

$$\text{نسبة التباين المفسر للمكون الأول} = \frac{\text{الجذر الكامن}}{\text{مجموع الجذور الكامنة}} * 100 = \frac{2.9}{5} = 58\% .$$

وأن المكون الثاني يفسر 22.2% من التباينات ويفسر المكونان نسبة 80.2% من هيكل التباينات للمتغيرات الخمسة ، وقد أهمل البرنامج بقية المكونات نظراً لكون جذورها الكامنة تقل عن الواحد .

Total Variance Explained

Component	Initial Eigenvalues			Extraction Sums of Squared Loadings		
	Total	% of Variance	Cumulative %	Total	% of Variance	Cumulative %
1	2.900	58.009	58.009	2.900	58.009	58.009
2	1.110	22.207	80.216	1.110	22.207	80.216
3	.523	10.457	90.673			
4	.393	7.864	98.537			
5	7.313E-02	1.463	100.000			

Extraction Method: Principal Component Analysis.

المخرج التالي يمثل مصفوفة المكونات Components Matrix التي تتضمن تشبعات Loadings المكونين الأول والثاني الذين تم استخلاصهما أن التشبع هو عبارة عن معامل الارتباط البسيط بين المكون (أو العامل) والمتغير . أن أقوى المتغيرات ارتباطاً بالعامل الأول هو متغير التعليم العالي HIGHEDEC حيث أن تشبع المتغير بالمكون الأساسي الأول هو 0.941 يليه متغير التعليم LITERACY ثم متغير عدد الأطباء DOCTORS وأن أضعف المتغيرات ارتباطاً بالعامل الأول هما متغيري GDP ومتغير عدد أسرة المستشفيات HOSPBED. أما أقوى المتغيرات ارتباطاً بالمكون الثاني فهو GDP ثم متغير HOSPBEDS ولكن باتجاه معاكس، ويرتبط المكون الثاني بعلاقة ضعيفة ببقية المتغيرات .

Component Matrix ^a

	Component	
	1	2
GDP	.477	.754
LITIRACY	.918	1.899E-02
HIGHEDEC	.941	.102
DOCTORS	.829	-.131
HOSPBED	.508	-.717

Extraction Method: Principal Component Analysis.

a. 2 components extracted.

ملاحظات :

1. أن اشتراكية المتغير هي مجموع مربعات تشبعات المتغير بالعوامل المستخلصة فاشتراكية المتغير GDP تساوي $0.477^2 + 0.754^2 = 0.797$ وهكذا لباقي المتغيرات .
2. أن مجموع مربعات تشبعات المتغيرات بالمكون (العامل) يساوي الجذر الكامن للمكون فمثلاً نحصل على الجذر الكامن للمكون (العامل) الأول كما يلي :

$$0.477^2 + 0.918^2 + 0.941^2 + 0.829^2 + 0.508^2 = 2.9$$
 لاحظ أنه قد يحصل اختلاف بسيط في النتائج بسبب عمليات التقريب .

المخرج التالي يمثل معاملات المكونات (العوامل) Component Scores Coefficients وتحسب هذه المعاملات من مصفوفة المكونات السابقة Martix Components فمثلاً المعامل GDP للمكون الأول يحسب كما يلي :

$$\frac{0.477}{0.477^2 + 0.918^2 + 0.941^2 + 0.829^2 + 0.508^2} = 0.165$$

Component Score Coefficient Matrix

	Component	
	1	2
GDP	.165	.679
LITIRACY	.316	.017
HIGHEDUC	.324	.092
DOCTORS	.286	-.118
HOSPBED	.175	-.645

Extraction Method: Principal Component Analysis.

Component Scores.

أما قيم المكونات (العوامل) فتحسب بموجب الدالة الخطية التالية للمكون الأول مثلاً :

fac_1 =

0.165*GDP+0.316*LITIRACY+0.324*HIGHEDUC+0.286*DOCTORS+0.175*HOSPBED

علماً أنه يتم استعمال المتغيرات المعيارية Standardized variables في التركيب الخطي (أي

القيم المعيارية لمتغيرات GDP ، Literacy ، ...) وكما ذكرنا تضاف العوامل الى يمين المتغيرات

الموجودة في شاشة Data Editor وكما يلي :

Region	gdp	literacy	higheduc	doctors	hospbed	fac_1	fac_2
Dohok	17.2	30.1	1.09	17	139	-.84873	.00862
Nineveh	24.0	44.2	1.85	14	172	.05265	-.01091
Arbil	22.2	35.2	1.18	13	163	-.65458	-.04095
Sulayman	16.2	33.5	1.01	10	115	-1.10939	.37657
Ta'meem	32.3	49.4	1.85	18	143	.37097	.54501
Salah AL-Deen	98.4	37.9	1.32	15	80	-.11362	3.32500
Diala	23.1	44.1	1.93	10	153	-.14014	.26391
Anbar	22.7	44.3	1.58	16	144	-.08303	.22314
Baghdad	75.0	61.6	4.04	36	280	3.22748	.21850
Wasit	19.5	36.7	1.11	28	199	.04721	-.80457
Babylon	22.8	44.1	1.82	18	145	.09229	.21066
Kerbala	21.5	47.7	1.53	24	173	.41839	-.29133
Najaf	18.7	46.2	1.59	27	190	.53708	-.62771
Qadisia	21.0	35.2	.95	9	195	-.81013	-.42906
Muthana	21.3	33.5	.84	18	178	-.62869	-.37423
Thi-Qar	18.1	33.8	.73	12	144	-1.03052	.02027
Maysan	20.4	34.4	.90	11	301	-.44139	-1.76886
Basrah	19.0	53.6	2.24	25	219	1.11415	-.84406

علماً أن المكونات الأساسية هي متغيرات وهمية (نظرية) ليس لها أي تفسير محدد ولكن لها

استعمالات خاصة مثلاً معالجة مشكلة التعدد الخطي في نماذج الانحدار ، كما ويمكن الاستفادة من المكونات

(العوامل) في تحديد الحالات الشاذة . أن المكونات (العوامل) هي متغيرات معيارية بمتوسط مساو للصفر وانحراف معياري مساو للواحد فإذا كانت هذه العوامل تتبع التوزيع الطبيعي فعليه تعتبر قيم العامل التي تقع خارج المدى (2,2-) قيمة شاذة .

أن رسم مخطط الانتشار Scatter Plot للمكونين المستخلصين في هذا المثال يساعد في تحديد الحالات الشاذة . لرسم هذا المخطط نتبع الخطوات التالية :

من شريط القوائم اختر Scatter (Simple) → Graphs فيظهر صندوق حوار Scatter Plot .

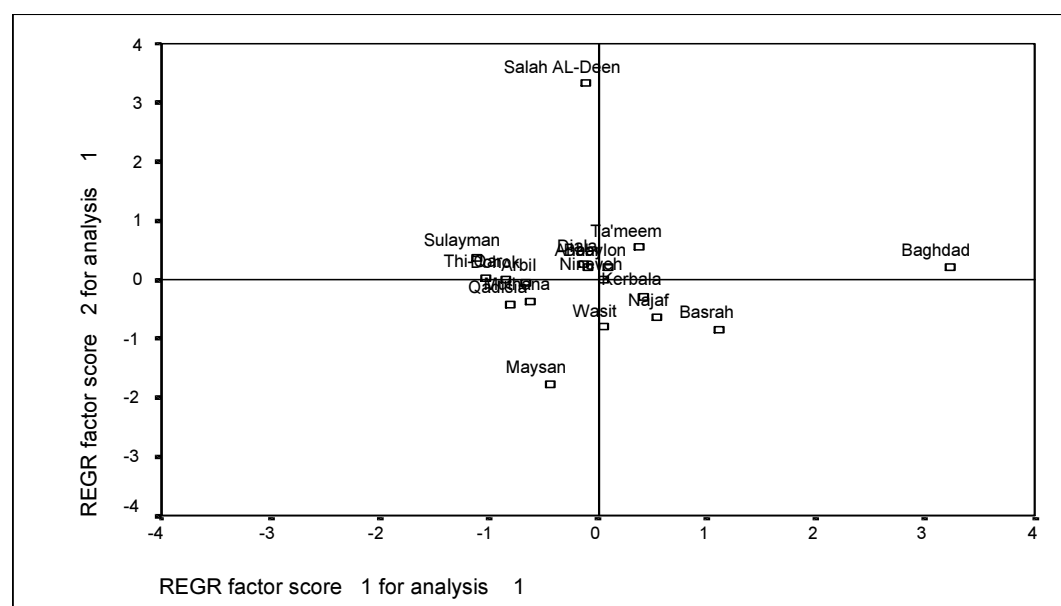
عرف Y Axis : fac_2 .

عرف X Axis : fac_1 .

Label Cases by : Region

أقر زر OK .

(عند ظهور الرسم أقره مرتين ثم اختر Reference Line → Chart من شريط قوائم شاشة Chart Editor لكل من X Axis و Y Axis عند الصفر) يظهر المخطط التالي :



في حالة أن التوزيع الاحتمالي للمكونات (العوامل) يقترب من التوزيع الطبيعي فإن نقاط شكل الانتشار تتوزع بشكل دائري حول النقطة (0,0). لرصد القيم الشاذة نلاحظ أن كافة نقاط المكونين تقع ضمن المدى (2,2-) عدا نقطتين الأولى لمحافظة بغداد التي لها قيمة شاذة للعامل الأول (تقرب من 3) والثانية لمحافظة صلاح الدين التي لها قيمة شاذة للعامل الثاني (تقرب من 3 أيضاً) ولغرض تحديد المتغيرات المسؤولة عن القيم الشاذة للمحافظتين (الحالتين) المذكورتين نقوم بمقارنة بيانات المحافظة مع متوسط كافة المحافظات للمتغيرات ذات التشعبات الكبيرة في العامل ، فبالنسبة لمحافظة بغداد فأن هناك ثلاثة متغيرات تنتشعب بدرجة عالية بالعامل الأول هي literacy و highedu و doctors وقد كانت قيمة المتغيرات لمحافظة بغداد أعلى بكثير من متوسط هذه المتغيرات لعموم محافظات القطر كما يبين ذلك الجدول التالي وهذا هو سبب الحصول على قيمة كبيرة لمحافظة بغداد للعامل الأول (3.227).

Case Summaries

Statistics: Mean

	LITIRACY	HIGHEDUC	DOCTORS
Baghdad	61.60	4.04	36.00
Total	41.42	1.53	17.83

اما سبب القيمة الشاذة (3.325) لمحافظة صلاح الدين فيرجع الى القيمة العالية لمتغير GDP (الذي يتشبع بدرجة عالية بالمكون الثاني) بالمقارنة مع متوسط المتغير لكافة المحافظات وعلى الرغم من انخفاض HOSBED في المحافظة بالمقارنة بمتوسط القطر الا أن هذا المتغير يتشبع بدرجة عالية بالمكون الثاني ولكن باشارة سالبة (أنظر مصفوفة المكونات Components Matrix) . المقارنة يوضحها الجدول التالي :

Case Summaries

Mean

	GDP	HOSPBED
Salah AL-Deen	98.40	80.00
Total	28.52	174.06

ملاحظات :

1. يمكن الحصول على المتجه الكامن (أو المميز) Eigen Vector المرافق للجذر الكامن بإيجاد طول المتجه NORM حسب المعادلة $\|V\| = \sqrt{X_1^2 + X_2^2 + \dots + X_n^2}$ ولكل مكون في المصفوفة Components Matrix وذلك بقسمة كل قيمة من قيم المكون (التشبع) على طول المكون Norm للحصول على المتجه المعياري Normalized Vector الذي يكون طوله واحد وهذا يكون متجه وحيد Unique Vector فمثلاً تحسب القيمة الأولى للمتجه الكامن الأول من تشبع المتغير GDP بالعامل الأول كما يلي :

$$\frac{0.477}{\sqrt{0.447^2 + 0.918^2 + 0.941^2 + 0.829^2 + 0.508^2}} = 0.280$$

الجدول التالي يمثل المتجهين الكامين الأول والثاني المرافقين للجذرين الكامين الأول (2.9) والثاني

(1.11) لمصفوفة الارتباطات .

	Eigen Vector 1	Eigen Vector 2
GDP	0.280278	0.715757
LITERACY	0.538845	0.018025
HIGHEDUC	0.552474	0.096474
DOCTORS	0.486656	-0.12474
HOSPBED	0.298378	-0.68007

2. أن المكونات الأساسية تكون متعامدة فيما بينها Orthogonal أي أن حاصل الضرب القياسي Dot Product لها يساوي صفر $\langle fac_1, fac_2 \rangle = \sum fac_1 * fac_2 = 0$ أي مجموع حاصل

ضرب القيم المتقابلة للمكونين . كذلك فأن حاصل الضرب القياسي لكل من متجهي معاملات المكونات Components Score Coefficients والمتجهين الكامنين يساوي صفر .

(11 - 3) طرق التحليل العاملي Factor Analysis Methods

أن طريقة المكونات الأساسية هي واحدة من طرق التحليل العاملي وتعتبر نقطة البداية في إجراء التحليل العاملي بأية طريقة أخرى وتتميز بالبساطة وبأنها تعتمد أسلوباً رياضياً في الاحتساب .

أما طرق التحليل العاملي بصورة عامة فيعبر عنها نموذج التحليل العاملي Factor Analysis Model الذي يعبر عن كل متغير من متغيرات الاستجابة $X_1, X_2, \dots, X_p$ كدالة في مجموعة من العوامل المشتركة (لمتغيرات الاستجابة) Common Factors وعامل وحيد خاص بذلك المتغير .

بافتراض التوزيع الطبيعي لمتغيرات الاستجابة يكون نموذج التحليل العاملي كما يلي :

$$X_1 = a_{11}Y_1 + \dots + a_{1m}Y_m + e_1$$

$$\dots$$

$$X_p = a_{p1}Y_1 + \dots + a_{pm}Y_m + e_p$$

حيث أن

Y_j : العامل المشترك j (عدد العوامل هو m أقل من عدد مكونات الاستجابة p) .

a_{ij} : التشبعات Loadings وهي معالم تعكس أهمية العامل j في تركيب متغير الاستجابة i .

e_i : العامل الوحيد الخاص بالمتغير X_i (Specific Factor) .

مثال 2 : (على طريقة الإمكان الأعظم Maximum Likelihood Method)

لبيانات المثال 1 سنحاول إجراء التحليل العاملي بطريقة الإمكان الأعظم التي تعتبر من الطرق المهمة في التحليل العاملي لمقارنة أوجه الاختلاف بين هذه الطريقة وطريقة المكونات الرئيسية .

بالنسبة لخطوات التنفيذ فهي مشابهة لما ورد في المثال الأول عدا أننا نختار نوع الطريقة

Maximum Likelihood في صندوق حوار Factor Analysis : Extraction حيث تظهر المخرجات التالية :

Factor Analysis

Communalities

	Initial	Extraction
GDP	.350	.999
LITIRACY	.852	.903
HIGHEDUC	.870	.935
DOCTORS	.503	.487
HOSPBED	.286	.227

Extraction Method: Maximum Likelihood.

أن التشبعات الأولية initial Communalities بالنسبة لطريقة المكونات الأساسية هي واحد دائماً أما باقي طرق التحليل العاملي فأن التشبع الأولي لأي متغير هو مربع معامل الارتباط المتعدد R^2 الذي نحصل عليه من أنحدار المتغير (Dependent Variable) على بقية المتغيرات (Independent Variables) . أن القيمة 0.999 لمتغير الاستجابة GDP تبين أن العاملين المستخلصين الأول والثاني يفسران 0.999 من الاختلافات الكلية للمتغير GDP .

Total Variance Explained

Factor	Initial Eigenvalues			Extraction Sums of Squared Loadings		
	Total	% of Variance	Cumulative %	Total	% of Variance	Cumulative %
1	2.900	58.009	58.009	1.413	28.254	28.254
2	1.110	22.207	80.216	2.138	42.756	71.009
3	.523	10.457	90.673			
4	.393	7.864	98.537			
5	7.313E-02	1.463	100.000			

Extraction Method: Maximum Likelihood.

الجدول أعلاه يبين الجذور الكامنة الأولية Initial Eigenvalues وهي نفسها التي حصلنا عليها لطريقة المكونات الأساسية وفي طريقة المكونات الأساسية نلاحظ أن مجموع مربعات التشعبات المستخلص Extraction Sums Squared Loadings هو نفسه Initial eigenvalues وكذلك تتساوى نسبة التباين المفسر أما في طريقة الإمكان الأعظم فيوجد اختلاف بين القيم الأولية والمستخلصة وفي نسبة التباين المفسر كما هو واضح في الجدول أعلاه فمثلاً يتم احتساب مجموع مربعات تشعبات العامل المستخلص الأول من مصفوفة العوامل أدناه كما يلي :

$$0.999^2 + 0.333^2 + 0.469^2 + 0.275^2 - 8.71E-02^2 = 1.413$$

وأن نسبة ما يفسره العامل الأول هي $1.413 / 5 \times 100 = 28.254\%$

Factor Matrix^a

	Factor	
	1	2
GDP	.999	-9.49E-03
LITIRACY	.333	.890
HIGHEDUC	.469	.846
DOCTORS	.275	.641
HOSPBED	-8.71E-02	.468

Extraction Method: Maximum Likelihood.

a. 2 factors extracted. 14 iterations required.

المصفوفة أعلاه هي مصفوفة العوامل المستخلصة ولها نفس خواص مصفوفة المكونات في المثال 1

Goodness-of-fit Test

Chi-Square	df	Sig.
1.337	1	.247

الجدول أعلاه يمثل اختباراً لفرضية العدم القائلة بأن العوامل المستخلصة كافية لتمثيل هيكل التباينات (أو الارتباطات) ضد الفرضية البديلة القائلة بعدم كفاية العوامل المستخلصة لتمثيل هيكل التباينات. وقد استخدمت لهذا الغرض أحصائية χ^2 حيث أن قيمة $P\text{-value} = 0.247$ تدعونا إلى قبول فرضية العدم

بمستوى دلالة 1% و 5% أي التسليم بكفاية العاملين المستخلصين لتمثيل هيكل التباينات لمتغيرات الاستجابة. ان هذا الاختبار غير موجود لطريقة المكونات الأساسية .

ملاحظة :

أن درجات العوامل Factor Scores وكما هو الحال بالنسبة للمكونات هي دالة خطية في متغيرات الاستجابة المعيارية وأن المعاملات Factor scores Coefficients تحتسب بطريقة واحدة لطريقة المكونات الأساسية ولهذا فهناك مكون (أو عامل) وحيد أما بالنسبة لطريقة الإمكان الأعظم فهناك 3 طرق لتقدير معاملات العوامل ولهذا نحصل على ثلاثة بدائل للعامل الواحد وهذه الطرق كما يلي (تظهر هذه الطرق عند نقر زر Scores في صندوق حوار Factor Analysis) :

Regression :العوامل الناتجة لها متوسط صفر وتباين مساوي لمربع معامل الارتباط المتعدد بين العوامل المقدرة والعامل الحقيقي .

Bartlett : العوامل الناتجة لها متوسط يساوي صفر وأن مجموع مربعات العوامل الخاصة اقل ما يمكن .
Anderson-Rubins :هي تحويل لطريقة Bartlett بحيث أن العوامل الناتجة لها متوسط يساوي صفر وانحراف معياري يساوي واحد وتكون العوامل مستقلة (غير مرتبطة) .

الفصل الثاني عشر

الاختبارات اللامعلمية

Non Parametric Tests

أن الاختبارات اللامعلمية هي اختبارات لا تعتمد إحصائية الاختبار فيها على معالم المجتمع كمعلمة المتوسط Mean أو التباين Variance كما أنها لا تفترض توزيع ما للبيانات ولهذا فهي تعرف أيضاً باختبارات التوزيع الحر Distribution – Free tests. أن سبب استعمال الاختبارات اللامعلمية يعود الى عدم توفر الفرضيات الخاصة بالاختبارات المعلمية فمثلاً يتوجب أن تتوزع بيانات المجتمع قريباً من التوزيع الطبيعي عند تطبيق اختبار T وهو أحد الاختبارات المعلمية شائعة الاستعمال وعند عدم توفر هذا الشرط نلجأ الى الاختبارات اللامعلمية علماً ان هذه الأخيرة لها شروط يجب توفرها ولكنها أسهل بكثير من شروط الاختبارات المعلمية كما أنه يجب استعمال الاختبارات المعلمية في حالة توفر الشروط الخاصة بها كونها أكثر دقة من الاختبارات اللامعلمية. يغطي برنامج SPSS عدداً كبيراً جداً من طرق الاختبار اللامعلمية وفيما يلي توضيح لكيفية تنفيذ بعض هذه الاختبارات .

(12 - 1) اختبار Chi-Square

يستعمل اختبار Chi-Square للمقارنة بين التكرار المشاهد للفئات Observed frequencies والتكرار المتوقع لها Expected Frequencies المحتسب على أساس فرضية العدم. فإذا كان لدينا فئات عددها K فإن أحصائية Chi-Square المستعملة في الاختبار تعطى كما يلي :

$$\chi^2 = \sum_{i=1}^k \frac{(O_i - E_i)^2}{E_i}$$

حيث أن O_i يمثل التكرار المشاهد وأن E_i يمثل التكرار المتوقع وأن درجة حرية الاختبار تساوي

k-1 .

مثال 1:

في تجربة لتهجين صنفين من الشعير تم الحصول على الصفات التالية :

التكرار المشاهد O_i	الصفات type	التسلسل
439	أسود بدون سفا	1
168	أسود ذو سفا	2
133	أبيض بدون سفا	3
60	أبيض ذو سفا	4
800	المجموع	

خاشع الراوي ، المدخل الى الأحصاء ، جامعة الموصل ، 1979، ص 375 .

المطلوب اختبار فرضية العدم التالية بمستوى دلالة 5% :

$$H_0 : P_1 = \frac{9}{16}, P_2 = \frac{3}{16}, P_3 = \frac{3}{16}, P_4 = \frac{1}{16}$$

اما الفرضية البديلة H_1 فتتص على أن النتائج تختلف عن هذه النسب (مثلاً P_1 تمثل نسبة الصفة الأولى أسود بدون سفا) وهكذا .

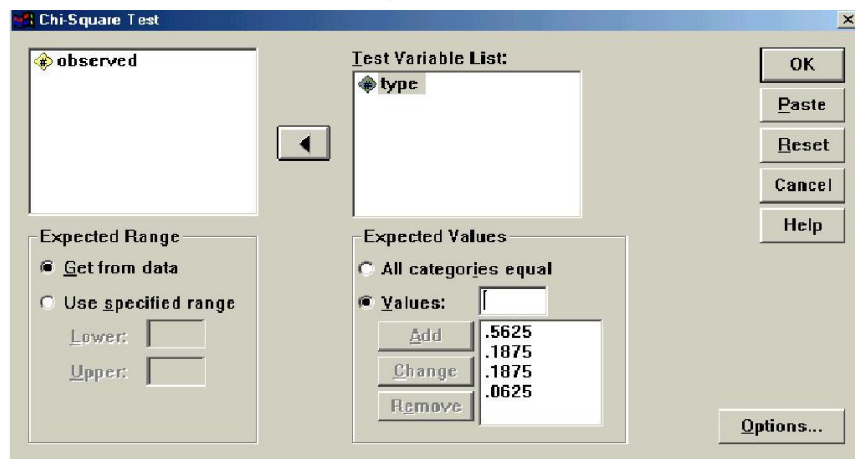
لاختبار الفرضية المذكورة نطبق اختبار Chi-Square وحسب الخطوات التالية :

◀ يتم ترتيب البيانات في ورقة Data Editor كما يلي :

type	Observed
1	439
2	168
3	133
4	60

لاحظ أننا استعصنا عن الصفات بالمتغير العددي Type لعدم إمكانية التعامل مع المتغيرات الرمزية أي عدم إمكانية ذكر أسماء الصفات مع العلم أنه بالإمكان إضافة عنوان القيمة Value Label الى المتغير Type لكي تعرض عناوين القيم في جداول المخرجات بدلاً من الأرقام التسلسلية. أما المتغير Observed فيمثل التكرار المشاهد في كل فئة وتستعمل قيم هذا المتغير كأوزان للحالات ويتم ذلك عن طريق اختيار Data → Weight Cases من شريط القوائم بعدها نؤشر الخيار Weight Cases by في صندوق حوار Weight Cases ثم إدخال المتغير Observed الى خانة الخيار الأخير (Weight Cases by) .

◀ من شريط القوائم اختر Chi-Square → Nonparametric Tests → Analyze فيظهر صندوق حوار Chi-Square Test الذي نرتبه كما يلي :



لاحظ أننا ادخلنا المتغير Type الى خانة Test Variable List .

في خانة Expected Values يوجد خيارين :

1. All categories equal : جميع الفئات لها نفس التكرار المتوقع أولها نفس النسبة .
2. Values : يتم تأشير هذا الخيار عندما يختلف التكرار المتوقع (وبالتالي تختلف النسب) من فئة الى أخرى كما هو الحال في هذا المثال فقد أدخلت النسب المحددة في فرضية العدم وكما يلي :

لأدخال القيمة الأولى $P_1 = \frac{9}{16} = 0.5625$ أنقر المستطيل المجاور لكلمة Value تم أكتب القيمة

0.5625 بعدها أنقر زر Add فتضاف أسفل الكلمة value (نفس الشيء لباقي القيم) .

◀ عند نقر زر OK نحصل على النتائج التالية :

TYPE

	Observed N	Expected N	Residual
1	439	np1 450	-11.0
2	168	np2 150	18.0
3	133	np3 150	-17.0
4	60	np4 50	10.0
Total	800		

الجدول أعلاه يمثل القيم المشاهدة والمتوقعة فمثلاً تحتسب القيمة المتوقعة للنوع الأول كما يلي
 $nP_1 = 800 * 0.5625 = 450$ لاحظ أن مجموع النسب يساوي 1 وعليه فإن حجم العينة للقيم المتوقعة هو نفسه للقيم المشاهدة .

Test Statistics

	TYPE
Chi-Square ^a	6.356
df	3
Asymp. Sig.	.096

a. 0 cells (.0%) have expected frequencies less than
 5. The minimum expected cell frequency is 50.0.

من الجدول أعلاه فإن قيمة إحصائية Chi-Square تساوي 6.356 وأن قيمة P-value المرافقة 0.096 تدعونا إلى قبول فرضية العدم بمستوى دلالة 5% أي قبول النسب الواردة في فرضية العدم
 → باعتبارها نسباً صحيحة
 يمكنك استخراج P-Value من الأمر Compute Transform ثم استعمال الدالة التالية لتوزيع مربع كاي :

$$1 - \text{CDF.CHISQ}(6.356, 3) = .096$$

مثال 2 :

رميت خمسة قطع نقود 1000 مرة وفي كل مرة حسبت عدد الصور Head وكانت النتائج كالتالي
 (بعد إدخالها في ورقة Data Editor) :

headno	observed
0	38
1	144
2	342
3	287
4	164
5	25

باستخدام اختبار Chi-Square أختبر الفرضية القائلة بأن عدد مرات ظهور الصورة يتبع توزيع Binomial وبمستوى دلالة 5% .

أن اختبار Chi-Square هنا يعرف باختبار حسن المطابقة Goodness of Fit. في البداية تقل (أوزن) الحالات بالمتغير Observed بنفس الطريقة المتبعة في المثال السابق ثم أتبع الخطوات التالية:
 ← من شريط القوائم اختر Chi-Square → Nonparametric Tests → Analyze فيظهر صندوق حوار Chi-Square Test (نفسه للمثال السابق) الذي نرتبه كما يلي :

- أنقل المتغير headno الى خانة Test Variable List .
- أستعمل دالة توزيع binomial لحساب احتمال ظهور الصورة من صفر الى خمسة وكانت النتائج كما يلي :

No. of heads	Probability
0	0.03125
1	0.15623
2	0.3125
3	0.3125
4	0.15623
5	0.03125

- أنقر الخيار values في خانة Expected values ثم أدخل النسب (الاحتمالات) أعلاه واحدة بعد الأخرى .
- عند نقر زر OK نحصل على النتائج التالية :

HEADNO

	Observed N	Expected N =1000*Pr	Residual
0	38	31.2	6.8
1	144	156.2	-12.2
2	342	312.6	29.4
3	287	312.6	-25.6
4	164	156.2	7.8
5	25	31.2	-6.2
Total	1000	1000.0	

يستخرج التكرار المتوقع للفئة الأولى مثلاً (عدد الصور يساوي صفر) كما يلي :

$$\text{Expected Frequency} = 1000 * 0.03125 = 31.25$$

Test Statistics

	HEADNO
Chi-Square ^a	8.920
df	5
Asymp. Sig.	.112

a. 0 cells (.0%) have expected frequencies less than 5. The minimum expected cell frequency is 31.2.

أن قيمة $P\text{-value} = .112$ تدعونا الى قبول فرضية العدم بمستوى دلالة 5% أي أن عدد الصور المشاهد يتبع توزيع Binomial .

ملاحظة : يمكن أخذ مدى من الفئات بدلاً من كافة الفئات ففي المثال السابق سنهمل الفئة الأولى (عدد الصور = 0) وسنكتفي بالفئات الخمس المتبقية (عدد الصور = 1-5) وفي هذه الحالة سنكتفي بتزويد الاحتمالات لهذه الفئات فقط (1-5) في خانة Expected Values في صندوق حوار Chi-square Test فما دمنا قد أدخلنا الاحتمالات المقابلة لعدد الصور (0-6) في المثال السابق نقوم الآن بحذف الاحتمال

المقابل لعدد الصور =0 بتأشير هذا الاحتمال ويساوي 0.03125 في خانة Expected Values ثم نقر زر Remove .

بما أن الفئات المتبقية هي خمسة فئات (1-5) نقوم باختيار هذا المدى من الفئات في خانة Expected Range بتأشير الخيار Use Specified Range ثم إدخال قيمة Lower وتساوي 1 وقيمة Upper وتساوي 5 ،وعند نقر زر OK نحصل على النتيجة التالية :

Frequencies				
	HEADNO			
	Category	Observed N	Expected N	Residual
1	1	144	155.1	-11.1
2	2	342	310.4	31.6
3	3	287	310.4	-23.4
4	4	164	155.1	8.9
5	5	25	31.0	-6.0
Total		962	962	

(12 - 2) اختبارات عينتين مستقلتين *Two Independent Samples tests*

هذه الاختبارات مقارنة لاختبار T لمقارنة متوسطي عينتين مستقلتين حيث يستعمل اختبار T في حالة أن إحصائية الاختبار تتبع توزيع T عدا ذلك لا يمكن استعمال هذا الاختبار لعدم توفر شرط التوزيع الطبيعي لمجموعي العينتين ولذلك نلجأ إلى الاختبارات اللامعلمية حيث يوفر برنامج SPSS الاختبارات اللامعلمية التالية :

1. اختبار Mann-Whitney U
2. اختبار Kolmogorov-Smirnov Z
3. اختبار Moses Extreme Reactions
4. اختبار Wald-Wolfwitz Runs

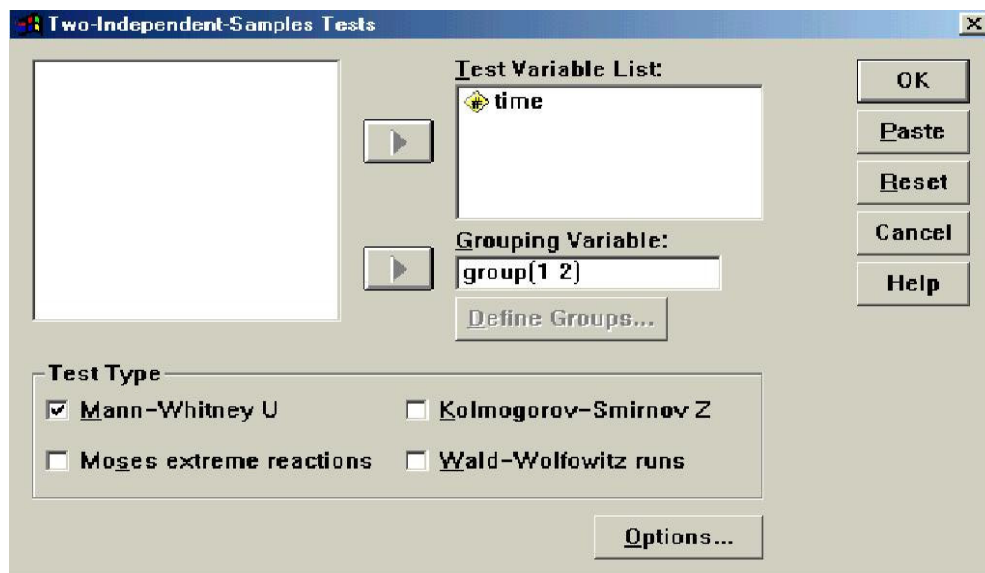
مثال 3: البيانات التالية تمثل عينتين مستقلتين الأولى بحجم 8 للأشخاص الأصحاء (يرمز لهم 1) والثانية بحجم 10 للأشخاص المرضى (الرمز 2) والمجموعتين من الأشخاص يضمها المتغير Group علماً أنه تم إضافة Value Label لهذا المتغير (القيمة 1 يقابلها العنوان Healthy والقيمة 2 يقابلها العنوان disease لكي تظهر العناوين بدلاً عن القيم في جداول الإخراج) وقد تم قياس الوقت المستغرق في فحص الجهد (المتغير Time) لكل شخص في العينتين وتظهر البيانات كما يلي في شاشة Data Editor لبرنامج SPSS :

time	group
1014	1
684	1
810	1
990	1
840	1
978	1
1002	1
1110	1
864	2
636	2
638	2
708	2
786	2
600	2
1320	2
750	2
594	2
750	2

يطلب اختبار فرضية عدم القائل بأن العينتين مسحوبتين من نفس المجتمع بمستوى دلالة 5% باستخدام اختبار Mann-Whitney (بمعنى آخر اختبار الفرضية عدم التالية $\mu_1 = \mu_2$ ضد البديلة $\mu_1 \neq \mu_2$).

