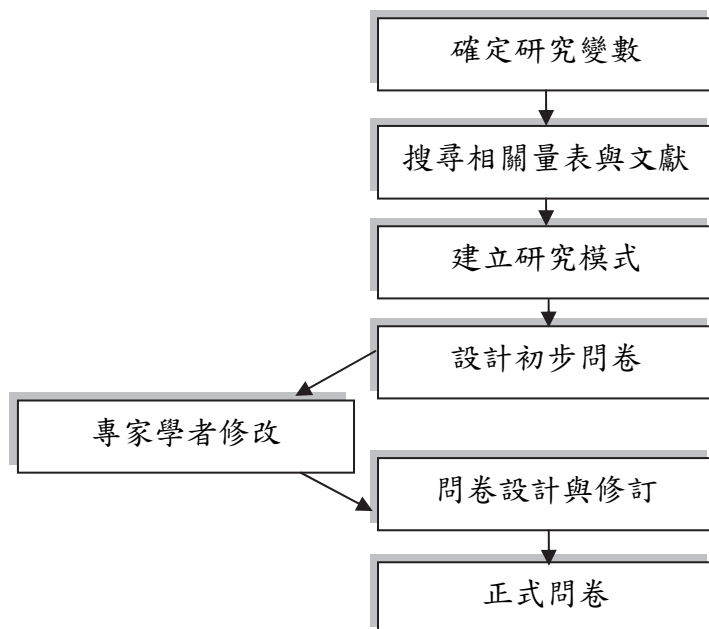


結構方程模式 實例

05 CHAPTER

這是一個結構方程模式實例，研究目的在探討高階主管支持、團隊合作、ERP 系統品質與使用者滿意度之關聯性，以及各因素間的因果關係。經由相關文獻之理論探討，建立初步的研究模式，再經由實證分析與研究，而獲得了本研究最後整體架構關係路徑圖之結果。在使用 Structural Equation Modeling(結構方程模式)時，需要有測量的工具——量表，量表對於社會科學研究中從事量化研究的人員而言，是相當重要的一環，少了量表，我們就無法作到量化的效果，從事社會科學研究的人員常常會遇到在進行問卷調查設計時，發現自行發展量表並不是一件容易的事，必須經過嚴謹的處理，才能發展出一份適當的、穩定的量表，因此，我們可以借用發展成熟的量表，來進行量測，我們整理問卷發展的步驟如下：



問卷發展的步驟：

問卷結構分為六個部分，以李克特五尺度法(Likert Scale)將程度分為五尺度衡量。在資料編碼上，依照答卷者勾選的程度強弱由 1~5 進行編碼，例如極不同意則編碼為 1，而非常同意則編碼為 5。綜合上述，本研究問項問卷結構以及操作化參考來源如下表所示：

表一：本研究問卷結構以及參考來源

構面	問項參考來源	問卷對應選項
高階主管支持	McDonald & Eastlack (1971)	ERP 專案團隊的運作-A 部分
團隊合作	Lee & Choi (2003)	ERP 專案團隊的運作-B 部分
系統品質	Wixom & Waston(2001)	大型 ERP 系統的開發/使用-A 部分
資訊品質	Rai et al. (2002)	大型 ERP 系統的開發/使用-B 部分
服務品質	Pitt et al.(1995)	大型 ERP 系統的開發/使用-C 部分
使用者滿意度	Bailey & Pearson (1983)	大型 ERP 系統的開發/使用-D 部分

我們發展的問卷內容如下表：

表二：問卷內容

【ERP 專案團隊的運作】					
	非常不同意	有些不同意	普通	比較同意	非常同意
A. 他們（她們）在參與專案時，您覺得：					
1. 對 ERP 系統開發給予明確的規範	☐	☐	☐	☐	☐
2. 參與 ERP 系統開發與建置團隊人選的指派	☐	☐	☐	☐	☐
3. 制定新 ERP 系統做與不做的標準	☐	☐	☐	☐	☐
B. 團隊合作方面，我們專案小組的成員					
4. 對於合作的程度是滿意	☐	☐	☐	☐	☐
5. 對專案是支持的	☐	☐	☐	☐	☐
6. 對跨部門的合作是很有意願	☐	☐	☐	☐	☐

【大型 ERP 系統的開發／使用】					
	非常不同意	有些不同意	普通	比較同意	非常同意
C. 對於系統的品質，您覺得					
7. ERP 系統可以有效地整合來自不同部門系統的資料	☐	☐	☐	☐	☐
8. ERP 系統的資料在很多方面是適用的	☐	☐	☐	☐	☐
9. ERP 系統可以有效地整合組織內各種型態的資料	☐	☐	☐	☐	☐
D. 對於資訊的品質，您覺得					
10. 提供精確的資訊	☐	☐	☐	☐	☐
11. 提供作業上足夠的資訊	☐	☐	☐	☐	☐
E. 對於資訊部門的服務，您覺得					
12. 會在所承諾的時間內提供服務	☐	☐	☐	☐	☐
13. 堅持作到零缺點服務	☐	☐	☐	☐	☐
14. 員工將會給使用者即時的服務	☐	☐	☐	☐	☐
15. 總是願意協助使用者	☐	☐	☐	☐	☐
16. 整體而言，資訊部門的服務良好	☐	☐	☐	☐	☐
F. 就使用者滿意而言，您覺得					
17. 滿意 ERP 系統輸出資訊內容的完整性	☐	☐	☐	☐	☐
18. ERP 系統是容易使用	☐	☐	☐	☐	☐
19. ERP 系統的檔是有用的	☐	☐	☐	☐	☐

本研究問卷共發出 957 份，回收 372 份，扣除填答不全與胡亂填答之無效問卷 22 份，有效問卷 350 份，有效回收率為 36.57 %。

一般來說，SEM 的分析至少須包含兩大階段(Anderson & Gerbing, 1988)。第一階段稱為量測模型(Measurement model)，第二階段稱為結構模型(Structural model)。

