 + \dots + \omega_{A_K}^2} = 1 \dots\dots\dots (3)$$

當所有決策者輸入其權值後，系統便可產生全體之平均權值 ($\bar{\omega}_{A_1}$ 、 $\bar{\omega}_{A_2}$ 、 $\dots$ 、 $\bar{\omega}_{A_K}$)；依據此平均權值可進一步計算得決策者 i 之期望落差如下：

$$G_i = \sqrt{(\bar{\omega}_{A_1}^2 \times d(A_{1i}, A_{1\phi})^2 + \dots + \bar{\omega}_{A_K}^2 \times d(A_{Ki}, A_{K\phi})^2)} \dots\dots\dots (4)$$

本研究採用 Han&Kamber 差距量化法，定義差距函數 $d(A_{ki}, A_{kj})$ 如下：

1. 若屬性之資料型態為名目資料 (nominal variable)，則其差距值計算如下：

if $A_{ki} = A_{kj}$ then $d(A_{ki}, A_{kj}) = 0$ else $d(A_{ki}, A_{kj}) = 1$

其中， A_{ki} 為方案 i 之第 k 個屬性； A_{kj} 為方案 j 之第 k 個屬性。兩名目資料間之比較，只能產生「相同」及「相異」兩種結果，而無法產生線性距離，如旅遊地之「型態」屬性，其將旅遊型態分為山林湖泊、濱海、健康休閒三類，若兩方案所選擇之旅遊地型態不同則此屬性之差距為 1，反之為 0。

2. 若屬性之資料型態為區間數值 (interval-based numerical variable)，則將其轉成比例值 (ratio data)。此種轉換之目的是為了使各屬性具有相同之比較基準，俾使每一個屬性差距必須介於 0 到 1 之間。其差距值計算如下：

$$d(A_{ki}, A_{kj}) = \frac{|A_{ki} - A_{kj}|}{\max(A_K) - \min(A_K)} \dots\dots\dots (5)$$

3. 若屬性之資料型態為次序值 (ordinal variable)，則須將其轉換為 Z_{ki} ，使其值成為介於 [0,1] 之區間數值，再利用上述運算式 (5) 計算得 $d(A_{ki}, A_{kj})$ ，其中， M_K 為最大次序值。

$$Z_{ki} = \frac{A_{ki} - 1}{M_K - 1} \dots\dots\dots (6)$$

決策者之滿意度之 Gaussian 函數如下：

$$Satisfactory = \exp\left(-\frac{G^2}{2\sigma^2}\right) \dots\dots\dots (7)$$

決策者之滿意度與期望落差率兩者呈現指數曲線之關係，當期望落差為 0 時，也就是決策者之期望與全體決策結果完全相同時，決策者具有最高滿意度，其值為 1；當期望落差為最大值 $G_{\max}$ 時，決策者具有最低滿意度。當期望落差率在 20% 以內，決策者可以獲得高滿意度；當期望落差率增至 20% 以上時，其滿意度呈快速下降之趨勢；當期望

落差率增至 80% 以後，其滿意度已漸趨於最低滿意度，亦即與漸近線 $f(x)=0$ 貼近。

「共識趨勢分析處理系統」主要功能在於提供使用者有關整體決策之共識資訊，若決策者能掌握這些資訊，將可逐步調整個人偏好使其與整體偏好漸趨一致，因此，整體期望落差將逐步縮小，而整體之決策滿意度也將逐步提昇，因此，研究者提出如下之假說一、二、三。

假說一：在決策者使用「多梯次共識趨勢資訊系統」的情況下，決策者之期望落差將隨著決策梯次增加而逐步降低，亦即兩者之間具有顯著而直接之負相關。

假說二：決策者之期望落差愈小，則決策滿意度會愈高，亦即兩者之間具有顯著而直接之負相關。

假說三：在決策者使用「多梯次共識趨勢資訊系統」的情況下，決策者對決策結果之滿意度將隨著梯次增加而提高，亦即兩者之間具有顯著而直接之正相關。

三、決策一致性程度之定義、衡量與假說

「共識程度」為群體達成共識之指標，有文獻將決策之「共識程度」等同於「一致性程度」(Kacprzyk 1997; Herrera 1996)，「一致性程度」意涵「整體決策者意見相近的程度」，關於一致性程度之計算，Kacprzyk 採用數值測量法，而 Herrera 採用模糊量詞 (fuzzy) 測量法來計算一致性程度，然而，兩者皆是透過所選擇方案間兩兩比較 (pair-wise comparison) 的差值總和為計量基礎。本文採用數值測量法，將一致性程度定義如下：

$$\begin{aligned}
 C_{ij} &= 1 - d(X_i, X_j) = 1 - \sqrt{\omega_{A1}^2 \times d(A_{i1}, A_{j1})^2 + \dots + \omega_{AK}^2 \times d(A_{Ki}, A_{Kj})^2} \\
 C_i &= \sum_j \frac{C_{ij}}{m} = 1 - \sum_{j=1}^m \frac{\sqrt{\omega_{A1}^2 \times d(A_{i1}, A_{j1})^2 + \dots + \omega_{AK}^2 \times d(A_{Ki}, A_{Kj})^2}}{m} \\
 C &= \sum_i \frac{C_i}{m} = 1 - \sum_{i=1}^m \sum_{j=1}^m \frac{\sqrt{\omega_{A1}^2 \times d(A_{i1}, A_{j1})^2 + \dots + \omega_{AK}^2 \times d(A_{Ki}, A_{Kj})^2}}{m^2} \quad \dots\dots (8) \\
 &= 1 - \frac{2}{(m-1)^2} \sum_{i=1}^{m-1} \sum_{j=i+1}^m \sqrt{\omega_{A1}^2 \times d(A_{i1}, A_{j1})^2 + \dots + \omega_{AK}^2 \times d(A_{Ki}, A_{Kj})^2} \quad \because d(X_i, X_j) = d(X_j, X_i)
 \end{aligned}$$

