

جمهورية مصر العربية



معهد التخطيط القومى

سلسلة مذكرات خارجية

مذكرة خارجية رقم (١٦٢٩)

التحليل الإحصائى باستخدام

برنامج SPSS

إعداد

أ.د. ماجدة ابراهيم

د. زلفى شلبى

أ. أحمد عبد الباقي

أكتوبر ٢٠٠٥

جمهورية مصر العربية - طريق صلاح سالم - مدينة نصر - القاهرة - مكتب بريد رقم ١١٧٦٥

A.R.E Salah Salem St. Nasr City, Cairo P.O. Box:11765

٢	مقدمة
٤	الباب الأول : تشغيل برنامج SPSS والتعامل مع البيانات والملفات
٤	١-١ مراحل التحليل الإحصائي على الحاسب الآلى
٥	٢-١ طريقة تشغيل البرنامج
٦	٣-١ التعامل مع البيانات والملفات
٢١	٤-١ تحويل البيانات
٤٠	الباب الثانى : الأسلوب الإحصائى وعرض وتحليل البيانات
٤٠	١-٢ تحديد الأسلوب الإحصائى المناسب لتحليل البيانات
٤٨	٢-٢ الرسوم البيانية باستخدام برنامج SPSS
٥٦	٣-٢ الجداول التكرارية
٥٩	٤-٢ التحليل الوصفى للبيانات
٦٩	٥-٢ استكشاف البيانات
٧٤	الباب الثالث : اختبارات الفروض الإحصائية
٧٤	١-٣ الخطوات الأساسية فى اختبارات الفروض الإحصائية
٧٩	٢-٣ إجراء اختبار الفروض الإحصائية باستخدام برنامج SPSS
٧٩	١-٢-٣ اختبار T للمقارنة بين المتوسطات
٨٥	٢-٢-٣ اختبار F لتحليل التباين
٩١	٣-٢-٣ اختبار كا ^٢
٩٩	الباب الرابع : الارتباط والانحدار
٩٩	١-٤ الارتباط
١٠٣	٢-٤ الانحدار الخطى البسيط
١١٧	٣-٤ الانحدار الخطى المتعدد
١٢٦	الباب الخامس : إسقاط الفروض عند تقدير معالم نموذج الانحدار
١٢٦	١-٥ الارتباط الخطى بين المتغيرات المستقلة
١٣١	٢-٥ الارتباط الذاتى للأخطاء
١٤١	٣-٥ وجود أخطاء القياس والملاحظة فى المتغيرات التفسيرية
١٤٣	المراجع

أولاً : مقدمة

يعتبر البرنامج الجاهز SPSS (وهو اختصار للحروف الأولى من Statistical Package of Social Sciences) من أكثر البرامج الإحصائية استخداماً من قبل شريحة واسعة من الطلبة والباحثين في المجالات التربوية والاجتماعية والفنية والهندسية والزراعية والطبية في إجراء التحليلات الإحصائية اللازمة .

وقد بدأت شركة SPSS بإعداد هذا النظام الذي كان يعمل تحت نظام تشغيل Ms-Dos وقد تم تطويره ليعمل في بيئة نظام التشغيل Windows في عام ١٩٩٣ متلافياً بذلك الصعوبات التي كانت تواجه العاملين في هذا النظام في بيئة Ms-Dos ويتميز هذا البرنامج بالميزات التالية :

- ١- السهولة في إدخال البيانات حيث يقدم لنا نافذة (Data Editor) عبارة عن Spreadsheet به صفحة لتعريف المتغيرات وصفحة أخرى لإدخال وتعديل البيانات.
- ٢- السهولة في استخراج النتائج ، حيث تظهر النتائج في نافذة (SPSS viewer) بطريقة تمكن من الاختيار والتعديل ، كما يمكن طباعة هذه النتائج بطريقة محسنة عن طريق استخدام أنواع وأحجام ومواصفات عديدة لخطوط وأشكال النتائج المطبوعة وفي شكل جداول محورية pivot Tables ذات خطوط طولية وعرضية.
- ٣- العرض البياني للبيانات عن طريق رسم بياني غاية في الدقة وبأشكال وألوان متعددة (أعمدة -خطوط-دائرة ... الخ)
- ٤- الاستفادة من بيئة Windows بما تقدمه لنا من رموز icons صناديق حوارية Dialog boxes تسهل من تنفيذ الأوامر.
- ٥- إمكانية استخدام بيانات من برامج أخرى تعمل تحت Windows كما يمكن استدعاء البيانات من برامج الجداول الالكترونية مثل Excel وجداول قواعد البيانات مثل Access كما يمكن الاستفادة من النتائج في برامج أخرى.
- ٦- يقدم لنا البرنامج نافذة (Syntax Editor) يمكن عن طريقها كتابة الأوامر كما يمكن ترجمة الاختيارات من الصناديق الحوارية إلى أوامر عن طريق لصقها (Paste) في نافذة الـ syntax وتخزينها لاستخدامها لاحقاً.
- ٧- وجود وسيلة مساعدة (Statistics Coach) تساعد على التحليل الإحصائي المناسب لكل نوع من أنواع البيانات المختلفة.

ولقد غير برنامج SPSS وغيره من البرامج المناظرة الحياة لكثير من المشتغلين بالإحصاء وتحليل البيانات ومنهم الطلاب الذين يتعلمون الإحصاء والمعلمون الذين يدرسونها والباحثون الذين يطبقونها ومع ذلك فما زال هناك الكثيرون الذين يجدون في التعامل مع البرامج الإحصائية عملية صعبة وشاقة وخالية من أى متعة

وما زالوا يواجهون بصعوبات جمة وعراقيل تمنعهم من إتقان هذه المهارات الهامة ومن هذه العراقيل:-

١- ضعف الخلفية الإحصائية اللازمة للتعامل مع البرنامج تجعل المستخدم يحارب للوصول إلى هدفه دون جدوى، لأنه غير قادر على تحقيق شئ. واستخدام برنامج SPSS أو غيره من البرامج الإحصائية دون الإلمام بالأساليب الإحصائية عملية خطيرة تقود إلى نتائج خاطئة وغير سليمة. ولذلك يحتاج العمل في برنامج SPSS إلى خلفية إحصائية قوية تمكن المستخدم من التعامل براحة ويسر مع المفاهيم الإحصائية المختلفة.

٢- يمكن لبعض الباحثين التفكير في طرق عدة يحلون بها بياناتهم ولكن يقفون مترددين حائرين لا يدرون انسبها وأصلحها للبيانات التي لديهم، وبخاصة عندما يتعلق الأمر بمسلمات يجب استيفاءها حتى يمكن قبول النتائج. وحتى إذا استطاعوا أن يصلوا إلى القرارات السليمة والنتائج المطلوبة، فما زال أمامهم أن يعدوا تقريراً يفسرون به نتائجهم، ويكون متمشياً مع الطرق المتعارف عليها أكاديمياً في كتابة التقارير.

٣- كمية النتائج التي يعطيها البرنامج قد تكون كبيرة جداً مما يدفع الباحثين إلى التراجع أمامها وبخاصة عندما يحاولون تفسيرها بل وقد يتركون الأمر برمته ويحاولون البحث عن يستطيع مساعدتهم أمام هذا الخضم من الطلاسم المكتوبة سلفاً غير مفهومة.

٤- رغم السهولة الكبيرة لبرنامج SPSS إلا أن المستخدمين الجدد يجدونه صعباً ومعقداً للغاية. وعليهم أن يتعلموا كيفية إدخال البيانات في محرر البيانات وحفظها واسترجاعها، بل وإجراء بعض التعديلات ولكن المثابرة والجهد المتواصل في دراسة البرنامج وفهمه كثيراً ما يؤدي ثماره في الحصول على مهارة من أهم المهارات وهي استخدام الحاسب الآلى في تحليل بيانات البحوث.

ولهذه المشاكل السابقة وغيرها تم إعداد هذه المادة العلمية حتى تكون عوناً للباحثين في مختلف التخصصات وللمساعدتهم على تخطي العقبات التي تصادفهم أثناء تحليل البيانات وأثناء تفسير النتائج، والإصدار الذي استخدم في أمثلة هذه المذكرة SPSS ver 11.1.0

الباب الأول

تشغيل برنامج SPSS والتعامل مع البيانات والملفات

1-1 : مراحل التحليل الإحصائي باستخدام برنامج SPSS على الحاسب الآلي

نقطة البداية لأي عمل نقوم به في برنامج SPSS هي عادة محرر البيانات وهو عبارة عن جدول إلكتروني يستخدم لإدخال البيانات فيه وتحديد أسماء المتغيرات . وبعد الانتهاء من هاتين العمليتين يتم الانتقال إلى خطوات تحليل البيانات التي نريدها واختبار النتائج.

الخطوة الأولى : إدخال البيانات

الخطوة الأولى بطبيعة الحال هي إدخال البيانات وأخبار SPSS ما تمثله هذه البيانات وأسهل طريقة للقيام بذلك هي استخدام محرر البيانات وفيه:-

١- تدخل البيانات في صفوف وأعمدة محرر البيانات.

٢- تسمية المتغيرات التي تستخدم في تحليل البيانات.

الخطوة الثانية : تحديد التحليل الإحصائي

الخطوة التالية بعد إدخال البيانات هي إصدار التعليمات لبرنامج SPSS للقيام بالعمليات الإحصائية المرغوبة . هناك طريقتان لتنفيذ هذه الخطوة :-

١- الطريقة الأولى : هي طريقة التأشير والضغط على زر الفأرة Point-and-click Method وفي هذه الطريقة نقوم بالتحليل الذي نريده باستخدام الفأرة لفتح قوائم مسددة ومربعات حوار واختيار ما نريد منها . وهذه الطريقة أسهل الطريقتين لأنها لا تتطلب معرفة أية لغة من لغات البرمجة (Syntax) . فالبرنامج مكتوب بالفعل وجاهز ولكنه موجود في الخلفية وغير ظاهر للعين . ولكنه رهن إشارتك في أي وقت . وهذه الطريقة مريحة ومألوفة لمستخدمي بيئة النوافذ حيث أن الواجهة المستخدمة شبيهة بالواجهات الأخرى المستخدمة في برامج النوافذ.

٢- الطريقة الثانية : وهي التي يمكن أن نطلق عليها الطريقة اللغوية Syntax Method وتتضمن استخدام برنامج SPSS بالطريقة التقليدية والتي كانت مستخدمة في الماضي عندما كان البرنامج يعمل في بيئة Dos .

وعند استخدام هذه الطريقة نبدأ بفتح نافذة جديدة يطلق عليها محرر اللغة Syntax Editor ويكتب فيها التعليمات بلغة البرمجة الخاصة ببرنامج SPSS . وتتطلب هذه الطريقة تعلم هذه اللغة ، وإن كان هذا الأمر يبدو صعباً بعض الشيء في البداية ، إلا أن هناك مزايا عدة لاستخدام هذه الطريقة أقلها أن المستخدم يستطيع عمل أشياء بما غير متاحة في طريقة التأشير والضغط .

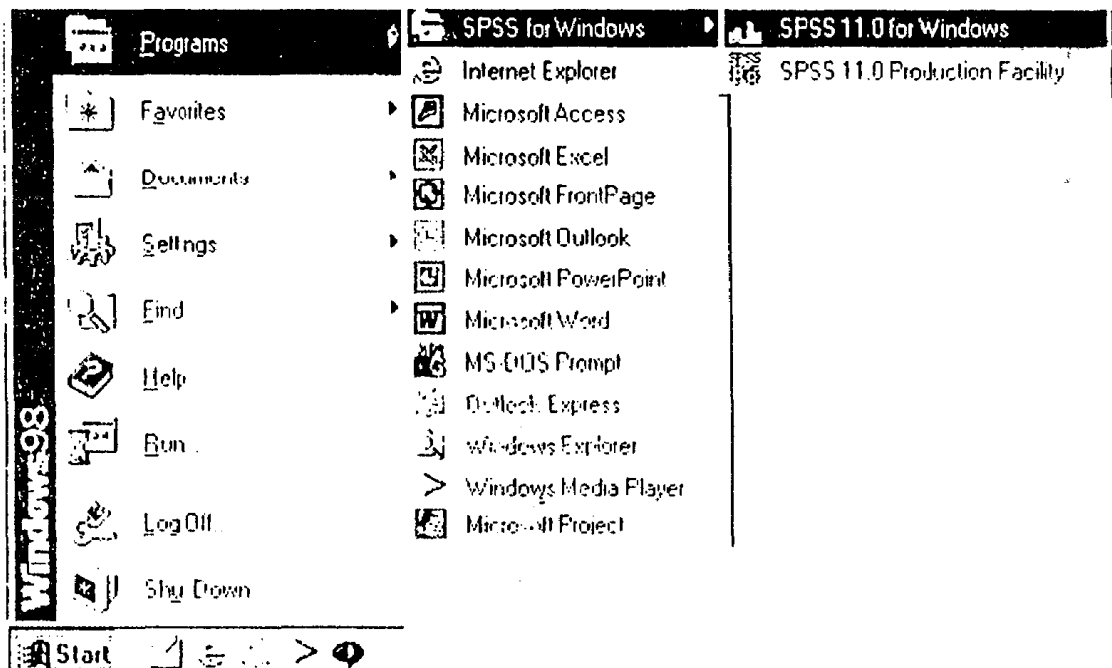
وتستخدم كلا الطريقتين عند مناقشة تحليل البيانات تحليلاً إحصائياً . ويمكن للمستخدم الاختصار على طريقة واحدة فقط منهما، وتمكنه أن ينتقل بين الطريقتين كيفما شاء.

الخطوة الثالثة: فحص المخرجات وتعديلها:-

بعد ادخال البيانات وتحليلها باستخدام إحدى الطريقتين السابقتين ، تظهر نافذة جديدة تحتوي على نتائج التحليل. ويمكن طباعة النتائج أو حفظها على القرص الصلب أو القرص المرن أو القرص المضغوط للعودة إليها في المستقبل .

٣-١ : طريقة تشغيل البرنامج

يتم تحميل النظام عن طريق نقر رمز Start ومنها يتم اختيار Programs ومنها الرمز SPSS for Windows ومنها SPSS 11.0 for Windows .



ويمكن إنشاء أيقونة خاصة بالبرنامج على سطح المكتب Desktop واستدعاء البرنامج منها مباشرة .

بعد تشغيل البرنامج تظهر (نافذة البيانات SPSS Data Editor) ، وهي عبارة عن نافذة على شكل صفوف وأعمدة Spreadsheet من النوع Tab Pages وهي تشمل على صفحتين:

صفحة تعريف المتغيرات (Variable View) و صفحة إدخال البيانات (Data View) .

تختلف محتويات هذه النافذة حسب الصفحة النشطة ، ففي صفحة تعريف المتغيرات (Variable View) تعبر الصفوف عن المتغيرات حيث يعبر كل صف عن متغير ، كما تعبر الأعمدة عن مواصفات هذه المتغيرات مثل اسم المتغير Name ، نوع المتغير Type ، حجم المتغير Width ... إلخ .

	Name	Type	Width	Deci	Label	Values	Missing	Columns	Align	Measure
1	id	Numeric	8	0	رقم الفرد	None	None	8	Right	Scale
2	name	String	21	0	الاسم	None	None	11	Left	Nominal
3	gender	Numeric	8	0	النوع	None	None	8	Right	Scale

أما إذا كانت الصفحة النشطة هي صفحة البيانات Data View فإن الأعمدة تعبر عن المتغيرات ، حيث يعبر كل عمود عن متغير Variable مثل رقم الفرد Id ، اسم الفرد Name ، نوعه Gender ... إلخ ، وكل صف يعبر عن حالة Case أى مفردة من مفردات التحليل كبيانات فرد معين .

	id	name	gender
1	1	1000000000	1
2	2	2000000000	2
3	3	3000000000	3

وتشتمل هذه النافذة على شريط القوائم ويشمل مجموعة القوائم التالية :

File, Edit, View, Data, Transform, Analyze, Graphs, Utilities, Window Help
وتستخدم هذه القوائم في تنفيذ أوامر النظام ، كما تشتمل على شريط الأدوات (Toolbar) وهو يستخدم لتنفيذ الأوامر شائعة الاستخدام في النظام .

١-٣ : التعامل مع البيانات والملفات

بعد تشغيل البرنامج ، ولإجراء التحليل الإحصائي على بيانات معينة بعد جمعها من مصادرها المختلفة ، يلزم تعريف هذه البيانات للبرنامج حتى يمكن إجراء التحليل الإحصائي عليها ، ويشمل ذلك ما يلي :

■ تعريف وتوصيف المتغيرات Define Variables

يستخدم في تعريف المتغيرات نافذة البيانات Data Editor ، ويشمل تعريف المتغيرات معلومات عن هذه المتغيرات ، ويمكن تعريف المتغيرات قبل أو بعد إدخال البيانات ، ولكن يفضل تعريفها قبل الإدخال .

وقبل البدء في تعريف المتغيرات يجب التأكد أن النافذة النشطة Data Editor وأن الصفحة النشطة هي صفحة تعريف المتغيرات Variable View حيث يظهر عنونها SPSS Data Editor ، وعنوان الملف Untitled مضاءً وهو عبارة عن Spreadsheet الصفوف تمثل المتغيرات Variables والأعمدة تمثل مواصفات هذه المتغيرات مثل النوع والحجم ... إلخ ويتم التعريف على النحو التالي :

File Edit View Data Transform Analyze Graphs Utilities Window Help										
	Name	Type	Width	Decimals	Label	Values	Missing	Columns	Align	Measure
1	id	Numeric	8	2		None	None	8	Right	Scale
2	name	String	30					8	Left	Nominal
3	b_date	Date				None		8	Right	Scale
4	gender	Numeric	8	2		1.00, 2.00, ...	None	8	Right	Scale

(١) اسم المتغير Variable name

أكتب اسم المتغير في عمود Name مع ملاحظة أن اسم المتغير لا يزيد عن ٨ خانات ، ولا بد أن تبدأ بحرف ، وغير مسموح بالمسافات والعلامات الخاصة مثل علامتي الجمع والطرح ، كما أن اسم المتغير لا يتكرر في الملف الواحد ويفضل أن يعبر اسم المتغير عن البيانات الداخلة فيه فمثلاً متغير اسم الفرد يسمى Name أو النوع يسمى Gender أو تاريخ الميلاد يسمى B_date ... إلخ ، أما إذا كانت البيانات عبارة عن أسئلة استبيان Questionnaire فيمكن أن تكون أسماء المتغيرات Q1, Q2, ...etc بأرقام أسئلة الاستبيان للتبسيط ويمكن شرح وتوصيف هذه المتغيرات في خانة Label .

(٢) نوع المتغير Variable Type

Type
Numeric ...

لتعريف نوع البيانات في هذا المتغير ، انقر النقاط الثلاث في خانة Type

حيث يظهر الصندوق الحوارى Variable Type التالى :

The dialog box titled 'Variable Type' has a list of options on the left: Numeric (selected), Comma, Dot, Scientific notation, Date, Dollar, Custom currency, and String. On the right, there are input fields for 'Width' (set to 8) and 'Decimal Places' (set to 0). At the bottom right are buttons for 'OK', 'Cancel', and 'Help'.

وهو يحدد الأنواع المختلفة للمتغيرات ، مع ملاحظة أن نوع المتغير يتحدد تلقائياً (by default) بأنه رقمى ، ويمكن تعريف نوع المتغير بأحد الأنواع التالية :

- Numeric رقمى: ويشمل التعريف العدد الكلى لخانات المتغير Width ، وعدد الخانات العشرية Decimal Places .

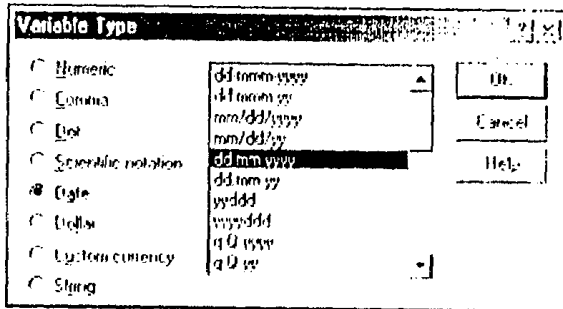
- Comma : متغير رقمى ، ولكن أرقامه تنفصل بفاصلة comma بين الملايين والألوف ، وبنقطة للأرقام العشرية مثلاً الرقم 134565.15 يمثل هكذا 134,565.15

- Dot : متغير رقمى ، ولكن أرقامه تنفصل بنقطة بين الملايين والألوف ، وبفاصلة comma للأرقام العشرية . مثلاً الرقم 134565.15 يمثل هكذا 134.365,15

- Scientific notation : متغير رقمى ، ولكن يظهر بالـ E Format أى الرقم مضروباً في ١٠ مرفوعة لأس معين (موجب في حالة البيانات الكبيرة ، وسالب في حالة البيانات الصغيرة) وذلك يستخدم في حالة تمثيل البيانات الكبيرة جداً أو الصغيرة جداً ، ويلاحظ أن البرنامج SPSS يمثل به الكثير من مخرجاته فمثلاً الرقم 14000000 يمثل 1.4 E+07 أى ١.4 مضروباً في 10^7 والرقم 0.000023 يمثل هكذا 2.3E-5 أى 2.3 مقسوماً على 10^5

- Custom currency : متغير تظهر به رموز عملات ، يمكن اختيارها

- Dollar : متغير رقمي قيمته تشتمل على علامة الدولار .



- Date متغير يمثل التاريخ ، وفي حالة

اختياره يسمح بعدة أشكال معينة للتاريخ

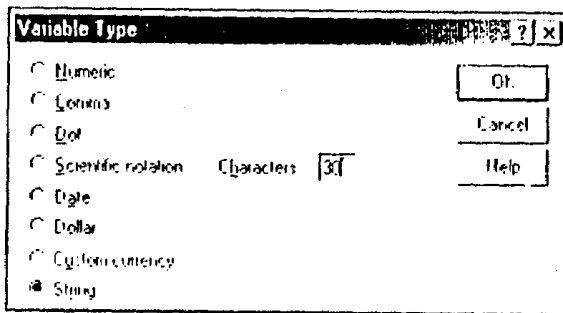
يمكن الاختيار من بينها ، حيث تمثل dd

اليوم ، تمثل mm الشهر ، تمثل yy السنة

وكل من اليوم والشهر والسنة يظهر في خانتي فمثلاً إذا اخترنا الشكل dd.mm.yy فإن

التاريخ ٤ أغسطس ٢٠٠٣ يسجل هكذا ٠٣/٨/٣ ويظهر بالشكل ٠٣,٠٨,٠٣ ويلاحظ

أن المد Date لا يحتاج إلى تحديد حجم - الأقل حيث أنه يخزن تلقائياً في ٨ خانات .



- String متغير حرفي: ويستخدم

للبيانات الحرفية مثل أسماء أفراد الأسرة

التي لا تجري عليها عمليات حسابية

وفيه يتحدد عدد حروف هذا المتغير في

خانة Character .

ملاحظة

الخانات Width , Decimal خاصة بالبيانات الرقمية وفيها يتم تحديد عدد خانات هذا

المتغير ، وعدد الخانات العشرية فمثلاً الرقم ١٣٣,٨ إلى Width=5 Decimal=1 ، ويلاحظ

دائماً أن يكون حجم المتغير Width بحيث يتسع لأكثر رقم فيه وإلا فإن الرقم سيمثل بالـ E

format ، كما يجب أيضاً أن تكون عدد الخانات العشرية حسب الدقة المطلوبة وإلا سيؤدي إلى

تقريب الأرقام .

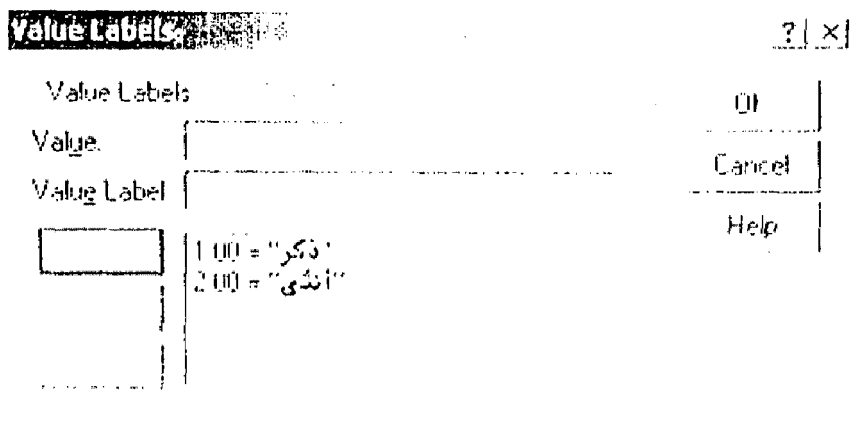
(٣) تعريف وصف المتغير Label

في خانة Label اكتب وصفاً يشرح كل متغير ، باللغة العربية أو الإنجليزية حسب البيانات ، وذلك لسهولة تفسير البيانات والنتائج فمثلا في استمارة الاستبيان السؤال الأول يعطى اسم المتغير q1 أما الـ Label فيكون "النوع" ، السؤال الثاني يعطى اسم المتغير q2 أما الـ Label فيكون "صلة القرابة برب الأسرة"

(٤) تعريف أدلة المتغير Values

كما نعلم فإن معظم إجابات أسئلة الاستبيان تأخذ فمثلاً q1 أى السؤال الأول وهو النوع يرمز الكود ١ إلى ذكر ، يرمز الكود ٢ إلى أنثى .

ولكى يتم تعريف هذه البيانات المكودة انقر الثلاث نقاط في عمود Values أمام المتغير المراد تعريفه وليكن متغير gender فيظهر الصندوق الحوارى Value Labels :



أمام خانة Value اكتب الرقم ١ ، أمام خانة Value Label : اكتب التعريف "ذكر" ثم انقر أمر Add ، أمام خانة Value اكتب الرقم ٢ أمام خانة Value Label : اكتب "أنثى" ثم انقر أمر Add .

ملاحظة : يمكن إعادة تعريف هذه الأكواد أو إلغاؤها عن طريق الأوامر

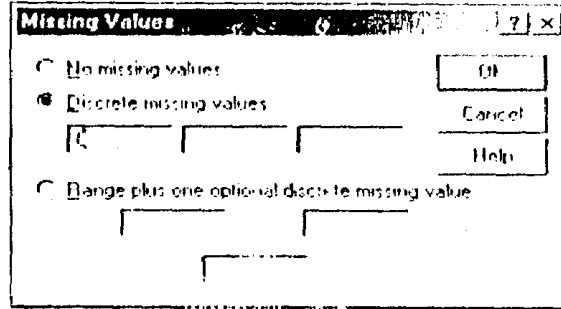
Remove , Change

(٥) تحديد القيم المفقودة Missing Values

يمكن تحديد عدة قيم للقيم المفقودة (وهي البيانات غير المبينة) والتي لم يتم الحصول عليها ، أو البيانات التي لا تنطبق .

لتعريف هذه البيانات المفقودة انقر الثلاث نقاط في عمود Missing فيظهر الصندوق

الحوارى Missing Values الآتى :



ويوجد به الاختيارات التالية :

No Missing Value

لا يوجد قيم مفقودة

Discrete missing values القيم المفقودة لها رموز مقطعة مثل 99 , 9 وهكذا

Discrete missing values Range plus one discrete

يسمح بتعريف مدى متصل مع تعريف قيمة واحدة متقطعة مثل من 100 إلى

105 ، 999 حيث تكتب 100 في الحانة Low ، 105 في الحانة High ، 999 في حانة

Discrete value.

ويعطى البرنامج هذه الإمكانيات في القيم المفقودة لأن أنواعها قد تتعدد فيلزم تحديد كود

لكل نوع ، فهناك قيم مفقودة Missing لأن مصادرها غير متاحة وأخرى لأن السؤال لا

ينطبق Not Applicable على هذه الحالة ، والثالثة لأن الفرد امتنع عن الإجابة Not

Answer

(٦) تحديد عرض العمود (المتغير)

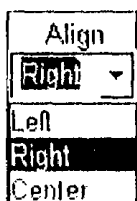
في العمود Column يتم تحديد عرض العمود في طريقة إظهاره فقط على الشاشة ، أو في

الطباعة ، أما السعة التخزينية للمتغير فتحدد في العمود Width.

(٧) تحديد محاذاة المتغير Align

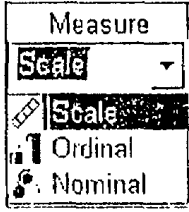
في العمود Align يتم محاذاة المتغير إما في يمين العمود أو إلى الوسط أو في

اليسار .



(٨) تحديد طريقة قياس المتغير Measure

حيث يحدد ثلاثة أساليب لقياس المتغير :



أ- اسمي Nominal ، حيث أن الأرقام لا تدل على كميات ، وإنما ترمز فقط إلى أكواد ، مثل متغير الحالة الزوجية (١ أتزوج ، ٢ متزوج ، ٣ مطلق ، ٤ أرمل) ومتغير النوع (١ ذكر ، ٢ أنثى) ، ويستخدم مع هذه البيانات أنواع معينة من المقاييس مثل المتوال Mode ، كما Chi Square .

ب- متغير ترتيبي Ordinal ، ويستخدم عندما تكون البيانات في شكل تسريبي من الأدنى إلى الأعلى مثل الحالة التعليمية (٠ أقل من السن ، ١ أمي ، ٢ يقرأ ويكتب ، ٣ إبتدائي ، ٤ إعلادي ، ٥ ثانوي ، ٦ فوق المتوسط وأقل من جامعي ، ٧ جامعي ، ٨ عليا) ، ويوجد أيضاً أنواع معينة من التحليل تناسب هذا النوع من البيانات مثل معامل ارتباط سبيرمان للرتب .

ج- رقمي Scale ، حيث أن الرقم يدل على قيمة للمتغير مثل العمر ، الوزن ، الدخل ... الخ

ملاحظة

في حالة تشابه خصائص المتغيرات في Type, Width, Values...etc يمكن استخدام أوامر النسخ Copy واللصق Paste لنسخ الخصائص المتشابهة.

▪ إدخال البيانات

كيفية إدخال البيانات عن طريق النافذة (Data Editor)

❖ يمكن إدخال البيانات قبل توصيف المتغيرات وفي هذه الحالة يحدد البرنامج للمتغيرات أسماء

var00001, var00002, ... etc إلى أن يتم توصيف هذه المتغيرات

❖ المتغيرات Variables تدخل أعمدة ، والحالات تدخل صفوف .

ملاحظة

يمكن إظهار Value Labels في الخلايا بدلاً من الأكواد عن طريق اختيار Value Labels

من القائمة View فمثلاً يمكن إظهار متغير النوع ذكر ، أنثى بدلاً من ١ ، ٢ .

مثال : أدخل البيانات التالية ، وهي عبارة عن بيانات افتراضية لمجموعة من العاملين بإحدى الجهات :

الجنس	الحالة الزوجية	الحالة التعليمية	العمر	النوع	رقم الفرد
أعمال فنية	متزوج	فوق المتوسط	48	ذكر	1
أعمال فنية	متزوج	فوق المتوسط	40	أنثى	2
أعمال كتابية	أعزب	ثانوى	17	ذكر	3
وظائف معاونة	أعزب	إعدادى	15	أنثى	4
أعمال تخصصية	متزوج	جامعى	56	ذكر	5
وظائف معاونة	متزوج	أمى	50	أنثى	6
أعمال تخصصية	مطلق	جامعى	22	ذكر	7
أعمال فنية	متزوج	ثانوى	20	أنثى	8
وظائف معاونة	متزوج	إعدادى	17	أنثى	9
أعمال كتابية	أرمل	فوق المتوسط	35	ذكر	10

■ يتم تسمية المتغيرات كالاتى :

ID رقم الفرد

Gender النوع

Age العمر

Educat الحالة التعليمية

marriage الحالة الزوجية

Work المهنة

■ يتم تكوين المتغيرات الوصفية كالتالى :

Gender النوع

١ ذكر

٢ أنثى

الحالة التعليمية Educat

- ١ أمي
- ٢ يقرأ ويكتب
- ٣ مؤهل متوسط
- ٤ فوق المتوسط
- ٥ جامعي

الحالة الزوجية marriage

- ١ أعزب
- ٢ متزوج
- ٣ مطلق
- ٤ أرمل

المهنة Work

- ١ أعمال تخصصية
- ٢ أعمال فنية
- ٣ أعمال كتابية
- ٤ وظائف معاونة

▪ بعد تكويد البيانات ، يظهر جدول البيانات بالشكل التالي :

ID رقم الفرد	Gender النوع	Age العمر	Educat الحالة التعليمية	marriage الحالة الزوجية	Work المهنة بالتفصيل
1	1	48	4	2	2
2	2	40	4	2	2
3	1	17	3	1	3
4	2	15	2	1	4
5	1	56	5	2	1
6	2	50	1	2	4
7	1	22	5	3	1
8	2	20	3	2	2
9	2	17	2	2	4
10	1	35	4	4	3

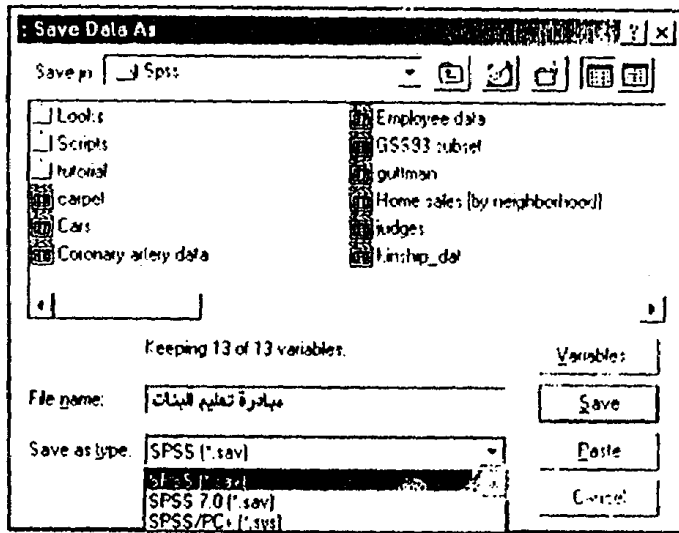
■ يتم توصيف المتغيرات كما سبق أن تعرفنا في صفحة Variable View بالتوصيف التالي :

Name	Type	Width	Decimals	Label	Values	Missing	Column	Align	Measure
ID	Numeric	3	0	رقم الفرد	None	None	3	Right	Nominal
Gender	Numeric	1	0	النوع	١- ذكر ٢- أنثى	None	3	Right	Nominal
Age	Numeric	3	0	العمر	None	None	3	Right	Scale
Educat	Numeric	1	0	الحالة التعليمية	١- أمي ٢- يقرأ ويكتب ٣- موهل متوسط ٤- فوق المتوسط ٥- جامعي	None	3	Right	Scale
Marriage	Numeric	1	0	الحالة الزوجية	١- أعزب ٢- متزوج ٣- مطلق ٤- أرمل	None	2	Right	Ordinal
Work	Numeric	1	0	المهنة	١- أعمال تخصصية ٢- أعمال فنية ٣- أعمال فنية ٤- وظائف معانة	None	2	Right	Nominal

■ بعد ذلك يتم إدخال البيانات في صفحة إدخال البيانات (Data View) فتظهر

النافذة بالشكل التالي :

test.sav - SPSS Data Editor							
File Edit View Data Transform Analyze Graphs Utilities Window Help							
1: id							
	id	gender	age	educat	marriage	work	
1	1	1	43	4	2	2	
2	2	2	40	4	2	2	
3	3	1	17	3	1	3	
4	4	2	15	2	1	4	
5	5	1	56	5	2	1	
6	6	2	50	1	2	4	
7	7	1	22	5	3	1	
8	8	2	20	3	2	2	
9	9	2	17	2	2	4	
10	10	1	35	4	4	3	



حفظ البيانات

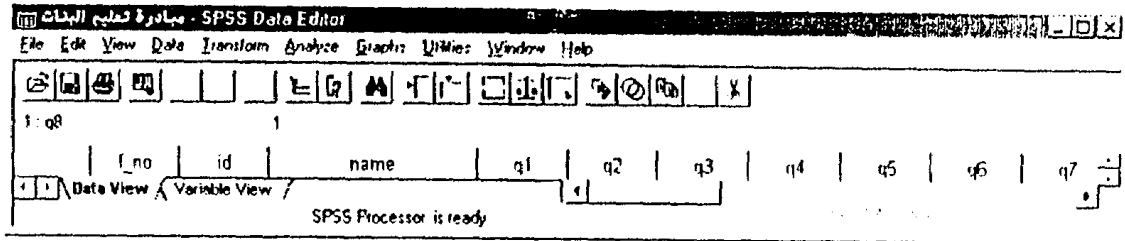
نجد أن البرنامج يخصص لهذه البيانات الاسم Untitled ، ولحفظ هذه البيانات وتخصيص اسم للملف ، نختار الأمر Save من القائمة File (أو الضغط على أيقونة  Save) فيظهر الصندوق الحوارى Save Data As كما في الشكل المقابل :

يتم كتابة اسم الملف في خانة File Name ، ويتم تحديد وسيط التخزين ، الدليل من خانة Save in .

ملاحظات

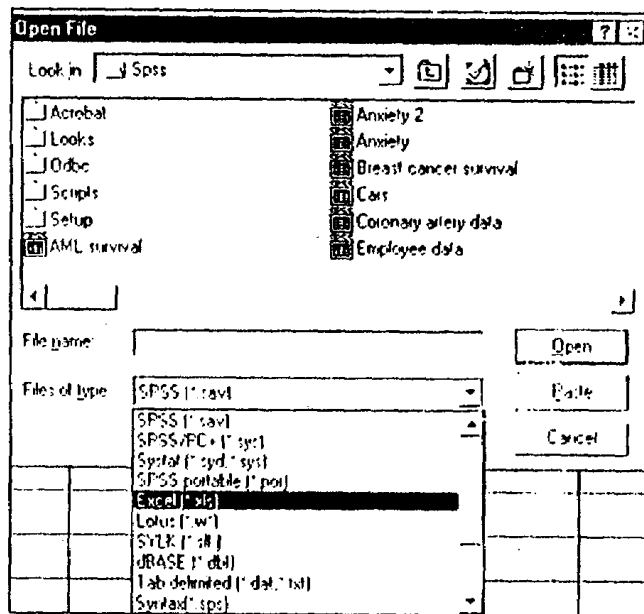
- عند الحفظ لأول مرة يفتح الصندوق الحوارى Save Data As ليسمح بكتابة اسم الملف ، ومكان حفظ الملف ، أما عند تكرار الحفظ باختيار أمر Save من القائمة File فيستمر الحفظ بدون فتح الصندوق الحوارى .
- إذا أردت أن تحتفظ بنسخة من الملف باسم آخر أو في مكان آخر اختر أمر File Save As → سيتم فتح الصندوق الحوارى Save Data As السابق .
- البرنامج يحفظ البيانات تلقائياً في الدليل Program Files\SPSS ، وهذا الدليل ي حذف إذا تم حذف برنامج SPSS من على الجهاز مما يؤدي إلى فقد البيانات ولذلك يفضل الحفظ في دليل آخر ينشئه مستخدم البرنامج .
- يفضل الاحتفاظ بنسخ من البيانات Data على وسيط تخزين خارجى مثل الأقراص المرنة Floppy diskettes أو الاسطوانات CD-Rom .

- ملفات بيانات البرنامج SPSS يعطى لها البرنامج الامتداد SAV ، وإذا أردنا أن نحفظ الملف بنوع آخر من أنواع البيانات كـ Excel فيمكن اختيار هذا النوع من القائمة Save as Type في الصندوق الحوارى .Save Data As
- بعد حفظ الملف ، يظهر اسمه في شريط العنوان كما يلي :



التعامل مع بيانات من أنظمة أخرى

على الرغم من سهولة استخدام Data Editor كنافذة للإدخال والتعديل والتعامل مع البيانات في نظام SPSS إلا أنه أيضاً يتيح فتح الملفات التي تم تسجيلها بأنظمة أخرى مثل dBase ، Excel ، على النحو التالى :



(١) من قائمة File اختر أمر Open ومنها اختر Data يتم فتح الصندوق الحوارى Open File.

(٢) من خانة Files of type تظهر قائمة بأنواع الملفات التي يتعامل معها البرنامج مثل Excel وهي الملفات التي امتدادها xls ، وملفات قواعد البيانات التي امتدادها dbf يتم تحديد نوع الملفات المطلوب فتحها ، وذلك بعد تحديد وسيط ودليل التخزين ، ثم نختار أمر Open يتم فتح الملف في نافذة البيانات Data Editor .

تعديل البيانات Editing Data

بعد إدخال البيانات أو نقلها من أى نظام آخر ، قد نحتاج إلى إجراء تعديلات على هذه البيانات ، وتتيح نافذة البيانات إمكانيات كبيرة في عمليات التعديل المختلفة كالآتي :

(١) استبدال البيانات

قد يوجد خطأ في البيانات المسجلة في بعض الخلايا ، لنختار الخلية ، ونجرى التصحيح

(٢) نسخ ونقل البيانات

يمكن قص Cut ، أو نسخ Copy البيانات ثم لصقها Paste طبقاً للحالات التالية :

- يمكن نقل أو نسخ خلية واحدة إلى خلية أو مجموعة خلايا .

- يمكن نقل أو نسخ صف (Case) إلى صف أو مجموعة صفوف .

- يمكن نقل أو نسخ عمود (Variable) إلى عمود أو مجموعة أعمدة .

ولتنفيذ ذلك يتم اختيار الخلايا المطلوب نسخها أو نقلها ويتم اختيار Cut من قائمة Edit في

حالة النقل ، أو Copy في حالة النسخ ، ثم تحديد الخلايا المراد نقل أو نسخ البيانات إليها ،

واختيار أمر Paste من القائمة Edit .