第一階段中的量測模型，也可稱為驗證性因素分析(Confirmatory factor analysis, CFA)，通常包含了信度(Reliability)檢測、建構效度(Construct validity)檢測與模型適配度(Model fit)檢測。首先，信度檢測是想瞭解研究者針對某一群固定的受測者，利用同一種特定的測量工具，在重複進行多次測量後，所得到的結果是否具有穩定性(Stability)與一致性(Consistency)的結果；常用來檢測信度的指標為組成信度(Composite reliability)。第二，量測模型中的建構效度檢驗通常至少需包含收斂效度(Convergent validity，也可翻譯成聚合效度)與區別效度(Discriminant validity，或稱為鑑別效度、區辨效度)。

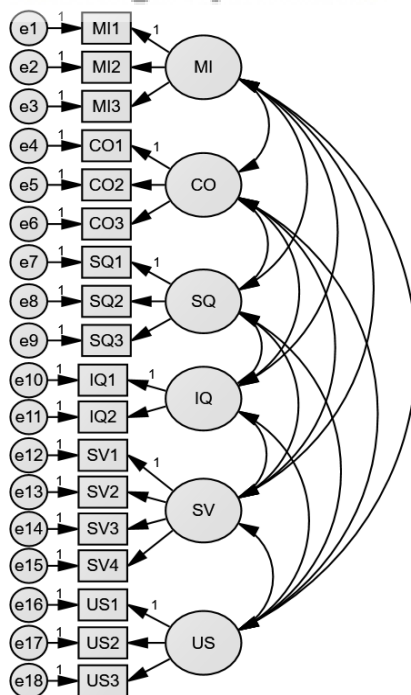
第二階段中的結構模型，通常包含模型式適配度檢測、路徑分析(Path analysis)與其他進階分析。

5-1 量測模式

本節提供兩種不同研究模型量測模式。5-1a 第一種將因素負荷量(factor loading)預設為 1 的量測模型，5-1b 第二種將構面的共變數預設為 1 的量測模型。原因是第一種將因素負荷量(factor loading)預設為 1 的量測模型之題項(指標變數)，未提供因素負荷量(factor loading)預設為 1 之題項(指標變數)得標準化值與顯著性估計，論文審查者可能會要求提供此數值。然而，這兩者的量測模型是等值模型，視需求擇一進行量測模型評估即可。

□ 5-1a 將因素負荷量(factor loading)預設為 1 的量測模型

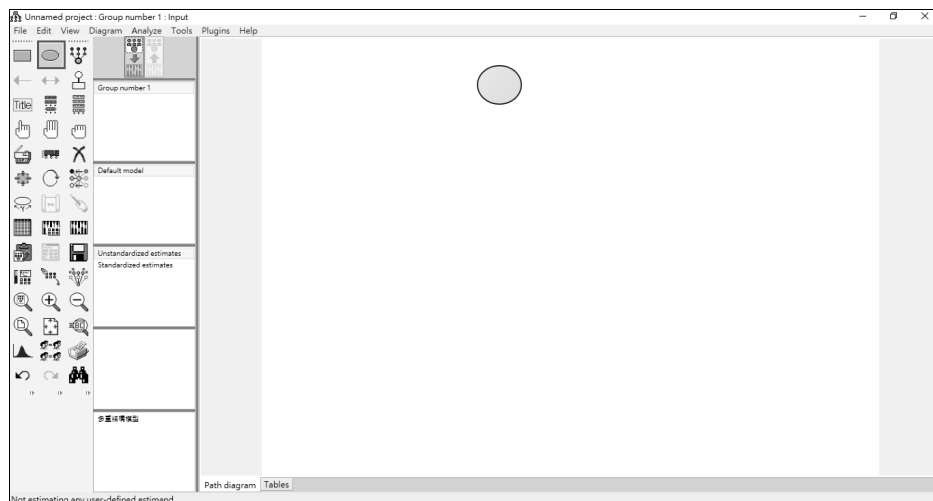
本研究模型量測模式，不包含 SV5 (16. 整體而言，資訊部門的服務良好)，範例如下：



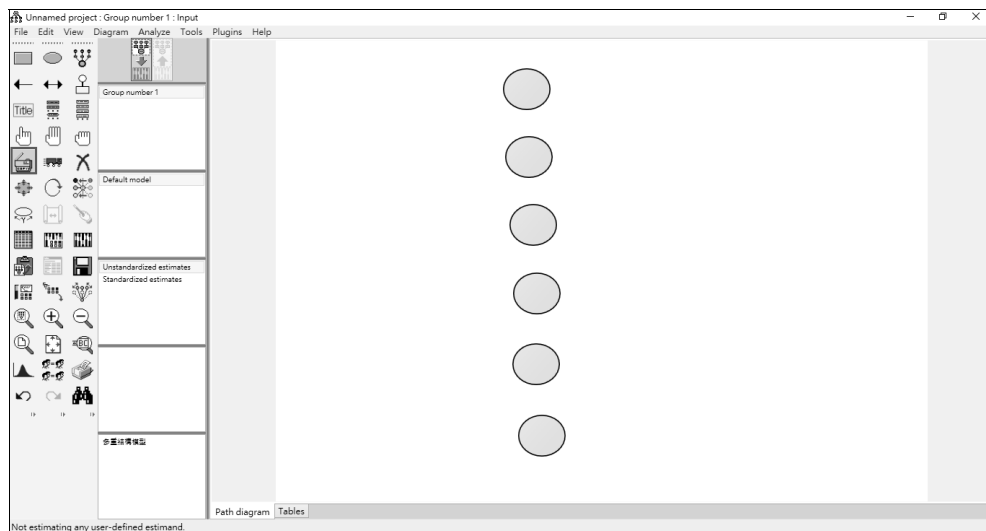
本章的量測模式操作如下(相關檔案請參閱 Ch 05 資料夾中的 5-1 資料夾內容)，
注意：不包含 SV5 (16. 整體而言，資訊部門的服務良好)：