其中， C_{ij} 為兩決策者間之一致性程度，若決策者 i 所選擇之方案與決策者 j 所選擇之方案相近時，則兩決策者間之一致性程度 C_{ij} 高，反之。 C_i 為決策者 i 相對於全體決策者之一致性程度； C 為全體決策者之一致性程度； $0 \leq C \leq 1$ ， m 為決策者人數。為探討當決策梯次增加時，決策者間之一致性程度是否受到影響，而提昇一致性程度對決策滿意度是否有幫助，因此，本研究提出假說四及假說五。

假說四：在決策者使用「多梯次共識趨勢資訊系統」的情況下，決策者間之一致性程度將隨著決策梯次增加而逐步提高，亦即兩者之間具有顯著而直接之正相關。

假說五：提昇決策者之一致性程度，將促使決策滿意度升高，亦即兩者之間具有顯著而直接之正相關。

四、期望落差與一致性程度間之互逆性因果關係與假說

一致性程度與期望落差皆為重要之共識指標，然而，兩者相互間之影響，鮮有文獻提及，就兩者之定義來分析，一致性程度可視為決策者之偏好與「全體平均偏好」之差值函數，當決策者偏好愈接近「全體平均偏好」時，該決策者之一致性程度會愈高；而期望落差則為決策者偏好與「最獲偏好方案」之差值，也就是說，一致性程度以「全體平均偏好」為比較基礎，而期望落差是以「最獲偏好方案」為比較基礎，因此，當「全體平均偏好」接近於「最獲偏好方案」時，一致性程度會大量提昇，且期望落差會大量降低，為探討兩者間之相互影響，本研究提出如下之假說六，此假說為一「互逆性」(reciprocal)或「非回溯性」(non-recursive)因果關係(1976; Heise, 1975; Namboodiri et al., 1975)，亦即兩變數間之影響並非單向性而是雙向性。

假說六：一致性程度與期望落差兩者具有負向之互逆性關係；也就是說，提昇（降低）決策者之一致性程度，將促使決策者之期望落差降低（提昇）。

肆、實驗設計與程序

為了能線上收集假說驗證所需之基本資料，研究者建置決策趨勢資訊分析系統後，進行本節所述之實驗程序，受測者為資管科系高年級學生，實驗範例為受測者所熟悉之旅遊活動決策。進行實驗之時間主要是利用研究者在技術學院資管系日夜間部教授資料庫課程，當講解資料倉儲章節有關線上資料分析處理(OLAP)時，以此實驗給予學生操作。

選擇資管科系高年級學生為受測者之原因是：(1)樣本容易取得且容易集中於電腦教室進行實驗，(2)可減少教育程度、年齡、職級及文化差異等外生變數之影響，(3)資管科系高年級學生之電腦操作能力較佳且一致，不致因不熟悉工具而造成實驗結果之偏差。

實驗範例選擇旅遊活動決策之原因是：(1)受測者參加過多次旅遊活動，對此種決策之目標非常熟悉，(2)決策項清楚明確，相關決策項之屬性不外乎旅遊地點名稱、費用、型態、日期、季節等，(3)此種決策主要取決於受測者之個人偏好且決策者不需特殊決策專長，因此接近公民決策型態，(4)國內旅遊之資訊容易取得。

本研究之實驗架構如圖 8，實驗變數為「決策梯次數」，此變數之操弄方式，主要是將決策分成數個梯次進行，每一梯次決策前，系統將提供受測者以前梯次之共識資訊，而於每一梯次決策後，系統將產生新的共識趨勢資訊，並計算決策滿意度以判定是否達成共識，因此，共識趨勢資訊的充分程度將隨著梯次數增加而增加，此實驗設計之目的在於探討增加決策梯次對於期望落差、一致性程度及決策滿意度等共識指標是否產生顯著影響。

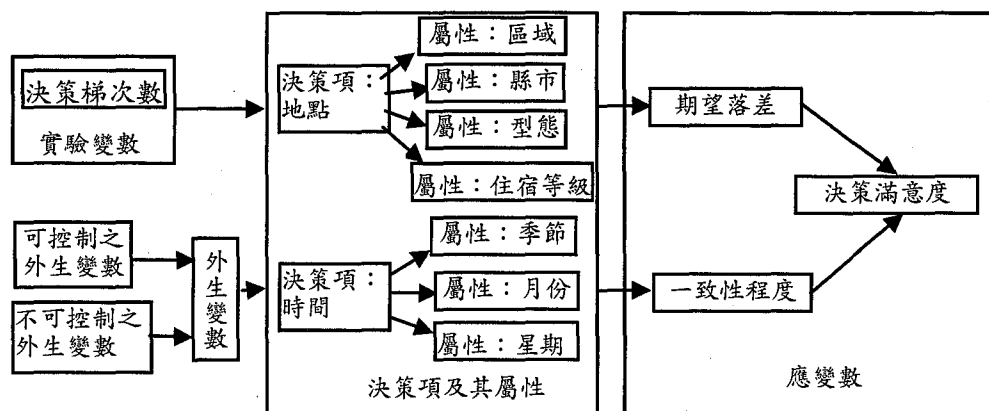


圖 8：實驗架構

一、外生變數之控制與隨機化（或均佈化）

Alavi 和 Joachimsthaler (1992) 應用彙總分析法來歸納出決策者滿意度影響因素，列出幾種主要因素，(1)人口統計變數：決策者之年齡、教育程度、性別，(2)決策者之人格屬性：武斷、好冒險、易於焦慮、追求成就等，(3)決策者之認知型態：分析型/直覺型、解決問題型，(4)情境變數：使用 DSS 之熟悉程度與經驗、參與時間。另外有文獻提出種族文化差異 (Landman 1987)、責任感 (Frijda et al., 1989) 等影響因素。

為了盡量減少外生變數對應變數之影響，本實驗從兩方面著手，(1)將可控制之外生變數保持為定值。其做法是挑選具有相近之年齡、教育程度、文化背景、電腦操作技能、決策責任、決策參與率的受測者，(2)將不可控制之外生變數影響均勻分布化（或隨機化）：不可控制之外生變數主要為人格屬性及認知型態，若受測者同屬某一類人格屬性或認知型態，將造成實驗結果之偏頗。為避免此種現象，本研究提高受測者人數，並使受測者廣含各類型之人格屬性及認知型態。本實驗對外生變數之控制如下：