(٣) إدراج أو حذف الصفوف

يمكن إدراج صفوف (Cases) في أى مكان في نافذة البيانات ، حيث يتم اختيار صف أسفل

الصف المراد إدراجه ثم من قائمة Data نختار الأمر Insert case ، ولكي نحذف صف نختار

الصف المراد حذفه عن طريق نقر رقم الصف ، ثم اختيار أمر Cut من قائمة Edit .

(٤) إدراج أو حذف أعمدة

بنفس الطريقة السابقة يمكن إدراج أعمدة (Variables) في أى مكان في نافذة البيانات حيث

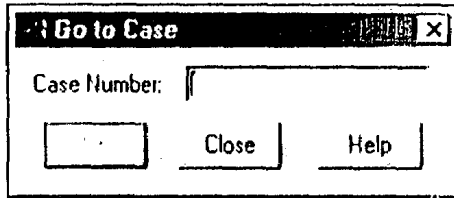
يتم اختيار متغير بعد المتغير المراد إدراجه ثم من قائمة Data نختار الأمر Insert Variable ،

ولكي نحذف متغير معين نختار المتغير المراد حذفه عن طريق نقر اسم المتغير ، ثم اختيار أمر Cut

من قائمة Edit

(٥) الانتقال إلى صف معين

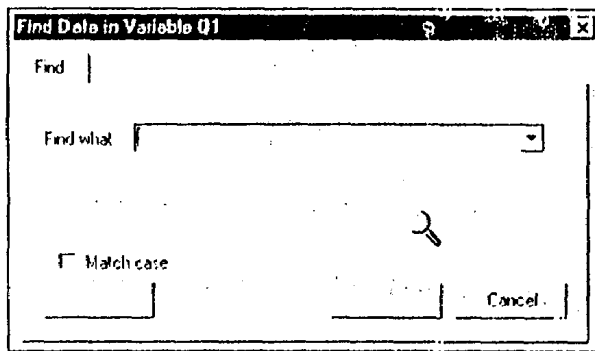
يمكن الانتقال خلال نافذة البيانات عن طريق الـ Scrolling bar فيمكن الانتقال صف لأعلى أو لأسفل عن طريق نقر أسهم التصفح ، كما يمكن الانتقال شاشة لأعلى أو لأسفل عن طريق النقر أعلى أو أسفل مربع التصفح ، وللانتقال إلى صف (أو حالة) معينة نتبع الآتي :



من قائمة Data اختر أمر Go to Case فيظهر الصندوق الحوارى Go to Case وأمام خانة Case Number أكتب رقم الصف المطلوب .

(٦) البحث عن بيانات

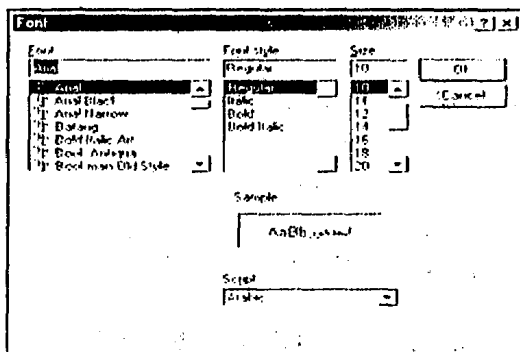
يمكن البحث عن بيانات معينة باتباع ما يلي :
يتم تحديد المتغير المراد البحث عن قيمة معينة في بياناته ، ومن قائمة Edit اختر أمر Find



أمام خانة Find what أكتب البيان المراد البحث عنه ، يتم تحديد أول قيمة تتفق مع قيمة البحث ثم يتم تحديد القيمة الثانية ، وهكذا حتى ينتهى البحث .

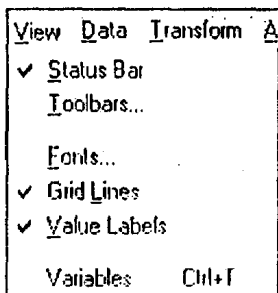
طرق العرض فى نافذة إدخال البيانات View

يمكن التحكم فى شكل البيانات المعروضة فى شاشة الإدخال ، وذلك لتحسين الشكل كالاتى



من القائمة View يتم اختيار :

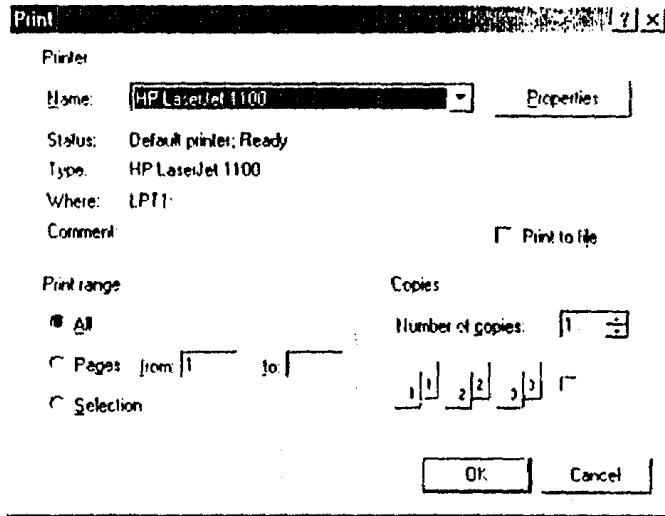
(١) Fonts وذلك للتحكم فى الخط المعروض من حيث حجمه ، ونوعه ومواصفاته . يتم اختيار Arabic من خانة script إذا كانت هناك أدلة باللغة العربية



(٢) يمكن تنشيط الاختيار Gridlines لإظهار الخطوط الطولية والعرضية ، أو عدم تنشيطه لإخفائها .

(٣) يمكن تنشيط الاختيار Value Labels لإظهار الوصيف Category للبيانات المكونة ، أو عدم تنشيطه لإخفاء الوصيف وإظهار الأكواد .

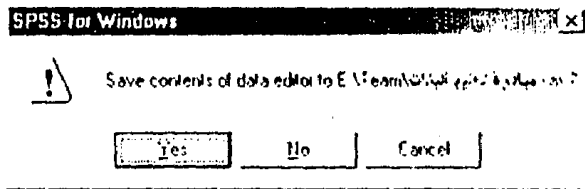
طباعة البيانات



يمكن طباعة البيانات بالكامل أو جزء منها ، أو صفحات معينة كما يمكن تحديد عدد النسخ عن طريق اختيار أمر Print من القائمة File فتظهر نافذة Print المقابلة . ويمكن تحديد مدى الطباعة من خانة Print range على النحو التالي :

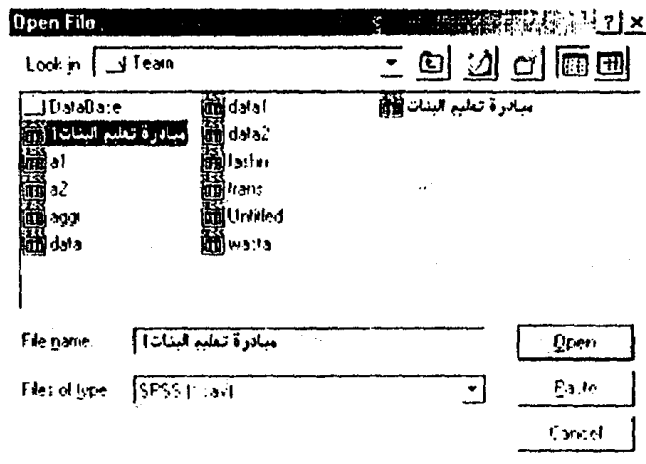
- All لطباعة كل المستند .
- Pages لطباعة صفحات متصلة ، يتم تحديد رقم صفحة البداية أمام خانة from : ورقم صفحة النهاية أمام خانة to .
- Selection لطباعة مدى تم تحديده من قبل .

الخروج من البرنامج



لكي يتم الخروج من البرنامج اختر أمر Exit من القائمة File ، تظهر رسالة تحذيرية ترشد إلى حفظ الملفات ، اختر أمر Yes للحفظ .

ملاحظة



- ١- بعد تشغيل البرنامج ، اختر الأمر File → Open → Data نافذة البيانات يظهر الصندوق الحوارى Open File .
- ٢- حدد مكان تخزين الملف ، واسم

الملف ، ثم اختر الأمر Open يتم فتح الملف في نافذة البيانات .

٤-١: تحويل البيانات

في بعض الأحيان يمكن إجراء التحليلات الإحصائية على البيانات الخام ، ولكن في الغالب نحتاج إلى إجراء بعض التحويلات على البيانات مثل إيجاد متغيرات جديدة عبارة عن علاقة بين المتغيرات الموجودة باستخدام صيغ ومعادلات معينة ، كما يمكن أن يكون هذا التحويل مقترناً بشرط أو غير مشروط ، كما نحتاج إلى اختيار حالات معينة لإجراء التحليل وليس كل الحالات .

حساب قيم متغير بناءً على دالة معينة Compute Variable :

يمكن حساب قيم متغير وإحلال القيم الجديدة محل القيم القديمة لنفس المتغير ، كما يمكن وضع نتيجة الحساب في متغير جديد ، ولإجراء ذلك نتبع الآتي :

- (١) من قائمة Transform اختر أمر Compute ، يظهر الصندوق الحوارى Compute .
- (٢) أمام خانة Target Variable أكتب المتغير الجديد المراد وضع نتيجة الحساب فيه ، أو اسم المتغير المراد تحويله إذا أردنا أن يتم الحساب في نفس المتغير .
- (٣) أكتب الدالة المراد إجراؤها (أى الطرف الأيسر للعملية الحسابية) مستعيناً بالـ Functions الموجودة في قائمة Functions والتعبيرات الحسابية والمنطقية الموجودة .
- (٤) اختر أمر OK .

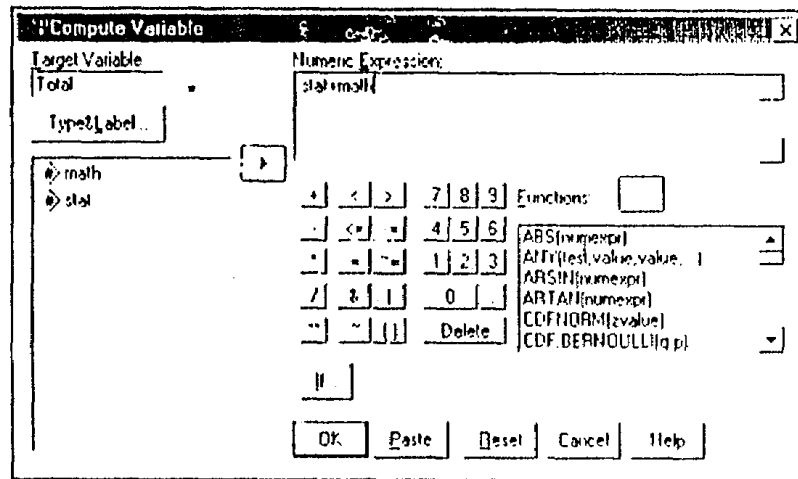
مثال : المثال التالى عبارة عن درجات عينة من الطلبة في مادتي الرياضة والإحصاء ، ونريد أن نضيف متغيراً آخر عبارة عن مجموع المادتين :

٦٠	٦٩	٦٧	٥٠	٦٣	٦١	٥٢	٥٧	٦٣	٦٧	الرياضة (Math)
٧٣	٦٥	٦٤	٦٨	٦٧	٢١	٧٤	٦١	٩٢	٥٨	الإحصاء (Stat)
١	٠	١	١	١	١	١	١	١	٠	النوع (Sex)

- (١) سجل المتغيرات في نافذة إدخال البيانات باسم sex , math , stat حيث ترمز stat لدرجة الإحصاء ، math لدرجة الرياضيات ، sex للنوع ، حيث الكود ٠ يشير إلى أنثى ، الكود ١ يشير إلى ذكر .

(٢) من القائمة Transform اختر أمر Compute فيظهر الصندوق الحوارى

Compute Variable أمام خانة Target Variable أكتب اسم المتغير الجديد Total



٣) أمام خانة Numeric Expression أكتب العملية المطلوب إجراؤها

Math + Stat

٤) يمكن تحديد Type & Label للمتغير الجديد عن طريق اختيار Type & Label

٥) يمكن استخدام Function من الدوال الموجودة تحت خانة Functions

Total = sum(math , stat)

بعد اختيار أمر OK تظهر نافذة البيانات ويضاف إليها متغير جديد على النحو التالي :

	math	stat	sex	total				
1	67	58	0	125				
2	63	92	1	155				
3	57	61	1	118				
4	52	74	0	126				
5	61	21	0	82				
6	63	67	0	130				
7	50	68	0	118				
8	67	64	1	131				
9	69	65	0	134				
10	60	73	1	133				

إذا أردنا أن نحفظ بالمتغيرات المحسوبة في ملف البيانات فينبغي حفظ هذا الملف ، أما إذا

أغلقنا بدون حفظ فلن يتم حفظ هذه المتغيرات المحسوبة بالأمر Compute .

ملحوظة :

إذا تم التعديل في البيانات ، يجب أن يعاد الحساب مرة ثانية ، حيث أنه لا يتم تحديث البيانات في المتغيرات المحسوبة تلقائياً .

يمكن استخدام الـ Arithmetic Operators التالية :

+ للجمع (Addition) ، - للطرح (subtraction) ، * للضرب (Multiplication) ، / للقسمة (Division) ، ** للأس (Exponential) ، كما يمكن استخدام الأقواس للتحكم في أولوية تنفيذ العمليات الحسابية

يراعى الأولوية التالية في تنفيذ العمليات الحسابية :

(١) فك الأقواس .

(٢) رفع الأس

(٣) الضرب أو القسمة أيهما أسبق

(٤) الجمع أو الطرح أيهما أسبق .

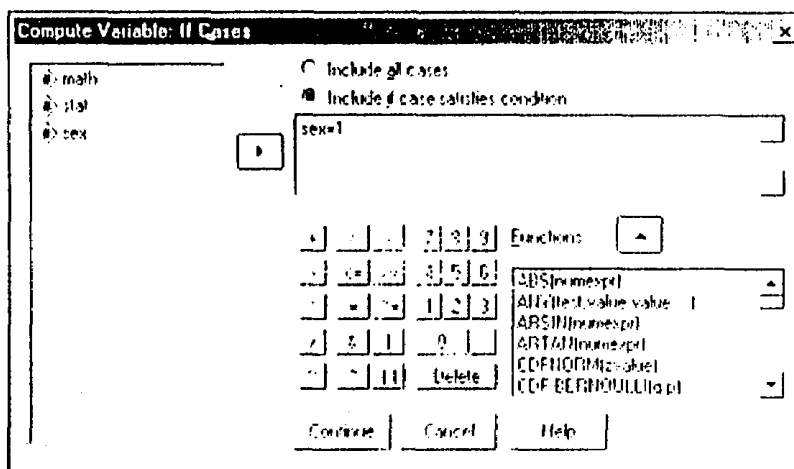
التحويل المشروط للمتغيرات Compute variable if case :

التحويل المشروط للمتغيرات يسمح باختيار حالات معينة لإجراء التحويلة المطلوبة عليها ، نتيجة شرط معين ، إذا كانت نتيجة الشرط True تجرى التحويلة ، إذا كانت نتيجة الشرط False لا تجرى التحويلة .

يتم تنفيذ ذلك بنقر If من الصندوق الحوارى Compute Variable .

في المثال السابق يمكن اختيار عينة الذكور (male) فقط دون الإناث لذلك يتم تنشيط If من

الصندوق الحوارى Compute



واختيار Include if cases satisfies condition ثم كتابة الشرط sex=1 ثم اختيار Continue ، يظهر هذا الشرط أمام خانة If ... في الصندوق الحوارى Compute Variable ثم اختيار أمر OK يتم حساب الـ total لبيانات الذكور فقط .

تستخدم الـ Relational Operators التالية :

< أكبر من

<= أكبر من أو يساوى

> أصغر من

>= أصغر من أو يساوى

= لا يساوى

كما تستخدم الـ Logical operators التالية :

Or | And & Not !

وتكون أولوية التنفيذ Not أولاً ، ثم And ، ثم Or ، وتستخدم في حالة تعريف أكثر من شرط. بعد تنفيذ Compute على متغير معين ، يمكن تعريف Type & Label للمتغير الذى يحتوى نتيجة الحساب .

الدوال Functions :

يقدم البرنامج العديد من الدوال Functions التى يمكن استخدامها فى الـ Compute منها :

الوظيفة	اسم الدالة
القيمة المطلقة	ABS
أخذ الرقم الصحيح وبتز الكسر	TRUNC
الجذر التربيعى	SQRT
اللوغاريتم للأساس ١٠	LG10
اللوغاريتم للأساس (E)	LN
جيب الزاوية (جا)	SIN
جيب تمام الزاوية (جتا)	COS
الرقم للأساس (E)	EXP

ويراعى القيم المفقودة Missing Values في العمليات الحسابية ، لأن الـ Missing Values حينما تدخل في العمليات الحسابية تكون النتيجة Missing .

إعادة توكويد المتغيرات : Recoding Variables

قد يوجد متغير ونحتاج إلى تقسيمه إلى فئات ، وذلك عن طريق إعطاء قيم جديدة لهذا المتغير تسمى Recoding ، وقد يكون ذلك التوكويد بتغيير قيم المتغير نفسه إلى الأكواد الجديدة ، أو يكون التوكويد داخل متغير جديد .

التوكويد داخل نفس المتغير : Recode into same variable

تحل الأكواد الجديدة محل القيم الفعلية :

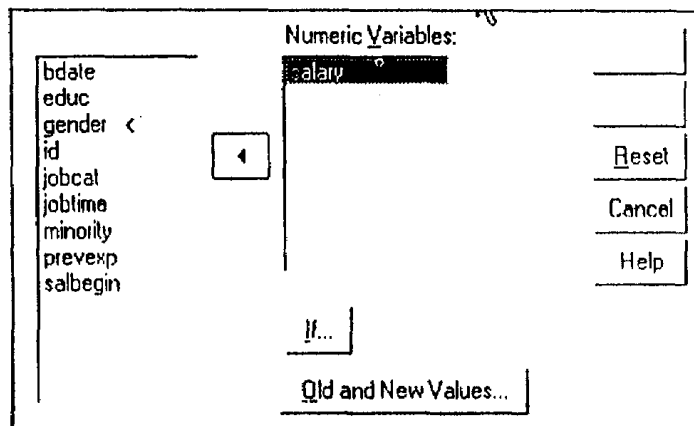
مثال : يمكن عمل Recoding للمتغير salary (الرواتب) في الملف Employee data.sav بتقسيمه إلى فئات .

ولتوكويد المتغير بتغيير القيم الفعلية يتم اتباع الآتى :

(١) افتح الملف Employee data.sav

(٢) من قائمة Transform اختر أمر Recode بنسدل اختياران ، اختر Into Same

Variables يظهر الصندوق الحوارى التالى :



(٢) انقر Old and New Values يظهر الصندوق التالى :

Recode into Same Variables: Old and New Values

Old Value

☐ Value: _____

☐ System-missing

☐ System- or user-missing

☐ Range: _____

☐ Range: _____

☒ Range: _____ through highest

☐ All other values

New Value

☒ Value: _____

☐ System-missing

Old -> New:

Lowest thru 9999	-> 1
10000 thru 19999	-> 2
20000 thru 29999	-> 3
30000 thru 39999	-> 4
40000 thru 49999	-> 5
50000 thru 59999	-> 6
60000 thru Highest	-> 7

Continue Cancel Help

(٢) سجل القيمة الفعلية للمتغير (أرالمدي المعين) المراد إعادة تكويده في خانة Old Value ،

ثم سجل الأكواد الجديدة في خانة New Value ، ثم انقر أمر Add .

أثناء تسجيل قيم المتغير القديمة Old يمكن الاستعانة بالآتي :

Range : وذلك في حالة تحديد مدة من البيانات مثل 10000 thru 19999

Lowest through : مثل Lowest thru 9999 أى من أقل قيمة إلى 9999 ، وذلك لتشمل

جميع القيم الدنيا

Through highest : مثل 6000 thru Highest أى من 6000 إلى أعلى قيمة ، وذلك

لتشمل جميع القيم العليا

(٣) بعد الانتهاء من تسجيل الأكواد انقر أمر Continue ، ثم اختر أمر OK يتم استبدال قيم

المتغير الأصلية بالقيم المكودة ، وفي حالة الحفظ تختفى القيم الأصلية للمتغير وتفقد نهائياً .

التكويد في متغيرات جديدة Recode into Different Variables :

أحياناً نحب أن نكود المتغيرات مع الاحتفاظ بالقيم الفعلية لها ، لأن تكويد المتغيرات رغم أنه

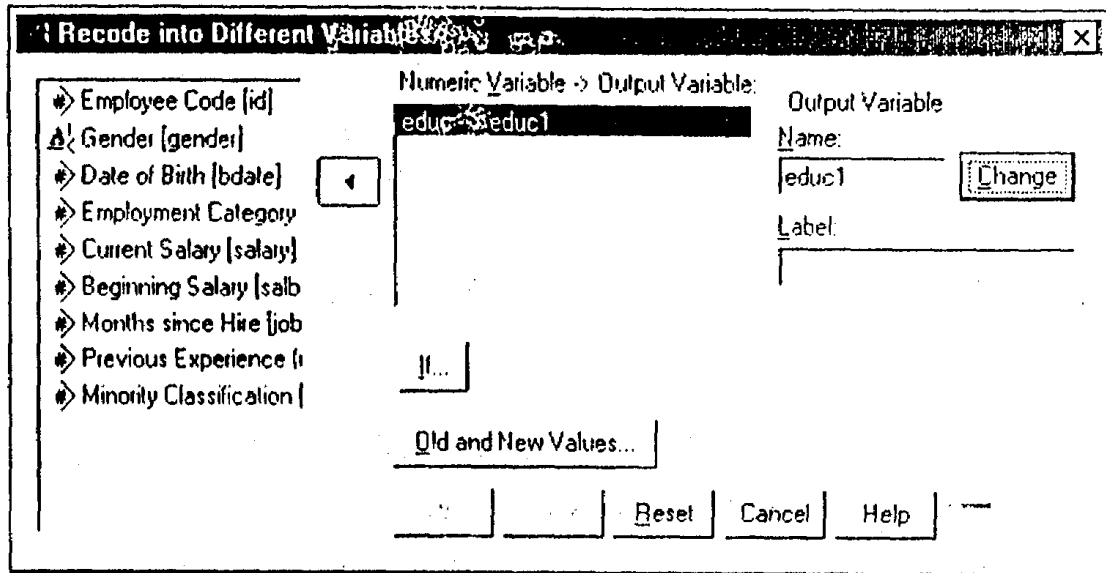
يؤدي إلى تلخيص البيانات وسهولة عرضها ، وسرعة استيعابها ، إلا أنه قد يؤثر على دقة

المقاييس الإحصائية المترتبة على هذه البيانات .

ولتكويد المتغيرات داخل متغيرات جديدة نتبع الخطوات التالية :

(١) افتح الملف Employee data.sav

(٢) من قائمة Transform اختر أمر Recode لتسجل قائمة بها اختيارات ، اختر



١) Recode into Different Variables ، يفتح الصندوق الحوارى التالى :

٢) فى خانة Input Variable اختر المتغير المراد تكويده .

٣) فى خانة Output Variable اكتب اسم المتغير الجديد الذى سيحتوى القيم المكودة ، ثم انقر أمر Change .

٤) انقر أمر Old and New Values يفتح الصندوق الحوارى

Recode into Different Variables: Old and New Values

٥) سجل القيم الفعلية للمتغير الاصلى فى خانة Old Value ، ثم سجل القيم الجديدة (الأكواد) فى خانة New Value ، بنفس الطريقة السابقة ثم اختر أمر Ok

إعطاء رتب للحالات Rank Cases :

يمكن إعطاء رتب للحالات بالنسبة لمتغير معين ، البرنامج يضع هذه الرتب داخل متغيرات مناظرة للمتغيرات الأصلية ، يمكن أن تكون هذه الرتب تصاعدية أو تنازلية .

ولإعطاء رتب للحالات وفقاً لمتغير معين تتبع الآتى :

١) افتح الملف Employee data.sav

٢) من قائمة Transform اختر أمر Rank Cases ، يفتح الصندوق الحوارى

Rank Cases

Variable(s):

Current Salary [salary]

By:

Employment Category []

Assign Rank 1 to:

☒ Smallest value

☐ Largest value

☒ Display summary tables

Rank Types... Ties...

OK Paste Reset Cancel Help

٣) أمام خانة Variable اختر المتغير المراد إعطاء رتب للحالات على أساسه ، وليكن المتغير salary .

٤) أمام خانة By اختر المتغير المراد تقسيم الرتب إلى مجموعات على أساسه وليكن Jobcat ، أى كلما تغيرت الـ Jobcat كلما بدأ الترتيب من رقم ١ .

٥) إذا أردنا أن تكون الرتب تصاعدية فإننا نختار Smallest value من خانة Assign Rank 1 to ، أما إذا أردنا أن تكون الرتب تنازلية فإننا نختار Largest Value

٦) إذا أردنا معرفة طريقة معالجة القيم المتشابهة في الرتب ، فإننا نختار Ties يفتح الصندوق الحوارى التالى :

Rank Assigned to Ties

☒ Mean ☐ Low ☐ High

☐ Sequential ranks to unique values

Continue Cancel Help

ولتوضيح الفرق بين هذه الأنواع في معالجة القيم المتشابهة نعرض الجدول التالى

Value	Mean	Low	High	Sequential
١٠	١	١	١	١
١٥	٣	٢	٤	٢
١٥	٣	٢	٤	٢
١٥	٣	٢	٤	٢
١٦	٥	٥	٥	٣
٢٠	٦	٦	٦	٤

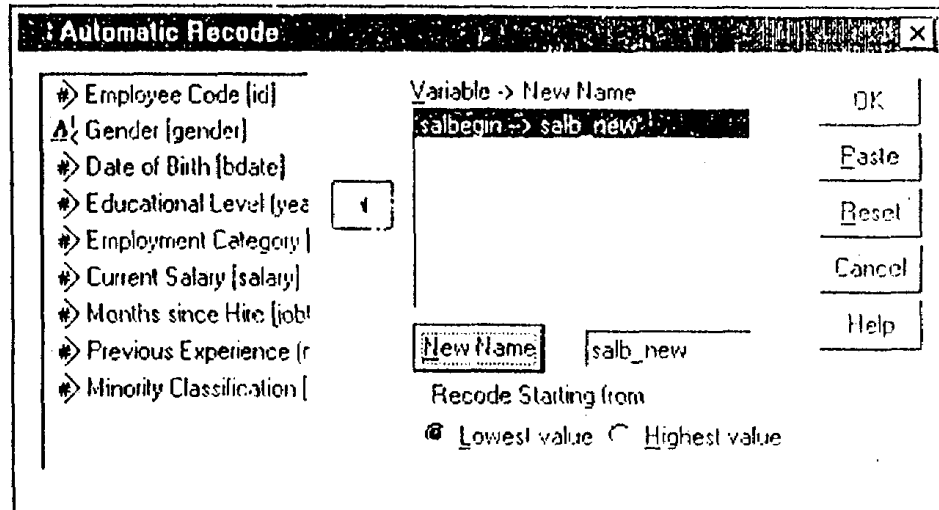
التكويد التلقائي Automatic Recode :

يقوم بترتيب البيانات وإعادة تكويدها في متغير آخر اعتباراً من ١ فصاعداً إذا اخترنا Lowest Value ، وهو أشبه بالرتب ، ولعمل Automatic Recode لمتغير معين نتبع الآتي :

١) افتح الملف Employee data.sav

٢) من قائمة Transform اختر أمر Automatic Recode يفتح الصندوق الحواري

٣) أمام خانة Variable اختر المتغير المراد تكويده .



٣) أمام خانة New Name أكتب اسم المتغير الجديد ، ثم انقر أمر New Name

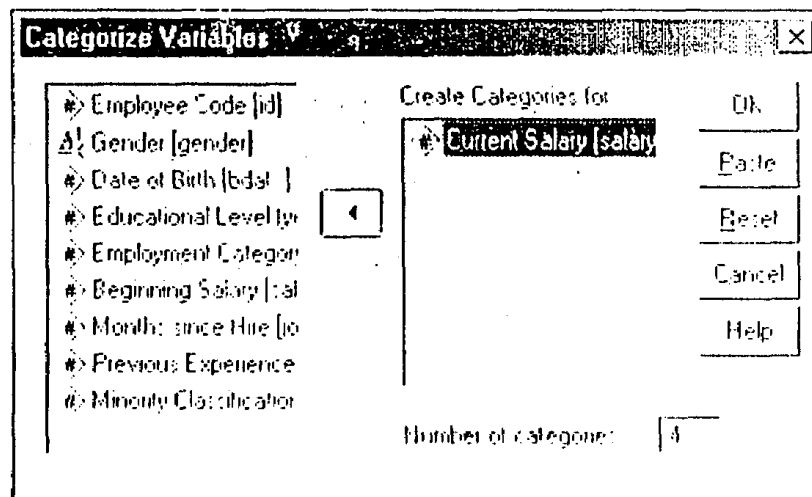
٣) اختر أمر Ok

تحويل متغير كمي الى فئات Categorize Variable

قد يوجد لدينا بيانات كمية مثل بيانات العمر أو الراتب ونريد تقسيمها إلى فئات Categories ، يمكن استخدام الأمر Recode السابق التعرف عليه ، كما يمكن استخدام أمر آخر وهو

Categorize Variable

١- افتح الملف Employee data.sav .



٢- من قائمة Transform اختر أمر Categorize Variable

٢- اختر المتغير المراد تقسيمه إلى فئات .

٣- أمام خانة Number of categories اختر عدد الفئات المراد تقسيم المتغير إليها .

تعريف سلسلة زمنية Define Dates :

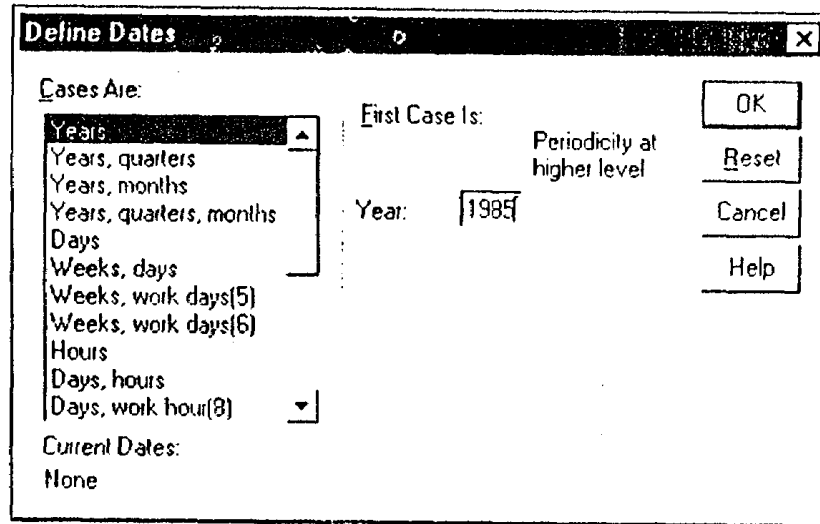
يمكن تعريف متغير على أنه سلسلة زمنية ، وذلك يفيد في تحليل السلاسل الزمنية ، كما يفيد في

أن هناك وظائف معينة Functions يمكن إجراؤها على المتغيرات الزمنية

ولتعريف متغير على أنه سلسلة زمنية تبدأ من ١٩٨٥ نتبع الآتي :

(١) افتح ملفاً جديداً .

(٢) من قائمة Data اختر أمر Define Dates يظهر المربع الحواري Define dates



(٣) أمام خانة Cases Are اختر وحدة الزمن المطلوبة (Years , Weeks , Days , ... etc)

(٤) أمام خانة First Case Is أكتب تاريخ البداية ، ثم انقر أمر OK يظهر متغير جديد بملف

البيانات عبارة عن سلسلة زمنية تبدأ بالرقم ١٩٨٥

إجراء عمليات على متغيرات زمنية Create Time Series :

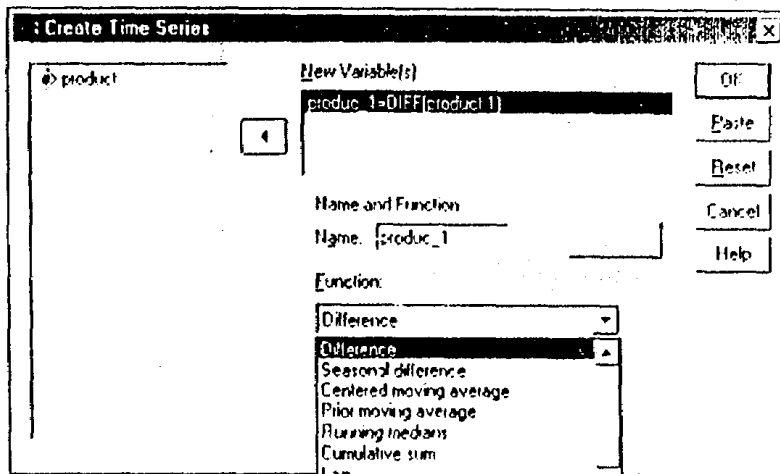
لإيجاد متغيرات تعتمد على متغيرات في سلسلة زمنية معينة مثل تطور قيمة الإنتاج الصناعي

خلال عشر سنوات .

(١) افتح الملف السابق .

(٢) من قائمة Transform اختر Create Time Series يظهر الصندوق الحواري

Create Time Series



٢) اختر الـ Function من خانة Function ، يلاحظ أن هناك عدة أنواع من الـ Functions نذكر منها :

Difference : وهو الفرق بين القيمة الحالية والقيمة السابقة للمتغير الأصلي.

Lag : وهو القيمة السابقة للمتغير الأصلي (في سنة سابقة)

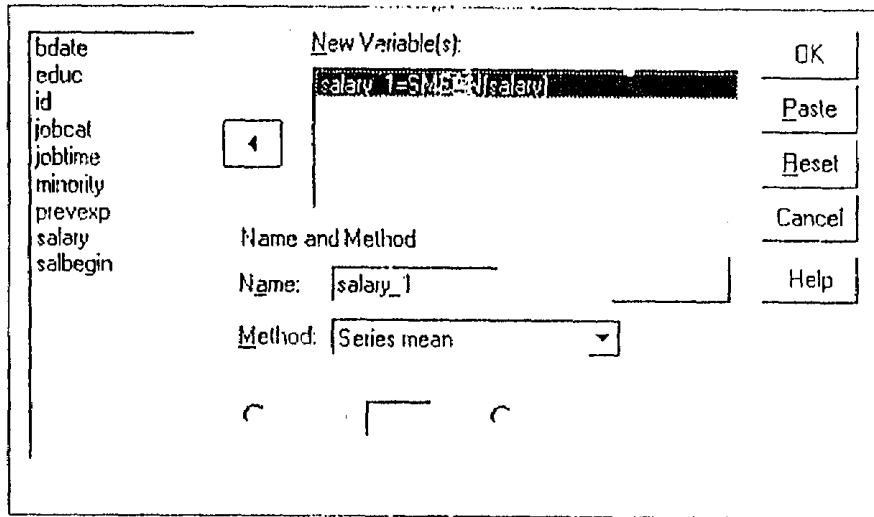
Led : وهو القيمة التالية للمتغير الأصلي (في سنة تالية) ... وهكذا .

استبدال القيم المفقودة : Replace Missing Values

القيم المفقودة قد تسبب بعض المشاكل في التحليلات الإحصائية خاصة في تحليل السلاسل الزمنية ، لذلك قد نرغب في استبدال هذه القيم الـ Missing باشتقاق متغيرات جديدة مع استبدال القيم المفقودة بقيم محسوبة من المتغيرات الأصلية بعدة طرق Methods ، ولتنفيذ ذلك نتبع الخطوات التالية:

١) افتح الملف Employee data.sav

٢) من قائمة Transform اختر Replace Missing Values يظهر الصندوق الحوارى التالى



٣) اختر طريقة التقدير من خانة Method .

٤) اختر المتغير المراد استبدال قيمه المفقودة .

٥) نجد أن البرنامج يحدد اسماً للمتغير الجديد مشتقاً من المتغير الأصلي ، يمكن تغيير هذا الاسم .

٦) اختر أمر Ok .

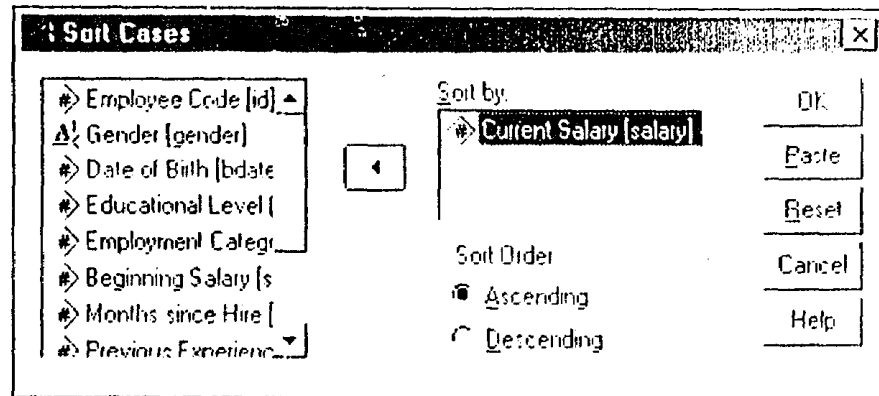
فرز البيانات Sort Cases :

قد نحتاج إلى فرز (ترتيب الحالات) تصاعدياً أو تنازلياً بناءً على قيم متغير معين ، أو عدة

متغيرات ، ولتنفيذ ذلك نتبع الخطوات التالية :

١) افتح الملف Employee data.sav .

٢) من قائمة Data اختر أمر Sort Cases



٣) أمام خانة Sort by ، اختر المتغير أو المتغيرات المراد الترتيب على أساسها

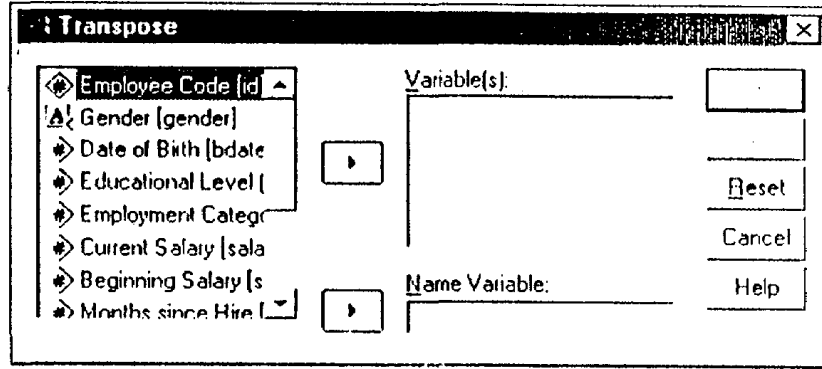
٤) أمام خانة Sort Order ، اختر نوع الترتيب (تصاعدي أو تنازلي) .

٥) اختر أمر OK يتم فرز البيانات بالترتيب المطلوب .

تدوير البيانات Transpose :

أى جعل الصفوف أعمدة ، والأعمدة صفوف .

قد تسجل البيانات بوضع لا يناسب التحليل الإحصائي ، لذلك تحتاج إلى قلب الصفوف أعمدة والأعمدة صفوف ، ولتنفيذ ذلك نتبع الآتى :



(١) افتح الملف Employee data.sav

(٢) من قائمة Data اختر أمر Transpose

(٣) أمام خانة Variables اختر المتغيرات المراد عمل Transpose لها ، مع ملاحظة أن المتغيرات التى لا يتم اختيارها سوف تفقد

(٤) اختر أمر OK .

(٥) يتم تدوير البيانات تصبح المتغيرات Cases وتأخذ أرقاماً متسلسلة ، وتصير الحالات متغيرات بأسماء etc ... var001, var002, ... ، إلى أن يتم إعادة تعريف هذه المتغيرات .

(٦) يتم حفظ البيانات على ملف آخر ، حتى تحتفظ بالبيانات الأصلية قبل تدويرها

دمج الملفات Merging Data Files :

قد يكون حجم البيانات التى نجرى عليها التحليل كبيراً ولذلك يتم تجزئة البيانات إلى عدة ملفات ، ثم نحتاج إلى دمجها بعد ذلك ، ومن الممكن أن يتم تجزئة البيانات من حيث الحالات وفى هذه الحالة يكون الدمج عن طريق Add Cases ، أو عن طريق تجزئة المتغيرات وفى هذه الحالة يكون الدمج عن طريق Add Variables .