Step 01 繪製六個潛在變數(構面)。

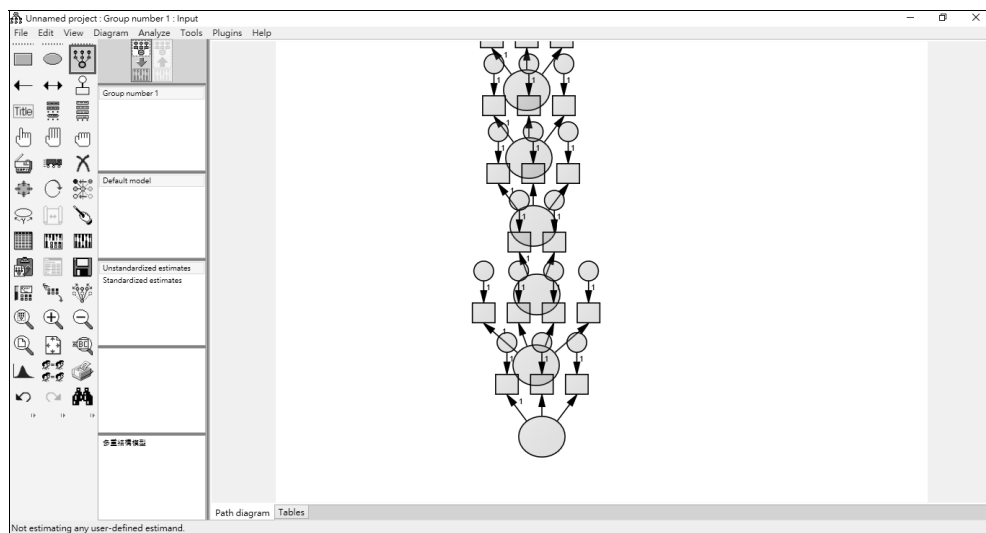
備註：一個研究需有幾個構面，由文獻探討與相關的研究理論基礎而定。



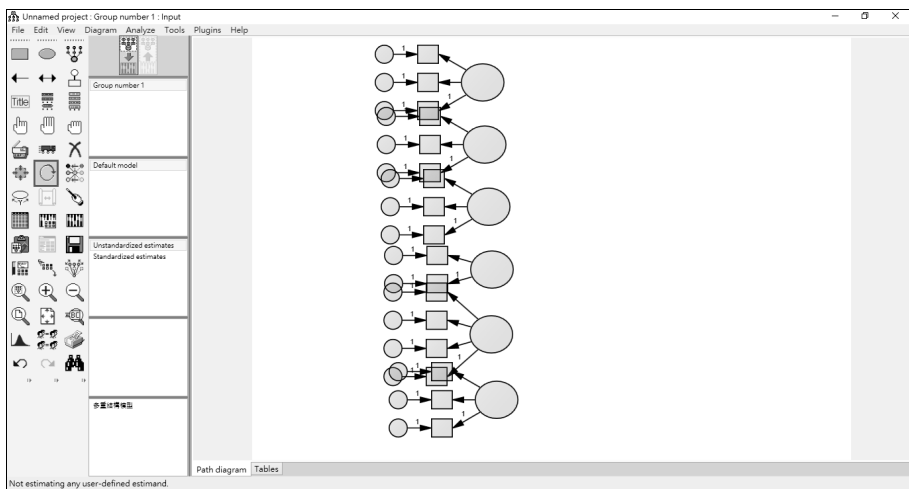
Step 02 用複製功能(Duplicate objects)確保每個構面的大小一致。



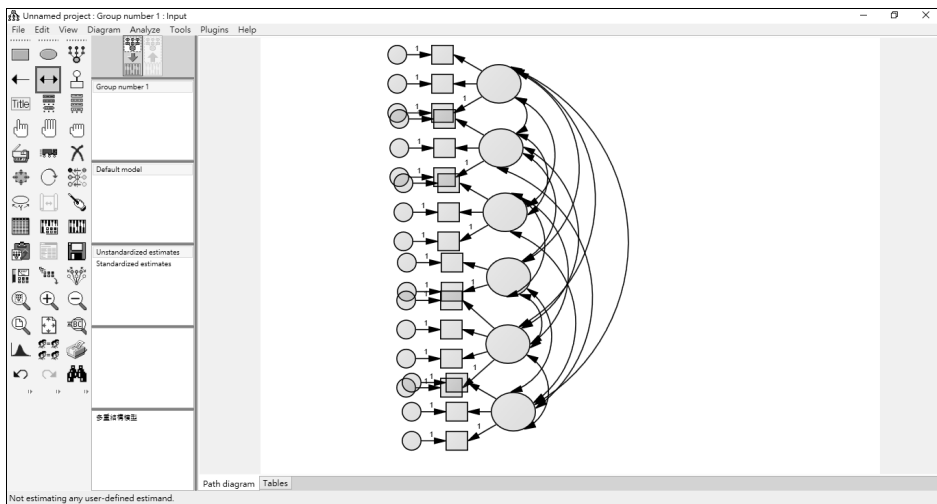
Step 03 在潛在變數上插入指標變數。按下增加指標變數的按鈕後，點選要增加指標變數的構面，即可增加指標變數於構面中。每個構面需要有幾個指標變數，取決於研究模型的設定。



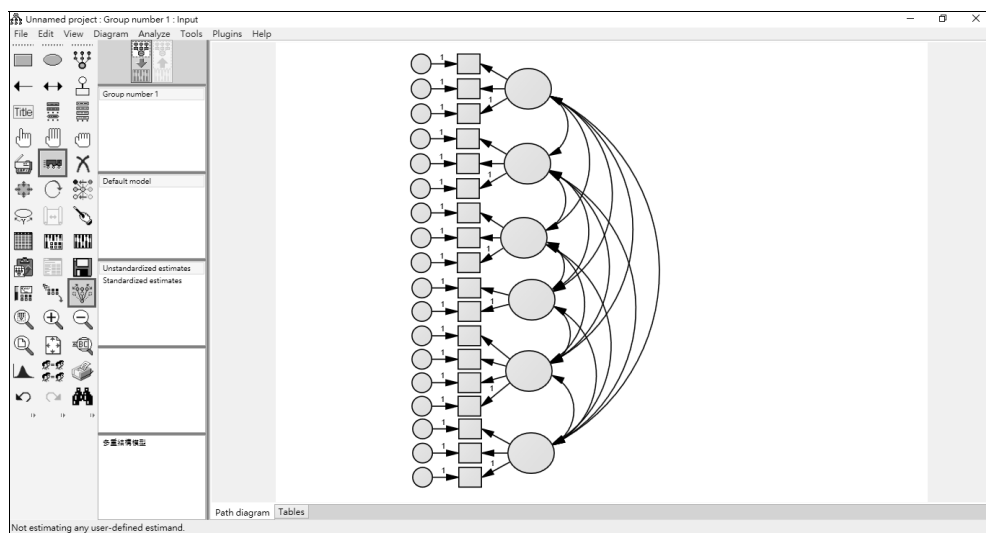
Step 04 使用旋轉潛在變數指標(Rotate the indicators of a latent variable)能把指標變數旋轉至到合適的位置。也可利用 Move objects 將構面移動至合適的位置。建議每個構面與指標變數不要重疊，以便分析後能直接從模型中直接觀察相關數據與分析結果。若需要潛在變數與指標變數一起移動，請先按 Preserve symmetries，再按 Move objects。