لتنفيذ ذلك نتبع الخطوات التالية :

← من شريط القوائم اختر Analyze → Non Parametric Tests → 2 Independent samples
فيظهر صندوق حوار 2 Independent samples Test الذي نرتبه كما يلي :



لاحظ أننا أدخلنا المتغير Time (الذي يمثل مشاهدات العينتين) في خانة Test Variable List أما المتغير group فقد أدخلناه في خانة Grouping Variable وهو متغير تجزئة لتعريف مجموعتي هذا المتغير انقر زر Define Group ثم عرف مجموعة الأصحاء 1 : Group1 ومجموعة المرضى 2 : Group2. أخيراً تم تأشير اختبار Mann-Whitney U .
← عند نقر زر OK تظهر النتائج التالية :

Ranks

	GROUP	N	Mean Rank	Sum of Ranks
TIME	healthy	8	12.63	101.00
	disease	10	7.00	70.00
	Total	18		

يعتمد هذا الاختبار على الرتب Ranks حيث يتم دمج العينتين (الأصحاء والمرضى) في عينة واحدة ثم ترتيب قيم العينة المدمجة تصاعدياً واعطاء رتبة لكل قيمة تم إيجاد مجموع رتب كل عينة (مجموع رتب الأصحاء 101 ومجموع رتب المرضى 70) كما في الجدول أعلاه ثم تحتسب إحصائية Mann-Whitney U بالصيغة التالية :

$$U = N_1 N_2 + \frac{N_1(N_1+1)}{2} - T_1$$

حيث أن N_1 يمثل حجم أي من العينتين (لنسمها الأولى $N_1=8$ و N_2 يمثل حجم العينة الأخرى (الثانية $N_2=10$) وأن T_1 يمثل مجموع رتب العينة الأولى (Sum of Ranks $T_1=101$). الجدول التالي يحتوي قيمة إحصائية U المحتسبة بالاعتماد على معطيات العينة الأولى بالإضافة الى إحصائية Wilcoxon W وهي مكافئة لإحصائية U . لاستخراج قيمة Z المرافقة يجب معرفة توزيع المعاينة لـ U حيث يستخرج المتوسط والانحراف المعياري للتوزيع كما يلي :

$$\mu = \frac{N_1 N_2}{2} = 40 \quad \sigma = \sqrt{\frac{N_1 N_2 (N_1 + N_2 + 1)}{12}} = 11.25$$

أن توزيع المعاينة لـ U يقارب التوزيع الطبيعي كلما كان حجم العينتين كبيراً وعليه يمكن استخراج القيمة المعيارية للتوزيع الطبيعي كما يلي $Z = (U - \mu) / \sigma = (15 - 40) / 11.25 = -2.222$ ويمكنك أستخراج قيمة P -Value المرافقة لإحصائية Z بالأمر Compute Transform وباستعمال دالة التوزيع الطبيعي القياسي وكما يلي $2 * \text{CDFNORM}(-2.222) = 0.0263$ (الاختبار من طرفين) حيث أن قيمة P -value=0.026 للتوزيع الطبيعي التقاربي تدعونا الى رفض فرضية العدم بمستوى دلالة 5% والأخذ بالفرض البديل أي أن العينتين لم تسحبا من نفس المجتمع أو أن متوسط الوقت المستغرق في فحص الجهد للأشخاص الأصحاء يختلف جوهرياً عن متوسط الوقت المستغرق للأشخاص المرضى .

Test Statistics^b

	TIME
Mann-Whitney U	15.000
Wilcoxon W	70.000
Z	-2.222
Asymp. Sig. (2-tailed)	.026
Exact Sig. [2*(1-tailed Sig.)]	.027 ^a

a. Not corrected for ties.

b. Grouping Variable: GROUP

ملاحظة : يمكنك الاعتماد على مجموع رتب العينة الثانية T_2 الكبيرة في حساب إحصائية U وفي هذه الحالة تحور إحصائية U كما يلي :

$$U = N_1 N_2 + \frac{N_2 (N_2 + 1)}{2} - T_2$$

N_2 يمثل حجم العينة الكبيرة وباستخراج القيمة المعيارية لإحصائية U وأجراء الاختبار يمكن

التوصل الى نفس النتيجة السابقة .

(12 - 3) أختبارات K من العينات المرتبطة K- Related Samples Tests

أن اختبار Mann-Whitney هو نسخة لأمعلمية لاختبار T لعينتين مستقلتين وأن اختبار Kruskal-Wallis هو اختبار لا معلمي لتحليل التباين لمعيار واحد One- Way ANOVA. أما في حالة اختبار عينتين غير مستقلتين فيستعمل اختبار T لذلك وفي حالة عدم تحقق الشروط اللازمة لاختبار T أو ان

البيانات عبارة عن رتب Ranks في هذه الحالة يمكن إجراء أحد الاختبارات اللامعلمية التالية التي يوفرها برنامج SPSS :

1. اختبار Friedman .
2. اختبار Kendall's W .
3. اختبار Cochran's Q .

مثال 4 :

قام تسعة مختصين Therapists بتخصيص رتب لثلاثة نماذج للمحفزات الكهربائية a و b و c (الرتبة 1 تشير الى درجة التفضيل الأولى تليها الرتبتين 2 و 3) وكانت النتائج⁵ كما تظهر في شاشة Data Editor

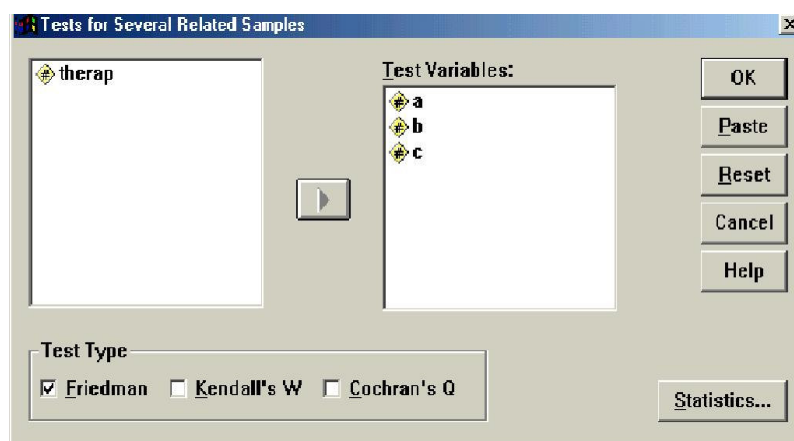
Therap	a	b	c
1	2	3	1
2	2	3	1
3	2	3	1
4	1	3	2
5	3	2	1
6	1	2	3
7	2	3	1
8	1	3	2
9	1	3	2

لبرنامج SPSS :

المطلوب اختبار فرضية عدم القائلة بعدم وجود فروقات معنوية في درجة التفضيل للنماذج الثلاثة باستخدام اختبار Friedman وبمستوى دلالة 5% .

لتنفيذ ذلك نتبع الخطوات التالية :

☐ من شريط القوائم اختر Analyze → Nonparametric Tests → K-Related samples فيظهر صندوق حوار Tests for several Related Samples الذي نرتبه كما يلي :



☐ عند نقر زر OK نحصل على النتائج التالية :

⁵Daniel W.W.(1978) , Biostatistics : A Foundation for analysis in The Health Sciences , 2nd Edition, pp.398 .

NPar Tests Friedman Test

Ranks

	Mean Rank
A	1.67
B	2.78
C	1.56

حيث أن Mean Rank تمثل متوسط الرتب في كل من النماذج الثلاثة . أن التجربة أعلاه تشبه تحليل التباين لمعيارين Two- way ANOVA أو ما يعرف بتصميم القطاعات العشوائية Randomized Blocks Experiment حيث تمثل الأعمدة (أو النماذج) المعالجات Treatments وان الصفوف تمثل القطاعات Blocks حيث يعتبر اختبار Friedman ملائماً لأجراء تحليل التباين لمعيارين حسب الرتب Ranks ويتطلب أن تكون البيانات ترتيبية Ordinal في الأقل وأن إحصائية Friedman المستخدمة تتبع توزيع χ^2 بدرجة حرية J-1 وتحتسب كما يلي:

$$\chi^2 = \frac{12}{KJ(J+1)} \left[\sum T_i^2 \right] - 3K(J+1) = 8.2$$

حيث أن K يمثل عدد الصفوف (القطاعات) وأن J يمثل عدد المعالجات (الأعمدة) وأن T_i يمثل مجموع الرتب لكل معالجة أي أن $K=9$ و $J=3$.
الجدول أدناه يبين قيمة إحصائية Friedman وقيمة P-Value المرافقة التي تدعونا الى رفض فرضية العدم بمستوى دلالة 5% أي أن هناك فروقاً معنوية في درجة التفضيل للنماذج الثلاثة .

Test Statistics^a

N	9
Chi-Square	8.222
df	2
Asymp. Sig.	.016

a. Friedman Test

الفصل الثالث عشر

المخططات البيانية

CHARTS

تعتبر المخططات البيانية أداة مهمة من أدوات الإحصاء الوصفي والتي يمكن بواسطتها عرض البيانات الإحصائية بطريقة مبسطة ومعبرة . من المخططات المهمة الأعمدة البيانية Bars والخطوط البيانية Lines والدوائر البيانية Pies... سنحاول في هذا الفصل أن نستعرض الإمكانيات والتسهيلات التي يوفرها برنامج SPSS في مجال معالجة الرسوم البيانية .

Bar Charts الأعمدة البيانية (1-13)

مثال 1 :

الجدول التالي يبين المبيعات Sales لإحدى المؤسسات حسب السنوات Year :

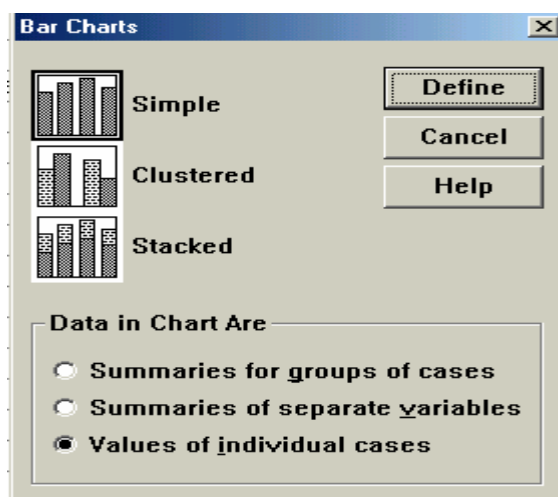
year	sales
1990	50
1991	52
1992	55
1993	60
1994	65

يطلب أعداد مخطط الأعمدة البيانية Bar Chart .

لتنفيذ ذلك نتبع الخطوات التالية :

من شريط القوائم اختر Bar → Graphs. فيظهر صندوق حوار Bar Chart وقد

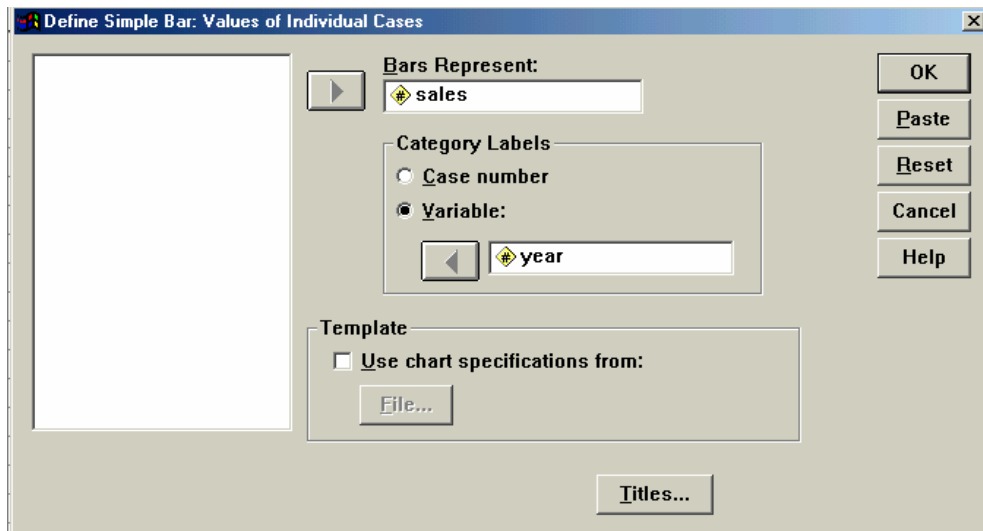
رتبناه كالتالي:



وقد اخترنا Values of individual cases في قائمة Data in Chart are ويقابل كل خيار في هذه القائمة ثلاثة خيارات وهي Simple ، Clustered ، Stacked وقد اخترنا Simple وهو أبسط الحالات .

عند نقر زر Define في صندوق الحوار السابق يظهر صندوق حوار Define Simple Bar

الذي نرتبه كالتالي :



. حيث أن :

Bars Represent : متغير عددي كل قيمة من قيمه تمثل بشريط في المخطط .

Category Labels : وهو عناوين الفئات للمخطط البياني (المحور السيني) ويتضمن ما يلي :

Case number : يعرض رقم الحالة كعنوان للقيمة على المحور السيني .

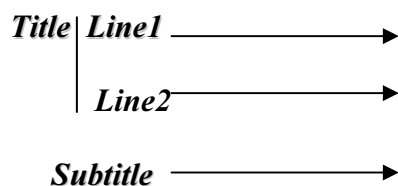
Variable : يعرض قيم المتغير الموجود في هذه القائمة كعناوين لقيم المتغير على المحور السيني (في هذا المثال المتغير Years) .

Title : لعرض عناوين المخطط Title , subtitle , Footnote .

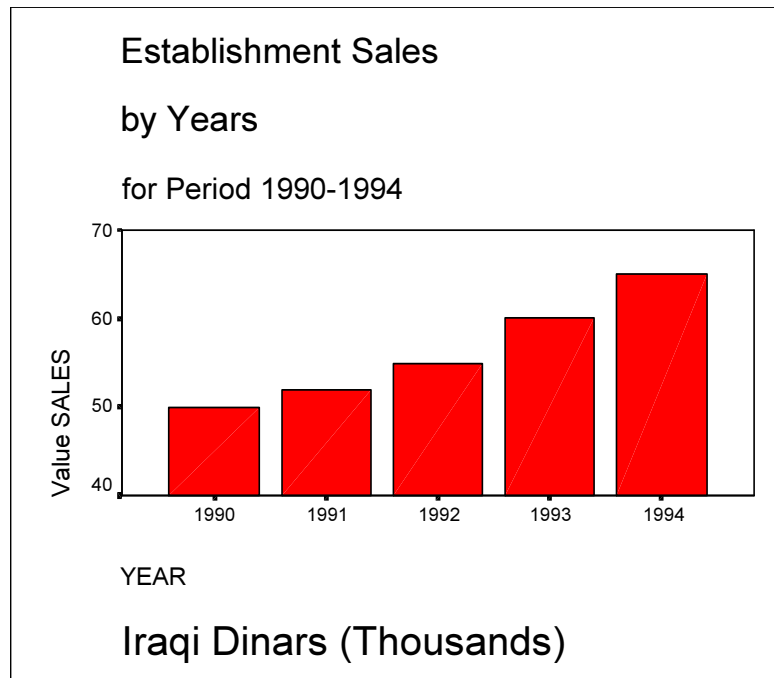
Template : لعمل قالب يحتوي مواصفات معينة يمكن تطبيقها مباشرة عند أعداد مخططات أخرى .

عند نقر زر OK يتم عرض المخطط في SPSS Viewer كما يلي :

مخطط رقم 1

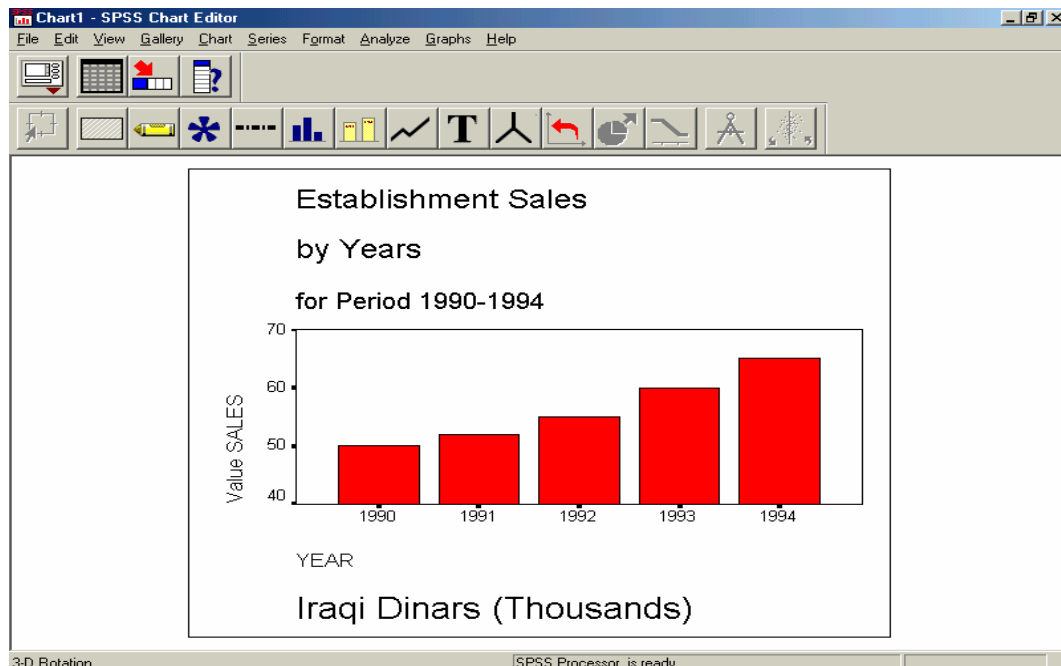


Footnote —————>



يمكن إجراء تعديلات على المخطط بنقره مرتين بزر الماوس الأيسر حيث يعرض المخطط في SPSS Chart Editor كما يلي :

شاشة تنقيح المخططات SPSS Chart Editor



حيث يمكن إجراء التعديلات التالية (كافة التعديلات التي ستجرى لاحقاً تتم في شاشة Chart Editor أي بعد نقر المخطط مرتين بزر الماوس الأيسر) :

1. تغيير لون الأعمدة : يتم تنفيذ ذلك كما يلي (كما ذكرنا أن التعديلات تجري في شاشة Chart Editor كما في الشكل أعلاه) :

➤ أنقر الأعمدة بزر الماوس الأيسر .

➤ من شريط القوائم في شاشة Chart Editor اختر Color → Format أو أنقر أيقونة

في شريط الأدوات فيظهر صندوق حوار Colors كما يلي :



حيث أن :

Fill : لتغيير لون إملاء الشريط .

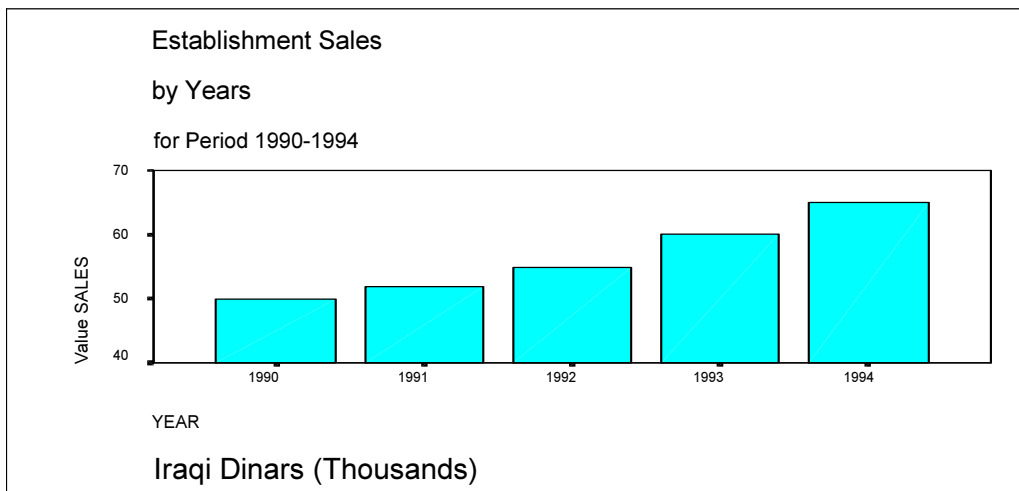
Border : لتغيير لون الحدود الخارجية للشريط .

◀ لتغيير لون إملاء الشريط الى اللون الأزرق مثلاً اختر اللون الأزرق بنقره بزر الماوس الأيسر بعد أن تتأكد من تأشير الخيار Fill ثم أنقر زر Apply فيتغير لون الأعمدة .

◀ أنقر زر Close لغلاق صندوق حوار Colors . ويظهر المخطط بعد تغيير لون الأعمدة كما

يلي :

مخطط رقم 2



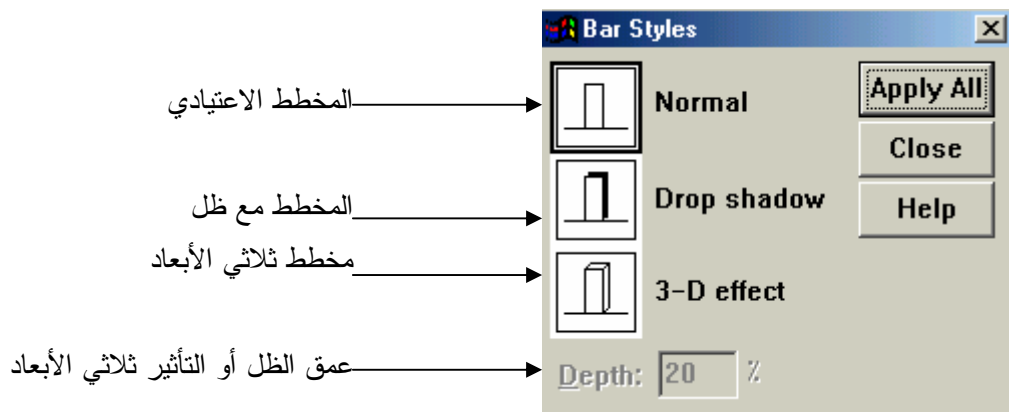
ملاحظة :

يمكن تغيير لون الخلفية مثلاً بنفس الخطوات السابقة (عدا الخطوة الأولى حيث نقوم بنقر خلفية المخطط بدلاً من نقر الأعمدة) .

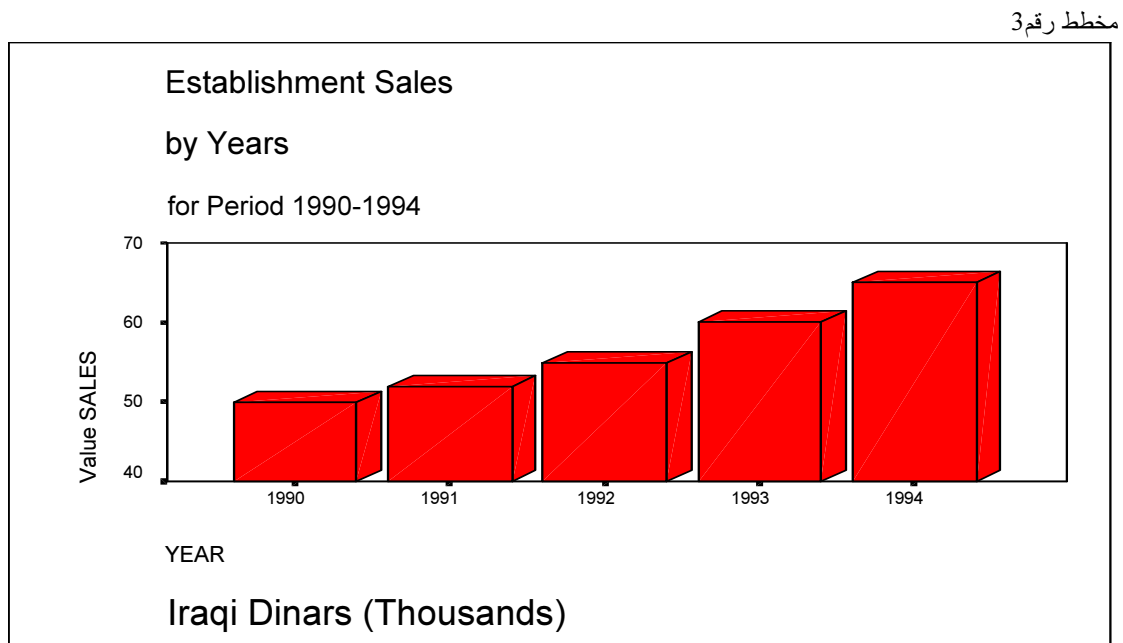
2. تغيير نمط الأعمدة Bar Style : لتنفيذ ذلك نتبع مايلي :

◀ من شريط القوائم في شاشة Chart Editor اختر Bar Style → Format أو أنقر

الأيقونة  في شريط الأدوات فيظهر صندوق الحوار التالي :



عند اختيار 3-D effect بنقره بزر الماوس الأيسر ثم نقر زر Apply All تظهر الأعمدة ثلاثية الأبعاد كما يلي :

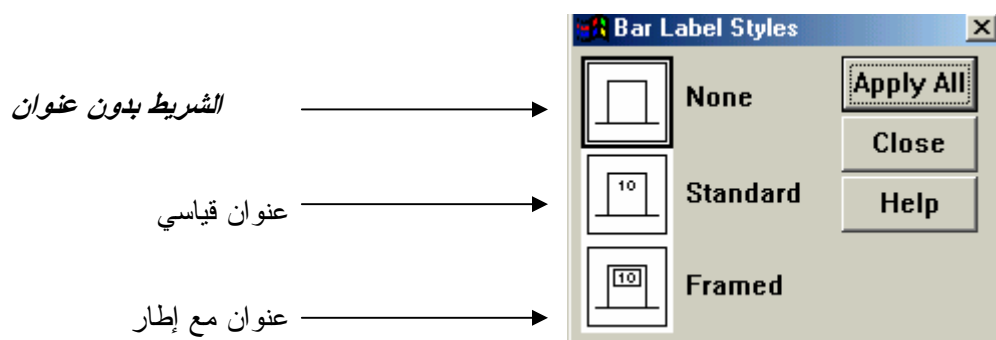


ملاحظة :

عند تحديد رقم موجب في Depth يتم إضافة التأثير الى يمين الشريط وعند تحديد رقم سالب يضاف التأثير الى يسار الشريط .

3. تغيير نمط عناوين الأعمدة Bar Label Style : لتنفيذ ذلك نتبع الخطوات التالية :

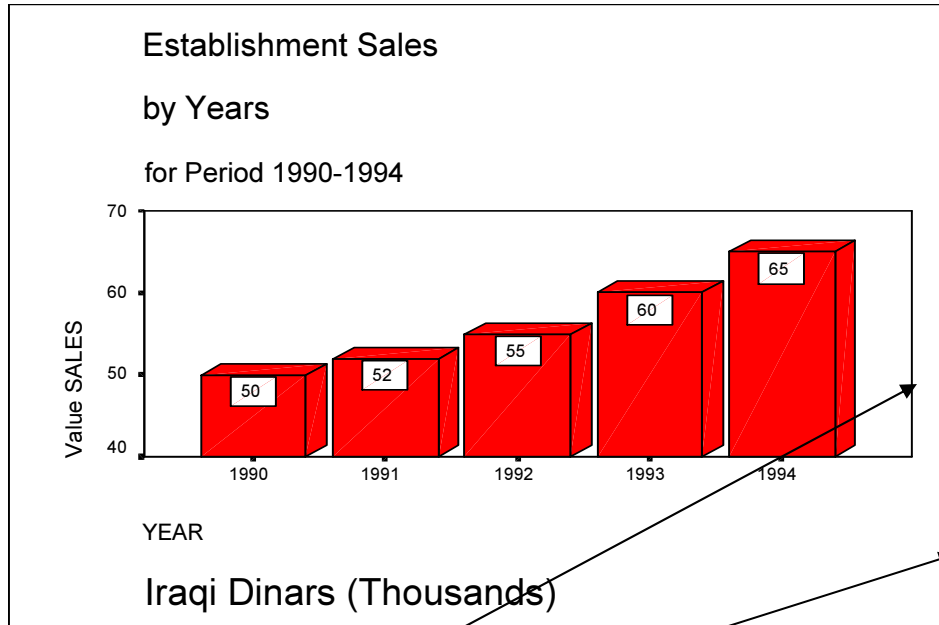
من شريط الأدوات اختر Bar Label Style → Format أو أنقر الأيقونة  فيظهر صندوق



الحوار التالي :

أختر Framed بنقره بزر الماوس الأيسر ثم أنقر زر Apply فنحصل على المخطط التالي :

مخطط رقم 4



الإطار الداخلي الإطار الخارجي

ملاحظات :

أ. لإزالة الإطار الخارجي للمخطط أختر Chart من شريط القوائم في شاشة Chart Editor ثم أنقر Outer Frame لإزالة العلامة منها .

ب. لإزالة الإطار الداخلي للمخطط أختر Chart من شريط القوائم في شاشة Chart Editor ثم أنقر Inner Frame لإزالة العلامة منها .

4. تغيير الإعدادات الخاصة بالمحور الرأسي Scale Axis

تكون خطوات تغيير الإعدادات كما يلي :

أنقر مرتين على عناوين المحور الرأسي (40،50،60...) للمخطط في شاشة Chart Editor فيظهر صندوق حوار Scale Axis كما يلي :

حيث أن :

- a. Display Axis Line: عرض أو إزالة خط المحور الرأسي (لإزالة المحور الرأسي يتوجب قبلها إزالة الإطار الداخلي بالأمر Chart ثم Inner Frame) .
- b. Title Justification: لتغيير موقع عنوان المحور الرأسي (Value Sales) .
- c. Scale: لتحديد نوع المقياس (خطي، لوغاريتمي) .
- d. Major divisions: التقسيمات الرئيسية .
- e. Minor divisions: التقسيمات الثانوية .
- f. Increment: المسافة بين تقسيمين علماً أن $\geq \text{Major Increment} \text{ Minor Increment}$
- g. Grid: عرض خطوط الشبكة .
- h. Ticks: لعرض Tick Marks .
- i. Display Labels: عرض أو إزالة عناوين التقسيمات (60, 50, 40, ...)

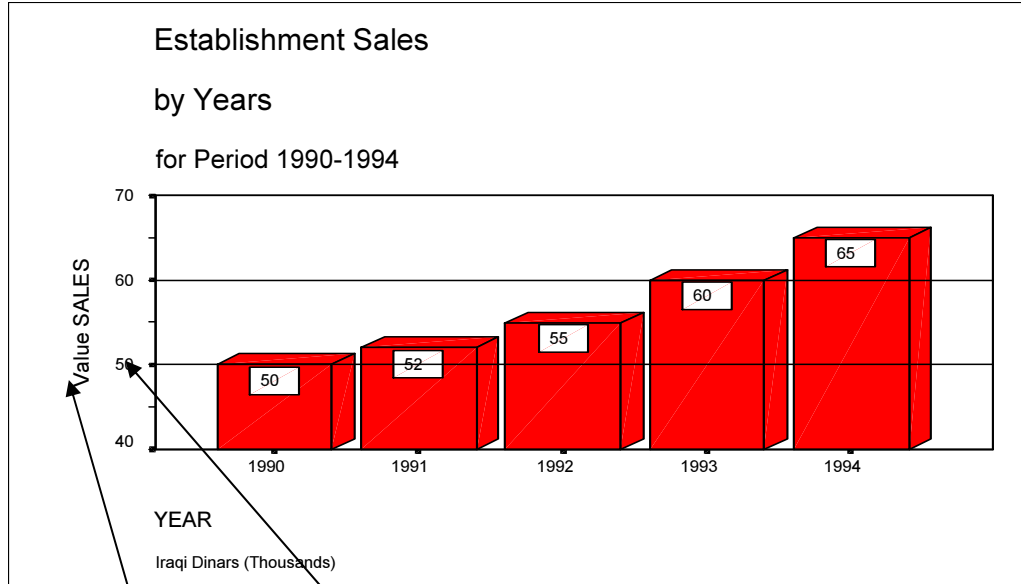
إذا أردنا إجراء الفعاليات التالية على المخطط رقم 4 :

- وضع عنوان المحور الرأسي (الذي هو Value Sales) في الوسط Center .
- جعل المسافة بين التقسيمات الثانوية Minor Increment هي 5 بدلاً من 10 مع إظهار علامات التقسيم الثانوية Tick Marks .
- إضافة خطوط الشبكة Grid Lines إلى التقسيمات الرئيسية .

فإن صندوق حوار Scale Axis يرتب كالتالي (نحصل عليه بنقر عناوين المحور الرأسي مرتين ثم نقوم بالتعديلات اللازمة) :

أما المخطط فيظهر كما يلي بعد نقر الزر OK :

مخطط رقم 5

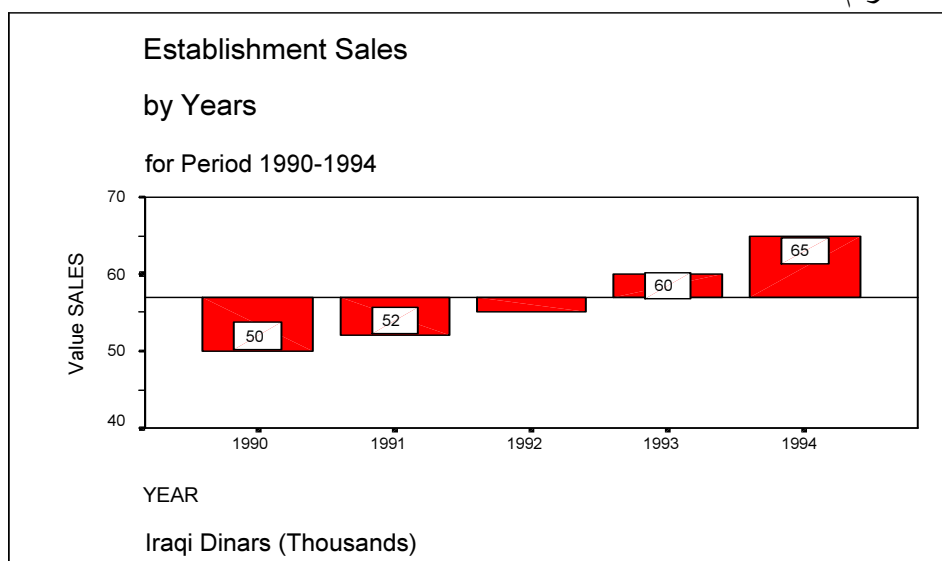


Axis Title

Axis Label

- J. Labels : عند نقر هذا الزر يظهر صندوق حوار Labels ومن خلاله يمكن القيام بالعمليات التالية:
- Decimal Places : لتحديد عدد المراتب العشرية لعناوين المحور الرأسي فإذا حددنا عدد المراتب العشرية يساوي واحد فستعرض العناوين كما يلي : 40.0، 50.0، 60.0، 70.0 .
 - Leading Character : وهو رمز يضاف الى بداية كل عنوان من عناوين المحور الرأسي فإذا أردنا إضافة رمز الدينار العراقي D فستظهر العناوين كما يلي : D40 ، D50 ، D60 ، D70 .
 - Trailing Character : وهو رمز يضاف الى نهاية كل عنوان من عناوين المحور الرأسي فإذا أردنا إضافة علامة % فستظهر العناوين كما يلي : 40%، 50%، 60%، 70% .
 - 1000 S Separator : لعرض القيم التي تزيد عن 1000 بفواصل (فارزة Comma أو فترة Period) .