١- الدمج عن طريق إضافة حالات Add Cases :

مثال : تم تسجيل بيانات الطلبة فى مادتى الرياضيات والإحصاء على ملف male كالتالى :

math	٥٨	٩٢	٦١	٧٤	٢١	٦٧	٦٨	٦٤	٦٥	٧٣
stat	٦٧	٦٣	٥٧	٥٢	٦١	٦٣	٥٠	٦٧	٦٩	٦٠
sex	١	١	١	١	١	١	١	١	١	١

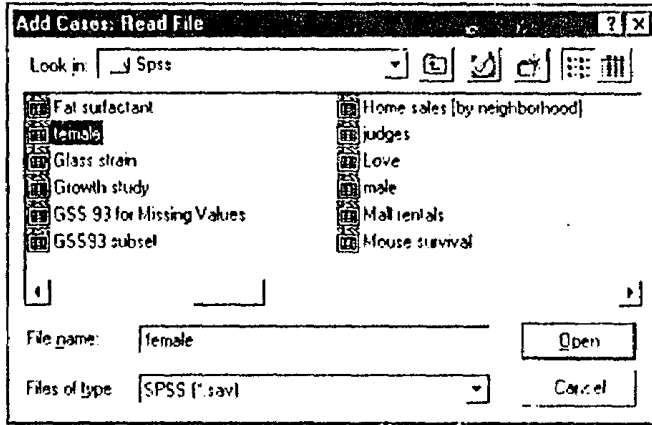
وتم تسجيل بيانات الإناث على ملف female كالتالى :

٥٢	٤٣	٧١	٨٠	٩٥	٥٧	٧٤	٥٥	٦٥	٣٠	math
٥٦	٦٧	٦٣	٥٩	٩٠	٧٠	٣٦	٥٩	٧١	٦١	stat
٢	٢	٢	٢	٢	٢	٢	٢	٢	٢	sex

ونريد دمج الملفين معاً ، ولدمج الملفين عن طريق Add Cases يفضل أن تكون متغيرات الملف الأول هي نفس متغيرات الملف الثانى ، كما يفضل ترتيب بيانات الملفين بطريقة واحدة قبل الدمج ، ولدمج حالات الملفين نتبع الآتى :

١- افتح الملف الأول male

٢- من القائمة الرئيسية اختر Data ومنها اختر Merge Files ومنها اختر Add Cases يظهر الصندوق الحوارى التالى :

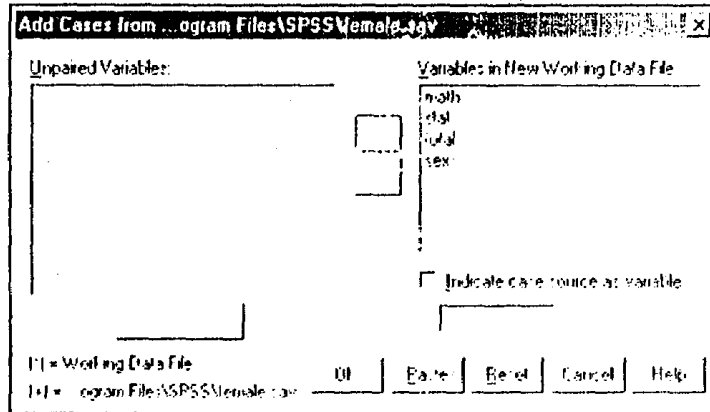


٣- اختر اسم الملف المطلوب دمج مع الملف الأول ، ثم اختر أمر Open

٤- يظهر الصندوق الحوارى وبه اسم الملف ، وبه أسماء المتغيرات فى الملف الجديد بعد الدمج فى

خانة Variables in New Working Data File

٥- يمكن تنشيط الاختيار Include case source as variable لإضافة متغير جديد به رقم يعبر عن الملف (مصدر البيانات) .



٦- اختر أمر OK يتم دمج بيانات الملفين معاً

٢- دمج البيانات : Add Variables

يفضل في هذه الحالة أن تتساوى عدد الحالات ، فإذا لم تتساوى عدد الحالات فإن النظام يضع للحالات الأقل (.) system missing حتى تتساوى مع عدد الحالات الأكبر
مثال :

١- أدخل البيانات التالية في ملف A1 وهي عبارة عن بيانات الطول والوزن :

No	١	٢	٣	٤	٥
height	١٦٠	١٨٠	١٧٠	١٧٥	١٦٥
weight	٦٢	٩٠	٨٠	٦٧	٧٢

٢- ثم أدخل بيانات العمر في ملف آخر A2 :

No	١	٢	٣	٤	٥
age	٣٠	٢٠	٢٢	٢٥	٢٦

ويراد دمج بيانات الملفين معاً ، ولتنفيذ ذلك نتبع الآتى :

١- افتح بيانات الملف الأول a1 .

٢- من قائمة Data اختر Merge Files ، ومنها اختر Add Variables .

٣- يفتح الصندوق الحوارى Add Variables : Read File

٤- اختر اسم الملف a2 واختر أمر Continue

٥- يظهر الصندوق الحوارى Add Variables from وبه اسم ومسار الملف الثانى a2 .

٦- اختر أمر Continue .

حساب مقاييس إحصائية للبيانات مقسمة إلى مجموعات : Aggregate Data

يمكن تقسيم البيانات إلى مجموعات (حسب متغير معين) وحساب مقاييس معينة لهذه المجموعات مثل (المتوسط ، الانحراف المعياري ... الخ) ، ويمكن أن يكون هذا التقسيم على نفس الملف أو على ملف آخر ، البرنامج يخصص ملف Aggr.sav تلقائياً ويمكن تغييره ، ولعمل ذلك نتبع الآتى :

١- افتح الملف Employee data.sav .

٢- من قائمة Data اختر أمر Aggregate يفتح الصندوق الحوارى التالى :

Aggregate Data

Break Variable(s):
 Employment Category [i]

Aggregate Variable(s):
 salary_1 = MEAN(salary)
 salary_2 = SUM(salary)

Save number of cases in break group as variable: ☐

Create new data file ☒ File... D:\...SPSS\AGGR.SAV

Replace working data file ☐

Buttons: OK, Paste, Reset, Cancel, Help

Buttons: Name & Label..., Function...

٢- أمام خانة Break يتم اختيار المتغير الذي يتم تجميع البيانات على أساسه ، وكل قيم من قيم هذا المتغير ستناظر حالة واحدة من البيانات بعد التجميع .

٣- أمام خانة Aggregate Variable(s) يتم اختيار المتغيرات التي سيتم تجميعها ويخصص البرنامج لها أسماء مشتقة من المتغيرات الأصلية ، ويمكن إعادة تسمية هذه المتغيرات عن طريق اختيار أمر Name & Label .

٤- انقر أمر Function يتم اختيار المقاييس الإحصائية المراد حسابها لهذه البيانات المجمعة .

٥- اختر أمر تظهر البيانات في ملف Aggr.sav يمكن فتحه فيظهر على النحو التالي

jobcat	salary_1	salary_2
Clerical	27838.54	10105390
Custodial	30938.89	835350.0
Manager	63977.80	5374135

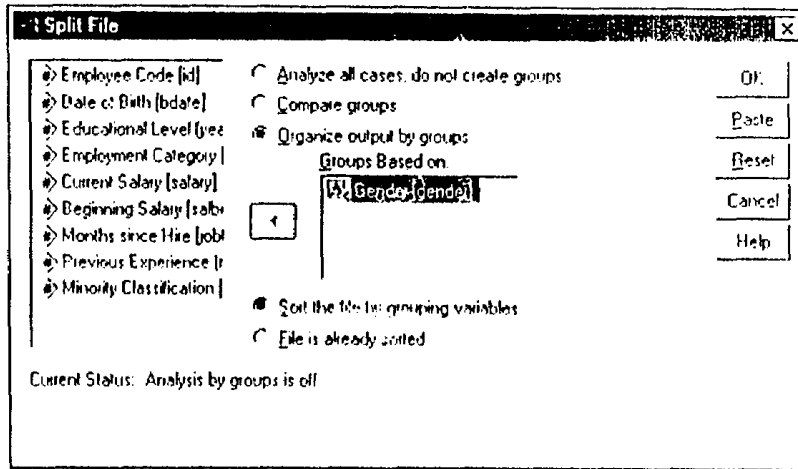
تقسيم الملفات Split Files :

يمكن تقسيم ملف البيانات إلى مجموعات منفصلة لأغراض التحليل بناءً على قيم متغير معين مثل النوع ولتنفيذ ذلك نتبع الآتي :

(١) افتح الملف Employee data.sav

٢) من قائمة Data اختر أمر Split File

يظهر الصندوق الحوارى Split File .



٣) نشط خانة Organize output by groups .

٢) اختر المتغير المراد التقسيم على أساسه أما خانة Groups based on ، ثم اختر أمر OK

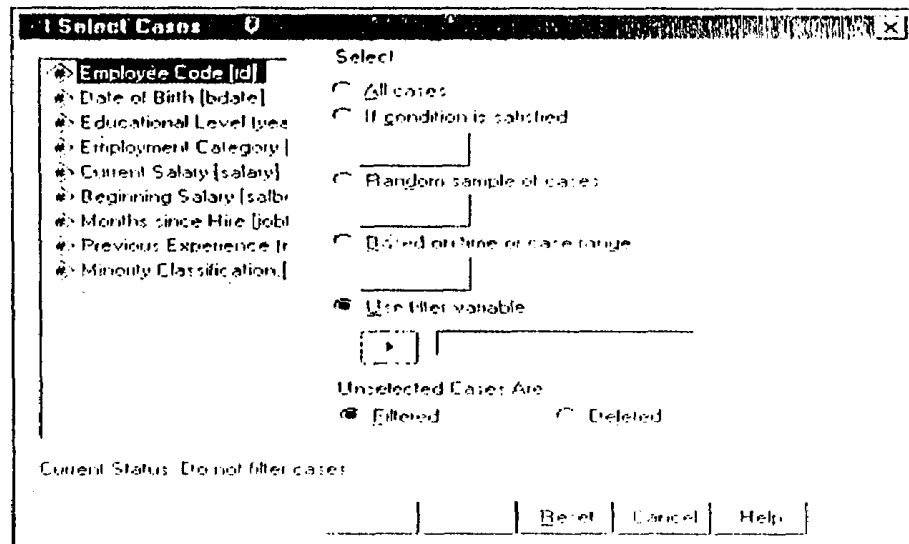
اختيار حالات Select Cases:

قد نكون مهتمين بعرض وتحليل حالات معينة من البيانات ، وليس كل البيانات ولتنفيذ ذلك

نتبع الآتى :

١) افتح الملف Employee data.sav

٢) من قائمة Data اختر أمر Select Cases



يظهر الصندوق الحوارى Select Cases .

ويظهر أمامه الاختيارات التالية :

١) الاختيار All Cases وذلك فى حالة اختيار كل الحالات

٢) الاختيار المشروط وفي هذه الحالة يتم تنشيط If condition ، ونقر if
وكتابة الشرط المطلوب اختيار الحالات على أساسه ، مثلاً في حالة اختيار الذكور يتم كتابة
select case if gender= 'm' في الصندوق الحوارى

يتم اختيار بيانات الـ Male فقط أما بيانات الـ Female فيظهر بها شطب على النحو التالى

	id	gender	bdate	educ	jobcat	salary	salbegin
1	2	Male	05/23/58	16	Clerical	\$40,200	\$18,750
2	3	Female	07/26/29	12	Clerical	\$21,450	\$12,000
3	4	Female	04/15/47	8	Clerical	\$21,900	\$13,200
4	5	Male	02/09/55	15	Clerical	\$45,000	\$21,000
5	6	Male	08/22/58	15	Clerical	\$32,100	\$13,500
6	7	Male	04/26/56	15	Clerical	\$36,000	\$18,750
7	8	Female	05/06/66	12	Clerical	\$21,900	\$9,750
8	9	Female	01/23/46	15	Clerical	\$27,900	\$12,750
9	10	Female	02/13/46	12	Clerical	\$24,000	\$13,500
10	11	Female	02/07/50	16	Clerical	\$30,300	\$16,500
11	12	Male	01/11/66	8	Clerical	\$20,350	\$12,000

- إذا تم اختيار Filtered فيكون هذا الاختيار اختياراً مؤقتاً ، ويمكن بعد ذلك اختيار كل الحالات ، ويضيف إلى البيانات متغيراً يأخذ القيمة ١ في حالة البيانات المختارة ، صفر في حالة البيانات غير المختارة. إما إذا اخترنا Deleted فيكون هذا الاختيار اختياراً دائماً ولا يمكننا بعد ذلك اختيار كل الحالات

اختيار عينة عشوائية :

١) لاختيار عينة عشوائية يتم اختيار sample من الصندوق الحوارى

select Cases وذلك بعد تنشيط الاختيار Random sample of cases كالتالى :

Select Cases: Random Sample

Sample Size

☒ Approximately % of all cases

☐ Exactly cases from the first cases

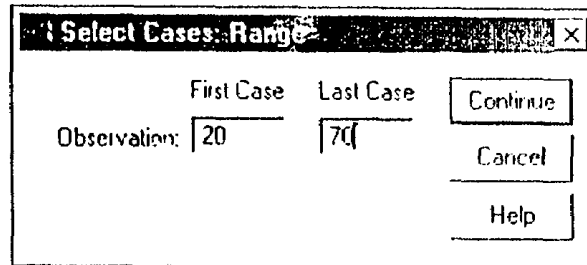
٢) يتم تحديد حجم العينة إما نسبة من البيانات ، أو عدد حالات من إجمالى الحالات .

تحديد مدى معين من البيانات

٣) لتحديد مدى معين من البيانات يتم اختيار Range وذلك بعد تنشيط المجموعة Based on

Select cases time or case range من الصندوق الحوارى

وأمام خانة First Case , Last Case يتم اختيار مدى البيانات المراد التحليل على أساسه .



يلاحظ أن التحليل سيقصر على البيانات المختارة فقط (والبيانات الغير مختارة سيظهر عليها

شطب أمام رقم الحالة) فإذا أردنا أن يتم التحليل على كل البيانات فلا بد من اختيار All

Cases .

ملحوظة

إذا أردنا التعامل مع البيانات بعد ذلك بدون اختيار بيانات معينة ، فيتم اختيار

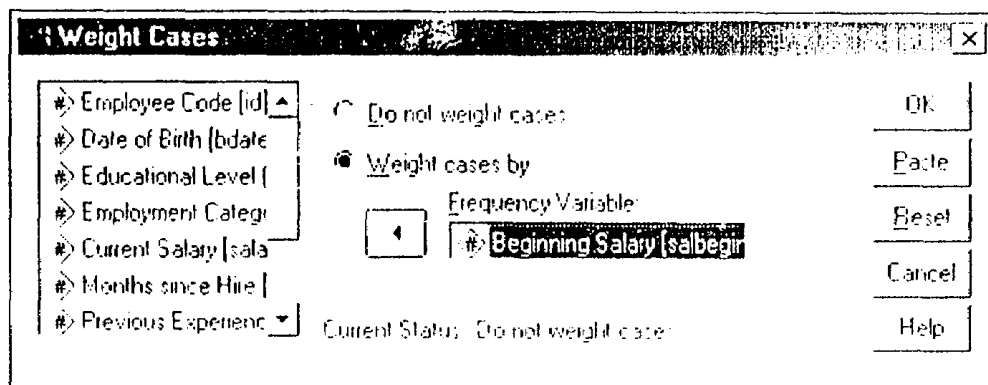
All Cases من الصندوق الحوارى Select Cases

ترجيح الحالات :

في بعض التحليلات نحتاج لترجيح البيانات خاصة إذا كانت البيانات موزعة في جداول

تكرارية فيتم ترجيح البيانات بأوزانها أى بتكراراتها ويتم ذلك عن طريق :

١) من قائمة Data اختر أمر Weight Cases يفتح الصندوق الحوارى Weight Cases



٢) نشط الاختيار Weight Cases by ، ثم اختر المتغير المراد الترجيح به في خانة

Frequency Variable .

٣) إذا أردنا بعد ذلك التعامل مع البيانات بدون ترجيح يتم اختيار

Do not weight cases

الباب الثانی الأسلوب الإحصائي وعرض وتحليل البيانات

١-٢ : تحديد الأسلوب الإحصائي المناسب لتحليل البيانات :-

إن أصعب مرحلة تواجه الباحث في بحثه هي مرحلة التحليل الإحصائي وذلك لكثرة تعداد الطرق الإحصائية . فالباحث وهو يستعرض مجموعة كبيرة من مختلف الطرق الإحصائية سيقف حائراً ولا يدري أيها منها يختار ما لم يكن لديه معايير ومعلومات مسبقة يستشير بها في هذا الاختيار ، وإن عدم الأخذ بهذه المعايير أو المعلومات تجعل الباحث في النهاية يسعى استخدام الإحصاء فيستخدم بالتالي طرقاً لا تتناسب مع طبيعة بيانات بحثه، ويمكن التغلب على هذه المشكلة بسهولة إذا استطاع الباحث مقدماً التعرف أي من الطرق الإحصائية يناسب بياناته ، أي الطرق العلمية أم الطرق اللامعلمية ؟ وفيما يلي تفصيل لكل منهما :-

أولاً : الطرق المعلمية (parametric)

الإحصاءات المعلمية هي تلك الطرق التي تتطلب الوفاء بافتراضات معينة حول المجتمع الذي تسحب منه العينة، ومن هذه الافتراضات أن تتخذ المشاهدات في المجتمع شكل التوزيع الطبيعي على سبيل المثال لا الحصر . وقد أشتق مصطلح " معلمية " من مفهوم " معلم " الذي يعنى صفة أو خاصية من خصائص مجتمع معين ، فكل قيمة من القيم التي تتعلق بخصائص المجتمع تسمى parameter أما تلك الخصائص المتعلقة بالعينة التي سحبت عشوائياً من المجتمع فتسمى كل منها تقدير Estimate أي أن القيمة المستخرجة لأي خاصية في العينة ماهي إلا تقدير لقيمة تلك الخاصية في المجتمع والتي على الأغلب تكون غير معروفة ، إذ أنها لو كانت معروفة لاتكون هناك حاجة لحساب تقديرها. ومن الأمثلة على هذه الخصائص الوسط الحسابي، والانحراف المعياري وشكل توزيع المجتمع .

والاختبارات الإحصائية المعلمية هي التي تعتمد على الافتراضات الخاصة بخصائص المجتمع مثل اختبار (ت) ، واختبار (ف) " وتحليل التباين الاحادي وتحليل التباين المتعدد.

بالإضافة إلى ما سبق يجب أن يكون اختيار العينة من المجتمع بصورة عشوائية ، وأن تكون إحصاءات العينة (مقاييس النزعة المركزية والتشتت) صورة مقربة للمعلومات الإحصائية للمجتمع، وكذلك تستخدم الطرق الإحصائية المعلمية لمعالجة وتحليل البيانات الكمية.

وكذلك يفترض أن يكون حجم العينة ليس صغيراً جداً لأن الحجم الصغير للعينة يؤثر على خصائص التوزيع التكرارى للعينة الصغيرة ، فتبعد بذلك عن اعتدالية التوزيع التكرارى للمجتمع.

وتكون كذلك الطرق المعلمية أكثر ملائمة لتحليل البيانات الفترية والبيانات النسبية فضلاً على أنها يمكن استخدامها (مع بعض التحفظ) لتحليل البيانات الرتبية فقط، وذلك بعد أن يتم تحويلها إلى بيانات فترية من خلال إعطاء كل رتبة درجات تتناسب مع قيمتها مع مراعاة أن تكون البيانات فى صورة رتب ذات درجات متتابعة مثل موافق بشدة ، موافق ، غير موافق ، غير موافق بشدة.

ثانياً : الطرق اللامعلمية (Non Parametic)

الإحصاءات اللامعلمية هى تلك الطرق التى تستخدم فى تحليل البيانات واختبار الفرضيات الخاصة بالبيانات الاسمية والرتبية والتى تكون من النوع المتقطع (المنفصل) عادة ويمكن استخدامها مع البيانات الفترية والنسبية بعد ان يتم تحويلها الى بيانات اسمية أو رتبية. وفى حالة الإحصاءات اللامعلمية يجب أن تكون العينة المختارة عشوائية بمعنى أن يكون هناك استقلال بين مفردات العينة. هذا فضلاً عن أن هذه الطرق لا تتقيد بالشروط الواجب توافرها لاستخدام الإحصاءات المعلمية حيث أنها تستخدم فى الحالات التى لا يكون فيها التوزيع النظرى للمجتمع الاصلى الذى اختيرت منه العينة معروفاً، وفى حالة عدم إمكانية الوفاء بافتراض أن التوزيع النظرى للمجتمع طبيعياً ويكون كذلك حجم العينة صغيراً جداً .

وتحتوى الحزمة SPSS على مجموعة متميزة من الاختبارات اللامعلمية. وسوف نستعرض هذه الاختبارات وفقاً إذا كان الاختبار يتعلق بعينة أو عيتين أو أكثر (مستقلتين أو غير مستقلتين) .

أولاً : اختبار عينة :

يوجد أربعة اختبارات لعينة واحدة الأول هى

- اختبار كا²
- اختبار ذى الحدين
- اختبار العشوائية

• ثم أخذنا كولو مجروف

ويمكن إجرائها كالتالى :-

١- اختبار كا^٢ يعد من أشهر الاختبارات اللامعلمية وهناك ثلاثة شروط لتطبيقه كا^٢ وهى :

١- أن تكون العينة عشوائية

٢- استقلال المشاهدات.

٣- أن حجم العينة أكثر من ٣٠

وفى هذا الاختبار نختبر هل هناك فرقاً معنوياً بين التكرارات المشاهدة والتكرارات المتوقعة التى تم الحصول عليها طبقاً لقاعدة أو نظرية معينة ويكون الفرض العدمى أن التكرار النظرى يتبع التكرار المشاهد والفرض البديل أنه لا يتبعه.

٢- اختبار ذى الحدين

يستخدم هذا الاختبار فى معرفة هل متغير ثنائي يتبع عشوائياً توزيع ذى الحدين أم لا؟
الفرض العدمى فى هذه الحالة أن البيانات تتبع توزيع ذى الحدين بمعلمة h أو بأي معلمة يحددها الاختبار و الفرض البديل أن البيانات لا تتبع ذى الحدين.

٣- اختبار العشوائية

يستخدم اختبار العشوائية لمتغير يقاس بمقياس اسمي، لإجراء ذلك الاختبار يحدد رقماً معيناً سوف يستخدمه برنامج SPSS فى تقسيم قيم المتغير إلى مجموعتين .
القيم الأقل منه تعامل معاملة واحدة وكذلك القيم التى تساوى أو تزيد عنه. الفرض العدمى يكون أن العينة عشوائية أو وسيط العينة يساوى الرقم المحدد.

٤- اختبار كولومجروف سيمنروف لعينة واحدة:

يهدف هذا الاختبار إلى معرفة هل البيانات المتاحة تتوزع حسب توزيع معين أم لا؟ يتولى البرنامج توفيق البيانات أولاً إلى أحد التوزيعات الأربعة الآتية ثم اختبار جودة التوفيق بمعنى هل اختبار التوزيع موفقاً أم لا :

١- التوزيع الطبيعي

٢- التوزيع المنتظم

٣- التوزيع الأسى

٤- توزيع بواسون

ويمكن اعتبار هذا الاختبار من اختبارات جودة التوفيق للتوزيعات الأربعة المشار إليها.

ثانياً : اختبار عينتين غير مستقلتين :-

إذا كان لدينا عينتين مستقلتين فإنه يمكن معرفة هل هناك فرقاً بينهما أم لا وذلك باستخدام مجموعة من الاختبارات اللامعلمية حيث كل اختبار يقيس الاختلاف من وجه نظر معينة. والاختبارات اللامعلمية الأربعة هي :-

- ١- اختبار مان وتيني. Mann-whithney
- ٢- اختبار كولومجروف سيمنوف. Kolomogorov-Smirnov
- ٣- اختبار موزيس للقيم الشاذة. Moses Extreme Reactions
- ٤- اختبار والد للدورة. Wald-Wolfowitz

ثالثاً : اختبار عينتين مستقلتين :-

إذا كان لدينا عينتين غير مستقلتين وأردنا إجراء اختبار لامعلمي وذلك لمعرفة هل هناك اختلاف بين العينتين أم لا؟ يمكن إجراء ثلاثة اختبارات في هذه الحالة وهي :-

- ١- اختبار ولكوكسن Wilcoxon Test
- ٢- اختبار الإشارة Sign Test
- ٣- اختبار ماكنمار McNemar Test

رابعاً : اختبار أكثر من عينتين مستقلتين :-

إذا كان لدينا أكثر من عينتين مستقلتين وكانت احاد الشروط اللازمة لتطبيق اختبار تحليل التباين غير مستوفاة فإن يمكن إجراء تحليل التباين لامعلمي للرتب وذلك باستخدام اختبار يطلق عليه أسم كيرسكال ويلز حيث يستخدم للفرق بين رتب أكثر من عينتين مستقلتين وهو يعتبر الصورة العامة لاختبار مان وتيني ويجري الاختبار تحت الفروض التالية :-

١. نفرض أن لدينا K من العينات المستقلة ذات الأحجام $n_1, n_2, \dots, n_K$ ونعني بالاستقلال انه بين وداخل العينات .

٢. المتغيرات محل الدراسة متغيرات مستمرة وأن وحدة القياس على الأقل ترتيبية.

٣. المجتمعات المسحوبة منها العينات متطابقة فيما عدا أن مجتمع واحد على الأقل مختلف في مقياس الموضع.

الفرض العدمي والفرض البديل كالآتي:-

الفرض العدمي : المجتمعات لها نفس الوسيط.

الفرض البديل : المجتمعات ليس لها نفس الوسيط.

خامساً : اختبار أكثر من عینتين غير مستقلتين :-

نفرض أن لدينا أكثر من عینتين غير مستقلة ونريد معرفة هل هناك فرقاً معنوياً بينهما أم لا؟

فيكون الفرض العدمي والفرض البديل كالآتي:-

الفرض العدمي : لا يوجد فرق بين العينات.

الفرض البديل : يوجد زوج واحد من العينات على الأقل الفرق بينهما معنوي.

ويوجد أكثر من اختبار في هذه الحالة وهي :-

١- اختبار كندال.

٢- اختبار كوكران.

٣- اختبار فريدمان.

وفيما يلي عرض تلخيص للمقارنة بين الطرق المعلمية واللامعلمية

الطرق الإحصائية المعلمية	الطرق الإحصائية اللامعلمية
- يتطلب استخدام الطرق المعلمية الوفاء بافتراضات معينة حول التوزيع الاساسي للمجتمع، من حيث أن يكون توزيع المجتمع طبيعياً.	- لا يتطلب استخدام الطرق اللامعلمية ايه افتراضات أو معلومات حول خصائص التوزيع الاساسي للمجتمع. (لذلك يطلق بعض الإحصائيين عليها إحصاءات التوزيعات الحرة).
- لا يمكن استخدام الطرق الإحصائية المعلمية لمعالجة وتحليل البيانات في المواقف التجريبية التي يكون فيها حجم العينة صغير جداً، لأن الحجم الصغير للعينة يؤثر على خصائص التوزيع التكراري للعينة الصغيرة ، فتنبتعد بذلك عن اعتدالية التوزيع التكراري للمجتمع الأب.	- يمكن استخدام بعض الطرق اللامعلمية لمعالجة وتحليل البيانات في المواقف التجريبية التي يكون فيها حجم العينة صغير جداً.
- تكون الطرق المعلمية أكثر ملائمة لتحليل البيانات الفترية والبيانات النسبية فضلاً عن أنها يمكن استخدامها (مع بعض التحفظ) لتحليل البيانات الرتببة فقط وذلك بعد أن يتم تحويلها إلى بيانات فترية من خلال إعطاء كل	- تكون الطرق اللامعلمية عادة أكثر ملائمة للاستخدام عندما تكون البيانات الخاصة بالبحث من النوعين الاسمي والرتبي فضلاً عن أنها يمكن استخدامها أحياناً عندما تكون البيانات فترية أو نسبية ، وذلك بعد أن

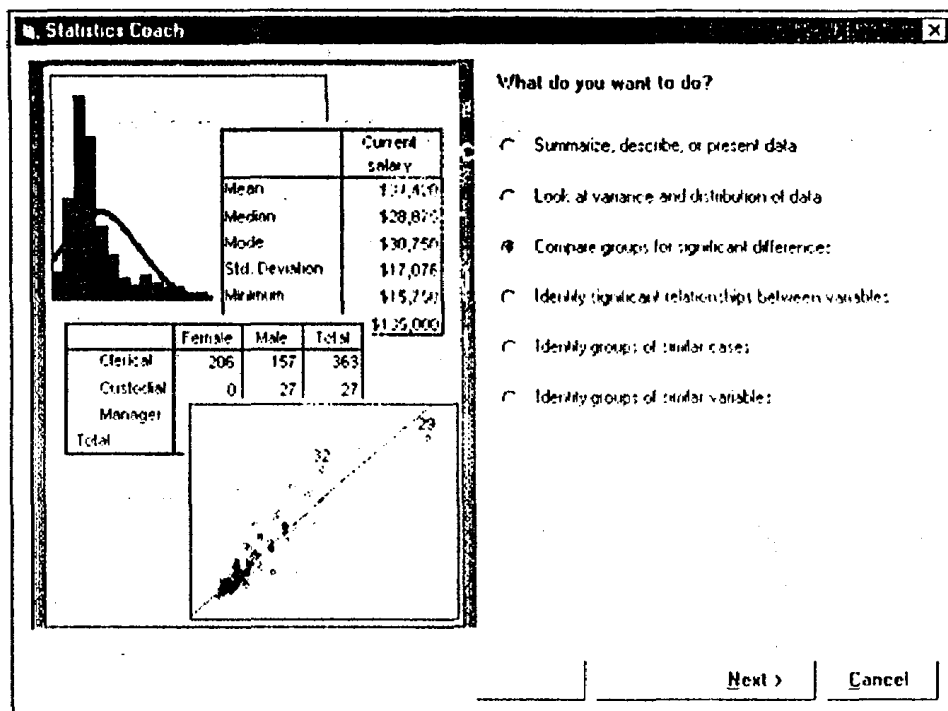
الطرق الإحصائية المعلمية	الطرق الإحصائية المعلمية
يتم تحويلها إلى بيانات اسمية أو رتبية (وهذا فيه إهدار لجزء كبير من البيانات التي نحللها ، لأن الطرق الإحصائية المعلمية تستخدم اجزاء يسيراً من تلك البيانات) لذلك أوضح بعض الإحصائيين انه عندما نستخدم الطرق الإحصائية المعلمية للبيانات الفترية والنسبية فإنها لا تستخدم إلا للتقدير المبدئي، على أن يتلوه بعد ذلك استخدام الطرق المعلمية.	رتبة درجات تتناسب مع قيمتها ، مع مراعاة أن تكون البيانات في صورة رتب ذات درجات متتابعة مثل : موافق بشدة، موافق، غير موافق، غير موافق بشدة.
-- بالرغم من تحرر الطرق الإحصائية المعلمية من الشروط والخصائص التي قد تعوق أحياناً استخدام الطرق الإحصائية المعلمية، إلا أن الطرق الإحصائية المعلمية بشكل عام أقل قوة من الطرق الإحصائية المعلمية	- تعتبر الطرق الإحصائية المعلمية بشكل عام أقوى من الطرق الإحصائية المعلمية، حيث أن الطرق المعلمية تميل إلى رفض الفرضية الصفرية أكثر من ميل الطرق المعلمية لرفض نفس الفرضية.
- تعتمد الطرق المعلمية في اغلب الأحيان على البيانات التي هي بشكل تكرارات أو رتب مما يؤدي إلى ضياع بعض المعلومات المفيدة.	- تعتمد الطرق الإحصائية المعلمية بشكل عام على الدرجات الأصلية والتي يتم تحليلها كما هي .
- تعتبر الطرق الإحصائية المعلمية أفضل وسيلة للتقدير المبدئي السريع، فهي بصورة عامة أسهل استخداماً من الطرق المعلمية وبالتالي فإن الوقت الذي يحتاجه الباحث لتحليل بياناته. يكون أقل، مما يؤدي إلى الإسراع في الحصول على النتائج والإفادة منها تطبيقياً.	- تعتبر الطرق الإحصائية المعلمية بصورة عامة أصعب في الاستخدام من الطرق الإحصائية المعلمية ، وبالتالي فإن الوقت الذي يحتاجه الباحث لتحليل بياناته يكون أطول ، مما يؤدي إلى تأخر الحصول على النتائج (نسبياً) والإفادة منها تطبيقياً.
- تعتبر الطرق الإحصائية المعلمية أكثر شيوعاً واستخداماً لدى غير المتخصصين في الإحصاء.	- تعتبر الطرق الإحصائية المعلمية أكثر شيوعاً واستخداماً لدى المتخصصين في الإحصاء.
- ليس هناك ضرورة أن يكون اختيار العينة من المجتمع بصورة عشوائية.	- ينبغي أن يكون اختبار العينة من المجتمع بصورة عشوائية، وأن تكون إحصاءات العينة (مقاييس التزعة المركزية والتشتت) صورة مقربة للمعلمات الإحصائية للمجتمع.

الطرق الإحصائية الالاعلمية	الطرق الإحصائية المعلمية
- تستخدم الطرق الإحصائية الالاعلمية لمعالجة وتحليل البيانات النوعية، والتي لا يمكن عادة استخدام أى طريقة إحصائية معلمية لتحليلها.	- تستخدم الطرق الإحصائية المعلمية لمعالجة وتحليل البيانات الكمية.
- من الأمثلة على الطرق الإحصائية الالاعلمية اختبار (كا ^٢) واختبار مان ويتنى واختبار كروسكال واليز.	- من الأمثلة على الطرق الإحصائية المعلمية اختبار (ز) واختبار (ت) واختبار (ف).

تحديد الأسلوب المناسب للتحليل الإحصائي باستخدام البرنامج الإحصائي SPSS

يمكن الاستعانة بالبرنامج في تحديد الأسلوب المناسب للتحليل كالتالى :

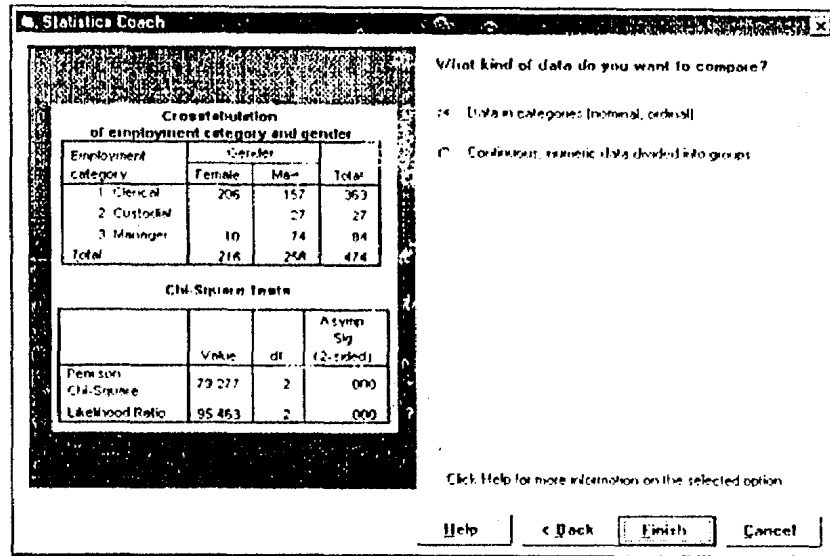
١- من قائمة Help اختر أمر Statistics Coach يظهر المعالج التالى



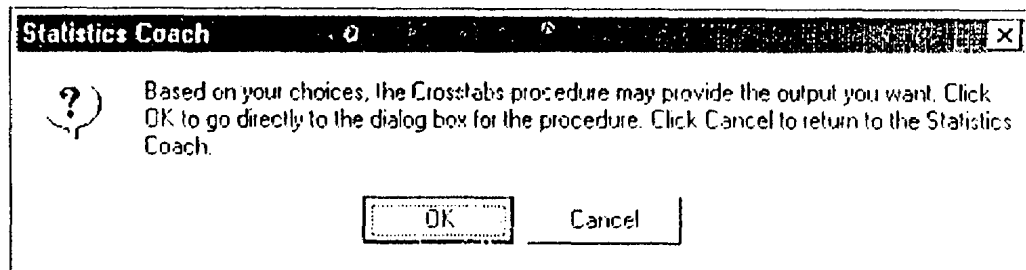
٢- اختر الموضوع الذى تريده وليكن

Compare groups for significant differences من قائمة الموضوعات ، ثم

اختر أمر Next يظهر الصندوق التالي

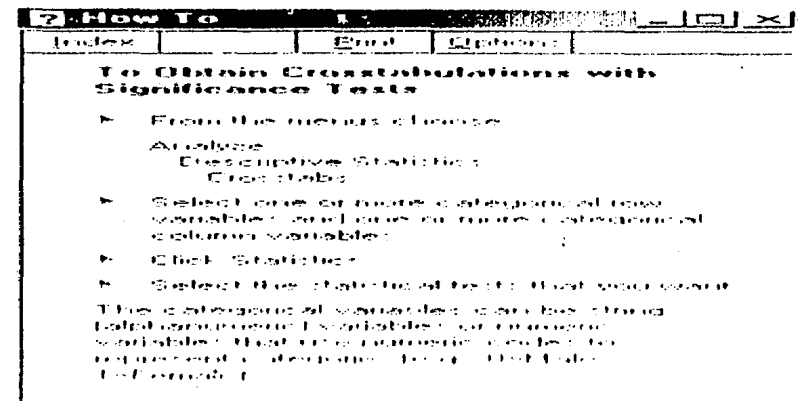


٣- يسأل هل البيانات وصفية (nominal , ordinal) Data in categories
 فبكون التحليل المناسب لها هو Chi square , Cross tabulation ويظهر نموذج من هذا التحليل في خانة المعاينة



٤- وإذا كانت البيانات غير وصفية يمكن اختيار Continuous , Numerical Data
 Divided Into groups ، واختيار التحليل المناسب .

٥- اختر أمر Finish يحدد البرنامج طريقة إجراء هذا التحليل على النحو التالي



٢-٢ الرسوم البيانية باستخدام SPSS :-

يتيح البرنامج إمكانيات هائلة في الرسوم البيانية ، حيث تعتبر وسيلة مناسبة لعرض واستيعاب البيانات لأن الرسم البياني أكثر قدرة على توضيح المعلومات ، وقد تعرضنا لبعض الرسوم البيانية الإحصائية مثل Histogram , Boxplot وسنعرض صورة مبسطة لبعض إمكانيات البرنامج في الرسم البياني .

١- إنشاء الرسم البياني Graph:

مثال : يوضح الجدول التالي درجة الحرارة العظمى ودرجة الحرارة الصغرى في بعض المدن

المدن	القاهرة	الاسكندرية	مطروح	السويس	شرم الشيخ	الأقصر
درجة الحرارة العظمى	٣٥	٢٨	٢٨	٣٦	٣٩	٤٣
درجة الحرارة الصغرى	٢٠	٢١	٢١	٢٢	٢٦	٢٣

مثل هذه البيانات بالأعمدة البيانية

لتمثيل البيانات بالرسم البياني نتبع الآتي :

١- سجل البيانات في ملف بيانات كالتالي

	city	max	min
1	القاهرة	35.00	20.00
2	الاسكندرية	28.00	21.00
3	مطروح	28.00	21.00
4	السويس	36.00	22.00
5	شرم الشيخ	39.00	26.00
6	الأقصر	43.00	23.00

٢- من القائمة الرئيسية اختر أمر Graph تظهر قائمة

بأنواع الرسوم البيانية المتاحة ، حدد نوع الرسم المطلوب وليكن Bar ،

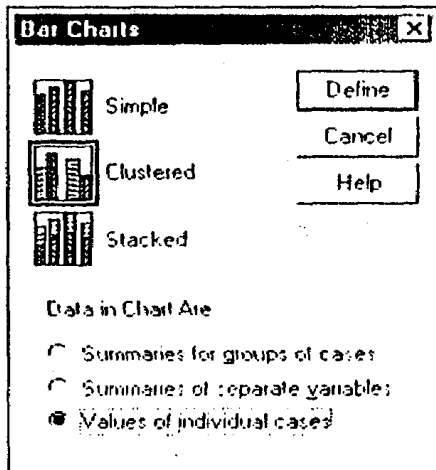
٣- بعد تحديد نوع الرسم يظهر الصندوق الحوارى

Bar Charts

٤- من هذا الصندوق حدد نوع الأعمدة البيانية المطلوبة

(Simple , Clustered, Stacked) ، اختر

Clustered



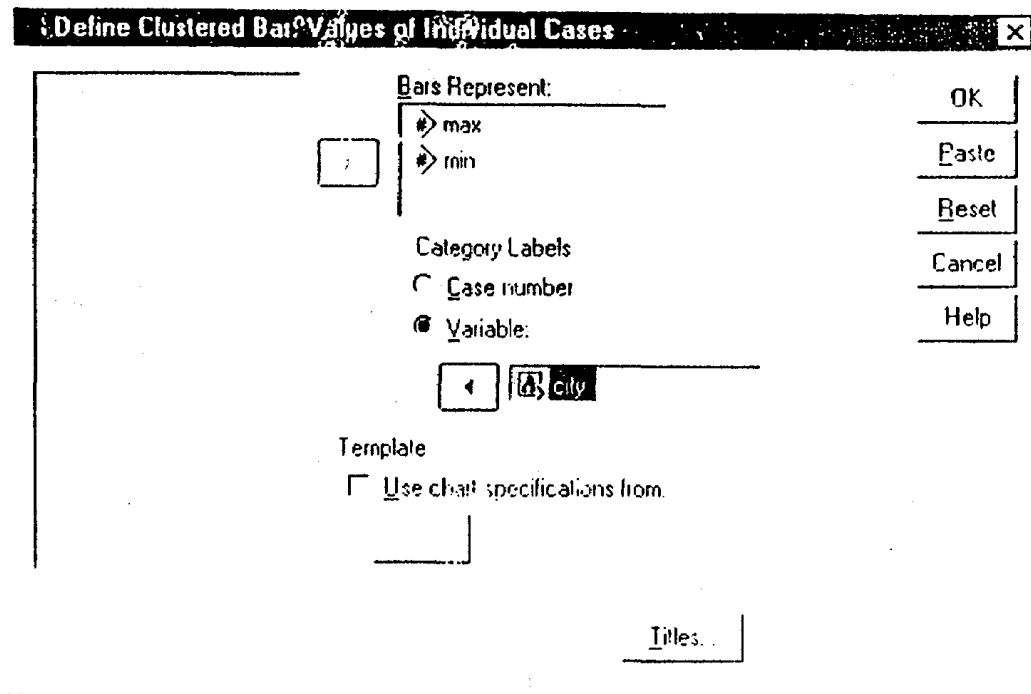
٥- نجد في خانة Data in Chart Are ثلاثة اختيارات وتعتمد هذه الاختيارات على طريقة تمثيل البيانات :

Summaries for group of cases : في حالة متغير أو متغيرات يراد تقسيمها إلى مجموعات

Summaries for separate variables : في حالة تمثيل أكثر من متغير بيانياً .