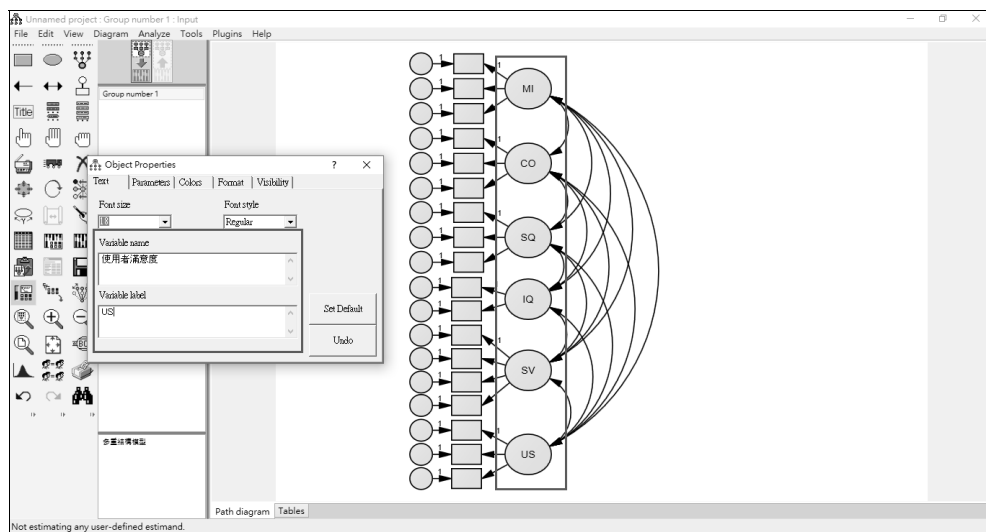
Step 05 利用共變數(雙箭頭符號)設定潛在變數之間的相關。本範例的測量模型中，共有 6 個構面，須設定 15 個相關。備註：若量測模型有兩個構面，則需有 1 個相關需要設定；若量測模型有三個構面，則需有 3 個相關需要設定；若量測模型有五個構面，則需有 10 個相關需要設定。



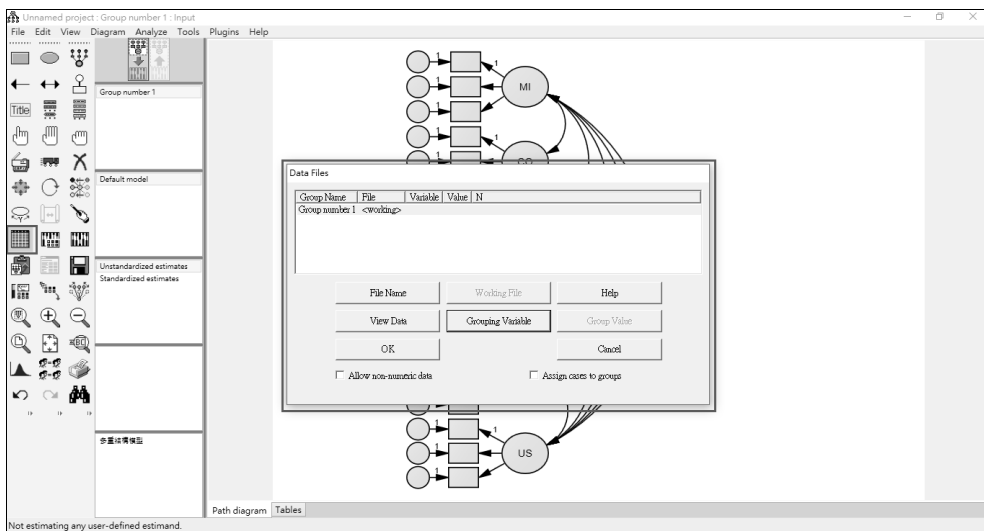
Step 06 利用 Move objects 將構面移動至合適的位置。建議每個構面與指標變數不要重疊，以便分析後能直接從模型中直接觀察相關數據與分析結果。



Step 07 在潛在變數點擊 2 下，在物件屬性(Object Properties)中 Variable name 輸入名稱。Variable name 是指分析報表(View text)中顯示變數名稱；而 Variable label 則是在模型(圖形)上顯示的名稱。此部分不影響分析或操作，使用者至少需輸入 Variable name。

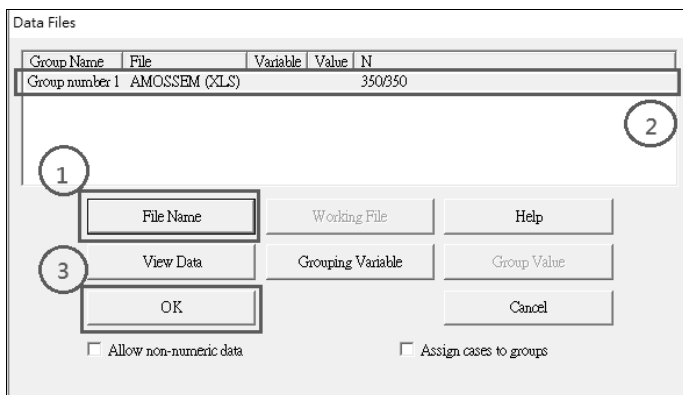


Step 08 選擇樣本來源或資料檔。點選 **File Name**，可指定資料檔在電腦或硬碟中的路徑，將要分析的資料檔匯入。按 **View Data** 則是打開指定的資料檔。

