

سلسلة الإحصاء والبحث للعلوم السلوكية والنفسية

باستخدام

حزم البرامج الإحصائية

الجزء الثالث

تحليل البيانات المتعددة والمتدرجة في ضوء حجم التأثير

"SPSS & LISREL"

**Multiple and Multivariate Data Analysis in light of
Effect size**

عبد الناصر السيد عامر

كلية التربية - جامعة قناة السويس

2020م

فهرس المحتويات

فهرس المحتويات

مقدمة الجزء الثالث

الفصل الأول: مفاهيم أساسية

الفصل الثاني: إعداد البيانات للتحليل

الفصل الثالث: تحليل التباين أحادى الاتجاه داخل الأفراد (تصميم القياسات المتكررة)

الفصل الرابع: التحليل العاملى للتباين أو تحليل التباين ذى اتجاهين بين المجموعات

الفصل الخامس: تحليل التباير (ANCOVA)

الفصل السادس: تحليل التباين المتدرج (MANOVA)

الفصل السابع: الانحدار المتعدد

الفصل الثامن: التحليل العاملى الاستكشافى

الفصل التاسع: التحليل العاملى التوكيدى

الفصل العاشر: نمذجة المعادلة البنائية

الفصل الحادي عشر: الاختبارات والنماذج الاحصائية في البحث النفسى والتربوي

المصري والعربي

المراجع

المقدمة

تتم المعالجة الإحصائية للظواهر الإنسانية والاجتماعية بصورة ذرية جزئية نظرًا لعدم إدراك وإلمام الكثير من الباحثين بالأساليب الإحصائية المتدرجة المتقدمة التي تناسبها حيث تناولت دراسة هذه الظواهر في ضوء معرفة الباحثين للمعالجات الإحصائية البسيطة وليس ما يجب أن يكون. والظاهرة الإنسانية والاجتماعية معقدة ومتشابكة ومتفاعلة ولدراستها تحتاج إلى أساليب إحصائية تراعى هذا التفاعل والتعقيد، من ثم فالأساليب الإحصائية الأحادية أو الثنائية البسيطة غير ملائمة بالدرجة الكافية لدراساتها وعليه كانت الحاجة لإصدار هذا الكتاب الذى تناول أساليب إحصائية متقدمة تتعامل مع متغيرات مستقلة عديدة ومتفاعلة ومتغيرات تابعة عديدة ومتفاعلة وهذا يعكس الواقع الحقيقى لدينامية السلوك الإنسانى.

وفى هذا الكتاب عرضت الأسلوب الإحصائى فى ضوء معالجة منهجية وليست إحصائية فقط وكيفية تحليل بيانات هذه الأساليب فى برنامج SPSS من إدخال وتنفيذ وتفسير نتائج المخرجات وكتابة نتائج الاختبارات فى تقرير البحث وفقًا للجمعية النفسية الأمريكية (APA)، كما حاولت بقدر الإمكان البعد عن التعقيدات الحسابية ويعتبر هذا الكتاب دليلًا عمليًا للباحثين فى كافة التخصصات وتناول ان الكتاب اساليب متدرجة فى غاية الاهمية مثل تحليل الانحدار المتعدد وتحليل التباين المتدرج و التحليل العاملي الاستكشافي والتحليل العاملي التوكيدى ونمذجة المعادلة البنائية، وهي اساليب لا يميل الباحث النفسي والسلوكي والتربوي في البيئة العربية الي توظيفها في دراساته حيث انها تمثل 5% في البحث النفسي التربوي المصري و 1.1% من اجمالي الاساليب الاحصائية في البحث النفسي التربوي العربي(عامر، 2012).

والهدف الأساسى من هذا الكتاب هى مساعدة الباحثين على إجراء هذه الاختبارات بدون قلق أو خوف حتى يتمكنوا من استخدامها فى دراساتهم.

وتضمن هذا الكتاب الموضوعات الآتية:

الفصل الأول تناول مفاهيم أساسية ضرورية لدراسة الأساليب الإحصائية المتدرجة والمتعددة مثل المتغيرات والمنهجية البحثية خاصة التصميمات التجريبية بين المجموعات (المستقلة) وداخل المجموعات (المرتبط)، والتصميمات العاملية. **الفصل الثاني** تضمن إعداد البيانات للتحليل متضمنًا المسلمات الواجب توافرها للبيانات مثل حجم العينة المناسب، الاعتدالية، استراتيجية التعامل مع البيانات المفقودة. **الفصل الثالث** تضمن تحليل التباين داخل المجموعات (المرتبطة) متضمنًا اختبارات فروض لقضية بحثية وكيفية تنفيذه في برنامج SPSS وتفسير مخرجاته وحساب القوة الإحصائية. **الفصل الرابع** تضمن تحليل التباين المتعدد مفهومه ومناسبته للتصميمات العاملية ومصادر التباين وعرض قضية بحثية وتنفيذها في SPSS وتقدير حجم التأثير والقوة الإحصائية. **الفصل الخامس** تضمن تحليل التباين (ANCOVA) بنية تحليل التباين باعتباره أسلوب للضبط الإحصائي للمتغيرات الدخيلة التي تقاس في التصميمات التجريبية وعرضت قضية بحثية وكيفية صناعة القرار باستخدام خطوات اختبارات الفروض الصفرية وتحليل بياناته في SPSS وتقدير حجم التأثير والقوة الإحصائية. **الفصل السادس** تحليل التباين المتدرج (MANOVA) باعتباره من انسب الأساليب الإحصائية لدراسة الظواهر النفسية حيث يتعامل مع متغير مستقل واحد فأكثر وأكثر من متغير تابع ويبين التفاعلات بين المتغيرات المستقلة على التفاعل بين المتغيرات التابعة وتناولت قضية بحثية واقعية وكيفية صياغة الأسئلة والفروض البحثية والإحصائية، وتحليل البيانات في SPSS وتفسير مخرجاته. **الفصل السابع** الانحدار المتعدد ومسلماته وطرائق تحليله، ومثالًا تطبيقًا للانحدار المتعدد وتفسير معالمه برنامج SPSS وكيفية تقدير القوة الإحصائية. **الفصل الثامن** التحليل العائلي الاستكشافي (EFA) مفهومه، مسلماته، طرائق الاستخلاص، طرائق التدوير، وكيفية تحديد عدد العوامل وتفسيرها مع إعطاء مثال لكيفية تنفيذه في برنامج SPSS. **الفصل التاسع** التحليل العامل التوكيدي، مفهومه، مميزاته، معالمه، وخطوات تنفيذه في برنامج LISREL. **الفصل العاشر** نمذجة المعادلة البنائية مفهومها، مميزاتها، معالمها، خطوات بنائها، وتنفيذها في برنامج LISREL .

وأخيراً أتمنى أن أكون قدمت عملاً مفيداً للباحثين لمساعدتهم على التعامل مع القضايا والظواهر والمشكلات النفسية والاجتماعية بما يناسبها من أساليب إحصائية متقدمة تراعى التفاعلات البينية بين مسببات حدوثها لتمدنا برؤية وتفسيراً أكثر عمقاً وشمولية لفهم هذه الظواهر التي حاولت الدراسات فى البيئة العربية التعامل معها بصورة جافة وسطحية لا تحاكي واقعها الفعلى مما جعل نتائجها غير قابلة للتطبيق للحياة الواقعية الفعلية.

والله الموفق والمستعان

أ.د. عبد الناصر السيد عامر

استاذ القياس والتقويم والإحصاء النفسى

كلية التربية- جامعة قناة السويس

adr.abdenasser@yahoo.com

الفصل الأول

مفاهيم أساسية

Basic Concepts

الظاهرة الإنسانية الاجتماعية من الظواهر المعقدة التى تتعدد فيها المتغيرات المكونة لها بل تتفاعل هذه المتغيرات والعوامل بحيث تشكلها وتكونها مما يصعب فصل هذه المتغيرات عن بعضها البعض، ويبدو أن العمل الإحصائى الذى يتعامل بها معها باعتبارها مكونات أو متغيرات منفصلة عن بعضها البعض غير مناسب لدراستها من ثم تحتاج دراسة هذه الظواهر درجة من التعقيد فى العمل الإحصائى ونماذج إحصائية أكثر تعقيداً من كونها نماذج إحصائية بسيطة تتعامل معها فى ضوء ثنائية من المتغيرات. وفى هذا الكتاب حاولت عرض بعض النماذج الإحصائية الأكثر ملائمة واستخداماً فى دراسة الظواهر النفسية والتربوية والاجتماعية.

الإحصاء هى مجموعة من الطرق والأساليب والاختبارات يطبقها الباحث لتحليل بياناته فى محاولة لفهمها وتفسيرها، وهى مجموعة من النظريات والإجراءات والطرق تطبق بهدف فهم البيانات، ويرى (Hinkle, Wiersma, & Jurs (1994 أن العلوم الاجتماعية والسلوكية ما كنت أن توجد بهذا الوضع الان بدون الإحصاء. ويرى (Aron, Coups, & Aron (2013 أن الإحصاء أحد فروع الرياضيات التى تهتم بتنظيم وتحليل وتفسير البيانات. ويشير (Duun (2001 إلى ان الإحصاء تركز على الطرق المناسبة لجمع وتبويب وتحليل وتفسير البيانات وهذه الطرق المستخدمة لفحص البيانات توضع فى مجموعة من القواعد والإجراءات يطلق عليها الإحصاء.

فرعى الإحصاء

ويوجد فرعين أساسيين من الأساليب الإحصائية حيث تصنف دراسة الإحصاء التصنيفى حسب الهدف من إجراء الدراسة هما:

الإحصاء الوصفى: يستخدم الإحصاء الوصفى عندما يكون هدفه وصف مجموعة من البيانات فقط وتستخدم لتصنيف وتلخيص البيانات فى صورة مبسطة كتقدير المتوسط أو الانحراف المعياري للبيانات. فعلى سبيل المثال المدرس فى الفصل يعد

اختبار تحصيلي لقياس مهارات اللغة العربية ثم يطبقه ويصححه ويعطى درجات للطلاب بعد ذلك يقوم بعرض بياني لدرجات الطلاب وتقدير متوسط التحصيل وتحديد نسبة النجاح والرسوب في الفصل.

والإحصاء الوصفي مجموعة من الإجراءات تهدف إلى وصف وتنظيم و تلخيص البيانات بهدف تيسير فهمها والوصول لانطباعات مبدئية عنها، بينما التعمق في فحص البيانات قبل استخدام الأساليب الإحصائية المتقدمة فإن هذه العملية تسمى تحليل البيانات الاستكشافي (Exploratory data analysis (EDA، و من اهم الأساليب الإحصائية التي تقع تحت الإحصاء الوصفي هي مقاييس النزعة المركزية ومقاييس التشتت والرسومات البيانية وغيرها.

الإحصاء الاستدلالي: بعد وصف البيانات وأصبحنا على وعى بماذا تقول على مستوى وصفي سطحي نهتم بعد ذلك بالإحصاء الاستدلالي وهي تتكون من مجموعة من الإجراءات لعمل تعميمات Generalization عن الحالة في المجتمع عن طريق دراسة العينات (مجموعات فرعية) التي يتم سحبها من هذه المجتمع، فمثلاً في تصميم التجارب نحن ندرك أنه من المستحيل تطبيق التجربة على المجتمع ككل ولظروف خاصة بعملية الضبط والتكلفة والوقت يتم سحب عينات من المجتمع الاصلى Population.

وكما نعلم أن العمل الإحصائي يعتمد على بيانات واهم شئ يتوفر في هذه البيانات هي جودتها، بمعنى ثباتها يكون مرتفع والتعامل مع بيانات ثباتها أقل من 0.70 دائماً يؤدي إلى نتائج مشوهة وذلك لان القياسات في مجملها كانت لأخطاء وليس للظاهرة الحقيقية المراد قياسها. والأساليب الإحصائية ادوات تكون مفيدة للباحثين الذين يحسنون استخدامها فالبداية بالسؤال وتصميم الدراسة الجيد. وفي تحليل البيانات فالإحصاء هي اداة مساعدة للوصول إلى استنتاجات واجابات سليمة بقدر الإمكان، وعلى ذلك فإن الإحصاء لا تكذب انما سوء استخدامها يؤدي إلى قرارات خاطئة.

التصميم البحثي أو منهج البحث Research Design

في العلوم السلوكية والاجتماعية لا بد من استخدام منهج البحث أو طريقة البحث Research method. وفيها يحدد الباحث طبيعة المنهج أو المدخل العام

الذى يتبعه للتحقق من فروض بحثه، وهذا يتحدد من خلال طبيعة أسئلة البحث المطروحة وبدوره منهج البحث يحدد الإجراءات والخطوات التى يتبعها الباحث من اختيار العينة وتطبيق أدوات القياس والإجراءات المتبعة لجمع البيانات. وتوجد العديد من طرائق البحث ولكننا فى هذا المقام نركز على اهم طرق البحث الشائعة فى البحث السلوكى والنفسى والاجتماعى لأن استخدام هذه الطرق يسمح بالحصول على بيانات وإجراء تحليلات إحصائية معينة وهى:

أولاً : المنهج التجريبى Experimental method : يهدف إلى محاولة الوصول إلى التحقق من وجود علاقة السبب والتأثير Cause-effect بين متغيرين من خلال تجربة أحدهما يطلق عليه المتغير المستقل أو المتغير التجريبى أو متغير المعالجة هو يخضع لسيطرة وتحكم الباحث ويجب أن يمتلك على الأقل مستويين وتسمى شروط أو عوامل Factors، أحدهما شرط المعالجة Treatment والآخر شرط الضبط Control، بينما المتغير الآخر يسمى بالمتغير التابع أو المتغير المحك Criterion أو المتغير الناتج Outcome ويشار إليه بنتائج أو مخرجات التجربة مثل التحصيل والابتكارية. ولدراسة السببية من خلال تجربة فلا بد أن تتضمن إجراءات صارمة منها:

- المعالجة لمتغير تجريبى مراد دراسته Manipulation.
- الضبط لكل المتغيرات الدخيلة التى يمكن أن تؤثر فى المتغير التابع.
- العشوائية Randomization فى اختيار عينة البحث وكذلك عشوائية فى توزيع المشاركين على مجموعات التجربة بطريقة تحدث درجة كبيرة من التكافؤ Equivalence أو التجانس.
- المقارنة بين درجات المعالجات المختلفة وهى المجموعات التجريبية (التى تتلقى المعالجة) والمجموعة الضابطة (لا تتلقى المعالجة أو تتلقى مستوى آخر من المعالجة)، فالفرق بين نواتج المعالجات دليل على أن المعالجة (المتغير المستقل) هى المسبب للتغيرات فى الدرجات المتغير التابع.
- وعلى ذلك فأى تجربة تتضمن أربع مكونات أساسية لاستنتاج السببية هى المعالجة، العشوائية، الضبط، والمقارنة. فمثلاً إستراتيجية التدريس هى (المعالجة) أو المتغير

المستقل (تقليدية-عصف ذهني) والتحصيل تابع أو ناتج، العلاج السلوكي(المعالجة) بينما خفض الاكتئاب التابع (ناتج)، إستراتيجية طرح التساؤل الذاتي (المعالجة) بينما الفهم القرائي التابع (الناتج).