Values of individual cases : في حالة متغير واحد ويراد تمثيل بيانات كل مفردة على حدة .

٦- حدد الاختيار المطلوب وليكن Values of individual cases ، ثم اختر أمر Define يظهر الصندوق الحوارى التالى :



٧- أمام خانة Category Labels اختر المتغير الذى يحدد عناوين الأرقام (أسماء المدن) وليكن City .

٨- أمام خانة Bar represent اختر المتغيرات المراد تمثيلها max , min

٩- انقر أمر Titles يظهر الصندوق الحوارى التالى :

Titles

Title

Line 1: **شكل بياني رقم (1)**

Line 2: درجات الحرارة العظمى والصغرى في بعض المدن

Subtitle:

Footnote

Line 1:

Line 2:

Continue

Cancel

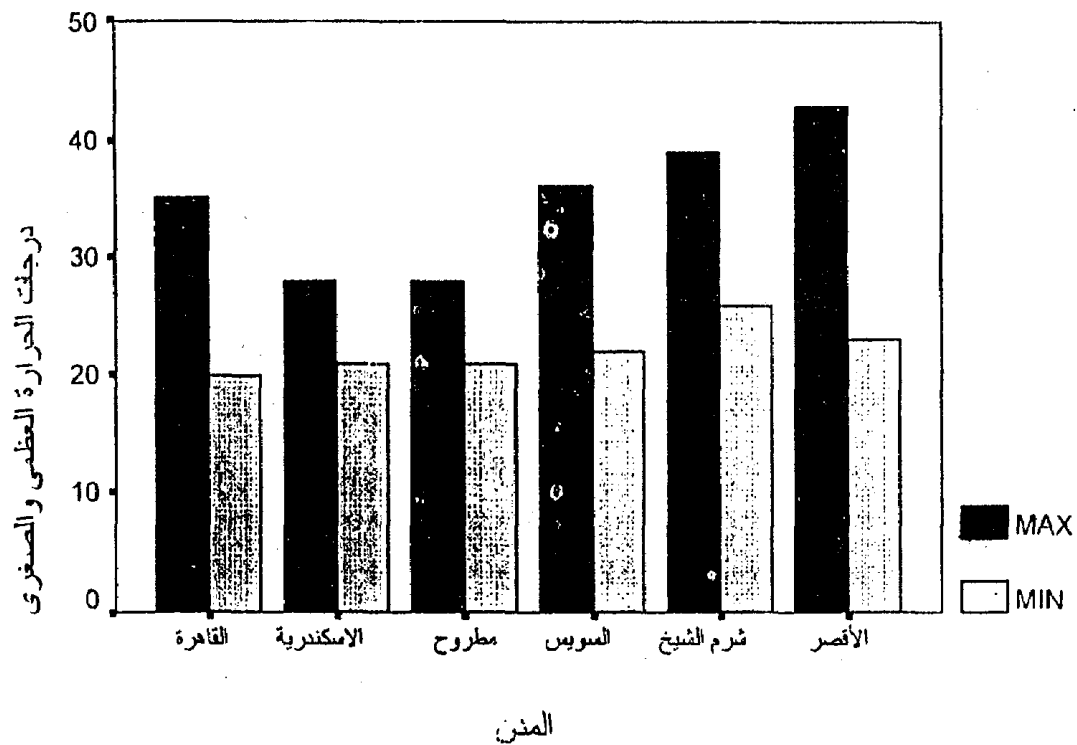
Help

١٠- بعد كتابة البيانات المطلوبة اختر أمر **Continue** ثم أمر **OK** يظهر الرسم البياني التالي:

Graph

شكل بياني رقم (1)

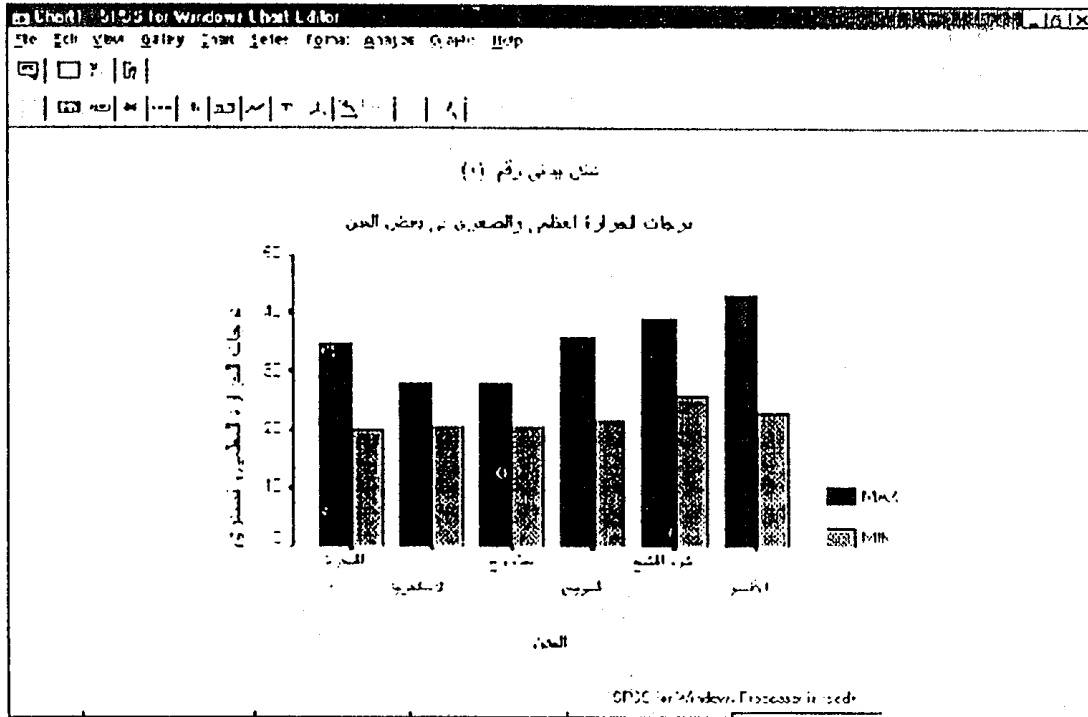
درجات الحرارة العظمى والصغرى في بعض المدن



يلاحظ أن شكل الرسم البياني غير مناسب وفي حاجة إلى تعديلات ، لذلك ستعرض إلى تعديل الرسم

تعديل الرسم Editing Chars :

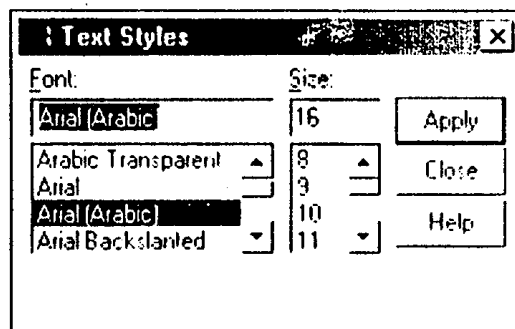
يلاحظ أن الرسم البياني يظهر داخل نافذة (Output Navigator) ، ولتعديل الرسم نضغط Double Click في منطقة الرسم فيظهر الرسم في نافذة مستقلة تسمى Chart Editor تتيح لنا التعامل مع الرسم على النحو التالي



١- تغيير حجم خط الكتابة :

يلاحظ في الرسم السابق أن الكتابة (عناوين ، توصيف المحاور) قد لا تظهر باللغة العربية لذلك يلزم تغيير خط هذه الكتابة :

– حدد الكتابة ، ثم من قائمة Format اختر أمر Text ، حيث يمكن اختيار شكل وحجم خط الكتابة ، ثم اختر أمر Apply .

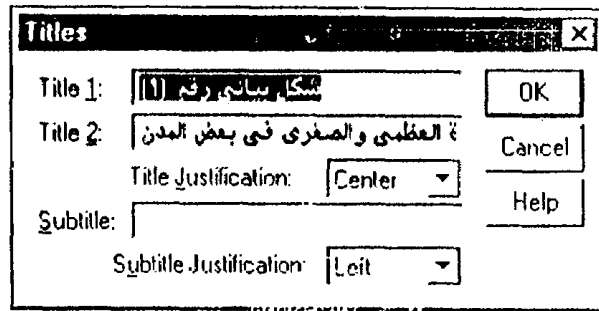


ويلاحظ أن الخط المناسب للعناوين العربية هو Arial Arabic أو Arabic Transparent

٢- محاذاة العناوين :

يلاحظ أن العناوين يتم محاذاتها نحو الشمال وأحياناً يكون من المناسب بالنسبة للعناوين أن يتم توسيطها في الرسم ولعمل ذلك :

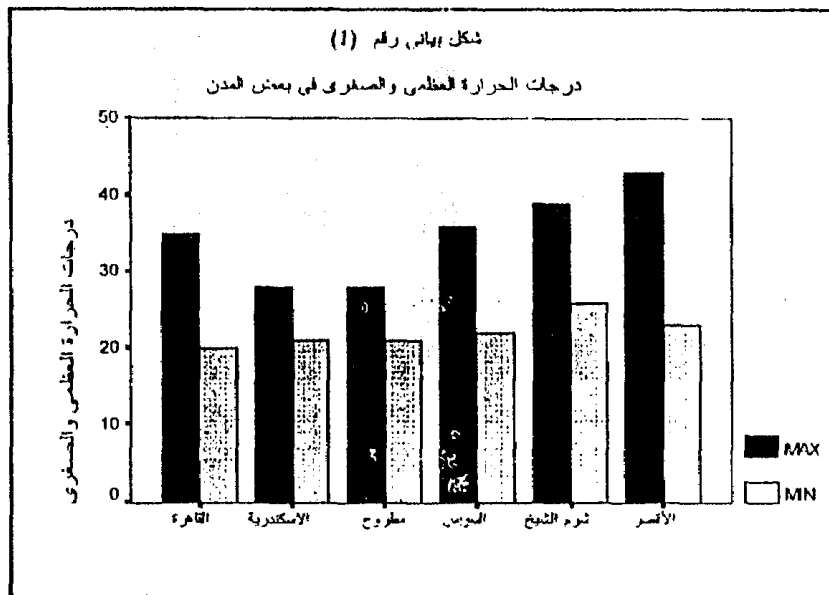
- ١- من قائمة Chart اختر أمر Titles ، وأمام خانة Title justification اختر أمر Center يتم توسيط العنوان



٢- بنفس الطريقة السابقة يتم محاذاة الـ Footnote

٣- وضع إطار خارجي للرسم :

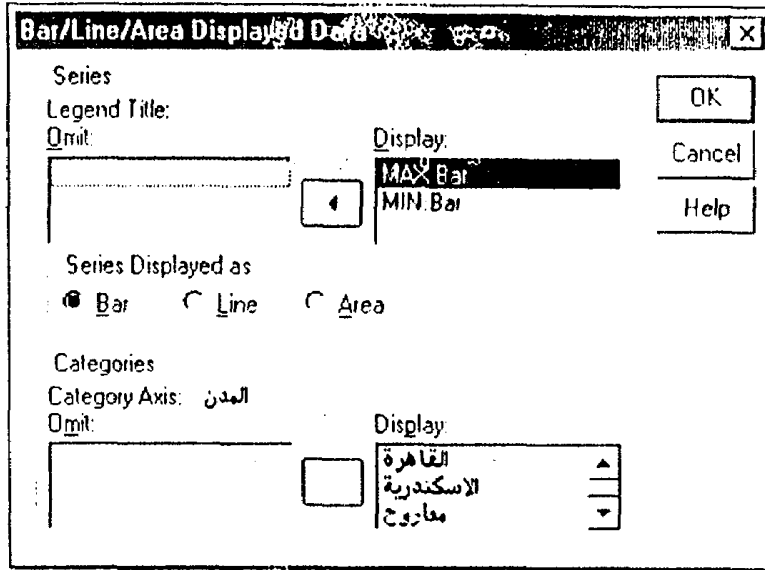
- ١- من قائمة Chart نشط أمر Outer frame يتم وضع إطار خارجي للرسم
- ٢- بعد الاختيارات السابقة يظهر الرسم البياني بالشكل المناسب التالي :



٤ - إظهار وإخفاء الأعمدة :

يمكن إظهار وإخفاء بعض الأعمدة على النحو التالي :

١ - من قائمة Series اختر أمر Display يظهر الصندوق الحوارى التالى :



٢ - يلاحظ أن الأعمدة التى تظهر فى خانة Display هى الأعمدة الممثلة ببياناً ، وإخفاء بعض

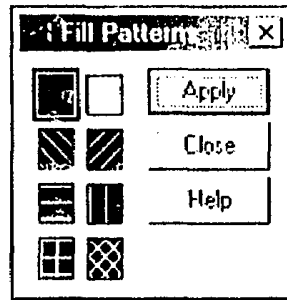
الأعمدة اختر العمود المراد إخفاؤه ثم انقر السهم المجاور ينتقل إلى خانة Omit أى يحدف

من الرسم البيانى .

٥ - قمشير الأعمدة :

لتهشير الأعمدة اختر الأعمدة المراد قمشيرها ، ثم من قائمة Format اختر أمر Fill

Pattern يظهر الشكل التالى :



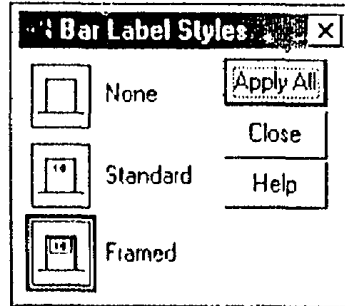
اختر الشكل المطلوب ، ثم اختر أمر Apply

٦- إظهار الأرقام على الرسم :

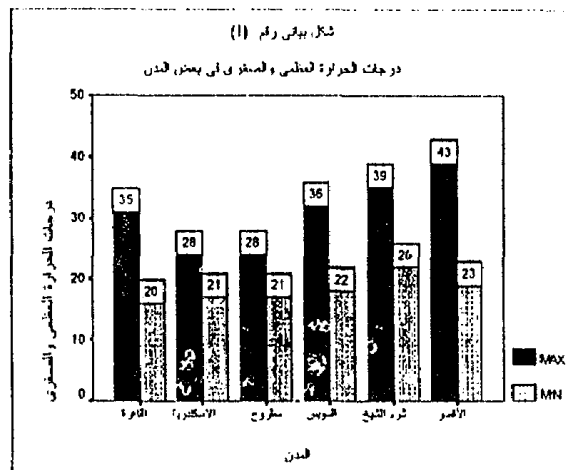
أ. لإظهار أرقام البيانات الممثلة على الأعمدة ، من قائمة **Format** اختر أمر **Bar**

Label Style

ب. يظهر الصندوق الحوارى التالى :



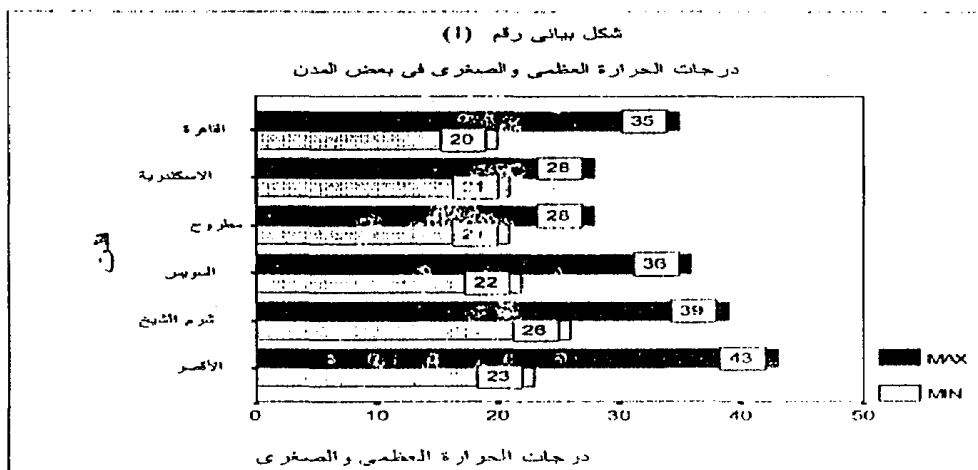
ج. اختر الشكل **Framed** يظهر الرسم البيانى بالشكل التالى :



٧- تدوير الرسم :

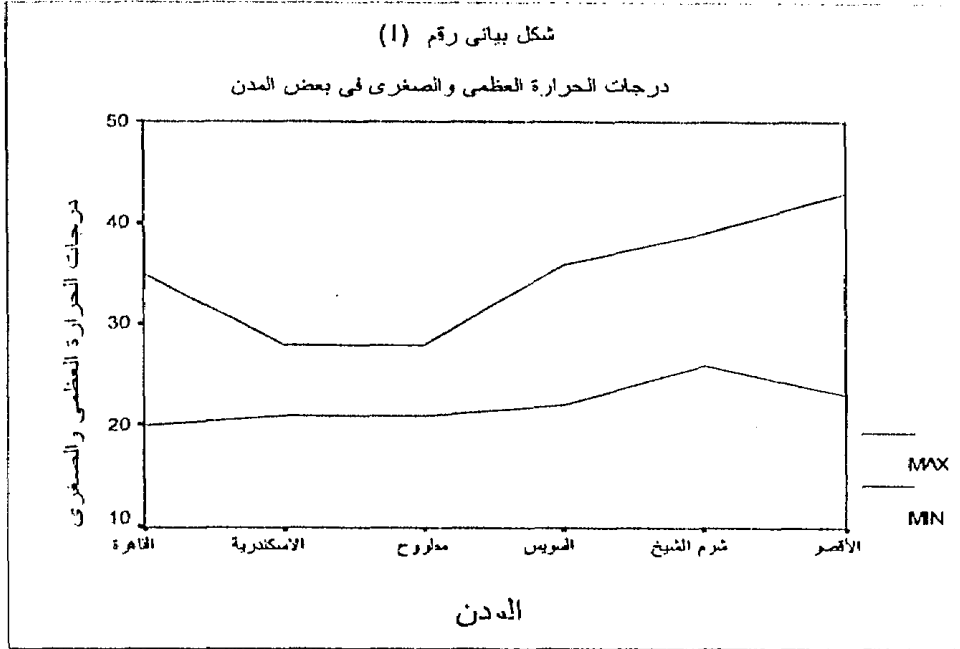
لتدوير الرسم اختر أمر **Format** ، ثم اختر أمر **Swap Axes** يظهر الرسم البيانى بالشكل

التالى



٨- تعديل نوع الرسم :

- لتعديل نوع الرسم (Bar Line Area etc) نتبع الآتي
- من قائمة Gallery اختر النوع المطلوب وليكن Line ، ثم اختر Multiple و اختر Replace يظهر الرسم البياني بالشكل التالي :



التمثيل الدائري للبيانات

إذا كانت البيانات نسباً مئوية ، ومجموع هذه النسب ١٠٠% فيصلح تمثيلها بالقطاعات الدائرية

تمرين : في إحدى السنوات وجد أن التركيب المحصولي في جمهورية مصر العربية كالتالي

مخاصل حقلية	الخضروات	فواكه	نباتات طبية وعطرية
٧٠%	١٠%	١٢%	٨%

والمراد تمثيل هذه البيانات بقطاعات دائرية

٢-٣ : الجداول التكرارية

يتطلب التحليل الاحصائي-في مرحلته الأولى-جمع البيانات المطلوب دراستها وتأتى المرحلة التالية بعد الحصول على البيانات ، وهى تنظيم وعرض بيانات العينة بمجرد الحصول عليها. ويعتبر تبويب البيانات وعرضها من الخطوات الأساسية ، لأنه من النادر استخلاص أية نتائج مفيدة أو التعرف على أية ملامح أساسية لبيانات غير مبنية . ويتم تبويب البيانات من خلال الجداول التكرارية.

خطوات بناء الجدول التكراري:-

أولاً : تحديد عدد الفئات : معظم الإحصائيين يفضلون أن لا يقل عدد الفئات عن ٤ وأن لا يزيد عن ١٥ (وهذه ليست قاعدة جامدة) وذلك لأن اختصار البيانات في عدد قليل من الفئات يرافقه إضاعة الكثير من حقيقة البيانات . واختيار عدد كبير من الفئات يعنى أن هذه البيانات لم تختصر.

ثانياً : تحديد طول الفئة : لتحديد طول الفئة لابد من معرفة مدى التوزيع، والمدى هو الفرق بين أكبر مشاهدة وأصغر مشاهدة ، ثم نحدد عدد الفئات ويكون

$$\text{طول الفئة} = \frac{\text{المدى}}{\text{عدد الفئات}}$$

ثالثاً : حدود الفئات : معظم الإحصائيين يميلون لأن يكون الحد الأدنى للفئة الأولى من مضاعفات طول الفئة بشرط أن يكون أقل من أو يساوى أصغر مشاهدة في التوزيع . أما الحد الأعلى فيكون عبارة عن الحد الأدنى مضافاً إليه طول الفئة

مثال : البيانات التالية توضح نتائج ٣٥ متدرب في دورة تدريبية لتحسين الأداء في إدارة ما، كون التوزيع التكرارى المناسب لهذه البيانات .

٥٠ ، ٥٤ ، ٥٨ ، ٤٨ ، ٤٩ ، ٤٦ ، ٧٨ ، ٤٨ ، ٣٧ ، ٦٢ ، ٤٠ ، ٥٨ ، ٣٣ ، ٥٩ ، ٤٨ ، ٥٤ ، ٣٩ ، ٦٨ ، ٤٤ ، ٦١ ، ٤٧ ، ٤٥ ، ٤١ ، ٥٧ ، ٤٣ ، ٦٣ ، ٥٢ ، ٦٨ ، ٥١ ، ٧٠ ، ٧٣ .

الحل : ١ - المدى = الفرق بين أعلى قيمة - أدنى قيمة = ٧٨ - ٣٣ = ٤٥

٢ - اختيار عدد الفئات وليكن ١٠

$$٣ - \text{طول الفئة بالتقريب} = \frac{٤٥}{١٠} = ٤.٥ \approx ٥$$

٤ - تعيين حدود الفئات، وحيث أن أصغر مشاهدة في التوزيع = ٣٣ وطول الفئة = ٥

فإن الحد الأدنى للفئة يفضل أن يكون أقرب ما يكون إلى اصغر مشاهدة وعلى ذلك فإن الحد الأدنى للفئة الأولى يكون ٣٠ ، والحد الأعلى يكون الحد الأدنى للفئة مضاف إليه طول الفئة أى يساوى ٣٥ وعلى ذلك نحصل على الجدول التكراري التالي :-

(التوزيع التكراري لفئات ٣٥ متدرب في دورة)
تدريبية لتحسين الأداء

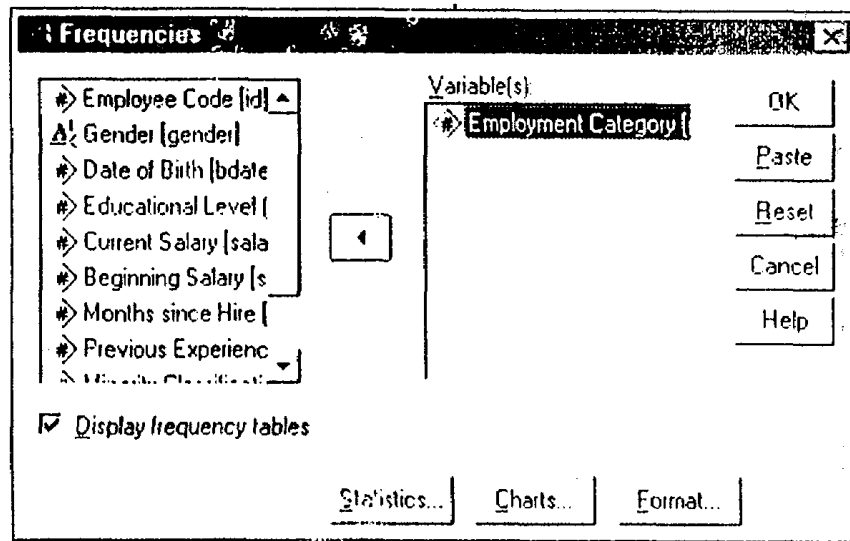
الفئات	العلاقات	التكرارات
٣٠ -	/	١
٣٥ -	//	٢
٤٠ -	////	٤
٤٥ -	//// ###	٩
٥٠ -	/ ###	٦
٥٥ -	###	٥
٦٠ -	///	٣
٦٥ -	//	٢
٧٠ -	//	٢
٧٥ - أقل من ٨٠	/	١

الجدول التكراري للمتغيرات باستخدام البرنامج SPSS

١. افتح الملف Employee Data.sav

٢. افتح قائمة Analyze ، ثم اختر منها Descriptive Statistics ، ومنها اختر

Frequencies ، يظهر الصندوق الحوارى Frequencies



٢. اختر المتغير أو المتغيرات المراد عمل جدول تكرارى لها فى خانة (Variables)

٣. يمكن إظهار الجدول التكرارى بتنشيط خانة Display frequency tables

٤. اختر أمر OK تظهر النتائج التالية :

Frequencies

Statistics

Employment Category

N	Valid	474
	Missing	0

Employment Category

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Clerical	363	76.6	76.6	76.6
	Custodial	27	5.7	5.7	82.3
	Manager	84	17.7	17.7	100.0
	Total	474	100.0	100.0	

ونجد بالنتائج اسم وتوصيف المتغير ، يتبعه جدول به Value Label ، وقيم المتغير Value ، التكرار Frequency ، النسب المئوية Percent ، النسب المئوية بعد استبعاد القيم المفقودة إن وجدت Valid Percent ، ثم التكرار المتجمع Cumulative Percent .

٢-٤ : التحليل الوصفي للبيانات

يتم وصف البيانات من خلال حساب مقاييس النزعة المركزية ومقاييس التشتت وكذا مقاييس الالتواء والتفلطح أو التفرطح (التدبب)

أولاً : مقاييس النزعة المركزية وتشمل :-

١- الوسط الحسابي : (Arithmetic Mean) الوسط الحسابي أو المتوسط لمجتمع

ما لمجموعة من القيم هو مجموع القيم مقسوماً على عددها.

الوسط الحسابي من بيانات غير مبوبة :-

الوسط الحسابي للمجتمع μ

$$\mu = \frac{x_1 + x_2 + \dots + x_n}{N} = \frac{\sum x_i}{N} \quad \text{حيث } N \text{ حجم بيانات المجتمع}$$

الوسط الحسابي للعينه $\bar{x}$:-

$$\bar{x} = \frac{x_1 + x_2 + \dots + x_n}{n} = \frac{\sum x_i}{n} \quad \text{حيث } n \text{ حجم بيانات العينة}$$

الوسط الحسابي للعينه $\bar{x}$:-

من بيانات مبوبة :-

الوسط الحسابي للمجتمع μ

$$\mu = \frac{\sum f_i m_i}{N} \quad \text{حيث } f_i = \text{تكرار الفئة } i$$

$$m_i = \text{مركز الفئة}$$

$$N = \text{حجم بيانات المجتمع}$$

الوسط الحسابي للعينه $\bar{x}$: يستخدم القانون السابق مع إحلال n ، مكان N في المعادلة .

٢- الوسيط : Median

يعرف الوسيط لمجموعة من البيانات هو قيمة الملاحظة الموجودة في منتصف البيانات مرتبة ترتيباً تصاعدياً أو تنازلياً

أولاً : البيانات غير المبوبة :-

لايجاد الوسيط للبيانات غير المبوبة يجب ترتيبها أولاً ترتيباً تصاعدياً أو تنازلياً، فإذا كان عدد المشاهدات فردياً فإن قيمة الوسيط هي قيمة المشاهدات رقم $\frac{(n+1)}{2}$ الموجودة في منتصف البيانات المرتبة.

وإذا كان عدد المشاهدات زوجياً فإن قيمة الوسيط هي متوسط المشاهدين اللتين

$$\text{ترتيبهما } \frac{n}{2} \text{ والتي تليها } \frac{n+1}{2}$$

ثانياً : البيانات المبوبة :-

يمكن إيجاد الوسيط للبيانات المبوبة في جدول تكرارى، باستخدام المعادلة التالية :-

$$\text{Median} = \left[\frac{\frac{n}{2} - F}{f_m} \right] w + L_m$$

حيث n = مجموع المشاهدات في العينة

F = مجموع التكرارات السابقة على الفئة الوسيطة.

f_m = تكرار الفئة الوسيطة.

L_m = الحد الأدنى للفئة الوسيطة.

w = طول الفئة.

٣- المنوال The Mode

المنوال هو القيمة أو الصفة الأكثر شيوعاً (تكراراً) في البيانات.

ويمكن أن يكون للبيانات منوال واحد فقط يسمى التوزيع في هذه الحالة أحادى المنوال

(unimodal) أو منولان (ثنائى المنوال Bimodal) أو متعدد المنوال (Multimodal)،

كما يمكن أن لا يوجد منوال للمشاهدات.

البيانات المبوبة: يمكن إيجاد المنوال للبيانات المبوبة في جدول تكرارى باستخدام المعادلة التالية :-

$$\text{Mode} = L_{m_0} + \left[\frac{d_1}{d_1 + d_2} \right] w$$

حيث d_1 = تكرار الفئة المنوالية - تكرار الفئة التالية (بعد) الفئة المنوالية.

d_2 = تكرار الفئة المنوالية - تكرار الفئة السابقة (قبل) الفئة المنوالية.

w = طول الفئة ، L_{m_0} = الحد الأدنى للفئة المنوالية.

ثانيا : مقاييس التشتت :

تستخدم مقاييس التشتت لوصف كيفية انتشار البيانات على جانبي القيمة المتوسطة ولا يقل ذلك أهمية عن تحديد مراكز البيانات ويعرف انتشار القيم حول القيمة المتوسطة بالتشتت فإذا زاد التشتت فإن هذا يدل على قلة تجانس البيانات أما إذا قل التشتت فهذا يدل على زيادة تجانس البيانات . وإذا كان التشتت مساوياً للصفر فهذا يعنى تطابق قيم جميع المفردات أى تجانسها الكامل وهذا بالطبع نادر الحدوث وإذا تساوى متوسطا عيّنتين فإن هذا لا يعنى أنهما متماثلتان فى تشتت مفرداهما . وترجع أهمية مقاييس التشتت إلى أنه من الممكن وجود عيّنتين فى المفردات لهما نفس القيمة المتوسطة ولكنهما تختلفان فى التشتت . ومن أكثر الطرق فى قياس التشتت التباين والانحراف المعياري والمدى ومعامل الاختلاف .

المدى The Range :

المدى هو عبارة عن الفرق بين أكبر القيم وأصغرها . وفى حالة البيانات المبوبة فإن المدى هو الفرق بين الحد الأعلى لأكبر فئة والحد الأدنى لأصغر فئة أو هو الفرق بين مركز الفئة (لأعلى فئة) ومركز الفئة (لأقل فئة) .

التباين The variance

أولاً : التباين من بيانات غير مبوبة يحسب من المعادلتين التاليتين:-

التباين للمجتمع σ^2 :-

$$\sigma^2 = \frac{\sum (x_i - \mu)^2}{N} = \frac{\sum x_i^2 - N \mu^2}{N}$$

μ = متوسط المجتمع

N = حجم بيانات المجتمع

حيث

التباين للعينة S^2 :

$$S^2 = \frac{\sum (x_i - \bar{x})^2}{n-1} = \frac{\sum x_i^2 - n \bar{x}^2}{n-1}$$

حيث $\bar{x}$ = متوسط العينة

n = حجم بيانات العينة حيث n ليست كبيرة

ثانياً : التباين من بيانات مبوبة:-

وتستخدم المعادلتين التاليتين لحساب التباين من البيانات المبوبة

$$\sigma^2 = \frac{\sum f_i m_i^2 - N \mu^2}{N} \quad \text{تباين المجتمع}$$

$$S^2 = \frac{\sum f_i m_i^2 - n \bar{x}^2}{n-1} \quad \text{تباين الفئة}$$

٣- الانحراف المعياري The standard Deviation

الانحراف المعياري لمجموعة من القيم هو الجذر التربيعي الموجب للتباين

$$S = \sqrt{S^2}, \quad \sigma = \sqrt{\sigma^2}$$

٣- معامل الاختلاف The coefficient of variation

يستخدم معامل الاختلاف (وهو مقياس نسبي) لمقارنة تشتت مجموعتين أو أكثر من البيانات.

ويحسب معامل الاختلاف من المعادلتين التاليتين:-

$$\text{معامل الاختلاف للمجتمع} = cv = \frac{\sigma}{\mu} \times 100$$

$$\text{معامل الاختلاف للعينة} = cv = \frac{s}{\bar{x}} \times 100$$

حساب المقاييس الإحصائية من برنامج SPSS

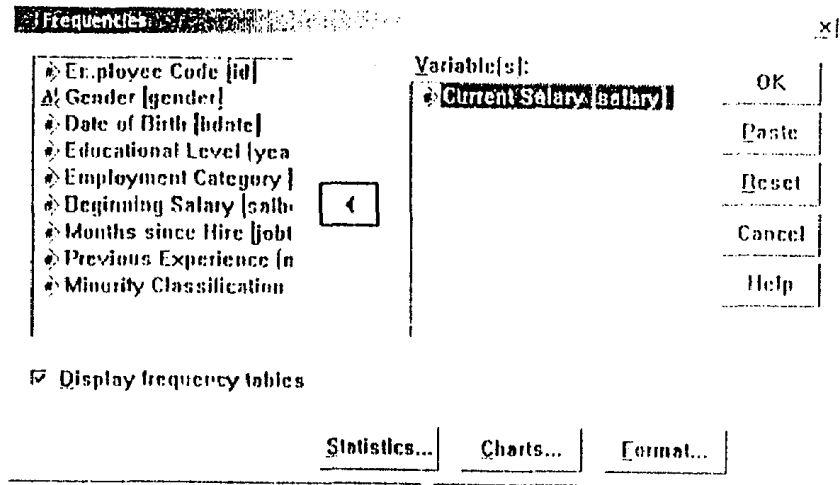
يلاحظ أن هذه المقاييس تصلح للبيانات الكمية

لحساب المقاييس الإحصائية نتبع الآتي :

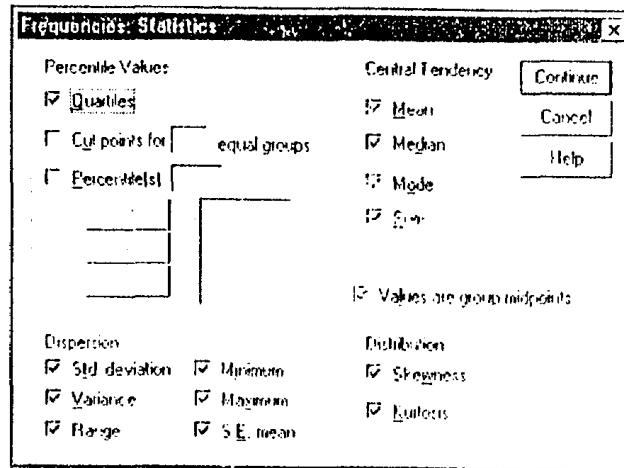
١. افتح الملف Employee Data.sav

٢. من قائمة Analyze اختر Descriptive Statistics ثم اختر Frequencies يظهر

الصندوق الحوارى Frequencies



٣. اختر المتغير Salary من الصندوق الحوارى Frequencies وأنقر أمر Statistics يظهر الصندوق الحوارى التالى :



- تشتمل هذه المقاييس الإحصائية على ما يلى :
- مقاييس التزعة المركزية Central Tendency : وتشمل الوسط الحسابى Mean ، الوسطى Median ، المنوال Mode ، المجموع Sum .
 - مقاييس التشتت Dispersion : وتشمل الانحراف المعياري Std. deviation ، التباين Variance ، المدى Range ، أقل قيمة للمتغير Minimum ، أكبر قيمة للمتغير Maximum ، الخطأ المعياري للمتوسط S.E. mean .
 - مقاييس التوزيع Distribution : وتشمل مقياس الالتواء Skewness ، ومقياس التفرطح Kurtosis .

- مقاييس المدى (الربيعات - العشرية) Percentile Values : وتشمل الربيعات Quartiles ، العشرية Cut points for 10 equal groups ، يمكن تقسيم المدى إلى أى مجموعات متساوية أخرى ، كما يمكن كتابة النسب المطلوبة في خانة (s) Percentile(s) .
- ٢. يمكن تنشيط الخانة Values are group mid points إذا كنا قد أخذنا القيم عن طريق مراكز الفئات وليست البيانات الخام .

٣- بعد اختيار المقاييس الإحصائية اختر أمر Continue تظهر النتائج التالية

Statistics		
Employment Category		
N	Valid	474
	Missing	0
Mean		1.41
Std. Error of Mean		3.55E-02
Median		1.28 ^a
Mode		1
Std. Deviation		.77
Variance		.60
Skewness		1.456
Std. Error of Skewness		.112
Kurtosis		.268
Std. Error of Kurtosis		.224
Range		2
Minimum		1
Maximum		3
Sum		669
Percentiles	25	b.
	50	1.28
	75	1.89

- a. Calculated from grouped data.
- b. The lower bound of the first interval or the upper bound of the last interval is not known. Some percentiles are undefined.
- c. Percentiles are calculated from grouped data.

Frequencies: Charts

Chart Type

☐ None

☐ Bar charts

☐ Pie charts

☒ Histograms

☒ With normal curve

Continue

Cancel

Help

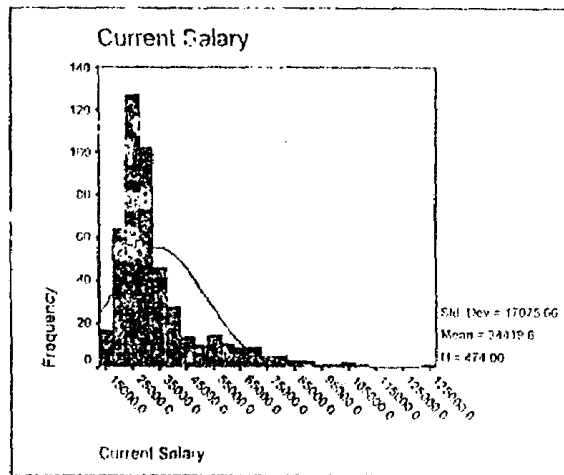
- ٣. لإظهار الرسوم البيانية التي تستخدم في عرض البيانات ، انقر خانة Charts من الصندوق الحوارى Frequencies يظهر الصندوق الحوارى التالى :

ويشتمل هذا الصندوق على مايلي :

- None : في حالة عدم إظهار أى رسوم بيانية
- الأعمدة البيانية Bar chart(s) : يستخدم في حالة إظهار رسم بياني (أعمدة) للبيانات المتقطعة

- المدرج التكرارى Histogram(s) : يستخدم في حالة إظهار الرسم البياني للمتغيرات المستمرة ، كما يمكن إظهار منحنى التوزيع المعتدل باختيار With normal curve .

بعد اختيار أمر Continue ، ثم OK يظهر الرسم البياني التالي :



للتحكم في شكل إظهار الجدول التكرارى ، انقر الأمر Format من الصندوق الحوارى Frequencies يظهر الصندوق الحوارى التالى :

Frequencies: Format

Order by	Multiple Variables	Continue
☒ Ascending values	☒ Compare variables	Cancel
☐ Descending values	☐ Organize output by variables	Help
☐ Ascending count	☐ Suppress tables with more	
☐ Descending count	than categories	

ويشتمل هذا الصندوق على ما يلي :

- Order by أى ترتيب البيانات بالجدول :
- Ascending values أى بناءً على قيم المتغير تصاعدياً ، Descending values أى بناءً على قيم المتغير تنازلياً ، Ascending counts أى بناءً على التكرارات تصاعدياً ، Descending counts أى بناءً على التكرارات تنازلياً .
- Suppress tables with more than 10 categories : لإهمال الجداول التى تزيد عدد القيم للمتغير بها عن ١٠ ، كما يمكن تحديد رقم آخر .
- وفى حالة حساب المقاييس الإحصائية مثل المتوسط الحسابى والوسيط ... الخ لأكثر من متغير Multiple Variables يمكن :
- Compare variables : أى إظهار المقاييس الإحصائية لكل المتغيرات فى جدول
- Organize output by variables : أى إظهار المقاييس الإحصائية لكل متغير فى جدول على حدة

حساب المقاييس الإحصائية من جدول تكرارى باستخدام برنامج SPSS :

قد لا تتوافر لدينا تفاصيل البيانات ، ولكن تعطى فى صورة جدول تكرارى ، ونريد حساب المقاييس الإحصائية لها .

مثال : لدينا الجدول التكرارى التالى :

الترددات (Fre)	قيم (X)
٢٥	٥
٦٤	٨
٤٩	٧
١١٠	١٠
١٢١	١١
١٦	٤
٨١	٩

لحساب المقاييس الإحصائية للمتغير نتبع الخطوات التالية :

- (١) يتم تسجيل بيانات المتغير الأسمى X والتكرارات fre
- (٢) يتم ترجيح البيانات بالمتغير Fre عن طريق اختيار الأمر weight cases من قائمة Data .