K. Bar Origin Line : لرسم خط الأصل Origin Line وهو خط يظهر في المخطط البياني عند قيمة محددة بحيث أن الأعمدة التي تمثل قيماً أعلى من قيمة خط الأصل سوف تمتد فوق الخط أما الأعمدة التي تمثل قيماً أقل من قيمة الخط فسوف تمتد تحته. إذا حددنا قيمة Bar Origin Line تساوي 57 في صندوق حوار Scale Axis للمخطط رقم 4 (بعد إزالة التأثير الثلاثي الأبعاد) فسوف نحصل على المخطط التالي :



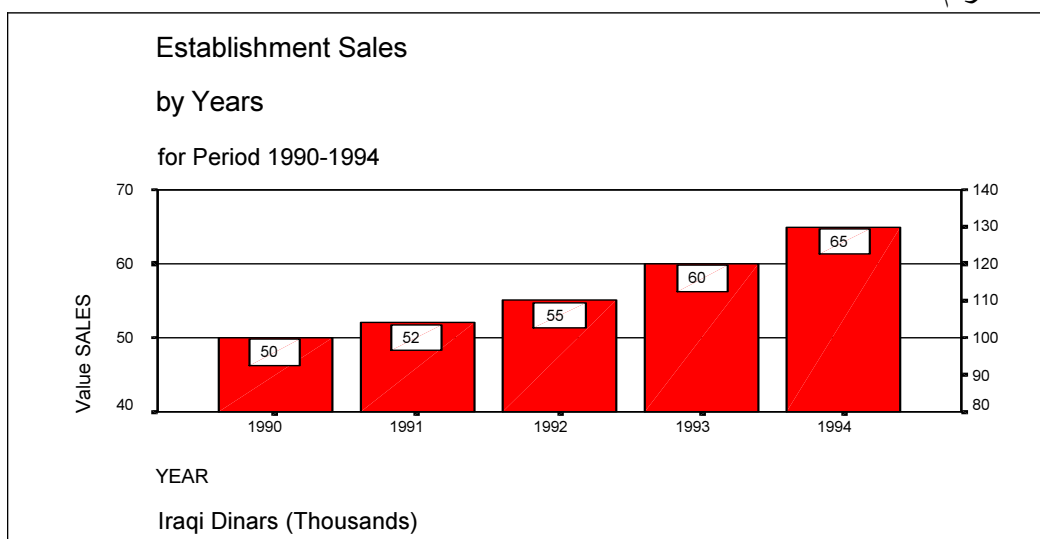
حيث أن الخط الوسطي Origin Line عند النقطة 57 .

L. Display Derived Axis : لعرض محور رأسي إضافي في الاتجاه المقابل للمحور الرأسي الأصلي ويكون له مقياس مختلف ويتم ذلك بتأشير Display Derived Axis في صندوق حوار Scale Axis ثم نقر زر Derived Axis فيظهر صندوق حوار Derived Axis كما يلي :

فاذا أردنا جعل كل وحدة من المحور الأصلي Scale Axis تقابل وحدتين من Derived Axis للمخطط رقم 1 فأن قائمة Definition في صندوق الحوار أعلاه ترتب كما يلي :

	Scale Axis		Derived Axis	
Ratio	1	Units Equal	2	Units
Match	0	Value Equal	0	Value

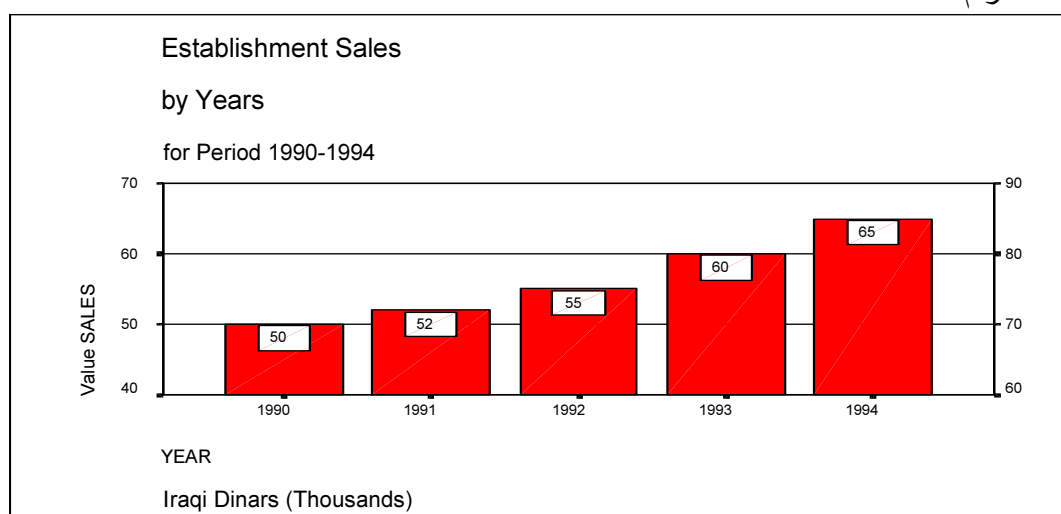
عند نقر زر Continue ثم زر OK نحصل على المخطط التالي :



وقد أضفنا العناوين Labels وكذا الشبكة لغرض التوضيح. نلاحظ أن كل قيمة على المحور الأصلي يقابلها الضعف على المحور المشتق Derived Axis .
 أما الخيار Match فيسمح بتحديد قيمتين مختلفتين (أو متساويتين) على كل من المحور الأصلي والمشتق بحيث تظهران في نفس الموقع. فإذا أردنا جعل القيمة 50 على المحور الأصلي تقابل القيمة 70 على المحور المشتق للمخطط 1 نرتب قائمة Definition في صندوق حوار Derived Axis كما يلي :

	Scale Axis		Derived Axis	
Ratio	1	Units Equal	1	Units
Match	50	Value Equal	70	Value

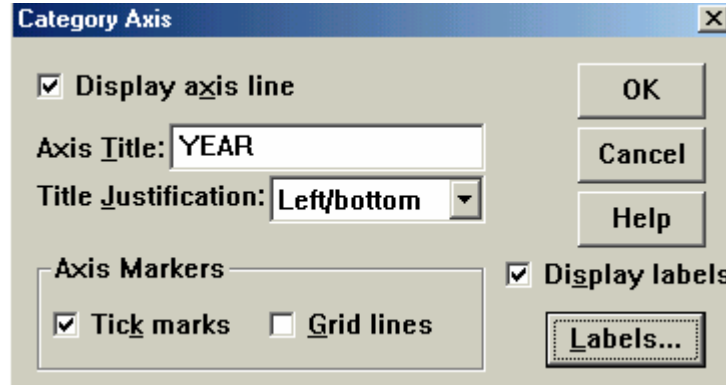
عند نقر زر Continue ثم زر OK نحصل على المخطط التالي :



نلاحظ أن القيمة 50 تقابل القيمة 70 والقيمة 60 تقابل القيمة 80 ... وهكذا .

5. تغيير الإعدادات الخاصة بالمحور الأفقي Category Axis

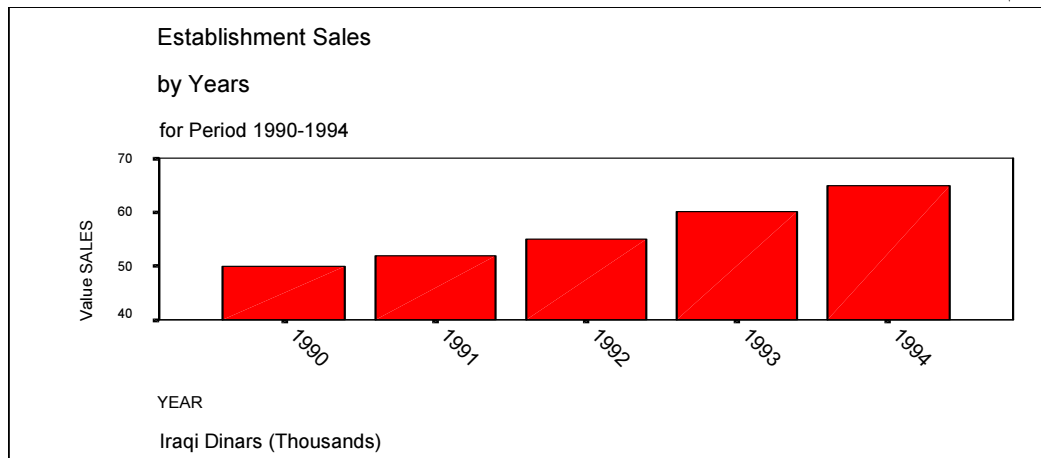
يمكن تغيير الإعدادات الخاصة بالمحور الأفقي بنقر التخطيط مرتين بزر الماوس الأيسر في SPSS Viewer ثم نقر عناوين المحور الأفقي للمخطط مرتين وهو في شاشة Chart Editor فيظهر صندوق الحوار التالي (للمخطط رقم 1 مثلاً) :



حيث يمكن القيام بالفعاليات التالية :

- عرض أو إزالة المحور الأفقي Display Axis Line (لإزالة المحور الأفقي يتوجب قبلها إزالة الإطار الداخلي بالأمر Chart ثم Inner Frame) .
- تغيير عنوان المحور بتظليل العنوان القديم في المستطيل المجاور لـ Axis Title ثم كتابة العنوان الجديد .
- موقع العنوان Title justification
- إظهار أو إزالة التقسيمات أو خطوط الشبكة .
- عند نقر زر Labels يظهر صندوق حوار Labels وفيه يمكن عرض عناوين المحور الأفقي جميعها أو جزء منها ، تغيير العناوين الموجودة ، التحكم بزاوية الميلان أو موقع العناوين Orientation حيث تتوفر الخيارات التالية: Horizontal ، Vertical ، Diagonal ، Staggered . فعند اختيار Diagonal لعرض عناوين المحور الأفقي للمخطط 1 فسيظهر المخطط كما يلي :

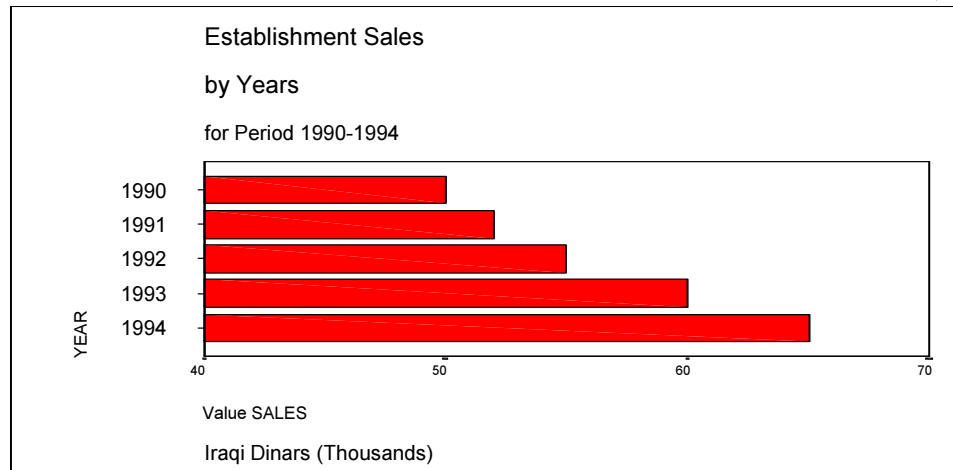
مخطط رقم 9



يكون من الضروري عرض العناوين الطويلة بهذه الطريقة تجنباً لتشابك العناوين فيما بينها كما يمكن الاستفادة من الخيارين Staggered و Vertical في هذه الحالة .

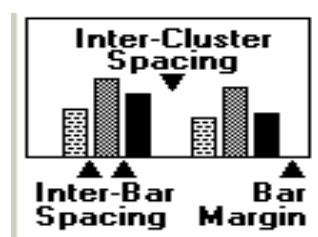
6. قلب المحاور Swap Axis : لقلب المحاور رأي تحويل المخطط من أعمدة الى أشرطة وبالعكس من شريط القوائم نختار Format → Swap Axis أو نقر الأيقونة  في شريط الأدوات فمثلاً عند قلب محاور المخطط رقم 1 فإنه يظهر كما يلي :

مخطط رقم 10



ملاحظات :

- يمكن تغيير نوع الخط Font وحجمه Size لأي عنوان في المخطط (عندما يكون في شاشة SPSS Chart Editor) بنقر العنوان ثم اختيار Text → Format من شريط القوائم أو نقر الأيقونة .
- في حالة عدم ظهور الكتابة باللغة العربية أنقر المخطط مرتين للتحويل الى شاشة SPSS Chart Editor ثم أنقر الأيقونة  بعد نقر النص المطلوب بزر الماوس الأيسر ومن قائمة Font اختر نوع الخط (Arial(Arabic) أو Andalus أو Arabic Transparent).
- يمكن تغيير نمط الإملاء للأشربة بنقر أشربة المخطط (عندما يكون في شاشة SPSS Chart Editor) ثم اختيار Fill Pattern  من شريط القوائم أو نقر الأيقونة .
- يمكن إضافة وسيلة إيضاح Legend باختبار Legend → Chart من شريط القوائم في شاشة Chart Editor ثم تأشير Display Chart Legend في صندوق حوار Legend.
- يمكن التحكم بالمسافات بين الأعمدة وبين كل من العمود الأول والأخير وبين الإطار الداخلي للمخطط Inner Frame وكذلك بين العناقيد Cluster باختبار Bar Spacing → Chart من شريط الأدوات كما في الشكل التالي :



الجدير بالذكر أنه يمكن التحكم بعرض الأشربة بواسطة هذا الأمر.

(2-13) **عمل قالب لمخطط بياني Chart Template :** من الواضح أننا ننفذ عمليات مختلفة على المخطط الى أن يأخذ شكله النهائي وفي بعض الأحيان يتطلب الأمر تنفيذ نفس العمليات ولكن على سلسلة من المتغيرات وهذه تعتبر عملية مطولة خاصة إذا كان هناك عدد كبير من المتغيرات وفي هذه الحالة يكون بالإمكان عمل قالب Template يتضمن كافة العمليات التي أجريت على المخطط الأول حيث يمكن تطبيقه على متغيرات أخرى. فمثلاً المخطط رقم 4 يتضمن عمليات مختلفة (ثلاثي الأبعاد ، عناوين الأعمدة بإطار ، العناوين الأصلية والفرعية ...). فلغرض تطبيق نفس هذه العمليات على متغير آخر x (الذي يأخذ القيم التالية 80،88،110،90،77 للسنوات 90-94) نتبع الخطوات التالية :

➤ **خزن المخطط 4 بصيغة قالب بالأمر Save Chart Template** → File في شاشة Chart

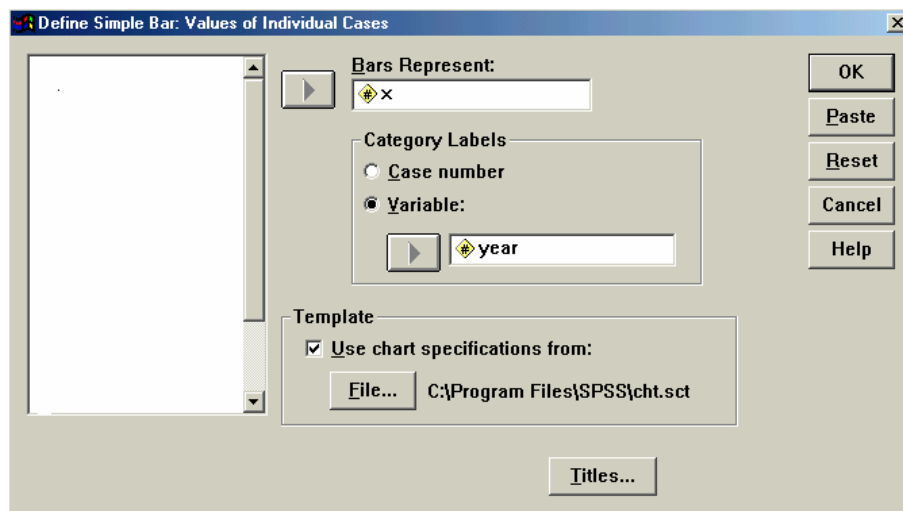
Editor حيث يأخذ القالب الاستطالة sct وليكن أسم القالب cht .

➤ **من شريط القوائم اختر Bar** → Graphs ثم اختر Values of Individual Cases

(simple) ثم أدخل المتغيرات x و year في صندوق حوار Values of Individual Cases أدناه .

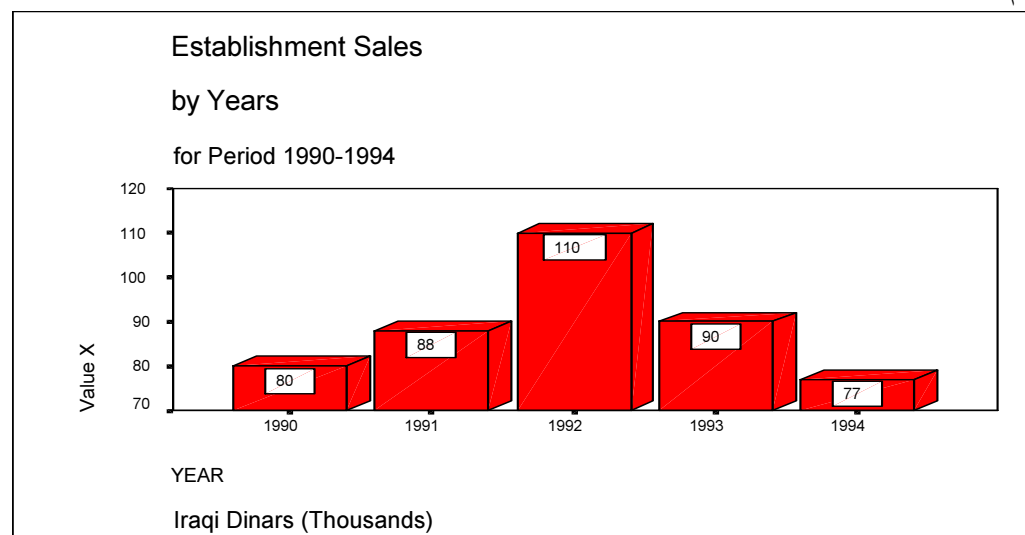
➤ **في خانة template قم بتأشير Chart Specification from** ثم انقر زر File لفتح القالب

cht حيث يتم ترتيب صندوق حوار Values of Individual Cases كما يلي :



➤ عند نقر زر Ok نحصل على المخطط التالي للمتغير x :

مخطط رقم 11

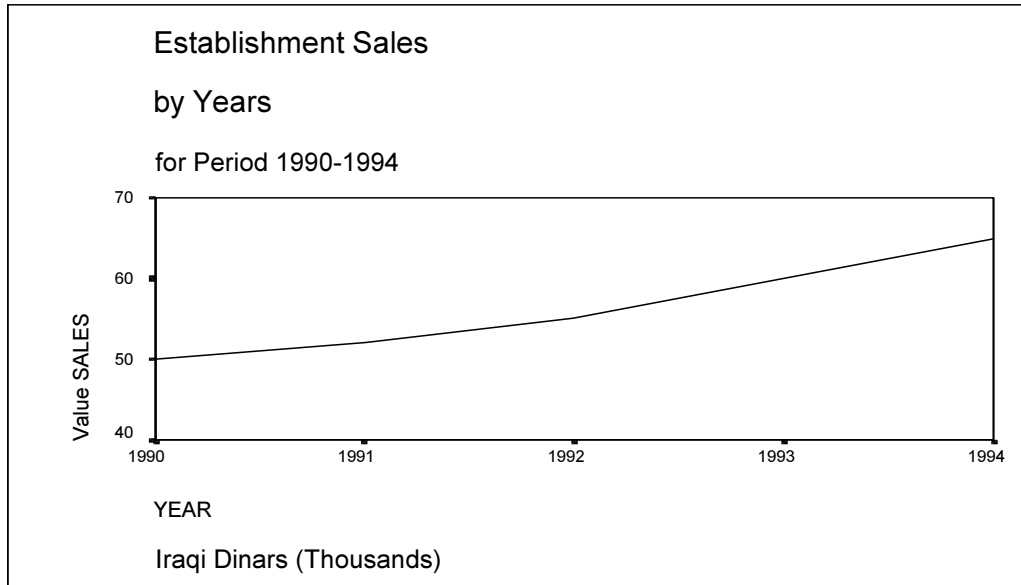


(13-3) تغيير مخطط Bar الى مخطط Line

لتغيير المخطط رقم 1 من أعمدة Bar الى خطوط Line نتبع الخطوات التالية :

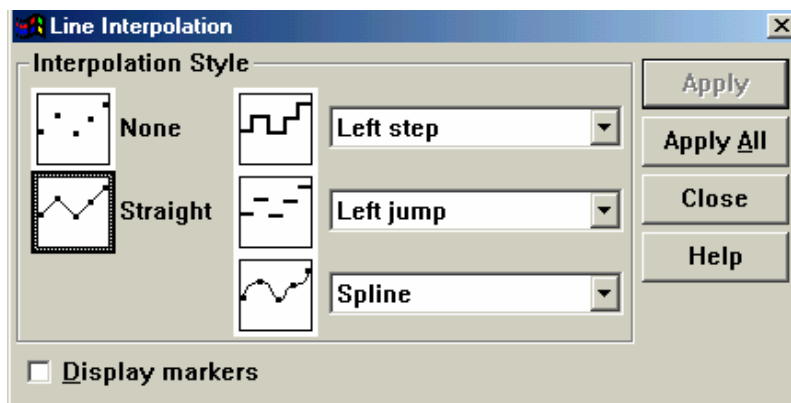
- انقر المخطط مرتين ليكون في شاشة Chart Editor .
- من شريط القوائم اختر Simple → Line → Gallery فيظهر صندوق حوار Line Charts (اخترنا Simple لان المخطط يتضمن سلسلة واحدة فقط Sales في حالة وجود أكثر من سلسلة (أو متغير) ونريد تضمينها جميعها نختار Multiple) .
- انقر Replace فيتحول المخطط الى Line كما يلي :

مخطط رقم 12

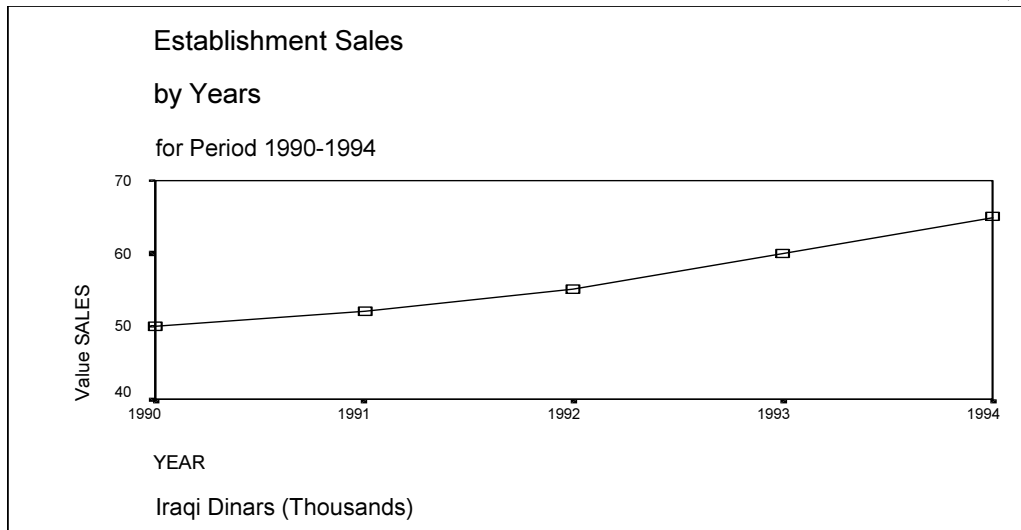


لعرض العلامات Markers على الخط البياني نتبع الخطوات التالية :

- من القوائم اختر Interpolation → Format أو انقر الأيقونة  في شريط الأدوات (عندما يكون المخطط في شاشة Chart Editor) فيظهر صندوق حوار Interpolation كما يلي :

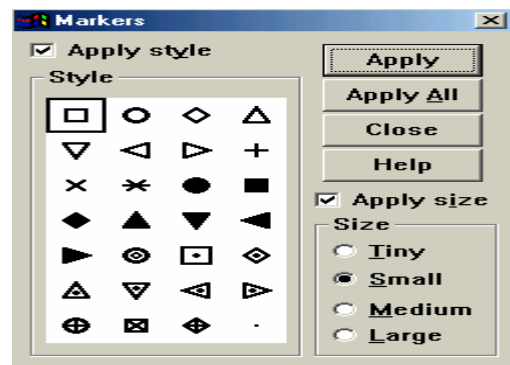


- انقر المربع الصغير Check Box المجاور لـ Display Markers لتأثيره ثم انقر زر Apply
- All فتضاف العلامات الى المخطط كما في الشكل التالي :



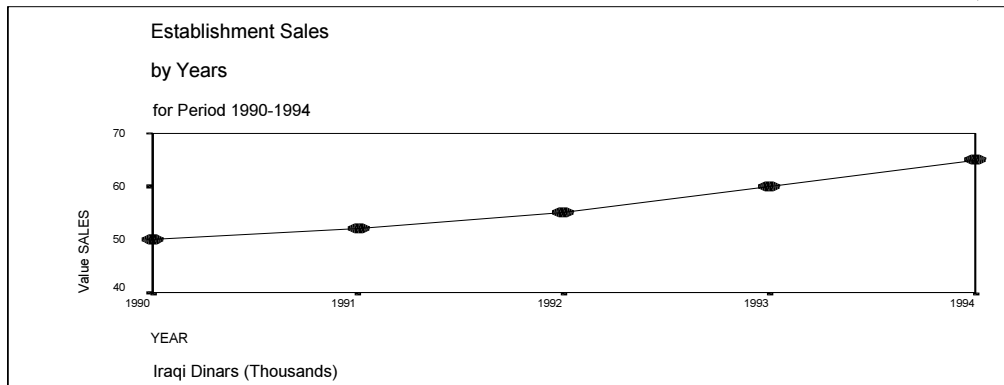
لتغيير شكل العلامة نتبع ما يلي :

- انقر المخطط مرتين ليكون في شاشة Chart Editor ثم انقر العلامات لتأشيرها .
- من القوائم اختر Markers → Format أو انقر الأيقونة  فيظهر صندوق حوار Markers التالي :



➤ اختر نوع العلامة وحجمها .

➤ انقر زر Apply فيتم تغيير العلامة كما يلي :



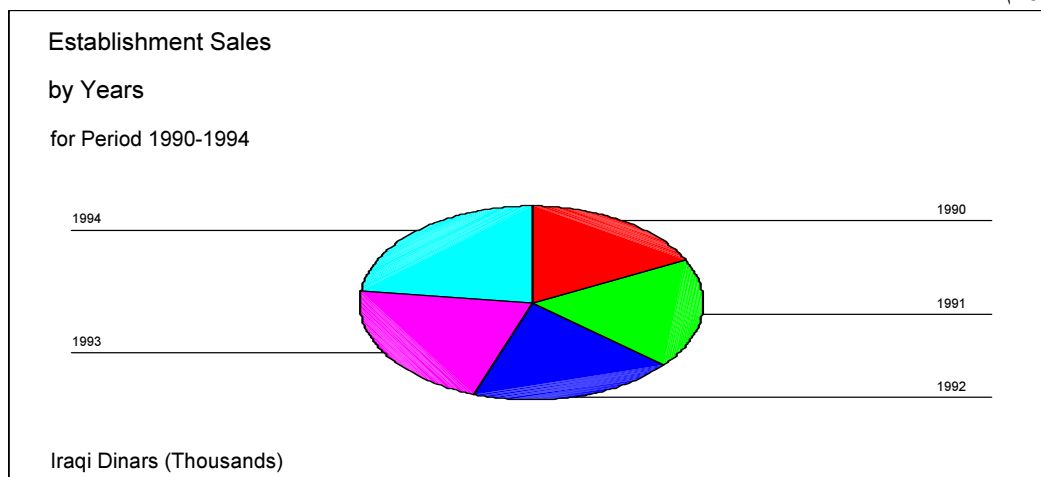
ملاحظة : في حالة وجود أكثر من سلسلة (متغير) ونريد تغيير شكل العلامة لها جميعاً انقر زر Apply All .

(13 - 4) تغيير مخطط Bar الى مخطط Pie

لتغيير المخطط رقم 1 من أعمدة Bar الى دائرة بيانية Pie نتبع الخطوات التالية :

- انقر المخطط مرتين ليكون في شاشة Chart Editor .
- من شريط القوائم اختر Pie → Gallery فيظهر صندوق حوار Pie Charts
- اختر Simple ثم انقر Replace فيتحول المخطط الى Pie كما يلي :

مخطط رقم 15

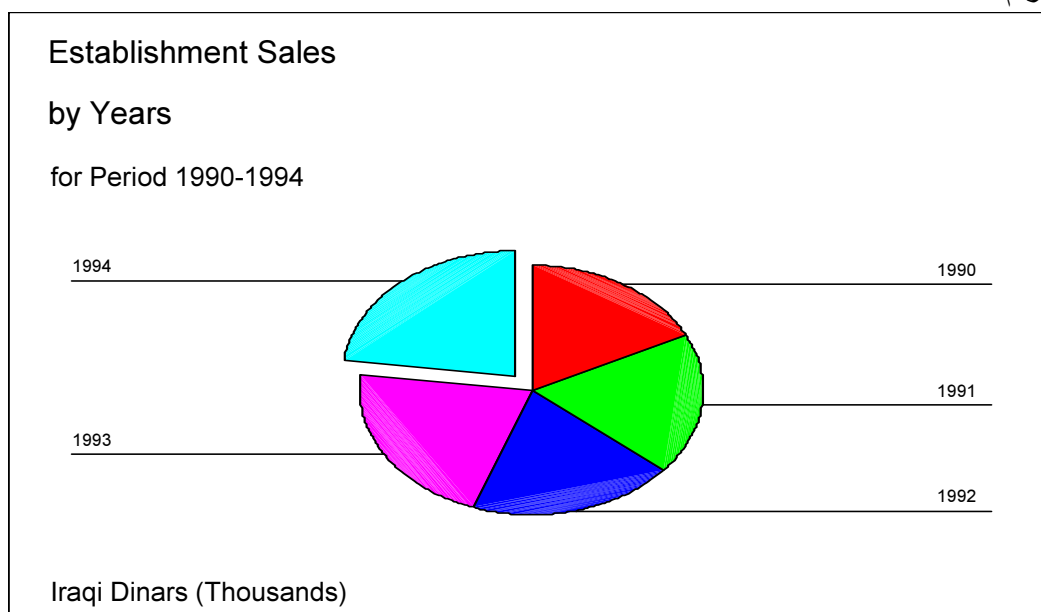


لفصل مقطع من الدائرة الخاص لسنة 1994 مثلاً نتبع الخطوات التالية :

- انقر المقطع الخاص بسنة 1994 لتحديده بعد أن تتأكد أن المخطط معروض في شاشة Chart Editor .

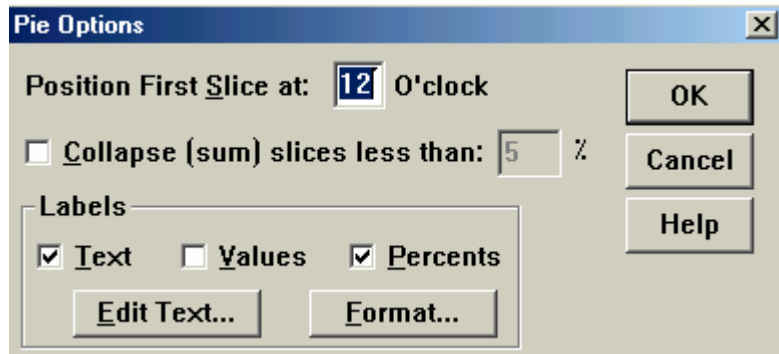
- من شريط القوائم اختر Explode Slice → Format أو انقر الأيقونة  فنحصل على المخطط التالي :

مخطط رقم 16

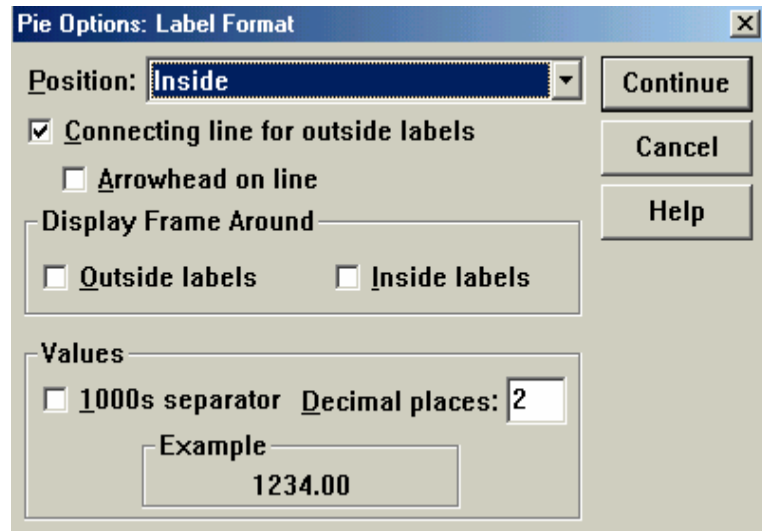


- لأعادة المقطع المفصول الى ما كان عليه أنقر المقطع بزر الماوس الأيسر ثم اختر Explode slice → Format أو أنقر الأيقونة  .

لعرض النسبة المئوية للمتغير في كل سنة في داخل كل مقطع Inside Slice نتبع الخطوات التالية :
☐ من القوائم أختَر Options → Chart (بعد التأكد أن المخطط في Chart Editor) فيظهر صندوق حوار Pie Options الذي نرتبه كما يلي :



لاحظ أننا أشرنا الخيار Percents في خانة Labels لعرض النسبة المئوية والسنة Text لكل مقطع .
 انقر زر Format فيظهر صندوق حوار Label Format الذي يرتب كما يلي :



حيث أن :

Position : يمثل موقع عرض النسبة المئوية وقيمة المتغير Value و Text (inside Outside Justified) ، Best fit، Outside، numbers inside Texts Outside) . وقد اخترنا Inside لأننا نرغب في عرض النسبة المئوية والسنة داخل المقطع .