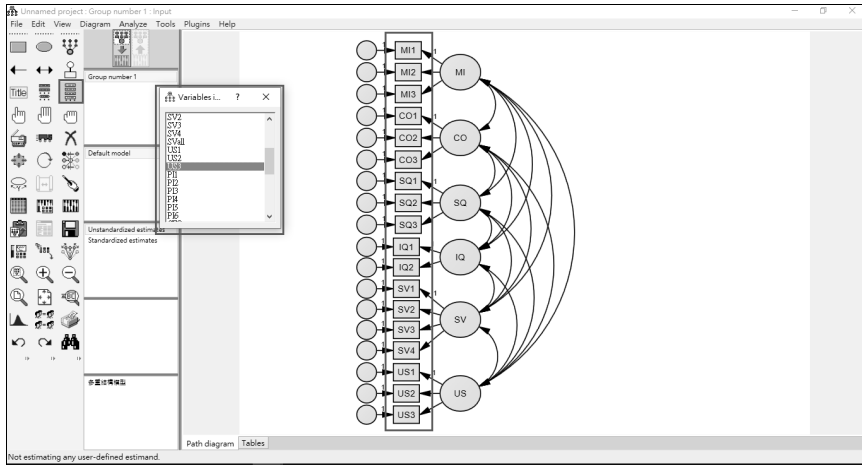


Step 09 點選 **File Name** 把檔案選進，**Data Files** 中就能看到選進的檔案跟樣本數，最後點選 **OK**。

(如果是 Excel 檔案，必須是「Excel 97-2003 活頁簿」的檔案)



Step 10 根據研究模型的設定，將列出變數名稱名稱拖曳至各所屬構面中。

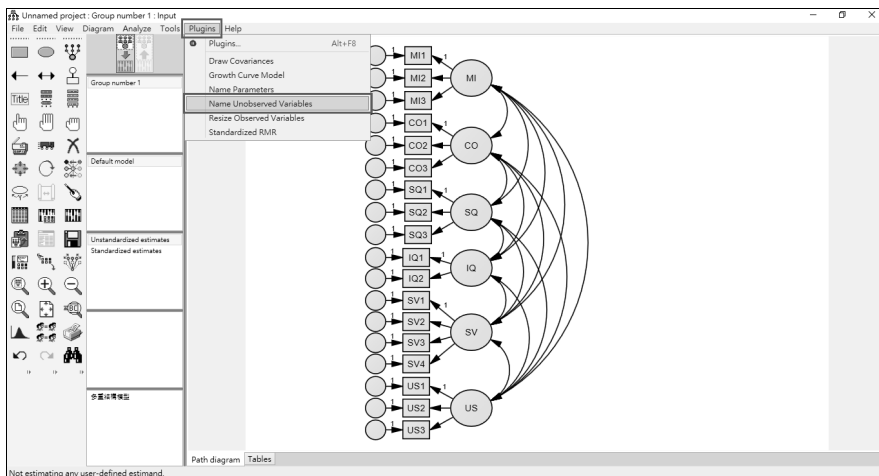


Step 11 在殘差上加入變數名稱。務必注意，在 SEM 中，每個變數都必須有名稱，不可空白，否則無法分析。

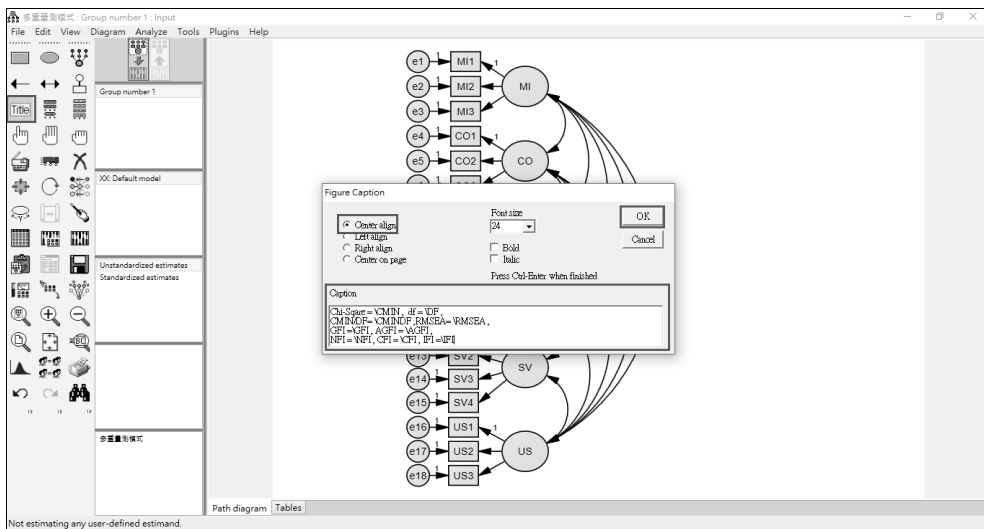
有兩種方式：

- (1) Plugins→Name Unobserved Variables [自動插入殘差]。
- (2) 點選殘差，在 Variables name 上輸入殘差[手動輸入殘差]。

建議使用自動插入殘差的功能即可。例如，使用者只需要知悉 MI1 的殘差變數為 e1、MI2 的殘差變數為 e2、MI3 的殘差變數為 e3 即可(其餘以此類推)。如果潛在變數與觀察變數很多，建議採用此功能，即無須再為每個殘差變數命名。



Step 12 點選 Title，可以將常用的適配度指標顯示在圖面上。此設定可讓使用者無須進入 View text (分析結果頁面)即可知悉量測模型的適配度(選取 Center align 可以隨時移動位置)。使用者可加入自己常觀察的適配度指標至本圖形介面上。