٣) بعد ذلك نختار أمر Summarize ، ومنها Frequencies ، ونختار المقاييس الإحصائية بالطريقة السابقة ، تظهر لنا النتائج التالية :

Frequencies

Statistics

X		
N	Valid	456
	Missing	0
Mean		9.0000
Std. Error of Mean		8.901E-02
Median		9.3702 ^a
Mode		11.00
Std. Deviation		1.9008
Variance		3.6132
Skewness		-.960
Std. Error of Skewness		.114
Kurtosis		.270
Std. Error of Kurtosis		.228
Range		7.00
Minimum		4.00
Maximum		11.00
Sum		4104.00
Percentiles	25	7.8584 ^b
	50	9.3702
	75	10.5158

a. Calculated from grouped data.

b. Percentiles are calculated from grouped data.

X

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	4.00	16	3.5	3.5	3.5
	5.00	25	5.5	5.5	9.0
	7.00	49	10.7	10.7	19.7
	8.00	64	14.0	14.0	33.8
	9.00	81	17.8	17.8	51.5
	10.00	100	21.9	21.9	73.5
	11.00	121	26.5	26.5	100.0
	Total	456	100.0	100.0	

التوصيف Descriptives

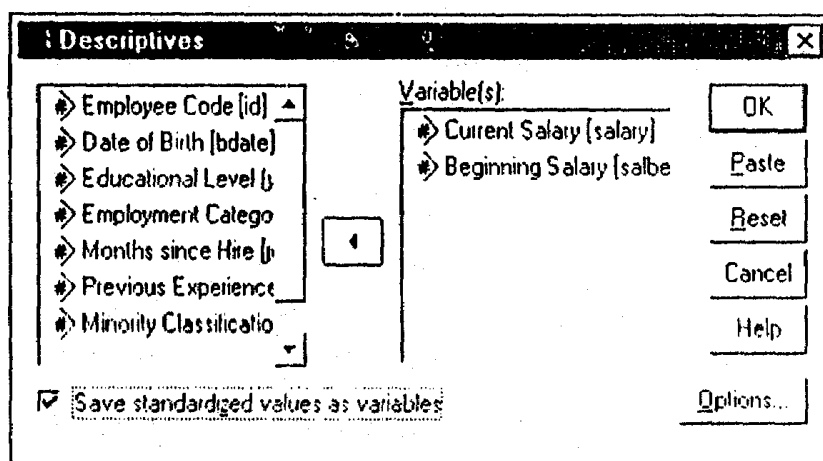
يفضل استخدام هذا الأمر مع المتغيرات الكمية

الأمر Descriptives يعطى توصيف ومقاييس إحصائية لعدة متغيرات في جدول واحد مما يسهل عملية المقارنة ، كما يمكن حساب متغير عياري (Z scores) يسهل عملية المقارنة إذا اختلفت وحدات القياس ، كما يمكن حساب العديد من المقاييس الإحصائية التي يتم حسابها عن طريق الـ Frequencies .

ولتوصيف Descriptives لمتغير أو لعدة متغيرات نتبع الآتي :

١. من قائمة Statistics اختر أمر Analyze ، ومنها اختر Descriptive Statistics ، يظهر الصندوق الحوارى التالى :
٢. أمام خانة Variable(s) يتم اختيار المتغير أو المتغيرات التي يراد توصيف Descriptive لها .

٣. إذا أردنا حساب متغير عياري (Z scores) فإننا نشط الخانة Save Standardized



٤. Values as variables ، في هذه الحالة تحسب قيم المتغير العياري لكل متغير وتدرج بملف البيانات ، إذا أردنا الاحتفاظ بهذه المتغيرات فينبغى حفظ ملف البيانات
٤. إذا أردنا حساب المقاييس الإحصائية للمتغيرات ، فإننا ننقر الأمر Options يظهر الصندوق الحوارى التالى :

Descriptives: Options

☒ Mean ☒ Sum

Dispersion

☒ Std deviation ☒ Minimum

☒ Variance ☒ Maximum

☒ Range ☒ S.E. mean

Distribution

☐ Kurtosis ☐ Skewness

Display Order

☒ Variable list

☐ Alphabetic

☐ Ascending means

☐ Descending means

من هذا الصندوق يتم تنشيط المقاييس الإحصائية المطلوبة ، فتظهر النتائج التالية :

Descriptives

Descriptive Statistics

	Statistic			Std. Error	
	Current Salary	Beginning Salary	Valid N (listwise)	Current Salary	Beginning Salary
N	474	474	474		
Range	\$119,250	\$70,980			
Minimum	\$15,750	\$9,000			
Maximum	\$135,000	\$79,980			
Sum	\$16,314,875	\$8,065,625			
Mean	\$34,419.57	\$17,016.09		\$784.31	\$361.51
Std. Deviation	\$17,075.66	\$7,870.64			
Variance	291578214.453	61946944.959			
Skewness	2.125	2.853		.112	.112
Kurtosis	5.378	12.390		.224	.224

٥-٢ استكشاف البيانات Explore

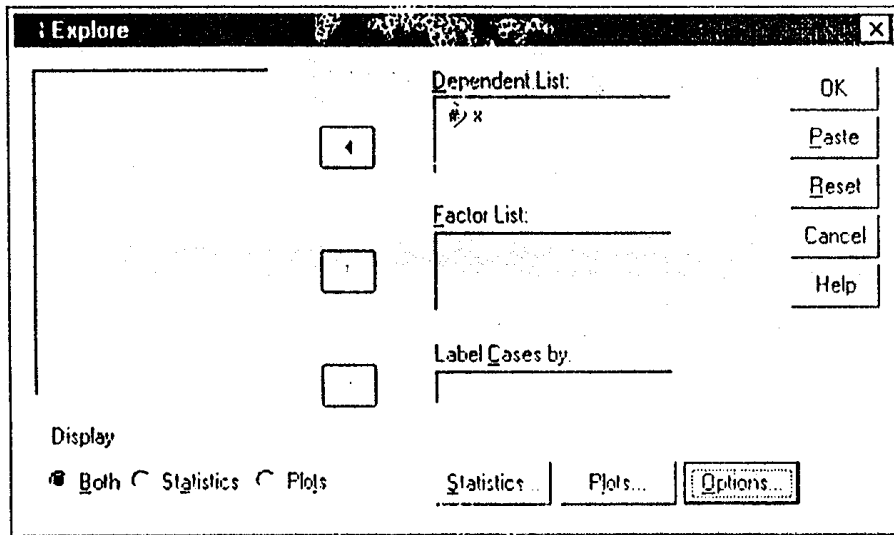
نحتاج لاختبار البيانات لمعرفة القيم الشاذة (خارج نطاق معين Outlire) ، معرفة توزيع البيانات لمعرفة إمكانية تطبيق المقاييس الإحصائية عليها أو أنها تحتاج إلى بعض التحويلات ، كما نحتاج إلى بعض الرسوم البيانية للمقارنة بين مجموعات البيانات .

مثال : في عينة لملاحظات عن المتغير العشوائى الذى يعبر عن شدة الزلازل التى حدثت في أحد المناطق مقاسة بمقياس ريختر ، اختبر هذه البيانات .

١,٣, ١,٤, ٦,٣, ١,٩, ٢,٤, ١,١, ١,٢, ٥,١, ١,١, ٣,١, ٨,٣, ١,٥, ٧,٧, ١,٢, ٢,٢, ٥,٤, ١,٣, ٢,٤, ٢,٧, ١,٤, ٢,١, ٢,١, ٢,٢, ٣,٣, ١,٥

ولاختبار البيانات نتبع الخطوات التالية :

١. من قائمة Statistics اختر أمر Analyze ، ومنها اختر Descriptive statistics ومنها اختر Explore يظهر الصندوق الحوارى التالى :

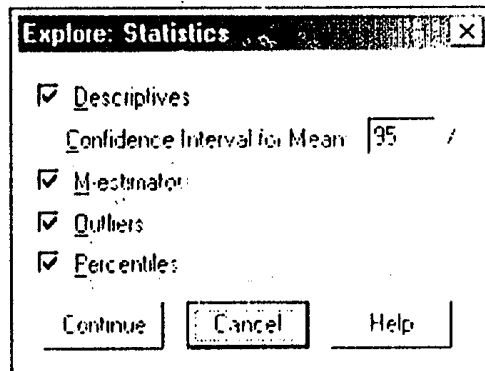


٢. أمام خانة Dependent List اختر المتغير أو المتغيرات المراد اختبارها .

٣. أمام خانة Factor List اختر المتغير الذى يقسم هذه المتغيرات إلى مجموعات

٤. أمام خانة Label Cases by اختر المتغير الذى يراد ترقيم الحالات على أساسه ، ولإظهار

المقاييس الإحصائية ، انقر أمر Statistics فيظهر الصندوق الحوارى التالى :



يشتمل هذا الصندوق على ما يلي :

- Descriptives : بحسب المقاييس الإحصائية السابق التعرف عليها (مقاييس التوزع المركزية - مقاييس التشتت - مقاييس التوزيع)
 - M-estimators حساب تقديرات الوسط الحسابي والوسيط للمجتمع من العينة ،
 - Outliers : يظهر أكبر ٥ قيم ، وأصغر ٥ قيم للمتغير مع إظهار أرقام الحالات .
 - Percentiles : مقاييس تقسيم البيانات إلى نسب مئوية .
- ولإظهار الرسوم البيانية الخاصة باختبار البيانات Explore ، نقر Plots من الصندوق الحوارى Explore فيظهر الصندوق الحوارى التالى

بعد هذه الاختيارات تظهر لنا النتائج التالية :

Explore

Case Processing Summary

	Cases					
	Valid		Missing		Total	
	N	Percent	N	Percent	N	Percent
X	29	100.0%	0	.0%	29	100.0%

Descriptives

			Statistic	Std. Error
X	Mean		2.8069	.3668
	95% Confidence Interval for Mean	Lower Bound	2.0556	
		Upper Bound	3.5582	
	5% Trimmed Mean		2.6125	
	Median		2.1000	
	Variance		3.901	
	Std. Deviation		1.9750	
	Minimum		1.00	
	Maximum		8.30	
	Range		7.30	
	Interquartile Range		2.3000	
	Skewness		1.529	.434
	Kurtosis		1.763	.845

M-Estimators

	Huber's M-Estimator ^a	Tukey's Biweight ^b	Hampel's M-Estimator ^c	Andrews' Wave ^d
X	2.2495	2.0547	2.2537	2.0552

a. The weighting constant is 1.339.

b. The weighting constant is 4.685.

c. The weighting constants are 1.700, 3.400, and 8.500

d. The weighting constant is $1.340 \cdot \pi$.

Percentiles

		Percentiles					
		5	10	25	50	75	95
Weighted Average(Definition 1)	X	1.0000	1.1000	1.3500	2.1000	3.8500	6.0000
Tukey's Hinges	X			1.4000	2.1000	3.3000	

Extreme Values

		Case Number	Value
X	Highest	1	2
		2	28
		3	13
		4	5
		5	25
	Lowest	1	1
		2	7
		3	8
		4	4
		5	27

a. Only a partial list of cases with the value 1 are shown in the table of lower extremes.

Tests of Normality

	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
X	.207	29	.003	.811	29	.010**

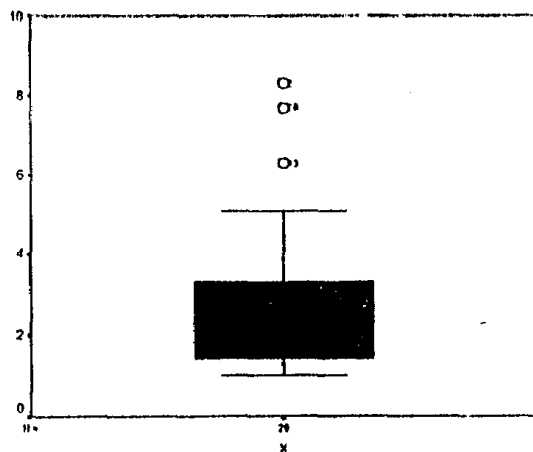
** . This is an upper bound of the true significance.

a. Lilliefors Significance Correction

X Stem-and-Leaf Plot

Frequency	Stem &	Leaf
9.00	1 .	001122344
3.00	1 .	599
6.00	2 .	011224
1.00	2 .	7
3.00	3 .	003
.00	3 .	
2.00	4 .	00
.00	4 .	
2.00	5 .	01
3.00	Extremes	(>=6.3)

Stem width: 1.00
Each leaf: 1 case(s)



الباب، الثالث

اختبارات الفروض الإحصائية

تعتبر اختبارات الفروض الإحصائية **Testing statistical Hypothesis** واحدة من أهم التطبيقات التي قدمها علم الإحصاء كحل للمشاكل العلمية المختلفة بشتى فروع العلم ، باستخدام نظرية الاحتمالات وخصائص التوزيعات العينية أمكن التعرف على ما يسمى باختبارات الفروض الإحصائية ، ومن خلالها يمكن لاي شخص أن يتخذ قرار برفض أو قبول فرض معين أو مجموعة من الفروض المتعلقة بمشكلة معينة، موجودة في الحياة العامة . وقد تزايد الاهتمام باختبارات الفروض في السنوات الأخيرة حتى أصبحت تدخل في شتى فروع العلم (الطب-الزراعة-الفلك ، الاقتصاد، الصناعة ، علم النفس... الخ) مما أدى تطورها تطوراً هائلاً وملحوظاً.

وتعتبر اختبارات الفروض الإحصائية أسلوباً لاتخاذ قرار بخصوص فرضين مختلفين بناء على بيانات العينة المتاحة والاختبارات الإحصائية المناسبة. والفرض الاحصائي هو عبارة عن صياغة مبدئية حول واحد أو أكثر من معالم المجتمع.

٣-١ : الخطوات الأساسية في اختبارات الفروض الإحصائية :-

تتضمن اختبارات الفروض عدة خطوات هي :-

١- التعرف على توزيع المجتمع الأصلي

٢- صياغة الفروض

أ - فرض العدم (H_0) ويسمى أيضا الفرض الصفري.

ب- الفرض البديل (H_1)

٣- تحديد مستوى المعنوية (α)

٤- صياغة قاعدة القرار وتتضمن :-

أ - تحديد إحصائية الاختبار $t, z, \dots$

ب- تحديد المنطقة الحرجة والتي تتضمن منطقى القبول والرفض.

٥- اتخاذ القرار بناء على المعلومات السابقة وذلك برفض أو عدم رفض (قبول) الفرض

الصفري (فرض العدم).

وستتناول كل منها بالتفصيل كما يلي

أولاً : التعرف على توزيع المجتمع الأصلي :-

يجب التعرف على توزيع المجتمع الأصلي وتحديد ما إذا كان يتبع التوزيع الطبيعي أو توزيع ذى الحدين أو ما إلى ذلك .

ثانياً : صياغة الفروض :-

هناك فرضان : الفرض الأول وهو فرض العدم ويرمز له بالرمز (H_0) (أو الفرض الصفري)
والفرض الثانى وهو الفرض البديل ويرمز له بالرمز (H_1)
ويصاغ الفرض الصفري أو فرض العدم كالتالى

$$H_0 : \mu = \mu_0$$

أى أنه صياغة مبدئية عن معلمة من معالم المجتمع المجهولة (وليكن وسط المجتمع μ) ،
 μ_0 قيمة محددة

أما الفرض البديل μ_1 فهو صياغة مبدئية تشير إلى نفس المعلمة المجهولة لها قيمة تختلف عن القيمة التى حددها فرض العدم . أى أنه يتم قبول الفرض البديل فى حالة رفض فرض العدم.
ويأخذ الفرض البديل إحدى الصور التالية:-

$$\underline{1} \quad H_1 : \mu \neq \mu_0$$

وذلك فى حالة عدم توافر معلومات عن معلمة المجتمع إذ قد تكون أكبر أو أصغر .

$$\underline{2} \quad H_1 : \mu < \mu_0$$

$$\underline{3} \quad H_1 : \mu > \mu_0$$

وتستخدم هذه الصيغ حينما تكون μ_0 هي مواصفات قياسية أو قيمة يجب أن تحقق وتشك فى تحقيقها.

ثالثاً : تحديد مستوى المعنوية (α) :-

يقيس مستوى المعنوية α درجة عدم التأكد ، فلو كان لدينا فترة ثقة ٩٥% فيكون لدينا عدم تأكد مقداره ٥% أى أن مستوى المعنوية هو احتمال رفض فرض العدم وهو صحيح ويسمى ذلك خطأ من النوع الأول كما يتضح من الجدول :

القرار	الخطأ	نوع الخطأ
رفض H_0 وهو صحيح	خطأ من النوع الأول	α
قبول H_0 وهو خطأ	خطأ من النوع الثاني	β

رابعاً : صياغة قاعدة القرار :-

أ- تحديد إحصائية الاختبار :- وتحدد بناء على نوع التوزيع الاجمالي لمجتمع الدراسة وعلى حجم العينة n بصفة عامة .

١- فإن كان التوزيع طبيعي وكانت $n \geq 30$ والانحراف المعياري للمجتمع σ معلوم . فتكون إحصائية الاختبار هي

$$Z = \frac{\bar{X} - \mu}{\sigma}$$

٢- وإذا كان توزيع المجتمع طبيعياً وكانت $n < 30$ والانحراف المعياري للمجتمع غير معلوم فتكون إحصائية باستخدام توزيع t حيث

$$t = \frac{\bar{X} - \mu}{\sigma}$$

٣- أما إذا كان توزيع المجتمع يتبع توزيع كاي تربيع χ^2 أو توزيع F فإن إحصائية الاختبار تكون مختلفة وسوف نتعرض له بالتفصيل في أجزاء تالية :

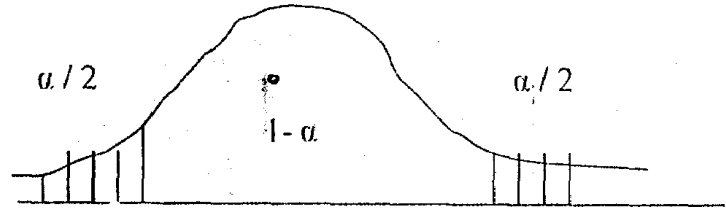
ب- تحديد المنطقة الحرجة : وهي أما منطقة رفض أو منطقتي رفض فرض العدم . وهي أما تكون على جانبي المنحنى من (ذيلين) أو في جانب واحد (من ذيل واحد) طبقاً لطبيعة الاختبار البديل .

ويعتمد تحديد المنطقة الحرجة على كل من حجم العينة (n) ومستوى المعنوية (α) وقيمة إحصائية الاختبار المستخدم وطبيعة الفرض البديل .

والإشكال التالية تبين منطقة الرفض (المنطقة المظللة) وتسمى المنطقة الحرجة ، منطقة عدم الرفض (الغير مظللة)

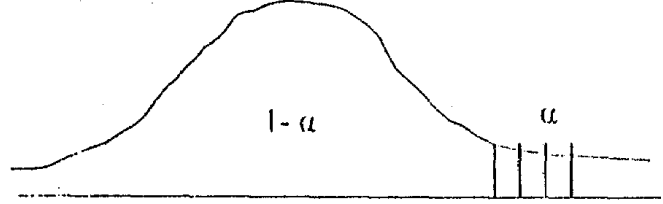
$$H_0 : \mu = \mu_0$$

$$H_1 : \mu \neq \mu_0$$



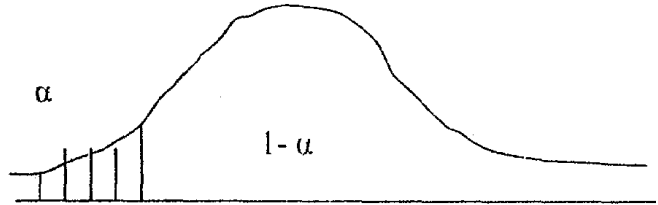
$$H_0 : \mu = \mu_0$$

$$H_1 : \mu > \mu_0$$



$$H_0 : \mu = \mu_0$$

$$H_1 : \mu < \mu_0$$



خامساً : اتخاذ القرار

توجد أربع قرارات ممكنة في أى مشكلة من مشاكل اختبارات الفروض وهى كما يلى :-

- ١- رفض فرض العدم وهو صحيح ويسمى هذا بخطأ النوع الأول وتمثل (α) احتمال وقوع هذا الخطأ.
- ٢- عدم رفض فرض العدم وهو صحيح وهذا قرار سليم واحتمال اتخاذ هذا القرار هو ($1 - \alpha$).
- ٣- عدم رفض العدم وهو خاطئ ويسمى هذا بخطأ من النوع الثانى وتمثل بيتا β احتمال وقوع هذا الخطأ.
- ٤- رفض فرض العدم وهو خاطئ قرار سليم واحتمال اتخاذه ($1 - \beta$). ويوضح الجدول التالي ماسبق:-

القرار	H_0 صائب	H_0 خاطئ
رفض H_0	قرار خاطئ خطأ من النوع الأول (α)	قرار سليم ($1 - \beta$)
عدم رفض H_0	قرار سليم ($1 - \alpha$)	خطأ من النوع الثاني β

والاختبارات التي سوف نتناولها هي اختبارات معلمية تتعلق بعينة واحدة أو عينتين .
وفي هذه الاختبارات الفروض تدور حول معلمة المجتمع المجهولة وهي الوسط الحسابي والنسبة والتباين.

ولإجراء الاختبارات السابقة نحتاج إلى ما يعرف باسم احصائي الاختبار **Test statistic** تحدد بمعرفة بعض النظريات الإحصائية وهي عبارة عن حقائق ونظريات تم برهانها ويمكن الرجوع إليها في مراجع الإحصاء .

وهناك طريقتان لاتخاذ القرار في خامساً وذلك أما برفض أو قبول الفرض العدمي ، الطريقة الأولى تعتمد على قيم محسوبة لاحصائي الاختبار وتقيم بتدولية له تستخرج من جداول أعدت خصيصاً لذلك الفرض ، وتوجد طريقة أخرى لرفض أو قبول الفرض العدمي لاتتطلب وجود جداول بل تعتمد على قيم تسجلها الحزمة SPSS أو أية حزمة احصائية نلخص الطريقتين كالآتي :-

الطريقة الأولى لرفض أو قبول الفرض العدمي البديل :-

أ- الاختبار الاحصائي قد يكون من طرف **One tail test** أو من طرفي **Two tail test** إذا كان الاختبار من طرفين توجد قيمتين حرجيتين نحصل عليهما من جداول التوزيع العيني الذي يتحكم في الاختبار أحدهما في الذيل الأيمن والآخر في الذيل الأيسر (طرفين) للتوزيع العيني تحت الفرض العدمي هذه القيم تحجز مساحة في كل ذيل مقدارها $\alpha/2$ ، وإذا وقعت القيمة المحسوبة لاحصائي الاختبار من بيانات العينة بين هاتين القيمتين نقبل الفرض العدمي خلاف ذلك نقبل البديل .

ب- إذا كان الاختبار من طرف واحد (طرف أيمن) فإن منطقة الرفض تكون من جهة اليمين فنحسب القيمة الجدولية التي تحجز مسافة في الذيل مقدارها α ، إذا كانت القيمة المحسوبة اقل من القيمة الجدولية نقبل العدمى والعكس صحيح . وإذا كانت منطقة الرفض جهة اليسار (طرف أيسر) نحسب القيم الجدولية التي تحجز مساحة في الذيل مقدارها α . إذا كانت القيمة المحسوبة اكبر من القيمة الجدولية نقبل العدمى والعكس صحيح .

الطريقة الثانية لرفض أو قبول الفرض العدمى البديل:-

هناك طريقة أخرى للرفض والقبول استخدمت كثيراً في الآونة الأخيرة وتعتمد على احتمال محسوب يسمى p -value ويرمز له في الحزمة SPSS وبالرمز sig ويعرف بأنه مستوى معنوية محسوب أو خطأ من النوع الأول محسوب ، هذه الطريقة تستخدم كثيراً في البرامج الجاهزة ، وتتلخص هذه الطريقة في الآتي :-

١- يجرى الاختبار بنفس الخطوات السابق الحديث عنها لانهسب قيم جدولية لاحصائى الاختبار ولكن نستخرج احتمال محسوب يرمز له بالرمز sig ، تعبر هذه القيم عن مساحة معينة في ذيل (أو ذيلين) للتوزيع العيى لاحصائى الاختبار وتسمى قيمة مستوى المعنوية المحسوب .

٢- إذا كان الاختبار من طرفين نقارن هذه القيمة بقيمة $\alpha/2$ ، إذا زادت القيمة المحسوبة عن $\alpha/2$ نقبل الفرض العدمى والعكس صحيح . إذا كان الاختبار من طرف واحد نقارن القيمة المحسوبة بقيمة α ، إذا ما زادت عنها نقبل الفرض العدمى والعكس صحيح .

٣-٢ : اجراء اختبارات الفروض الاحصائية باستخدام برنامج SPSS

٣-٢-١ : اختبار ت (T test) : للمقارنة بين المتوسطات :

يستخدم اختبار ت (T test) للمقارنة بين متوسطى عيتين لمعرفة هل الفرق بين هذين المتوسطين معنوياً أو غير معنوى ، ويستخدم هذا الاختبار في البيانات الكمية ، قد تكون هاتان العيتان مستقلتان ، أو غير مستقلتين ، أو قد تكون عينة واحدة ويراد مقارنتها بقيمة معينة .

٩ - اختبارات عينة واحدة

لإجراء اختبار عينة واحدة ، أى مقارنة متوسط عينة واحدة بقيمة معينة ، فعلى سبيل المثال لاختبار أن متوسط الراتب السنوى للعاملين فى الملف Employee Data.sav يعادل ٣٠٠٠٠ دولار نتبع الآتى :

- افتح الملف Employee Data.sav

- من قائمة Analyze اختر أمر Compare Means ، ثم One Sample T Test فيظهر

الصندوق الحوارى One Sample T Test

One-Sample T Test

Test Variable(s):

- Current Salary [salary]

Test Value: 30000

Buttons: OK, Paste, Reset, Cancel, Help, Options...

- اختر المتغير Salary وهو المتغير المراد اختبار المتوسط بالنسبة له فى خانة

Test Variable (s)

- اكتب المتوسط المراد المقارنة به هو ٣٠٠٠٠ فى خانة Test Value

- اختر أمر Ok فتظهر النتائج التالية :

One-Sample Statistics

	N	Mean	Std. Deviation	Std. Error Mean
SALARY Current Salary	474	\$34,419.57	\$17,075.661	\$784.311

One-Sample Test

	Test Value = 30000					
	t	df	Sig. (2-tailed)	Mean Difference	95% Confidence Interval of the Difference	
					Lower	Upper
SALARY Current Salary	5.635	473	.000	\$4,419.57	2,878.40	5,960.73

من النتائج نلاحظ ما يلى :

- ١- الفرض العدمي : متوسط الراتب السنوي لا يختلف معنويًا عن ٣٠٠٠٠ دولار .
 - ٢- الفرض البديل : يوجد فرق معنوي بين متوسط الراتب السنوي والقيمة ٣٠٠٠٠ دولار .
- القرار المتخذ: نلاحظ أن Sig (2-tailed) أقل من مستوى المعنوية (٠,٠٥) لذلك نرفض الفرض العدمي ، ونقبل الفرض البديل أي أنه هناك اختلاف معنوي بين متوسط المرتبات والقيمة ٣٠٠٠٠ دولار .

٢- اختبار عينتين مستقلتين Compare Means

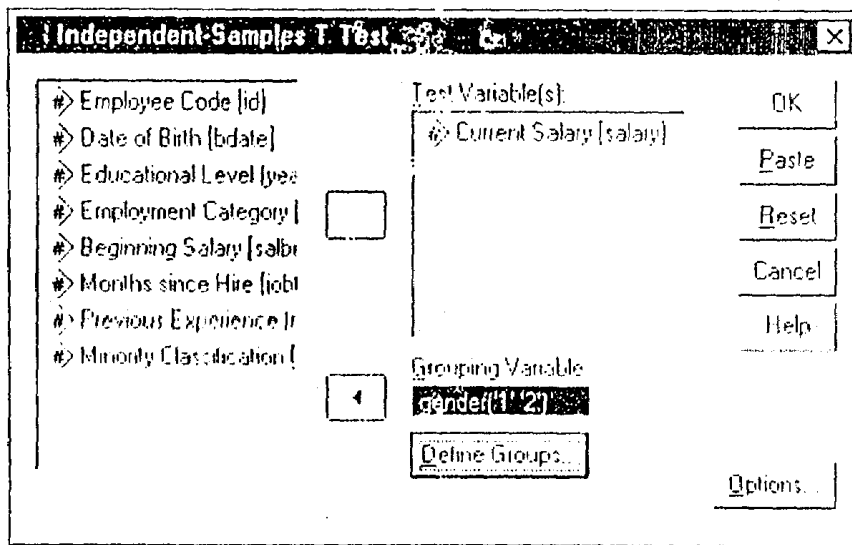
المقارنة بين عينتين مستقلتين :

في هذه الحالة يكون لدينا متغير تابع يتم تقسيمه إلى مجموعتين بناءً على متغير آخر مستقل ، ويراد معرفة معنوية الفرق بين هذين المتوسطين ، ولتنفيذ ذلك تابع الخطوات التالية :

١- افتح الملف Employee Data.sav

٢- من قائمة Analyze اختر أمر Compare Means ، ومنها اختر Independent-

samples T Test يظهر الصندوق الحواري التالي :



٢- أمام خانة Test Variable(s) اختر المتغير أو المتغيرات المراد المقارنة بين متوسطاتها ، وليكن salary

Define Groups

Group 1:

Group 2:

٣- أمام خانة Variable(s) Grouping اختر المتغير الذى يقسم المتغيرات السابقة إلى مجموعتين ، وليكن gender.

٤- اختر Define Groups لتحديد قيمتي هذا المتغير .

٥- اختر Options لتحديد درجة الثقة Confidence Interval ، إمكانية إظهار Labels ، وطريقة معاملة الـ Missing Values .

Independent-Sample T-Test

Confidence Interval: %

Missing Values

☒ Exclude cases analysis by analysis

☐ Exclude cases listwise

بعد تنفيذ الاختيارات السابقة تظهر النتائج التالية :

T-Test

Group Statistics

	Gender	N	Mean	Std. Deviation	Std. Error Mean
Current Salary	Male	258	\$41,441.78	\$19,499.21	\$1,213.97
	Female	216	\$26,031.92	\$7,558.02	\$514.26

Independent Samples Test

		Equal variances assumed	Equal variances not assumed
Levene's Test for Equality of Variance Sig.		119.669	
t-test for Equality of Means		10.945	11.688
	df	472	344.262
	Sig. (2-tailed)	.000	.000
	Mean Difference	\$15,409.86	\$15,409.86
	Std. Error Difference	\$1,407.91	\$1,318.40
	95% Confidence Interval of the Difference Lower	\$12,643.32	\$12,816.73
	Upper	\$18,176.40	\$18,003.00

نلاحظ ما يلي

البرنامج يقوم باختبار للتجانس أولاً ، وبحسب قيمة T في الحالتين ، في حالة التجانس Equal variances assumed ، وفي حالة عدم التجانس Not equal variances assumed ،

ولإجراء اختبار التجانس يصاغ الفرض العدمي والبديل ، والقرار المتخذ كالآتي

١- الفرض العدمي : العينتان متجانستان ، أى لا يوجد فرق معنوي بين تباينهما

٢- الفرض البديل : العينتان غير متجانستين .

القرار المتخذ : قيمة $F = 119.669$ بمستوى معنوية أقل من 0.0005 ، لذا فإننا نرفض

الفرض العدمي ، ونقبل الفرض البديل ، أى أن العينتين غير متجانستين ، ولإجراء اختبار T

وهو الاختبار الأساسي للفرق بين المتوسطين ، يكون اختبارنا بناءً على الاختبار Equal

variance not assumed ، ويصاغ الفرض العدمي والبديل والقرار المتخذ كالتالي :

١- الفرض العدمي : ليس هناك اختلاف بين متوسطي مرتبات الذكور والإناث أى (المتوسطان

متساويان) .

٢- الفرض البديل : يوجد اختلاف بين متوسطي أعمار الذكور والإناث أى (المتوسطان

مختلفان) .

٣- القرار المتخذ : نلاحظ أن قيمة $T = 11.668$ ، بدرجات حرية $= 334,262$ ،

ومستوى المعنوية (Sig (tailed) أقل من (0.05) لذلك نرفض الفرض العدمي ، ونقبل

الفرض البديل أى أنه هناك اختلاف بين متوسطي مرتبات الذكور والإناث ، أى أن متوسط

مرتبات الذكور $41,441.78$ أعلى من متوسط مرتبات الإناث $26,03,92$ حيث يوجد

فرق معنوي بين المتوسطين .

ملحوظة : في حالة Equal variance not assumed يجرى البرنامج تحويله على قيمة T ، وعلى درجات الحرية لتخليصهما من اثر عدم التجانس ، لذلك من المحتمل أن نجد درجات الحرية بها قيمة كسرية

٣- المقارنة بين عيتين غير مستقلتين :

قد يراد المقارنة بين عيتين مرتبطتين ، كما لو أننا قراءات معينة لمغبر معين قبل إجراء معالجة معينة ، وبعد إجراء هذه المعالجة ، كمتوسط المبيعات لمجموعة معينة قبل إجراء دورة تدريبية في فن البيع ، وبعد إجراء هذه الدورة ، ويراد معرفة هل الفرق بين هذين المتوسطين معنوياً أم لا .
مثال :

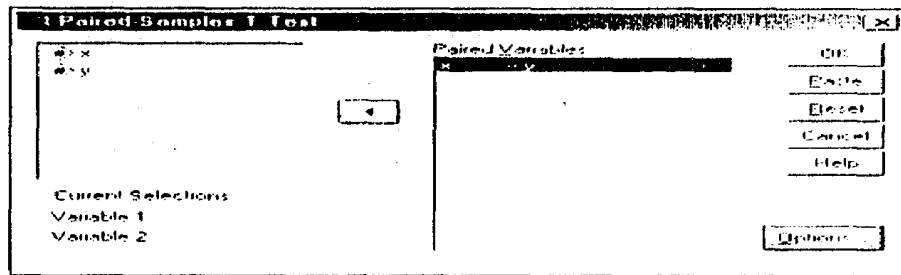
أراد طبيب أن يقرر أثر نوع معين من الدواء على تخفيض ضغط الدم فاختار عينة من ١٥ فرداً ، وأخذ قياس ضغط الدم قبل تعاطي الدواء (X) ، وبعد تعاطي الدواء (y) .

٨٤	٦٨	٧٤	٩٢	٧٤	٦٤	٨٢	٧٨	٧٢	٧٦	٧٦	٧٦	٧٢	٨٠	٧٠	x
٧٤	٧٢	٧٤	٦٠	٧٤	٧٢	٦٤	٥٢	٦٨	٦٦	٥٧	٧٠	٦٢	٧٢	٦٨	y

ولتنفيذ ذلك تابع الخطوات التالية :

١- سجل البيانات في ملف Data ، ولتكن اسماء المتغيرات x , y .

٢- من قائمة Analyze اختر أمر Compare Means ، ومنها اختر Paired-samples T Test يظهر الصندوق الحوارى التالى :



٢- اختر المتغيرين المراد المقارنة بين متوسطيهما ، وهما المتغيران X, Y ، وانقلهما إلى القائمة Paired Variables ، ثم اختر أمر OK تظهر النتائج التالية :

T-Test

Paired Samples Statistics

		Mean	N	Std. Deviation	Std. Error Mean
Pair 1	X	75.2000	15	6.7103	1.7326
	Y	67.0000	15	6.7718	1.7485

Paired Samples Correlations

	N	Correlation	Sig.
Pair 1 X & Y	15	-.270	.330

Paired Samples Test

		Paired Differences				t	df	Sig. (2-tailed)
		Mean	Std. Deviation	Std. Error Mean	95% Confidence Interval of the Difference			
					Lower Upper			
Pair 1	X - Y	8.2000	10.7451	2.7744	2.2496 14.1504	2.956	14	.010

يصاغ الفرض العدمى والبديل كالتالى :

الفرض العدمى : ليس هناك فرق معنوى بين المتوسطين (أى قبل تعاطى الدواء ، وبعد تعاطى الدواء) .

الفرض البديل : يوجد فرق معنوى بين المتوسطين

القرار المتخذ : حيث أن الـ (Sig. (2-tailed) أقل من ٠.٥ ، لذلك نرفض الفرض العدمى القائل بتساوى المتوسطين (قبل وبعد الدواء) ونقبل الفرض البديل القائل بأن هناك فرق فى متوسط ضغط الدم ، حيث أنه انخفض بعد الدواء إلى ٦٧ فى حين كان قبل الدواء ٧٥,٢ أى أن الدواء له أثر على تخفيض ضغط الدم

٣-٢-٢ اختبار F لتحليل التباين :

اختبار تحليل التباين يعتمد أساساً على احصائى اختبار يطلق عليه اسم F نسبة إلى توزيع احتمالى شهير يسمى F Distribution هذا التوزيع له تطبيقاته العديدة فى اختبار الفروض ، ويرتبط تحليل التباين بعلم تصميم التجارب الذى يستخدم بكثرة فى مجال البحوث الزراعية ومجال اختبار فعالية طرق تدريب وتدریس مختلفة أو أنواع مختلفة من الأدوية ... وهو يعتبر اختبار للمتوسطات ولكن بين أكثر من عيتين.

وبصفة عامة في اختبار تحليل التباين يحسب تقديراً إجمالياً لتباين المجتمع ثم يقسم إلى جزئين الأول يسمى التباين بين المجموعات **Between groups** والثاني يسمى التباين داخل المجموعات **Within groups** ثم بعد ذلك يتم حساب احصائي الاختبار يعتمد على النسبة بين هذين التباينين .

ويكون الفرض العدمي والبديل لهذا الاختبار كالتالي:-

الفرض العدمي : متوسطات المجتمعات المحسوبة منها العينات متساوية

الفرض البديل : يوجد زوج واحد على الأقل من المجتمعات المحسوبة منها العينات متوسطاته مختلفة معنوية .

ويشترط لاستخدام اختبار تحليل التباين توفر عدة شروط أهمها:-

١- أن تكون العينات عشوائية مستقلة ، ويتم التحقق من هذه الشروط عند سحب العينات.

٢- أن تكون العينات مسحوبة من مجتمعات موزعة توزيعاً طبيعياً.

٣- أن تكون وحدة القياس بفترة (بيانات مستمرة) .

٤- أن تكون تباينات المجتمعات متساوية . ويستخدم في ذلك عدد من الاختبارات

مثل اختبار " بار تليت **Bartlett** " أو اختبار " هارتلي **Hartly** "

لإجراء اختبار تحليل التباين يكون لدينا متغير مستقلاً ، المتغير الذي حدث في ذلك المتغير يمكن أن يكون في اتجاه واحد طبقاً لشروط معينة، والاختبار المستخدم في هذه الحالة يسمى اختبار في اتجاه واحد . وإذا كنا نهتم بالتغير الذي حدث على المتغير المستقل في اتجاهين فإن الاختبار في هذه الحالة يسمى تحليل التباين في اتجاهين.

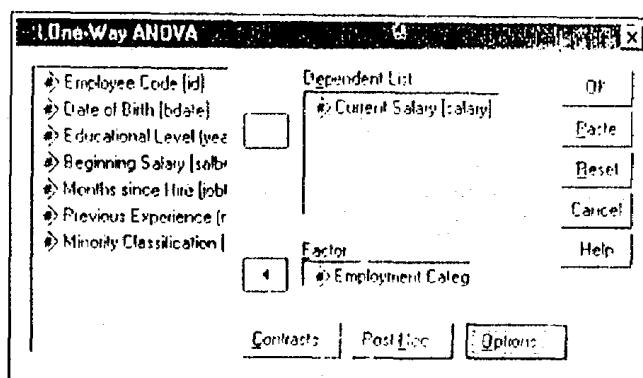
تحليل التباين في اتجاه واحد باستخدام برنامج SPSS:

تحليل التباين في اتجاه واحد يعني دراسة أثر عامل واحد (متغير مستقل واحد) على المتوسطات (متغير تابع) .

مثال : في بيانات employee data.sav نريد أن نعرف هل هناك فرق بين متوسط الرواتب salary بين مجموعات الوظائف jobcat ، ولتفصيل ذلك تابع الخطوات التالية :

١- افتح الملف Employee Data.sav .

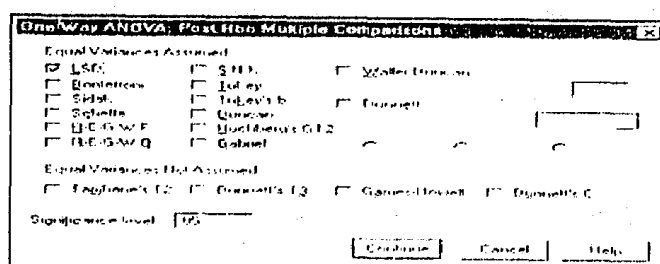
٢- من قائمة Analyze اختر أمر Compare Means ومنها اختر أمر One-way ANOVA يظهر الصندوق الحوارى التالى :



٢- أمام خانة Dependent list اختر المتغير التابع salary .

٣- أمام خانة Factor اختر المتغير الذى يقسم المتغير التابع إلى مجموعات jobcat .

٤- أنقر أمر Post Hoc لتحديد الاختبارات الإحصائية التى تظهر الأزواج من المتوسطات التى بينها اختلاف ، وأشهر هذه الاختبارات LSD .