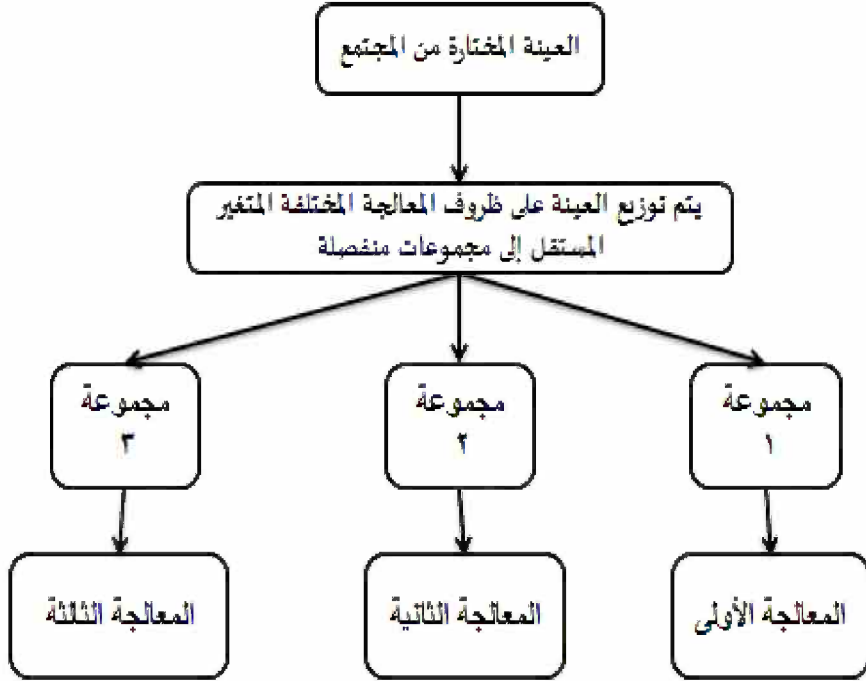
وعليه فإن معالجة مستويات المتغير المستقل (عالي- منخفض أو معالجة-لامعالجة) وتوزيع أفراد العينة عشوائيًا على مجموعات أو مستويات المعالجة وضبط المتغيرات الدخيلة هي سمات البحث التجريبي. والمنهج التجريبي يتضمن تصميمات عديدة منها(Shadish & Cook, 2002):

1. تصميم بين الأفراد أو المجموعات **Between-subjects design**: السمة

الأساسية لهذه التصميمات هو المقارنة بين مجموعات مختلفة من الأفراد وتوزيعهم عشوائيًا على مستويات ومعالجات المتغير المستقل وبالتالي توجد مجموعة أو مجموعات تجريبية **Experimental group(s)** ومجموعة ضابطة **Control group** ويطلق عليه التصميم التجريبي للقياسات المستقلة **Independent measures** **experimental Design** وفيه تتم المقارنة بين المجموعات و تحديد الفروق بينهما.

خصائص تصميمات بين الأفراد أو المجموعات

السمة الأساسية لهذا التصميم هي المقارنة بين مجموعات مختلفة من الأفراد. وفي هذا النوع من التصميمات فالباحث يعالج المتغير المستقل لخلق ظروف معالجة مختلفة ومجموعات مختلفة من المشاركين يتم توزيعهم على الظروف المختلفة من المتغير المستقل، ثم يقاس المتغير التابع على كل فرد يقوم الباحث بفحص البيانات للبحث عن وجود فروق بين المجموعات.



الشكل (1.1): بنية تصميم بين المجموعات أو الأفراد.

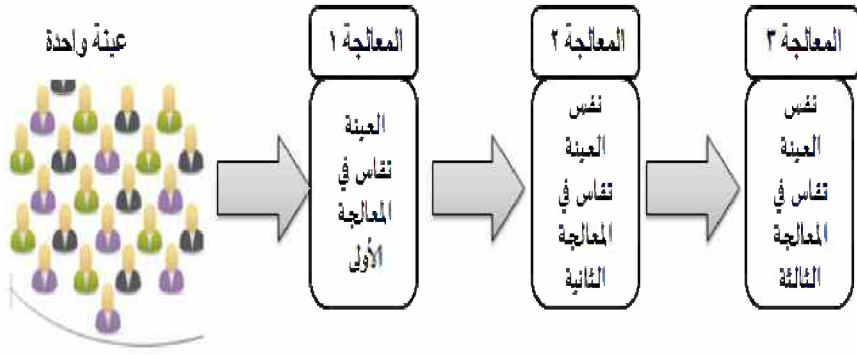
والهدف من هذه النوعية من التصميمات تحديد ما إذا كانت توجد فروق بين مجموعتين أو معالجتين فأكثر. مثلاً أراد الباحث المقارنة بين طريقتين للعلاج (معالجتين) لتحديد أيهما أكثر فعالية من الأخرى فى هذه الحالة يقوم بأخذ مجموعتين منفصلتين من الأفراد؛ كل مجموعة تأخذ أو تتلقى معالجة. ويجب التأكيد أن تصميم بين المجموعات يستخدم بصورة عامة فى استراتيجيات بحثية أخرى مثل التصميمات غير التجريبية (السببى المقارن) والتصميمات شبه التجريبية.

وتُعد الدرجات المستقلة Independent scores هى الخاصية المميزة لهذا التصميم وتستحق التنويه عنها وهى أن درجة الفرد تعكس مشارك منفصل فلو أن تصميم بين المجموعات أعطى 30 درجة فى المعالجة الأولى A و 30 درجة فى المعالجة الثانية B بالتالى فهذا يعنى أن أفراد المعالجة الأولى مستقلين أفراد المعالجة الثانية، ويكون العدد الكلى من المشاركين 60 فرد وهذا يعنى أن هذه النوعية من التصميم تستخدم مجموعة مختلفة من الأفراد فى كل معالجة. وكل مجموعة تتعرض لمعالجة واحدة فقط من معالجات المتغير المستقل، ويمكن للباحث فى هذا التصميم استخدام مجموعة من

القياسات ودمجها في درجة واحدة إذا كان المتغير التابع واحد ويمكن استخدام أكثر من متغير تابع. فإذا كان المتغير التابع مثالاً هو وقت ما مستغرق، فإن الباحث يقيس المتغير التابع عدة مرات. ثم يأخذ متوسط القياسات لإعطاء درجة وحيدة أكثر ثباتاً واستقراراً ويكون لكل فرد درجة واحدة. وهذا التصميم يجب أن يلتزم بالمبادئ الأساسية لأي تجربة مثل المعالجة للمتغير المستقل وضبط المتغيرات الخارجية.

2. تصميم داخل الأفراد أو تصميم القياسات المتكررة Repeated measures designs within-subject: في هذا التصميم تستخدم مجموعة واحدة من الأفراد ويطبق عليها كل معالجات أو مستويات المتغير المستقل ثم قياس أو ملاحظة أداء كل فرد في كل المعالجات ويطلق عليه تصميم القياسات المتكررة لأنه يحدث إعادة للقياسات على نفس الأفراد وتحت ظروف معالجة مختلفة وهذا التصميم سائد أيضاً في الدراسات غير التجريبية مثل الدراسات الطولية التي تدرس التغيرات عبر فترة زمنية طويلة.

والسمة الأساسية لهذا التصميم أنه يستخدم مجموعة وحيدة من الأفراد وقياس وملاحظة كل فرد في كل المعالجات المراد المقارنة بينهما. وعلى ذلك لا تقسم العينة إلى مجموعات متعددة ولكن تستخدم نفس المجموعة من الأفراد في كل مستوى من مستويات المتغير المستقل ولذلك تتعرض المجموعة إلى كل المعالجات المختلفة للمتغير المستقل، وبلغة إحصائية يطلق على هذا التصميم بتصميم القياسات المتكررة لأن الدراسة تعيد تكرار القياس على نفس الأفراد تحت الأفراد تحت شروط المعالجات المختلفة، وعلى ذلك فالمقارنة لا تكون بين المجموعات المختلفة بل تكون بين المعالجات المختلفة على نفس الأفراد. وهذا التصميم يكون شائع في الدراسات غير التجريبية مثل الدراسات الطولية التي تفحص تأثيرات التغيرات عبر الزمن. وهذا الوضع شائع في مجال علم النفس النمو حيث يدرس ظاهرة ما على مجموعة واحدة من الأفراد وعبر أعمار مختلفة لتحديد درجة النمو عبر الزمن:



الشكل (2.1): بنية التصميم التجريبي داخل الأفراد (المجموعات).

ويهدف هذا النوع من التصميم بالبحث عن الفروق بين المعالجات المختلفة على نفس الأفراد وهذا التصميم يلتزم بكل متطلبات المنهج التجريبي من معالجة المتغير المستقل وضبط المتغيرات الدخيلة.

وهذه النوعية من التصميمات نادرة الاستخدام في البحوث النفسية والتربوية والسلوكية، والمعالجة الإحصائية لهذا النوع من خلال تحليل التباين للقياسات المتكررة Repeated ANOVA أو الاختبار اللابارامترى المقابل وهو اختبار فريدمان .Friedman

مميزات تصميم داخل المجموعات

1. يتطلب عدد محدد من الأفراد (حجم عينة صغير) مقارنة بتصميم بين المجموعات، فعند مقارنة ثلاثة معالجات مختلفة، وكل مجموعة تضم 30 فرد فهذا يستلزم 90 فرد بينما في تصميم داخل المجموعات يستلزم 30 فرد فقط في كل المعالجات وهذا مناسب في المواقف التي يصعب فيها الحصول على عينة مثل مجال الإعاقات أو مجال دراسة التوائم.

ومن مميزات هذا التصميم استبعاد كل المشكلات الناتجة عن الفروق الفردية هي التي تحدث في تصميم بين المجموعات وهي أن تصبح الفروق الفردية متغير دخيل وتخلق تباين عالي داخل المجموعة وبدوره يخفي أي فروق بين المعالجات وعلى ذلك فالتصميم داخل المجموعات يستبعد هذه المشكلات مثل التباين والمتغير الدخيل. وفي

تصميم بين المجموعات فإن كل درجة تعكس شخص مختلف بينما فى تصميم داخل المجموعات فإن نفس الفرد يقاس عبر المعالجات.

وعيوب تصميم داخل المجموعات أنه يطبق على الفرد يتم تطبيقه عليه سلسلة من المعالجات ويتم إدارته وتطبيق كل معالجة فى أوقات مختلفة وهذا ينشأ عنه فرصة وجود عوامل مرتبطة بالوقت أو الزمن مثل التعب أو الطقس وهذا يؤثر فى درجات الأفراد. فإذا حدث تغير فى الأداء أو ضمور أو انخفاض فلا نستطيع تحديد هذا التغير ما إذا كان سببية المعالجات المختلفة أو مؤشر لحدوث حالات التعب والانهك وهذا تغير تهديداً للصدق الداخلى للتجربة وعلى ذلك يوجد تفسير يوجد تفسير بديل للنتائج.

والمشكلة الأخرى فى هذا التصميم تأكل العينة حيث تبدأ مجموعة من الأفراد أو الدراسة البحثية وقد تتسحب من الدراسة قبل اكتمالها فيحدث فقد بين القياس الأول والقياس النهائى وفى هذه المشكلة تكون خطيرة عندما تمتد الدراسة لفترة زمنية طويلة. وربما ينسى الأفراد مواعيد المعالجات ويفتر الاهتمام لديهم وينسحبوا أو ربما يموتوا، وهذا بدوره يحدث تقليص لحجم العينة ولذلك ينصح بالبداية فى الدراسة بعدد كبير من الأفراد أكثر مما هو مطلوب بالفعل.

التحليلات الإحصائية لتصميمات داخل المجموعات (الأفراد)

فى هذا التصميم يتم تقدير المتوسط وتقييم فروق المتوسطات بين ظروف المعالجات، وفى هذا التصميم يتم خلق معالجتين فأكثر وملاحظتها على نفس الأفراد وفى هذا التصميم تزداد احتمالية الحصول على فروق دالة إحصائية. ومع البيانات الفترية والنسبية فالإجراء الإحصائى هو حساب متوسط درجات كل معالجة لوصف المعالجات على حدة ويستخدم اختبار T للقياسات المتكررة فى حالة استخدام معالجتين لتقويم الدلالة الإحصائية لفروق المتوسطات، أو إذا كانت البيانات من مستوى القياس الترتيبى تستخدم الاختبار اللابارامترى وهو ويلكوكسون Wilcoxon test.

وفى حالة تصميم المعالجات المتعددة Multiple treatment design حيث يوجد أكثر من معالجتين. وتوجد مشكلة أساسية فى هذا التصميم وهو أن التمييز بين

المعالجات ربما يكون صغيراً للوصول إلى فروق دال إحصائياً. ويحتاج هذا التصميم إلى وقت كبير حتى يكمل أفراد العينة انجاز المعالجات الثلاثة. وفي هذا التصميم يتم حساب متوسط كل معالجة ولتقويم الفروق بين المتوسطات خاصة للبيانات الفترية والنسبية يستخدم تحليل التباين للقياسات المتكررة Repeated measures ANOVA. وإذا كانت البيانات ترتيبيه يستخدم الاختبار الإحصائي اللابارامترى اختبار فريدمان.

مقارنة بين تصميمات داخل المجموعات وتصميمات بين المجموعات

توجد ثلاثة عوامل تميز هذه الفروق بين التصميم وهي كالآتي:

أ. الفروق الفردية: تعتبر العيب الرئيسى ومن محددات تصميم بين الأفراد لأنها متغيرات دخيلة تزيد التباين وهذا يمكن تجنبه فى تصميم داخل المجموعات. ولأن تصميم داخل المجموعات يقلل التباين بالتالى فهو أكثر احتمالاً وقدرة للكشف عن تأثير المعالجة الحقيقى مقارنة بتصميم بين المجموعات. وبالتالي لو أن المشاركين فى التجربة يوجد بينهما فروق فردية كبيرة فمن الأفضل استخدام تصميم داخل المجموعات.

ب. العوامل المرتبطة بالزمن وتأثيرات الترتيب: يوجد احتمالية أن العوامل المرتبطة بالزمن تشوه نتائج تصميم داخل المجموعات وهذه الإشكالية يمكن تجنبها فى تصميم بين المجموعات حيث يوجد كل فرد داخل مجموعة واحدة (أو معالجة واحدة) ويقاس مرة واحدة. ولذلك عندما نتوقع أحد المعالجات لها تأثير عالى وطويل الأمد فمن الأفضل استخدام تصميم بين المجموعات.