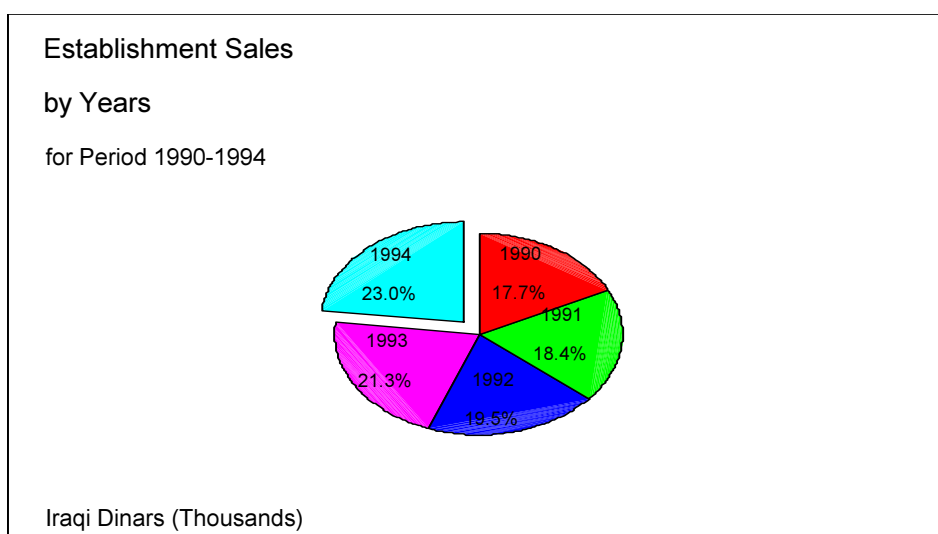
Connecting Line for Outside Labels : لعرض خطوط تربط بين المقاطع والعناوين الخارجية فقط .

Arrowhead on Line : استعمال خطوط سهمية الرأس .

Display Frame Around : لعرض إطار حول العناوين الداخلية أو الخارجية .

Values : لإضافة فواصل للأعداد التي تزيد عن 1000 وتحديد عدد المراتب العشرية للقيم فقط.

☐ عند نقر زر Continue ثم زر Ok نحصل على المخطط التالي :



مثال 2 :

الجدول التالي يمثل صادرات وأستيرادات بلد ما للسنوات 1998-2002 .

year	Exports	Impots
1998	190	180
1999	220	200
2000	219	215
2001	245	230
2002	250	244

يطلب ما يلي :

1. أعداد مخطط الأعمدة (الاشريطة) القطاعية Clustered Bars .

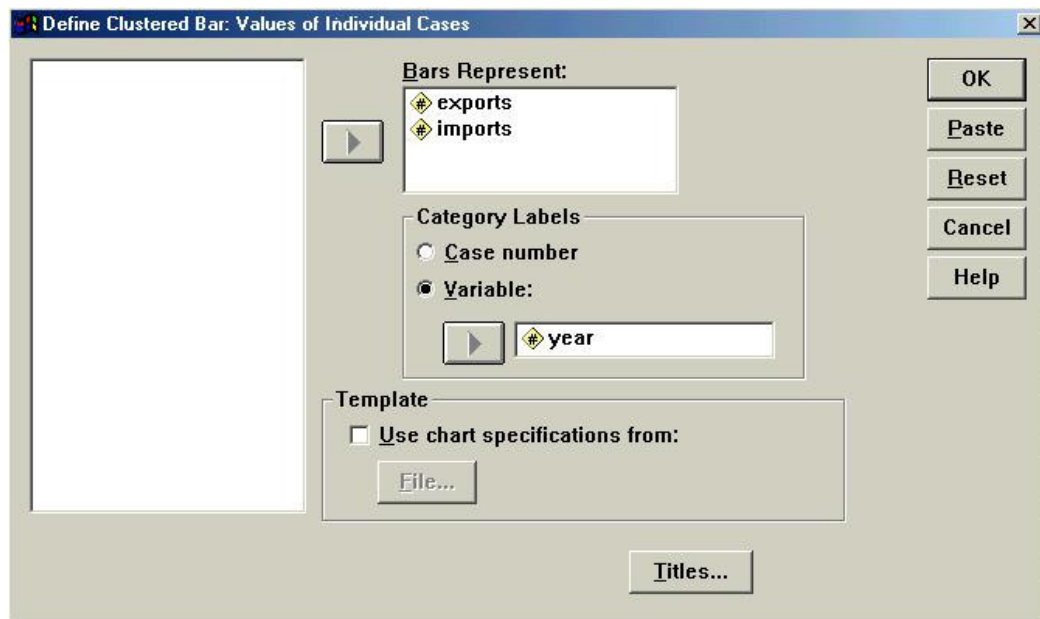
2. أعداد مخطط الأعمدة التراكمية Stacked Bars .

1. لأعداد الأعمدة القطاعية تتبع الخطوات التالية :

□ من شريط القوائم اختر Bar → Graph فيظهر صندوق Bar Chart ومنه اختر

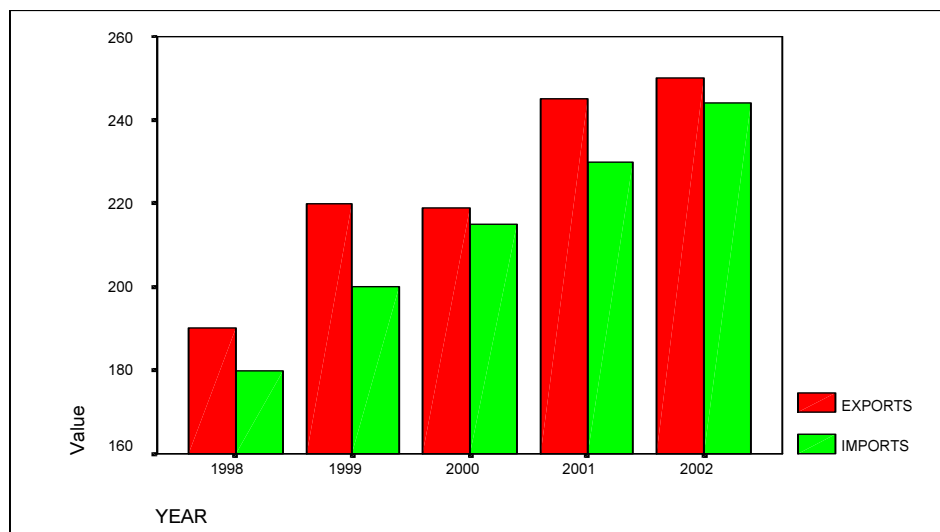
Values of Individual Cases / Clustered ثم انقر زر Define فيظهر صندوق حوار

Clustered Bar / Values of Individual Cases الذي نرتبه كما يلي :

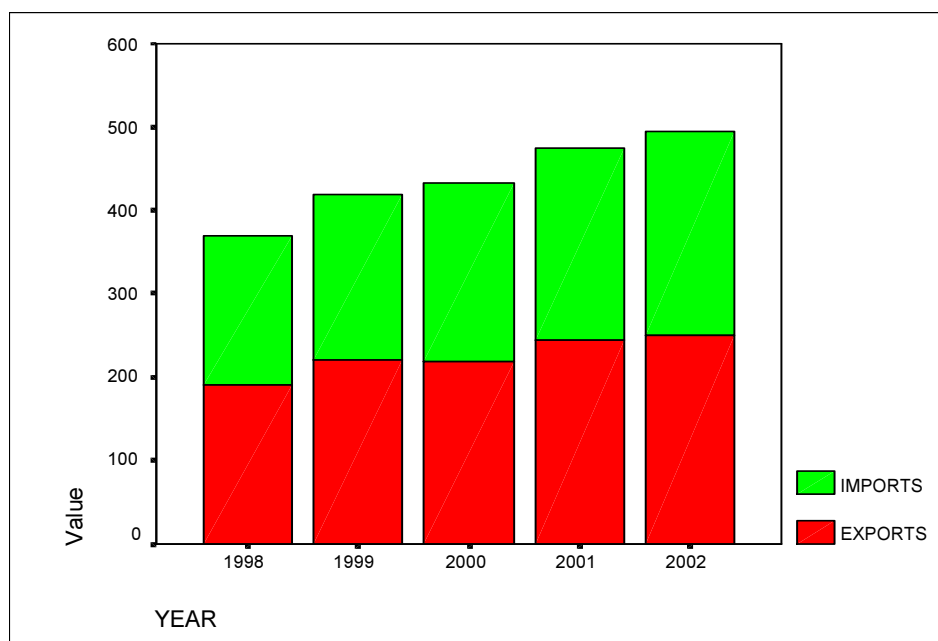


□ عند نقر زر Ok نحصل على المخطط التالي :

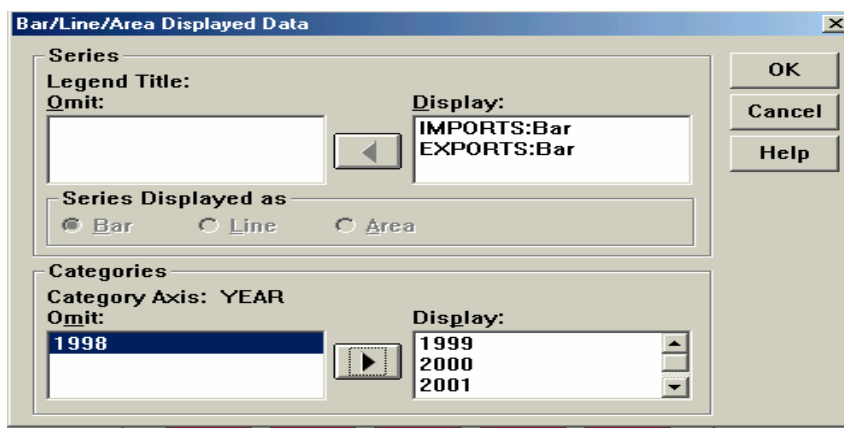
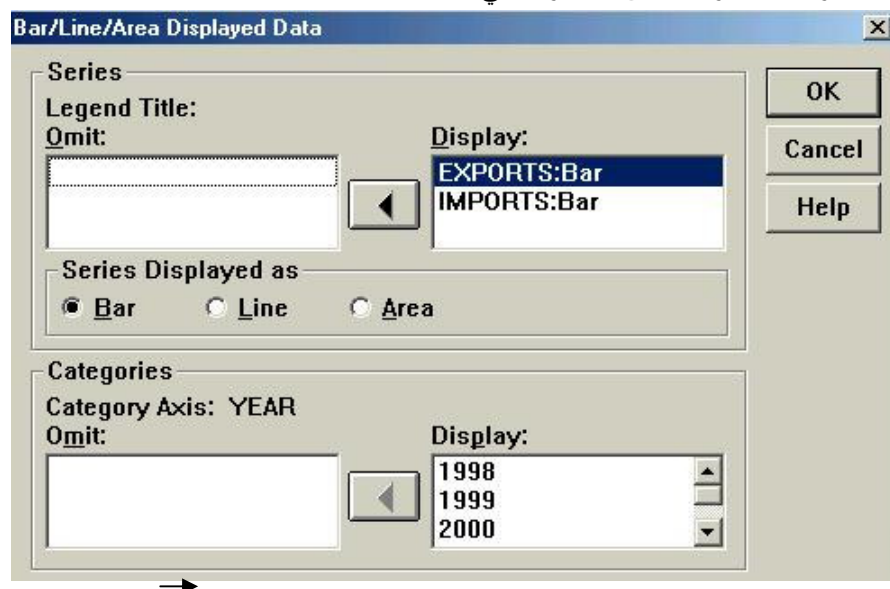
مخطط رقم 18



2. لتنفيذ مخطط Stacked Bar نتبع الخطوات السابقة نفسها عدا الخطوة رقم 2 حيث نختار Stacked بدل Clustered في صندوق حوار Bar Charts حيث نحصل على المخطط التالي :



بفرض أننا نريد جعل السلسلة Imports تعرض أولاً (في الأسفل) و السلسلة Exports تعرض ثانياً (في الأعلى) بالإضافة الى ذلك نرغب في إزالة فئة 1998 من المخطط لتنفيذ ذلك انقر المخطط مرتين ليكون في شاشة Chart Editor ثم اختر من شريط القوائم Series → Displayed أو انقر أي من السلسلتين مرتين فيظهر صندوق الحوار التالي :

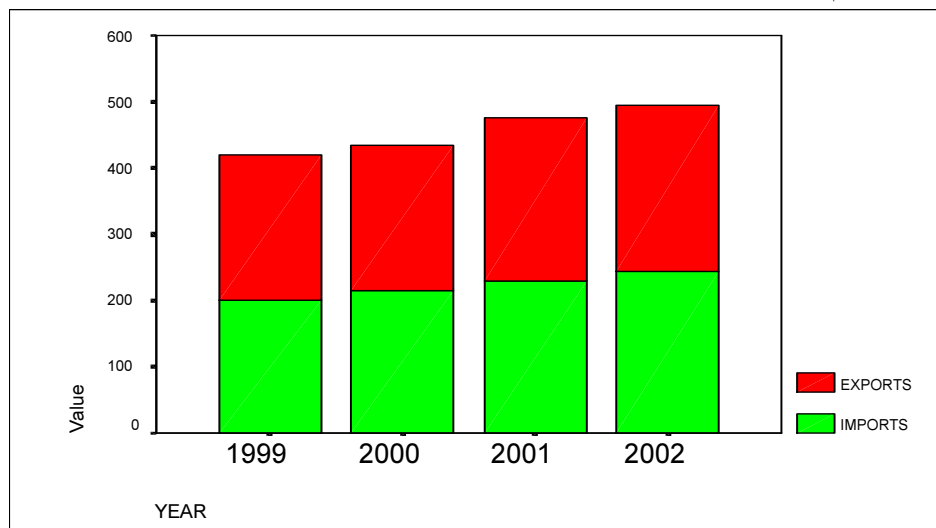


- أنقر السلسلة Exports في خانة Display (في الأعلى) لتحديد لها (إذا لم تحدد مسبقاً) ثم أنقر الزر  لنقلها الى خانة Omit (السلاسل المحذوفة من المخطط) في اليسار ثم أنقر الزر  لإرجاع السلسلة Exports من خانة Omit الى خانة Display (السلاسل المعروضة) وبهذا نكون قد غيرنا ترتيب السلاسل في خانة Display .
- أنقر العنوان 1998 في خانة Display (في الأسفل) لتحديده ثم أنقر الزر  لنقله الى خانة الحذف Omit .

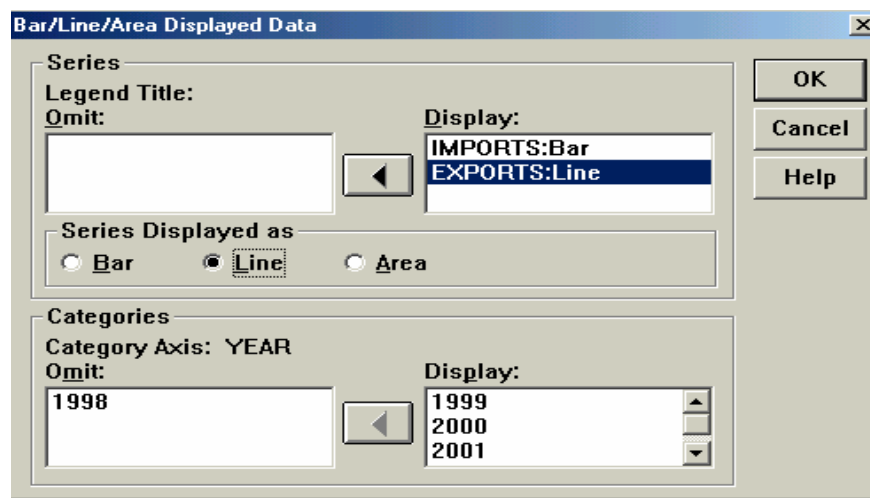
يظهر صندوق الحوار بعد التحوير كما يلي :

- عند نقر زر OK يظهر يعرض المخطط التالي :

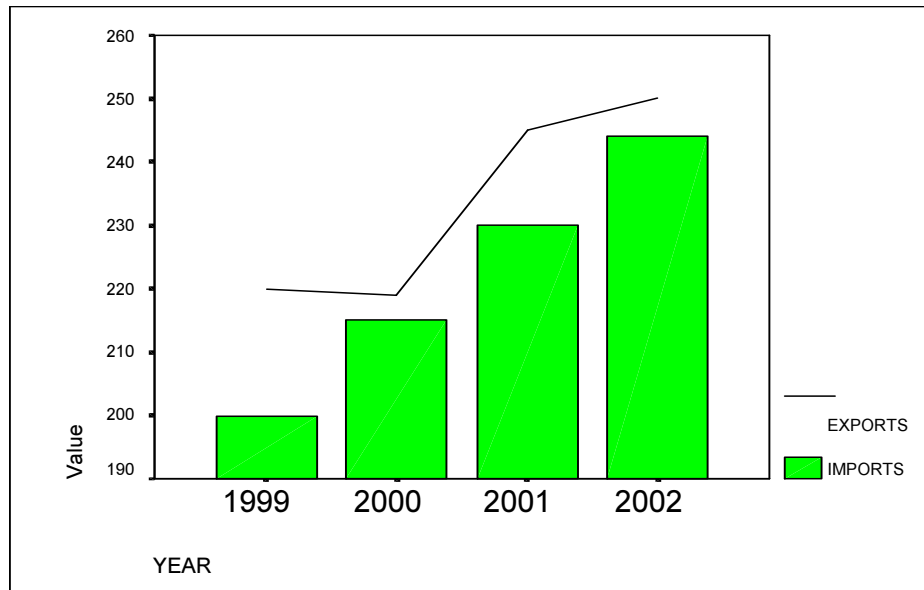
مخطط رقم 20



لعرض السلسلة Exports على شكل خط Line والسلسلة Imports على شكل أعمدة Bar نقوم بتحويل صندوق حوار Bar/Line/Area Displayed Data كما يلي :



لاحظ أننا أختارنا Bar للسلسلة Imports و Line للسلسلة Exports وعند نقر زر OK يعرض المخطط التالي :



ملاحظة :

لقد حصلنا على مخططي Line و Pie من مخطط Bar وبطبيعة الحال يمكن أن نحصل على هذه المخططات بصورة مستقلة كما حصلنا على مخطط Bar وكما يلي :

Line	Graph	Line	→
Pie	Graph	Pie	→
Area	Graph	Area	→

مثال 3 (على الخيار Summaries of Separate Variables)

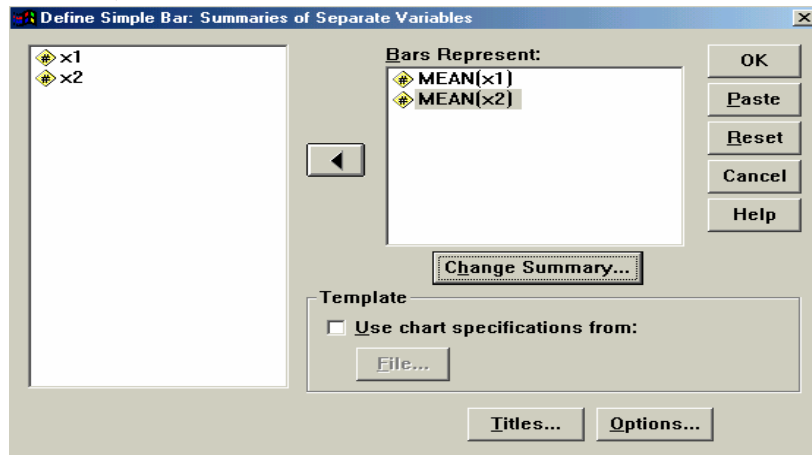
الجدول التالي يمثل المتغيرين x1 و x2 والمتغير g وقد أدخلت البيانات في شاشة Data Editor لبرنامج SPSS كما يلي :

x1	x2	g
100	10	a
200	20	b
300	30	a
400	40	b
500	50	b

يطلب مايلي :

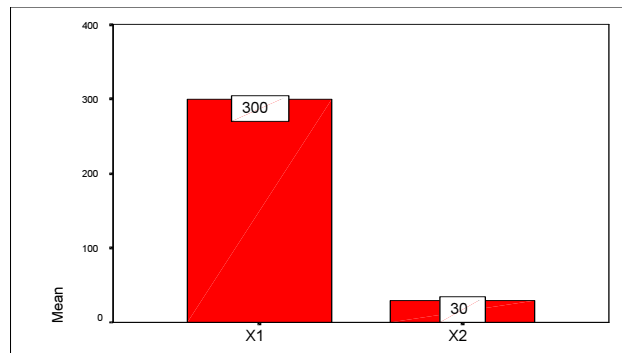
1. أعداد مخطط Bar لمتوسطي المتغيرين x1 و x2 .
 2. أعداد مخطط قطاعي Clustered Bar لمتوسطات فئتي (a و b) للمتغيرين x1 و x2 .
1. لتنفيذ مخطط Bar لمتوسطي المتغيرين x1 و x2 نتبع الخطوات التالية:
- ☐ من شريط القوائم نختار Bar → Graph فيظهر صندوق حوار Bar Charts ومنه نختار Summaries of Separate Variables / Simple .
 - ☐ عند نقر زر Define في هذا الصندوق الأخير يظهر صندوق حوار

Define Simple Bar : Summaries of Separate Variables الذي نرتبه كما يلي :



نلاحظ أن الأعمدة ستمثل متوسط كل من المتغيرين x_1 و x_2 حيث أن المتوسط هو الاختيار الافتراضي لنوع الدالة ويمكن تغيير المتوسط بنقر الزر Change Summary الذي يتيح اختيار مؤشرات أخرى مثل ... Median, Mode, No. of Cases, Min. Value, Max. Value .

□ عند نقر زر OK نحصل على المخرج التالي :



مخطط رقم 22

أن ارتفاع العمود x_1 هو الوسط الحسابي للمتغير ويساوي 300 وأن ارتفاع العمود x_2 هو الوسط الحسابي للمتغير ويساوي 30 .

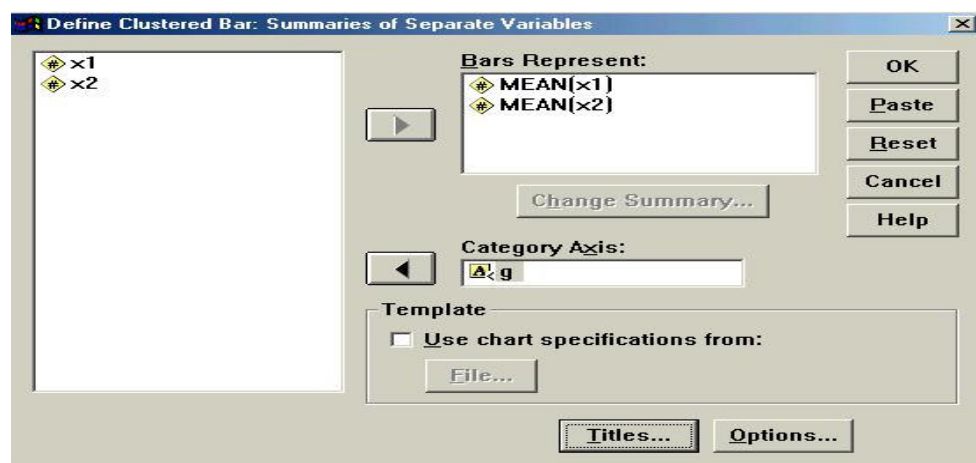
2. لتنفيذ المخطط القطاعي Clustered Bar لمتوسطات فئتي (a و b) للمتغيرين x_1 و x_2 نتبع

الخطوات التالية :

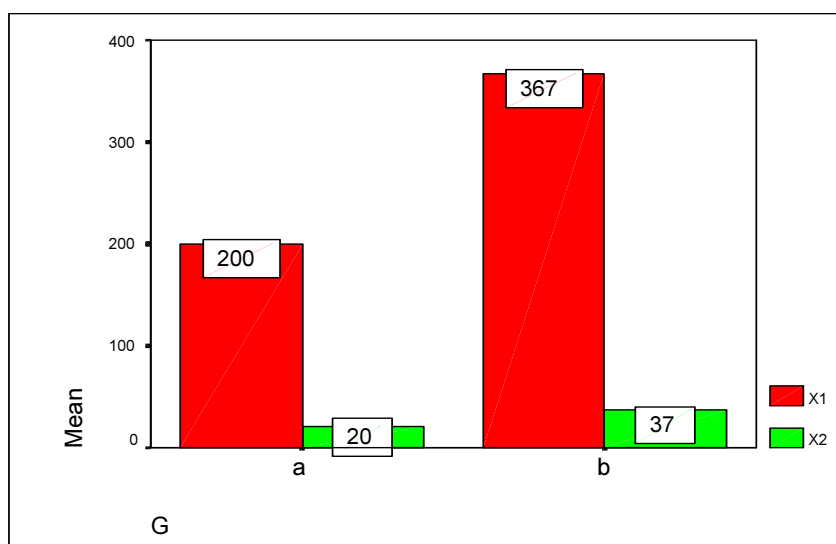
من شريط القوائم نختار Bar → Graph فيظهر صندوق حوار Bar Charts ومنه نختار Summaries of Separate Variables / Clustered .

□ عند نقر زر Define في هذا الصندوق الأخير يظهر صندوق حوار

Define Clustered Bar : Summaries of Separate Variables الذي نرتبه كما يلي :



□ عند نقر زر Ok نحصل على المخطط التالي :
مخطط رقم 23



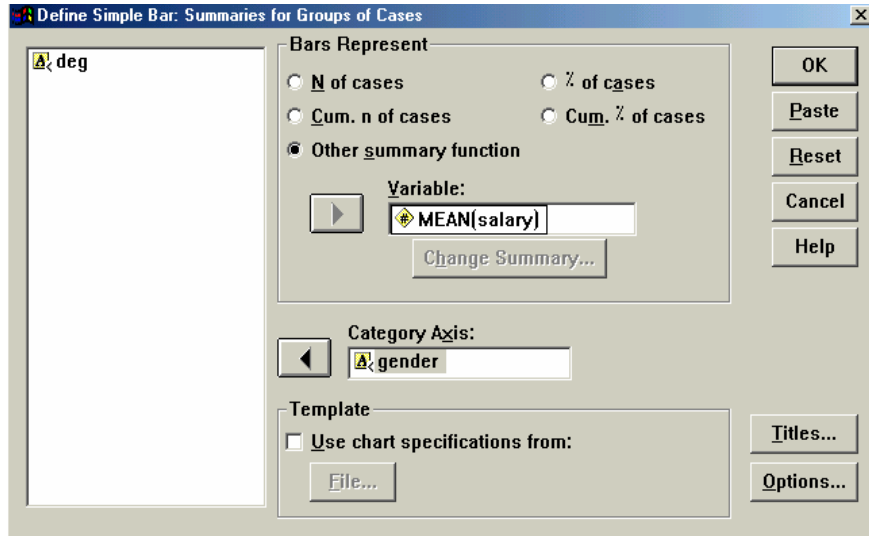
مثال 4 : (على الخيار Summaries of Groups of Cases)

الجدول التالي يمثل رواتب salary مجموعة من الموظفين حسب الدرجة الوظيفية deg والجنس gender .

deg	gender	salary
First	Male	90
First	Female	70
Third	Male	56
Second	Male	65
First	Male	85
Second	Female	60
Second	Male	69
Third	Male	57
Third	Female	50
First	Female	75
Second	Female	62
Third	Female	51
First	Male	85

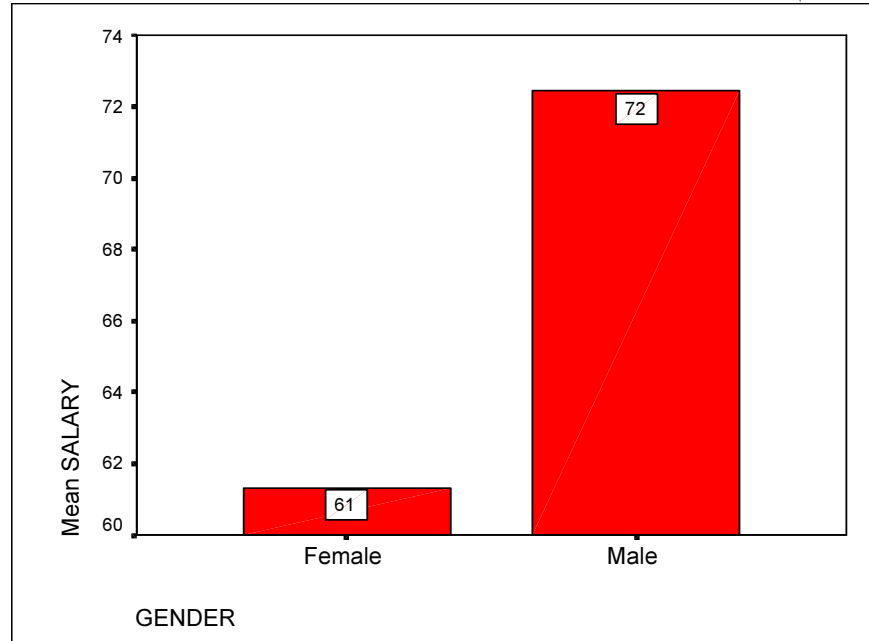
المطلوب مايلي :

1. أعداد مخطط بالأعمدة البيانية Bars يمثل متوسط الراتب لفئتي الذكور والإناث .
 2. أعداد مخطط بالأعمدة البيانية Bars يمثل متوسط الراتب لكل الذكور والإناث حسب كل درجة وظيفية .
1. لتنفيذ المطلوب الأول نتبع الخطوات التالية:
- ☐ من شريط القوائم نختار Bar → Graphs فيظهر صندوق حوار Bar Charts ومنه نختار Summaries for Group of Cases .
- ☐ عند نقر زر Define في صندوق الحوار أعلاه يظهر صندوق حوار Define Simple Bar : Summaries for Group of Cases الذي يرتب كما يلي :



☐ عند نقر زر OK نحصل على المخطط التالي :

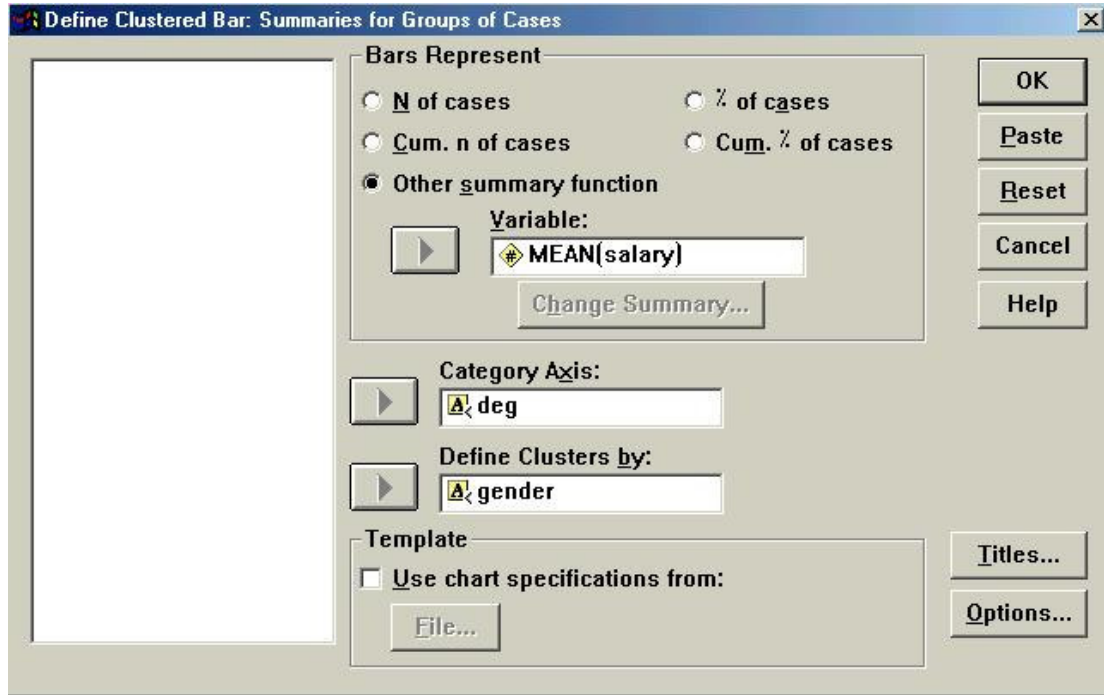
مخطط رقم 24



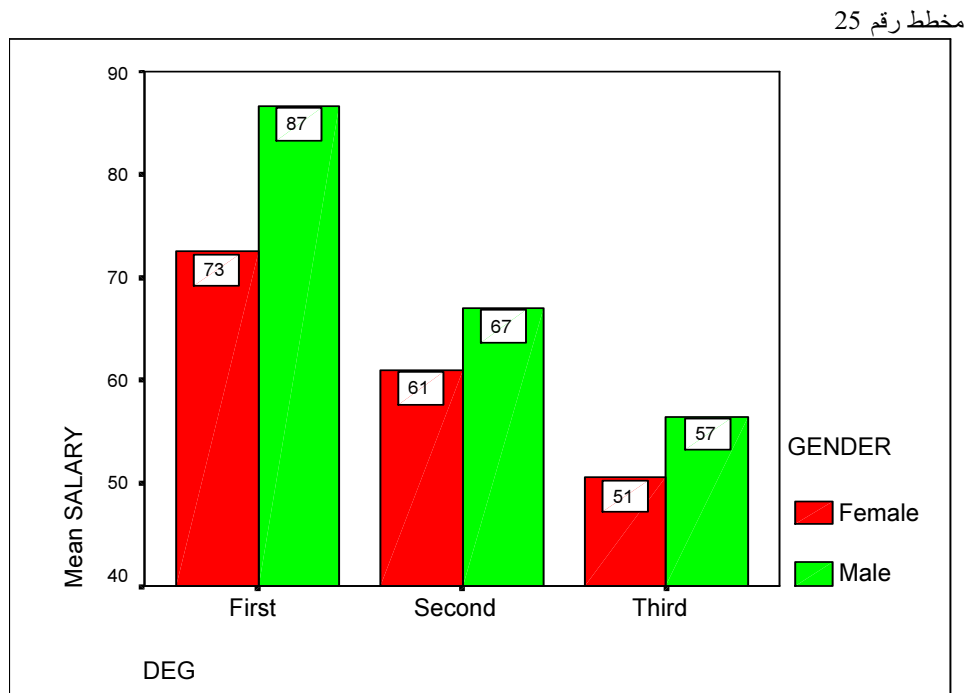
2. لتنفيذ المطلوب الثاني نتبع الخطوات التالية :

من شريط القوائم نختار **Bar Charts** → **Graphs** فيظهر صندوق حوار **Bar Charts** ومنه نختار **Summaries for Group of Cases/Clustered** حيث أننا سنمثل الدرجة الوظيفية على المحور الأفقي يقابل كل درجة وظيفية عمودين أحدهما لمتوسط رواتب الذكور والآخر لمتوسط رواتب الإناث **Clusters**. عند نقر زر **Define** في صندوق الحوار أعلاه يظهر صندوق حوار

Define Clustered Bar: Summaries for Group of Cases الذي يرتب كما يلي :



□ عند نقر زر **OK** نحصل على المخطط التالي :



من خلال المثالين الأخيرين نلاحظ الميزة المهمة التي تتمتع بها مخططات برنامج SPSS والتي لا تتوفر في مخططات تطبيقات أخرى مثل EXCEL وهي المعالجة الإحصائية للبيانات حيث يمكن استخراج المؤشرات الإحصائية كالمتوسطات الحسابية مثلاً لفئات متغير معين وتمثيلها بيانياً وهذا يعد اختصاراً للجهود والوقت اللازمين لترتيب البيانات وخاصة إذا كان حجم البيانات كبيراً كما في استمارات الاستبيانات الإحصائية .