- (一) 決策者之文化、年齡、教育差異之控制：選擇之受測者，國籍相同，年齡相近，接受教育之過程皆類似。
- (二) 決策者職位或責任感之控制：本研究選擇之受測者職位相同，皆為學生，因此無法利用職位影響他人之決策。受測者皆非主辦單位或執行決策單位，故無須承擔決策失敗之責任或懲罰，因此可大量降低 Frijda 所提之責任感變異。
- (三) 使用決策工具的熟悉程度之控制：受測者皆為資管系高年級學生，對於各種網路連線之資訊系統操作皆相當熟練，可輕易操作本實驗所用之共識趨勢資訊系統，因此，上述之「情境變數」所可能產生之變異可大量降低。
- (四) 參與程度之控制：在參與決策領域裡，有學者研究證實參與程度高者，其決策滿意度亦相對提高 (Robey 1989; James 1994)。為控制此外生變數，本研究於各梯次給予受測者之參與決策時間、參與方式皆相同。

二、實驗資料之取得方式

- (一) 決策項資料之取得：研究者選擇決策項資料時，力求擴大屬性變異，其目的是藉著擴大決策者之偏好選擇範圍，使偏好產生較顯著之差異，如此將有助於擴大應變數之變異。所選擇之旅遊地點由研究者選出 33 處，挑選原則：(1) 住宿、遊憩空間及交通等必須符合參加人數之容量需求，(2) 為了達到擴大屬性變異之目的，因此相關屬性如旅遊型態、住宿等級、縣市及所屬區域皆力求均佈，如旅遊地之型態中，山林湖泊類有 11 個，濱海類有 9 個，健康休閒中心類有 12 個；住宿等級屬性值是依住宿費用排序，含括簡陋至高級五星級飯店；縣市屬性含括北、中、南各區域。可選擇之時間屬性為 1~12 月、1~4 季。
- (二) 各屬性偏好權值之取得：受測者決策前須先上線輸入個人之基本資料〈身分 ID、姓名、性別、年齡〉及屬性偏好權值。由於每一位決策者對決策項屬性之重視程度不一，因此產生不同之屬性偏好權值 (individual preference-weighting)，個人偏好權值有七個，分別為 ω_{quarter} 、 ω_{month} 、 ω_{week} 、 ω_{city} 、 ω_{region} 、 ω_{type} 、 $\omega_{\text{fee-rank}}$ 。為了方便決策者輸入，每一個輸入權值 W_i 介於 0 到 10，但為了滿足各權值平方和必須為 1 之限制，因此系統會進行如下之正規化處理。

$$\omega_i = \frac{W_i}{\sqrt{W_1^2 + \dots + W_n^2}}$$

個人權值再經過彙整後，可計算得平均權值 ω_{quarter} 、 ω_{month} 、 ω_{week} 、 ω_{city} 、 ω_{region} 、 ω_{type} 、 $\omega_{\text{fee-rank}}$ 。

3. 決策資料之取得：受測者於各梯次指定時間範圍內線上輸入決策資料，系統將其存入決策紀錄表〈決策者#、梯次別、選擇方案〉。

三、實驗程序

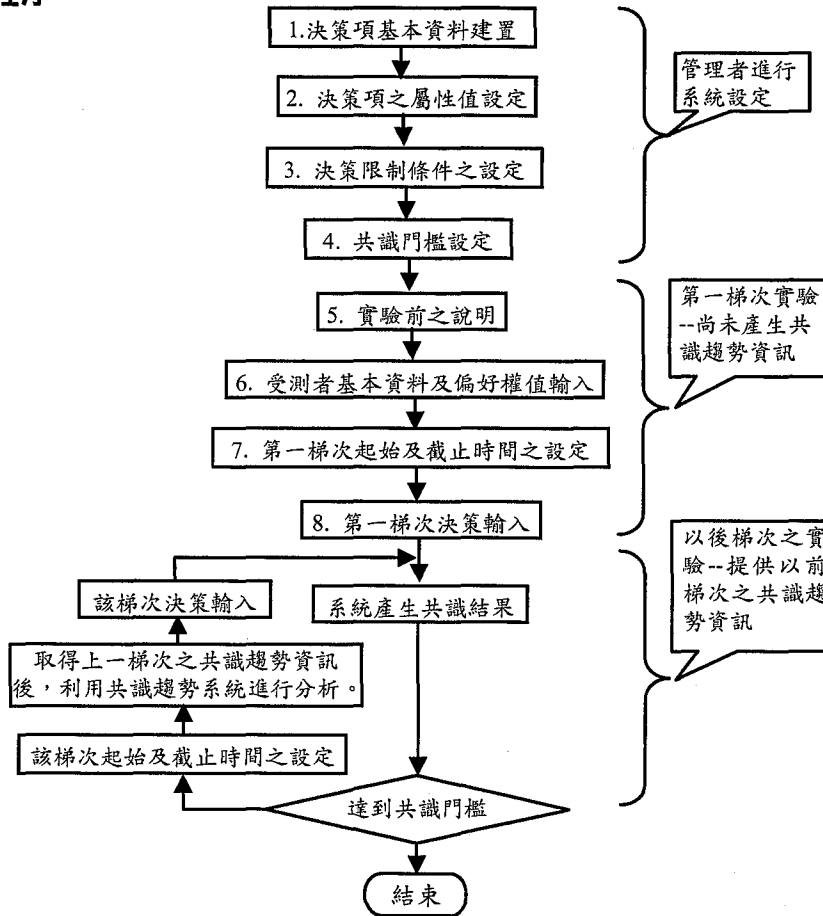


圖 9：實驗程序

圖九之實驗程序，說明如下：

- (一) 決策項基本資料檔建置：兩決策項包含旅遊地點與旅遊時間。研究者挑選 33 個旅遊定點，旅遊時間為 1 到 12 月。
- (二) 決策項之屬性值設定：輸入各旅遊地點相關屬性值，如旅遊地點之型態、區域、住宿等級等
- (三) 決策限制條件之設定：針對各決策項基本資料進行欄位限制 (constraints)。如總費用限制、特定時間之限制，如排除農曆年時間，或限定必須為週末等非上課期間。
- (四) 共識門檻設定：如平均決策滿意度必須達到 70% 才算完成決策，否則繼續下一梯次決策。
- (五) 實驗前之說明：實驗目的、程序及注意事項等之解說，為時約 10 分鐘。
- (六) 受測者基本資料及偏好權值輸入，為時約 10 分鐘。