ج. المشاركون: يتطلب تصميم داخل المجموعات حجم عينة محدود وعلى الجانب الآخر فإن تصميم بين المجموعات يحتاج إلى حجم كافى من المشاركين للحصول على البيانات وبالتالي عندما توجد صعوبة فى ايجاد عدد كافى من المشاركين فى التجربة فيفضل استخدام تصميم داخل المجموعات واختبار أى تصميم يتأثر بنوعية الأسئلة التى يطرحها الباحث.

التصميمات العاملية أو التصميمات المعقدة Factorial or Complex designs

فى الدراسات النفسية دائماً يتم التركيز على التصميمات التجريبية الأساسية (أحادية العامل) حيث يهدف الباحث إلى دراسة تأثير متغير مستقل واحد بأكثر من مستوى وكيف يتم تنفيذ المتغير المستقل على مجموعات منفصلة من الأفراد فى كل معالجة (تصميمات المجموعات المستقلة) أو يتعرض مجموعة من الأفراد لكل ظروف أو مستويات المعالجة (تصميمات القياسات المتكررة) وعلى ذلك فإن هذه التصميمات تتناول معالجة متغير مستقل وحيد وهذا النوع من التصميمات هو الأكثر شيوعاً فى مجال علم النفس والتربية. ولكن أحياناً يستخدم الباحثون التصميمات المعقدة حيث يتناول معالجة متغيرين مستقلين فأكثر متزامنين فى نفس التجربة.

والتصميمات المعقدة يطلق عليها بالتصميمات العاملية لأنها تتضمن خليط عاملى من المتغيرات المستقلة. وهذا الخليط العاملى يتضمن مزوجة كل مستوى من أحد المتغيرات المستقلة مع كل مستوى من المتغير المستقل الآخر وهذا يجعلنا نحدد تأثير كل متغير مستقل على حدة (تأثير رئيسى Main effect) وتأثير المتغيرات المستقلة معاً أو متفاعلة (تأثير التفاعل Interaction effect) ولذلك فعندما تتضمن الدراسة متغيرين مستقلين فأكثر فى دراسة واحدة فإن المتغيرات المستقلة تسمى عوامل Factors. والدراسة التى تتضمن عاملين فأكثر تسمى تصميم عاملى. وهذا النوع من التصميمات يشار إليها بعدد العوامل فى الدراسة مثل تصميم ذو عاملين Two factorial design أو تصميم ذو ثلاثة عوامل (ثلاثة عوامل مستقلة) بينما الدراسة التى تتضمن متغير مستقل وحيد يطلق عليها تصميم العامل الوحيد Single factor designs.

وصف التأثيرات فى التصميم المعقد أو العاملى

1. يستخدم الباحثون التصميمات المعقدة لدراسة تأثير متغيرين مستقلين فأكثر فى تجربة واحدة.
2. ويمكن فى التصميمات المعقدة دراسة كل متغير مستقل فى ضوء تصميم المجموعات المستقلة أو تصميم القياسات المتكررة.

3. من أبسط التصميمات المعقدة أو العاملية هو تصميم 2×2 وهو يعنى وجود متغيرين مستقلين وكل متغير بمستويين.

4. عدد الشروط أو الظروف المختلفة للمعالجات يتحدد بضرب مستويات كل متغير مستقل مثل $2 \times 2 = 4$ أربع مستويات أو معالجات أو مجموعات.

5. يعد التصميم العاملى الأكثر كفاءة وقوة هو ذلك الذى يتضمن المستويات العديدة للمتغير المستقل أو المتضمن متغيرات مستقلة كثيرة فى التصميم.

وكل متغير مستقل فى التصميمات العاملية يتم تضمينه إما من خلال تصميم مجموعات مستقلة أو تصميم قياسات متكررة ولكن إذا كان التصميم العاملى يتضمن كلاً من متغير المجموعات المستقلة ومتغير القياسات المتكررة فإنه يسمى التصميم المختلط Mixed design. وتتضمن التجربة البسيطة متغير مستقل واحد بمستويين وهكذا التصميم العاملى البسيط يتضمن متغيرين مستقلين وكل متغير بمستويين، وتتحدد التصميم العاملى بعدد مستويات المتغير المستقل فالتصميم العاملى 2×2 هو أبسط التصميمات العاملية. ولا يوجد عدد محدد للتصميمات العاملية لأن أى عدد من المتغيرات المستقلة يمكن دراستها وأى متغير مستقل يمكن أن يكون لديه أى عدد من المستويات. وفى الممارسات البحثية من غير المعتاد أن توجد تجربة تتضمن أكثر من أربعة أو خمسة متغيرات مستقلة ولكن الشائع أن توجد تصميمات عاملية تتضمن متغيرين مستقلين أو ثلاثة. وفى التصميم العاملى 2×2 يوجد أربعة مستويات أو ظروف للمعالجات. وفى التصميم 3×3 يوجد متغيرين مستقلين كلاهما بثلاثة مستويات بالتالى توجد تسعة ظروف أو شروط للمعالجة أو تسعة مجموعات فى تصميم المجموعات المستقلة.

وفى تصميم $2 \times 4 \times 3$ حيث يوجد ثلاثة متغيرات مستقلة بثلاثة وأربعة ومستويين وبالتالى توجد 24 شرط أو ظرف للمعالجة أو 24 مجموعة. وتعتبر الميزة الأساسية للتصميمات العاملية هى أنها تمدنا بفرصة لتحديد التفاعلات بين المتغيرات المستقلة.

وفيما يلي مثال لتصميم عاملي حيث يدرس الباحث أثر شرب الكحول (عامل A) وشرب القهوة (عامل B) على الحالة المزاجية:

الجدول (1.1): بنية تصميم عاملي ثنائي 3×2 .

لا قهوة	200 مليجرام قهوة	400 مليجرام قهوة	
لا كحول	(الحالة المزججة)	الحالة المزججة	الحالة المزججة
كحول	الحالة المزججة	الحالة المزججة	الحالة المزججة

وعلى ذلك فإن المتغيرين المستقلين يكونوا مصفوفة 3×2 حيث يوجد ستة مجموعات، حيث المستويات المختلفة من القهوة يخلقوا ثلاثة أعمدة والمستويات المختلفة من الكحول يخلقوا صفوف وبالتالي يوجد ستة شروط معالجة وهي ستة تكوينات لتفاعل المتغيرين المستقلين ولذلك فهذا تصميم 3×2 .

وعلى ذلك كما لاحظنا فإن التصميم العاملي يخلق موقف أكثر طبيعية يشابه الحياة العادية وذلك مقارنة بالموقف التجريبي الذي يتضمن عامل وحيد في عزله عن بقية العوامل الأخرى المرتبطة بالمتغير التابع المراد دراسته وهذا مدعاة لكثير من التساؤلات وهو لماذا يجرى الباحث دراسات بسيطة ومنفصلة حيث يبحث بتأثير كل عامل بمعزل عن أهم العوامل أو المتغيرات الأخرى.

وتسمح التصميمات العاملية بالجمع بين عاملين فأكثر داخل دراسة واحدة تتيح للباحثين الفرصة ليروا كيف أن لكل عامل على حدة التأثير في السلوك وكيف لمجموعة من العوامل معاً أن تؤثر في السلوك وهذا يتناسب مع طبيعة الظاهرة النفسية حيث إنها متفاعلة بمتغيراتها. فبدراسة أثر التعزيز (الإيجابي والسلبي) وطريقة التدريس (مناقشة، محاضرة) على التحصيل فالباحث يدرس أثر التعزيز وطريقة التدريس على حدة على التحصيل ثم يدرس أثر العاملين معاً على التحصيل.

التأثيرات الرئيسية وتأثيرات التفاعل

التأثير الكامل لكل متغير مستقل فى التصميمات العاملية يسمى تأثير رئيسى ويعرض الفروق بين متوسطات العامل أو المعالجة كل على حدة. وتأثير التفاعل بين المتغيرات المستقلة يحدث عندما يكون تأثير أحد المتغيرات المستقلة تختلف اعتماداً على مستويات المتغير المستقل الآخر. وكما سبق فى جدول (1-1) يتضح أن مستويات تناول القهوة تمثل الأعمدة بينما مستويات تناول الكحول تمثل الصفوف وكل خلية فى المصفوفة تناظر اتحاد أو تفاعل العوامل. والباحث يهدف إلى ملاحظة وقياس مجموعة الأفراد تحت كل شروط أو ظروف معالجة موجودة فى الخلية.

وفى أى تصميم عاملى من المحتمل أن نختبر تنبؤات أو فروض فيما يخص التأثير الكامل لكل متغير مستقل فى التجربة وهذا يسمى تأثير رئيسى والبيانات من دراسة تجريبية عاملية تمدنا بثلاث مجموعات من المعلومات المنفصلة والمتمايزة وهو كيف العاملين المستقلين (مجموعتين) ومتفاعلين (مجموعة) يؤثروا على السلوك أو المتغير التابع. ولتوضيح الأنواع الثلاثة من المعلومات يمكن عرض بناء دراسة تأثير الكحول والقهوة على الحالة المزاجية فى الجدول الآتى:

الجدول (2.1): بيانات لمتوسطات المعالجات لدراسة ذو عاملين (متغيرين مستقلين) على الحالة المزاجية.

المتوسط الاجمالى	400 مليجرام قهوة	200 مليجرام قهوة	لا قهوة	
M= 600	M= 575	M= 600	M= 625	لا كحول
M= 650	M= 625	M= 650	M= 675	كحول
--	M= 600	M= 625	M= 650	المتوسط العام

والبيانات فى الجدول تمدنا بالمعلومات حيث كل عمود من المصفوفة يناظر مستوى محدد من استهلاك القهوة وبتقدير المتوسط الاجمالى لكل عمود يمثل كل مستوى من استهلاك القهوة وهذا يؤدى إلى ثلاثة متوسطات تناظر ثلاثة أعمدة وهذا يعطى مؤشراً لكيفية تأثير استهلاك القهوة على السلوك. والفرق بين متوسطات التأثير وبين الأعمدة

الثلاثة يسمى تأثير رئيسى للقهوة بمعنى الفرق هنا يحدد التأثير الرئيسى للعامل الأول، لاحظ أن حساب المتوسط لكل عمود يتضمن تقدير متوسطات كل مستويات الاستهلاك للكحول (نصف الدرجات مع عدم استهلاك الكحول والنصف الآخر لاستهلاك الكحول). ولذلك فإن استهلاك الكحول متكافئ أو متوازن عبر المستويات الثلاثة من استهلاك القهوة وهذا يعنى أن أى فرق متحصلاً عليه بين الأعمدة لا نستطيع تفسيرها عن طريق الفروق فى الكحول.

ثانيًا: المنهج شبه التجريبي Quasi experimental method : فى التصميمات شبه التجريبية يهدف الباحث إلى دراسة أثر المتغير المستقل على المتغير التابع فى تجربة تحدث مواقف الحياة الطبيعية مثل المدرسة أو المستشفى حيث إن إجراءات الضبط فى التصميم التجريبي الحقيقى غير متوفرة وكذلك لا يتوفر فيها التوزيع العشوائى للأفراد على ظروف المعالجات المختلفة، وهذه التصميمات سائدة فى العلوم النفسية والتربوية والسلوكية حيث إجراءات الضبط تبدو عملية صعبة إلى حد ما ولذلك فإن هذه التصميمات تعاني من نقص الصدق الداخلى Internal validity وهو الادعاء بأن التغير فى المتغير التابع يعود و يعود فقط إلى المتغير التجريبي و بالتالى فالتجارب التى تجرى فى واقعنا التربوى فى المدارس يغلب عليها التصميم شبه التجريبي و يطلق عليها تجارب الميدان Field experimental وعلى الرغم من ذلك فإن التصميمات شبه التجريبية لها درجة صدق خارجى External validity عالية حيث يعايش الباحث الظاهرة كما تحدث فى الواقع بالتالى فإن نتائج التجربة لها قدرة تعميمية عالية.

ثالثًا: المنهج الارتباطى Correlational method or design: الهدف من إستراتيجية البحث الارتباطى هو دراسة ووصف العلاقة أو الارتباط بين المتغيرات ولا يحاول تفسير العلاقة ولا يهتم بأى محاولة للضبط أو التجريب أثناء دراسة المتغيرات كما هو سائد فى البحث التجريبي. فالبحث الارتباطى هو أحد الطرق للوصف فى صورة كمية موضحاً إلى أى درجة واتجاه العلاقة بين المتغيرات، ويعتبره البعض أحد أشكال البحث أو المنهج الوصفى ويجب التأكيد أن اكتشاف العلاقة الارتباطية لا يعنى

وجود ارتباط سببي بالتالى فالارتباطية لا تعنى بالضرورة سببيه Correlation doesn't mean Causality ولكن وجود علاقة ارتباطية يمكن أن تكون سبباً أساسياً لوجود دراسات تنبئية ودراسات تحليل المسار Path analysis. والبحث الارتباطى يعطى الوضع عن المتغيرات التى نبحث عنها والباحث لا يؤثر على الاحداث والقياسات تكون غير متحيزة.

المتغيرات Variables

فى اى دراسة يتعامل الباحث مع مجموعة بيانات Data set ويحصل عليها من مرحلة جمع البيانات Data collection وفيها يمتلك أى فرد أو حالة فى عينة الدراسة مجموعة بيانات تعكس مجموعة متغيرات مثل الجنس، التحصيل، الذكاء، وغيرها، ويطلق على بيانات متغيرات مثل الجنس، موقع السكن، التخصص، الوظيفة، والحالة الاجتماعية بالبيانات الكيفية Qualitative data ويسمى البعض بالبيانات الأساسية أو الديموجرافية وهى تقاس بصورة كيفية تصنيفية. ويمكن أن يكون المتغير الكيفى ثنائى Dichotomous مثل الجنس (ذكر - أنثى) أو متعددة Polychotomous مثل الديانة (مسلم - مسيحى - يهودى) ، ويطلق على بيانات متغيرات مثل الابتكارية، الاتجاه، الامن النفسى، ومفهوم الذات بالبيانات الكمية Quantative data وهى تقاس بصورة كمية أو عددية و يتم تحليلها بإجراءات إحصائية معينة.

وتصنف المتغيرات حسب اتصالياتها إلى:

1. متغيرات متصلة (مستمرة) Continuous variable: يعنى أن عدد القيم الواقعة بين اى قيمتين (نقطتين) على متصل السمة لا نهائى وغير قابل للعد. المتغيرات المتصلة غير قابله للعد نظرياً ودائماً المتغيرات الفترية والنسبية هى متغيرات متصلة. وكأمثلة على المتغيرات المتصلة: العمر، الوزن، الطول، المسافة ولكن عدد أفراد الاسرة ليس متغير متصل على الرغم أنه نسبى.