(5 - 13) المدرج التكراري Histogram

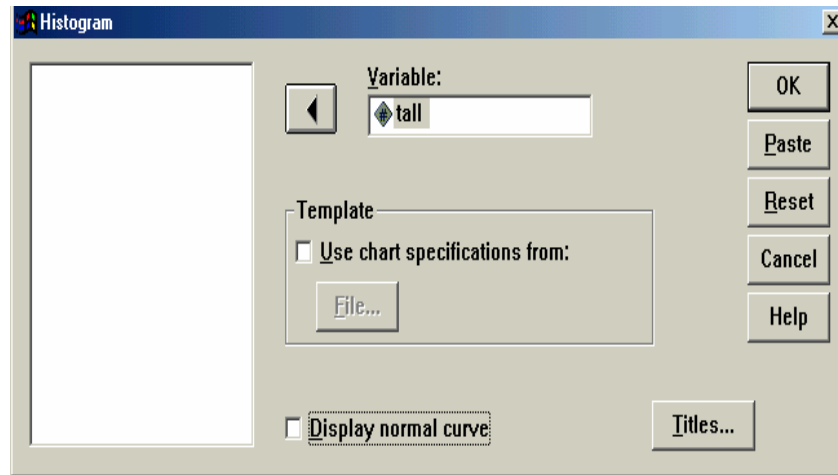
أن المخططات السابقة Bar و Line و pie تستعمل لعرض البيانات الخام Row Data أما المدرج التكراري فيستعمل لعرض البيانات المبوبة Grouped Data على الرغم من أننا نقوم بإدخال البيانات بصيغتها الخام Ungrouped الى برنامج SPSS حيث يقوم البرنامج بتصنيف البيانات الى جدول توزيع تكراري Frequency Table تلقائياً ومن ثم رسم المدرج التكراري لهذا الجدول حيث يمكن تغيير إعدادات المخطط (عدد الفئات ، طول الفئة ، ...) كما في المثال التالي .

مثال 5

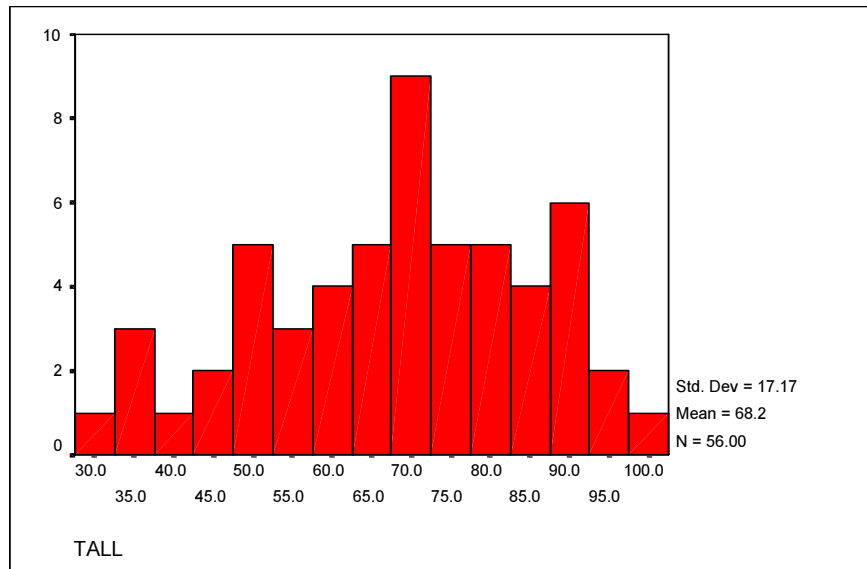
يطلب رسم المدرج التكراري للمتغير tall والذي أدخل مسبقاً في شاشة Data Editor (نفس المتغير الوارد في المثال التابع للبند (4-1) من الفصل الرابع) .

لرسم المدرج التكراري نتبع الخطوات التالية :

□ من شريط القوائم اختر Histogram → Graph فيظهر صندوق حوار Histogram الذي نرتبه كما يلي :

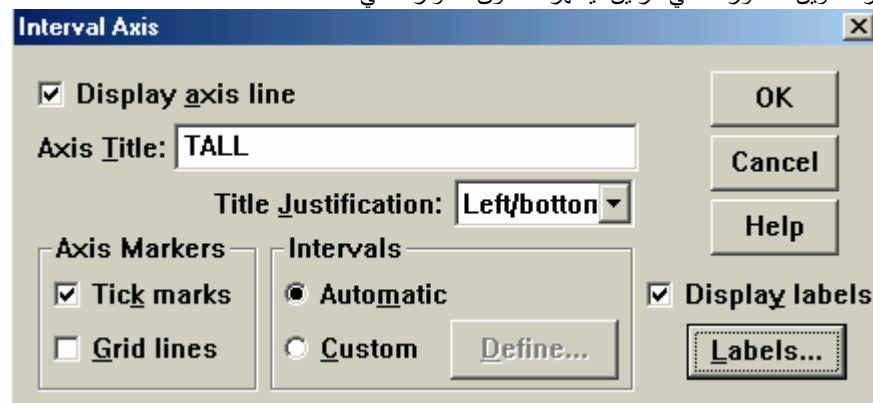


□ عند نقر زر OK يتم عرض المدرج التكراري التالي :
مخطط رقم 26



حيث نلاحظ أنه تم تمثيل مراكز الفئات Midpoints على المحور الأفقي نرغب بتمثيل الفئات Classes على المحور الأفقي عوضاً عن مراكز الفئات كما نود أيضاً تحديد طول الفئة Interval Width بـ 10. لتنفيذ ذلك نتبع الخطوات التالية:

- انقر مخطط Histogram مرتين لينتقل الى شاشة Chart Editor .
- انقر عناوين المحور الأفقي مرتين فيظهر صندوق الحوار التالي :



أقر زر Custom في خانة Intervals فيتحفز الزر Define وعند نقره يظهر صندوق الحوار Define Custom Intervals الذي نرتبه كما يلي :

لقد حددنا طول الفئة بـ 10 في Interval width وجعلنا اقل قيمة معروضة 30 وأكبر قيمة معروضة 100 وعليه فأن عدد الفئات يساوي 7 الذي يحتسب من العلاقة التالية :

$$No. of Classes = \frac{Range}{Class (Interval) Width} = \frac{70}{10} = 7$$

• أنقر زر continue للرجوع الى صندوق حوار Interval axis ثم أنقر زر Labels في هذا الصندوق فيظهر صندوق حوار Labels الذي نرتبه كما يلي :

حيث قمنا بالعمليات التالية :

1. اخترنا Range بدلاً من Midpoint في خانة Type لأننا نرغب في عرض الفئات على المحور الأفقي بدلاً من مراكز الفئات .
2. حددنا عدد المراتب العشرية بصفر للحدود العليا والدنيا للفئات في خانة Decimal Places .
3. في خانة Orientation اخترنا Diagonal كي لا تتشابك حدود الفئات فيما بينها .

أقر زر Continue للرجوع الى صندوق حوار Interval axis ثم أنقر زر OK فيعرض المخطط التالي :

مخطط رقم 27

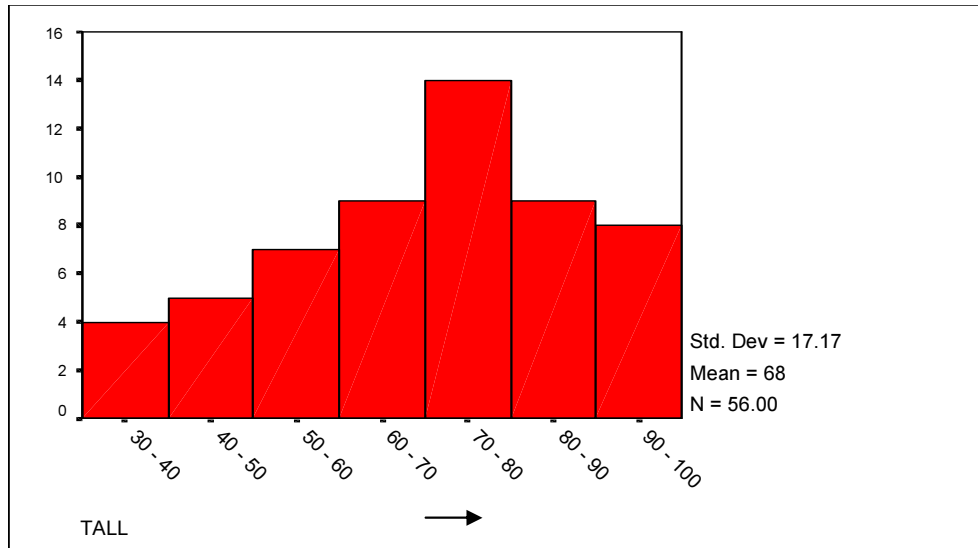


Chart Options عندما يكون المخطط في شاشة Chart Editor ثم انقر Normal curve في خانة Display أو يمكنك تأشير Display Normal Curve في صندوق حوار Histogram عند ترتيب هذا الصندوق في بداية العمل .

Box Plot (6 - 13) مخطط

يستعمل هذا المخطط لوصف توزيع مجموعة من المشاهدات حول الوسيط Median (راجع البند 6 - 1 لمزيد من التفاصيل حول مكونات المخطط) .

مثال 6 : الجدول التالي يتضمن مجموعة من المتغيرات التي أدخلت الى شاشة Data Editor في برنامج SPSS وكما يلي :

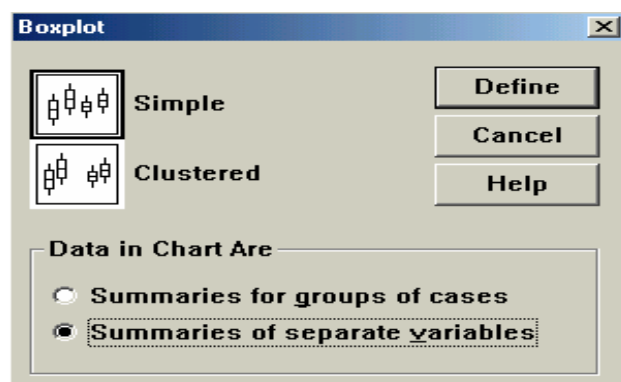
x	y	cat	z	g	gender
30	47	a	20	a	f
20	66	a	60	b	f
31	77	a	50	a	m
35	90	a	80	b	m
100	55	a	100	b	m
80	62	a	70	a	m
79	80	b	89	b	m
55	99	b	35	a	f
50	43	b	65	b	f
60	87	b	40	a	f
95	92	b	55	b	f
45	70	b	69	b	m
39	40	b	40	a	m

1. يطلب رسم مخطط Box Plots للمتغير x .
2. رسم مخطط Box plot للمتغيرين x و y حسب الفئتين a و b التي يحددها المتغير cat .
3. رسم مخطط Box Plot للمتغيرين x و y حسب الفئتين a و b التي يحددها المتغير cat .
4. رسم مخطط Box Plot لفئتي المتغير z (b و a) التي يحددها المتغير g .

5. رسم مخطط Box Plot لفئتي المتغير z (a و b) التي يحددها المتغير g علماً أن كل من الفئتين تتضمن بدورها فئتين إحداهما للذكور m والأخرى للإناث f .

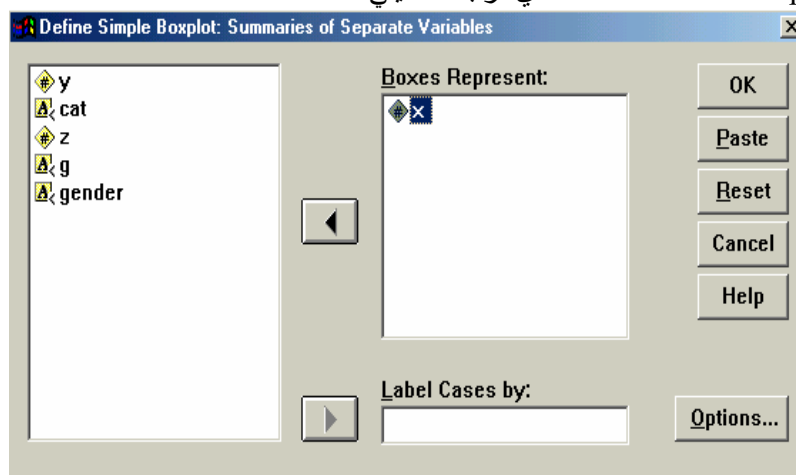
1. لتنفيذ المطلوب الأول نتبع مايلي :

من شريط القوائم اختر Box plot → Graphs فيظهر صندوق حوار Boxplot التالي :



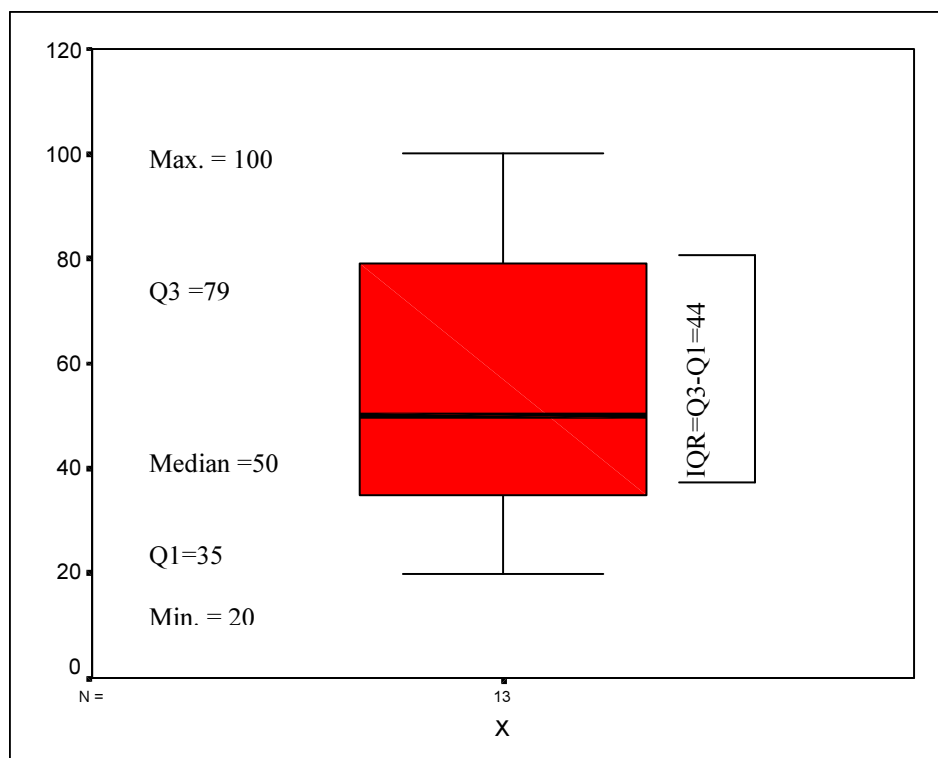
لاحظ أننا اخترنا summaries of Separate Variables مع الخيار Simple (حيث لايمكن اختيار Clustered لعدم وجود فئات للمتغير x) .

عند نقر زر Define في صندوق الحوار أعلاه يظهر صندوق حوار Define Simple boxplot: Summaries of Separate Variables الذي نرتبه كما يلي :



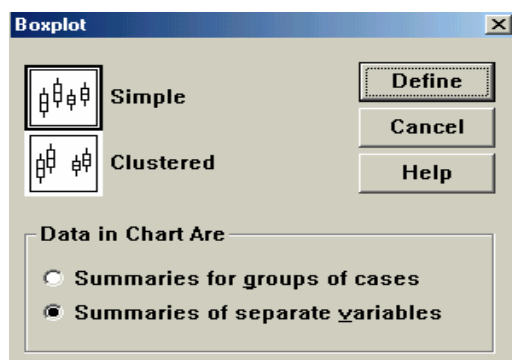
لقد تم إدخال أسم المتغير في Boxes Represent كما يتم إدخال متغير في خانة Label Cases by عند الرغبة في تحديد عنوان Label للقيم المتطرفة والشاذة (تظهر العناوين في المخطط) أما في حالة عدم إدخال متغير فيستعمل رقم الحالة التي تقع ضمنها القيمة المتطرفة أو الشاذة كعنوان لها .

عند نقر زر OK نحصل على المخطط التالي :



2. لتنفيذ المطلوب الثاني نتبع مايلي :

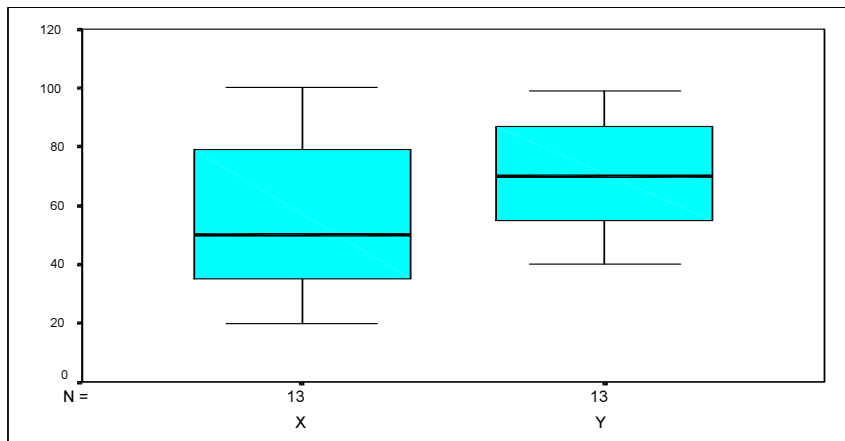
من شريط القوائم اختر Box plot → Graphs فيظهر صندوق حوار Boxplot الذي نرتبه كما يلي :



عند نقر زر Define يظهر صندوق حوار Define Simple Box plot : Summaries of

Separate Variables الذي نرتبه كالتالي :

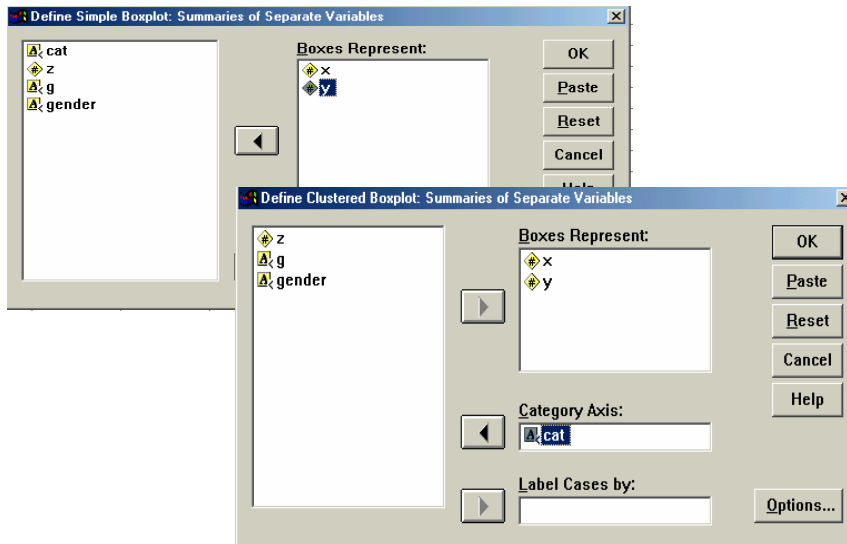
عند نقر زر OK يعرض المخطط التالي :



3. تنفيذ المطلوب الثالث نتبع الخطوات التالية :

Graphs

من القوائم نختار

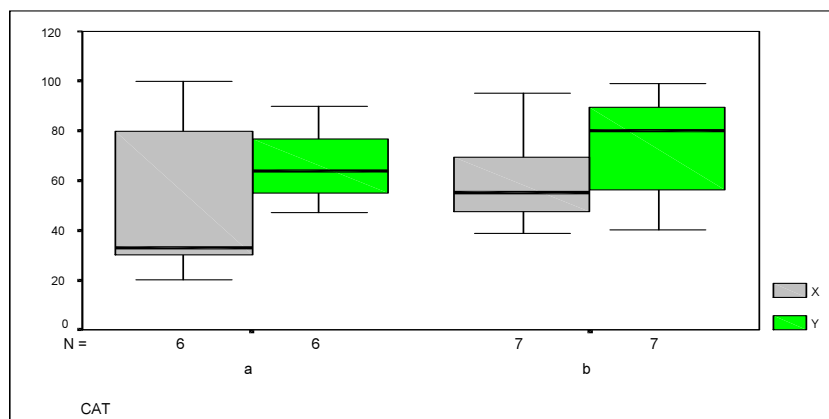


Boxplot فيظهر صندوق حوار Box Plot ومنه نختار Summaries of Separate Variables مع اختيار Clustered وعند نقر زر Defined في صندوق الحوار الأخير يظهر صندوق حوار Define Clustered Box Plot : Summaries of Separate Variables الذي نرتبه كما يلي :

لاحظ أن خانة Box Represent يجب أن تتضمن متغيرين في الأقل وأن Category Axis هو المتغير الذي يتم تمثيل فئاته على المحور الأفقي .

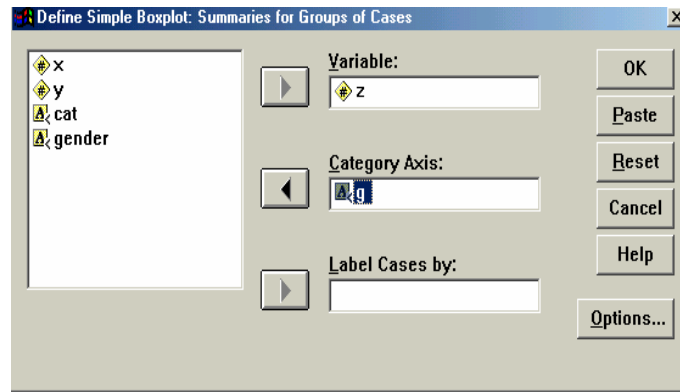
عند نقر زر OK يعرض المخطط التالي :

مخطط رقم 30



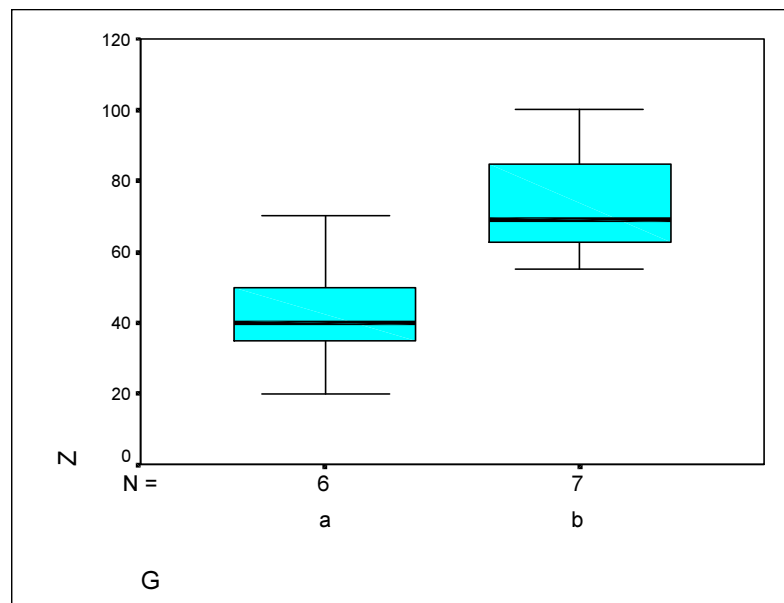
4. لتنفيذ المطلوب الرابع نتبع الخطوات التالية :

من القوائم نختار Boxplot → Graphs فيظهر صندوق حوار Box Plot ومنه نختار Summaries for Group of Cases مع اختيار Simple وعند نقر زر Defined في صندوق الحوار الأخير يظهر صندوق حوار Define Simple Box Plot : Summaries for Group of Cases الذي نرتبه كما يلي :



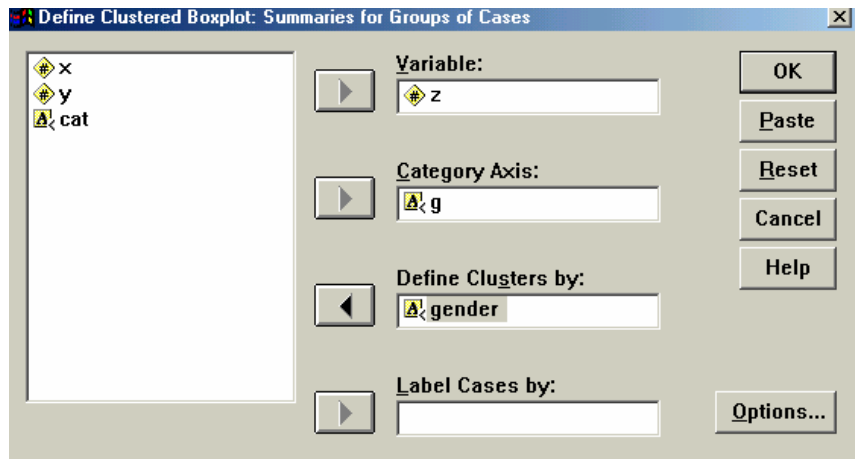
عند نقر زر OK يعرض المخطط التالي :

مخطط رقم 31



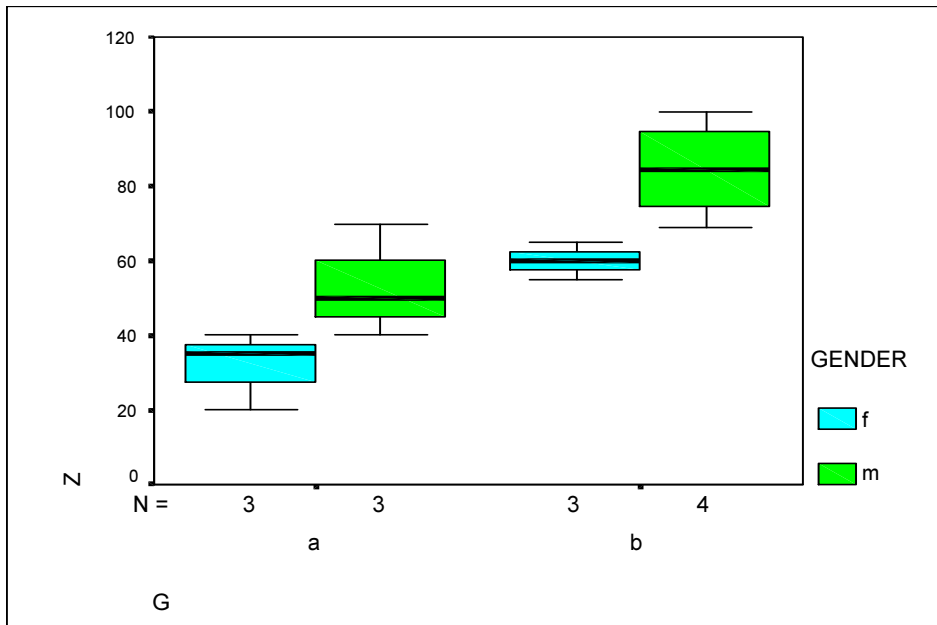
5. لتنفيذ المطلوب الخامس نتبع الخطوات التالية :

من القوائم نختار Boxplot → Graphs فيظهر صندوق حوار Box Plot ومنه نختار Summaries for Group of Cases مع اختيار Clustered وعند نقر زر Defined في صندوق الحوار الأخير يظهر صندوق حوار Define Clustered Box Plot : Summaries for Group of Cases الذي يتم ترتيبه كالتالي :



عند نقر زر OK يعرض المخطط التالي :

مخطط رقم 32



(7-13) مخطط شكل الانتشار Scatterplot

يستعمل هذا النوع من المخططات لتمثيل العلاقة بين متغيرين أو أكثر بمجموعة من النقاط المتناثرة ويتضمن أربعة أنواع لشكل الانتشار .

مثال 7 :

الجدول التالي يتضمن عدد من المتغيرات والتي أدخلت في ورقة Data Editor لبرنامج SPSS

وكما يلي :

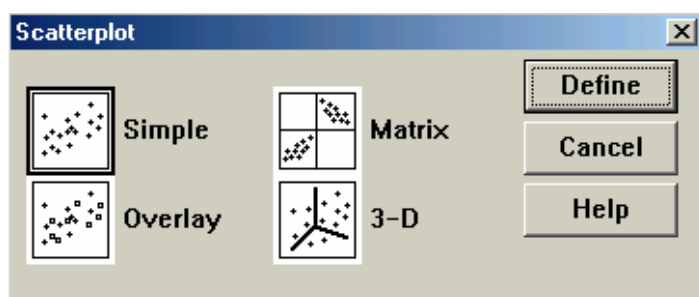
y	x1	x2	x3	marker	label
79	26	6	60	a	c1
74	29	15	52	a	c2
104	56	8	20	a	c3
87	31	8	47	a	c4
95	52	6	33	a	c5
109	55	9	22	a	c6
102	71	17	6	b	c7
72	31	22	44	b	c8
93	54	18	22	b	c9
115	47	4	26	b	c10
83	40	23	34	b	c11
113	66	9	12	b	c12
109	68	8	12	b	c13

سنحاول فيما يلي توضيح كيفية أعداد الأنواع الأربعة لشكل الانتشار :

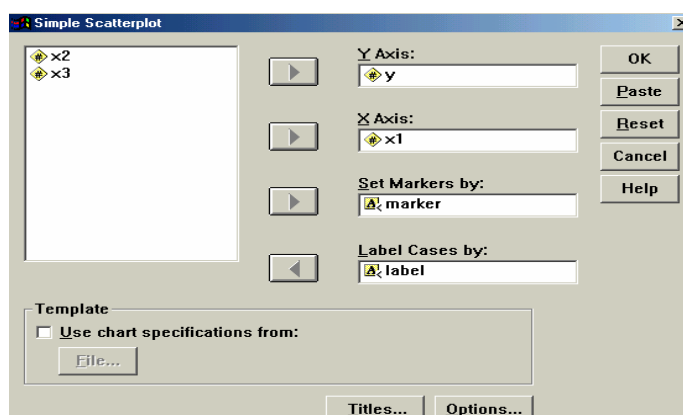
1. شكل الانتشار البسيط Simple

يستعمل لرسم العلاقة بين متغيرين فقط أحدهما يمثل المحور الأفقي للشكل والآخر المحور العمودي له وأن إحداثيات كل نقطة في الشكل تعتمد على هذين المتغيرين ، لأعداد شكل الانتشار البسيط للعلاقة بين y و x1 نتبع الخطوات التالية:

◀ من شريط القوائم اختر Scatter → Graphs فيظهر صندوق حوار Scatterplot التالي :



◀ في صندوق الحوار أعلاه أختَر Simple بنقره ثم أنقر Define فيظهر صندوق حوار Simple scatterplot الذي نرتبه كما يلي :



حيث أن :

Y Axis : تتضمن هذه الخانة المتغير الذي يمثل المحور الرأسي لشكل الانتشار .

X Axis : تتضمن هذه الخانة المتغير الذي يمثل المحور الأفقي لشكل الانتشار .

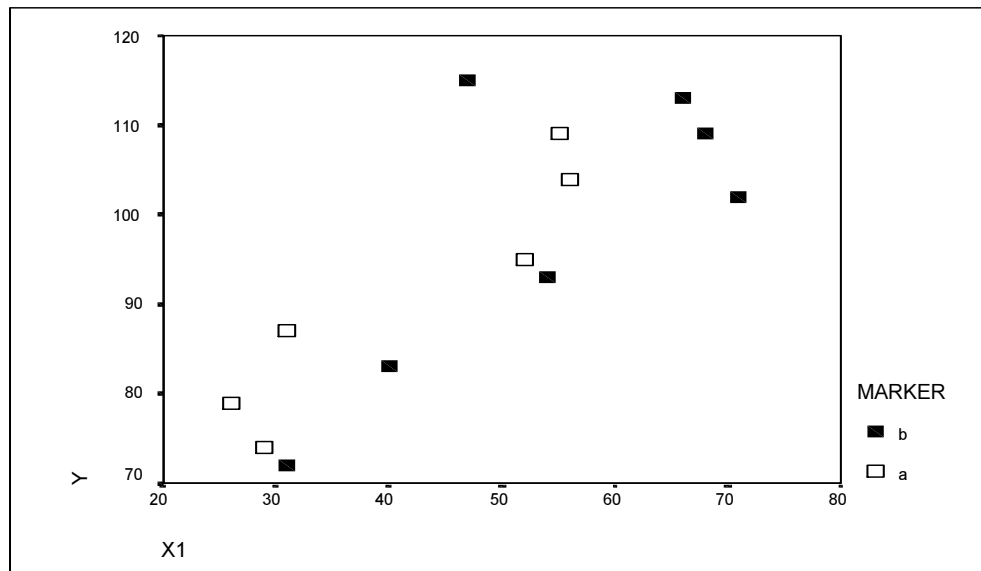
Set Marker by : تتضمن هذه الخانة متغير له فئتين أو أكثر وكل فئة تمثل بعلامة Marker خاصة بها في شكل الانتشار .

Label Cases by : تتضمن هذه الخانة متغير تستعمل قيمه كعناوين Labels لنقاط شكل الانتشار حيث أنه بالإمكان إظهار العناوين أو إخفائها .

علماً أن الخانتين Set Markers by و Label Cases by اختياريّتين Optional أي بالإمكان عدم إدخال متغير في أي منهما .