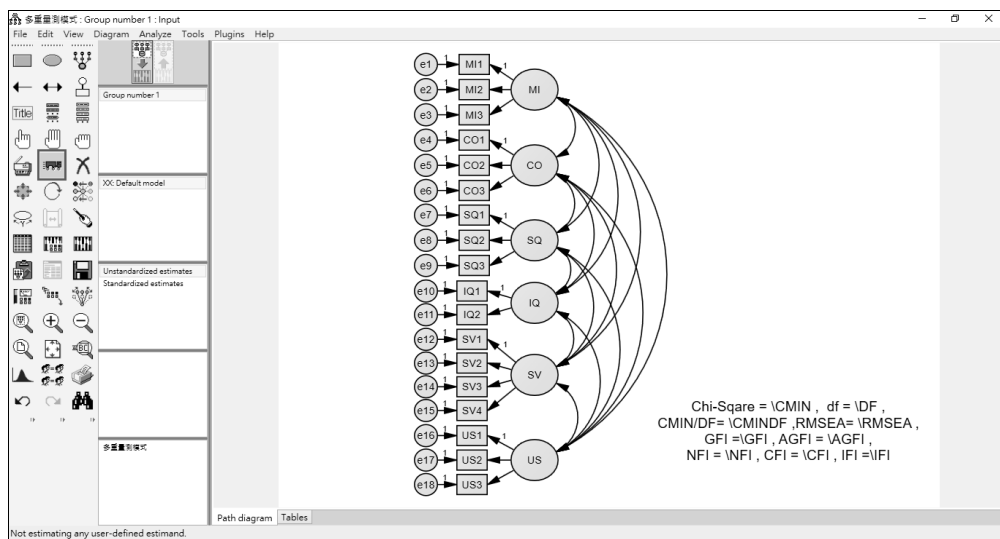
配適度統計量的代碼如下：

配適度指標	代碼	配適度指標	代碼	配適度指標	代碼
AGFI	= \AGFI	F0 HI90	= \F0HI	P 值	= \P
AIC	= \AIC	F0 LO90	= \F0LO	PARTIO	= \PARTIO
BCC	= \BCC	FMIN	= \FMIN	PCFI	= \PCFI
BIC	= \BIC	GFI	= \GFI	PCLOSE	= \PCLOSE
CAIC	= \CAIC	HOELTER (A = 0.01)	= \HONE	PGFI	= \PGFI
CFI	= \CFI	HOELTER (A = 0.05)	= \HFIVE	PNFI	= \PNFI
CMIN	= \CMIN	IFI	= \IFI	RFI	= \RFI
CMIN/DF	= \CMINDF	MECVI	= \MECVI	RMR	= \RMR
DF	= \DF	NCP	= \NCP	RMSEA	= \RMSEA
ECVI	= \ECVI	NCP HI90	= \NCPHI	RMSEA HI90	= \RMSEA HI
ECVI HI 90	= \ECVIHI	NCP LO90	= \NCPLO	RMSEA LO90	= \RMSEA LO
ECVI LO 90	= \ECVILO	NFI	= \NFI	TLI	= \TLI
F0	= \F0	NPAR	= \NPAR		

整理配適度指標值與理想數值如下：

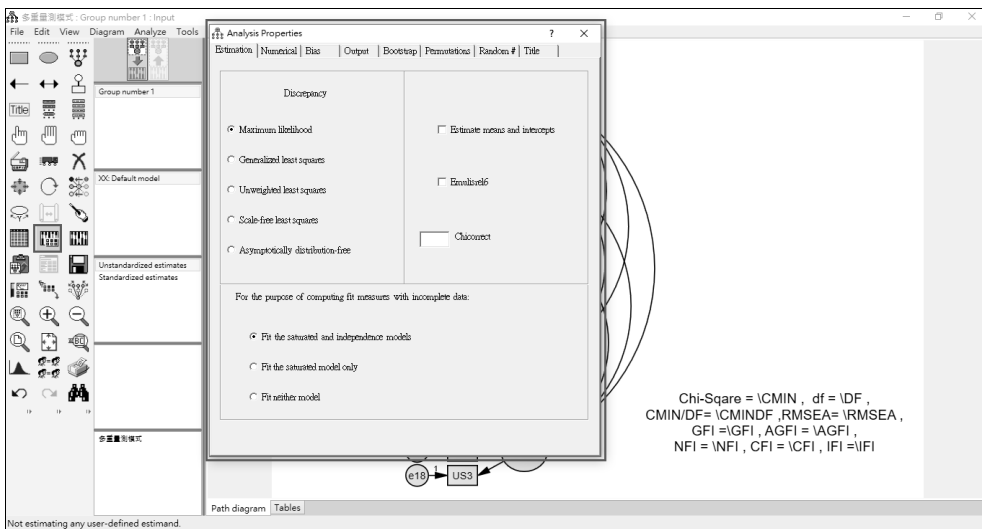
配適度指標	理想數值	配適度指標	理想數值	配適度指標	理想數值
CMIN/DF	<3	IFI	>0.9	NFI	>0.9
GHI	>0.9	CFI	>0.9		
AGFI	>0.8	RMSEA	<0.08		