(七) 各梯次起始及截止時間之設定：為了方便不同班級學生在不同之上課時段進行同一梯次實驗，故每一梯次以一週為期限，受測者於每一梯次之實際決策時間為時 15 分鐘。

(八) 各梯次之決策輸入：每一位決策者輸入個人所期望之旅遊地點及時間。

四、各梯次之偏好分布圖

一般線上資料分析處理所產生之圖表主要為資料方塊圖、樞紐分析表及彙總表，然而，對於決策者而言，藉著這些圖表工具仍不易清楚了解每一梯次之「主流意見群」分布概況，因此本研究採用 DENCLUE 密度叢集演算法，將每一梯次之決策結果轉換成「偏好叢集圖」來呈現。本研究以「住宿等級」及「月份」兩屬性為座標軸屬性，其原因有二：(1) 由表 4 可知「住宿等級」及「月份」分別為地點及時間決策項之最重要屬性，因此所產生之偏差率低，(2)「住宿等級」及「月份」兩屬性均為次序型態之屬性。圖 10 至圖 14 為本研究於各梯次實驗所產生之偏好叢集圖。

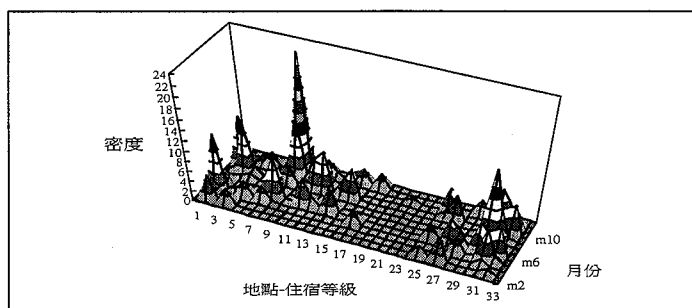


圖 10：梯次一之偏好密度叢集圖

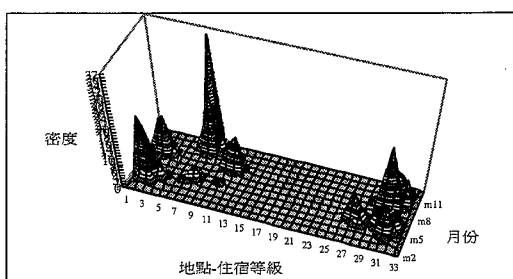


圖 11：梯次二之叢集圖

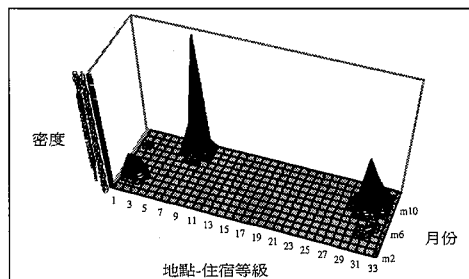


圖 12：梯次三之叢集圖

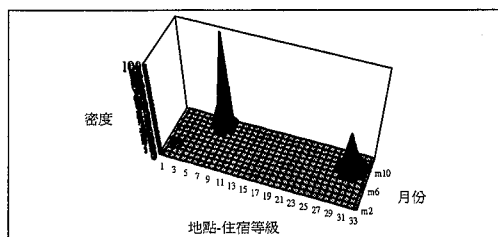


圖 13：梯次四之叢集圖

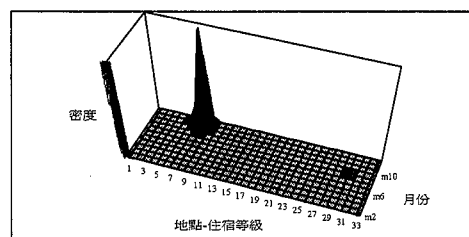


圖 14：梯次五之叢集圖

各梯次之偏好叢集圖之分析如表 3，其中，「最大叢集之引力中心」即為「最獲偏好方案」所在，如第一、二梯次之最獲偏好方案為 (m9, 30)，亦即旅遊時間為九月，旅遊之住宿等級為 30；第三~五梯次之最獲偏好方案為 (m12, 8)。「決策質心」即為全體之平均偏好。由表 3 可知實驗隨者梯次增加發生幾種效應，(1) 被選擇方案數及叢集數遞減效應，(2) 最大叢集之群聚量遞增效應，(3) 最大叢集之平均密度遞增效應，(4) 最大叢集及決策質心發生移轉效應。

表 3：各梯次之最大叢集變動分析表

梯次別	方案數	叢集數	最大叢集之 引力中心	最大叢集 之群聚量	最大叢集 之密度	決策質心
梯次一	68	25	(m9,30)	38	4.75	(m8.38,15.98)
梯次二	35	9	(m9,30)	62	6.89	(m8.58,15.79)
梯次三	21	5	(m12,8)	117	39	(m9.98,16.48)
梯次四	11	3	(m12,8)	149	49.67	(m10.69,16.59)
梯次五	5	3	(m12,8)	245	81.67	(m11.61,8.14)

伍、資料分析與假說檢定

一、基本資料敘述統計

由於本實驗在課程中進行，因此資料之回收率高且空值率低，受測者平均年齡為 20.2（標準差為 0.9）。經五個梯次實驗後，資料回收率 0.95，有效資料佔 99%。受測者輸入各屬性權值後，經正規化處理後之平均權值如表 4。