2. المتغيرات المتقطعة (المنفصلة) Discrete variables : هى المتغيرات القابلة للعد مثل متغير الجنس (ذكر وانثى) أو متغير موقع السكن (الريف - البدو - الحضر)

بمعنى لا توجد قيم بين اى قيمتين للمتغير فمثلاً بين الذكور والريف لا يوجد شىء اخر بينهما. والمتغيرات الاسمية مثل الجنس والديانة والتخصص فى الكلية والوظيفة هى متغيرات متقطعة فى طبيعتها وقابلة للعد ويطلق عليها أيضاً متغيرات تصنيفية

. Categorical variable

وظيفتها فى التصميم التجريبي أو التصميم البحثى كالآتى:

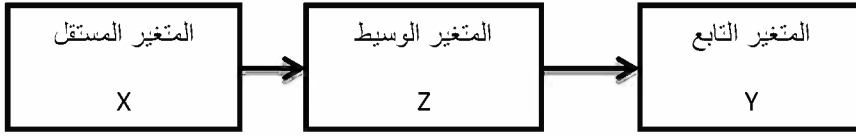
1. المتغير المستقل Independent variable: هو المتغير الخاضع لسيطرة الباحث أو سيطرة المجرب، وهو خاصية أو سمة تؤثر فى الناتج أو التأثيرات ، ويمكن للمتغير المستقل أن يؤثر فى المتغير التابع من خلال متغير وسيط أو ناقل بالتالى يكون التأثير غير مباشر، وتسمى المتغيرات المستقلة متغيرات منبئة، معالجات، ومتغيرات سابقة، فمثلاً طريقة التدريس أو التعزيز، ما وما وراء المعرفة كلها متغيرات يسيطر عليها الباحث ومستقلة للتحصيل، فالباحث الذى يدرس أثر المتغير (X) على متغير اخر (Y) فإن المتغير (X) مستقل ولا بد أن يكون سابق فى الحدوث زمنياً عن المتغير (Y) ويسمى المتغير المستقل أحياناً فى التصميمات التجريبية بمتغير المعالجة Treatment variable، وفى الدراسات التنبؤية يسمى المتغير المنبئ. وإذا أراد الباحث أن يتنبأ بالنجاح الاكاديمى فى الجامعة من درجات الثانوية العامة فى هذه الحالة فإن امتحانات الثانوية العامة متغير مستقل، أى أن المتغير المستقل هو السبب لحدوث تغير فى متغير اخر تالى له. الدراسة أو البرنامج هو متغير مستقل. ومتغيرات المعالجة هى نوع من المتغيرات المستقلة وتستخدم فى البحوث التجريبية والمعالجات التربوية، فيعطى المعالجة لمجموعة بينما يستبعداها من المجموعة الأخرى. ويصبح السؤال ما إذا كانت المجموعة التى تلقت المعالجة تتفوق أو تختلف على المجموعة التى لم تتلق المعالجة. بالتالى فإن هذا المتغير تصنيفى يمكن أن يكون بمستويين فأكثر. وفى التراث البحثى تسمى هذه المتغيرات بالمعالجات أو المستويات أو ظروف المعالجة. والمتغيرات المستقلة فى التربية وعلم النفس كثيرة مثل برنامج علاجى أو تدريبي، مهارات ما وراء المعرفة، أثر عادات العقل، أثر عادات الاستذكار، عدد ساعات الاستذكار، العلاج السلوكى، العلاج العقلانى، وغيرها.

2. **المتغير التابع Dependent variable**: هو خاصية أو سمة تعتمد على أو تتأثر بالمتغير المستقل ويكتب في تقارير البحوث تحت مسمى الناتج، التأثير، المحك، ويمكن للدراسة الواحدة أن تتضمن أكثر من متغير تابع ويمكن أن يكون قياسات كمية متصلة أو قياسات كيفية منفصلة. هو المتغير الناتج ويعتمد على سلوك المستجيبين ويتغير حدوثه نتيجة تغير المتغير المستقل أي أنه النتيجة، فالنجاح الأكاديمي في الجامعة هو المتغير التابع وهو المتغير الذي يأتي بعد المتغير المستقل أو يحدث معه في نفس اللحظة، ونواتج الدراسة (البيانات) هي متغيرات ناتجة تابعة. وأمثلة للمتغيرات التابعة في التربية وعلم النفس التحصيل، تنمية الابتكارية، تنمية الاتجاه، خفض القلق، تنمية الذكاء الاجتماعي، المهارات الإدارية، المناخ الأسري، خفض الاكتئاب.

3. **المتغير الضابطة Control variable**: هي متغيرات مستقلة ولكن يضبط أثرها أثناء التجربة، فمثلا أراد باحث دراسة أثر التماسك الأسري على الاكتئاب، فلدراسة هذا فإنه توجد متغيرات أخرى تؤثر على الاكتئاب مثل المستوى الاقتصادي للأسرة والوحدة النفسية وهي متغيرات تؤثر على النواتج فلا بد أن يضبطها الباحث ويستبعد أثرها عند تصميمه للتجربة من خلال قياسها. ومن أمثلة هذه النوعية من المتغيرات الحالة الاقتصادية الاجتماعية، النوع، الجنس، والذكاء. ويمكن ضبط هذه المتغيرات إحصائياً وإذا كانت متغيرات متصلة يطلق عليها تغايرات Covariates من خلال أساليب إحصائية مثل تحليل التغاير Analysis of Covariance

4. **المتغير الوسيط Mediating variable**: هي متغيرات تتأثر بالمتغير المستقل وتؤثر على المتغير التابع، بمعنى أنها تنقل أثر متغير مستقل إلى المتغير التابع ويقسها الباحث وتدخل في المعالجة الإحصائية، وهي متغيرات شائعة في النماذج السببية وتلعب دور المتغير المستقل والتابع في نفس الوقت. وكأمثلة للمتغيرات الوسيطة البناء المعرفي والذكاء وغيرها وهي متغيرات فترية. ويرى (Creswell 2012) أن المتغيرات الوسيطة هي متغيرات متداخلة Intervening حيث إنها تختلف عن المتغيرات التابعة وأي نوع من المتغيرات المستقلة، فباستخدام تفكير السبب والنتيجة، فإن بعض العوامل تتداخل بين المتغير المستقل والتابع وتؤثر في المتغير التابع، بل

تتقل أثر المتغير المستقل إلى المتغير التابع، بالتالى تتوسط العلاقة بين المستقل والتابع. وفى بعض الدراسات الكمية يمكن ضبطها باستخدام أساليب إحصائية.



5. المتغير الدخيل Confounding: هى متغيرات غير مقاسة وغير ملاحظة وتلعب دوراً فى التصميم التجريبي أو فى العلاقة السببية بين المتغير المستقل والتابع ويطلق عليها بالمتغيرات المشوشة أو Spurious variables ولا يمكن استبعادها من هذه المتغيرات وتؤثر فى طبيعة العلاقة بين المتغيرات المستقلة والتابعة.

ويمكن أن يسأل الباحث عدة تساؤلات عن طبيعة هذه المتغيرات كالتالى:

- ما هو النواتج فى دراستى التى تحتاج لتفسير؟ (متغيرات تابعة).
- ما هى المتغيرات أو العوامل التى تؤثر فى الناتج؟ (متغيرات مستقلة).
- ما هى المتغيرات التى احتاج لقياسها والتى من المحتمل أن تشكل عوامل أساسية فى تأثيرها على النواتج وليس متغيرات أخرى؟ (المتغيرات الضابطة).
- ما هى المتغيرات التى ربما تؤثر على نواتج الدراسة ولا يمكن قياسها؟ (المتغيرات الدخيلة).

6. المتغيرات المتفاعلة Moderating variables: هذه المتغيرات تستحق الانتباه وغالباً ما تستخدم فى التصميمات والتجارب التربوية، هذه المتغيرات هى جديدة وتولد نتيجة التفاعل مع متغيرات أخرى فمثلاً عن طريق تفاعل مستويات المتغيرات المستقلة المعالجة ويطلق عليها تأثير التفاعل Interaction effect ومثال لهذا: الطلاب الذكور مع طريقة المناقشة يحصلوا على درجات تحصيلية أكثر من الطالبات الإناث مع المحاضرة.

الفصل الثانى

إعداد البيانات للتحليل

Data preparation

ويتضمن هذا الفصل عدة قضايا أساسية أهمها:

أولاً: **حجم العينة sample size**: لحجم العينة تأثير كبير على ثبات تقديرات معالم الأسلوب الإحصائي، والقوة الإحصائية. فكلما زادت حجم العينة زادت دقة التباينات بين المتغيرات، بدوره يؤدي إلى الحصول على نتائج ثابتة وصادقة. وأن استخدام أحجام عينات صغيرة يؤدي إلى قوة إحصائية منخفضة للحصول على دلالة إحصائية لتقديرات المعالم وقضية تحديد حجم العينة تلعب فيها عوامل عديدة مثل عدد المتغيرات المقاسة وخصائص التوزيع فإذا كانت المتغيرات متصلة وذات توزيع اعتدالي والعلاقة بينهم خطية فيمكن الاعتماد على أحجام عينات معقولة نسبياً.

مدخل تحليل القوة الإحصائية

اشار (1979) Cook & Campbell إلى أن صدق الاستنتاج الإحصائي Statistical conclusion validity للدراسة يتحدد بعدة عوامل أهمها تصميم البحث، مستوى الدلالة الإحصائية (الخطأ من النوع الأول)، تباين المجتمع في المتغير المحك (التابع)، حجم الفروق بين القيم الحقيقة للمعلم في المجتمع والقيمة المحددة عن طريق الفرض الصفرى (حجم التأثير)، نوع الفرض ذو ذيل واحد (موجهة) مقابل ذو اتجاهين (غير موجهة)، نوع الاختبار المستخدم، وحجم العينة. وكلما زاد حجم العينة فإن النتائج تكون أكثر مصداقية لتمثيلها للمجتمع وتقل أخطاء المعاينة.

وتعتبر القوة الإحصائية دالة وظيفية للعوامل الآتية:

- مستوى الدلالة الإحصائية (α) أو الخطأ من النوع الأول
- الخطأ من النوع الثانى (بيتا β)
- مقدار حجم التأثير المفترض
- الانحراف المعياري أو التباين لحجم التأثير المفترض

• حجم العينة

والمنطق وراء تحليل القوة هو أن الدراسة لو اعتمدت على حجم عينه كبيراً كبيراً كافياً فإنه توجد احتمالية مسبقة للحصول على تأثير أو فروق إذا كانت موجودة بالفعل وإذا كان حجم العينة صغيراً فإن الدراسة تفشل في الكشف عن الفروق الموجودة بالفعل في المجتمع وعلى ذلك فالدراسة مضيعة للوقت والجهد والموارد.

ويفيد تحليل القوة الإحصائية في انتقاء حجم العينة في ضوء أسس منطقية بدلاً من انتقائها في ضوء القواعد المتعارف عليها Rules of thumb التي ليس لها أسس علمية موضوعية وعلى ذلك ففي تحليل القوة قبل إجراء الدراسة لا بد من تحديد حجم التأثير وقيمتي α ، β ، تحديد اتجاه الاختبار، وعلى ذلك يمكن تقدير حجم العينة المناسب. وفي هذا التحليل يرغب الباحث في تحديد حجم العينة في ضوء مستويات مقبولة من حجم التأثير و α و القوة وقد يرتضى الباحث $\alpha=0.05$ والحد المقبول المناسب من القوة هو 0.80.

والقضية هنا هي كيفية تقدير حجم التأثير وفي هذا الشأن يوجد ثلاثة مداخل لتقدير حجم التأثير هي:

1. إجراء دراسة استطلاعية ($n=50$) مثلاً وتقدير حجم التأثير للاختبار المستخدم.
2. مراجعة الدراسات السابقة وتقدير متوسط حجم التأثير للدراسات التي تناولت نفس متغيرات الدراسة ويمكن إجراء ما وراء التحليل لعدد محدود يتراوح من 5 إلى 10 دراسات للوصول إلى حجم التأثير في الدراسات السابقة.
3. تبني معايير حجم التأثير للاختبار المستخدم، وهذه المعايير وضعها (Cohen 1988) ، فمثلاً في حالة اختبار T للعينات المستقلة القيمة $d = 0.2$ تعبر عن حجم تأثير ضعيف، 0.50 حجم تأثير متوسط، 0.80 كبير.

مثال: حساب حجم العينة لاختبار T المستقلة

وفيما يلي مثال لتحديد حجم العينة في حالة استخدام اختبار T باستخدام تحليل القوة القبلي باستخدام برنامج G-POWER وحيث أراد باحث دراسة فعالية طريقة تدريس ما على تحسين القدرة اللغوية. حدد المعالم الآتية:

1. الاختبار هو T لعينات مستقلة ذو ذيلين
 2. مستوى الدلالة الإحصائية $\alpha=0.05$
 3. القوة الإحصائية المناسبة $=0.80$
 4. تبني الباحث حجم تأثير متوسط $= 0.50$
- وفيما يلي الشاشة الافتتاحية للبرنامج حين يتم تحديد Test family وتتضمن اختبارات عديدة وهي كما هو موضح :

- اختار التحليل القبلي Apriori تحت مربع Type of power analysis حيث يهدف إلى تقدير حجم العينة المناسب.