Step 13 點選移動，可以移動標題。

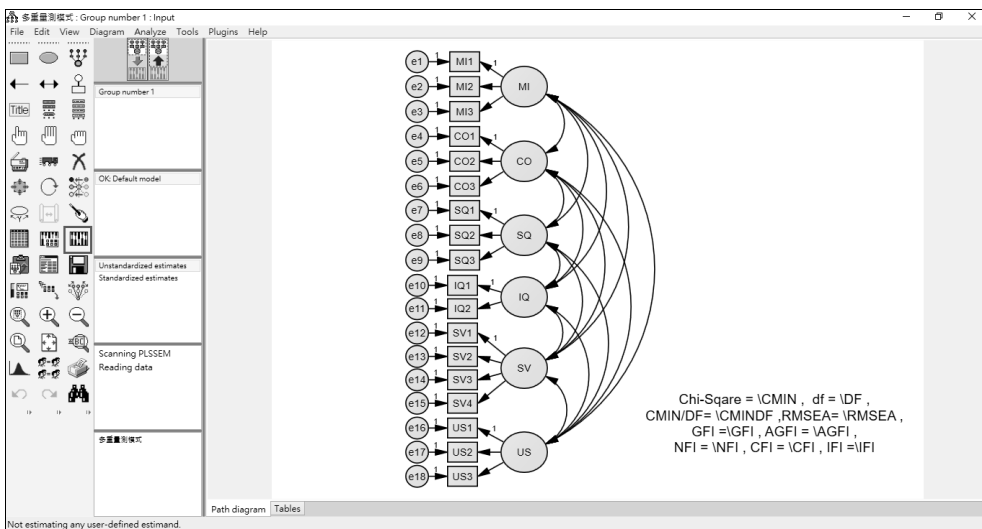


Step 14 選擇分析屬性設定。

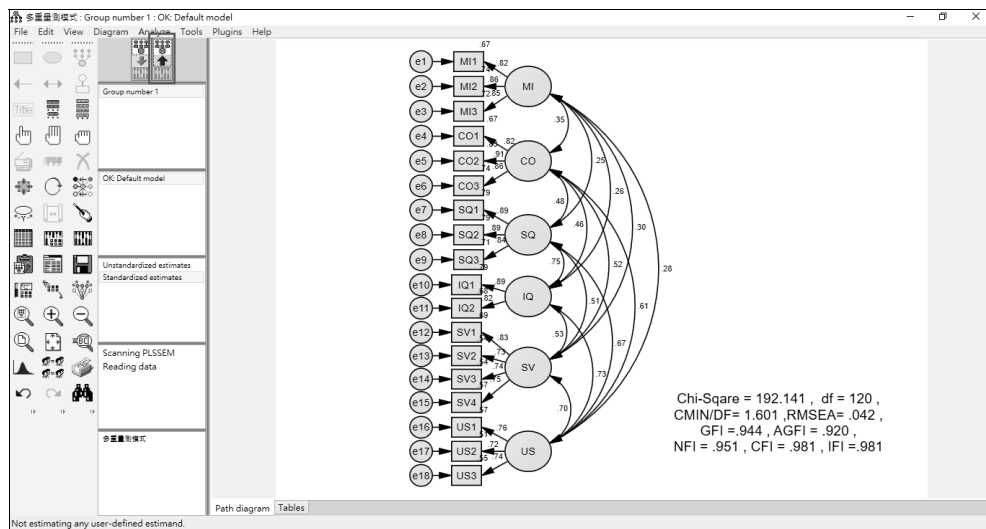
(勾選：Minimization history(模型收斂的統計量)、Standardized estimates(標準化估計值)、Squared multiple(多元相關平方)、All implied moments(含潛在變數的模型矩陣)和 Modification indices(模型修正))。



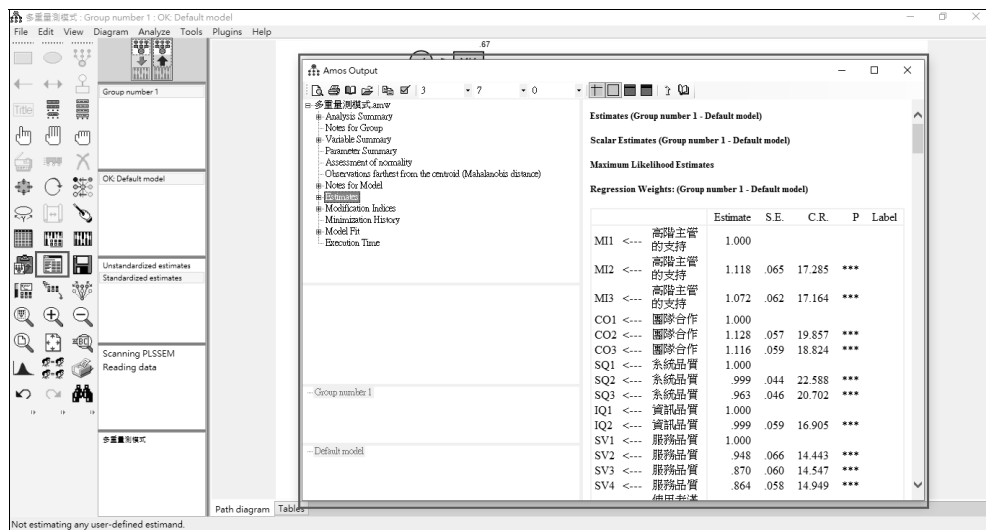
Step 15 按 Calculate(執行)即可分析量測模型之結果。建議將分析的模型儲存後分析再行分析。



Step 16 按紅色箭頭，可以看到結果。



Step 17 點選 View Text，看到報表結果。



報表結果：

◆ Model Fit

整理如下：

	基本模型		基本模型
卡方(CMIN)	117.271	NFI	0.962
自由度(DF)	67	GFI	0.956
標準卡方檢定(CMIN/DF)	1.750	RMR	0.026
CFI	0.983	RMSEA	0.046

卡方(CMIN)：卡方差異值越小越好

自由度(DF)：自由度(愈小表示愈精簡)

標準卡方檢定(CMIN/DF)：在 1~3 之間為理想的模型適配度

◆ Regression Weights: (Group number 1 - Default model)

	Estimate	S.E.	C.R.	P	Label
MI1 <--- 主管階級的支持	1.000				
MI2 <--- 主管階級的支持	1.113	.065	17.151	***	par_1
MI3 <--- 主管階級的支持	1.068	.063	17.013	***	par_2
CO1 <--- 團隊合作	1.000				
CO2 <--- 團隊合作	1.140	.059	19.341	***	par_3
CO3 <--- 團隊合作	1.124	.061	18.565	***	par_4
SQ1 <--- 系統品質	1.000				
SQ2 <--- 系統品質	.998	.045	22.231	***	par_5
SQ3 <--- 系統品質	.963	.046	20.758	***	par_6
IQ1 <--- 資訊品質	1.000				
IQ2 <--- 資訊品質	1.001	.061	16.436	***	par_7
US1 <--- 使用者滿意度	1.000				
US2 <--- 使用者滿意度	1.048	.093	11.279	***	par_8
US3 <--- 使用者滿意度	1.168	.096	12.172	***	par_9

非標準化因素負荷量，可以看到點估計值、標準差 S.E.和 C.R.以 P 值是否顯著，顯著表示假設成立。

◆ Standardized Regression Weights: (Group number 1 - Default model)

			Estimate
MI1	<---	主管階級的支持	.819
MI2	<---	主管階級的支持	.856
MI3	<---	主管階級的支持	.848
CO1	<---	團隊合作	.813
CO2	<---	團隊合作	.914
CO3	<---	團隊合作	.862
SQ1	<---	系統品質	.878
SQ2	<---	系統品質	.877
SQ3	<---	系統品質	.828
IQ1	<---	資訊品質	.882
IQ2	<---	資訊品質	.816
US1	<---	使用者滿意度	.767
US2	<---	使用者滿意度	.685
US3	<---	使用者滿意度	.759

標準化因素負荷量，介於 0.5~0.95 之間。

(>0.7 是理想；>0.6 為可接受。)

◆ Variances: (Group number 1 - Default model)