表 4：各屬性之平均權值

屬性 vs. 權值	住宿等級	月份	型態	季節	區域	縣市	週
平均權值	0.58	0.50	0.36	0.33	0.28	0.27	0.11
平方值(%)	0.34	0.25	0.13	0.11	0.08	0.07	0.01

由表 4 可知「住宿等級」、「月份」及「型態」三屬性之受重視程度佔 72%（因各權值平方和為 1），因此，此三屬性可視為決定性屬性。當進入決策後，受測者於各梯次分別輸入所偏好之決策方案，系統便可得出該梯次之最獲偏好方案，將決策者所選方案之屬性值與最獲偏好方案之屬性值比對，便可計算得個別屬性之偏好差距，表 5 可用以觀察決策者於不同梯次決策後，各決定性屬性之平均偏好差距的變化，由表 5 可得知，「住宿等級」、「月份」兩屬性之偏好差距於一、二梯次並無明顯變異，而於第三梯次以後產生較明顯之變異；而「型態」屬性之偏好差距則有逐梯次降低之趨勢。

表 5：決策者相對於各屬性之平均偏好差距及標準差

	一梯次後	兩梯次後	三梯次後	四梯次後	五梯次後
住宿等級	0.48 (0.35)	0.49 (0.38)	0.33 (0.35)	0.29 (0.36)	0.01 (0.08)
月份	0.38 (0.28)	0.37 (0.25)	0.30 (0.31)	0.21 (0.22)	0.06 (0.09)
型態	0.57 (0.49)	0.49 (0.50)	0.24 (0.42)	0.11 (0.31)	0.06 (0.24)

依據決策者所選方案之屬性及其屬性權值，將其代入前述公式(4)(8)(7)可得整體之期望落差、一致性程度及決策滿意度等應變數值，各應變數於各梯次決策後之平均值與標準差如表 6。

表 6：各應變數於各梯次決策後之平均值與標準差

操控方式	一梯次後	兩梯次後	三梯次後	四梯次後	五梯次後
期望落差	0.53 (0.22)	0.52 (0.26)	0.36 (0.29)	0.28 (0.29)	0.06 (0.11)
一致性程度	0.53 (0.15)	0.52 (0.14)	0.60 (0.15)	0.66 (0.12)	0.94 (0.13)
決策滿意度	0.14 (0.27)	0.20 (0.34)	0.43 (0.47)	0.55 (0.47)	0.91 (0.22)

二、檢定流程及檢定工具

本研究之檢定分成兩階段，第一階段主要進行多變量變異數分析 (MANOVA)，總檢定在檢驗實驗變數對應變數群是否產生顯著影響，邊際檢定在於檢驗實驗變數對於個別應變數是否產生顯著影響，成偶檢定將實驗變數分群，成對檢定在於驗證不同之梯次群對應變數之變率影響是否顯著。此階段所使用之檢定工具為 SAS 之 1-way MANOVA 檢定法。

第二梯次之檢定，可稱之為因果關係檢定。第一階段檢定只能得知實驗變數對應變數之影響是否顯著，無法得知各變數間之因果關係，為了得知所有變數之整體影響關係，如直接及間接影響之共變量，必須進行第二階段之因果關係檢定。此階段本研究採用 SAS 之「CALIS 程序」作為因果分析工具。因果分析是探討一群變數間之因果關係模式的方法 (Blalock 1971; James 1982; Lewis-Beck 1995)，CALIS 與 LISREL 兩者都是運用 SEM (Structural Equation Model) 技術來分析因果關係之一種工具。雖然過去許多研究常使用 OLS 回歸法 (Ordinary Least Squared regression)，然而 LISREL (Linear Structural Relation) 已經被證實更具有強健可靠的因果分析能力 (Hatcher 1998)，因此，本研究採用 SAS 中具有 LISREL 功能之 CALIS 分析工具，另外，由於本研究之因果模型出現「互逆性」的因果關係，故必須仰賴具有互逆性之路徑分析能力之工具，此亦為本研究選擇 CALIS 之重要原因。

三、第一階段檢定——多變量變異數分析

(一) 檢定

為檢定實驗變數與應變數間之顯著關係，本研究執行 1-way MANOVA 後之結果，

發現 Wilk's Lambda 之 F 統計量為 105.32，相對應的 p 值為 0.0001 小於顯著水準 0.05，因此可宣稱模式顯著，棄卻虛無假說，換言之，即表示可以從梯次別進一步進行邊際檢定。

(二) 邊際檢定

執行 1-way MANOVA 對期望落差、滿意度、一致性程度之邊際檢定結果如表 7、8、9，發現 F 統計量分別為 157.28、172.03、376.59，相對應的 p 值小於顯著水準 0.05，表示梯次別對期望落差、滿意度、一致性程度有顯著影響，因此棄卻虛無假說，可以進一步進行後續之成偶檢定及對比檢定。

表 7：期望落差之邊際檢定結果

來源	自由度	平方和	均方	F 統計量	P 值
模 式	4	37.8477258	9.4619314	157.28	<.0001
誤 差	1234	74.2393741	0.0601616		
總 和	1238	112.0870999			

表 8：滿意度之邊際檢定結果

來源	自由度	平方和	均方	F 統計量	P 值
模 式	4	94.0459093	23.5114773	172.03	<.0001
誤 差	1234	168.6474506	0.1366673		
總 和	1238	262.6933599			

表 9：對一致性程度之邊際檢定結果

來源	自由度	平方和	均方	F 統計量	P 值
模 式	4	29.03910297	7.25977574	376.59	<.0001
誤 差	1234	23.78877661	0.01927778		
總 和	1238	52.82787958			