- أسفل عائلة الاختبارات اختار t-test

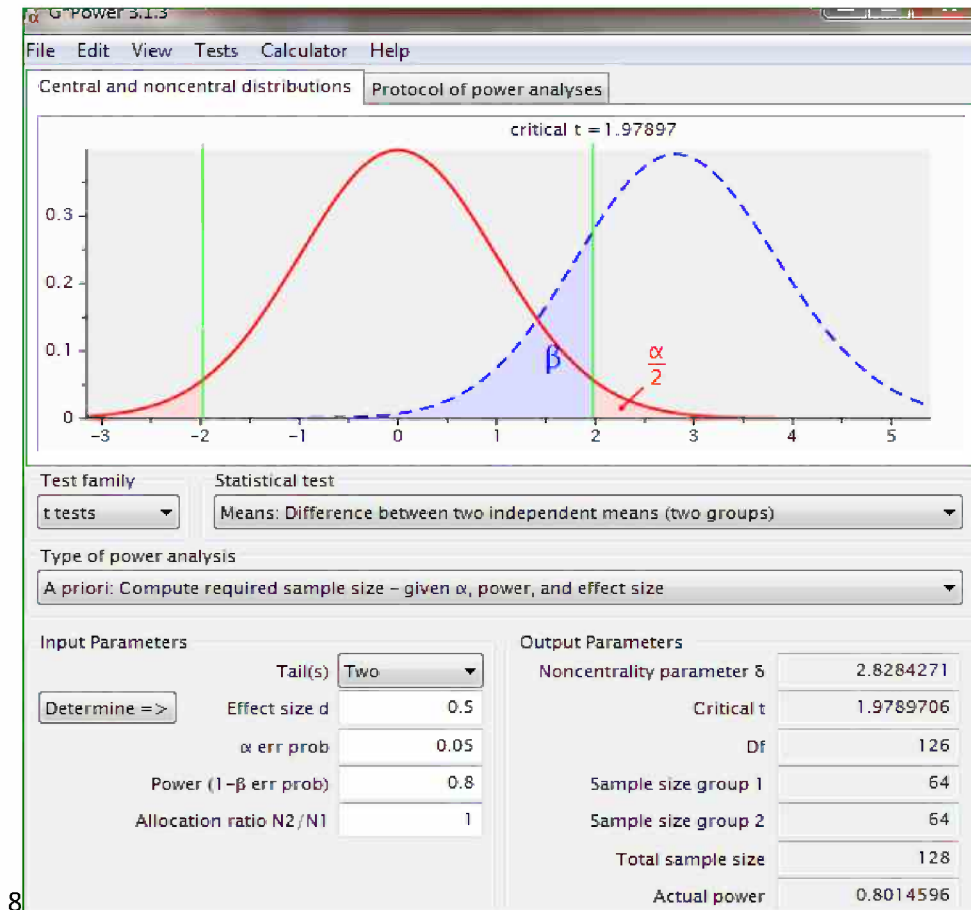
- أسفل Statistical test اختار نوع الاختبار المراد حجم العينة له وفي هذه الحالة هو:

Means: Difference between two independent means (two groups)

- ويتضمن البرنامج معالم الإدخال Input parameters قم بإدخالها وهي: تحديد طبيعة

اتجاهية الاختبار ذو ذيل واحد أو ذو ذيلين، حجم التأثير، قيمة α ، قيمة القوة P ، وتحديد نسبة حجم العينة الأولى $N1$ إلى الثانية وغالبًا تكون النسبة $N1:N2$ وهي 1:1 وتم إدخال المعالم السابقة في البرنامج.

- اضغط Calculate تظهر الشاشة الآتية :



تم اعطاء قيمة T الحرجة = 1.97 و $Df = 126$ وحجم العينة الكلي = 128 مقسمة إلى 64 فردًا في المجموعة الضابطة و 64 فردًا في المجموعة التجريبية وذلك لتحقيق قوة فعلية 0.8014

مدخل القواعد المتعارف عليها Rules of thumbs

قدم (2007) Wilson Vanvoorhis & Morgan إرشادات عملية واطار لتحديد حجم العينة في ضوء القواعد المتعارف عليها:

- لابد أن تتضمن الدراسة حجم عينة مناسب لهدفها، ويجب أن تكون كبيرة كبرًا مناسبًا وذلك في محاوله الحصول على دلالة إحصائية. فحجم عينة غير مناسب هو

مضيعة للموارد والوقت وحجم عينة كبير غير ضرورى وذلك لأنها تحمل الباحث عبء شديد. غالبا ما يتم تجاهل تحديد حجم العينة فى ضوء مدخل القوة الإحصائية، ولكن لا بد على الأقل الانتباه إلى القواعد المتعارف عليها، وفيما يلى جدول يوضح ذلك:

الجدول (1.2): حجم العينة فى ضوء القواعد المتعارف عليها WilsonVanvoorhis& Morgan (2007)

الاختبار	حجم العينة المناسب
T , ANOVA	تتضمن الخلية 30 فردا لتحقيق $P=0.8$ واذا قلت ليس أقل من 7 . 50 فردًا على الأقل 300 فاكتر جيد
الارتباط والانحدار التحليل العاملي	

وفيما يلى توضيح لذلك:

أولاً: **حجم العينة لاختبارات الفروق**: حجم الخلية لاختبارات الفروق مثل اختبار T المستقلة والمرتبطة وتحليل التباين الأحادى ANOVA وتحليل التباين المتعدد وتحليل التباين المتدرج MANOVA ولتحقيق حجم تأثير متوسط إلى كبير يجب أن يكون 30 فردًا فى كل خلية لان هذا يقود إلى قوة إحصائية 80%. وفى بعض الظروف ولبعض الأسباب فإن أدنى حد لحجم العينة 7 أفراد وهذا يؤدي إلى الحصول على قوة 50%، و14 فردًا لكل خلية ولحجم تأثير 0.50 يؤدي إلى قوة تقريباً 80%. ويجب الإشارة هنا إلى أن الخلية تمثل المجموعة، فى هذه الحالة فإن اختبار T يتضمن خليتين وF يتضمن خليتين فاكتر. ويجب التأكيد على عدة امور منها:

1. المقارنة بين مجموعات قليلة (2 مثلاً) يتطلب حجم عينة مناسب وذلك للحصول على قوة مناسبة.

2. إذا كان حجم التأثير المتوقع صغير أو منخفض فإن هذا يتطلب حجم عينة أكبر لتحقيق مستوى قوة مناسب.

3. لاستخدام تحليل التباين المتعدد أو المتدرج فمن الضروري توافر حجم عينة كبيراً في كل خلية (Tabachnick & Fidell, 2007).

ثانياً. **حجم العينة للتحليل العاملي:** القاعده الجيدة للعينة في التحليل العاملي هو 300 فرداً (Tabachnick & Fidell, 2007) أو أكثر من 50 فرد لكل عامل (Pedhazur & Schmelkin, 1991). وقدم Comrey & Lee (1992) القاعدة الآتية: 50 فقير جداً، 100 حجم عينة فقير، 200 مناسب، 300 جيد، 500 جيد جداً، و1000 ممتاز. وأشار Guadagnoli & Velicer (1988) ان الحلول العاملية مع وجود تشبعات عالية للمتغيرات (أكبر من 0.80) لا تتطلب استخدام عينات كبيرة. ويرى البعض أن يمثل كل معلم من 5 إلى 10 أفراد إذا كان توزيع المتغيرات اعتدالي (Bentler & Chou, 1987)، أو عشرة أفراد، أو من 10 إلى 20 فرداً (Kline, 2016)، أو كحد أدنى 20 فرداً (Costello & Osbrne, 2005). ويجب التأكيد على أن هذه القواعد تحت شرط مسلمة الاعتدالية للبيانات.

ثالثاً: **حجم العينة اللازمة لدراسة العلاقات والانحدار:** يوجد العديد من الصيغ المعقدة في هذا الشأن وعموماً يجب استخدام حجم عينة لا تقل عن 50 فرداً للارتباط وأيضاً للانحدار ولكن يجب أن يزيد العدد في حالة زيادة عدد المتغيرات التابعة. وقدم Green (1991) اطاراً شاملاً للإجراءات المستخدمة لتحديد أحجام عينات الانحدار واقترح الصيغة الآتية: $N > 50 + 8m$

حيث m عدد المتغيرات المستقلة أو المنبئات وذلك لاختبار الدلالة الإحصائية لمعامل الارتباط المتعدد (اختبار F).

ولاختبار معالم المنبئات أو المتغيرات المستقلة (اختبار T) تستخدم الصيغة الآتية:

$$N > 104 + m$$

وقدم Harris (1985) صيغة تستخدم لتحديد عدد أفراد العينة في تحليل الانحدار لعدد منبئات 5 فأقل فاقترح حجم العينة 50 على الأقل وهو أن العدد الكلي من أفراد العينة يساوى عدد المنبئات بالإضافة الى 50 فرداً:

$$N > m + 50$$

فمثلاً إذا كان تحليل انحدار يتضمن ستة متنبئات (متغيرات مستقلة) أو أكثر فإن الحد الأدنى المطلق هو تمثيل المتغير المستقل بعشرة أفراد وهذا مناسب.

مدخل فترات الثقة

الاستراتيجية الأخرى لتحديد حجم العينة هي على أساس الحدود المرغوبة من فترات الثقة (Confidence interval (CI وهو مدى من القيم حول ما إذا كان معلم المجتمع (المتوسط) من المحتمل أن يقع فيه. فعلى سبيل المثال إذا كانت القيمة 0.0 (صفر) تقع داخل الفترة فإن التقدير غير دال إحصائياً لمستوى $(1-\alpha)$. وتقدر حدود الثقة باستخدام برامج مثل SPSS, SAS, Minitab. والمؤيدين لاستخراجية فترات الثقة يؤكدون على أنها تمدنا بمعلومات عن دقة التقديرات وحجم مسافة فترات الثقة تشير إلى الدقة (درجة الخطأ العشوائى المرتبط بها) ويمكن أن تكون التقديرات عالية الدقة إذا وقعت قيمتها داخل هذا المدى الضيق. واتساع مدى فترات الثقة يشير إلى دقة ضعيفة. واختيار حدود ثقة عند 99% يزيد من دقة حدود الثقة (يمتلك فرص كبيره لتضمين معالم المجتمع) ولكن يقلل من مستوى الدقة.

ويمكن تقدير CI قبل إجراء الدراسة وتستخدم كمرشد لاختيار حجم العينة، وتقدر حتى لو لم يوجد فرض إحصائي صفري. بعد الدراسة فإن CI يقدم معلومات عن الدقة وتكون أكثر فائدة واستخداماً من قيمة القوة الإحصائية. وعلى الرغم أن CI حول الفروق تكون أداة لفحص حجم الفروق ويوجد دعم قوى لاستخدام CI اما مكملًا أو بديلاً لاختبارات الفروض ويمكن أن تستخدم لتقدير حجم العينة وقارن Orme & Hudson (1995) بين مدخل القوة الإحصائية ومدخل فترات الثقة لتقدير حجم العينة وتوصلا إلى أن استخدام مدخل CI يمدنا بحجم عينة تؤدي إلى تقديرات أكثر دقة من مدخل القوة. ويوجد معادلات وجداول لتقدير حجم العينة في ضوء حدود فترات الثقة.

ثانياً: البيانات الغائبة أو المفقودة **Missing data**: تعتبر قضية البيانات الغائبة من أهم المشكلات فى تحليل البيانات وتحدث عندما يرفض أو ينسى المفحوص الاستجابة على مفردة أو أكثر وهذا يؤثر على تعميم النتائج، والإجراءات الإحصائية لأى أسلوب تتطلب أن تكون وحدة التحليل تامة أو كاملة البيانات. ويتعامل الباحثون

دائماً مع البيانات الكاملة، والمشكلة التي تواجهه الباحثين أثناء التحليل الإحصائي هو كيفية معالجة البيانات المفقودة (غياب بيانات على متغير واحد فأكثر لفرد واحد فأكثر) وتظهر البيانات المفقودة نتيجة أخطاء الإدخال، أو مشاكل أثناء جمع البيانات أو الرفض للاستجابة على مفردة ما فأكثر وقضية كيفية تحليل قواعد بيانات بها درجات غائبة هي معقدة.

وتوجد عدة استراتيجيات للتعامل مع البيانات الغائبة منها كما اشار اليها (Hair et al., 1998; Kline, 2016)

• **طريقة الحالة المتاحة Available case method:** وهي تحليل البيانات المتاحة وهذا يتضمن نوعين أساسيين هما:

1. **طريقة الحذف List-wise deletion:** الحالات التي لها درجات غائبة على أى متغير تستبعد من كل التحليلات، وهذه الطريقة تتعامل مع الحالات كاملة البيانات فى التحليل **Complete case analysis**، وعلى ذلك يحدث نقص لحجم العينة، وهي من الإجراءات البسيطة والمباشرة ويمثل هذا **بالاتجاه المحافظ**، ويطلق على البيانات فى هذه الحالة **Complete data only**. ومن المتوقع أن تكون الأخطاء المعيارية بعد تطبيق هذه الإستراتيجية أكبر من مثيلاتها على قاعدة البيانات الكلية قبل الحذف، ولكنها تعطى تقديرات دقيقة لمعالم النموذج، ومن اهم مميزاتها هي أن كل التحليلات تكون لنفس العدد من الحالات وعلى ذلك تقدر مصفوفة التغاير (الارتباط) للحالات كاملة البيانات فقط.

2. **طريقة الحذف Pair-wise deletion:** وتسمى أيضاً تحليل الحالات المتاحة فقط **Available case analysis** ففي الطريقة السابقة يحذف الفرد الذى لديه بيانات غائبة على أحد المتغيرات من التحليل، ولكن فى **طريقة Pair-wise** يبقى الفرد فى التحليل للاستفادة منه بالبيانات الموجودة له على بقية المتغيرات، وربما يكون هذا حل فعال ولا يسبب فقد الحالة كلها. ولكن إذا كانت البيانات الغائبة على المتغير التابع فإن الحالة كلها تحذف من التحليل، وهذه الطريقة تعطى معالجات إحصائية بأحجام عينات مختلفة من تحليل إلى آخر. بافتراض أن حجم العينة = 300 ، وتتضمن حالات

غائبة، إذا حالة بدون درجات غائبة على المتغيرات فأحياناً يتم تقدير معامل الارتباط ($r_{x1\ x2} = 0.30, N = 290$) وأحياناً ($r_{x1\ y} = 0.85, N = 255$)، وأحياناً ($r_{x2y} = 0.43, N = 280$) أى تختلف حجم العينة من تحليل لآخر وهذا يقود إلى عدم الاتساق الرياضى. وإذا كان حجم العينة كبيراً كافياً ويوجد عدد قليل من المفردات بها بيانات غائبة فلا توجد فروق عملية بين الطريقتين ولكن إذا وجد عدد كبير من أفراد العينة لديهم بيانات مفقودة على مفردات المقياس ففى هذه الحالة فإن الطريقة المفضلة هى pair wise ، وفيما يلى مثال للبيانات المفقودة أو الغائبة:

الجدول (2.2): مثال لمجموعة بيانات غير كاملة.

Y	X ₂	X ₁	CNo
13	8	42	1
12	10	34	2
-	12	22	3
8	14	مفقود	4
7	16	24	5
10	مفقود	16	6
10	مفقود	30	7

- الطرق التعويضية الانفرادية Single imputation Methods: وعلى الرغم من انتشار طريقتى pairwise و listwise إلا أنه توجد بعض الطرق التعويضية للتعامل مع البيانات الغائبة وأهمها:

1. استبدال البيانات الغائبة بالمتوسط العام لاستجابات الأفراد على المتغير وعلى ذلك تستخدم العينة كاملة العدد فى التحليل group – mean substitution. فمثلا فى جدول (1.2) نلاحظ أن المتغير X_1 به الحالة 4 ليس لها قيمة على هذا المتغير، ويتقدير المتوسط لـ X_1 يتبين أن $X_1 = \frac{148}{7} = 21.14$ بذلك تكون القيمة الغائبة للحالة 4 على المتغير X_1 هى 21.14، ولكن استبدال البيانات الغائبة بالمتوسط يغير أو يشوه خصائص التوزيع للبيانات عن البيانات قبل التعويض بالمتوسط، وذلك لأنها تقلل من التباينات للمتغيرات التى لها بيانات غائبة. ويفضل

عدم استخدامها إذا كانت البيانات الغائبة كثيرة وتستخدم إذا كان حجم البيانات الغائبة صغيراً.