٦- اختر أمر Continue لرجع إلى الصندوق الحوارى الأصلى ، ثم اختر أمر OK تظهر النتائج التالية :

Oneway

ANOVA

Current Salary					
	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	89438483926	2	44719241963.0	434.481	.000
Within Groups	48478011510	471	102925714.459		
Total	137916495436	473			

Post Hoc Tests

Multiple Comparisons

Dependent Variable: Current Salary
LSD

(I) Employment Cate	(J) Employment Cate	Mean Difference (I-J)	Std. Error	Sig.	5% Confidence Interval	
					Lower Bound	Upper Bound
Clerical	Custodial	3,100.35	,023.76	.126	\$7,077.06	\$876.37
	Manager	6,139.26*	,228.35	.000	38,552.99	33,725.53
Custodial	Clerical	3,100.35	,023.76	.126	-\$876.37	\$7,077.06
	Manager	3,038.91*	,244.41	.000	37,449.20	28,628.62
Manager	Clerical	6,139.26*	,228.35	.000	33,725.53	38,552.99
	Custodial	3,038.91*	,244.41	.000	28,628.62	37,449.20

The mean difference is significant at the .05 level.

تحليل النتائج

- ١- الفرض العدمي : المتوسطات متساوية بين مختلف الوظائف
- ٢- الفرض البديل : يوجد زوج واحد على الأقل بين أنواع الوظائف مختلف في المتوسط .
- ٣- القرار المتخذ من تحليل التباين اتضح أننا نرفض الفرض العدمي القائل بأن المتوسطات بين الرواتب salary حسب طبيعة الوظيفة jobcat متساوية ، ونقبل الفرض البديل القائل بأن هناك فرق بين المتوسطات ، حيث أن قيمة $F = 484.481$ ، ومستوى المعنوية أقل من

٠,٠٠٠٥

- ٤- من تحليل المقارنات المتعددة Multiple Comparisons يوجد فرق معنوي بين متوسط راتب الـ Manager و كل من الـ Clerical والـ Custodial ولا يوجد فرق معنوي بين الـ Clerical والـ Custodial في متوسط الراتب

تحليل التباين في اتجاهين أو أكثر :

في حالة تحليل التباين في اتجاهين أو أكثر سنقارن الفرق بين المجموعات لمغير تابع (كمي) بعد تقسيم الحالات طبقاً لخصائص متغيرين مستقلين أو أكثر في نفس الوقت (متغيرات وصفية) .
مثال : أجريت تجربة لاختبار تأثير اختلاف التربة (٣ قطع من الأرض) ، واختلاف نوع القمح (٤ سلالات) على محصول الحبوب ، وكان ناتج المحصول كالتالي (بالطن)

السلالة				قطعة الأرض
(٤)	(٣)	(٢)	(١)	
٢٣	١٩	٢٤	١١	(١)
١٥	١٨	٢٥	١٢	(٢)
٢٧	١٧	٢١	١٠	(٣)

يتم تسجيل البيانات على ملف ، وليكن Anova2 كالتالي :

land	seed	crop
١	١	١١
١	٢	٢٤
١	٣	١٩
١	٤	٢٣
٢	١	١٢
٢	٢	٢٥
٢	٣	١٨
٢	٤	١٥
٣	١	١٠
٣	٢	٢١
٣	٣	١٧
٣	٤	٢٧

لاختبار أثر عاملى (التربة ، وسلالة القمح)
على الإنتاج نتبع الآتى :

١ - افتح الملف Anova2.sav .

٢ - من قائمة Analyze اختر أمر
General Linear Model ، ومنها
اختر Univariate يظهر الصندوق
الحوارى Univariate :

٣- أمام خانة Factors Fixed اختر المتغيرات المستقلة land , seed

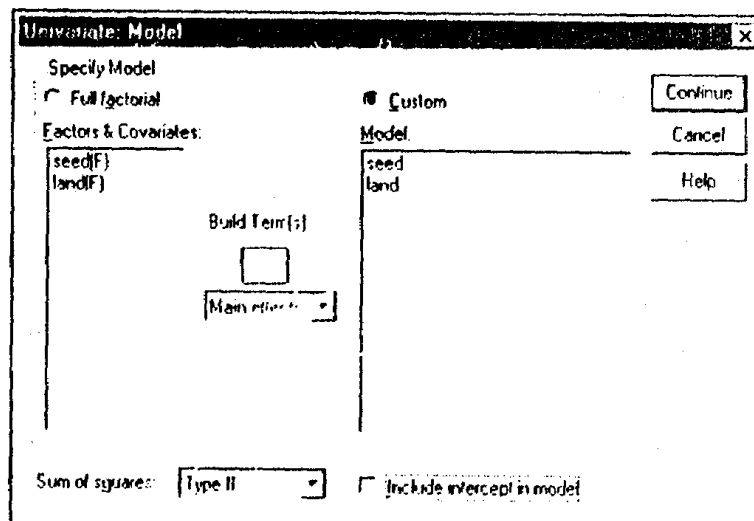
٤- أمام Dependent Variable اختر المتغير الكمي التابع Crop

٥- اختر Model فيظهر الصندوق الحوارى Model ، أمام خانة Build Term(s) اختر

Main Effects ، وأمام خانة Model اختر المتغيرين land , seed واحداً تلو الآخر .

٦- أمام خانة Sum of squares اختر Type II ، ولا تنشط خانة Include intercept in

model .



٢- اختر أمر Continue فترجع إلى الصندوق الحوارى الأصلى ، ثم اختر أمر OK فتظهر النتائج

التالية :

Univariate Analysis of Variance

Between-Subjects Factors

		N
SEED	1.00	3
	2.00	3
	3.00	3
	4.00	3
LAND	1.00	4
	2.00	4
	3.00	4

Tests of Between-Subjects Effects

Dependent Variable: CROP

Source	Type II Sum of Squares	df	Mean Square	F	Sig.
Model	4383.167 ^a	6	730.528	54.225	.000
SEED	269.667	3	89.889	6.672	.024
LAND	6.500	2	3.250	.241	.793
Error	80.833	6	13.472		
Total	4464.000	12			

a. R Squared = .982 (Adjusted R Squared = .964)

تصاغ الفروض العدمية والبديلة كما يلي

الفروض العدمية

١- لا يوجد أثر لسلالة البذور على المحصول

٢- لا يوجد أثر لنوع قطعة الأرض

الفروض البديلة

١- يوجد أثر لسلالة البذور على المحصول

٢- يوجد أثر لنوع قطعة الأرض

القرار المتخذ

١- قيمة F بالنسبة لسلالة البذور 6.672 ومستوى المعنوية 0.024 ، لذلك فإننا نرفض الفرض

العدمي ، ونقبل الفرض البديل ، أى يوجد تأثير للبذور على كمية المحصول .

٢- قيمة F بالنسبة لقطع الأرض 0.241 ، ومستوى المعنوية 0.793 ، لذلك فإننا نقبل الفرض العدمي

، أى لا تأثير لقطع الأرض على كمية المحصول .

٣-٢-٣ اختبار χ^2 - Chi-Square Test

يعتبر اختبار χ^2 من الاختبارات التي قد تستخدم كاختبار معلمي واختبار لامعلمي. في حالة اختبار الفروض حول تباين المجتمع فإننا نطبق اختبار χ^2 ، وفي هذه الحالة يكون الاختبار معلمي. ويوجد اختبارين معلمين يستخدم فيهما الاحصائي χ^2 ، أولهما اختبار جودة التوفيق لمتغير واحد غير مقيد وثانيهما اختبار الاستقلال بين المتغيرات الغير مقيدة ، ويبني الاختبار على مقارنة التكرارات المشاهدة والتكرارات المتوقعة لمعرفة هل هناك فرقاً معنوياً بينهما أم لا ؟ والشروط اللازمة لتطبيق اختبار χ^2 :-

١- عشوائية العينة Random sampling : العينة يجب أن تكون مختارة عشوائياً من مجتمع يارجاع ، أى أن مفردات العينة الواحدة مستقلة عن بعضها البعض .

٢- استقلال المشاهدات independence of observations في هذا نشترط أن كل مشاهدة مستخدمة في الاختبار يجب أن تكون مأخوذة من مصدر مستقل عن المشاهدة الأخرى.

٣- حجم المشاهدات المتوقعة Size of Expected Frequencies : حجم العينة المستخدم في التحليل يفضل أن يكون أكبر من ٣٠ ، عندما يكون حجم العينة صغيراً ، وعدد الخلايا الداخلة في الاختبار أقل من عشرة خلايا فإن أقل تكرار متوقع مطلوب للاختبار هو خمسة تكرارات للخلية الواحدة.

٤- حجم التكرارات المشاهدة Size of observed frequencies : التكرارات المشاهدة ممكن أن تكون صفر أو أقل من خمسة تكرارات في الخلية فلا شروط عليها.

ملاحظة : في حالة عدم تحقيق بعض هذه الشروط يمكن إجراء عمليات ضم للتكرارات المشاهدة والمتوقعة بقدر المستطاع.

١- اختبار كاي^٢ لجودة التوفيق Chi-square Goodness-of-fit Test في اختبار جودة التوفيق يكون لدينا تكرارا مشاهد تم جمعه مسبقاً ولدينا نظرية إحصائية (أو فرض معين) يستخدم في حساب التكرار المتوقع المناظر. ويكون الفرض العدمي والبديل بصفة عامة يكون كالآتي :-

الفرض العدمي : لا يوجد فرق بين التكرار المشاهد والمتوقع

الفرض البديل : يوجد فرق بين التكرار المشاهد والمتوقع.

٢- اختبار مربع كاي^٢ للاستقلال Chi-square Test for independence

نحتاج في حالات كثيرة إلى التعرف عما إذا كانت هناك علاقة بين صفتين من صفات مجتمع ما . مثلاً قد نحتاج لمعرفة هل توجد علاقة بين مستوى الدخل والمستوى التعليمي ؟ أو هل توجد علاقة بين لون العينين ولون الشعر في مجتمع ما ؟ للإجابة على مثل هذه الأسئلة يجب أن نختار عينة عشوائية من المجتمع ثم تصنف مشاهدات العينة حسب مستويات كل صفة من الصفتين ووضعهما في جدول يسمى جدول التوافق Contingency Table كما يلي :-

	B_1	B_2		B_j		B_s	
A_1	O_{11}	O_{12}		O_{1j}		O_{1s}	u_1
A_2	O_{21}	O_{22}		O_{2j}		O_{2s}	u_2
A_i	O_{i1}	O_{i2}		O_{ij}		O_{is}	u_i
A_r	O_{r1}	O_{r2}		O_{rj}		O_{rs}	u_r
المجموع	v_1	v_2		v_j		v_s	n

حيث مستويات الصفة A هي $A_1, A_2, \dots, A_r$

مستويات الصفة B هي $B_1, B_2, \dots, B_s$

وترمز O_{ij} إلى عدد المشاهدات التي يتوفر فيها المستويان A_i من الصفة A ، B_j من الصفة B

ويرمز u_i إلى عدد المشاهدات التي يتوفر فيها المستوى A_i من الصفة A

ترمز v_j إلى عدد المشاهدات التي يتوفر فيها المستوى B_j من الصفة B

ويكون فرض العدم والبديل كمايلي :-

فرض العدم H_0 : الصفتان A ، B مستقلتان.

الفرض البديل H_1 : الصفتان A ، B غير مستقلتين .

إجراء اختبار كاي^٢ باستخدام البرنامج الإحصائي SPSS

العرض الإحصائي للمتغيرين

الجدول التكرارية المزدوجة Crosstabs

تمثل العلاقة بين متغيرين ، عن طريق تمثيل هذه العلاقة في صورة متغير صفى ، ومتغير عمودى (Crosstab) ، كما يعطى عدة مقاييس إحصائية لقياس العلاقة بين هذين المتغيرين ، طبيعة المتغيرات (اسمية ، ترتيبية ، رقمية) هى التى تحدد المقياس المناسب لقياس هذه العلاقة ، ولكى يتم تمثيل البيانات المستمرة Continuous ينبغي إعادة تكويدها وتحويلها إلى فئات .

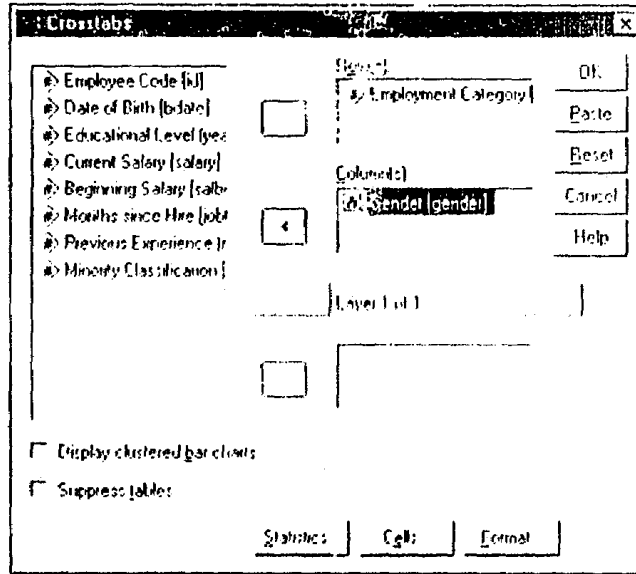
ولعمل جدول تكرارى مزدوج نتبع الآتى :

١. افتح الملف Employee Data.sav

٢. من قائمة Statistics اختر أمر Analyze ، ومنها اختر Descriptive Statistics

ومنها اختر Crosstabs.

يظهر الصندوق الحوارى التالى :



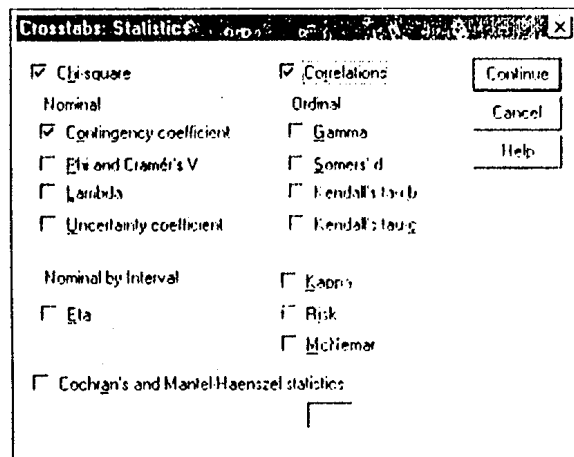
٢. اختر المتغيرات المراد تمثيلها كصفوف أمام خانة Row(s) ، وليكن Jobcat

٣. اختر المتغيرات المراد تمثيلها كأعمدة أمام خانة Column(s) ، وليكن gender

٤. اختر الـ Control Variables أمام خانة Layer ، ويلاحظ أنه سيتم إنشاء جدول

تكرارى مزدوج لكل قيم من قيم الـ Control Variable

٥. لحساب المقاييس الإحصائية المناسبة ، اختر أمر Statistics من الصندوق الحوارى Crosstabs ، واختر أمر Chi-Square , Correlation , Contingency Coefficient



٦. بالنسبة للجداول 2×2 تسمى جداول الاقتران ، الجداول أعلى من ذلك $C \times R$ بمعنى أى عدد من الصفوف وأى عدد من الأعمدة تسمى جداول التوافق .

٧. بالنسبة للجداول الاقتران 2×2 يستخدم لقياس معنوية العلاقة Chi-Square حيث يحسب Person chi-square , Likelihood chi-square , Fisher's exact test and Yate's corrected chi-square (continuity correction) .

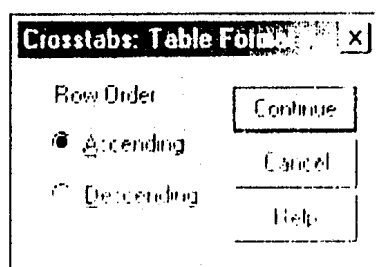
٨. بالنسبة للبيانات المقاسة بمقياس اسمى Nominal Data أى البيانات الوصفية مثل النوع أو الديانة توجد بالإضافة إلى Chi-Square مقاييس معامل التوافق Contingency coefficient ، الخ .

٩. بالنسبة للبيانات المقاسة بمقياس ترتيبى Ordinal بالإضافة إلى Chi-square يمكن اختيار Gamma , Somer's d

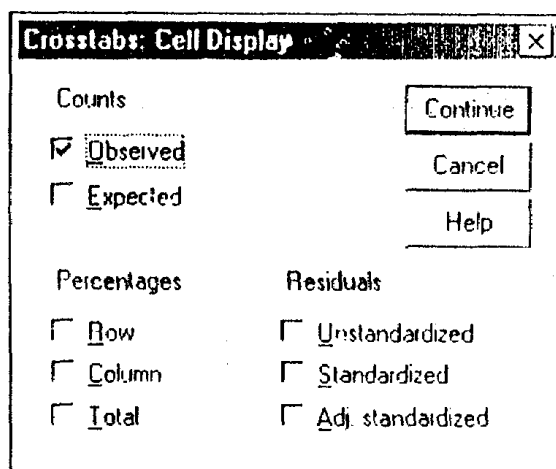
١٠. حينما يكون أحد المتغيرين وصفيا Category والآخر رقمى Numerical (Numeric by Interval) يمكن اختيار Eta , Kappa

١١. فى المتغيرات الرقمية يمكن اختيار معامل الارتباط الخطى Correlation ، وفى جداول التوافق عندما يكون عدد الصفوف يساوى عدد الأعمدة نختار Cappa

١٢. يمكن اختيار Format للتحكم فى شكل المخرجات كالتالى :



١٣. عادة تظهر في الخلايا القيم المشاهدة ، ولكن باختيار الأمر cell يظهر الصندوق الحوارى التالى :



ومنه يمكن إظهار في خلايا الجدول البيانات التالية :

(١) أعداد Counts :

– القيم المشاهدة Observed

– القيم المتوقعة Expected

(٢) البواقي Residual : وهو الفرق بين القيم المشاهدة والمتوقعة :

– غير قياسى Unstandardized.

– قياسى Standardized.

– قياسى معدل Adj-standardized

(٣) نسب مئوية Percentages

– Row : نسبة لإجمالى الصف .

– Column : نسبة لإجمالى العمود .

– Total : نسبة للإجمالى العام .

بعد اختيار أمر Continue ، فترجع إلى الصندوق الحوارى الأسمى ، فنختار أمر OK فتظهر لنا النتائج التالية :

Crosstabs

Warnings

CORR statistics are available for
numeric data only.

Case Processing Summary

	Cases					
	Valid		Missing		Total	
	N	Percent	N	Percent	N	Percent
Gender * Employment Category	474	100.0%	0	.0%	474	100.0%

Gender * Employment Category Crosstabulation

			Employment Category			Total
			Clerical	Custodial	Manager	
Gender	Female	Count	206		10	216
		% within Employment Category	56.7%		11.9%	45.6%
	Male	Count	157	27	74	258
		% within Employment Category	43.3%	100.0%	88.1%	54.4%
Total		Count	363	27	84	474
		% within Employment Category	100.0%	100.0%	100.0%	100.0%

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)
Pearson Chi-Square	79.277 ^a	2	.000
Likelihood Ratio	95.463	2	.000
N of Valid Cases	474		

a. 0 cells (.0%) have expected count less than 5. The minimum expected count is 12.30.

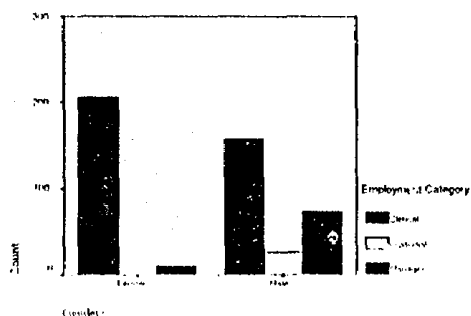
Symmetric Measures^a

	Value	Approx. Sig.
Nominal by Nominal Contingency Coefficient	.379	.000
N of Valid Cases	474	

a. Not assuming the null hypothesis.

b. Using the asymptotic standard error assuming the null hypothesis.

c. Correlation statistics are available for numeric data only.



تحليل البيانات من جدول تكرارى مزدوج :

قد تعطى البيانات فى جدول تكرارى مزدوج ، ونريد حساب المقاييس الإحصائية المناسبة :
مثال : لدينا التوزيع المشترك لخاصيتى الوفاة والتدخين لعينة من كبار السن ممثلة فى الجدول التالى :

الوفاة	التدخين		المجموع
	يدخن	لا يدخن	
١ توفى	٥٤	١١٧	١٧١
٢ حى	٣٤٨	٩٥٠	١٢٩٨
المجموع	٤٠٢	١٠٦٧	١٤٦٩

والمطلوب اختبار الفرض القائل بأن هناك علاقة بين التدخين والوفاة ، تمثل البيانات بالطريقة

التالية :

x1	x2	fre
1	1	54
1	2	117
2	1	348
2	2	950

	x1	x2	x
1	توفى	يدخن	54.00
2	توفى	لا يدخن	117.00
3	حى	يدخن	348.00
4	حى	لا يدخن	950.00

وبعد تخصيص Value Label للبيانات وتسجيلها تظهر فى نافذة Data Editor بالشكل

- (١) نقوم بترجيح البيانات بالمتغير 'fre' عن طريق اختيار weight cases من قائمة Data .
- (٢) نقوم بإجراء Crosstab ونختار X1 كمتغير صفى ، X2 كمتغير عمودى ونختار المقاييس الإحصائية المناسبة فتظهر لنا النتائج التالية :

Crosstabs

Case Processing Summary

	Cases					
	Valid		Missing		Total	
	N	Percent	N	Percent	N	Percent
X1 * X2	1359	100.0%	0	.0%	1359	100.0%

X1 * X2 Crosstabulation

			X2		Total
			يدخن	لا يدخن	
X1	توفى	Count	54	7	61
		% within X2	13.4%	.7%	4.5%
	حى	Count	348	950	1298
		% within X2	86.6%	99.3%	95.5%
Total	Count	402	957	1359	
	% within X2	100.0%	100.0%	100.0%	

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)	Exact Sig. (2-sided)	Exact Sig. (1-sided)
Pearson Chi-Square	106.526 ^b	1	.000	.000	.000
Continuity Correction ^a	103.584	1	.000		
Likelihood Ratio	97.860	1	.000		
Fisher's Exact Test					
Linear-by-Linear Association	106.447	1	.000		
N of Valid Cases	1359				

a. Computed only for a 2x2 table

b. 0 cells (.0%) have expected count less than 5. The minimum expected count is 18.04.

Symmetric Measures

	Value	Approx. Sig.
Nominal by Nominal Contingency Coefficient	.270	.000
N of Valid Cases	1359	

a. Not assuming the null hypothesis.

b. Using the asymptotic standard error assuming the null hypothesis.

الباب الرابع

الارتباط والانحدار

٤-١ الارتباط

لقياس العلاقة بين متغيرين ، ويستخدم معامل الارتباط الخطي ليرسون لقياس العلاقة الخطية بين متغيرين (كميين) وتتراوح قيمته بين -١ ، ١ ، وكلما اقترب من الواحد كلما دل ذلك على قوة العلاقة ، وكلما اقترب من الصفر دل ذلك على ضعف العلاقة ، كما أن إشارة المعامل (موجبة أو سالبة) تدل على اتجاه العلاقة ، فإذا كانت الإشارة موجبة كانت العلاقة طردية ، وإذا كانت الإشارة سالبة كانت العلاقة عكسية .

هناك عدة معاملات للارتباط ، ففي حالة العلاقة بين متغيرين ترتبيين يستخدم معامل كندال تاو B ، سبيرمان للرتب .

وللحصول على معامل الارتباط باستخدام البرنامج الإحصائي SPSS نتبع الآتي :

إذا أتاحت لك البيانات التالية ، أوجد معامل الارتباط الخطي بين المتغيرات

الأجور الزراعية X4	العمالة الزراعية بالألف X3	المساحة الزراعية X2	الاستثمارات الزراعية الثابتة X1	النتائج المحلى الزراعي الثابت y	السنوات Year
٢٦٧,٤٠	٤٢١٨,٠٠	٥,٧٩	٨٣,٤٠	٩٩٦,٩٠	١٩٧٥
٢٦٢,١٠	٤٠٦٨,٠٠	٥,٨٦	٩٨,٠٠	١٠٦٣,٧٠	١٩٧٦
٢٨٧,٧٠	٤١٠٤,٠٠	٥,٨٥	١١٤,٥٠	١٦١٧,٨٠	١٩٧٧
٣١٦,٢٠	٤١٣٥,٠٠	٥,٨٦	١٤٤,١٠	١٥١٥,٧٠	١٩٧٨
٣٤٧,٨٠	٤١٦٥,٠٠	٥,٨٦	١٨٧,٢٠	١٦٩٠,٨٠	١٩٧٩
٤٤٠,٣٠	٤٢٠٠,٠٠	٥,٨٧	٢٠٤,٤٠	١٧٠٧,٢٠	١٩٨٠
٩٧١,١٠	٤٢٢٥,٠٠	٥,٨٨	٤٢٣,٦٠	١٨٨٢,١٠	١٩٨١
١٠٥٥,٦	٤٢٦٤,٠٠	٥,٩٢	٤١٦,٣٠	٣٤٥٥,٦٠	١٩٨٢
١١٤١,١	٤٣٠٤,٠٠	٥,٨٤	٤٠٩,٨٠	٤٢٠٤,٩٠	١٩٨٣
١٢٢٨,٩	٤٣٤٤,٠٠	٥,٨٣	٤٥٢,٨٠	٤٠٦١,١٠	١٩٨٤
١٢٧١,٩	٤٣٨٤,٠٠	٥,٩٧	٤٩٥,٨٠	٤٠١٨,٧٠	١٩٨٥
١٥٨٦,١	٤٣٣٠,٠٠	٥,٩٧	٧٢٣,٥٠	٤٥٤٠,٠٠	١٩٨٦

السنوات Year	الناتج المحلي الزراعي الثابت y	الاستثمارات الزراعية الثابتة x1	المساحة الزراعية X2	العمالة الزراعية بالألف X3	الأجور الزراعية X4
١٩٨٧	٥٦١٠,٠٠	١٥١٥,٢٠	٦,٠٠	٤٣٨١,٠٠	١٨٥٥,٩
١٩٨٨	٥٧٨١,٠٠	١٦٦٧,٩٠	٦,١٨	٤٤٣٢,٠٠	٢١٠٥,٤
١٩٨٩	٥٩٣٥,٠٠	١٨٢١,٣٠	٦,٢٧	٤٤٧٨,٠٠	٢٣٨٢,٤
١٩٩٠	٦١٣٠,٠٠	١٩٠٥,٢٠	٦,٩٢	٤٥٣٣,٠٠	٢٦٩٨,٠
١٩٩١	٧٩٣٩,٠٠	١٩٩٥,٥٠	٧,٠٢	٤٥٨٥,٠٠	٣٠١٢,٠
١٩٩٢	٨١٣٦,٠٠	٢٢٨٨,٥٠	٧,١٢	٤٦٢٤,٠٠	٣٤٢٣,٠
١٩٩٣	٨٤٤٨,٠٠	٢٧٠٦,١٠	٧,١٧	٤٦٨٢,٠٠	٣٨٩٥,٠
١٩٩٤	٨٦٩٣,٠٠	٣٣٨٧,٦٠	٧,١٨	٤٧٤٤,٠٠	٤٣٦٩,٠
١٩٩٥	٨٩٦٠,٠٠	٣٧٢٩,٧٠	٧,١٨	٤٨١٢,٠٠	٤٩٤٦,٠

١- سجل البيانات في ملف وليكن Corr.sav

٢- من قائمة Analyze اختر أمر Correlate ، ومنها اختر Bivariate يظهر الصندوق الحوارى التالى :

Bivariate Correlations

Variable:

#> y
#> x1
#> x2
#> x3
#> x4

Correlation Coefficients

☒ Pearson ☐ Kendall's tau-b ☐ Spearman

Test of Significance

☒ Two-tailed ☐ One-tailed

☒ Flag significant correlations

Options

OK
Paste
Reset
Cancel
Help

٣- اختر المتغيرات المراد إيجاد معاملات الارتباط بينها وليكن y , x1 , x2 , x3 , x4 .

٤- إذا كانت المتغيرات كمية اختر **Person** من خانة معاملات الارتباط ، أما إذا كانت البيانات ترتيبية

اختر **Kendall's tau-b , spearman**

٥- اختر أمر **OK** تظهر النتائج التالية :

Correlations

Correlations

		Y	X1	X2	X3	X4
Y	Pearson Correlation	1.000	.943**	.898**	.971**	.968**
	Sig. (2-tailed)		.000	.000	.000	.000
	N	21	21	21	21	21
X1	Pearson Correlation	.943**	1.000	.925**	.961**	.986**
	Sig. (2-tailed)	.000		.000	.000	.000
	N	21	21	21	21	21
X2	Pearson Correlation	.898**	.925**	1.000	.917**	.936**
	Sig. (2-tailed)	.000	.000		.000	.000
	N	21	21	21	21	21
X3	Pearson Correlation	.971**	.961**	.917**	1.000	.984**
	Sig. (2-tailed)	.000	.000	.000		.000
	N	21	21	21	21	21
X4	Pearson Correlation	.968**	.986**	.936**	.984**	1.000
	Sig. (2-tailed)	.000	.000	.000	.000	
	N	21	21	21	21	21

** . Correlation is significant at the 0.01 level (2-tailed).

نلاحظ ما يلي :

- مصفوفة الارتباط مصفوفة متماثلة ، حيث أن القيم على القطر = ١ ، حيث أنها تمثل (ارتباط المتغير بنفسه) ، والقيم أعلى القطر = القيم أسفل القطر ، حيث أن ارتباط المتغير x_1 بالمتغير x_2 هو نفسه ارتباط المتغير x_2 بالمتغير x_1 .
- يوجد بكل خلية من خلايا الجدول ثلاثة قيم ، القيمة الأولى تمثل قيمة معامل الارتباط (**Person Correlation**) ، القيمة الثانية تمثل معنوية معامل الارتباط (**Sig. (2-tailed)**) ، إذا كانت هذه القيمة أقل من ٠,٠٥ كان الارتباط معنوياً ، أما إذا كانت أكبر من ٠,٠٥ كان الارتباط غير معنوي ، القيمة الثالثة تمثل عدد القيم (N) .
- يظهر البرنامج الرمز (*) بجوار معامل الارتباط إذا كان معنوياً عند ٠,٠٥ ، إما إذا كان معنوياً عند ٠,٠١ فإنه يظهر الرمز (**) .

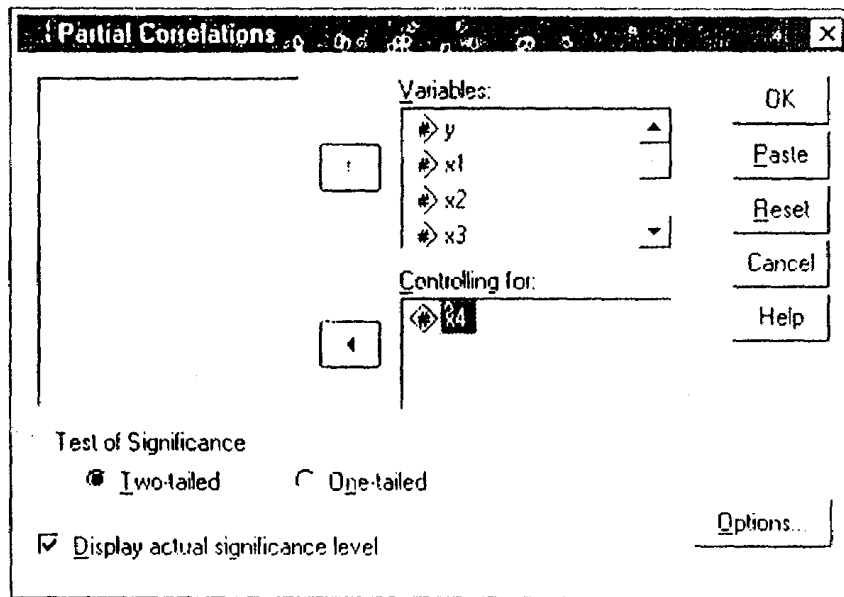
الارتباط الجزئي Partial Correlation :

يقيس معامل الارتباط الجزئي العلاقة بين متغيرين مع وجود متغير ثالث يؤثر عليهما ، ولإيجاد معامل الارتباط الجزئي (أى بعد عزل تأثير المتغير المؤثر على كل من المتغيرين) نتبع الآتى :

١- افتح الملف السابق corr.sav .

٢- من قائمة Analyze اختر أمر Correlate ، ومنها اختر امر Partial يظهر الصندوق الحوارى التالى :

٢- أمام خانة variables اختر المتغيرات المراد إيجاد معامل الارتباط الجزئي لها ، ولتكن y, x1, x2, x3



٣- أمام خانة controlling for اختر المتغيرات التى تؤثر على المتغيرات المرتبطة ، وليكن X4

٤- اختر أمر OK تظهر النتائج التالية :

Partial Corr

- - - PARTIAL CORRELATION COEFFICIENTS - - -

Controlling for.. X4

	Y	X1	X2	X3
Y	1.0000 (0) P= .	-.3054 (18) P= .190	-.0959 (18) P= .688	.4090 (18) P= .073
X1	-.3054 (18) P= .190	1.0000 (0) P= .	.0312 (18) P= .896	-.3451 (18) P= .136
X2	-.0959 (18)	.0312 (18)	1.0000 (0)	-.0787 (18)

	P= .688	P= .896	P= .	P= .741
x3	.4090	-.3451	-.0787	1.0000
	(18)	(18)	(18)	(0)
	P= .073	P= .136	P= .741	P=

(Coefficient / (D.F.) / 2-tailed Significance)

" . " is printed if a coefficient cannot be computed

نلاحظ من النتائج ما يلي :

- النتائج السابقة تمثل مصفوفة الارتباط بين المتغيرات x_3 , x_2 , x_1 , y بعد عزل تأثير المتغير x_4 ، كما يشير الرقم بين القوسين إلى درجات الحرية (عدد المشاهدات - ٣) ، وتشير الـ P إلى مستوى المعنوية .

- تتميز بنفس خصائص مصفوفة الارتباط السابق التعرف عليها من حيث أن معاملات الارتباط على القطر تساوى الواحد الصحيح ، والمعاملات، أعلى القطر تساوى المعاملات أسفل القطر .

٤-٢ الانحدار الخطى البسيط

يمثل تقدير العلاقة بين متغيرين أهمية كبيرة في مجال العلوم الاجتماعية وخاصة المجال الاقتصادي حيث تظهر الحاجة إلى قياس كمى بين بعض المتغيرات كالعلاقة بين الدخل الممكن التصرف فيه والاستهلاك مثلاً وذلك لتحديد شكل العلاقة ومدى قوتها وتقدير المعاملات الخاصة بها وقياس درجة معنويتها الإحصائية. لذلك فإن تحليل الانحدار يستخدم لاختبار العلاقة بين متغيرين أحدهم تابع أو مفسر وليكن (y) والآخر مستقل أو مفسر (x) ، كما يمكن استخدام في التنبؤ بقيم المتغير التابع عند حدوث تغير في المتغير المستقل،

ولتحديد شكل العلاقة بين المتغيرين المعنيين غالباً ما نبدأ بالتعرف على شكل الانتشار للقيم الخاصة بهم وما إذا كان يتقارب مع معادلة الخط المستقيم أو بتعبير رياضي على الصورة .

$$y_i = b_0 + b_1 x_i$$

وحيث أنه في العلوم الاجتماعية من الصعب أن تقع قيم y بالدقة على الخط المستقيم لذلك فإن الشكل السابق يعدل ليحتوى على الأخطاء العشوائية والذي يجعل المعادلة السابقة على الصورة

$$y_i = b_0 + b_1 x_i + u_i$$

حيث يلاحظ أن الأخطاء العشوائية u_i تعبر عن الجزء غير المنتظم في العلاقة نتيجة السلوك البشرى ، كما يمكن أن يعبر عن الأخطاء الناجمة عن حذف بعض المتغيرات في المعادلة ، أو أخطاء القياس والتجميع في العلاقة ، كما قد يرجع إلى تبسيط العلاقة واختزالها على صورة خط مستقيم.

ويخضع هذا المتغير العشوائى لمجموعة من الافتراضات

١. المتغير العشوائى له توزيع احتمالى اى أنه عند كل نقطة من نقط المتغير x يوجد توزيع اجمالى للقيم المتناظرة للمتغير y وبالتالي للمتغير u وهذه التوزيعات تتركز حول الخط المراد قياسه .
٢. له توزيع طبيعي والوسط الحسابى لقيم الأخطاء يساوى صفر ومعنى هذا أن هناك قيما مختلفة للأخطاء بعضها اكبر من الصفر وبعضها اقل ، لكن متوسطها يساوى صفر
٣. التباين بالنسبة للأخطاء العشوائية حول متوسطها والذي يساوى صفر ثابت عند كل القيم الممكنة للمتغير x بمعنى آخر أن انتشار الأخطاء حول متوسطها ثابت لا يختلف باختلاف القيم المتناظرة للمتغير x وليكن σ^2
٤. قيم الأخطاء العشوائية غير مرتبطة فيما بينها، أى أن القيمة التى يأخذها المتغير u عند مستوى معين للمتغير x لا تتأثر بالقيمة التى يأخذها عند مستوى آخر للمتغير x
٥. المتغير المستقل x غير مرتبط بالأخطاء العشوائية، أى أنه ليس هناك علاقة بين تغير x وتغير الأخطاء العشوائية.

ولتوفيق أفضل خط مستقيم لأزواج القراءات، الخاصة بالمتغيرين y و x يمكن استخدام طريقة المربعات الصغرى حيث أنها تتضمن تقليل مجموع مربعات الانحرافات عن خط الانحدار أو بصيغة رياضية

$$\text{Min } \sum (y_i - \hat{y}_i)^2$$

حيث y_i هى لقيم الحقيقية ، أما لقيم $\hat{y}_i$ هى القيم على خط الانحدار أو بمعنى آخر فان

$$y_i - \hat{y}_i = e_i$$

حيث e_i تقدير لقيم الأخطاء العشوائية

ومجموع مربعات الأخطاء اقل ما يمكن تعنى

$$\text{Min } \sum e_i^2 = \text{Min } \sum (y_i - \hat{y}_i)^2$$

أى أن

$$\sum e_i^2 = \sum (y_i - \hat{y}_i)^2 = \sum (y_i - b_0 - b_1 x)^2$$

وعلى ذلك فإننا نبحث عن قيم b_0 , b_1 التى تجعل هذه البواقي اقل ما يمكن.

ويمكن استخراج هذه القيم بالتفاضل جزئيا بالنسبة لـ b_0 , b_1

تنتج المعادلات الأساسية التالية

$$\sum y_i = nb_0 + b_1 \sum x$$

$$\sum x_i y_i = b_0 \sum \bar{x} + b_1 \sum x^2$$

حيث n عدد القراءات

من هذه المعادلات يمكن استنتاج كل من قيمة b_0 , b_1 -حيث

$$b_0 = \bar{y} - b_1 \bar{x}$$

حيث $\bar{x} = (\sum x)/n$, $\bar{y} = (\sum y)/n$ وتمثل الأوساط الحسابية لكل من y,x

$$\hat{b}_1 = \frac{\sum xy - n \bar{x} \bar{y}}{\sum x^2 - n \bar{x}^2} = \frac{\sum (x - \bar{x})(y - \bar{y})}{\sum (x - \bar{x})^2}$$

حيث أن القيم $(x - \bar{x})$, $(y - \bar{y})$ هي قيم المتغيرات مقاسة من أوساطها الحسابية.

بتقدير قيم b_0 , b_1 يكون قد تم تقدير معالم معادلة الانحدار ويتبقى ضرورة تقدير معنوية معادلة الانحدار ومعالمها.

وتتميز طريقة المربعات الصغرى بأنها تقديرات خطية غير متحيزة وكفاءة أى لها أصغر تباين ، كافية

➤ تقدير معنوية معادلة الانحدار.

لتقدير معنوية هذه المعادلة يتطلب تكوين جدول تحليل التباين والمتعارف عليه باسم ANOVA

لعلاقة الانحدار وبالتالي اختبار القوة التفسيرية لهذا الخط.

يلاحظ أن يمكن تقسيم الفرق بين القيمة الحقيقية y ووسطها الحسابي $\bar{y}$ إلى جزئين.

١- جزء يرجع إلى الأخطاء العشوائية وهو يمثل الفرق بين القيم الحقيقية والقيم المقدرة $(y - \hat{y})$

٢- جزء آخر يرجع إلى اختلاف القيم المقدرة عن وسطها $(\hat{y} - \bar{y})$ وعلى ذلك فإنه يمكن كتابة (*)

$$\sum (y - \bar{y})^2 = \sum (y - \hat{y})^2 + \sum (\hat{y} - \bar{y})^2$$

$$SST = \sum e^2 + SSR$$

$$SST = SSE + SSR$$

ويكون جدول تحليل التباين كما يلي.

(*) هناك إثبات رياضي يمكن الرجوع إليه لمن يرغب في أى من المراجع الإحصائية.

جدول تحليل التباين

فى حالة الانحدار الخطى البسيط

مصدر التغير	درجات الحرية	مجموع المربعات	متوسط مجموع المربعات	$F_{1, n-2}$ المحسوبة
SSR الانحدار البواقي (الانحراف SSE عن خط الانحدار)	1 n-2	$SSR = \sum (\hat{y} - \bar{y})^2$ $SSE = \sum (y_i - \hat{y})^2$	$\frac{SSR}{1} = M_{xR}$ $\frac{SSE}{n-2} = \sigma_E^2$	M_{xR} / σ_E^2
الكلى	n-1	$SS_T = \sum (y_i - \bar{y})^2$		

لاختبار القدرة التفسيرية للعلاقة الخطية تستخدم قيمة F المحسوبة بدرجات حرية البسط 1 والمقام (n-2) وتقارن بقيمتها الجدولية عند مستوى معنوية وليكن α ويمكن عدم رفض فرض العدم (قبول لفرض العدم) وهو أن لا يوجد مقدرة تفسيرية للخط المقدر وذلك إذ كانت قيمة F المحسوبة اقل من القيمة الجدولية.