(三) 成偶檢定及成對檢定

當進一步將不同梯次間所產生之決策滿意度進行成偶檢定（周文賢 2002），依據 SAS 之執行結果，所產生之差異檢查表如表 10，由表可知，梯次一、二之決策滿意度之變率小，可視為低決策滿意度之梯次群；而梯次三及梯次四之決策滿意度逐漸上升且變率逐漸增大，可視為中度決策滿意度之梯次群；而梯次五可視為高度決策滿意度之梯次群。當進一步將此三群組進行對比檢定。結果顯示 P 值皆小於 0.0001，小於 α ，顯示此三梯次群具有差異性。類似的三梯次群之變異情形亦發生在期望落差及一致性程度之成偶及成對檢定，限於篇幅不另贅述。因此，本實驗顯示，當決策者應用多梯次共識趨勢分析系統時，第 3 梯次以後才產生顯著提高決策滿意度之效益。

表 10：滿意度——各梯次間之差異檢查表

梯次別	滿意度	1	2	3	4	5
1	0.13772856		0.0573	<.0001	<.0001	<.0001
2	0.20071254	0.0573		<.0001	<.0001	<.0001
3	0.42849361	<.0001	<.0001		0.0004	<.0001
4	0.54648640	<.0001	<.0001	0.0004		<.0001
5	0.91000866	<.0001	<.0001	<.0001	<.0001	

四、第二階段檢定——因果關係檢定

前述圖 7 之因果模型經過 SAS 之 CALIS 分析後所產生之結果如圖 15。

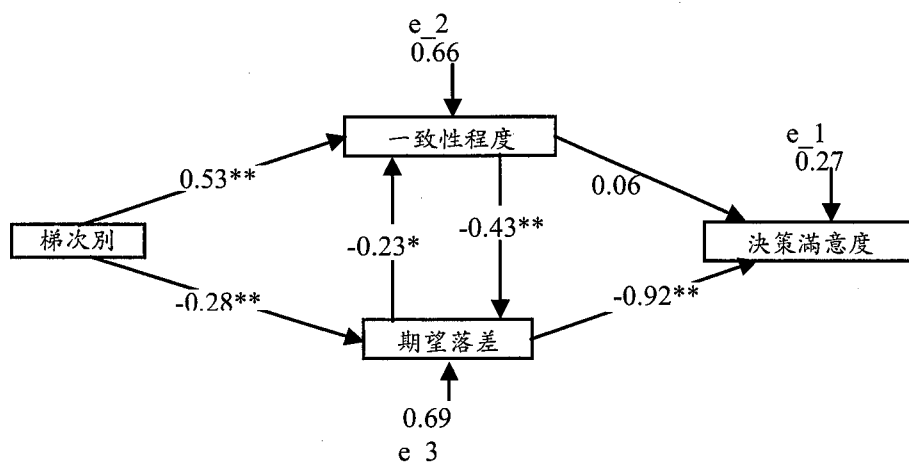


圖 15：經分析後之因果模型圖

Hatcher (1998) 提出良好之因果模型必須通過以下五項檢驗，而本研究之因果模型除了滿足收斂條件外，也都能通過這些檢驗，因此，可視為檢驗優良之模型，說明如下：

檢驗一：檢驗常態化殘差矩陣 (normalized residual matrix)，所有矩陣內之元素，其殘差的絕對值必須小於 2。圖 15 之因果模型，其最大殘差值發生在「梯次別」與「滿意度」之間，其值為 0.79，絕對值小於 2。

檢驗二：卡方檢定 (Chi-square test)，此為因果模型是否適應之總檢定，Chi-square 值愈趨近於 0 愈好，p 值愈趨近於 1 愈好，而本研究模型之 p 值趨近於 1，因此，可視為良好適應。

檢驗三：NNFI (Non-Normed Fit Index) 及 CFI (Comparative Fit Index) 檢定，Bentler & Bonett's (1980) NNFI 及 Bentler's CFI 皆須超過 0.9。本研究模型之 CFI 及 NNFI 皆為 0.99。

檢驗四：內因變數 (endogenous variables) 之可解釋變量 (R-squared value) 愈大愈好。在本研究模型中，期望落差之可解釋變量為 53%，一致性程度之可解釋變量為 56%，滿意度之可解釋變量為 91%，因此，本模型之可解釋性極高。

檢驗五：各標準化路徑係數 (standardized path coefficients) 之顯著性程度標示，一般將顯著程度依次分為四種層次，(1) $t > 3.3, p < 0.001$ ，(2) $t > 2.58, p < 0.01$ ，(3) $t > 1.96, p < 0.05$ ，(4) 不顯著，此四種層次之顯著程度分別以「***」、「**」、「*」及「」標示。本研究模型之標準化路徑係數及顯著性程度標示如圖 15。

五、影響效果及假說檢定

各變數間之總體影響效果包含直接效果及間接效果，圖 15 之路徑係數為直接影響效果，採用 Fox (1981) 之計算法，產生之總體效果及間接效果如下：

1. 梯次別對一致性程度之直接影響效果為 0.53，總體影響 ϵ_{21} 為 0.66，因此，間接影響效果為 0.13。

$$\text{令 } r = (-0.23) \times (-0.43)$$

$$\epsilon_{21} = 0.53(1 + r + r^2 + \dots) + 0.28 \times 0.23(1 + r + r^2 + \dots) = \frac{0.53}{1-r} + \frac{0.28 \times 0.23}{1-r} = 0.66$$

2. 梯次別對期望落差之直接影響效果為 -0.28，總體影響 ϵ_{31} 為 -0.56，因此，間接影響效果為 -0.28。

$$\epsilon_{31} = (-0.28) \times (1 + r + r^2 + \dots) + 0.53 \times (-0.43) \times (1 + r + r^2 + \dots) = -0.56$$

3. 梯次別對決策滿意度之直接影響效果不顯著，總體影響 ϵ_{41} 為 0.55，主要來自於間接影響效果。

$$\epsilon_{41} = 0.66 \times 0.06 + (-0.56) \times (-0.92) = 0.55$$

4. 一致性程度對期望落差之直接影響效果為 -0.43，總體影響 ϵ_{32} 為 -0.48，因此，間接影響效果為 -0.05。

$$\epsilon_{32} = -0.43 \times (1 + r + r^2 + \dots) = \frac{-0.43}{1-r} = -0.48$$