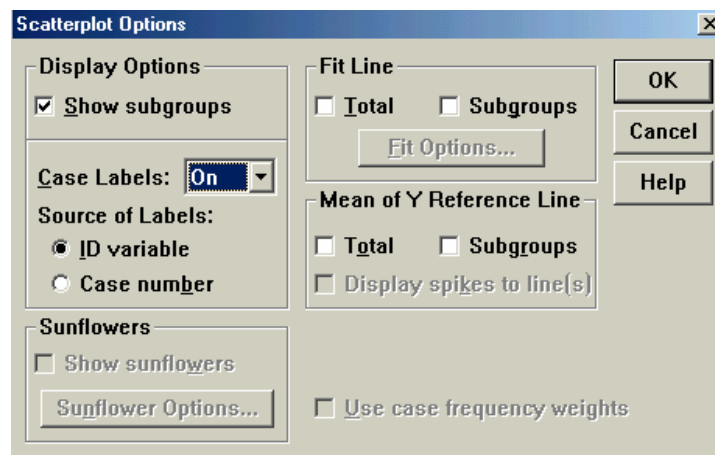
➤ عند نقر زر OK يعرض المخطط التالي :

مخطط رقم 33



لاحظ أن عناوين الحالات لم تظهر في المخطط على الرغم من أننا استعملنا قيم المتغير Label كعناوين للحالات ، ولإظهار العناوين نتبع مايلي :

- انقر المخطط مرتين بزر الماوس الأيسر ليتحول المخطط الى شاشة Chart Editor .
- في شاشة Chart Editor اختر من شريط القوائم Options → Chart فيظهر صندوق حوار Options الذي نرتبه كما يلي :



حيث أن خانة Case Labels تتظم عرض عناوين الحالات Case Labels وتحتوي ثلاثة خيارات :

off : لإزالة عناوين الحالات .

on : لعرض عناوين الحالات كافة .

As is : لعرض عناوين جزء من الحالات أنقر أيقونة تعريف النقطة  في شاشة Chart Editor فيصبح شكل الماوس مشابهاً للأيقونة وعند نقر أي نقطة في المخطط يعرض العنوان الخاص بها في جوار النقطة (لإنهاء عمل الأيقونة تنقر مرة ثانية) .

وبما أننا نرغب في عرض كافة عناوين الحالات فقد اخترنا on .

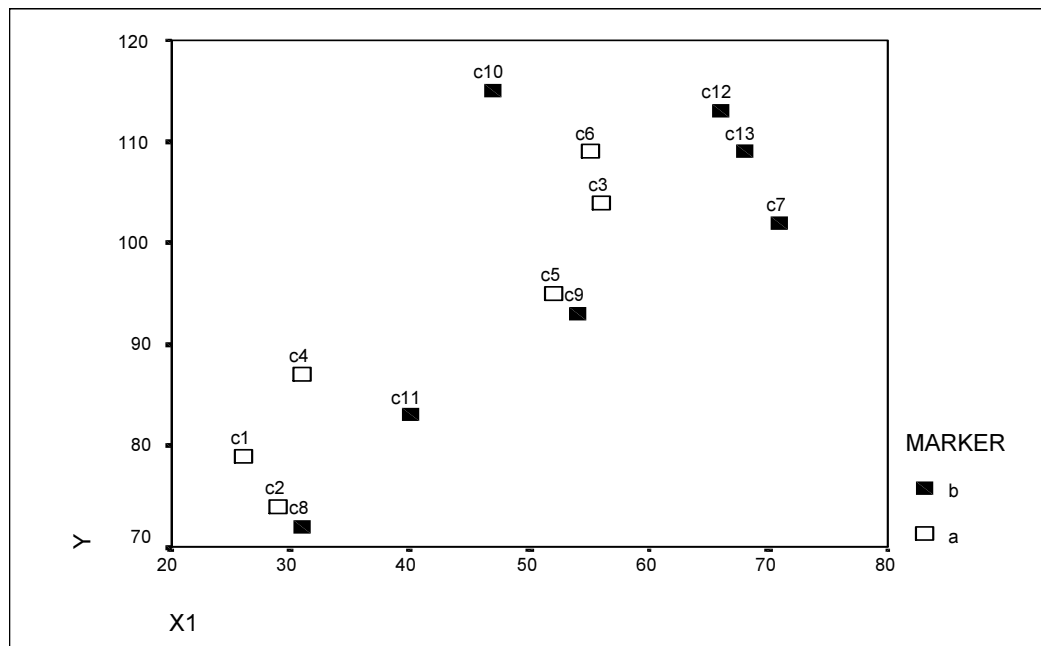
أما خانة Source of Labels فتبين مصدر عناوين الحالات حسب الخيارين التاليين :

ID Variable : يستعمل في حالة وجود متغير يمثل العناوين كما هو الحال في هذا المثال (متغير العناوين هو Label وقد أدخلناه في صندوق حوار Simple Scatterplot .

Case Number : يستعمل رقم الحالة كعنوان في حالة عدم وجود متغير لعناوين الحالات .

• عند نقر زر OK في صندوق الحوار أعلاه تعرض عناوين الحالات كما يلي :

مخطط رقم 34

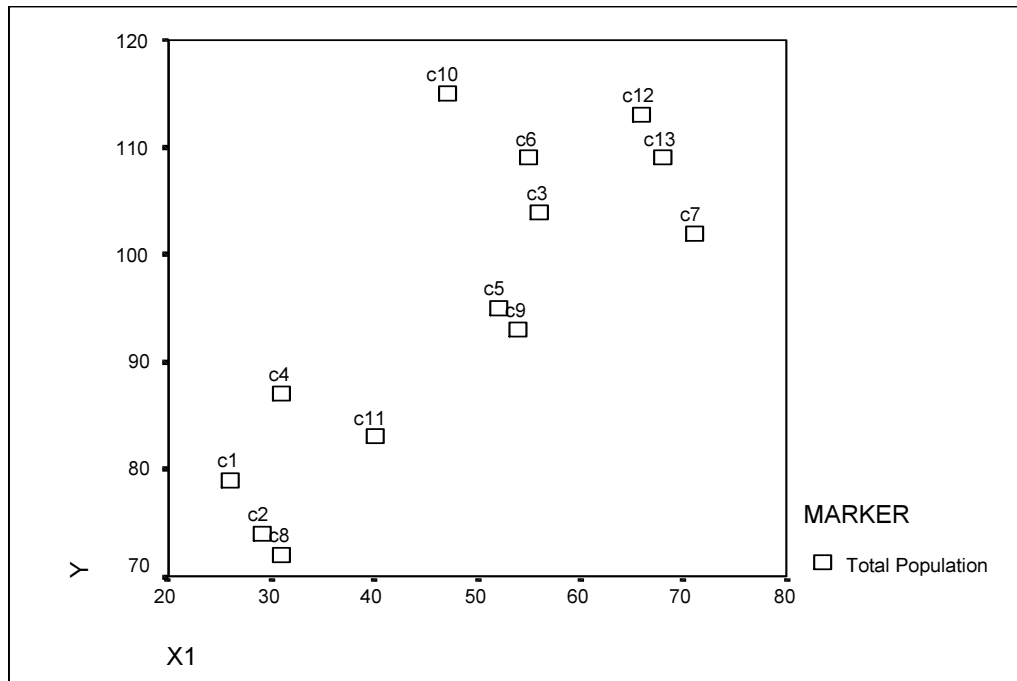


ملاحظة 1 : يمكن عرض عناوين الحالات مباشرة بنقر الزر Options في صندوق حوار Simple Scatterplot ثم تأشير الخيار Display Chart With Case Labels في صندوق حوار Options .

ملاحظة 2 : في صندوق حوار Scatterplot : Options إذا كان الخيار Show Subgroups (في خانة Display Options) مؤشراً يتم عرض المجاميع الفرعية (في هذا المثال المجموعتين a

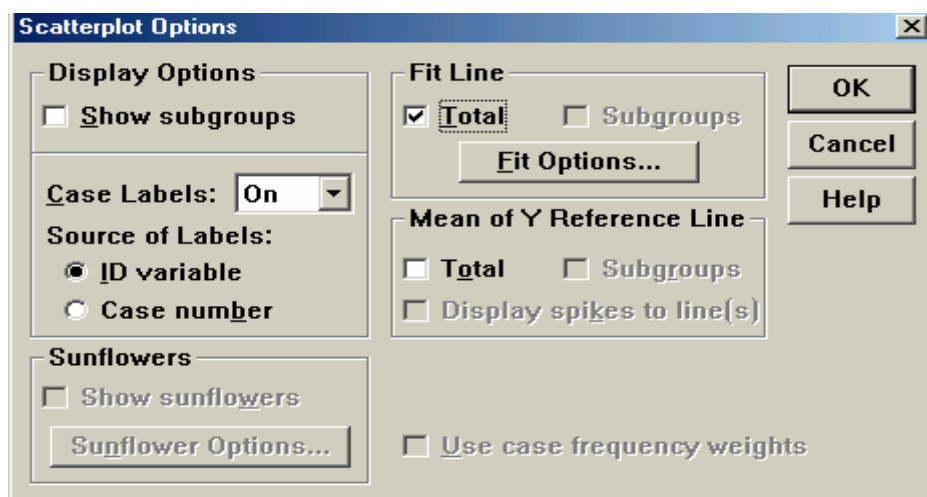
و b) بعلامات مختلفة لكل منهما كما في المخطط أعلاه (المربع الأسود لـ b والمربع الأبيض لـ a) وعند تأشير هذا الخيار يتم عرض المجموعتين بعلامة واحدة كما في المخطط التالي .

مخطط رقم 35



ملاحظة 3 : لإضافة خط الانحدار من الدرجة الأولى $\hat{Y} = B_0 + B_1X_1$ باعتبار أن Y هو المتغير المعتمد و X1 هو المتغير المستقل مع فترة ثقة 95% لمتوسط المتغير Y $(E(Y_0) = B_0 + B_1X_0)$ للمخطط السابق نتبع الخطوات التالية :

□ من شريط القوائم (شاشة Chart Editor) اختر Options → Chart فيظهر صندوق حوار Scatterplot Options الذي نرتبه كما يلي :



لاحظ أننا قمنا بتأشير Total في Fit Line الذي يعني إضافة خط انحدار واحد لمجمل نقاط شكل الانتشار أي أننا لم نأخذ بالاعتبار المجموعتين الجزئيتين a و b من النقاط التي يحددها المتغير Marker .

◀ عند نقر Fit Options في صندوق الحوار السابق يظهر صندوق حوار Fit Line الذي نرتبه كما يلي :

Scatterplot Options: Fit Line

Fit Method

☒ Linear regression ☐ Lowess

☐ Quadratic regression

☐ Cubic regression

% of points to fit: 50

of iterations: 3

Regression Prediction Line(s)

☒ Mean ☐ Individual

Confidence Interval: 95 %

Regression Options

☒ Include constant in equation

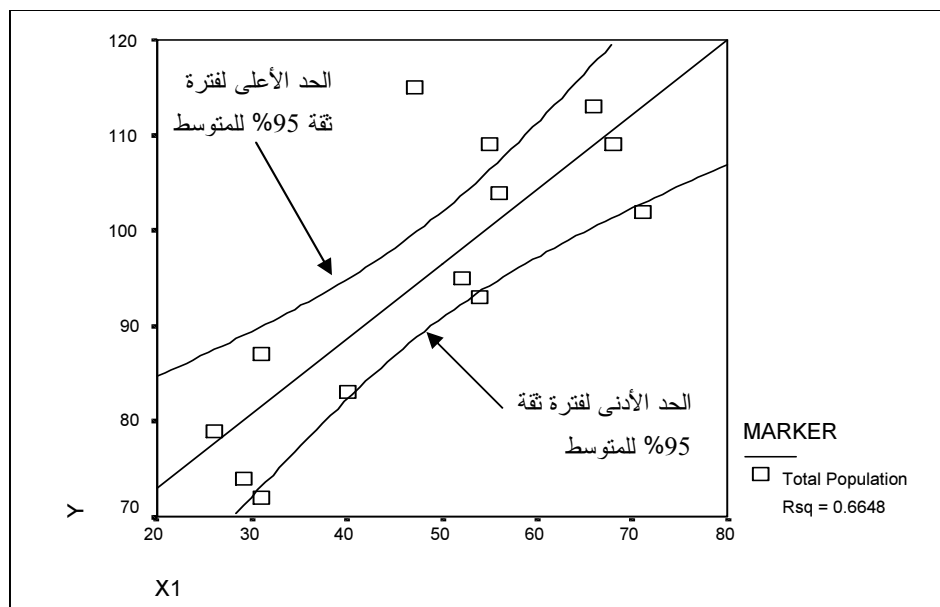
☒ Display R-square in legend

Continue Cancel Help

لاحظ أننا اخترنا Linear Regression انحدار خطي (يمكن اختيار انحدار تربيعي ، تكعيبي أو استخدام طريقة المربعات الصغرى الموزونة لتقدير منحنى الانحدار) وبما أننا نرغب في رسم فترة ثقة 95% للمتوسط فقد اخترنا Mean من خانة Regression Prediction Line (عند اختيار Individual يعرض البرنامج فترة ثقة للقيمة الكلية $Y_0 = B_0 + B_1X_0 + e_0$).

□ عند نقر زر continue ثم زر OK يعرض المخطط التالي (بعد إزالة Case Labels) :

مخطط رقم 36

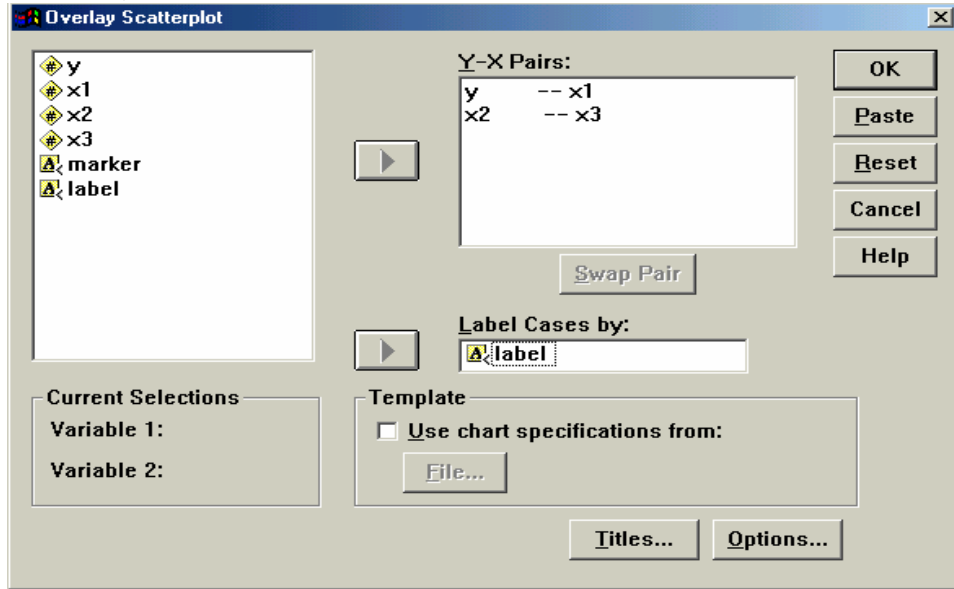


ملاحظة: إذا تم تأشير الخيار Subgroups في خانة Fit Line في صندوق حوار Scatterplot : Options فأن البرنامج يعرض خطي انحدار (خط انحدار لكل فئة من فئتي المتغير Marker) وفي هذه الحالة يجب أن يكون الخيار Show Subgroups في خانة Display Options مؤشراً .

2. شكل انتشار Overlay

يستعمل لتمثيل نقاط الانتشار لزوجين من المتغيرات على الأقل لتمثيل شكل انتشار للمتغيرين Y و X1 وكذلك للمتغيرين X2 و X3 في مخطط واحد تتبع الخطوات التالية :

□ من القوائم نختار Scatter → Graph فيظهر صندوق حوار Scatterplot ومنه نختار Overlay وعند نقر زر Define يظهر صندوق حوار Overlay Scatterplot الذي نرتبه كما يلي :



أن إدخال ازواج المتغيرات في خانة Y-X Pairs يتم كالتالي (مثلا Y و X1) :

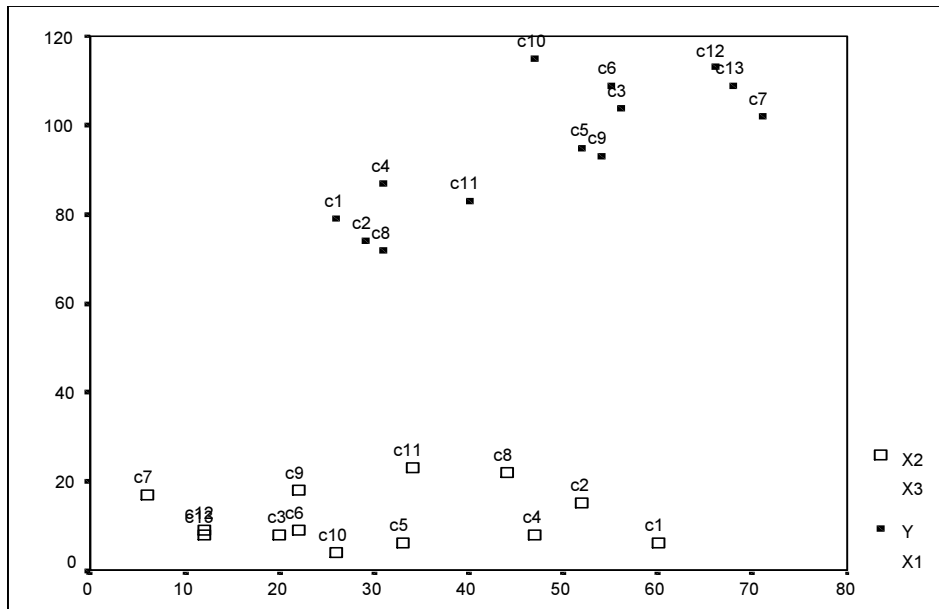
- أنقر المتغير Y لاختياره .
 - أنقر المتغير X1 لاختياره .
 - أنقر [] لنقل المتغيرين الى خانة Y-X Pairs .
- حيث أن المتغير Y يتم تمثيله على المحور العمودي أما المتغير X1 فيمثل على المحور الأفقي ، وب نفس الطريقة يتم إدخال المتغيرين X2 و X3 .

Swap Pair : تستعمل لعكس المحاور لأي زوج من المتغيرات بعد اختياره .

□ أنقر زر Options ثم أنقر Display Chart with Case Labels لعرض عناوين الحالات .

□ عند نقر زر OK في صندوق حوار Overlay Scatterplot يعرض المخطط التالي :

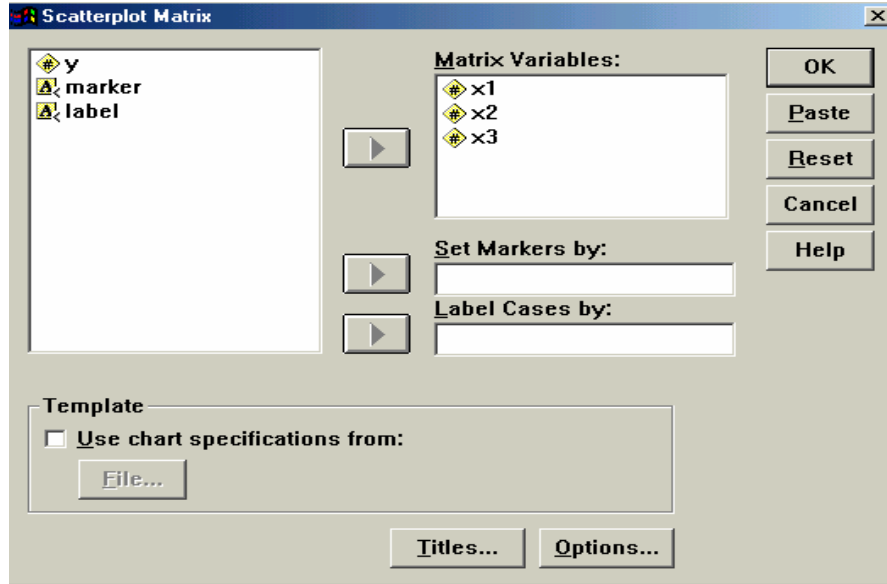
مخطط رقم 37



3. شكل الانتشار Matrix

يستعمل لرسم كل التوافيق الممكنة من متغيرين أو أكثر . لو أردنا مثلاً رسم مخطط Matrix للمتغيرات X1 و X2 و X3 نتبع الخطوات التالية :

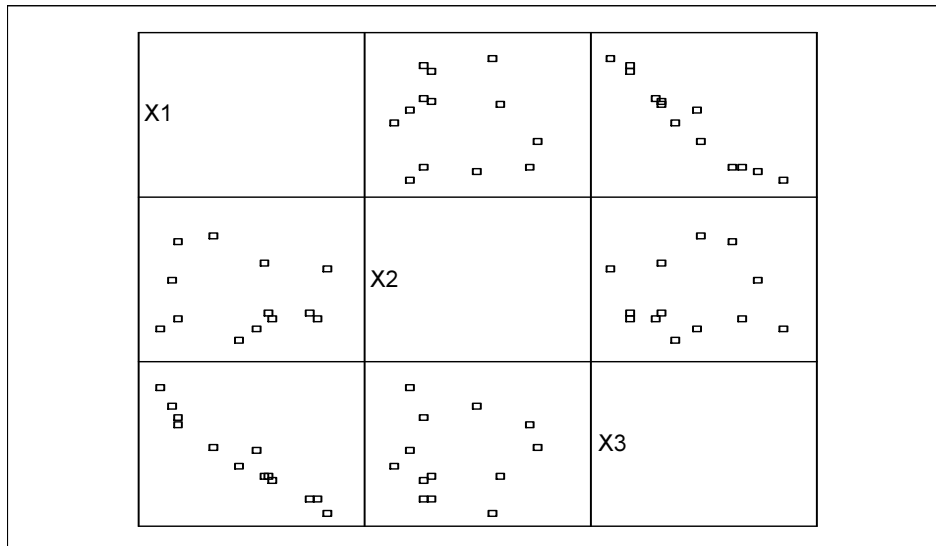
□ من القوائم نختار Scatter → Graph فيظهر صندوق حوار Scatterplot ومنه نختار Matrix فيظهر صندوق حوار Scatterplot Matrix الذي نرتبه كما يلي :



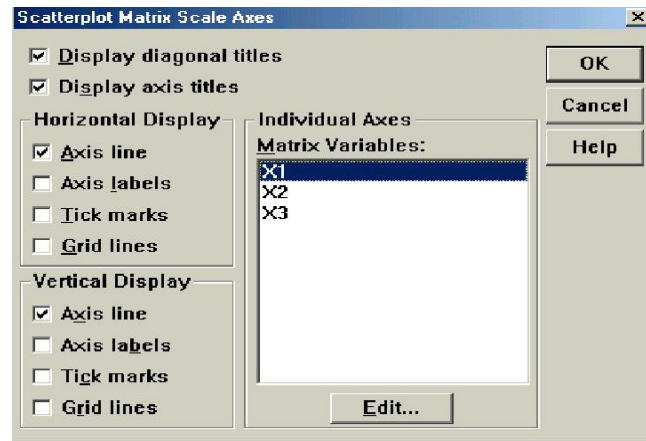
علماً أننا لم ندخل المتغيرين marker و Label اختياريّاً حيث أننا قد أوضحنا تأثير إدخالهما على المخطط.

□ عند نقر زر OK يظهر المخطط التالي :

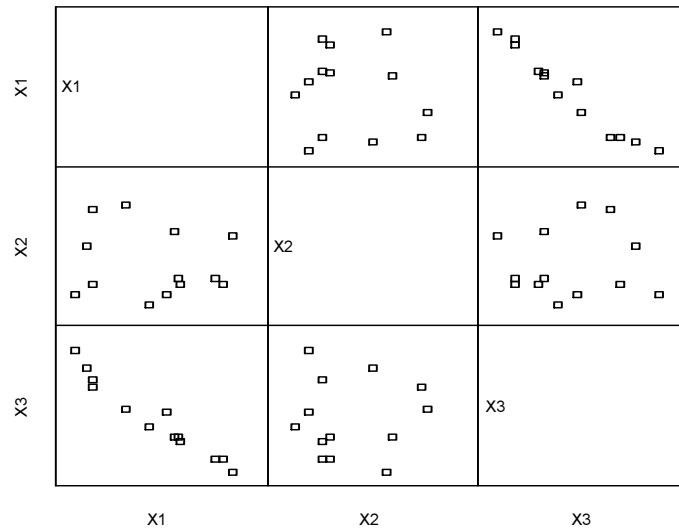
مخطط رقم 38



لإضافة عناوين المحاور Axis Titles أنقر المخطط مرتين بزر الماوس الأيسر ليتحول الى شاشة Chart Editor ثم اختر من شريط القوائم Axis → Chart فيظهر صندوق حوار Scatterplot Matrix Scale Axes حيث نقوم بتأشير الخيار Display Axis Titles ويظهر الصندوق كما يلي :



لتوسيط العناوين انقر زر Edit في صندوق الحوار السابق فيظهر صندوق حوار Edit Selected Axis ومنه اختر Center في خانة Justification. وبنقر زر Continue ثم OK نحصل على المخطط التالي :

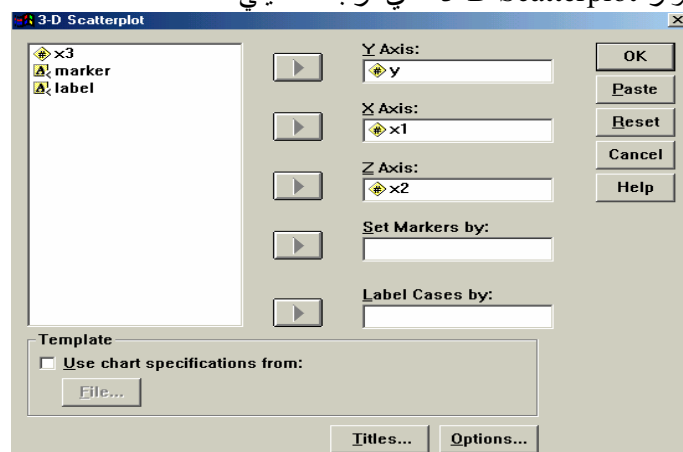


مخطط رقم 39

4. شكل الانتشار ثلاثي الأبعاد 3-D

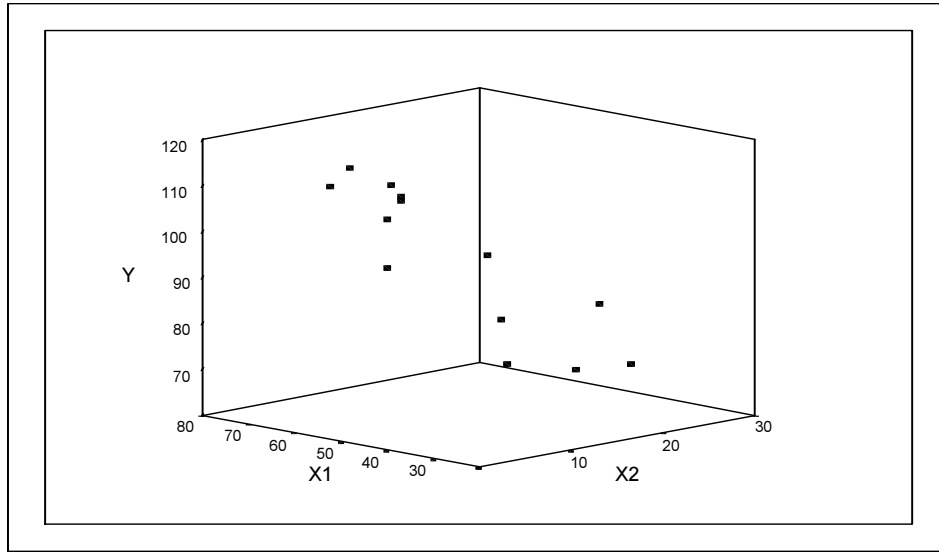
يستعمل لرسم ثلاثة متغيرات في مخطط ثلاثي الأبعاد . لو أردنا مثلاً رسم مخطط 3-D للمتغيرات Y و X1 و X2 نتبع الخطوات التالية :

□ من القوائم نختار Scatterplot → Graph فيظهر صندوق حوار Scatterplot ومنه نختار 3-D فيظهر صندوق حوار 3-D Scatterplot الذي نرتبه كما يلي :



□ عند نقر زر OK يعرض المخطط التالي :

مخطط رقم 40



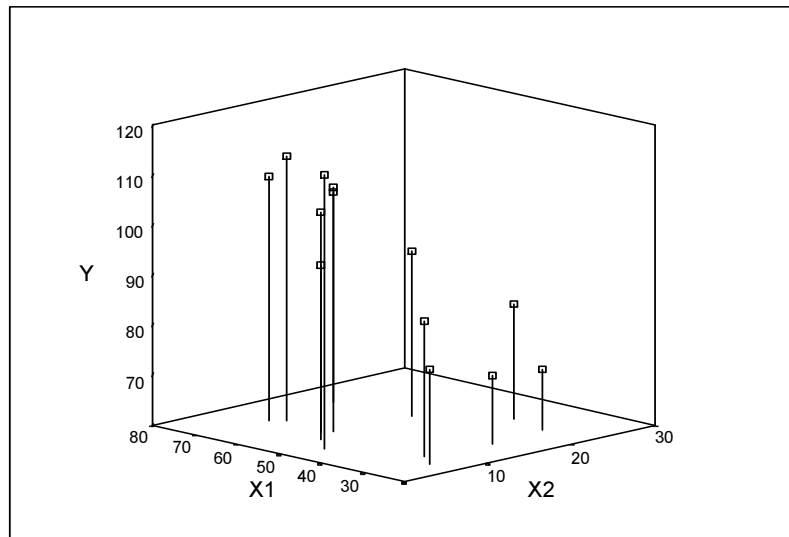
ملاحظة 1:

يمكن تدوير المخطط ثلاثي الأبعاد بنقر الأيقونة  في شريط الأدوات لشاشة Chart Editor أو اختيار Format → 3D-Rotation من شريط القوائم .

ملاحظة 2 :

يمكن إضافة خطوط تمتد من نقاط شكل الانتشار والى أرضية الشكل باختيار Options → Chart من شريط قوائم شاشة Chart Editor فيظهر صندوق حوار 3-D Scatterplot Options ومنه نختار Floor في خانة Spikes وبنقر زر OK نحصل على المخطط التالي :

مخطط رقم 41



علماً أنه يتوفر الخيارين التاليين لخانة Spikes في صندوق حوار 3-D Scatterplot Options :

Centroid : لرسم خط من كل نقطة من نقاط شكل الانتشار الى مركز كافة النقاط في الشكل Centroid .

Origin : لرسم خط من كل نقطة من نقاط شكل الانتشار الى نقطة الأصل Origin .

الفصل الرابع عشر

تبادل البيانات

Data Exchange

أن ملفات البيانات تأتي بصيغ مختلفة ويمكن لبرنامج SPSS أن يتعامل مع الأنواع التالية من الملفات:

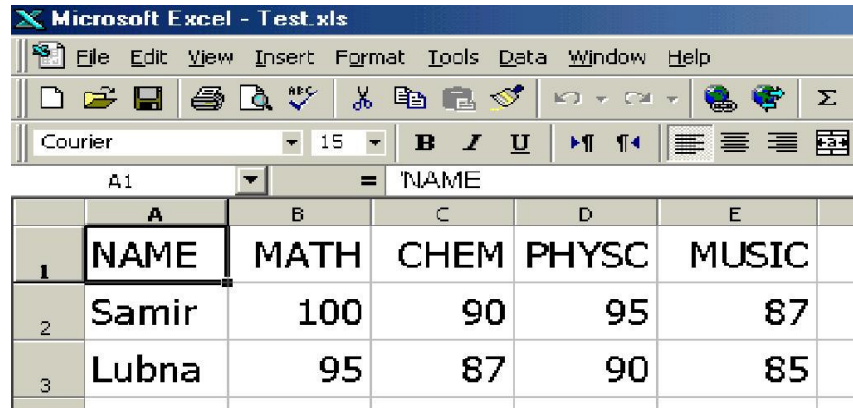
- ☐ أوراق النثر الخاصة ببرنامجي Lotus 1-2-3 و Excel .
 - ☐ ملفات dBase .
 - ☐ ملفات Tab-delimited, الأنواع الأخرى من ملفات نصوص ASCII .
 - ☐ ملفات SPSS المكونة بواسطة أنظمة تشغيل أخرى .
 - ☐ ملفات SYSTAT .
- حيث يمكن استيراد (فتح) ملفات من تطبيقات أخرى وتصدير (خزن) ملفات الى تطبيقات أخرى .

(14 - 1) استيراد الملفات *Importing Data files*

مثال 1 : (استيراد ملف من برنامج EXCEL 97)

الملف التالي Test المخزون في المجلد merge يحتوي قيود مجموعة من الأشخاص في برنامج

Excel97 كما يظهر في الشكل التالي :



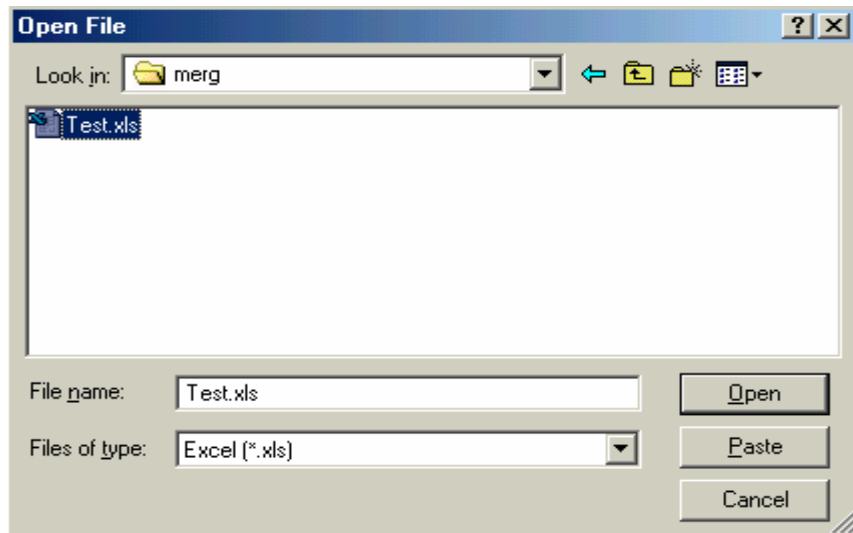
	A	B	C	D	E
1	NAME	MATH	CHEM	PHYSC	MUSIC
2	Samir	100	90	95	87
3	Lubna	95	87	90	85

يطلب استيراد الملف Test من صيغة Excel الى برنامج SPSS .