2. التعويض عن طريق حساب الانحدار Regression: تستبدل البيانات الغائبة بالدرجة المنبئة عن طريق حساب الانحدار المتعدد لدرجات المتغير الذى يتضمن بيانات غائبة من خلال المتغيرات الأخرى فى ملف البيانات ففى البيانات الموضحة بالجدول السابق يعتبر X_1 و X_2 منبئين بالمتغير Y :

$$\hat{y} = b_1 X_1 + b_2 X_2 + a$$

وهكذا بالنسبة لـ X_1 تابع والمتغيرات الأخرى مستقلة:

$$X_1 = b_1 y + b_2 x_2 + a$$

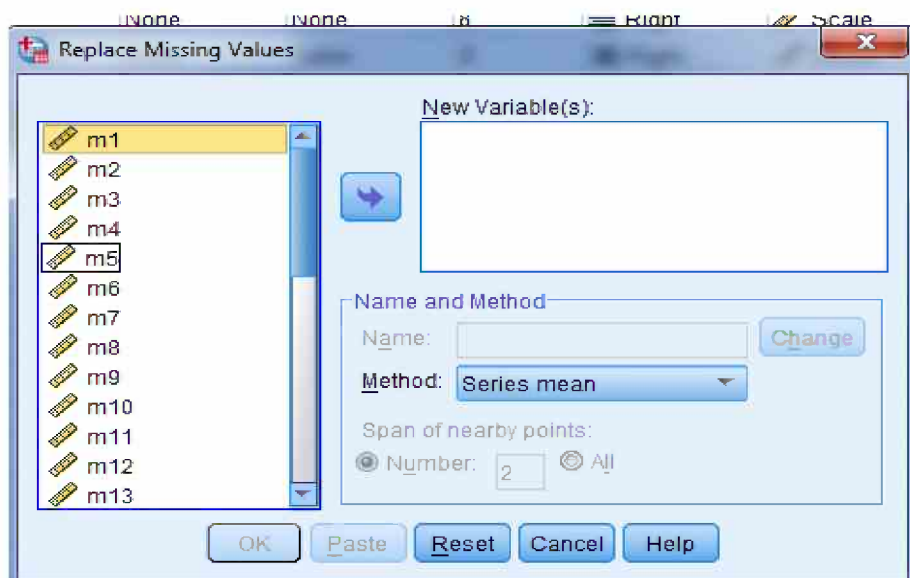
واستخدام هذه الطريقة يتطلب معلومات أكثر من طريقة التعويض بالمتوسط. ويمكن استخدام الانحدار اللوجستى Logistic regression إذا كانت المتغير التابع تصنيفى. وعموماً هى أقل مرغوبة مقارنة بطرق التعظيم المتوقع.

وإذا تواجدت البيانات الغائبة لدى عدد معين من الأفراد فإنه يمكن حذفهم من قاعدة البيانات (Tabachnick & Fidell, 2007)، ويفضل إذا استخدمت إحدى الطرق التعويضية للتعامل مع البيانات الغائبة أن تعيد التحليل مرة أخرى مع مصفوفة الارتباطات بدون استخدام الطرق التعويضية بمعنى أن البيانات كاملة. وإذا كانت النتائج متماثلة فتوجد ثقة فى الطريقة المستخدمة لمعالجة البيانات الغائبة ولو حدث اختلاف فالباحث بحاجة إلى معرفة أسباب هذا التناقض ويجب عرض نتائج التحليلين (قبل وبعد) معالجة البيانات الغائبة. حيث يعوض عن البيانات الغائبة من اقرب قيمة للحالة فى ملف البيانات قبله أو بعده.

استبدال البيانات الغائبة فى SPSS

ويتم إجراء الامر من خلال: 1. اضغط قائمة Transform واختار replace

missing data تظهر النافذة التالية:



2. انقل المتغيرات المراد معالجتها من أثر البيانات الغائبة إلى مربع **New variables**

3. اضغط على وظيفة **Method** فتظهر قائمة تحتها الطرق التى تقوم بمعالجة البيانات المفقودة أو الغائبة احدد طريقة معالجة البيانات الغائبة وليكن استبدالها بمتوسط المتغير.

4. اضغط **Ok** وبالرجوع إلى ملف البيانات يلاحظ استبدال البيانات المفقودة بمتوسط درجات المتغير.

ثالثاً: القيم المتطرفة Outliers: فى بعض الأحيان تتضمن بيانات العينة قيمة متطرفة أو أكثر على متغير واحد وهى **Univariate outlier** حيث تختلف تمامًا عن بقية البيانات، وهى بدورها تؤثر على قيمة العلاقات بين المتغيرات وهذا ينعكس سلباً على تقديرات معالم النموذج نتيجة تأثيرها على العلاقات بين المتغيرات وكذلك على مؤشرات الالتواء والتفرطح. وتوجد القيم المتطرفة نتيجة عدة أسباب أهمها وجود أخطاء إدخال أو ترميز أو نتيجة وجودها بصورة طبيعية. وإذا كانت القيم المتطرفة موجودة بصورة خاطئة تُحذف أما إذا كانت حقيقية فيجب التعامل معها، مثل إدخال عمر شخص 210 عاماً ولكنها خطأ إدخال فعمره الحقيقى 21 عاماً ويطلق على القيمة المتطرفة فى متغير واحد بـ **Univariate Outlier**. ولا يوجد تعريف واحد للقيمة المتطرفة ولكن القاعدة هى الدرجات التى لها أكثر من ثلاثة أضعاف الانحرافات

المعيارية عن المتوسط. ويمكن تشخيصها من خلال فحص التوزيعات التكرارية للدرجات المعيارية Z فإذا كانت $|Z| < 3.0$ فإن هذا يشير إلى وجود قيم متطرفة، بينما القيم المتطرفة المتدرجة Multivariate Quilters وهى درجات متطرفة على متغيرين فأكثر وهذا بدوره يزيد من الوقوع فى الخطأ من النوع الأول أو الخطأ من النوع الثانى، وهذا مثال على متغير يتضمن درجة متطرفة (1, 2, 5, 5, 7, 11, 50) وتعتبر 50 قيمة متطرفة ويوجد عدة طرق لتشخيص القيم المتطرفة المتدرجة أهمها:

1. إحصاء (D) Mahalanobis distance : حيث يشير إلى المسافة بوحدات الانحراف المعيارى بين مجموعة من الدرجات للحالة الواحدة ومتوسطات كل المتغيرات ، وفى العينات الكبيرة ذات التوزيعات الاعتدالية فإن توزيع مؤشر D^2 يعامل مثل توزيع χ^2 . فقيمة D^2 مع القيمة المنخفضة لمستوى الدلالة الإحصائية (P) يؤدي إلى رفض الفرض الصفرى وعلى ذلك يوجد تطرف فى التوزيع، وينصح دائماً باستخدام مستوى دلالة إحصائية (α) عند ادنى قيم مثل $\alpha = 0.001$ وتقوم بعض البرامج مثل SPSS و SAS بطباعة مؤشر D^2 .

2. إذا كان الدرجات المعيارية كبيرة جداً على متغير واحد فأكثر حيث المتغيرات التى لها درجات معيارية تزيد عن 3.29 (اختبار ذو ذيلين $0.001 < P$) فيدل هذا على وجود قيم متطرفة.

3. العرض البيانى لبيانات المتغير مثل المدرج التكرارى أو Boxplot. واحد الطرق للتعامل مع الدرجات المتطرفة هو إجراء تحويل للمتغير الذى يتضمن القيم المتطرفة. حيث تؤدي هذه الاستراتيجية إلى جعل توزيع المتغير أكثر اعتدالية

اختبارات الفروض للتحقق من الاعتدالية Normality

توجد مداخل عديدة للتحقق من مسلمة الاعتدالية منها ما هو وصفى من خلال:

1. عرض مقاييس النزعة المركزية: تساوى قيم المتوسط و الوسيط والمنوال يفيد بوجود الاعتدالية.

التحقق من اعتدالية البيانات فى SPSS

التحقق من الاعتدالية من العرض البيانى:

أولاً: إدخال البيانات: يتم إدخال متغير X في البرنامج وذلك من خلال تسميته من خلال الضغط على Variable view ثم كتابة مسمى المتغير في عمود Name ثم اضغط على Dataview

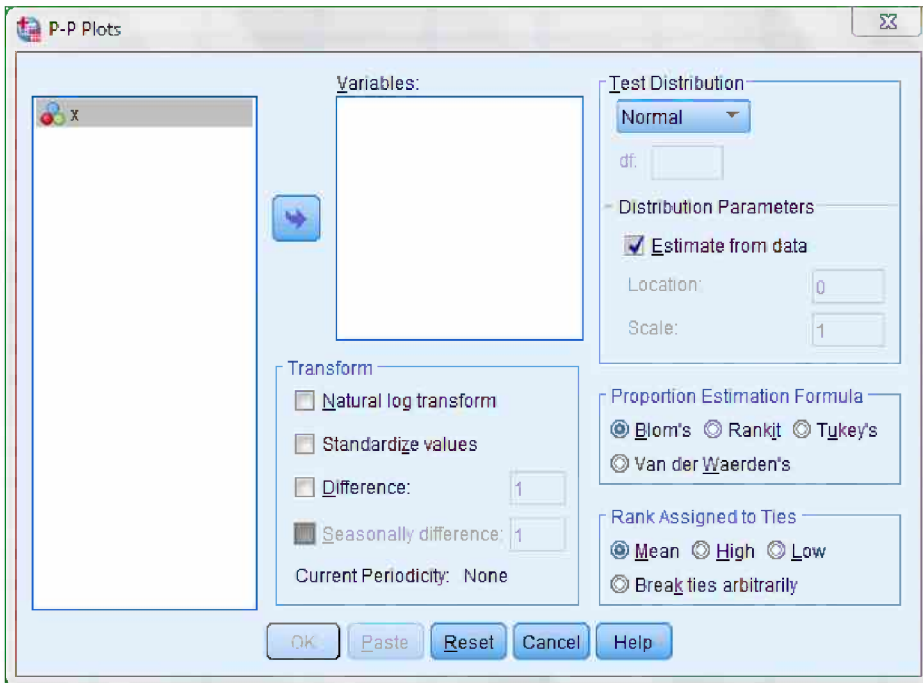
• شكل plots p-p (probability-Probability): وفي هذا المنحنى

يتم حساب الرتب المناظرة للدرجات ولكل رتبة يتم حساب قيمة Z الفعلية المناظرة وهذه هي القيمة المتوقعة ثم المطابقة بين الدرجات الخام بالدرجات المعيارية المناظرة لرتب الدرجات فإذا كانت الدرجات لها توزيع اعتدالي فإن الدرجات Z الفعلية سوف تكون على خط قطري مستقيم. والفكرة في هذا المنحنى هو مقارنة نقاط أو أحداثيات البيانات بالخط المستقيم القطري وإذا وقعت الدرجات على القطر فإن بيانات المتغير اعتدالية والابتعاد عن القطر يدل على الابتعاد عن الاعتدالية.

لتنفيذه اتبع الخطوات الآتية:

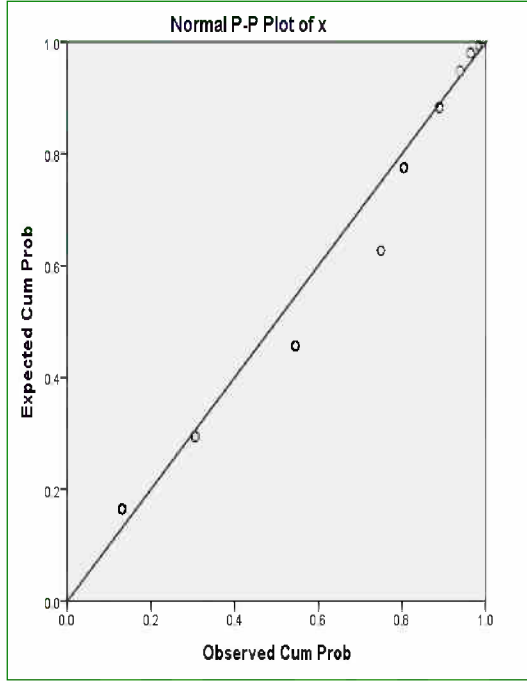
1 . اضغط Analyze ثم اختار Descriptive Statistics ثم اضغط على pp.plots

تظهر الشاشة الآتية:



2 . انقل متغير X إلى مربع variables. ثم اضغط على OK.

ثالثاً: المخرج :



كما هو ملاحظ أن الإحداثيات أو النقاط لا تقع تماماً عن الخط القطري وعليه فإن البيانات تبعد عن الاعتدالية ولا يتوافر فيها هذه المسلمة. وكما هو ملاحظ أن القيم الملاحظة على المحور السيني والقيمة المتوقعة على المحور الصادي. ويوضح تماماً أن الاحداثيات أو نقاط البيانات لا تقع على الخط المستقيم وعليه فإن البيانات غير اعتدالية.