5. 期望落差對一致性程度之直接影響效果為 -0.23，總體影響 ϵ_{23} 為 -0.26，因此，間接影響效果為 -0.03。

$$\epsilon_{23} = -0.23 \times (1 + r + r^2 + \dots) = \frac{-0.23}{1-r} = -0.26$$

6. 一致性程度對決策滿意度之直接影響效果為 0.06，總體影響 ϵ_{42} 為 0.49，因此，間接影響效果為 0.43。

$$\epsilon_{42} = 0.06 + (-0.43) \times (-0.92) \times (1 + r + r^2 + \dots) = 0.06 + \frac{(-0.43) \times (-0.92)}{1-r} = 0.49$$

7. 期望落差對決策滿意度之直接影響效果為 -0.92，總體影響 ϵ_{43} 為 -0.94，因此，間接影響效果為 -0.02。

$$\epsilon_{43} = -0.92 + (-0.23 \times 0.06) \times (1 + r + r^2 + \dots) = -0.94$$

假說一之檢定：

梯次別對期望落差之「總體影響」為-0.56，其中，直接影響為-0.28， p 值小於0.001，影響顯著，因此，假說一成立，也就是說，決策者之期望落差將隨者梯次別增加而逐步降低，兩者之間具有顯著而直接之負相關。

假說二之檢定：

期望落差對決策滿意度之直接影響高達-0.92，因此，兩者間存在顯著而直接之負相關，假說二成立。

假說三之檢定：

雖然梯次別對決策滿意度之「總體影響」高達0.55，然而，主要影響路徑是「梯次別-期望落差-決策滿意度」，也就是說，其主要來自於間接影響，經因果關係分析後，決策滿意度與梯次別之間的路徑係數極小（或不顯著），證實決策滿意度並未直接受到梯次別之影響，因此，拒絕假說三。

假說四之檢定：

梯次別對一致性程度之「總體影響」為0.66，其中，直接影響為0.53， p 值小於0.001，兩者之間具有顯著而直接之正相關，因此，假說三成立。

假說五之檢定：

一致性程度對決策滿意度之直接影響為0.06， p 值大於0.05，影響不顯著，因此，拒絕假說五，也就是說，提昇決策者之一致性程度並不能直接促使決策滿意度提昇。

假說六之檢定：

期望落差對一致性程度之直接影響效果為-0.23，一致性程度對期望落差之直接影響效果為-0.43，顯示兩者存在顯著而直接之互逆性影響，因此，假說六成立，也就是說，提昇決策者之一致性程度，將促使決策者之期望落差降低；反之，降低決策者之期望落差，將促使決策者之一致性程度提昇。

陸、結論與建議

本研究應用叢集技術所建構之「多梯次共識趨勢分析系統」，主要應用於「線上公民決策」，系統目的是協助決策者提昇共識程度，鑒於公民決策不易凝聚共識，因此，本系統建構時從兩方面加以改善共識凝聚問題，(1)一般線上分析處理系統主要提供樞紐分析表、方塊資料圖及彙總表給使用者進行分析，然而，當使用者在共識凝聚過程中欲了解各主流群體之分布與趨勢動向時，一般線上分析處理系統便無法滿足此需求，因此，本研究利用「密度叢集法」建構「共識趨勢分析系統」以滿足此需求。(2)因共識凝聚難以在單一梯次決策完成，故本系統以多梯次決策來逐步提昇共識程度。

為探討「多梯次共識趨勢分析系統」對於共識程度指標之改善是否產生顯著影響，本研究以實驗法進行研究，結果顯示：

1. 決策一致性程度隨者梯次增加而逐漸升高，兩者產生直接而顯著之正相關。
2. 決策者之期望落差隨者梯次增加而降低，兩者產生直接而顯著之負相關。
3. 決策滿意度隨者梯次增加而逐漸增加，主要來自於間接影響，其影響路徑為「梯次別-期望落差-決策滿意度」。

因此，多梯次共識趨勢分析系統對於共識程度指標之改善，經研究證實具有顯著效益。

就共識程度指標相互間之影響而言，研究亦證實決策滿意度與期望落差產生顯著而直接之正相關，而決策滿意度與決策一致性程度之間並未呈現直接相關性，也就是說，以「一致性程度為基礎」及「以期望落差為基礎」所計算得之共識程度有相當差異，而後者具有與決策滿意度產生直接正相關之特性，而前者則否。因此，「以期望落差為基礎」之共識程度計算法較能符合線上公民決策之需求。

就決策梯次間之差異比較而言，經成偶檢定可知，初期之第一、二梯次間無顯著差異，第三梯次以後，各梯次間產生顯著差異，最大差異出現在第四梯次與第五梯次之間。由於在共識凝聚過程中，潛藏著不同主流群體間之對立、衝突及拉鋸現象，以致造成欲達成共識所需之決策梯次數增加，因此，若系統能協助決策者找出衝突點並計算衝突程度，將有助於衝突之化解，如此決策者便能以較少之梯次數取得高共識、高滿意度之決議，此為本研究可進一步研究之方向。

參考文獻

1. 行政院研考會，民 87，八十七年度政府業務電腦化報告書—邁向二十一世紀電子化政府，台北：行政院研考會。
2. 葉俊榮，民 86，邁向電子化政府：資訊公開與行政程序的挑戰與回應，台北，網路與法律研討會。
3. 項靖，民 88，地方政府網路公共論壇與民主行政之實踐，東海社會科學學報，第十八期。
4. 周文賢，民 91，多變量統計分析-SAS/STAT 之應用，智勝文化。
5. Alavi, M. and Joachimsthaler, E. A. "Revisiting DSS Implementation Research: A Meta-Analysis of the Literature and Suggestions for Researchers," MIS Quarterly, Vol.16, 1992, pp.95-116
6. Ankerst, M., Breunig, M., Kriegel, H.-P. and Sander, J. "OPTICS: Ordering Points to Identify the Clustering Structure," In Proc. 1999 ACM-SIGMOD Int. Conf. Management of Data (SIGMOD'99), Philadelphia, PA, June 1999, pp.49-60
7. Asch, S. E. *Effects of Experimental Social Psychology*, Academic Press, New York, 1980
8. Asch, S.E. *Group Pressure on the Modification and Distortion of Judgment*, Groups,