	Estimate	S.E.	C.R.	P	Label
主管階級的支持	.646	.074	8.778	***	par_18
團隊合作	.463	.052	8.972	***	par_19
系統品質	.502	.048	10.424	***	par_20
資訊品質	.398	.043	9.277	***	par_21
使用者滿意度	.319	.042	7.559	***	par_22
e1	.317	.035	9.132	***	par_23
e2	.292	.037	7.919	***	par_24
e3	.288	.035	8.256	***	par_25
e4	.238	.023	10.317	***	par_26
e5	.119	.020	5.859	***	par_27
e6	.203	.023	8.631	***	par_28
e7	.149	.018	8.228	***	par_29

	Estimate	S.E.	C.R.	P	Label
e8	.151	.018	8.289	***	par_30
e9	.213	.021	10.147	***	par_31
e10	.114	.021	5.518	***	par_32
e11	.199	.024	8.291	***	par_33
e12	.223	.026	8.652	***	par_34
e13	.397	.038	10.440	***	par_35
e14	.321	.036	8.877	***	par_36

變異數與殘差應為正值並且顯著。

- ◆ Squared Multiple Correlations: (Group number 1 - Default model) 。在 SEM 中， $SMC = R^2$ 。

	Estimate
US3	.576
US2	.469
US1	.588
IQ2	.667
IQ1	.778
SQ3	.686
SQ2	.768
SQ1	.771
CO3	.743
CO2	.835
CO1	.661
MI3	.719
MI2	.733
MI1	.671

量測模型：

$$R^2 > 0.5$$

結構模型：

$$R^2 = 0.19 \text{ small}$$

$$R^2 = 0.33 \text{ medium}$$

$$R^2 = 0.67 \text{ large}$$

◆ 信效度分析(Word 報表呈現)

構面	指標	模型參數估計值				收斂效度			
		非標準化因素負荷量	S. E.	C. R.	P	標準化因素負荷量	SMC	CR	AVE
高階主管的支持	MI1	1.000				0.817	0.667	0.879	0.708
	MI2	1.118	0.065	17.285	***	0.858	0.736		
	MI3	1.072	0.062	17.164	***	0.849	0.721		
團隊合作	CO1	1.000				0.819	0.671	0.899	0.748
	CO2	1.128	0.057	19.857	***	0.911	0.830		
	CO3	1.116	0.059	18.824	***	0.862	0.744		
系統品質	SQ1	1.000				0.888	0.789	0.906	0.762
	SQ2	0.999	0.044	22.588	***	0.887	0.788		
	SQ3	0.963	0.046	20.702	***	0.842	0.709		
資訊品質	IQ1	1.000				0.889	0.790	0.847	0.735
	IQ2	0.999	0.059	16.905	***	0.824	0.679		
服務品質	SV1	1.000				0.833	0.694	0.850	0.586
	SV2	0.948	0.066	14.443	***	0.733	0.538		
	SV3	0.870	0.060	14.547	***	0.738	0.544		
	SV4	0.864	0.058	14.949	***	0.755	0.570		
使用者滿意度	US1	1.000				0.757	0.574	0.782	0.545
	US2	1.108	0.088	12.532	***	0.715	0.511		
	US3	1.156	0.089	12.980	***	0.742	0.550		

- ◆ 組成信度的公式，可參閱第三章。研究中每個潛在變數(構面)的組成信度(C.R.) 需 >0.7 ，代表具有研究的構面有良好的一致性。
本範例的研究構面(高階主管的支持)C.R.為 0.879，具有良好的一致性；
本範例的研究構面(團隊合作)C.R.為 0.899，具有良好的一致性；
本範例的研究構面(系統品質)C.R.為 0.906，具有良好的一致性；
本範例的研究構面(資訊品質)C.R.為 0.847，具有良好的一致性；
本範例的研究構面(服務品質)C.R.為 0.850，具有良好的一致性；
本範例的研究構面(使用者滿意度)C.R.為 0.782，具有良好的一致性。
- ◆ 除了組成信度的檢驗之外，研究中的每個構面的平均變異數萃取量(AVE)需高於 0.5，且每個指標變數的標準化因素負荷須高於 0.6，代表具有收斂效度。
- ◆ 根據 Hair et al. (2009)與 Fornell & Larcker (1981)的建議，量測模型要達到收斂效度，須滿足以下三個條件：
- ◆ 每個直接觀察變數的標準化因素負荷量需高於 0.7，可接受的標準為 0.6 以上。

- ◆ 組成信度(CR)須高於 0.7 以上。
- ◆ 平均萃取變異量(AVE)須高於 0.5 以上。
 本範例的研究構面(高階主管的支持)AVE 為 0.708，具有收斂效度；
 本範例的研究構面(團隊合作)AVE 為 0.748，具有收斂效度；
 本範例的研究構面(系統品質)AVE 為 0.762，具有收斂效度；
 本範例的研究構面(資訊品質)AVE 為 0.735，具有收斂效度；
 本範例的研究構面(服務品質)AVE 為 0.586，具有收斂效度；
 本範例的研究構面(使用者滿意度)AVE 為 0.545，具有收斂效度。

附帶一提，同一構面內的量測指標(直接觀察變數)應為正相關(同向)。若有反向題，應重新編碼後再納入分析(如果有需要的話)。

此外，因本章節所提出的模型無須修正。未來若讀者的量測模型無法通過上述的信度或收斂效度的檢驗，可考慮進行以下的模型修正：

1. 刪除標準化因素負荷量低於 0.5 的指標變數。
2. 刪除可能有共線性(collinearity)疑慮的直接觀察變數(例如直接觀察變數間相關係數高於 0.85)。
3. 刪除 Modification Index (修正指標，簡稱 MI)較高的觀察變數。

◆ 區別效度(Word 呈現)

	AVE	使用者滿意度	服務品質	資訊品質	系統品質	團隊合作	高階主管的支持
使用者滿意度	0.545	0.738					
服務品質	0.586	0.695	0.766				
資訊品質	0.735	0.731	0.528	0.857			
系統品質	0.762	0.670	0.506	0.751	0.873		
團隊合作	0.748	0.607	0.523	0.458	0.479	0.865	
高階主管的支持	0.708	0.283	0.305	0.256	0.245	0.348	0.841