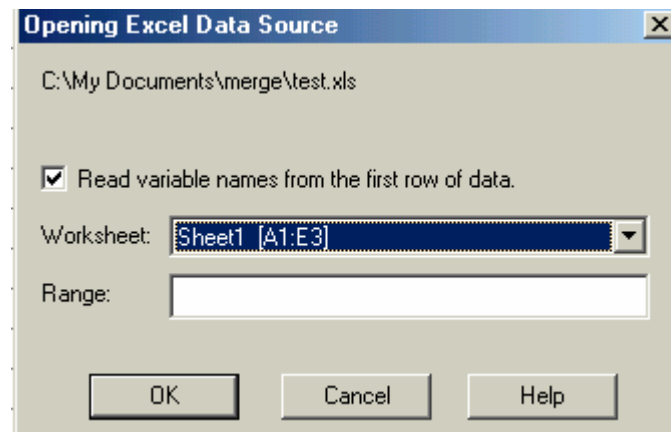
لتنفيذ ذلك نتبع الخطوات التالية :

- ☐ من شريط القوائم في برنامج SPSS اختر Data → Open → File فيظهر

صندوق حوار Open الذي يرتب كما يلي :



لاحظ أننا كتبنا اسم الملف Test في خانة File Name ، أما في خانة Files of Type فقد اخترنا ملفات Excel ذات الاستطالة xls حيث أنه بالإمكان فتح أنواع عديدة من الملفات مثل ملفات SPSS(*.SAV) (وهي ملفات بيانات برنامج SPSS للنظام Windows ذات الاستطالة SAV) وملفات SPSS للنظام DOS وهي SPSS/PC+(*.SYS) وملفات Lotus(*.W*) وملفات dBase(*.dbf) وملفات Text(*.txt) وغيرها من الملفات .
☐ عند نقر زر open يظهر صندوق الحوار التالي :



حيث أن :

Read Variable Names From The First Row of data : لقراءة أسماء المتغيرات من الصف الأول في ملف Excel .

Worksheet : لتحديد الورقة التي تحتوي البيانات المراد نقلها الى برنامج SPSS .

Range : لتحديد مدى معين من البيانات التي يراد نقلها الى برنامج SPSS .

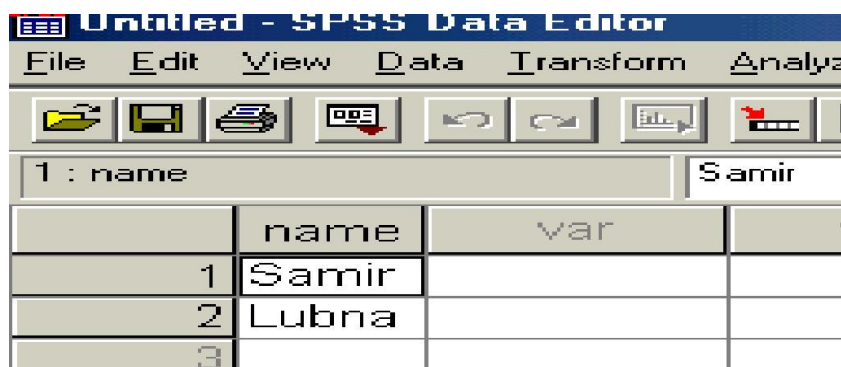
☐ عند نقر زر OK يتم نقل ملف Excel الى SPSS وتظهر البيانات في شاشة Data Editor كما يلي :

Untitled - SPSS Data Editor						
File Edit View Data Transform Analyze Graphs Utilities Window Help						
1 : name Samir						
	name	math	chem	physc	music	
1	Samir	100	90	95	87	
2	Lubna	95	87	90	85	
3		.	.	.	.	

ويمكن خزن الملف السابق كملف SPSS ذو الاستطالة SAV .

ملاحظات :

1. عند تحديد Range معين من البيانات في حقل Range في صندوق حوار Opening Excel Data Source (مثلاً A1:A3) يظهر الناتج التالي :



	name	var
1	Samir	
2	Lubna	
3		

2. يمكن تحديد أي ورقة عمل في الملف مثلاً Sheet2 بدلاً من Sheet1 في خانة Worksheet في صندوق حوار Opening Excel Data Source لاستيراد البيانات الموجودة في هذه الورقة أو أي Range من بيانات هذه الورقة .

3. عند عدم تأشير الخيار Read Variable Names From The First Row of data في المثال السابق فإن أسماء المتغيرات في ملف Excel سوف تضاف كأول قيمة لكل متغير في ملف SPSS ويتم إضافة أسماء متغيرات افتراضية بدلاً عنها .

4. عند وجود أكثر من متغير بنفس الاسم في ملف Excel فإن كل متغير سوف يكون له أسم وحيد Unique عند استيراد الملف الى برنامج SPSS كما أن أسم كل متغير في برنامج Excel له رموز Characters يزيد عددها عن 8 فسيتم قطع الاسم لغاية 8 رموز ويتم إضافة الاسم الكامل كعنوان للمتغير Label عند استيراد الملف الى برنامج SPSS .

5. يمكن استيراد بيانات ملف Excel بطريقة ثانية بتظليل البيانات المطلوب نقلها في برنامج Excel ثم اختيار Copy → Edit ثم الانتقال الى برنامج SPSS واختيار Paste → Edit من شريط القوائم في شاشة Data View فيتم نقل البيانات الى برنامج SPSS ولكن يجب ملاحظة فقدان أسماء المتغيرات الحقيقية وظهور أسماء افتراضية بدلاً عنها كما نلاحظ فقدان قيم المتغيرات الرمزية كون النوع الافتراضي لمتغيرات SPSS هو النوع العددي Numeric ويتوجب في هذه الحالة اختيار النوع String للمتغيرات الرمزية في شاشة Data Editor قبل إنجاز عملية النسخ .

مثال 2 : (استيراد ملفات النصوص Text Files)

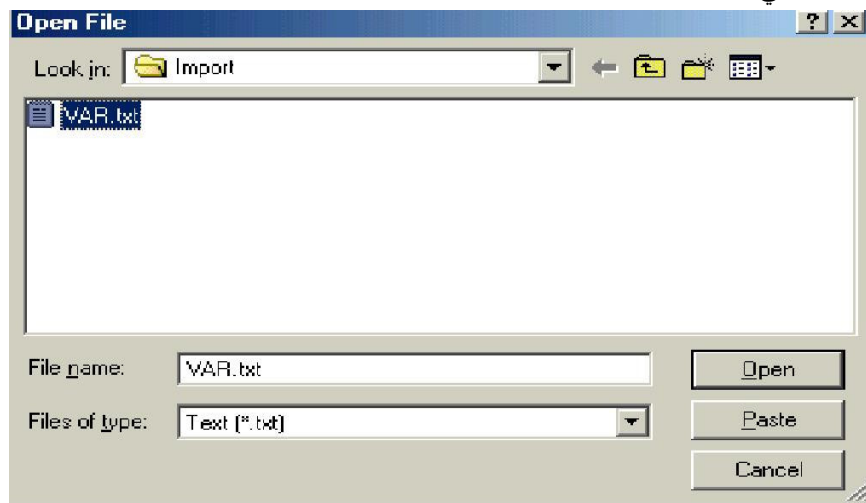
الملف التالي يحتوي بيانات تم إدخالها عن طريق معالج النصوص MS-DOS Editor لبرنامج MS-DOS والمخزون بأسم Var.txt وقد استخدمت الفراغات Spaces للفصل بين المتغيرات وكما يلي :

x1	x2	x3
12	30	33
16	45	60
29	64	88

لأستيراد الملف أعلاه الى برنامج SPSS نتبع الخطوات التالية:

□ من شريط القوائم اختر **File → Read Text Data** فيظهر صندوق حوار **Open File**

الذي نرتبه كما يلي :



كما يمكن انجاز هذه الخطوة باختيار **Data → Open → File** من شريط القوائم وفي هذه الحالة يتوجب تحديد نوع الملف **Text** في حقل **Files of type** في صندوق الحوار أعلاه .

□ عند نقر زر **Open** في صندوق الحوار أعلاه يظهر صندوق حوار **Text Import wizard** (step 1 of 6) وكما يلي :

Text Import Wizard - Step 1 of 6

Welcome to the text import wizard!

This wizard will help you read data from your text file and specify information about the variables.

Does your text file match a predefined format?

☐ Yes ☒ No

Text file:

	var1	var2	var3
1			
2			
3			
4			

1 2 3 4

0 10 20 30 40 50

1	x1	x2	x3
2	12	30	33
3	16	45	60
4	29	64	88

< Back Next > Finish Cancel Help

□ عند نقر زر Next يتم الانتقال الى الخطوة رقم 2 وكما يلي :

Text Import Wizard - Step 2 of 6

How are your variables arranged?

☒ Delimited - Variables are delimited by a specific character (i.e., comma, tab).
☐ Fixed width - Variables are aligned in fixed width columns.

Are variable names included at the top of your file?

☒ Yes
☐ No

Text file:

0 10 20 30 40 50

1	x1	x2	x3
2	12	30	33
3	16	45	60
4	29	64	88

< Back Next > Finish Cancel Help

في حقل How are your Variables arranged? فقد أشرنا الخيار Delimited الذي يستعمل للمتغيرات التي يفصل بينها فاصل من نوع معين (فارزة، فارزة منقوطة، فراغ، الخ) .
 في حقل Are Variables Names included at the top of your file? فقد أشرنا الخيار yes .
 عند نقر زر Next يتم الانتقال الى الخطوة 3 وكما يلي :

The first case of data begins on which line number? 2

How are your cases represented?

☒ Each line represents a case

☐ A specific number of variables represents a case: 3

How many cases do you want to import?

☒ All of the cases

☐ The first 1000 cases.

☐ A random percentage of the cases (approximate): 10 %

Data preview

1	12	30	33
2	16	45	60
3	29	64	88

< Back Next > Finish Cancel Help

في حقل the first Case of data Begins on Which Line Number ? فقد حددنا ابتداء الحالة الأولى بالسطر الثاني من الملف لأن السطر الأول مخصص لأسماء المتغيرات .
 في حقل How are Your Cases represented? فقد أشرنا الخيار Each Line Represents a Case أما في حالة وجود أكثر من حالة في السطر الواحد فنستخدم الخيار الثاني حيث يتوجب (في هذه الحالة) تحديد عدد المتغيرات التي تكون حالة واحدة .
 في حقل How Many Cases Do you want to Import ? فقد تم تأشير الخيار الأول - كافة الحالات .

□ عند نقر زر Next ننتقل الى الخطوة الرابعة من ال Wizard وكما يلي :

Text Import Wizard - Delimited Step 4 of 6

Which delimiters appear between variables?

☐ Tab ☒ Space
☐ Comma ☐ Semicolon
☐ Other:

Data preview

x1	x2	x3
12	30	33
16	45	60
29	64	88

< Back Next > Finish Cancel Help

الغاية من هذه الخطوة هي تحديد نوع الفاصل Delimiter وكما ذكرنا فأنا قد أيسر عملنا الفراغات كفاصل بين المتغيرات عند الإدخال في معالج النصوص Edit وقد تم تأشير الخيار Space (يرجى ملاحظة أن معظم الخيارات تؤثر تلقائياً من قبل البرنامج) .

□ عند نقر زر Next يتم الانتقال الى الخطوة الخامسة من ال wizard

Text Import Wizard - Step 5 of 6

Specifications for variable(s) selected in the data preview

Variable name:

Data format:

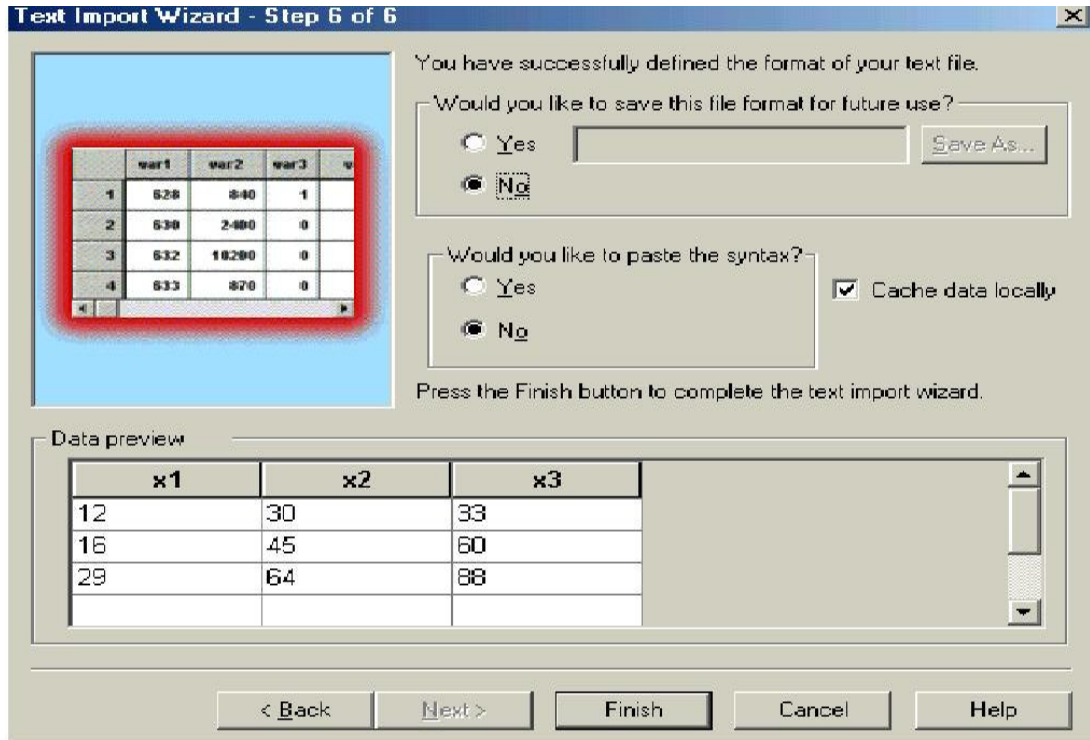
Data preview

x1	x2	x3
12	30	33
16	45	60
29	64	88

< Back Next > Finish Cancel Help

يمكن في هذه الخطوة تغيير أسم ونوع Format & Name أي من المتغيرات بنقر خلية أسم المتغير في Data Preview أسفل الصندوق ثم كتابة الأسم الجديد ونوع المتغير في حقل Variable Name و Data Format .

□ عند نقر زر Next يتم الانتقال الى الخطوة السادسة والأخيرة من الـ Wizard



في حالة الرغبة في تخزين صيغة البيانات التي تم استيرادها لغرض تطبيقها في المستقبل على بيانات مستوردة مباشرة نؤشر الخيار Yes كإجابة عن السؤال Would You Like to save This Format for Future Use? وفي هذه الحالة انقر Save As لحزن الصيغة في ملف معين (نوع استضافة tpf) حيث يمكن استرجاع الصيغة المخزونة مسبقاً وتطبيقها على البيانات المستوردة عند بداية استيراد الملف في الخطوة الأولى من الـ Wizard حيث يتم الإجابة بـ Yes على السؤال Does Your Text File Match a Predefined Format? ثم نقر الزر Browse لاسترجاع ملف الصيغة. لعدم الرغبة في تخزين صيغة البيانات فقد أشرنا الخيار NO .

□ عند نقر زر Finish يتم إضافة البيانات الى شاشة Data Editor وكما يلي :

Untitled - SPSS Data Editor				
File Edit View Data Transform Analyze Graphs Utilities Wi				
11 :				
	x1	x2	x3	va
1	12.0	30.0	33.0	
2	16.0	45.0	60.0	
3	29.0	64.0	88.0	
4	.	.	.	

(2-14) تصدير الملفات Exporting Data Files

يمكن تصدير ملفات SPSS الى تطبيقات أخرى وذلك عن طريق تخزين الملف بصيغ عديدة منها مايلي :

1. SPSS(*.SAV) وهي ملفات البيانات للإصدار الحالي للبرنامج .
2. SPSS 7.0(*.SAV) وهي ملفات البيانات للإصدار السابع للبرنامج .
3. SPSS/PC+(*.SYS) وهي ملفات بيانات SPSS للنظام DOS .
4. Tab-delimited(*.dat) وهي ملفات النصوص Text .
5. Fixed ASCII(*.dat)
6. Excel(*.xls)
7. Lotus1-2-3(*.w*) للإصدارات 1,2,3 .
8. dBase(*.dbf) للإصدارات II و III و IV .

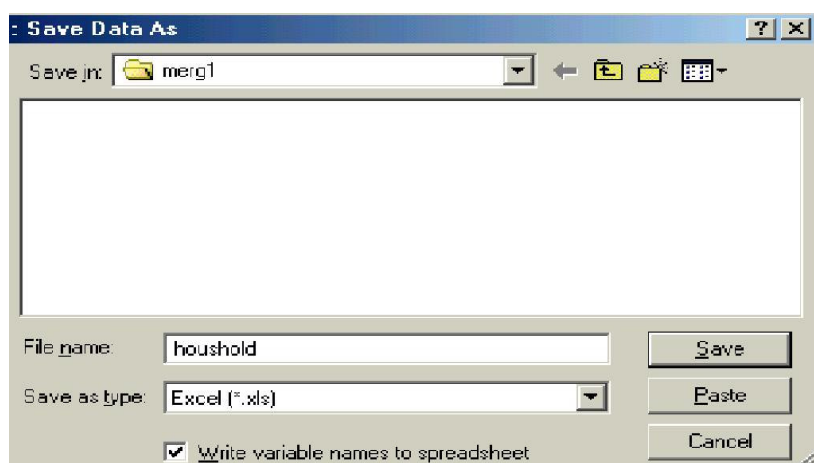
مثال 3: (على تصدير البيانات)

لدينا الملف household.sav الذي يظهر كما يلي في شاشة Data Editor لبرنامج SPSS :

household.sav - SPSS Data Editor					
File Edit View Data Transform Analyze Graphs Utilities Window Help					
1 : name Ahmad					
	name	age	edu	housno	
1	Ahmad	20	Sec	10	
2	Zeki	35	Bsc	10	
3	Sabah	30	sec	10	
4	Zainab	15	Prim	10	
5	Ibrahim	17	Sec	12	
6	Samir	40	Ma	12	
7	Selma	36	Bs	12	

المطلوب مايلي :

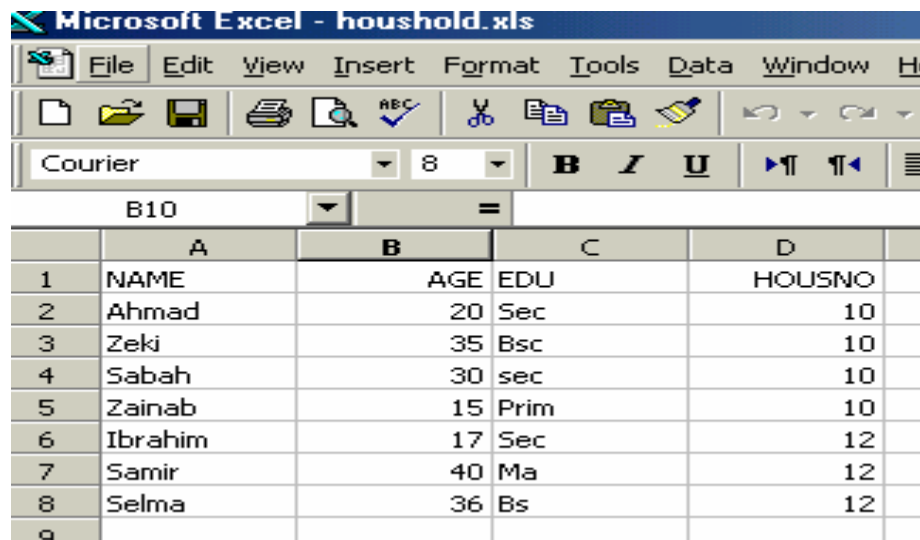
1. تصدير الملف الى برنامج Excel .
 2. تصدير الملف الى معالج النصوص MS-Editor .
1. لتصدير الملف الى برنامج Excel نتبع الخطوات التالية :
- من شريط القوائم اختر Save As → File فيظهر صندوق حوار SAVE AS الذي نرتبه كما



يلي :

لاحظ أنه تم إعطاء أسم للملف في حقل File Name وتم تحديد نوع الملف في حقل Save as Type كملف من ملفات Excel ذو الاستطالة xls .

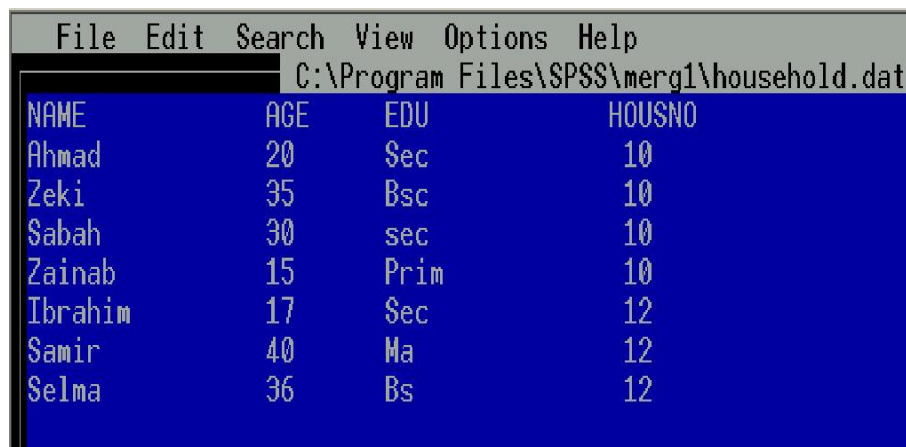
عند تأشير write variable names to spreadsheet يتم إضافة أسماء المتغيرات الى ورقة عمل
نشر Excel في السطر الأول منها أما عدم التأشير فيؤدي الى إزالة أسماء المتغيرات من السطر الأول .
□ عند نقر زر Save يتم تخزين الملف بصيغة ملفات Excel أي سيكون لدينا ملفين ضمن المجلد
merg1 أحدهما هو household ذو الاستطالة sav والآخر بنفس الاسم ولكن بالاستطالة xls .
يمكن فتح الملف household في برنامج Excel بأختيار Open → File ثم تحديد موقع الملف Path
حيث يظهر في ورقة Excel كما يلي :



The screenshot shows the Microsoft Excel interface with the file 'household.xls' open. The spreadsheet contains the following data:

	A	B	C	D
1	NAME	AGE	EDU	HOUSNO
2	Ahmad	20	Sec	10
3	Zeki	35	Bsc	10
4	Sabah	30	sec	10
5	Zainab	15	Prim	10
6	Ibrahim	17	Sec	12
7	Samir	40	Ma	12
8	Selma	36	Bs	12

2. لتصدير الملف الى MS-Editor نكرر نفس الخطوة الأولى (عند التصدير الى Excel) ولكن نختار
Tab-delimited في حقل Save as Type في صندوق حوار Save Data As ثم نعطي اسم
الملف household. وعند نقر زر Save يتم تخزين الملف بصيغة Tab-delimited أي يضاف ملف جديد
للملفين السابقين داخل المجلد merg1 باسم household أيضاً ولكن باستطالة dat ، ويظهر الملف
household بفواصل ثابتة عند فتحه في MS-Editor كما في الشكل التالي.



The screenshot shows the MS-Editor interface with the file 'C:\Program Files\SPSS\merg1\household.dat' open. The data is displayed in a text format with columns separated by tabs:

NAME	AGE	EDU	HOUSNO
Ahmad	20	Sec	10
Zeki	35	Bsc	10
Sabah	30	sec	10
Zainab	15	Prim	10
Ibrahim	17	Sec	12
Samir	40	Ma	12
Selma	36	Bs	12

الفصل الخامس عشر

كتابة الأوامر

Syntax Commands

أن معظم الأوامر يمكن الحصول عليها من القوائم أو في صناديق الحوار ولكن هناك أوامر لا تتوفر إلا باستخدام لغة الأوامر Command Language حيث يمكن من خلال هذه اللغة تنفيذ فعاليات لا تتوفر ضمن أوامر SPSS وهذه اللغة تسمح بخزن الفعاليات Jobs في ملف الأوامر Syntax File حيث يمكن استرجاعها في أي وقت وتنفيذها لأي مجموعة من المتغيرات .

(1-15) ملف الأوامر Syntax File : وهو عبارة عن ملف نصوص Text File يحتوي على أوامر مكتوبة Commands (لغة برمجة) ويجب مراعاة القواعد التالية عند كتابة هذه الأوامر :

- يجب أن يبدأ كل أمر بسطر جديد وينتهي بالفترة (.) .
- الأوامر الفرعية Sub Commands تفصل بعلامة Slash (/) ، علماً أن وضع علامة Slash قبل أول أمر فرعي يكون اختياريًا .
- يجب كتابة الأسماء الكاملة للمتغيرات .
- النص المحصور بين علامتي اقتباس Quotation Marks يجب أن يكتب في سطر واحد .
- لا يزيد طول السطر في ملف الأوامر عن 80 رمز .
- تستخدم الفترة (.) للتعبير عن الفاصلة العشرية بغض النظر عن إعدادات Windows الموجودة في Regional Setting .

للقوف على تفاصيل كتابة البرامج نحيل القارئ إلى SPSS Syntax reference Guide .

(2- 15) الطرق المساعدة في بناء ملفات الأوامر Command Syntax

أن كتابة الأوامر تحتاج إلى معرفة تفصيلية وممارسة طويلة في كتابة البرامج مما لا يتوفر لدى الكثيرين من مستخدمي برنامج SPSS . عوضاً عن ذلك يمكن الاستعانة بأحد الوسائل الثلاثة المبسطة التالية في تكوين ملف الأوامر :

- كتابة الأوامر من صناديق الحوار Dialog boxes .
- نسخ الأوامر من Log في مخرجات البرنامج .
- نسخ الأوامر من Journal File .

(1 - 2 - 15) كتابة الأوامر من صناديق الحوار Dialog boxes : وهذه أسهل طريقة لبناء ملف

الأوامر وذلك يفتح صندوق الحوار لأمير معين ثم استعمال الزر Paste في لصق الأوامر في ملف الأوامر syntax File كما هو واضح في المثال التالي :

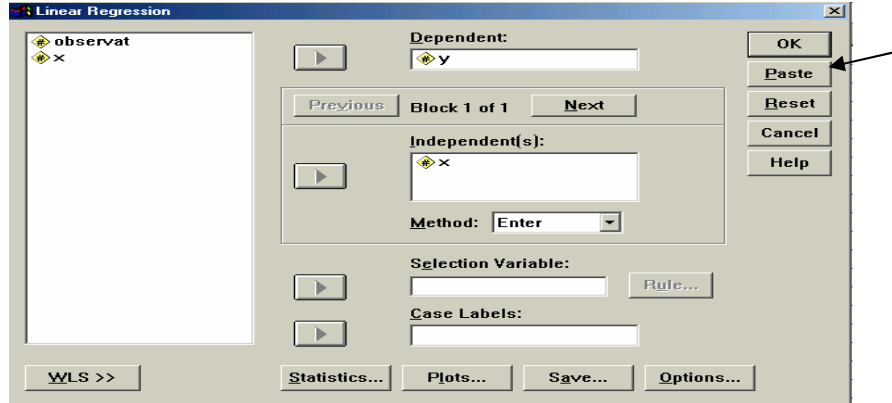
مثال 1 : لنفس بيانات المثال 3 الوارد في البند (10 - 4 - 1) حيث تم إدخال المتغيرين X و Y في شاشة Data Editor ونرغب في بناء ملف أوامر بطريقة الكتابة من صندوق الحوار يتضمن الفعاليات التالية :

1. استخراج معاملات Coefficients انحدار Y/X .
2. استخراج جدول تحليل التباين ANOVA .

3. استخراج معامل Durbin-Watson .

لتنفيذ ذلك نتبع الخطوات التالية :

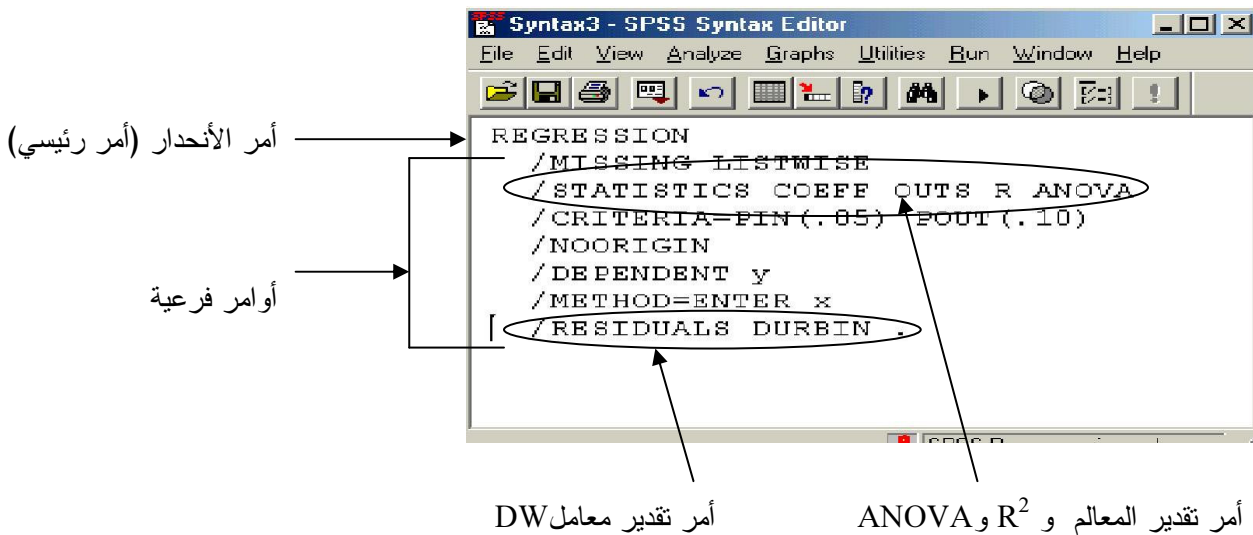
1. من شريط القوائم اختر Linear Regression → Analyze فيظهر صندوق حوار Linear Regression الذي نرتبه كما يلي :



2. في صندوق الحوار أعلاه انقر Statistics ثم أشر الخيارات التالية في صندوق حوار Statistics :

- Estimate لتقدير معالم النموذج .
- Model Fit لتكوين جدول تحليل التباين ANOVA واستخراج R^2 .
- Durbin - Watson لتقدير معامل DW .

3. انقر زر Paste في صندوق حوار Linear Regression فيتم فتح ملف الأوامر Syntax File تلقائياً (بأسم افتراضي) ولصق أوامر الانحدار المختارة فيه وكما يلي :



لحزن ملف الأوامر بأسم معين من شريط القوائم في نافذة Syntax Editor اختر File → Save As ثم أعطي أسم الملف ونوعه (SPSS Syntax File) الذي يكون ذو الاستطالة SPS .

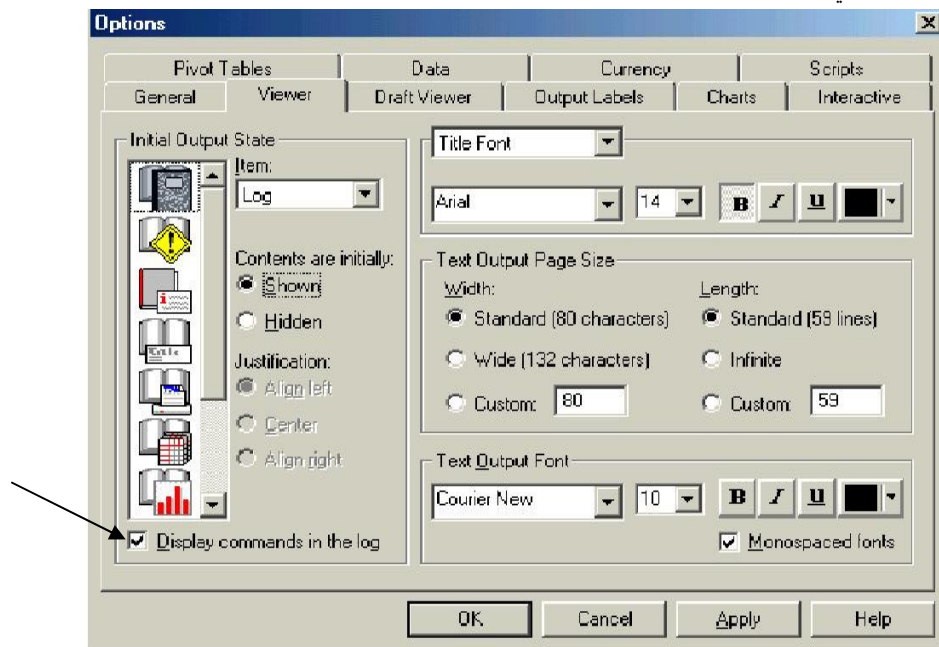
(2 - 2 - 15) نسخ الأوامر من Log في مخرجات البرنامج

يمكن نسخ أوامر من Log في مخرجات البرنامج عند تنفيذ أمر أي أسلوب إحصائي كالانحدار مثلاً وقبل ذلك يجب إظهار log في مخرجات البرنامج كما يلي :

□ من شريط القوائم لنافذة Data Editor اختر Options → Edit فيظهر صندوق حوار Options .

□ أنقر عروة Viewer في صندوق حوار Options للانتقال الى صندوق حوار Viewer.

□ أشر الخيار Display Commands in the Log إذا لم يكن مؤشراً .حيث يظهر صندوق حوار Viewer كما يلي :

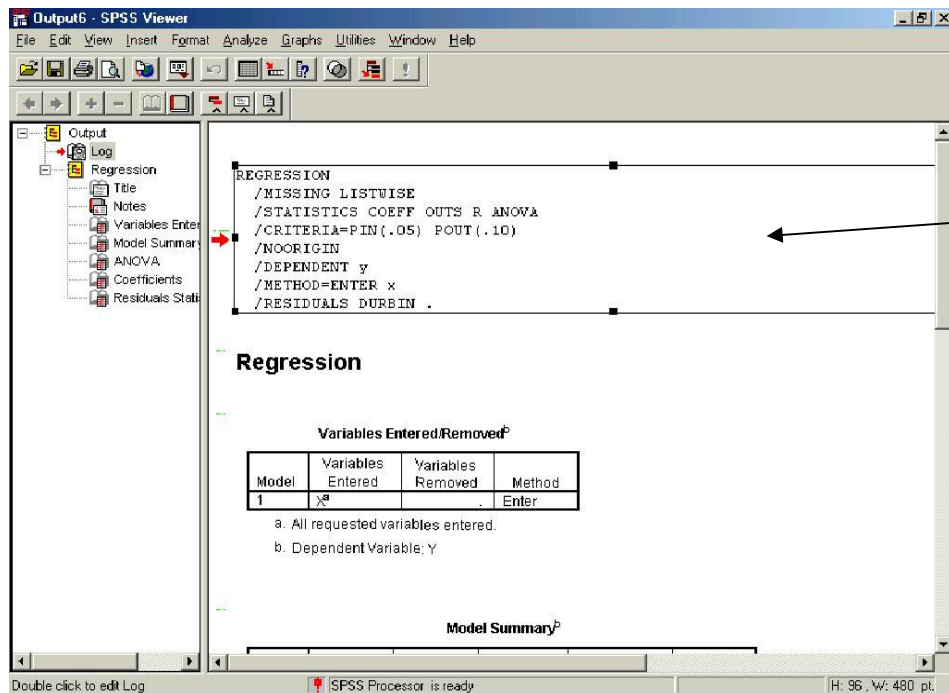


□ أنقر زر OK . أن هذا سيجعل الأوامر تظهر مع مخرجات أي أسلوب إحصائي في شاشة المخرجات Viewer.