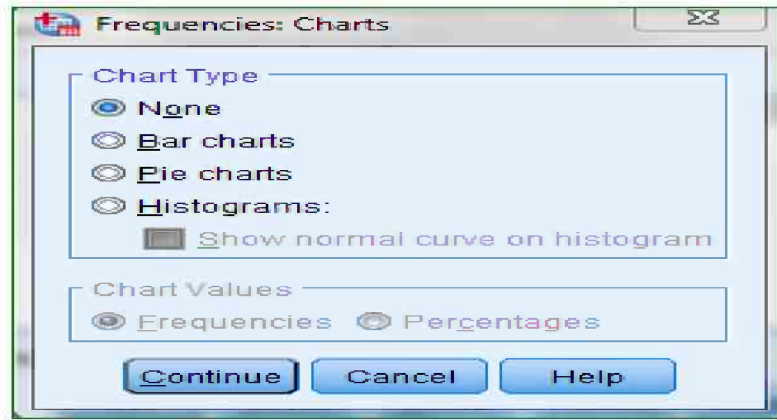
• المدرج التكرارى Histogram: تنفيذ الأمر :

1 . اضغط Analyze ثم Descriptive Statistics ثم اضغط Frequencies تظهر الشاشة :



2. اضغط x ثم انقله إلى مربع Variables

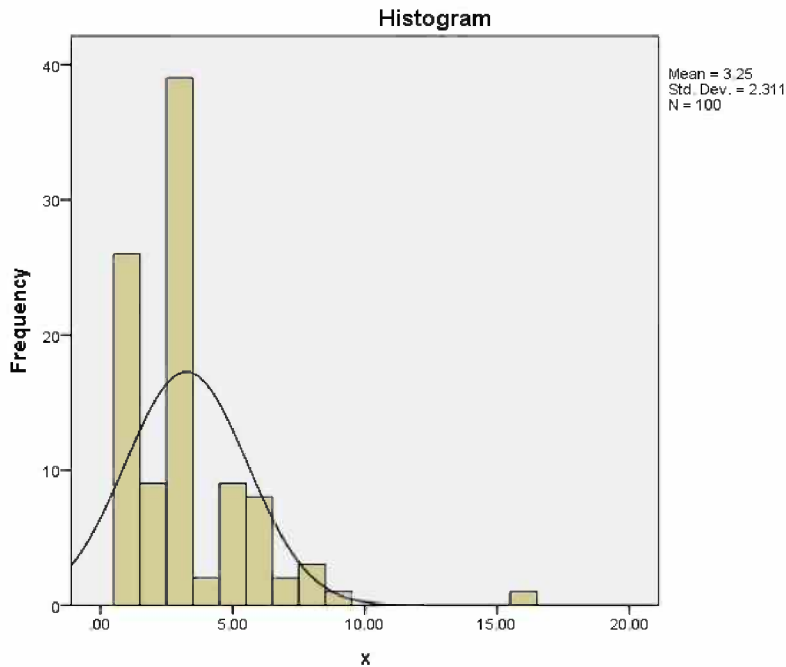
3 . اضغط على اختيار Charts تظهر الشاشة الآتية:



4. اضغط على Histograms و Show with normal curve

5. اضغط على Continue ثم اضغط Ok

المخرج :

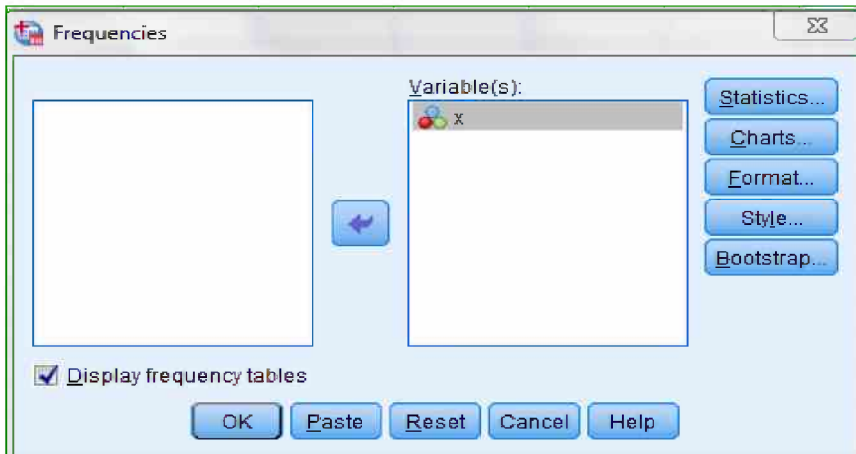


يظهر من المخرج أن المنحنى ملتوى ناحية اليمين بمعنى التواء موجب .

حساب مؤشري التفرطح والتواء باعتبارهما من خصائص التوزيع ويمكن تنفيذها في

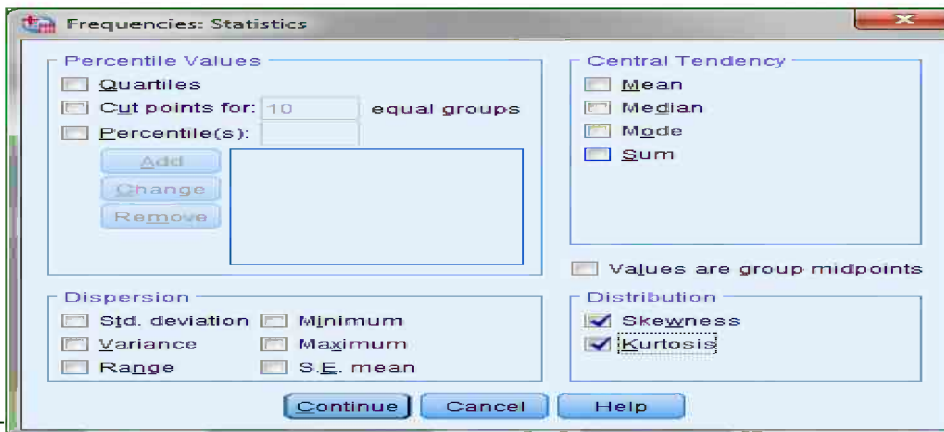
البرنامج:

اضغط Analyze ثم Descriptive Statistics ثم Frequencies



2. انقل المتغير x إلى مربع variables

3. اضغط على statistics تظهر الشاشة:



4. اضغط على Skewness, Kurtosis، واضغط Continue ثم OK

• المخرج :

Statistics		
N	Valid	100
	Missing	0
Skewness		2.164
Std. Error of Skewness		.241
Kurtosis		8.447
Std. Error of Kurtosis		.478

يتضح أن قيمة الالتواء = (2.164) أي

زادت عن الواحد الصحيح بما يدل على وجود التواء وبما أن قيمته موجب إذن هو التواء موجب وقيمة التفرطح = (8.447) أي زادت عن الواحد الصحيح بل قيمتها عالية جداً بما يدل على أن التوزيع ليس اعتدالي بل يوجد تفرطح وعليه فإن البيانات غير اعتدالية ويمكن تحويل قيم الالتواء

والتفرض إلى قيمة Z لاختبار دلالتها الإحصائية لها حيث إن :

$$Z_{skew} = \frac{S}{SE_{skew}}$$

$$k - 0$$

$$Z_{kurt} = \frac{K}{SE_{kurt}}$$

حيث S قيمة الالتواء و SE_{skew} الخطأ المعياري للالتواء و K قيمة التفرض و SE_{kurt} الخطأ المعياري للتفرض على حدة وعلية فإن :

$$Z_{skew} = \frac{2.164}{0.241} = 8.979$$

$$Z_{kurt} = \frac{8.447}{0.478} = 17.671$$

لاختبار الدلالة الإحصائية للالتواء والتفرض يتم مقارنة Z_{skew} و Z_{kurt} بـ 1.96 وهى قيمة Z لاختبار ذو ذيلين عند 0.05 أو مقارنتها بـ 2.56 قيمة Z لاختبار ذو ذيل واحد عند 0.05 وعليه فأن:

$$1.96 < (8.979) \rightarrow Z_{skew}$$

$$1.96 < (17.671) \rightarrow Z_{kurt}$$

إذا توجد دلالة إحصائية لقيمتى الالتواء والتفرض وعليه فإن التوزيع غير اعتدالى. ولكن عليك أن تكون حذراً عند استخدام اختبار Z لأنه يعطى دلالة للعينات الكبيرة و ينصح (2013) Field باستخدام Z فى حالة العينات الصغيرة والمتوسطة (50,100, 150) ولكن إذا زادت حجم العينة عن 200 فيفضل عدم الاعتماد على الدلالة الإحصائية لمؤشرى الالتواء و التفرض والاعتماد على قيمتهما المطلقة.

اختبارات الدلالة الإحصائية للتحقق من الاعتدالية

يوجد اختبارين للتحقق من الدلالة الإحصائية للاعتدالية وهما اختبار كولوموجروف - سميرونوف لعينة واحدة كذلك اختبار شابيرو ويلك Shapiro - wilk ولكن من محدداتهما تأثرهما

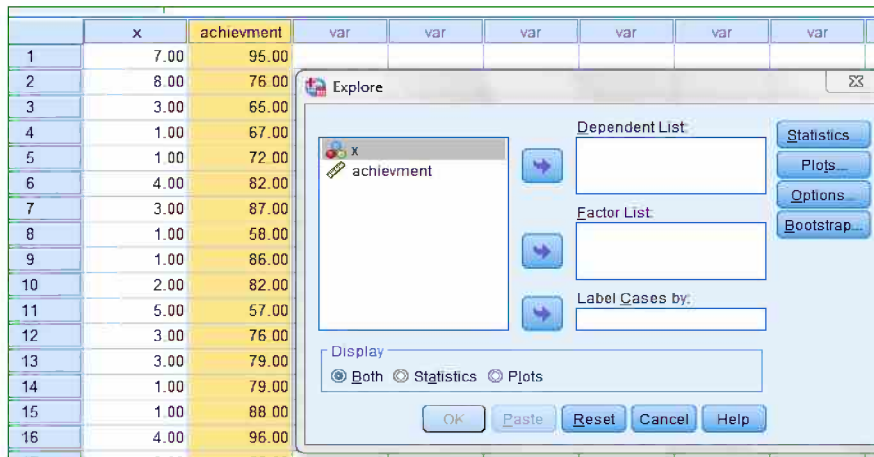
بأحجام العينات الكبيرة حيث من المتوقع مع حجم عينة كبيران نحصل على دلالة إحصائية حتى لو ابتعد توزيع درجات المتغير قليلا عن الاعتدالية.

الفروض الإحصائية: الفرض الصفري (H_0): توزيع درجات المتغير المتصل x اعتدالى .
الفرض البديل (H_A) : توزيع درجات المتغير غير اعتدالى.
وعدم الدلالة الإحصائية للاختبارين يعنى توافر الاعتدالية.

لتنفيذ الاختبارين فى برنامج الـ SPSS اتبع الآتى:

أولاً: إدخال البيانات: 1. اكتب مسمى المتغير X تحت عمود Name ، ثم اضغط على dataview

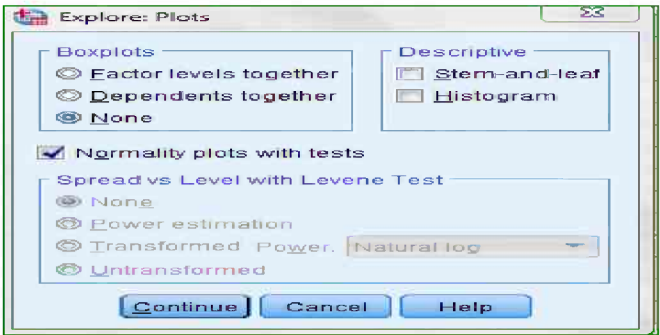
ثانياً : تنفيذ الامر : 1 . اضغط على Analyze ثم اضغط Descriptive Statistics ثم اضغط على Explore يعطى الشاشة الآتية:



2. انقل المتغير X إلى المربع Dependentlist عن طريق الضغط على السهم →

3. اضغط على اختيار Plots تظهر

الشاشة الآتية:



4. اضغط على اختيار Normality

Plots with tests وهذا يعطى رسم Q- QPlot

5. اضغط على Continue ثم اضغط Ok

ثالثاً: المخرج: يعطى إحصائيات:

			Statistic	Std. Error
x	Mean		3.2500	.23110
	95% Confidence Interval for Mean	Lower Bound	2.7914	
		Upper Bound	3.7086	
	5% Trimmed Mean		3.0111	
	Median		3.0000	
	Variance		5.341	
	Std. Deviation		2.31104	
	Minimum		1.00	
	Maximum		16.00	
	Range		15.00	
	Interquartile Range		3.00	
	Skewness		2.164	.241
	Kurtosis		8.447	.478

ثم عرض البرنامج الجدول الاتي للاختبارات :

Tests of Normality

	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	Df	Sig.
x	.283	100	.000	.788	100	.000

a. Lilliefors Significance Correction

وبالنسبة لاختبار K.S يتضح أن قيمته : $0.283 = K.S (100)$ وكانت درجات الحرية $df=100$ واتضح أن قيمة الاختبار دالة إحصائيا حيث: $P(Sig) (0.00) < \alpha(0.05)$ وتم رفض الفرض الصفري على ذلك توجد دلالة إحصائية وعليه فإن توزيع المتغير X غير اعتدالي. وهكذا بالنسبة لاختبار Shapiro-wilk حيث إن إحصائية أو قيمة الاختبار 0.788 و $df = 100$. وبالنسبة للقرار بما ان:

$$P(Sig) (0.00) < (0.05)$$

وعليه يتم رفض الفرض الصفري وبالتالي توجد دلالة إحصائية وعلى ذلك فإن توزيع درجات X غير اعتدالي.

ويمكن فحص الاعتدالية لكل متغير على حدة **Univariate normality** من خلال مؤشرى **الالتواء والتفرطح**. والالتواء وهو تراكم التوزيع على جانبي المنحنى وقد يكون الالتواء موجب **Positive skewness** وهو أن معظم درجات الأفراد تقع على الجانب الأيسر من المنحنى الاعتدالي، وعليه فإن معظم درجات الأفراد تقع تحت المتوسط بينما الالتواء السالب **Negative skew** يشير إلى أن معظم درجات الأفراد تقع على الجانب الأيسر من المنحنى الاعتدالي، وعليه فإن معظم درجات الأفراد تقع تحت المتوسط.

ولتفسير قيمة الالتواء والتفرطح فى ضوء قيم مطلقة لهم أو حدود قطع ، أشارت دراسات المحاكاة أن قيمة الالتواء (West, Finch & Curran, 1997) $SK > 3$ فإنه يوجد التواء شديد بينما التفرطح أكثر مرونة فالقيمة المطلقة أكبر من 7.0 تشير إلى تفرطح شديد (Chou & Bentler, 1995) وتوجد آراء أقل تشددًا فيما يخص التفرطح فإذا كانت قيمته أكبر من 8.0 وأحيانًا أكبر من 20 فإن التوزيع يوصف بأنه شديد التفرطح. ولكن القاعدة العامة هي إذا كان قيمة التفرطح أكبر من $< (10.0)$ فإنه توجد مشكلة فيما يخص التوزيع ، والقيمة 20 فأكثر تشير إلى توزيع يعانى بشدة من عدم توافر الاعتدالية (Kline 2016). ولمحاولة جعل البيانات اعتدالية تجرى عملية التحويل أو التعديل للبيانات من خلال تحويلات رياضية ولكن هذا ليس ضمانًا للحصول على الاعتدالية للمتغيرات . وكذلك استخدام إجراء Bootstrapping. وهو قائم على فكرة بسيطة وهو اخذ عينات متكررة من قاعدة البيانات الأصلية للتحقق من مدى وجود الاختلاف أو التباين الامبريقي فى النتائج. والباحث المحافظ يعتبر أن قيم الالتواء والتفرطح أقل من الواحد الصحيح حتى يكون التوزيع اعتدالى.

الفصل الثالث

تحليل التباين أحادى الاتجاه داخل الأفراد

(العينات المرتبطة)

Analysis Of Variance – One Way within Subjects

توجد مواقف بحثية تتطلب قياس نفس الأفراد على كل مستوى من مستويات المتغير المستقل بمعنى قياس نفس الفرد مرتين أو أكثر فى المتغير التابع وعلى ذلك فالباحث يختبر نفس العينة من الأفراد تحت شروط أو معالجات تجريبية مختلفة ويسمى هذا تصميم داخل الأفراد أو تصميم القياسات المتكررة، وفى تصميم القياسات المتكررة الدرجات على نفس الأفراد فى معالجات مختلفة أو مواقف مختلفة هى المتغير التابع بينما الدرجات للأفراد المختلفين هى المتغير المستقل.