كذلك فان يمكن حساب الانحراف المعيارى لهذا النموذج

$$\hat{\sigma} = \sqrt{\hat{\sigma}^2}$$

اختبار معنوية المعالم المقدرة

تم تقدير معالم النموذج ومن ثم من الضرورة اختبار المعنوية الإحصائية لهذه المعالم ، لذلك يجب تقدير تباين هذه المعالم.

كما يمكن استخدام خاصية التوزيع الطبيعى أو توزيع t حسب حجم العينة المستخدم فى التقدير ففى حالة العينات الصغيرة أى عدد القراءات اقل من ٣٠ دنفردة يستخدم توزيع t

➤ تباين المعالم المقدرة

لو فرضنا أن $S^2_{b_1}$ تباين المعلمة b_1 أى ميل خط الانحدار فإنها تساوى

$$S^2_{b_1} = \frac{\hat{\sigma}_E^2}{\sum (x - \bar{x})^2} = \frac{\hat{\sigma}_E^2}{\sum x^2 - n\bar{x}^2}$$

ذلك لأن

$$\hat{b}_1 = \frac{\sum (x - \bar{x})(y - \bar{y})}{\sum (x - \bar{x})^2}$$

ومنها نجد

$$\hat{b}_1 = \frac{\sum (x - \bar{x})y - \frac{\sum (x - \bar{x}) \sum y}{n}}{\sum (x - \bar{x})^2}$$

$$\hat{b}_1 = \frac{\sum (x - \bar{x})(b_0 + b_1 x + u)}{\sum (x - \bar{x})^2}$$

ويكون تبين $\hat{b}_1$ مساويا لـ

$$\hat{v}(b_1) = \frac{v(u) \sum x^2}{n \sum (x - \bar{x})^2}$$

وتستخدم القيمة $\hat{\sigma}_E^2$ السابقة بياها بالجدول كتقدير لتباين u

كذلك فان تقدير تباين b_0 يكون

$$\hat{v}(b_0) = v(u) \frac{\sum x^2}{n \sum (x - \bar{x})^2}$$

باستخراج قيم التباينات يمكن إجراء الاختبار الاحصائي اللازم عن قيم هذه الملمات فتكون :

نفرض أن $b_1 = 0$ أي أن الفرض العدمي

$$H_0 \Rightarrow b_1 = 0$$

$$H_1 \Rightarrow b_1 \neq 0$$

وأن الفرض البديل

وتكون

$$t_{\alpha, (n-2)} = \frac{\hat{b}_1 - b_1}{\sqrt{\hat{v}(b_1)}}$$

وبنفس الطريقة يمكن اختبار المعنوية الإحصائية للمعلمة b_0 حيث تكون

$$t_{\alpha, n-2} = \frac{\hat{b}_0 - b_0}{\sqrt{\hat{v}(b_0)}}$$

وذلك على اعتبار أن

$$H_0 \Rightarrow b_0 = 0, H_1 \Rightarrow b_0 \neq 0$$

ولاحظ أن الغرض من تحليل الانحدار هو قبول الفرض البديل أى أن كل من b_0 و b_1 لا يساوى الصفر ، وفى هذا النطاق نستخدم اختبار التيلين.

معامل التحديد R^2

معامل التحديد عبارة عن النسبة بين مجموع مربعات الانحدار ومجموع المربعات الكلى ويستخدم لمعرفة القوة التفسيرية للنموذج المقدر أى نسبة التغير الكلى للمتغير y بسبب علاقة الانحدار بين x و y ومن الجدول قد لاحظنا أن

$$\begin{aligned} SST &= SSR + SSE \\ 1 &= \frac{SSR}{SST} + \frac{SSE}{SST} \end{aligned}$$

$$\begin{aligned} &= \frac{\sum (y - \hat{y})^2}{\sum (y - \bar{y})^2} + \frac{\sum e_i^2}{\sum (y - \bar{y})^2} \\ 1 &= \frac{\sum e_i^2}{\sum (y - \bar{y})^2} = R^2 \end{aligned}$$

من هذا نجد أن قيمة R^2 تقع بين صفر ، 1 أى أن $0 \leq R^2 \leq 1$

$R^2 = 0$ = صفر عندما تكون الأخطاء كبيرة وتساوى مجموع المربعات الكلية أى عندما تقع نقط العينة

على خط مستقيم $y = \bar{y}$. أو على دائرة ولا توجد علاقة بين المتغيرين x , y

$R^2 = 1$ = عندما تقع كل نقط العينة على خط الانحدار المقدر .

الانحدار الخطى البسيط Linear Regression باستخدام البرنامج الإحصائى SPSS

يمكن توضيح ذلك من خلال المثال التالى :

مثال : فيما يلى بيانات الناتج المحلى الزراعى الثابت (y) ، الاستثمارات الزراعية الثابتة (x) ، والمطلوب إيجاد نموذج للعلاقة الخطية بين المتغيرين

السنوات	الناتج المحلي الزراعي الثابت	الاستثمارات الزراعية
١٩٧٥	٩٩٦,٩٠	٨٣,٤٠
١٩٧٦	١٠٦٣,٧٠	٩٨,٠٠
١٩٧٧	١٦١٧,٨٠	١١٤,٥٠
١٩٧٨	١٥١٥,٧٠	١٤٤,١٠
١٩٧٩	١٦٩٠,٨٠	١٨٧,٢٠
١٩٨٠	١٧٠٧,٢٠	٢٠٤,٤٠
١٩٨١	١٨٨٢,١٠	٤٢٣,٦٠
١٩٨٢	٣٤٥٥,٦٠	٤١٦,٣٠
١٩٨٣	٤٢٠٤,٩٠	٤٠٩,٨٠
١٩٨٤	٤٠٦١,١٠	٤٥٢,٨٠
١٩٨٥	٤٠١٨,٧٠	٤٩٥,٨٠
١٩٨٦	٤٥٤٠,٠٠	٧٢٣,٥٠
١٩٨٧	٥٦١٠,٠٠	١٥١٥,٢٠
١٩٨٨	٥٧٨١,٠٠	١٦٦٧,٩٠
١٩٨٩	٥٩٣٥,٠٠	١٨٢١,٣٠
١٩٩٠	٦١٣٠,٠٠	١٩٠٥,٢٠
١٩٩١	٧٩٣٩,٠٠	١٩٩٥,٥٠
١٩٩٢	٨١٣٦,٠٠	٢٢٨٨,٥٠
١٩٩٣	٨٤٤٨,٠٠	٢٧٠٦,١٠
١٩٩٤	٨٦٩٣,٠٠	٣٣٨٧,٦٠
١٩٩٥	٨٩٦٠,٠٠	٣٧٢٩,٧٠

ولإيجاد ذلك النموذج الخطي نتبع الآتي :

١- من قائمة Analyze اختر أمر Regression ، ومنها اختر أمر Linear

يظهر الصندوق الحوارى التالى :

Linear Regression

Dependent: #> y

Block 1 of 1 Next

Independent(s): #> x

Method: Enter

Selection Variable:

Case Labels:

WLS >> Statistics... Plots... Save... Options...

OK Paste Reset Cancel Help

٢- أمام خانة **Dependent** اختر المتغير التابع **Y**

٣- أمام خانة **Independent(s)** اختر المتغير المستقل **X1**

٤- أمام خانة **Method** اختر **Enter**.

٥- يمكن اختيار **Statistics** لإظهار بعض المقاييس الإحصائية الأخرى كالتالي :

Linear Regression: Statistics

Regression Coefficients: ☒ Model fit ☐ R squared change

☒ Estimates ☒ Descriptives

☐ Confidence intervals ☐ Part and partial correlations

☐ Covariance matrix ☐ Collinearity diagnostics

Residuals

☒ Durbin-Watson

☒ Casewise diagnostics

☐ Outliers outside standard deviation:

☒ All cases

Continue Cancel Help

٦- اختر أمر **OK** تظهر النتائج التالية :

Regression

Descriptive Statistics

	Mean	Std. Deviation	N
Y	4589.8333	2747.4347	21
X	1179.5429	1147.9951	21

Correlations

		Y	X
Pearson Correlation	Y	1.000	.943
	X	.943	1.000
Sig. (1-tailed)	Y		.000
	X	.000	
N	Y	21	21
	X	21	21

Variables Entered/Removed^b

Model	Variables Entered	Variables Removed	Method
1	X ^a		Enter

a. All requested variables entered.

b. Dependent Variable: Y

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Durbin-Watson
1	.943 ^a	.888	.883	941.3460	.552

a. Predictors: (Constant), X

b. Dependent Variable: Y

ANOVA^b

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	1.3E+08	1	1.3E+08	151.367	.000 ^a
	Residual	1.7E+07	19	886132.2		
	Total	1.5E+08	20			

a. Predictors: (Constant), X

b. Dependent Variable: Y

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	1928.961	298.232		6.467	.000
	X	2.256	.183	.943	12.303	.000

a. Dependent Variable: Y

Casewise Diagnostics^a

Case Number	Std. Residual	Y	Predicted Value	Residual
1	-1.190	996.90	2117.0992	-1120.20
2	-1.154	1063.70	2150.0346	-1086.33
3	-.605	1617.80	2187.2561	-569.4561
4	-.784	1515.70	2254.0293	-738.3293
5	-.702	1690.80	2351.2564	-660.4564
6	-.725	1707.20	2390.0571	-682.8571
7	-1.065	1882.10	2864.5394	-1002.44
8	.624	3455.60	2668.0717	587.5283
9	1.436	4204.90	2853.4087	1351.4913
10	1.180	4061.10	2950.4103	1110.6897
11	1.032	4018.70	3047.4118	971.2882
12	1.040	4540.00	3561.0689	978.9311
13	.279	5610.00	5347.0256	262.9744
14	.095	5781.00	5691.4939	89.5061
15	-.109	5935.00	6037.5414	-102.5414
16	-.103	6130.00	6226.8072	-96.8072
17	1.602	7939.00	6420.5105	1508.4895
18	1.110	8136.00	7091.4746	1044.5254
19	.440	8448.00	8033.5177	414.4823
20	-.933	8693.00	9570.8796	-877.8796
21	-1.469	8960.00	10342.6060	-1382.61

a. Dependent Variable: Y

Residuals Statistics^a

	Minimum	Maximum	Mean	Std. Deviation	N
Predicted Value	2117.0991	10342.61	4589.8333	2589.7050	21
Residual	-1382.61	1508.4895	6.496E-14	917.5105	21
Std. Predicted Value	-.955	2.221	.000	1.000	21
Std. Residual	-1.469	1.602	.000	.975	21

a. Dependent Variable: Y

نلاحظ من النتائج ما يلي :

يظهر البرنامج معامل الارتباط (R) وهو ٩٤. وهذا يعنى أن هناك ارتباط قوى وطردي بين المتغير التابع (Y) والمتغير المستقل (x1) ، وبحسب (R square) وهو مربع معامل الارتباط ويسمى معامل التحديد وهى نسبة التغيرات فى المتغير التابع التى يشرحها المتغير المستقل .

كما يظهر البرنامج جدول تحليل التباين والذي يختبر معنوية النموذج ككل ، ومنه نجد أن قيمة Signif F قيمة صغيرة جدا ، وهذا يعنى أننا نرفض الفرض العدمي القائل بأذ معاملات الانحدار = صفر ، أى أن نموذج الانحدار معنوى .

كما يظهر البرنامج معاملات الانحدار b , a (constant) واختبارات معنوية هذه المعاملات (T) و (T Sig)

كما يظهر بيانات عن القيم الفعلية والقيم المتوقعة من معادلة الانحدار Predicated ، كما يظهر بيانات الفرق بين القيم الفعلية والمتوقعة Residual

توفيق المنحنيات Curve Fitting :

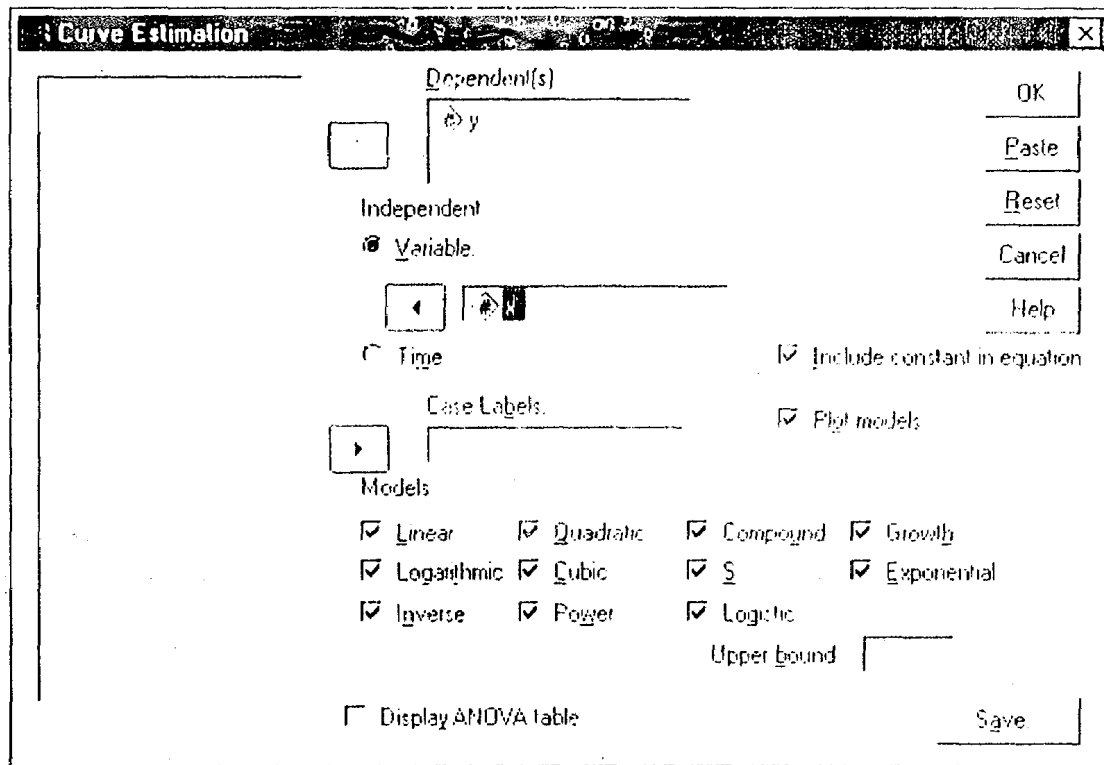
فى حالة العلاقة بين متغيرين قد لا نستطيع أن نحدد شكل العلاقة بين المتغيرين (Linear,Cubic,Power,....etc) ونريد معرفة أفضل نموذج لتمثيل هذه العلاقة ، لذلك يتيح البرنامج عدة أشكال للعلاقة بين المتغيرين ، ويعطى لنا الفرصة لاختيار أفضل النماذج ، وهذه النماذج هى :

شكل العلاقة	النموذج
$Y = b_0 + b_1 * X$	Linear
$Y = b_0 + b_1 * \ln(X)$	Logarithmic
$Y = b_0 + b_1 / X$	Inverse
$Y = b_0 + b_1 * X + b_2 * X^{**2}$	Quadratic
$Y = b_0 + b_1 * X + b_2 * X^{**2} + b_3 * X^{**3}$	Cubic
$Y = b_0 * X^{**}(b1)$	Power
$Y = b_0 * b_1^{**X}$	Compound
$Y = e^{**}(b_0 + b_1 / X)$	S

النموذج	شكل العلاقة
Logistic	$Y = 1/(1/u + b_0 * b_1 **X)$
Growth	$Y = e**(b_0 + b_1 * X)$
Exponential	$Y = b_0 + c**(b_1 * X)$

يمكن أن تكون العلاقة بين متغير تابع (y) ومتغير مستقل (x) ، ويمكن أيضاً أن تكون بيانات المتغير التابع (y) عبارة عن سلسلة زمنية ونريد إيجاد علاقة بين المتغير (y) والزمن (t) في المثال السابق نريد توفيق أفضل نموذج للعلاقة بين الناتج المثلّي الزراعي الثابت (y) والاستثمارات الزراعية الثابتة (x1) ، ولإيجاد ذلك نتبع الآتي :

١- من قائمة Statistics اختر أمر Regression ومنها اختر أمر Curve Estimation يظهر



الصندوق الحوارى التالى :

٢- أمام خانة Dependent(s) اختر المتغير التابع (y)

٣- أمام خانة Independent يوجد اختياران :

Variable _ فى حالة ما إذا كانت العلاقة بين المتغير التابع (y) ومتغير آخر (x) .

Time _ في حالة ما إذا كان المتغير التابع (y) عبارة عن سلسلة زمنية ، ونريد إيجاد علاقة بينه وبين الزمن T .

٤- إذا أردنا ظهور الثابت (a) في المعادلة فإننا ننشط الاختيار Include constant in equation
أما إذا أهملنا هذه الخانة فإن خط الانحدار يمر بنقطة الأصل (origin) .

٥- اختر النماذج المطلوبة من خانة Models ، وأشكال هذه النماذج موضحة بالجدول السابق .

٦- إذا أردنا عرض جدول تحليل التباين نشط الاختيار Display ANOVA table

٧- اختر أمر OK تظهر النتائج التالية :

Curve Fit

MODEL: MOD_1.

—

Independent: X

Dependent	Mth	Rsq	d.f.	F	Sigf	Upper bound	b0	b1	b2
b3									
Y	LIN	.888	19	151.37	.000		1928.96	2.2559	
Y	LOG	.927	19	240.93	.000		-9080.8	2116.11	
Y	INV	.657	19	36.34	.000		6530.27	-625782	
Y	QUA	.935	18	128.99	.000		1291.28	3.9497	-.0005
Y	CUB	.940	17	88.16	.000		940.267	5.6007	-.0017
2.1E-07									
Y	COM	.721	19	49.18	.000		1940.86	1.0005	
Y	POW	.934	19	268.05	.000		94.5684	.5667	
Y	S	.835	19	96.48	.000		8.7946	-188.35	
Y	GRO	.721	19	49.18	.000		7.5709	.0005	
Y	EXP	.721	19	49.18	.000		1940.86	.0005	
Y	LGS	.721	19	49.18	.000		.0005	.9995	

واضح من البيانات أن أفضل نموذج لتمثيل البيانات هو النموذج QUB حيث أن له أعلى Rsq أى مربع معامل الارتباط 0.940

بعد ذلك يتم اختيار ذلك النموذج ، وتحديد معاملاته على النحو التالي :

Curve Fit

MODEL: MOD_2.

—

Dependent variable.. Y

Method.. CUBIC

Listwise Deletion of Missing Data

Multiple R .96933
R Square .93961
Adjusted R Square .92895
Standard Error 732.34120

Analysis of Variance:

	DF	Sum of Squares	Mean Square
Regression	3	141850449.4	47283483.1
Residuals	17	9117501.9	536323.6

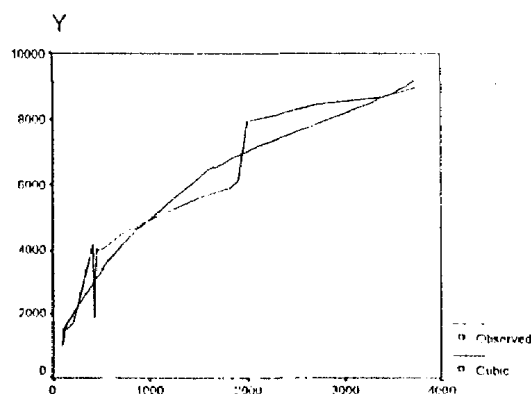
F = 88.16222 Signif F = .0000

----- Variables in the Equation -----

Variable	B	SE B	Beta	T	Sig T
X	5.600722	1.498362	2.340220	3.738	.0016
X**2	-.001709	.001027	-2.446629	-1.655	.1162
X**3	2.12499876E-07	1.8223E-07	1.068165	1.166	.2597
(Constant)	940.266994	419.118963		2.243	.0385

The following new variables are being created:

Name	Label
FIT_1	Fit for Y with X from CURVEFIT, MOD_2 CUBIC
ERR_1	Error for Y with X from CURVEFIT, MOD_2 CUBIC



يتم تمثيل القيم المتوقعة Predicated والبواقي Residual في ملف البيانات باسم FITT_1,

ERR_1 على النحو التالي

Y	X1	FITT_1	ERR_1
996,90	83,40	1395,67	398,77-
1063,70	98,00	1473,01	409,31-
1617,80	114,50	1559,58	58,22
1515,70	144,10	1712,67	196,97-

Y	X1	FITT 1	ERR 1
١٦٩٠,٨٠	١٨٧,٢٠	١٩٣٠,٥٤	٢٣٩,٧٤-
١٧٠٧,٢٠	٢٠٤,٤٠	٢٠١٥,٨٤	٣٠٨,٦٤-
١٨٨٢,١٠	٤٢٣,٦٠	٣٠٢٣,٨٤	١١٤١,٧٤-
٣٤٥٥,٦٠	٤١٦,٣٠	٢٩٩٢,٥٦	٤٦٣,٠٤
٤٢٠٤,٩٠	٤٠٩,٨٠	٢٩٦٤,٥٧	١٢٤٠,٣٣
٤٠٦١,١٠	٤٥٢,٨٠	٣١٤٧,٤٥	٩١٣,٦٥
٤٠١٨,٧٠	٤٩٥,٨٠	٣٣٢٥,١١	٦٩٣,٥٩
٤٥٤٠,٠٠	٧٢٣,٥٠	٤١٨٢,٩٩	٣٥٧,٠١
٥٦١٠,٠٠	١٥١٥,٢٠	٦٢٦٢,٧٦	٦٥٢,٧٦-
٥٧٨١,٠٠	١٦٦٧,٩٠	٦٥٣٨,٤٥	٧٥٧,٤٥-
٥٩٣٥,٠٠	١٨٢١,٣٠	٦٧٨٥,٥٢	٨٥٠,٥٢-
٦١٣٠,٠٠	١٩٠٥,٢٠	٦٩٠٩,٦٣	٧٧٩,٦٣-
٧٩٣٩,٠٠	١٩٩٥,٥٠	٧٠٣٥,٥٨	٩٠٣,٤٢
٨١٣٦,٠٠	٢٢٨٨,٥٠	٧٤٠١,٠٧	٧٣٤,٩٣
٨٤٤٨,٠٠	٢٧٠٦,١٠	٧٨٥٨,٣٠	٥٨٩,٧٠
٨٦٩٣,٠٠	٣٣٨٧,٦٠	٨٦٦٥,٢٩	٢٧,٧١
٨٩٦٠,٠٠	٣٧٢٩,٧٠	٩٢٠٦,٠٦	٢٤٦,٠٦-

٣-٤ الانحدار الخطى المتعدد

يوضح الانحدار الخطى المتعدد العلاقة بين متغير واحد تابع وأكثر من متغير مستقل، وهناك العديد من الأمثلة في هذا الصدد كالعلاقة بين الكمية المطلوبة من سلعة معينة وتأثرها بسعر السلعة وأسعار السلع الأخرى والدخل كمتغيرات مستقلة ، والعلاقة الخطية تعني أن التغير في أحد المتغيرات المستقلة بوحدة واحدة يؤدي إلى تغير ثابت بالمتغير التابع. كما ذكر سابقاً فإن نموذج الانحدار الخطى في تقدير العلاقات وخاصة الاجتماعية يتطلب إدخال ما يسمى بالخطأ العشوائي حيث أن الأفراد لا يتصرفون بنفس الطريقة أو أن تفضيلاتهم متماثلة. وبذلك يمكن كتابة النموذج في هذه الحالة على الصورة.

$$y = b_0 + b_1x_1 + b_2x_2 + b_3x_3 + \dots + u_i$$

حيث u_i تشير إلى الأخطاء العشوائية ويخضع لنفس الافتراضات السابق ذكرها بالانحدار الخطى البسيط.

لتقدير معالم الدالة السابقة... b_0, b_1, b_2 فيمكن تطبيق طريقة المربعات الصغرى العادية إلا أنه يلاحظ إضافة شرط هام في هذه الحالة وهو .

— استقلال المتغيرات $x_1, x_2, x_3 \dots$ عن بعضها البعض أي لا يرتبط أى زوجين منها.

نبحث عن b_1, b_0 التي تجعل مجموع مربعات الأخطاء أقل ما يمكن أى

$$\begin{aligned} \text{Min } \sum (y - \hat{y})^2 \\ \sum e_i^2 &= \text{Min } \sum (y - \hat{y})^2 \\ &= \text{Min } \sum (y - b_0 - b_1 x_1 - b_2 x_2 \dots)^2 \end{aligned}$$

لذلك فيتم التفاضل الجزئى بالنسبة لكل معلمة وينتج لدينا مجموعة المعادلات التالية:—

$$\sum y = b_0 n + b_1 \sum x_1 + b_2 \sum x_2 + b_3 \sum x_3 + \dots)$$

$$\sum x_1 y = b_0 \sum x_1 + b_1 \sum x_1^2 + b_2 \sum x_1 x_2 + b_3 \sum x_1 x_3 + \dots)$$

$$\sum x_2 y = b_0 \sum x_2 + b_1 \sum x_2 x_1 + b_2 \sum x_2^2 + b_3 \sum x_3 x_2 + \dots)$$

$$\sum x_n y = b_0 \sum x_n + b_1 \sum x_n x_1 + b_2 \sum x_n x_2 + b_n \sum x_n^2 + \dots)$$

ويمكن حل هذه المجموعة من المعادلات إما بأسلوب المحددات أو المصفوفات وسنستخدم هنا

أسلوب المصفوفات حيث يمكن التعبير عنها بالشكل التالى :-

$$M_{xy} = (X' X) B$$

حيث المصفوفة $(X' X)$ عبارة عن مصفوفة مربعة تتكون عناصرها من حواصل ضرب القراءات

للمتغيرات أو مربعاتها أى أن

$$X' X = \begin{bmatrix} 1 & 1 & 1 & \dots & 1 \\ x_{11} & x_{12} & x_{13} & \dots & x_{1n} \\ x_{21} & x_{22} & x_{23} & \dots & x_{2n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ x_{n1} & x_{n2} & x_{n3} & \dots & x_{nn} \end{bmatrix} \begin{bmatrix} 1 & x_{11} & x_{21} & \dots & x_{n1} \\ 1 & x_{12} & x_{22} & \dots & x_{n2} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_{13} & x_{23} & \dots & x_{n3} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_{1n} & x_{2n} & \dots & x_{nn} \end{bmatrix} = \begin{bmatrix} n & \sum x_{1i} & \sum x_{2i} & \dots & \sum x_{ni} \\ \sum x_{1i} & \sum x_{1i}^2 & \sum x_{1i} x_{2i} & \dots & \dots \\ \sum x_{2i} & \sum x_{1i} x_{2i} & \sum x_{2i}^2 & \dots & \dots \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \sum x_{ni} & \dots & \dots & \dots & \sum x_{ni}^2 \end{bmatrix}$$

وبالتالى فإن تقدير المعالم B يكون

$$\hat{B} = (X'X)^{-1} X'Y$$

وبالتالى يمكن الحصول على متجه المعالم ، كما يمكن استنتاج تباين المعالم حيث

$$V(\hat{B}) = (X'X)^{-1} \hat{\sigma}_E^2$$

حيث $(X'X)^{-1}$ هو مقلوب المصفوفة $(X'X)$

$\hat{\sigma}_E^2$ هو تقدير تباين الأخطاء ، والذي يمكن من اختبار المعنوية الإحصائية لهذه المعالم حيث يلاحظ أن

جميع العناصر على القطر الرئيسى لهذه المصفوفة تمثل تباين التقديرات $\hat{b}_0$ ، $\hat{b}_1$ ، $\hat{b}_2$ ،

والعناصر الأخرى على جانب القطر الرئيسى تمثل تباير هذه التقديرات مع بعضها البعض وهذه البيانات

تعطى ضمن نتائج برنامج SPSS

ويمكن كتابة جدول تحليل التباين فى صورته العامة كالتالى:

جدول تحليل التباين

فى حالة الانحدار الخطى المتعدد

المحسوبة	متوسط مجموع المربعات	مجموع المربعات	درجات الحرية	مصدر التغير
$(M_{XR}) \hat{\sigma}_E^2$	$\frac{SSR}{k} = M_{XR}$	$SSR = \sum(\hat{y} - \bar{y})^2$	k	الانحدار
	$\frac{SSE}{n-k-1} = \hat{\sigma}_E^2$	$SSE = \sum(y - \hat{y})^2$	n-k-1	البواقي (الانحراف عن خط الانحدار)
		$SST = \sum(y - \bar{y})^2$	n-1	الكلى

حيث تستخدم قيمة F بدرجات حرية k' , (n-k-1) لاختبار قوة العلاقة بين المتغيرات

المستقلة والمتغير التابع اى المعنوية الكلية للانحدار ، كما يمكن أن يختبر R^2 معامل التحديد وذلك

بحساب قيمة F وبدرجات حرية k , (n-k-1) حيث

$$F_{k,(n-k-1)} = \frac{R^2 / k}{(n-k-1) / (1-R^2)}$$

وتقارن F المحسوبة بنظيرتها الجدولة وبدرجات الحرية للبسط K والمقام $(n-k-1)$ ومستوى المعنوية المطلوب ٥% أو ١% . أى بدرجات ثقة ٩٥% ، ٩٩% على التوالي
كما يمكن اختبار المعنوية الإحصائية للمعالم المقدرة باستخدام توزيع t وبدرجات حرية $(n-k-1)$
فلو فرض أننا نريد اختبار المعنوية الإحصائية للمعلمة b_i

$$H_0 : \hat{b}_i = 0 \quad \text{فإننا نفترض فرض العدم أى}$$

$$H_0 : \hat{b}_i \neq 0 \quad \text{مقابل الفرض البديل أى}$$

$$t_\alpha = \frac{\hat{b}_i}{\sqrt{\text{var}(\hat{b}_i)}}$$

كما يمكن اختبار الفرض الخاص بالقيمة الحقيقية b_i وإنشاء فترة ثقة لهذه القيمة أى

$$t_\alpha = \frac{\hat{b}_i - b_i}{\sqrt{\text{var}(\hat{b}_i)}}$$

أى

$$t_\alpha \sqrt{\text{var}(\hat{b}_i)} + \hat{b}_i > b_i > \hat{b}_i - t_\alpha \sqrt{\text{var}(\hat{b}_i)}$$

كما يمكن اختبار ما إذا كان زوجين من هذه المعالم يساوى بعضهم البعض فإن نفترض :

$$H_0 : b_2 = b_3 \quad \text{فرض العدم}$$

$$\hat{H}_1 : b_2 \neq b_3 \quad \text{مقابل الفرض البديل}$$

وتكون t_α مساوية للمقدار

$$t_\alpha = \frac{\hat{b}_2 - \hat{b}_3}{\sqrt{v(\hat{b}_2) + v(\hat{b}_3) - 2\text{cov}(\hat{b}_2, \hat{b}_3)}}$$

حيث $\text{cov}(\hat{b}_2, \hat{b}_3)$ استخراج من مصفوفة التباين سالفة الذكر

يلاحظ أن إضافة متغيرات مستقلة من شأنه أن يرفع قيمة المربعات الخاصة بالانحدار وبالتالي فإن قيمة R^2 تزيد وإذا أخذ في الاعتبار نقص عدد درجات الحرية مع إضافة متغيرات إضافية مستقلة فيمكن حساب المعدلة أو R^2 من الصيغة

$$R^2 = 1 - \frac{(1 - R^2)(n-1)}{n-k}$$

حيث n هي عدد القراءات، k عدد المعالم المقدرة

ويمكن استنتاج عندما

$$\bar{R}^2 = R^2 \text{ فإن } 1 = k \text{ ويكون } (n-1)/(n-k) = 1$$

عندما

$$\bar{R}^2 < R^2 \text{ فإن } 1 < k \text{ فتكون } (n-1)/(n-k) > 1$$

وعندما تكون n كبيرة لقيمة معينة k فإن المقدار $(n-1)/(n-k)$ يقترب من

الواحد ولن تختلف كثيرا $\bar{R}^2$ عن R^2 أما عندما تكون n صغيرة وتكون k كبيرة بالنسبة إلى

n فإن $\bar{R}^2$ تكون اصغر كثيرا من R^2 وقد تكون $\bar{R}^2$ سالبة بالرغم من أن $0 \leq R^2 \leq 1$

صور أخرى لبعض دوال تحليل الانحدار

قد تكون العلاقة بين المتغير المستقل والمتغير التابع غير خطية إلا أنه في بعض الأحوال يمكن تحويل

مثل هذه العلاقات إلى علاقات خطية، وهذا يمكن استخدام أسلوب المربعات الصغرى OLS (مع تحقق الشروط) لتقدير معاملها، ومن هذه الدوال الأمثلة التالية.

$$1- y = b_0 + b_1 x^b + e^u \Rightarrow \bar{y}^* = b_0^* + b_1 x^* + u$$

حيث $y^* = \ln y$ ، $x^* = \ln x$ ، u الأخطاء العشوائية. وهي دالة لوغاريتمية من الطرفين

$$2- \ln y = b_0 + b_1 x + u \Rightarrow y^* = b_0 + b_1 x + u$$

وهي دالة نصف لوغاريتمية

$$3- y = b_0 + \frac{b_1}{x} + u \Rightarrow y = b_0 + b_1 z + u$$

$z = \frac{1}{x}$

$$z = \frac{1}{x} \quad \text{حيث}$$

$$4- y = b_0 + b_1 x + b_2 x^2 + u \Rightarrow y = b_0 + b_1 x + bx^2 + u$$

$$w = x^2 \quad \text{حيث}$$

الانحدار المتعدد Multiple Regression من خلال البرنامج الإحصائي SPSS
ويمكن توضيح ذلك من خلال المثال التالي :

مثال : البيانات التالية تمثل العلاقة بين الناتج المحلي الإجمالي الثابت (y) ، الاستثمارات الزراعية الثابتة (x1) ، المساحة الزراعية (x2) ، العمالة الزراعية بالألف (x3) ، الأجور الزراعية (x4) ، وذلك في سلسلة زمنية من سنة ١٩٧٥ إلى سنة ١٩٩٥ ، ونريد تحديد شكل للعلاقة بين المتغير التابع (y) والمتغيرات الزراعية (x1,x2,x3,x4) .

السنوات Year	الناتج المحلي الزراعي الثابت y	الاستثمارات الزراعية الثابتة x1	المساحة الزراعية X2	العمالة الزراعية بالألف X3	الأجور الزراعية X4
١٩٧٥	٩٩٦,٩٠	٨٣,٤٠	٥,٧٩	٤٢١٨,٠٠	٢٦٧,٤٠
١٩٧٦	١٠٦٣,٧٠	٩٨,٠٠	٥,٨٦	٤٠٦٨,٠٠	٢٦٢,١٠
١٩٧٧	١٦١٧,٨٠	١١٤,٥٠	٥,٨٥	٤١٠٤,٠٠	٢٨٧,٧٠
١٩٧٨	١٥١٥,٧٠	١٤٤,١٠	٥,٨٦	٤١٣٥,٠٠	٣١٦,٢٠
١٩٧٩	١٦٩٠,٨٠	١٨٧,٢٠	٥,٨٦	٤١٦٥,٠٠	٣٤٧,٨٠
١٩٨٠	١٧٠٧,٢٠	٢٠٤,٤٠	٥,٨٧	٤٢٠٠,٠٠	٤٤٠,٣٠
١٩٨١	١٨٨٢,١٠	٤٢٣,٦٠	٥,٨٨	٤٢٢٥,٠٠	٩٧١,١٠
١٩٨٢	٣٤٥٥,٦٠	٤١٦,٣٠	٥,٩٢	٤٢٦٤,٠٠	١٠٥٥,٦
١٩٨٣	٤٢٠٤,٩٠	٤٠٩,٨٠	٥,٨٤	٤٣٠٤,٠٠	١١٤١,١
١٩٨٤	٤٠٦١,١٠	٤٥٢,٨٠	٥,٨٣	٤٣٤٤,٠٠	١٢٢٨,٩
١٩٨٥	٤٠١٨,٧٠	٤٩٥,٨٠	٥,٩٧	٤٣٨٤,٠٠	١٢٧١,٩
١٩٨٦	٤٥٤٠,٠٠	٧٢٣,٥٠	٥,٩٧	٤٣٣٠,٠٠	١٥٨٦,١
١٩٨٧	٥٦١٠,٠٠	١٥١٥,٢٠	٦,٠٠	٤٣٨١,٠٠	١٨٥٥,٩
١٩٨٨	٥٧٨١,٠٠	١٦٦٧,٩٠	٦,١٨	٤٤٣٢,٠٠	٢١٠٥,٤
١٩٨٩	٥٩٣٥,٠٠	١٨٢١,٣٠	٦,٢٧	٤٤٧٨,٠٠	٢٣٨٢,٤

الأجور الزراعية X4	العمالة الزراعية بالألف X3	المساحة الزراعية X2	الاستثمارات الزراعية الثابتة x1	الناتج المحلي الزراعي الثابت y	السنوات Year
٢٦٩٨,٠	٤٥٣٣,٠٠	٦,٩٢	١٩٠٥,٢٠	٦١٣٠,٠٠	١٩٩٠
٣٠١٢,٠	٤٥٨٥,٠٠	٧,٠٢	١٩٩٥,٥٠	٧٩٣٩,٠٠	١٩٩١
٣٤٢٣,٠	٤٦٢٤,٠٠	٧,١٢	٢٢٨٨,٥٠	٨١٣٦,٠٠	١٩٩٢
٣٨٩٥,٠	٤٦٨٢,٠٠	٧,١٧	٢٧٠٦,١٠	٨٤٤٨,٠٠	١٩٩٣
٤٣٦٩,٠	٤٧٤٤,٠٠	٧,١٨	٣٣٨٧,٦٠	٨٦٩٣,٠٠	١٩٩٤
٤٩٤٦,٠	٤٨١٢,٠٠	٧,١٨	٣٧٢٩,٧٠	٨٩٦٠,٠٠	١٩٩٥

ولإجراء ذلك نتبع الآتي :

١- من قائمة Analyze اختر أمر Regression ومنها اختر أمر Linear يظهر الصندوق الحواري التالي :

The screenshot shows the 'Linear Regression' dialog box. On the left, a list of variables includes x1, x2, x3, and x4. The 'Dependent' field is set to 'y'. The 'Independent(s)' field contains 'x1' and 'x2'. The 'Method' dropdown is set to 'Enter'. The 'WLS' checkbox is checked. The 'Statistics', 'Plots', 'Save', and 'Options' buttons are visible at the bottom.

٢- اختر المتغير التابع (y) في خانة Dependent

٣- اختر المتغيرات المستقلة (x1,x2,x3,x4) في خانة Independent(s)

٤- اختر Enter من خانة Method ، ثم اختر أمر OK تظهر النتائج التالية :

Regression

Variables Entered/Removed^a

Model	Variables Entered	Variables Removed	Method
1	X4, X2, X3, X1 ^a		Enter

a. All requested variables entered.

b. Dependent Variable: Y

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Durbin-Watson
1	.975 ^a	.950	.938	685.2465	1.024

a. Predictors: (Constant), X4, X2, X3, X1

b. Dependent Variable: Y

ANOVA^b

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	1.4E+08	4	3.6E+07	76.377	.000 ^a
	Residual	7513004	16	469562.7		
	Total	1.5E+08	20			

a. Predictors: (Constant), X4, X2, X3, X1

b. Dependent Variable: Y

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients		Sig.
		B	Std. Error	Beta		
1	(Constant)	-22902.9	18529.063		-1.236	.234
	X1	-.677	.865	.283	-.783	.445
	X2	-222.844	788.384	.045	-.283	.781
	X3	6.138	4.310	.483	1.424	.174
	X4	1.549	1.129	.814	1.372	.189

a. Dependent Variable: Y

Casewise Diagnostics^a

Case Number	Std. Residual	Y	Predicted Value	Residual
1	-1.544	996.90	2054.6376	-1057.74
2	-.053	1063.70	1100.2525	-36.5525
3	.388	1617.80	1351.9403	265.8597
4	-.071	1515.70	1564.1097	-48.4097
5	-.113	1690.80	1768.0372	-77.2372
6	-.591	1707.20	2112.2900	-405.0900
7	-1.540	1882.10	2937.4455	-1055.35
8	.222	3455.60	3303.7497	151.8503
9	.731	4204.90	3703.9410	500.9590
10	.004	4061.10	4058.6028	2.4972
11	-.426	4019.70	4310.4391	-291.7391
12	.333	4540.00	4311.6402	228.3598
13	1.619	5610.00	4500.2437	1109.7563
14	1.058	5781.00	5056.3427	724.6573
15	.425	5935.00	5643.9338	291.0662
16	-.203	6130.00	6268.7946	-138.7946
17	1.383	7939.00	6990.9958	948.0032
18	.714	8136.00	7646.5125	489.4875
19	.012	8448.00	8439.9776	8.0224
20	-.581	8683.00	8991.4431	-398.4431
21	-1.767	8960.00	10171.1696	-1211.17

a. Dependent Variable: Y

Residuals Statistics^a

	Minimum	Maximum	Mean	Std. Deviation	N
Predicted Value	1100.2526	10171.17	4589.8333	2678.1985	21
Residual	-1211.17	1109.7563	-5.37E-12	612.9031	21
Std. Predicted Value	-1.303	2.084	.000	1.000	21
Std. Residual	-1.767	1.619	.000	.894	21

a. Dependent Variable: Y

يلاحظ من النتائج السابقة عدم معنوية معاملات الانحدار على الرغم من ارتفاع قيمة معامل التحديد في معادلة الانحدار ، ويرجع ذلك غالباً إلى وجود ارتباط بين المتغيرات المفسرة (المستقلة) مما يعنى عدم ملائمة طريقة المربعات الصغرى العادية لتقدير تلك المعامل ، وسوف نناقش هذه المشكلة فيما بعد .