مثال 2 :

لنفس المثال السابق يطلب تكوين ملف الأوامر بطريقة النسخ من Log .

□ لتنفيذ ذلك نتبع الخطوتين 1 و 2 للمثال السابق ثم أنقر زر OK في صندوق حوار Linear Regression فتظهر مخرجات الانحدار في شاشة viewer كما يلي :



Log

□ أفتح ملف أوامر جديد باختيار الأمر **File → New → Syntax** من شريط القوائم في شاشة **SPSS Viewer**.

□ انقر في شاشة **SPSS Viewer** مرتين بزر الماوس الأيسر لتفعيله.

□ ظلل الأوامر التي ترغب بنسخها (ترغب مثلاً بنسخ كافة الأوامر).

□ من قوائم **Viewer** اختر **Edit → Copy**.

□ من القوائم في نافذة **Syntax** اختر **Edit → Paste**.

□ أخرج ملف الأوامر **Syntax** الناتج للاستفادة منه فيما بعد.

(15 - 2 - 3) نسخ الأوامر من **Journal File**

أن كافة الأوامر التي تنفذ خلال الجلسة يتم تسجيلها في **Journal File** يسمى **spss.jnl** وعليه يمكن استدعاء هذا الملف (يتواجد ضمن الموقع الافتراضي **C:\WINDOWS\TEMP**) حيث يمكن تنقيح الملف بحذف الأوامر غير المطلوبة ورسائل الخطأ ثم خزن الملف الناتج بأسم مختلف كملف أوامر **Syntax** باستطالة **SPS** حيث يمكن استدعائه وتمشيته في أي وقت.

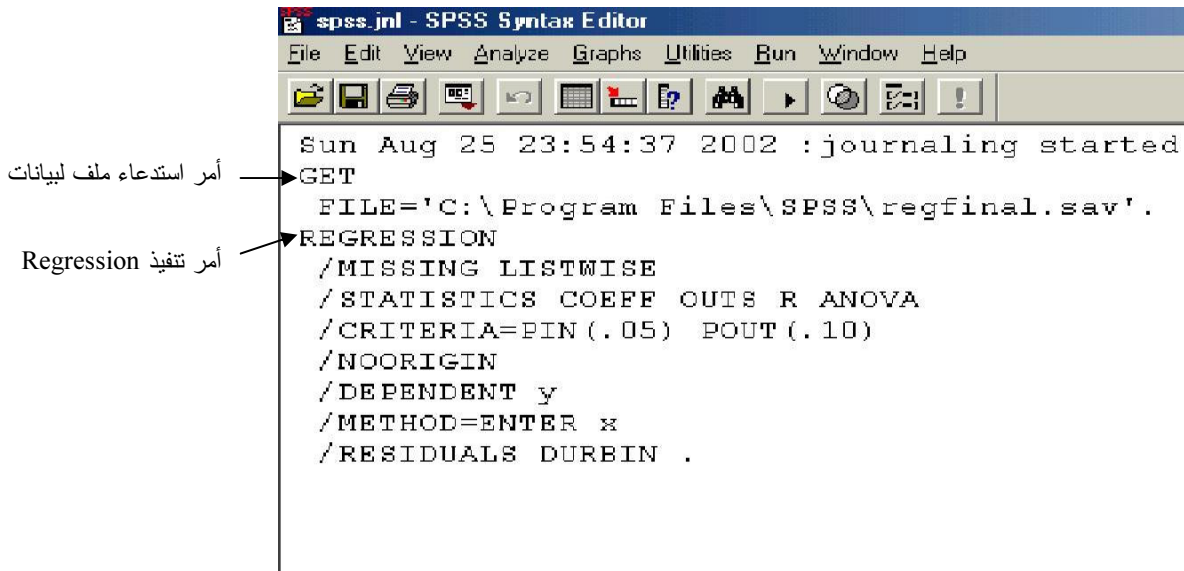
ملاحظة : يمكن التحكم بإعدادات **Journal File** باختبار الأمر **Edit → Options → General** في شريط القوائم لشاشة **Data Editor** حيث يمكن التحكم بموقع الملف. لغرض جعل البرنامج يقوم بتخزين الأوامر المنفذة في **Journal File** يجب تأشير الخيار **Record Syntax in Journal** (إذا لم يكن مؤشراً) حيث تتوفر طريقتين لخزن الأوامر :

Append : تخزن الأوامر المنفذة في الجلسات المتعاقبة ابتداء من تنصيب البرنامج.

Overwrite : تخزن الأوامر المنفذة خلال الجلسة الحالية فقط لبرنامج **SPSS**.

مثال 3 : للمثال الأول يطلب تكوين ملف الأوامر بطريقة النسخ من Journal File .

- لتتفيذ ذلك نتبع الخطوتين 1 و 2 للمثال الأول ثم أنقر زر OK في صندوق حوار Linear Regression فتظهر مخرجات الانحدار في شاشة SPSS Viewer .
- من شريط القوائم في شاشة SPSS Viewer أختار File → Open → Other ثم افتح الملف spss.jnl في الموقع C:\WINDOWS\TEMP فيظهر الملف كما يلي :



- يمكن تنقيح الملف spss.jnl في هذا المثال نرغب بالاحتفاظ به كما هو عدا أننا نقوم بحذف السطر الأول (تاريخ الخزن) لأنه ليس من أوامر البرمجة . الآن نقوم بخزن الملف بأسم مختلف كملف Syntax ذو الاستطالة SPS لأستدعائه وتمشيته عند الحاجة .
- ملاحظة : نوع الخزن المستخدم لـ journal file في هذا المثال هو Overwrite (أي أوامر الجلسة الحالية فقط) .

(15-3) تمشية ملف الأوامر To Run command syntax

- لتمشية ملف الأوامر (أو تمشية البرنامج) أو أي جزء منه نتبع الخطوات التالية :
- أفتح ملف الأوامر Syntax File المخزون مسبقاً بالأمر Syntax Open من File من شريط القوائم .
- من شريط القوائم في شاشة SPSS Syntax Editor أختار Run لتمشية البرنامج حيث تتوفر الخيارات التالية :
- All : لتمشية كافة أوامر الملف .

Selection : لتمشية الأوامر المختارة (المظللة) Highlighted .

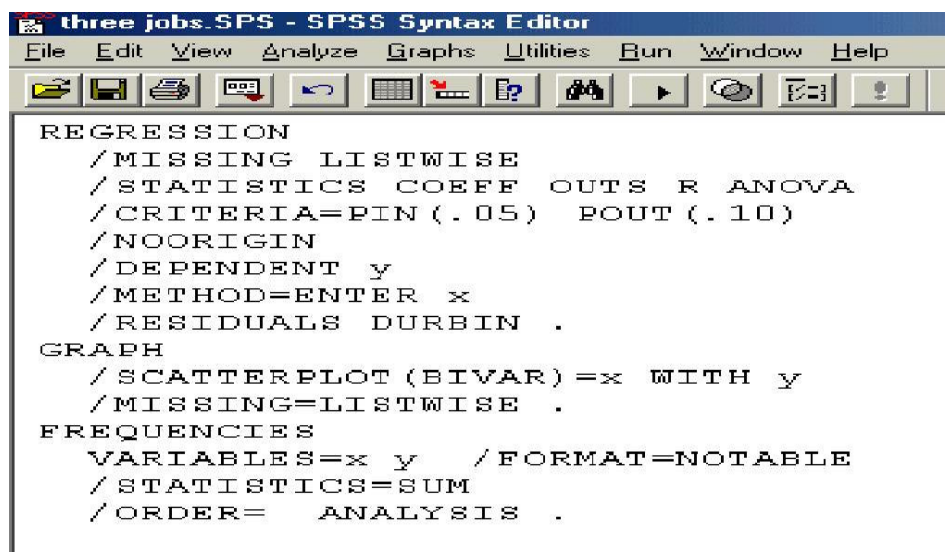
Current : لتمشية الأوامر في موقع المؤشر Cursor .

To End : لتمشية الأوامر من موقع المؤشر والى نهاية الملف .

مثال 4 : لنفس بيانات المثال الأول (يتضمن المتغيرين X و Y) يطلب تكوين ملف أوامر يقوم بتنفيذ الفعاليات التالية :

1. أستخراج معاملات Coefficients انحدار Y/X و أستخراج جدول تحليل التباين ANOVA وأستخراج معامل Durbin-Watson .

2. تكوين شكل الانتشار Scatterplot الذي يوضح العلاقة بين X و Y .
 3. استخراج مجموع قيم المتغير X ومجموع قيم المتغير Y .
- لتكوين ملف الأوامر للمثال أعلاه (بافتراض أن المتغيرين X و Y قد تم إدخالهما في شاشة Data Editor لبرنامج SPSS) سنستعمل طريقة اللصق من صناديق الحوار بموجب الخطوات التالية :
- ☐ أفتح ملف أوامر جديد Syntax File بالأمر File → New → Syntax .
 - ☐ نفذ الخطوتين 1 و 2 نفسها للمثال الأول ثم انقر زر Paste في صندوق حوار Linear Regression فيتم لصق أوامر الانحدار في ملف الأوامر المفتوح حديثاً .
 - ☐ لتكوين أوامر مخطط Scatterplot للمتغيرين X و Y من شريط القوائم في شاشة Data Editor اختر Simple → Scatter → Graphs ثم انقر زر Define فيفتح صندوق حوار Simple Scatterplot .
 - ☐ في صندوق حوار Simple Scatterplot انقر المتغير Y وأدخله في قائمة : Y Axis ثم انقر المتغير X وأدخله في قائمة : X Axis .
 - ☐ انقر زر Paste فيتم إضافة أمر Scatterplot الى ملف الأوامر .
 - ☐ لجمع قيم كل من المتغيرين X و Y من شريط القوائم اختر Frequencies → Descriptive Statistics → Analyze فيفتح صندوق حوار Frequencies .
 - ☐ في صندوق حوار Frequencies انقل المتغيرين X و Y الى قائمة Variable(s) ثم ألغ تأشير Display Frequency Tables. ثم انقر زر Statistics فيفتح صندوق حوار Statistics .
 - ☐ في صندوق حوار Statistics أشر الخيار Sum.
 - ☐ انقر زر Paste في صندوق حوار Frequencies فيتم إضافة أوامر Frequencies الى ملف الأوامر .
 - ☐ في نافذة Syntax Editor ستحصل على ملف أوامر بأسم افتراضي هو Syntax1 يتضمن الأوامر الثلاثة المطلوبة ، لخرن هذا الملف من شريط القوائم اختر File → Save As ثم أعط أسم الملف (مثلاً three jobs) بالاستطالة SPS ، حيث يظهر هذا الملف كما يلي :



```

three jobs.SPS - SPSS Syntax Editor
File Edit View Analyze Graphs Utilities Run Window Help
[Icons]
REGRESSION
  /MISSING LISTWISE
  /STATISTICS COEFF OUTS R ANOVA
  /CRITERIA=PIN(.05) POUT(.10)
  /NOORIGIN
  /DEPENDENT y
  /METHOD=ENTER x
  /RESIDUALS DURBIN .
GRAPH
  /SCATTERPLOT(BIVAR)=x WITH y
  /MISSING=LISTWISE .
FREQUENCIES
  VARIABLES=x y /FORMAT=NOTABLE
  /STATISTICS=SUM
  /ORDER= ANALYSIS .

```

□ عند اختيار All → Run من شريط القوائم في النافذة أعلاه يتم تنفيذ الأوامر الثلاثة ونحصل على النتائج التالية :

Regression

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Durbin-Watson
1	.900 ^a	.810	.787	8.82	1.885

a. Predictors: (Constant), X

b. Dependent Variable: Y

ANOVA^b

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	2661.050	1	2661.050	34.174	.000 ^a
	Residual	622.950	8	77.869		
	Total	3284.000	9			

a. Predictors: (Constant), X

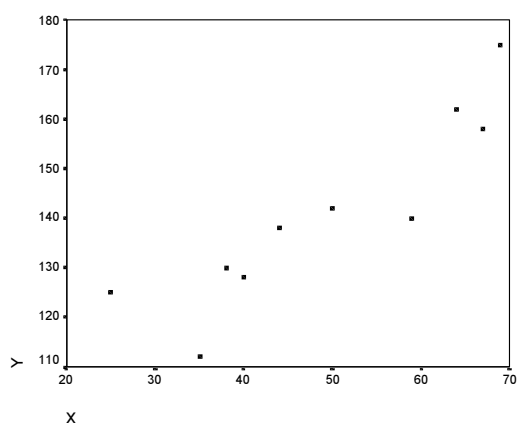
b. Dependent Variable: Y

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	85.044	9.970		8.530	.000
	X	1.140	.195	.900	5.846	.000

a. Dependent Variable: Y

Graph



Frequencies

Statistics

		X	Y
N	Valid	10	10
	Missing	0	0
Sum		491	1410

□ يمكن تنفيذ جزء من الأوامر فمثلاً لتنفيذ Regression نقوم بتظليل الأوامر الخاصة به ثم اختيار Run → Selection علماً أنه لا يمكن تنفيذ أمر فرعي دون غيره بل يتم تنفيذ كافة الأوامر الفرعية العائدة للأمر الرئيسي .

□ إذا كان موقع المؤشر على الأمر الرئيسي Graph فعند اختيار الأمر Run → Current من شريط القوائم أو نقر الأيقونة  في شريط الأدوات فيتم عرض المخطط البياني Scatterplot فقط. أما عند اختيار الأمر Run To End فيتم عرض المخطط البياني وناتج الأمر Frequency للمجموع Sum .

ملاحظة :

إذا رغبتنا بتطبيق الأوامر الثلاثة المذكورة آنفاً على المتغيرين X1 و Y1 بدلاً من المتغيرين X و Y فأن ذلك يتم حسب الخطوات التالية :

- إدخال المتغيرين X1 و Y1 في نافذة Data Editor .
- فتح ملف الأوامر الذي يحتوي الفعاليات الثلاث والمخزون مسبقاً باسم three jobs .
- تنقيح الملف بإبدال كل X بـ X1 وكل Y بـ Y1 .
- تمشية البرنامج Run .

الفصل السادس عشر

تحليل الاستجابات المتعددة

Multiple Response Analysis

في حالة إجابة المستجيب عن سؤال يمثل رأيه مثلاً في برامج التلفزيون عند تحديد الفئات التالية (جيد ، متوسط ، رديء) فإن بإمكانه اختيار فئة واحدة فقط من الفئات الثلاث المذكورة . أما إذا سئل المستجيب عن الصحف الأثيرة لديه فإن هناك احتمال أن يختار أكثر من صحيفة واحدة وفي هذه الحالة فإن إيجاد التكرارات Frequencies يتم عن طريق تحليل الاستجابات المتعددة وهناك أسلوبين في التحليل :

1. أسلوب التكرارات متعددة الاستجابة Multiple response Frequencies : ومنه نحصل على جدول تكراري بسيط .
2. أسلوب جداول التقاطع متعددة الاستجابة Multiple Response Crosstabs : ومنه نحصل على جداول تكرارية ببعدين وثلاثة أبعاد .

(16 - 1) أسلوب التكرارات متعددة الاستجابة Multiple response Frequencies

مثال 1

في استبيان لمعرفة تفضيل المواطن لواحدة أو أكثر من الصحف المحلية التالية (بابل ، الجمهورية ، العراق ، القادسية ، الثورة) فلغرض معرفة التكرارات لكل صحيفة (عدد الأشخاص الذين يفضلون كل صحيفة من الصحف) فإنه يمكن الاختيار بين أسلوبين :

1. أسلوب المتغيرات المتعددة ذات الفئتين Multiple dichotomy method : بموجب هذه الطريقة فأنا نخصص متغير لكل صحيفة من الصحف الخمسة يأخذ قيمتين (فئتين) فقط (مثلاً صفر وواحد) فإذا كان المستجيب يفضل الصحيفة المعنية فأنا نعطي القيمة واحد لمتغير الصحيفة عدا ذلك يأخذ المتغير القيمة صفر . المثال المبسط التالي يبين طريقة ترتيب البيانات في شاشة Data Editor .

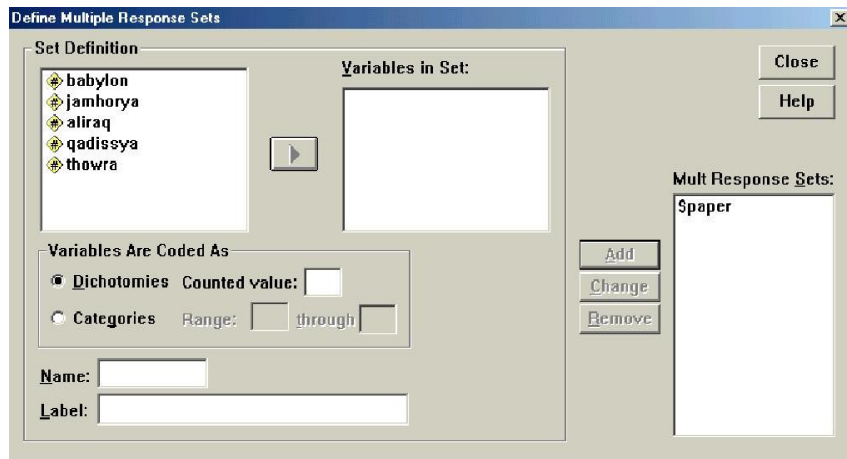
Babylon jamhorya aliraq qadissya thowra

1	0	1	1	0
1	1	0	0	0
0	0	1	0	0
1	0	0	1	1
0	1	1	0	0
1	0	0	1	1
1	0	1	0	0

لإيجاد التكرارات المقابلة لكل صحيفة نتبع الخطوات التالية :

من شريط القوائم لشاشة Data Editor اختر Define → Multiple Response → Analyze
sets

فيظهر صندوق حوار Define Multiple Response Sets الذي نرتبه كما يلي :



وقد قمنا بالخطوات التالية لترتيب صندوق الحوار وكما يظهر أعلاه :

- اختيار المتغيرات الخمسة من خانة Set definition ونقلها الى خانة Variables in set في اليمين .

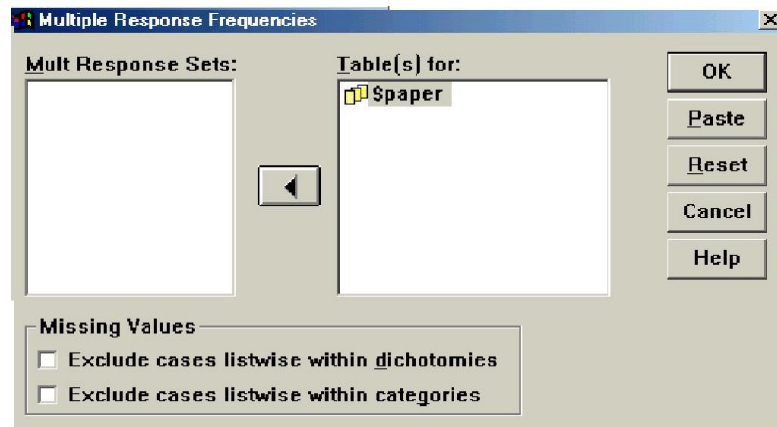
- كتابة فئة العد (واحد) بعد اختيار Dichotomies Counted value .

- إعطاء أسم لمجموعة الاستجابة المتعددة يتكون من سبعة رموز كحد أعلى في خانة Name وليكن paper (ويمكن إعطاء عنوان Label للمجموعة اختياريا) .

- نقر زر Add لإضافة المجموعة الى خانة Mult response sets

◀ عند نقر زر Close في صندوق الحوار يتم خزن مجموعة الاستجابة المتعددة بأسم paper للاستفادة منها فيما بعد .

◀ من القوائم اختر Frequencies → Multiple Response → Analyze فيظهر صندوق حوار Multiple Response Frequencies الذي نرتبه كالتالي :



حيث قمنا بنقل مجموعة الاستجابة \$paper الى خانة Tables for .

◀ عند نقر زر OK في صندوق الحوار أعلاه نحصل على الجدول التكراري البسيط التالي :

Multiple Response

Group \$PEPER

(Value tabulated = 1)

Dichotomy label Name	Pct of Count	Pct of Responses	Pct of Cases
BABYLON	5	31.3	71.4
JAMHORYA	2	12.5	28.6
ALIRAQ	4	25.0	57.1
QADISSYA	3	18.8	42.9
THOWRA	2	12.5	28.6
Total responses	16	100.0	228.6

0 missing cases; 7 valid cases

ملاحظات :

1. أن المتغيرات التي تظهر في شاشة Data Editor وهي Babylon و Jamhorya ... الخ تعرف بالمتغيرات الأولية Elementary Variables أما المجموعة التي تم تكوينها بأسم \$paper فتعرف بأسم مجموعة الاستجابات المتعددة Multiple response Set وهي تتضمن عدة متغيرات أولية.

2. في حالة وجود عناوين للمتغيرات المستخدمة في التحليل فأن عناوين هذه المتغيرات Labels تظهر في الجدول التكراري بدلاً من أسماء المتغيرات وفي هذا المثال فأن أسماء المتغيرات ظهرت في الجدول لأننا لم نعط عناوين للمتغيرات .

2. أسلوب الفئات المتعددة Multiple Category Method : لتطبيق هذا الأسلوب فانه يتوجب أولاً معرفة أكبر عدد من الإجابات على السؤال (أي أكبر عدد من الصحف المفضلة لدى المستجيب) لنفترض أن هذا العدد هو ثلاثة صحف كحد اعلى وفي المقابل نكون نفس العدد من المتغيرات لنسمها X1 و X2 و X3 حيث أن قيم كل منها الكودات التالية :

1 = Babylon , 2 = Jamhorya , 3 = Aliraq , 4 = Qadissya , 5 = Thowra
فمثلاً إذا كان المستجيب يفضل ثلاثة صحف هي بابل والعراق والقادسية فان المتغير X1 سيأخذ الكود 1 والمتغير X2 له الكود 3 والمتغير X3 له الكود 4 أما إذا كان المستجيب يفضل صحيفتي بابل والجمهورية فقط فأننا سنخصص الكود 1 للمتغير X1 والكود 2 للمتغير X2 أما X3 فستكون له قيمة مفقودة Missing Value . فبالنسبة للمثال السابق فأننا سترتب المتغيرات X1 و X2 و X3 في شاشة Data Editor

كما يلي :

X1	X2	X3
1	3	4
1	2.	
3.	.	
1	4	5
2	3.	
1	4	5
1	3.	

لإيجاد التكرارات المقابلة لكل صحيفة (أي لكل كود) نتبع الخطوات التالية وهي مشابهة تقريباً

لأسلوب المتغيرات المتعددة ذات الفئتين :

من شريط القوائم لشاشة Data Editor اختر Define sets → Multiple Response → Analyze

- فيظهر صندوق حوار Define Multiple Response Sets الوارد أنفاً الذي نرتبه كما يلي :
- اختيار المتغيرات الثلاثة X1 و X2 و X3 من خانة Set definition ونقلها الى خانة Variables in set في اليمين .
 - اختيار Categories في خانة Variables Are Coded As وتعيين فئات الصحف من 1 الى 5 (Range 1 through 5) .
 - إعطاء أسم لمجموعة الاستجابة المتعددة يتكون من سبعة رموز كحد أعلى في خانة Name وليكن paper (ويمكن إعطاء عنوان Label للمجموعة اختيارياً) .
 - نقر زر Add لإضافة المجموعة الى خانة Mult response sets
 - ◀ عند نقر زر Close في صندوق الحوار يتم خزن مجموعة الاستجابة المتعددة بأسم paper للاستفادة منها فيما بعد .
 - ◀ من القوائم اختر Frequencies → Multiple Response → Analyze فيظهر صندوق حوار Multiple Response Frequencies حيث قمنا بنقل مجموعة الاستجابة \$paper الى خانة Tables for .
 - ◀ عند نقر زر OK في صندوق الحوار المذكور نحصل على الجدول التكراري البسيط التالي :

Group \$PAPER

Category label	Count	Pct of Responses	Pct of Cases
Code			
1	5	31.3	71.4
2	2	12.5	28.6
3	4	25.0	57.1
4	3	18.8	42.9
5	2	12.5	28.6
	-----	-----	-----
Total responses	16	100.0	228.6

0 missing cases; 7 valid cases

أن هذه النتيجة مشابهة تماماً لما توصلنا إليه بأسلوب المتغيرات المتعددة ذات الفئتين .

ملاحظة : في حالة إعطاء قيم المتغيرات الثلاثة X1 و X2 و X3 عناوين للقيمة Value Label فأن عنوان القيمة يظهر بدلاً من القيمة نفسها في الجدول التكراري .

(16 - 2) أسلوب جداول التقاطع متعددة الاستجابة Multiple Response Crosstabs

أن هذا الأسلوب يتيح تكوين جداول تقاطع للمتغيرات الأولية Elementary Variables أو للمجاميع المعرفة Multiple response Sets أو لكليهما معاً .

مثال 2 :

لنفس بيانات المثال الأول لنفترض أن هناك سؤالاً يتعلق برغبة المستجيب في أن تكون الصحيفة بالألوان أو اسود وابيض وفي هذه الحالة سوف نكون المتغير Colored الذي يتضمن قيمتين (القيمة 1 إذا كان المستجيب يفضل الألوان والقيمة 0 إذا كان لا يفضل الألوان) بالإضافة إلى المتغيرات التي تمثل تفضيل المستجيب لكل صحيفة من الصحف. حيث تظهر المتغيرات في شاشة Data Editor كما يلي :

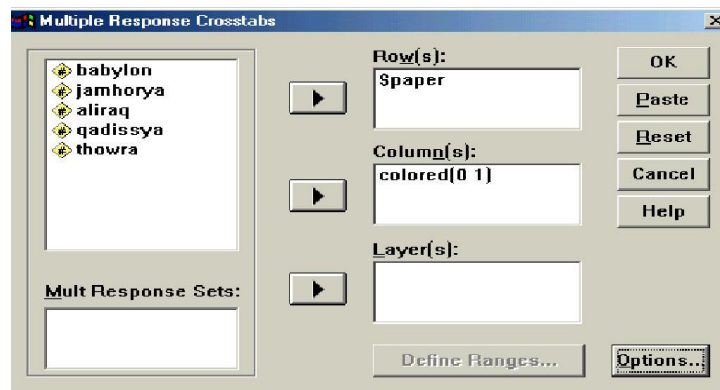
Babylon	jamhorya	aliraq	qadissya	thowra	colored
1	0	1	1	0	1
1	1	0	0	0	0
0	0	1	0	0	0
1	0	0	1	1	1
0	1	1	0	0	0
1	0	0	1	1	1
1	0	1	0	0	1

نرغب في تكوين جدول تكراري مزدوج (جدول تقاطع) بين تفضيل الصحف وبين السؤال الذي يتضمن تفضيل الألوان .

لتكوين هذا الجدول نحتاج للمتغير colored الذي يمثل تفضيل الألوان وهو متغير أولي Elementary Variable كما نحتاج إلى مجموعة الاستجابات المتعددة \$paper التي تمثل تفضيل الصحف ، يجب أولاً تكوين هذه المجموعة بنفس الطريقة المذكورة في المثال الأول بأسلوب المتغيرات المتعددة ذات الفئتين Multiple dichotomy method (علماً أنه يمكن اعتماد أسلوب الفئات المتعددة Multiple Category Method).

لتكوين جدول التقاطع بحيث تكون فئات تفضيل الصحف في الصفوف وفئات تفضيل الألوان في الأعمدة نتبع الخطوات التالية :

من شريط القوائم اختر Crosstabs → Multiple Response → Analyze فيظهر صندوق حوار Multiple Response Crosstabs الذي نرتبه كما يلي :



علماً أنه يتوجب تحديد الحدين الأعلى والأدنى لفئات المتغير الأولي colored وهي (0,1) ويتم ذلك بنقر الزر Define Ranges في صندوق الحوار أعلاه .

عند نقر زر OK نحصل على الجدول التالي .

```

* * * C R O S S T A B U L A T I O N * * *
$PAPER (tabulating 1)
by COLORED

```

		COLORED		
\$PAPER	Count			Row Total
		0	1	
BABYLON	1	1	4	5 71.4
JAMHORYA	2	2	0	2 28.6
AL IRAQ	2	2	2	4 57.1
QADISSYA	0	0	3	3 42.9
THOWRA	0	0	2	2 28.6
Column Total		3 42.9	4 57.1	7 100.0

Percents and totals based on respondents

ملاحظات :

1. يمكن الحصول على تكرارات كل صف أو عمود في جدول التقاطع إما كنسبة مئوية من الاستجابات responses أو كنسبة مئوية من المستجيبين respondents (أو الحالات Cases) ويتم ذلك بنقر زر Options في صندوق حوار Multiple Response Crosstabs ثم تأشير أحد الحالتين التاليتين في خانة Percentages Based on :

- Cases : لغرض الحصول على تكرار الصف أو العمود كنسبة مئوية من الحالات Cases أو المستجيبين Respondents (عددهم في هذا المثال هو 7) .
- Responses : لغرض الحصول على تكرار الصف أو العمود كنسبة مئوية من الاستجابات Responses (عدد الاستجابات في هذا المثال هو 16) .

2. يمكن الحصول على النسبة المئوية لكل خلية من خلايا جدول التقاطع بنقر زر Options في صندوق حوار Multiple Response Crosstabs ثم تأشير أحد الخيارات التالية في خانة Cell Percentages :

- Row : لعرض تكرار الخلية كنسبة مئوية من مجموع تكرارات الصف .
- Column : لعرض تكرار الخلية كنسبة مئوية من مجموع تكرارات العمود .
- Total : لعرض تكرار الخلية كنسبة مئوية من مجموع التكرارات الكلية .

مثال 3

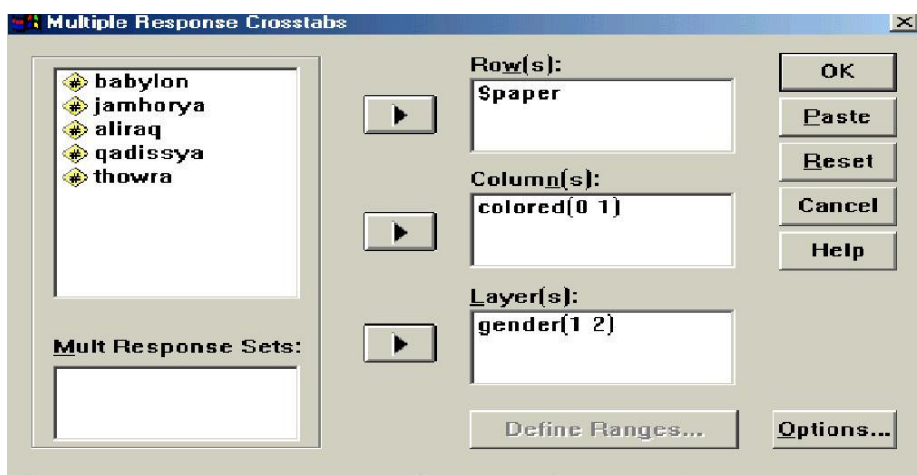
لنفس المثال السابق يطلب تكوين جدول تقاطع لكل من المستجيبين الذكور والإناث بصورة منفصلة. في هذه الحالة يتوجب إضافة متغير الجنس gender الى شاشة Data Editor للمثال السابق حيث أنه متغير ذو فئتين (1 يمثل الذكور و 2 يمثل الإناث وقد أعطينا عناوين القيم Value Label بحيث أن 1 = Male و 2 = Female) وتظهر الشاشة كما يلي :

Babylon	jamhorya	aliraq	qadissya	thowra	colored	gender
1	0	1	1	0	1	2
1	1	0	0	0	0	1
0	0	1	0	0	0	2
1	0	0	1	1	1	1
0	1	1	0	0	0	1
1	0	0	1	1	1	1
1	0	1	0	0	1	2

لتكوين هذا الجدول نحتاج الى مجموعة الاستجابات المتعددة \$paper التي تمثل تفضيل الصحف ، يجب أولاً تكوين هذه المجموعة بنفس الطريقة المذكورة في المثال الأول بأسلوب المتغيرات المتعددة ذات الفئتين Multiple dichotomy method (علماً أنه يمكن اعتماد أسلوب الفئات المتعددة Multiple Category Method).

لتكوين جدول التقاطع بحيث تكون فئات تفضيل الصحف في الصفوف وفئات تفضيل الألوان في الأعمدة وحسب فئتي الجنس نتبع الخطوات التالية :

من شريط القوائم اختر Crosstabs → Multiple Response → Analyze فيظهر صندوق حوار Multiple Response Crosstabs الذي نرتبه كما يلي :



لاحظ أننا حددنا الحدين الأعلى والأدنى لفئات المتغير Colored وكذلك فئات المتغير Gender بواسطة الزر Define Ranges بعد نقر أسم المتغير المعني . وعند نقر زر OK في صندوق الحوار أعلاه نحصل على المخرج التالي :

		COLORED		Row Total
Count		0	1	
\$PAPER	BABYLON	1	2	3 75.0
	JAMHORYA	2	0	2 50.0
	ALIRAQ	1	0	1 25.0
	QADISSYA	0	2	2 50.0
	THOWRA	0	2	2 50.0
	Column Total	2 50.0	2 50.0	4 100.0

\$PAPER (tabulating 1)

by COLORED

by GENDER

Category = 1 Male

Percents and totals based on respondents

لاحظ أن النسب المئوية مبنية على المستجيبين (الحالات) وأن عدد المستجيبين الذكور هو 4 (الجدول السابق) وأن عدد المستجيبين الإناث هو 3 (الجدول التالي) .

\$PAPER (tabulating 1)
 by COLORED
 by GENDER
 Category = 2 Female

		COLORED		
\$PAPER	Count.			Row Total
		0	1	
BABYLON		0	2	2 66.7
ALIRAQ		1	2	3 100.0
QADISSYA		0	1	1 33.3
Column Total		1 33.3	2 66.7	3 100.0

Percents and totals based on respondents