- Leadership and Men, Camegie Press, Pittsburgh, 1951, pp.177-190
9. Baker, S.M. and Petty, R.E. "Majority and Minority Influence: Source-Position in Balance as a Determinant of Message Scrutiny," *Journal of Personality and Social Psychology*, 67, 1994, pp.5-19.
10. Blalock, H. M. *Causal Models in the Social Science*, Chicago: Aldine, 1971
11. Blalock, H. M., *Causal Inference in Non-Experimental Research*, New York: W. W: Norton, 1971
12. Calabrese, A. and Borchert, M, "Prospects for electronic democracy in the United States: Rethinking communication and social policy", *Media, Culture, and Society*, 18, 1996, pp.249-268
13. Ester, M., Kriegel, H.-P., Sander, J. and Xu, X. "A Density-Based Algorithm for Discovering Clusters in Large Spatial Database," In Proc. 1996 Int. Conf. Knowledge Discovery and Data Mining (KDD'96), Portland, Aug., 1996, pp:226-231
14. Fox, J. "Effect Analysis in Structural Equation Models: Extensions and Simplified Methods of Computation," pp.178-203 in *Linear Models in Social Research*, SAGE Publications, Beverly Hill, CA, 1981.
15. Frijda, N. H., Kuipers, P. and Schure, E. "Relations among Emotion, Appraisal and Emotional Action Readiness," *Journal of Personality and Social Psychology*, Vol.57, 1989, pp.212-228
16. Hatcher, H. "A Step-by-Step Approach to Using the SAS System for Factor Analysis and Structural Equation Modeling," SAS Institute Inc., NC, USA, 1998
17. Heise, D. R. *Causal Analysis*, New York: Wiley-Interscience, 1975
18. Herrera, F., Herrera V.E. and Verdegay, J.L. "A Linguistic Decision Process in Group Decision Making", *Group Decision and Negotiation*, 5, 1996, pp.165-176
19. Herrera, F., Herrera-Viedma E. and Verdegay, J.L. "A Model of Consensus in Group Decision Making under Linguistic Assessments", *Fuzzy Sets and Systems*, 78, 1996, pp.73-88
20. Hinneburg, A. and Keim, D.A. "An Efficient Approach to Clustering in Large Multimedia Databases with Noise," In Proc. 1998 Int. Conf. Knowledge Discovery and Data Mining (KDD'98), New York, Aug. 1998, pp.58-65
21. James, D. "The Relationship between Participation and Satisfaction," *MIS Quarterly*, Vol.18(4), 1994, pp.427-445
22. James, L.R., Mulaik, S.A., and Brett, L.M. *Causal Analysis: Assumptions, Models, and Data*, Beverly Hills, CA: Stage, 1982
23. Kacprzyk, J. and Roubens, M. *Non-Conventional Preference Relations in Decision Making*, Springer-Verlag, Heidelberg, 1988
24. Kacprzyk, J., Nurmi, H. and Fedrizzi, M. *Consensus under Fuzziness*, Kluwer Academic Publishers, Boston, 1997.

25. Kaufman, L. and Rousseeuw, P.J. *Finding Groups in Data: An Introduction to Cluster Analysis*, John Wiley & Sons, New York, 1990
26. Landman, J. "Regret: A Theoretical and Conceptual Analysis," *Journal for the Theory of Social Behavior*, Vol.17, 1987, pp.135-160
27. Lewis-Beck, M. S. *Data Analysis: an Introduction*, Thousand Oaks: Sage Publications, 1995
28. Macpherson, "Citizen politics and the renewal of democracy," On-line paper, URL: <http://www.snafu.de/~mim/CP/cp2.html>, 1997
29. MacQueen, J. "Some Methods for Classification and Analysis of Multivariate Observations," *Proc. 5th Berkeley Symp., Math. Stat., Prob.*, 1967, pp.281-297
30. Moscovici, S. & Personnaz, B. "Studies in Social Influence V. Minority Influence and Conversion Behavior in a Perceptual Task," *Journal of Experimental Social Psychology*, 16, 1980, pp.270-282
31. Nambodiri, N. K., Carter, L. F. and Blalock, H. M. *Applied Multivariate Analysis and Experimental Design*, New York, McGraw-Hill, 1975
32. Robey, D., Farrow, D., and Franz, C. R. "Group Process and Conflict in System Development," *Management Science*, Vol.35(10), 1989, pp.1172-1191
33. Tears, R.K. "Expectations, Performance Evaluation and Consumer's Perception of Quality," *Journal of Marketing*, 57(4), 1993, pp.18-34
34. Tsiros, M. "Effect of Regret on Post-choice Valuation: The Case of More Than Two Alternatives," *Organizational Behavior and Human Decision Processes*. Vol.76, 1998, pp.48-69.
35. Voss, G.B., Parasuraman, A. and Grewal, D. "The Roles of Price, Performance, and Expectations in Determining Satisfaction in Service Exchanges," *Journal of Marketing*, New York, Oct 1998
36. Yager, R.R. "On Ordered Weighted Averaging Aggregation Operators in Multi-criteria Decision-Making", *IEEE Transactions on Systems, SMC-18*, 1988, pp.183-190
37. Yi, Y. "A Critical Review of Consumer Satisfaction," In V.A. Zeithaml (Ed.) *Review of Marketing*, Chicago, IL: American Marketing Association, 1990.
38. Zadeh, L.A. "A Computational Approach to Fuzzy Quantifiers in Natural Languages," *Computers and Mathematics with Applications*, 9, 1983, pp.149-184
39. Zadrozny, S. *Computer Based Consensus Reaching Support Employing the Elements of Fuzzy Logic*. Ph.D. Thesis, Polish Academy of Sciences, Warsaw, Poland, 1993