الباب الخامس،

مشاكل إسقاط الفروض عند تقدير معالم نموذج الانحدار

عرضنا فيما سبق لأسلوب تقدير معلمات نموذج الانحدار الخطى بطريقة المربعات الصغرى العادية وما تنطوي عليه من فروض أساسية إلا أنه إذا لم يتحقق أحاد هذه الشروط سوف يؤدي إلى عدة مشاكل، وفيما يلي بعض هذه المشاكل وإمكانية التغلب عليها :-

١-٥ الارتباط الخطى بين المتغيرات المستقلة Multicollinearity

تنشأ مشكلة الارتباط الخطى عندما يكون واحد على الأقل من المتغيرات المستقلة مرتبط بعلاقة خطية مع متغير أو متغيرات أخرى وغالباً ما يكون هذا الارتباط غير تام ، مما يصعب معه تقدير معالم المتغير بمعزل عن المتغيرات الأخرى ويمكن التعرف على مشكلة الارتباط الخطى من خلال :-

١. ارتفاع قيمة معامل التحديد R^2 مصاحب بتقديرات غير معنوية لمعظم المعالم . وعلى الرغم من حساسية ذلك إلا أن لها عدم ميزة (عيب) يتمثل في أنها شديدة جداً بمعنى أن الارتباط الخطى يعتبر ضار فقط عندما يكون كل تأثيرات المتغيرات المفسرة على المتغير y غير محددة disentangled .

٢. ارتفاع قيمة معامل الارتباط بين أزرانج المتغيرات ، ٨ ، ٠ ، فأكثر .

٣. كبر تباينات المعالم المقدرة ومن ثم اتساع فترات الثقة الخاصة بها وهناك عدة طرق لاختبار هذه المشكلة مثل اختبار كلاين Klein test ، واختبار فارار جلوبير Farrar Glauber Test ويمكن التغلب على مشكلة وجود الارتباط الخطى للمتغيرات المستقلة بمجموعة من الحلول التي يمكن أن تخفف من تأثير هذا الارتباط :-

١- استخدام عدد أكبر من القراءات أى تكبير حجم العينة لأن الارتباط الخطى مشكلة عينة .

٢- استخدام معلومات مسبقة في حالة توافرها فمثلاً إذا كنا بصدد تقدير معالم دالة الاستهلاك

$$y = b_0 + b_1 x_1 + b_2 x_2$$

حيث y قيمة الاستهلاك : x_1 تمثل الدخل ، x_2 تمثل الثروة ، وإذا كان هناك

علاقة بين x_1 ، x_2 بحيث أن $b_2 = c b_1$ حيث c نسبة ثابتة معينة، في هذه الحالة

٣- يمكن استخدام بيانات مركبة من سلاسل زمنية ومقطعية فلو فرض أننا بصدد تقدير دالة الطلب على سلعة معينة

$$y = b_0 + b_1 x_1 + b_2 x_2$$

حيث y الكمية المباعة ، x_1 متوسط السعر ، x_2 الدخل وإذا كان هناك ارتباط جوهري بين x_1 ، x_2 ، فيمكن التغلب على هذه المشكلة عن طريق تقدير العلاقة بين الكمية المباعة والدخل من بيانات مقطعية [على مستوى المحافظات-المبانيات ... الخ] حيث يكون السعر في هذه الحالة ثابت ثم تستخدم المعلمة الخاصة بالدخل b_2 في تقدير b_1 من بيانات السلسلة الزمنية.

٤- يمكن إسقاط بعض المتغيرات التفسيرية ذات الارتباط الخطي، خاصة إذا كانت قليلة الأهمية في التأثير على الظاهرة موضوع البحث ، ففي المثال الموجود بالنقطة رقم (٢) يمكن إسقاط الثروة من العلاقة -إلا أن ذلك قد يؤدي في بعض الأحيان إلى مشاكل خاصة بالنموذج وما قد يتبعه من مشاكل الارتباط الذاتي.

٥- تحويل المتغيرات قد يساعد على التخفيف من مشكلة الارتباط الخطي وذلك باستخدام الفروق الأولى للمتغيرات.

٦- هناك طرق أخرى للتغلب على مشكلة الارتباط الخطي وذلك باستخدام أسلوب تحليل العوامل Factor analysis وأسلوب المكونات الرئيسية Principal Component وغيرها من الطرق.

٧- برنامج SPSS يمكن من استخدام أسلوب الانحدار المتدرج في هذا الشأن ، جدير بالذكر أن نموذج الانحدار الخطي الذي يتصف بمشكلة الارتباط الخطي يعتبر مقبولا إذا ما كان الهدف منه التنبؤ.

معالجة الازدواج الخطي باستخدام البرنامج الإحصائي SPSS (الانحدار المتدرج Stepwise) :

يبني الانحدار المتعدد على عدة فروض أحدها ألا يكون هناك علاقة بين المتغيرات المستقلة ، ولكن في الواقع العملي قد يوجد ارتباط بين المتغيرات المستقلة ، ولحل هذه المشكلة نستخدم الانحدار التدريجي Stepwise حيث يتم حساب معاملات الارتباط لجميع المتغيرات الداخلة في النموذج ، ثم يتم اختيار أكثر المتغيرات المستقلة تأثيراً على معامل التحديد وعلى اختبار F ، ثم يضيف متغيراً مستقلاً آخر يلى المتغير الأول في الارتباط بالمتغير التابع وهكذا مع الأخذ في الاعتبار عدم وجود علاقة بين المتغيرات المستقلة الداخلة في النموذج

ولتنفيذ ذلك على تابع الخطوات التالية :

١- سجل البيانات التالية على ملف وليكن step.sav

الأجور الزراعية X4	العمالة الزراعية بالألف X3	المساحة الزراعية X2	الاستثمارات الزراعية الثابتة x1	الناتج المحلي الزراعي الثابت y	السنوات Year
٢٦٧,٤٠	٤٢١٨,٠٠	٥,٧٩	٨٣,٤٠	٩٩٦,٩٠	١٩٧٥
٢٦٢,١٠	٤٠٦٨,٠٠	٥,٨٦	٩٨,٠٠	١٠٦٣,٧٠	١٩٧٦
٢٨٧,٧٠	٤١٠٤,٠٠	٥,٨٥	١١٤,٥٠	١٦١٧,٨٠	١٩٧٧
٣١٦,٢٠	٤١٣٥,٠٠	٥,٨٦	١٤٤,١٠	١٥١٥,٧٠	١٩٧٨
٣٤٧,٨٠	٤١٦٥,٠٠	٥,٨٦	١٨٧,٢٠	١٦٩٠,٨٠	١٩٧٩
٤٤٠,٣٠	٤٢٠٠,٠٠	٥,٨٧	٢٠٤,٤٠	١٧٠٢,٢٠	١٩٨٠
٩٧١,١٠	٤٢٢٥,٠٠	٥,٨٨	٤٢٣,٦٠	١٨٨٢,١٠	١٩٨١
١٠٥٥,٦	٤٢٦٤,٠٠	٥,٩٢	٤١٦,٣٠	٣٤٥٥,٦٠	١٩٨٢
١١٤١,١	٤٣٠٤,٠٠	٥,٨٤	٤٠٩,٨٠	٤٢٠٤,٩٠	١٩٨٣
١٢٢٨,٩	٤٣٤٤,٠٠	٥,٨٣	٤٥٢,٨٠	٤٠٦١,١٠	١٩٨٤
١٢٧١,٩	٤٣٨٤,٠٠	٥,٩٧	٤٩٥,٨٠	٤٠١٨,٧٠	١٩٨٥
١٥٨٦,١	٤٣٣٠,٠٠	٥,٩٧	٧٢٣,٥٠	٤٥٤٠,٠٠	١٩٨٦
١٨٥٥,٩	٤٣٨١,٠٠	٦,٠٠	١٥١٥,٢٠	٥٦١٠,٠٠	١٩٨٧
٢١٠٥,٤	٤٤٣٢,٠٠	٦,١٨	١٦٦٧,٩٠	٥٧٨١,٠٠	١٩٨٨
٢٣٨٢,٤	٤٤٧٨,٠٠	٦,٢٧	١٨٢١,٣٠	٥٩٣٥,٠٠	١٩٨٩
٢٦٩٨,٠	٤٥٣٣,٠٠	٦,٩٢	١٩٠٥,٢٠	٦١٣٠,٠٠	١٩٩٠
٣٠١٢,٠	٤٥٨٥,٠٠	٧,٠٢	١٩٩٥,٥٠	٧٩٣٩,٠٠	١٩٩١
٣٤٢٣,٠	٤٦٢٤,٠٠	٧,١٢	٢٢٨٨,٥٠	٨١٣٦,٠٠	١٩٩٢
٣٨٩٥,٠	٤٦٨٢,٠٠	٧,١٧	٢٧٠٦,١٠	٨٤٤٨,٠٠	١٩٩٣
٤٣٦٩,٠	٤٧٤٤,٠٠	٧,١٨	٣٣٨٧,٦٠	٨٦٩٣,٠٠	١٩٩٤
٤٩٤٦,٠	٤٨١٢,٠٠	٧,١٨	٣٧٢٩,٧٠	٨٩٦٠,٠٠	١٩٩٥

٢- من قائمة Analyze اختر أمر Regression ، ومنه اختر أمر Linear يظهر الصندوق الحوارى

Linear Regression

Linear Regression

Dependent: # y

Block 1 of 1 Next

Independent(s): # x1, # x2, # x3

Method: Stepwise

Selection Variable:

Case Labels:

WLS >> Statistics... Plots... Save... Options...

OK Paste Reset Cancel Help

١- أمام خانة **Dependent** اختر المتغير التابع **Y**

٢- أمام خانة **Independent(s)** اختر المتغيرات المستقلة ، ولتكن **x1 , x2 , x3 , x4**

٣- أمام خانة **Method** اختر **Stepwise** ، ثم اختر أمر **Ok** ، فتظهر النتائج التالية :

Regression

Variables Entered/Removed^a

Model	Variables Entered	Variables Removed	Method
1	X3		Stepwise (Criteria: Probability-of-F-to-enter <= .050, Probability-of-F-to-remove >= .100).

a. Dependent Variable: Y

Model Summary^a

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Durbin-Watson
1	.971 ^a	.943	.940	671.7121	1.230

a. Predictors: (Constant), X3

b. Dependent Variable: Y

ANOVA^b

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	1.4E+08	1	1.4E+08	315.594	.000 ^a
	Residual	8572745	19	451197.1		
	Total	1.5E+08	20			

a. Predictors: (Constant), X3

b. Dependent Variable: Y

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	-49538.3	3050.428		-16.240	.000
	X3	12.354	.695	.971	17.765	.000

a. Dependent Variable: Y

Excluded Variables^b

Model		Beta In	t	Sig.	Partial Correlation	Collinearity Statistics
						Tolerance
1	X1	.121 ^a	.601	.555	.140	7.641E-02
	X2	.049 ^a	.348	.732	.082	.159
	X4	.396 ^a	1.298	.211	.293	3.092E-02

a. Predictors in the Model: (Constant), X3

b. Dependent Variable: Y

Casewise Diagnostics^a

Case Number	Std. Residual	Y	Predicted Value	Residual
1	-2.341	996.90	2569.7039	-1572.80
2	.517	1063.70	716.6440	347.0560
3	.679	1617.80	1161.3784	456.4216
4	-.043	1515.70	1544.3441	-28.6441
5	-.334	1690.80	1914.9561	-224.1561
6	-.953	1707.20	2347.3367	-640.1367
7	-1.152	1882.10	2656.1800	-774.0800
8	.473	3455.60	3137.9756	317.6244
9	.853	4204.90	3632.1249	572.7751
10	-.097	4061.10	4126.2742	-65.1742
11	-.896	4018.70	4620.4235	-601.7235
12	.873	4540.00	3953.3220	586.6780
13	1.528	5610.00	4583.3623	1026.6377
14	.845	5781.00	5213.4027	567.5973
15	.228	5935.00	5781.6744	153.3256
16	-.493	6130.00	6461.1297	-331.1297
17	1.244	7939.00	7103.5238	835.4762
18	.820	8136.00	7585.3193	550.6807
19	.218	8448.00	8301.8358	146.1642
20	-.558	8693.00	9067.7673	-374.7673
21	-1.411	8960.00	9907.8211	-947.8211

a. Dependent Variable: Y

Residuals Statistics^a

	Minimum	Maximum	Mean	Std. Deviation	N
Predicted Value	716.6440	9907.8213	4589.8333	2668.2879	21
Residual	-1572.80	1026.6377	-4.31E-12	654.7039	21
Std. Predicted Value	-1.452	1.993	.000	1.000	21
Std. Residual	-2.341	1.528	.000	.975	21

a. Dependent Variable: Y

يلاحظ أن البرنامج اكتفى بالمتغير X3 نظراً لوجود ارتباط خطي بين المتغيرات المستقلة ، مما يعني أن الناتج المحلي الزراعي الثابت يتأثر بالدرجة الأولى بالعملة الزراعية X3

٢-٥ الارتباط الذاتي للأخطاء Autocorrelation

استقلال الأخطاء عن بعضها من الافتراضات الهامة في نموذج الانحدار الخطي أي أن

$$E(u_i - u_j) = 0 \quad i \neq j$$

وعدم تحقق هذا الشرط يعنى وجود ما يسمى بالارتباط الذاتي للأخطاء أو الارتباط السلسلى **Serial correlation** ، بمعنى أن قيمة u في فترة معينة مرتبطة بقيمتها في فترات أخرى ، ويعرف الارتباط الذاتي للأخطاء من الدرجة الأولى بأنه ارتباط الأخطاء بالفترة التي تليها.

هذا وقد يكون الارتباط الذاتي سالبا عندما تعبر الأخطاء المتتالية إشاراتها كثيرا ، ولكن عندما تكون لعدة بواقي متتالية لها نفس الإشارة فيكون الارتباط الذاتي موجبا.

ويؤدي إسقاط فرض استقلال الأخطاء أو ظهور مشكلة الارتباط الذاتي إلى تقليل كفاءة المعالم المقدرة ، فمثلاً إذا وجد أن هناك ارتباط سلسلى موجب فإن الانحرافات المعيارية للمعالم المقدرة في هذه الحالة تكون أقل من المفروض أن تكون عليه مما يترتب عليه مبالغة في الدقة وفي المعنوية الإحصائية لهذه المعالم ، كما يؤدي بالتبعية إلى ارتفاع قيمة معامل التحديد R^2 . كما يؤدي هذا الارتباط في الأخطاء إلى أن التنبؤات المحسوبة بطريقة المربعات الصغرى يكون لها تباين اكبر من التنبؤات التي يمكن الحصول عليها بطريقة أخرى كطريقة المربعات الصغرى العامة .

ويجدر الإشارة في هذا المقام أن الارتباط الذاتي لا يؤثر على تحيز المعالم المقدرة.

➤ الكشف عن الارتباط الذاتي

يمكن استنتاج وجود ارتباط سلسلى وذلك بتتبع نمط تغير الاتجاه العام للأخطاء المقدرة e_t من معادلة الانحدار ، كما يمكن إيجاد معامل الارتباط بين السلسلتين e_t و e_{t-1} بالصيغة التالية

$$Re_t e_{t-1} = \frac{\sum e_t e_{t-1}}{\sqrt{\sum e_t^2 \sum e_{t-1}^2}}$$

وقد يكون هذه الطرق غير عالية الدقة ومن ثم لابد من التأكد أن عنصر الارتباط الذاتي لا يرجع إلى عامل الصدفة ، ولذلك لابد من إجراء اختبار المعنوية ، وأهم هذه الاختبارات هو اختبار ديربن-واتسون **Durbin-Watson test** ، حيث يفترض في هذا الشأن أن الارتباط السلسلى يأخذ النمط البسيط ويرمز له بالرمز DW أو DW وهو على الصورة .

$$D = DW = \frac{\sum_{t=1}^T (\hat{e}_t - e_{t-1})^2}{\sum_{t=1}^T e_t^2}$$

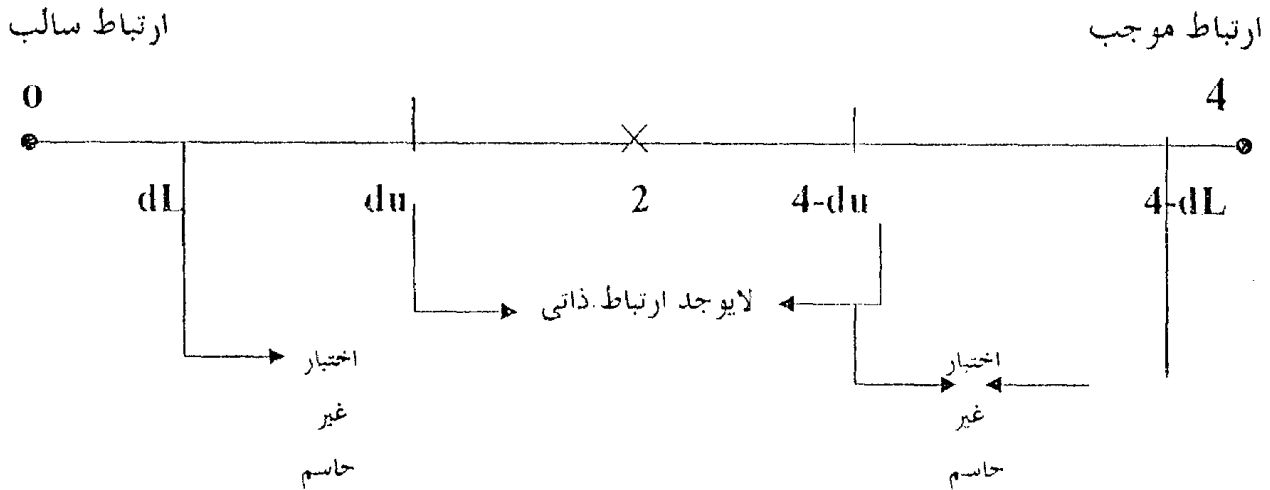
ويأخذ هذا الاختبار القيم بين (0 , 4) وذلك لقيم معامل الارتباط المختلفة حيث

$$DW \approx 2 \Rightarrow \rho = 0$$

$$DW \approx 0 \Rightarrow \rho = 1$$

$$DW \approx 4 \Rightarrow \rho = -1$$

ويوضح الشكل التالي قيم d التي تشير إلى وجود أو غياب ارتباط ذاتي من الدرجة الأولى سالب أو موجب ، كذلك يوضح المناطق التي تجعل هذا الاختبار غير حاسم



حيث dL, du هي القيمة العليا والدنيا في جداول ديربن واتسون وهي مجدولة طبقاً لدرجات الحرية ومستوى المعنوية ١% أو ٥% وبالتالي كلما اقتربت قيمة d من 2 فإن هذا يعني عدم وجود ارتباط ذاتي ويظهر هذا الاختبار في برنامج الحاسب SPSS

معالجة مشكلة الارتباط الذاتي

١- عند وجود الارتباط يمكن تطبيق طريقة المربعات الصغرى العادية بصورة معدلة ، ويكون نموذج الانحدار على الصورة

$$y_t - \rho y_{t-1} = b_0 (1-\rho) + b_1 (x_t - \rho x_{t-1}) + (e_t - \rho e_{t-1})$$

حيث ρ معروفة ونجد في هذه الحالة أن الأخطاء $(e_t - \rho e_{t-1})$ لا ترتبط ويمكن تقدير المعالم بالطريقة المعتادة.

أما إذا كانت ρ غير معروفة فإنها تقدر باستخدام المعادلة

$$\hat{e}_t = \hat{\rho} e_{t-1} + v_t$$

ثم يتم استخدامها في نموذج الانحدار لتقدير المعالم الخاصة به.

يلاحظ أن تحويل العلاقة الأصلية إلى علاقة يمكن قياسها بأسلوب المربعات الصغرى العادية تؤدي

إلى فقد الملاحظة الأولى ويمكن تعويض هذه الملاحظة عن طريق التعويض التالي ،

$$y_1 = \sqrt{1 - \rho^2} \quad x_1 = \sqrt{1 - \rho^2}$$

٢- يمكن افتراض أن قيمة $\rho = 1$ أى وجود ارتباط ذاتي موجب تام وتكون العلاقة على الصورة

$$y_t - y_{t-1} = b_1 (x_t - x_{t-1}) + (e_t - e_{t-1})$$

أى يمكن استخدام الفروق الأولى للمتغيرات الداخلة في نموذج الانحدار لتقدير معالم الدالة باستخدام أسلوب المربعات الصغرى.

٣- إذ كان الارتباط السلسلى يرجع إلى حذف بعض المتغيرات التفسيرية^(*) ، فإن الحل يكون بإدراج هذا المتغير المحذوف صراحة في النموذج ، كما يمكن أن يرجع الارتباط الذاتي إلى خطأ في صياغة نموذج الانحدار ذاته كأن يكون من الدرجة الثانية مثلاً.

٤- يمكن التغلب على مشكلة الارتباط الذاتي بإختيار قيم متتالية لمعامل ارتباط الخطأ ρ والتي قد تأخذ قيمة تتراوح بين صفر ، ١ وتستخدم في تقدير النموذج المعتمد على ρ ويقارن مجموع مربعات الأخطاء الناتجة من كل حالة وإختيار قيمة ρ التي تؤدي إلى أصغر قيمة لمجموع هذه المربعات.

الارتباط الذاتي بين البواقي ، وكيفية معالجته باستخدام البرنامج الإحصائي SPSS

مثال : الجدول التالي يتضمن متغيرين Y كمتغير تابع ، X1 كمتغير مستقل

Years	Y	X1
1962	188	460
1963	192	486
1964	200	561
1965	191	553
1966	232	682
1967	280	1010
1968	297	793
1969	306	840

^(*) يمكن التعرف على ذلك من خلال تقدير انحدار البواقي على كل متغير من المتغيرات التفسيرية المحذوفة.

Years	Y	X1
1970	396	851
1971	420	1117
1972	227	1102
1973	439	1262
1974	564	3532
1975	759	3711
1976	1030	4281
1977	1368	4563
1978	1474	4977
1979	1671	7597
1980	2196	8757
1981	2445	8875
1982	3287	7612
1983	3179	7789
1984	2780	7893
1985	2774	7322
1986	2575	7164

والمطلوب

١- حساب معادلة المحدار Y على $X1$

٢- اختبار مشكلة الارتباط الذاتي باستخدام إحصائية DW

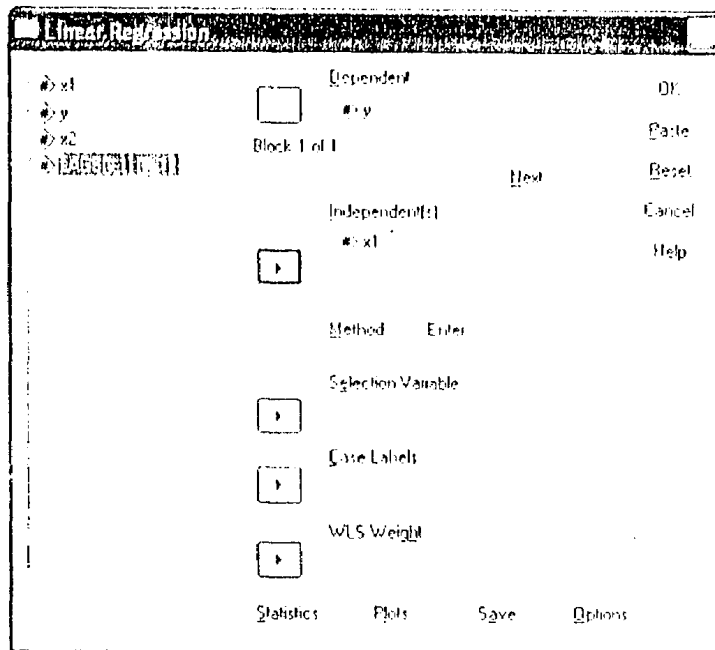
٣- معالجة مشكلة الارتباط الذاتي في حالة وجودها

لتففيذ ذلك باستخدام البرنامج الإحصائي SPSS تابع الخطوات التالية

١- سجل البيانات في Data Editor و احفظ الملف .

٢- من شريط القوائم اختر Analyze ، ومنها اختر أمر Regression ، ومنها Linear

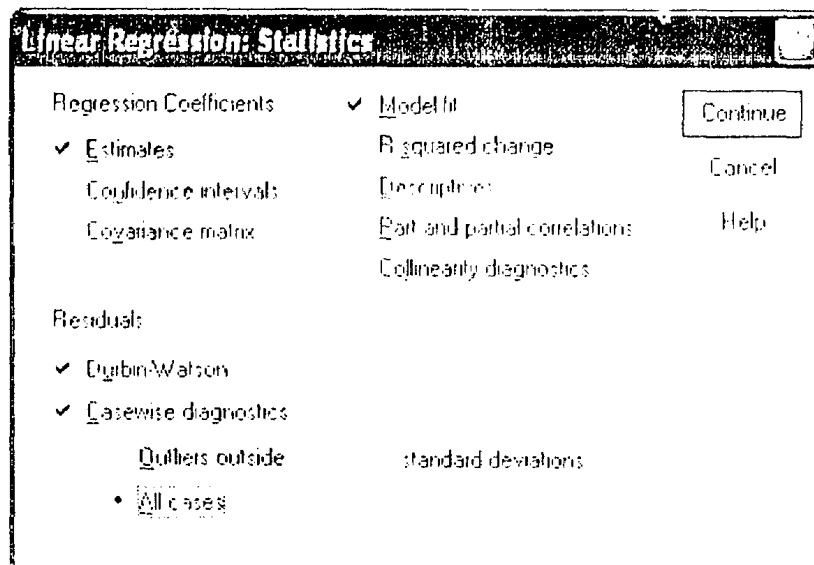
Regression فيظهر الصندوق الحوارى Linear Regression



٣- انقل المتغير Y إلى خانة Dependent

٤- انقل المتغيرين X1 إلى خانة Independent ، وفي خانة Method اختر Enter

٥- اختر أمر Statistics فيظهر الصندوق الحوارى الفرعى Statistics



٦- قم بتشيط الخيارات التالية

- Model fit : وذلك لتقدير R^2 ANOVA,

- Estimates ، لتقدير معالم النموذج

- Durbin-Watson وذلك لحساب إحصائية DW

- Casewise diagnostics ، وذلك لحساب القيم المتوقعة ، والبقاى .

٧- بعد اختيار امر Ok من الصندوق الفرعى ، والصندوق الرئيسى Linear Regression تظهر النتائج التالية

Regression

Variables Entered/Removed^a

Model	Variables Entered	Variables Removed	Method
1	x1 ^a		Enter

a. All requested variables entered.

b. Dependent Variable: y

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Durbin-Watson
1	.941 ^a	.886	.881	378.876	.865

a. Predictors: (Constant), x1

b. Dependent Variable: y

ANOVA^b

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	25555771	1	25555771.07	178.030	.000 ^a
	Residual	3301587	23	143547.258		
	Total ^a	28857358	24			

a. Predictors: (Constant), x1

b. Dependent Variable: y

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	-30.437	118.133		-.258	.799
	x1	.322	.024	.941	13.343	.000

a. Dependent Variable: y

Casewise Diagnostics^a

Case Number	Std. Residual	y	Predicted Value	Residual
1	.185	188	117.83	70.167
2	.174	192	126.21	65.787
3	.131	200	150.39	49.613
4	.114	191	147.81	43.191
5	.112	232	189.39	42.611
6	-.040	280	295.11	-15.112
7	.190	297	225.17	71.833
8	.173	306	240.32	65.684
9	.402	396	243.86	152.138
10	.239	420	329.60	90.399
11	-.258	227	324.77	-97.766
12	.165	439	376.34	62.662
13	-1.436	564	1108.02	-544.017
14	-1.073	759	1165.71	-406.714
15	-.843	1030	1349.44	-319.439
16	-.191	1368	1440.34	-72.335
17	-.263	1474	1573.78	-99.778
18	-1.972	1671	2418.27	-747.271
19	-1.574	2196	2792.17	-596.169
20	-1.017	2445	2830.20	-385.204
21	2.280	3287	2423.11	863.894
22	1.845	3179	2480.16	698.842
23	.703	2780	2513.68	266.320
24	1.173	2774	2329.63	444.368
25	.782	2575	2278.70	296.296

a. Dependent Variable: y

Residuals Statistics^a

	Minimum	Maximum	Mean	Std. Deviation	N
Predicted Value	117.83	2830.20	1178.80	1031.903	25
Residual	-747.271	863.894	.000	370.899	25
Std. Predicted Value	-1.028	1.600	.000	1.000	25
Std. Residual	-1.972	2.280	.000	.979	25

a. Dependent Variable: y

تفسير النتائج

يمكن كتابة معادلة الانحدار كما يلي :

$$Y = -30.437 + 0.332 X1$$

اختبار معنوية النموذج

يلاحظ من جدول تحليل التباين ANOVA أن قيمة $F = 178,03$ بمستوى معنوية أقل من $0,0005$ مما يدل على النموذج معنوي

اختبار معنوية المعاملات المقدرة

- قيمة T بالنسبة للثابت $-0,258$ بمستوى معنوية $0,799$ مما يدل على أن الثابت غير معنوي
- قيمة T بالنسبة للمتغير $X1 = 13,343$ بمستوى معنوية أقل من $0,0005$ مما يدل على أن المتغير $X1$ معنوي

القدرة التفسيرية للنموذج

بلغت قيمة $R^2 = 0,886$ مما يدل على أن المتغيرات المستقلة تفسر حوالي 87% من التغيرات التي تحدث في المتغير التابع .

اختبار وجود الارتباط الذاتي للأخطاء العشوائية

بلغت قيمة ديربن واتسون $0,865$ ، وبالكشف في جداول ديربن واتسون بمتغير واحد $K=1$ ودرجات حرية الخطأ 23 نجد أن

$$d_L = 1.018 \quad d_U = 1.187$$

وحيث أن $0 < DW < d_L$

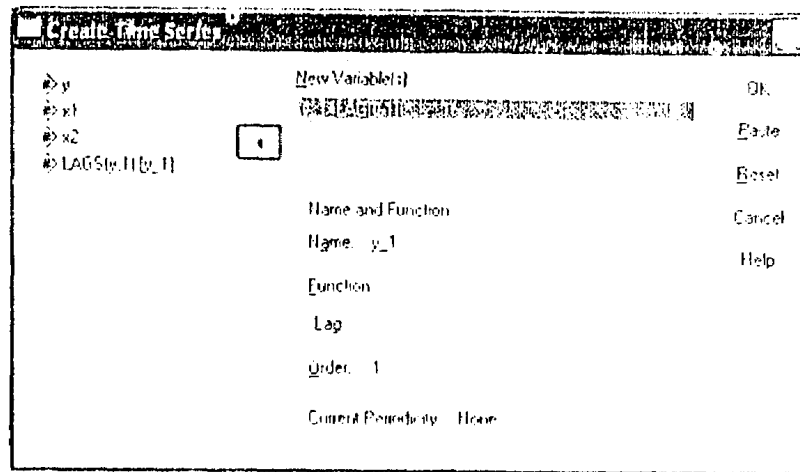
لذا يوجد ارتباط ذاتي ، ويمكن حسابه كما يلي
معامل الارتباط الذاتي

$$\rho = 1 - (DW/2) = 1 - (0.865/2) = 0.568$$

علاج مشكلة الارتباط الذاتي

يمكن علاج مشكلة الارتباط الذاتي عن طرق اختبار المتغير التابع بفترة إبطاء ، ضمن المتغيرات المستقلة ، ولتنفيذ ذلك تابع الخطوات التالية

١- من القائمة Transform اختر امر Create Time Series فيظهر الصندوق الحوارى Create Time Series



٢- أنقل المتغير Y إلى قائمة New Variable(s)

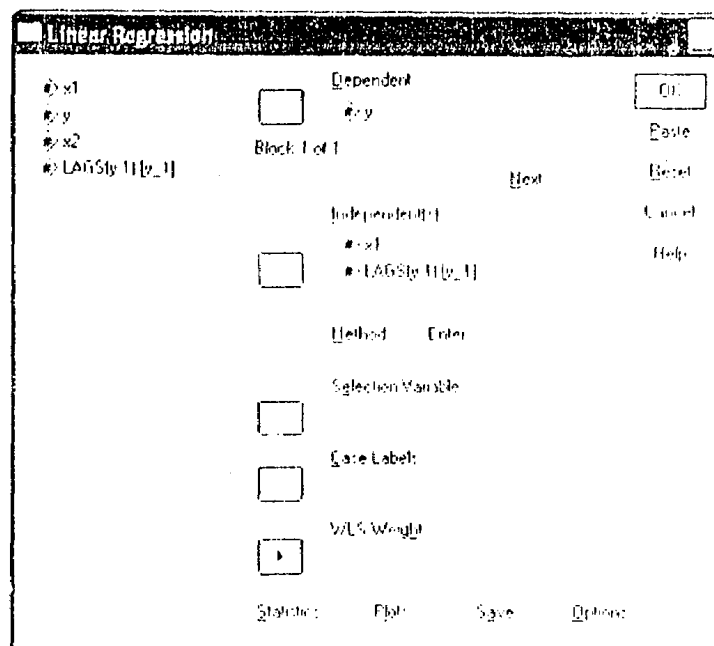
٣- من خانة Name يظهر اسم المتغير الجديد Y_1 ، ويمكن إعادة تسميته

٤- من خانة Function اختر Lag ، ومن خانة Order اكتب فترة الإبطاء ١

٥- اختر أمر Ok يتم إضافة المتغير الجديد Y_1 إلى المتغيرات في نافذة Data Editor

٦- نفذ الانحدار المتعدد بالطريقة السابق ذكرها بعد اختيار المتغير التابع $Y = \text{Dependent}$ ، والمتغيرات المستقلة

X1 , Y_1



٧- بعد إظهار نتائج الانحدار يمكن النظر إلى إحصائية DW حيث نجدها = ٢,٢٣٦

Model Summary^c

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Durbin-Watson
1	.974 ^a	.950	.947	252.710	
2	.985 ^b	.971	.968	197.601	2.236

a. Predictors: (Constant), y_1 LAGS(y,1)

b. Predictors: (Constant), y_1 LAGS(y,1), x1

c. Dependent Variable: y

ومنه نجد أن $2 < dw < (4-du)$

أى $2 < 2.236 < 2.723$ مما يدل على عدم وجود ارتباط ذاتي

٣-٥ وجود أخطاء القياس والمشاهدة فى المتغيرات التفسيرية.

تشتط طريقة المربعات الصغرى فى التقدير عدم وجود أخطاء مشاهدة أو قياس بالمتغيرات التفسيرية ، وهو شرط يضمن إستقلال المتغيرات التفسيرية عن عنصر الخطأ العشوائى . وترجع أخطاء القياس إلى أن المتغيرات المستخدمة غالبا ما تكون تجميعية أو أن تستخدم متغيرات ليست هى المتغيرات المفروض إدخالها فى النموذج كاستخدام الناتج القومى الإجمالى كمتغير محل الدخل الممكن التصرف فيه عند تقدير دالة الاستهلاك ، واستخدام الأرقام القياسية للأسعار لتحويل القيم النقدية إلى قيم حقيقية ، ومعروف أن هذه الأرقام عرضه للأخطاء . ولا يوجد اختبارات معينة للكشف عن وجود أخطاء فى المتغيرات ولكن يمكن أن تعطى النظرية الاقتصادية أو المعرفة بالطريقة التى جمعت بها البيانات إشارة إلى مدى خطورة المشكلة .

ويترتب على وجود هذه الأخطاء واستخدام أسلوب المربعات الصغرى العادية فى التقدير تحيز التقديرات.

علاج أخطاء القياس

يمكن علاج هذه الأخطاء بأسلوبين

١- أسلوب الانحدار المرجح Weighted Regression

٢- المتغيرات المساعدة

يلاحظ شيوع عدم ثبات التباين فى البيانات المقطعية ففي إلتفاق الاستهلاكى مثلا نجد أن تقلباته فى فئات الدخل العليا أعلى من تقلبات فئات الدخل المنخفضة .

ويؤدى عدم ثبات التباين إلى عدم صلاحية استخدام تباينات المعالم المقدرة بأسلوب المربعات الصغرى لإجراء اختبارات المعنوية الإحصائية وتحديد فترات الثقة لهذه التقديرات ، وذلك بالرغم من عدم تحيزها .

ويمكن اختبار عدم ثبات التباين بمعامل ارتباط الرتب سبيرمان وذلك عن طريق تقدير علاقة الانحدار بالطريقة العادية وإيجاد قيم البواقي e التي هي تقدير للأخطاء u ثم ترتب مع قيم e ترتيباً تصاعدياً أو تنازلياً (مع إهمال الإشارة) معامل ارتباط الرتب من المعادلة

$$r_{c,n} = 1 - \frac{6 \sum D_i^2}{N(N^2 - 1)}$$

حيث D_i هي الفرق بين رتب القيم المتناظرة من e , x كلما كبرت قيمة r كلما دل ذلك على عدم ثبات التباين لعنصر الخطأ .

علاج عدم ثبات التباين

يمكن معالجة مشكلة عدم ثبات التباين بتعديل النموذج بحيث يصبح عنصر الخطأ له تباين ثابت . فالعلاقة

$$y_i = b_0 + b_1 x_i + u_i$$

يمكن أن تكون

$$y_i / x_i = b_1 + b_0 (1/x_i) + (u_i / x_i)$$

وهذه العلاقة تسمح بأن يكون عنصر الخطأ العشوائي له تباين ثابت حيث

$$\text{Var}(u_i) = c x_i^2 \quad \text{إذا كان}$$

$$v\left(\frac{u_i}{x_i}\right) = \frac{c}{x_i^2} x_i^2 = c$$

$$\frac{u_i}{x_i} = u^* \quad \text{فان تباين } u^* \text{ يكون ثابت}$$

وبالتالى يمكن استخدام المربعات الصغرى العادية لتقدير معالم الدالة والتي سوف تنسم بعدم التمييز وأصغر تباين .

٢- قد تكون هناك معلومات عن قيمة تباين الخطأ من معلومات سابقة كأن يكون

$$\text{Var}(e_i) = \sigma_i^2$$

في هذه الحالة يمكن استخدام هذا التباين في تعديل علاقة الانحدار وتكون على الصورة

$$y_i / \sigma_i = b_0 + \sigma_i + b_1 x_i / \sigma_i + e_i / \sigma_i$$

وتقدر المعالم بطريقة المربعات الصغرى العادية.

أولا : المراجع العربية

- ١- دومينيك سالفاتور (ترجمة د. سعدية منتصر) الإحصاء والاقتصاد القياسي سلسلة ملخصات شوم، دار ماكجرو هيل للنشر - ١٩٨٢.
- ٢- ربيع ذكى عامر، تحليل الانحدار أساليبه وتطبيقاته العملية باستخدام البرنامج الجاهز SPSS/pc - جامعة الكويت ١٩٨٩م.
- ٣- رجاء محمد أبو علام ، التحليل الإحصائي للبيانات باستخدام برنامج SPSS القاهرة - دار النشر للجامعات، ٢٠٠٣م.
- ٤- سمير كامل عاشور ، سامية أبو الفتوح ، العرض والتحليل الإحصائي باستخدام SPSS/pc + معهد الدراسات والبحوث الإحصائية ١٩٩٣م.
- ٥- سمير كامل عاشور ، سامية أبو الفتوح ، العرض والتحليل الإحصائي باستخدام SPSSwin الجزء الأول المدخل والأساسيات-معهد الدراسات والبحوث الإحصائية ٢٠٠٣م.
- ٦- سعد زغلزل بشر، دليلك إلى البرنامج الإحصائي SPSS الإصدار العاشر version 10 - المعهد العربي للتدريب والبحوث الإحصائية بغداد ٢٠٠٣م.
- ٧- عبد الحميد عبد اللطيف : استخدام الحاسب الآلى فى مجال العلوم الاجتماعية (استخدام برنامج SPSS من خلال windows) جامعة عين شمس ٢٠٠١م.
- ٨- عبدالله بن عمر النجار، استخدام حزمة البرامج الإحصائية (SPSS) فى تحليل البيانات مؤسسة شبكة البيانات - الرياض - ٢٠٠٣م.
- ٩- عبد القادر محمد عبد القادر : طرق قياس العلاقات الاقتصادية ، مع تطبيقات الحاسب الالكترونى، ١٩٩٠.
- ١٠- عبد الرحيم عبد الحميد الساعاتى ، مدحت فهمى صالح ، مبادئ التحليل الإحصائي للاقتصاد والإدارة باستخدام برنامج EXCEL ، ٢٠٠٤ ، دار جدة .
- ١١- لنكون شاو- تعريب د.عبد المنعم حامد عزام- الإحصاء فى الإدارة - دار المريخ للنشر ١٩٩٦م.
- ١٢- محمد الدسوقي حبيب ، جلال السيد ، مقدمة فى الطرق الإحصائية، ١٩٩٧.

١٣- محمد علي حسن. عناف عبد الجبار سعيد . الاقتصاد القياسي النظرية والتطبيق - دار وائل للنشر والتوزيع عمان الأردن ١٩٩٨م.

ثانيا : المراجع باللغة الأجنبية

- 1- Damodar Gujarati, Basic Econometrics, 3rd ed., Mc Graw-Hall, Inc., 1995.
- 2- Sheridan J.coa Kes, Lyndall G.steed, SPSS Analysis without Anguish, version 10.0 for windows, 2000.