ويرى (2012) Nolan & Heinzen أن تحليل التباين داخل المجموعات يستخدم عندما يوجد متغير مستقل اسمى أو ترتيبى وله مستويين فأكثر والمتغير التابع هو فترى درجات، وفيه يتعرض نفس الأفراد لكل مستويات المعالجة (المتغير الكيفى) ويقاس المتغير الكمى التابع فى كل معالجة وهذا المتغير الكيفى يشار اليه بعامل القياسات المتكررة أو عامل داخل المجموعات ويسمى المتغير الكمى بالمتغير التابع. ويطلق على تحليل التباين فى هذه الحالة تحليل التباين أحادى الاتجاه داخل الأفراد- One-Related Way Within Subject ANOVA. أو تحليل التباين المرتبطة Related ANOVA أو تحليل التباين داخل المجموعات Within – groups ANOVA. وأحادى الاتجاه تعنى وجود عامل واحد أو متغير مستقل وحيد وداخل الأفراد يشير إلى أن نفس المشاركين (الأفراد) تتم قياسات عليهم فى كل مستوى من مستويات المتغير المستقل.

وتصميم القياسات المتكررة ليست مرتبطة بالمنهج التجريبى فقط بل قد يطبق على دراسات غير تجريبية حيث يتم ملاحظه الفرد فى اوقات مختلفة، ونقص الاستقلالية بين

المعالجات يسبب مشكله خطيرة لو أنهم كما يفترض مسبقاً منفصلين عن بعضهم البعض.

وعلى ذلك فإن البيانات فى كل الخلايا غير مستقلة ويوجد صعوبة فى تحليل بياناتها هذه التصميمات وتصبح المعادلات أكثر تعقيداً لـ ANOVA .

معالجه 1	معالجه 2	معالجه 3
X_{11}	X_{21}	X_{31}
X_{12}	X_{22}	X_{32}
X_{1n}	X_{2n}	X_{3n}

وعلى ذلك فإن نفس الأفراد يتعرضون للمعالجات الثلاثة فالفرد الأول تكون درجاته على المعالجات الثلاثة هى X_{11} ، X_{21} ، X_{31} فبدلاً من وجود ثلاث عينات $3n$ فى حاله ANOVA المستقلة توجد عينه واحده تتعرض لكل المعالجات وفى هذه الحالة من الصعب أن نعتقد أن الارتباطات بين الثلاثة معالجات تساوى صفر فعلى العكس فأفضل الأفراد أداءً فى المعالجة الأولى سوف يكون أداؤه أفضل فى المعالجة 2 وفى المعالجة 3. والأفراد الذين اداؤهم ضعيفاً فى المعالجة 1 من المحتمل أن يكون اداؤهم ضعيفاً فى المعالجة 2 ، 3 وعلى ذلك فإنه توجد ارتباطات دالة إحصائياً بين المعالجات.

وتنشأ الفروق داخل المجموعات وليس فروق بين المجموعات حيث إن الفرد عالى الاداء فى معالجه ما يتوقع أن يكون كذلك فى معالجه اخرى وعلى ذلك فالفروق بين الأفراد ينتج عنها معظم الفروق داخل المعالجات أو المجموعات.

نسبه F لـ ANOVA للقياسات المتكررة:

نسبه F فى تحليل التباين للقياسات المتكررة هى نفسها كما فى حالة ANOVA المستقلة :

$$F = \frac{\text{التباين بين المعالجات} - F}{\text{التباين المتوقع إذا لم يوجد تأثير لمعالجه (الخطأ)}}$$

فمقام F يقيس فروق المتوسطات الفعلية فى حالة عدم وجود معالجات (الخطأ) وهو يحدد الفروق الموجودة فى حاله عدم وجود المعالجة والقيمة الكبيرة لـ F تشير إلى أن

الفروق بين المعالجات أكبر من الفروق المتوقعة بدون أى معالجة. كما نعلم أنه يوجد فروق فردية وهى تشير إلى خصائص الأفراد مثل الجنس- العمر - الشخصية حيث تختلف من شخص إلى آخر وهو بدوره يؤثر على القياسات التى نحصل عليها من كل فرد. والفروق الفردية هى جزء من البسط والمقام فى نسبة F المستقلة ولكن هذه الفروق تستبعد أو تحذف من التباينات فى F المرتبطة لأن نفس الأفراد يشتركوا فى نفس المعالجات المختلفة وعلى ذلك فإنه إذا وجدت فروق بين متوسطات المعالجات فإن هذا لا يرجع إلى الفروق الفردية وذلك لأنه يتم استبعادها من بسط F المرتبطة. وعلى ذلك فإن تصميم القياسات المتكررة يسمح باستبعاد الفروق الفردية من التباين فى بسط ومقام نسبة F .

$$F = \frac{\text{التباين بين المعالجات (بدون فروق فردية)}}{\text{التباين داخل المجموعات (بدون فروق فردية)}}$$

وعلى ذلك فإن صيغة F المرتبطة مثل F المستقلة ماعدا أنها لا تتضمن التباين الناتج عن الفروق الفردية تصبح:

$$F = \frac{\text{تباين بين المجموعات (المعالجات)}}{\text{تباين الخطأ}}$$

$$F = \frac{\text{تأثيرات المعالجات + فروق عشوائية أو غير منتظمة}}{\text{فروق عشوائية أو غير منتظمة}}$$

وعندما لا يوجد تأثير للمعالجة فإن البسط والمقام هو نفسه وعليه فإن : $F = 1.00$ وعندما تعطى نتائج البحث أو التجربة نسبة F قريبة من الواحد الصحيح فهذا دليل على عدم وجود تأثير للمعالجة ويفشل الباحث فى رفض H_0 . بينما القيم الكبيرة لـ F تشير إلى تأثير واضح للمعالجة أى رفض H_0 .

مكونات التباين فى ANOVA المرتبطة

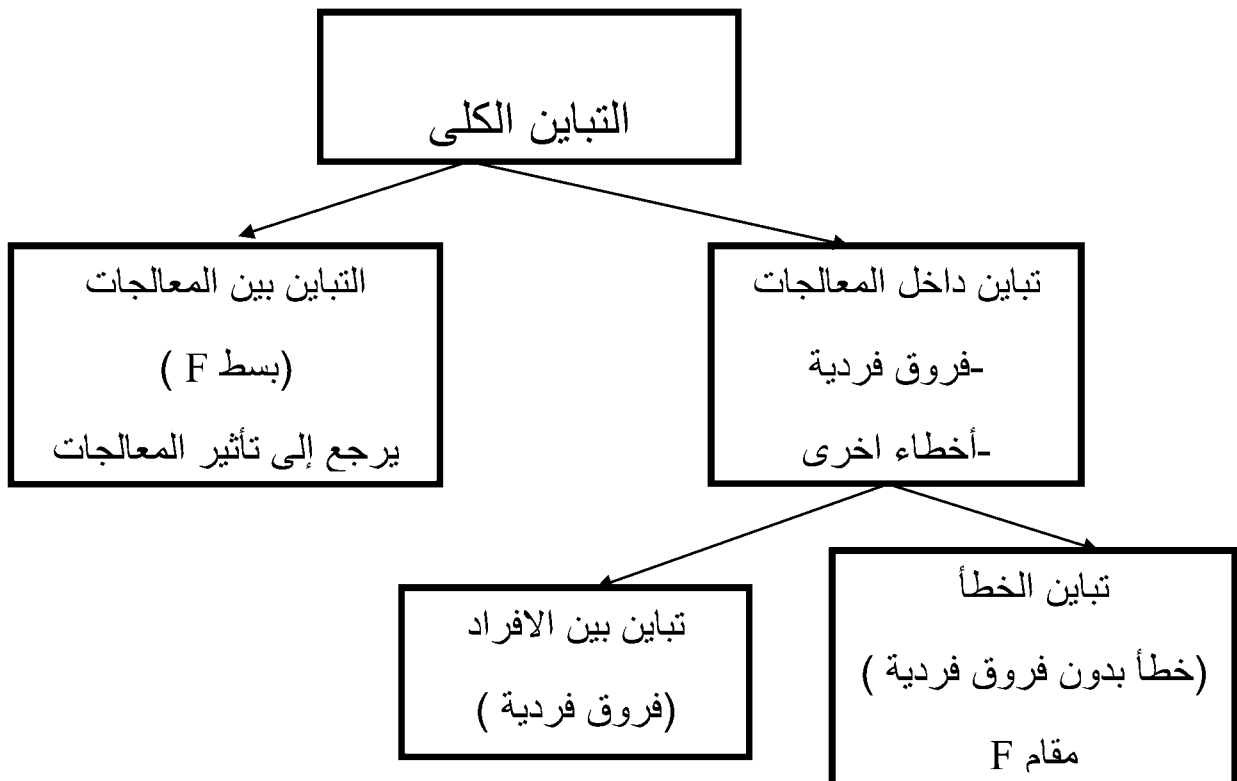
يمكن عرض التباين الكلى فى ANOVA المرتبطة كما هو موضح فى الجدول الآتى:

الجدول (1.3): مخطط لتصميم تحليل التباين المرتبط.

المعالجات				
تباين بين الأفراد ويتم حسابه من خلال متوسط درجه كل فرد فى الصف ويتم حذفه من مقام F				
المشاركين	A	B	C	متوسط الأفراد X
1	2	4	6	4
2	1	3	5	3
3	0	2	4	2
4	5	7	9	1
	$\bar{X} = 2$	$\bar{X} = 4$	$\bar{X} = 6$	

*التباين بين المجموعات:

ويمكن توضيحه بالآتى:



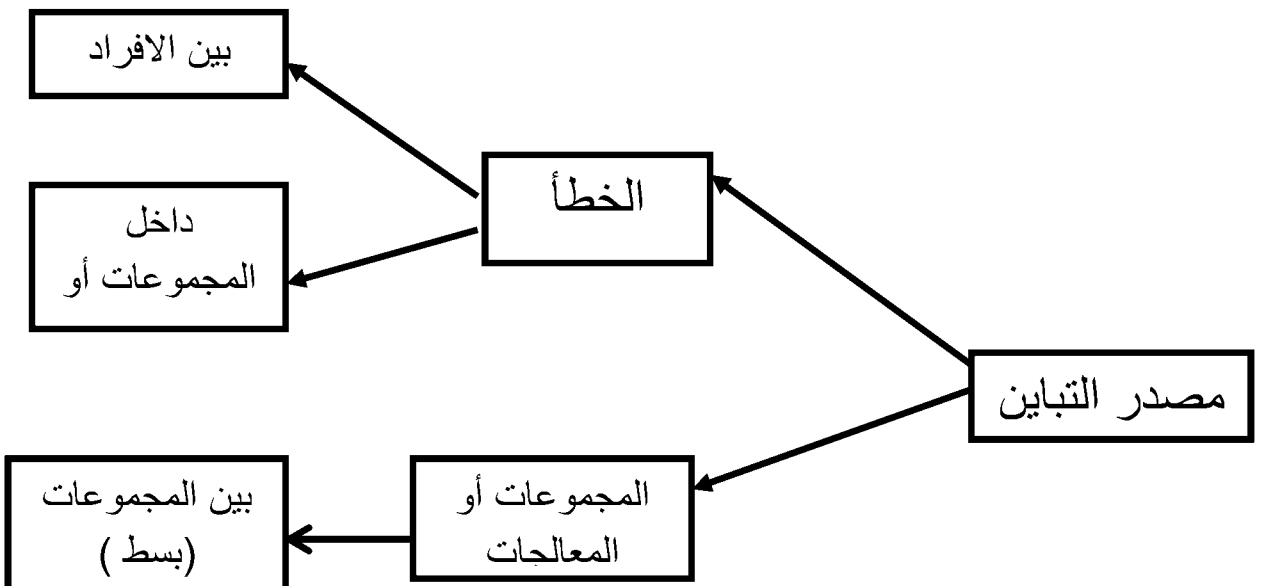
الشكل (1.3): مكونات التباين فى ANOVA قياسات متكررة.

1. التباين بين المعالجات: كما هو واضح فى جدول (1.3) وهو بين المعالجات A , B , C وهذا مشابه للتباين بين المجموعات فى ANOVA لقياسات مستقلة ويوضح فى بسط معادله F .

2. تباين الخطأ (داخل المعالجات): يوجد مصدرين من تباين الخطأ هما:

- تباين بين الاشخاص أو الأفراد: وهذا مرتبط بالفروق الفردية بين متوسطات الأفراد (الصف) عبر المجموعات، وعندما يتم قياس أفراد مختلفين فى كل مجموعة فإن الفروق الفردية أو خصائص الأفراد المختلفة تتغير عبر المجموعات وفى تصميم داخل المجموعات لا يوجد فروق فردية عبر المجموعات لأنه نفس الأفراد ولهذا السبب فإن التباين بين الأفراد ويتم حذفه من مقام F .

- تباين داخل المجموعات: يلاحظ نفس الأفراد فى كل مجموعة فإن الأفراد المختلفين مازال يتم ملاحظتهم داخل المجموعة ويعتبر هذا التباين مقام F . وعليه فإن التباين يرجع إلى تباين بين المجموعات أو المعالجات ومصدرين مرتبطين بالخطأ هما لنفس الأفراد عبر المجموعات وللأفراد داخل المجموعة الواحدة. وعلى ذلك يوجد مصدر تباين للخطأ اضافى عن مصدر تباين الخطأ فى تحليل التباين المستقلة ولكن هذا لا يعنى أن تحليل التباين للقياسات المتكررة يتضمن خطأ تباين أكبر من تحليل التباين للقياسات المستقلة. وفى تحليل تباين القياسات المرتبطة يتم حذف التباين بين الأفراد (المصدر الإضافى للخطأ) من مقام F . وعليه فإن :



الشكل (2.3): مصادر مختلفة من التباين في تصميم داخل المجموعات.

ويطلق على مقام F التباين داخل المجموعات (الأفراد) بتباين الخطأ أو تباين البواقي Residual Variance ويقيس كم نسبة التباين المتوقعة إذا لم يوجد تأثير منظم للتجربة. أو يطلق عليه التباين غير المفسر Unexplained Variance .

درجات الحرية في ANOVA للقياسات المتكررة

كما أوضحنا يوجد ثلاث مصادر للتباين وعليه فكل مصدر له درجة حرية خاصة به كما نعلم أن K عدد المجموعات أو المعالجات و n عدد أفراد عينة في كل مجموعة وعليه فإن:

- درجات الحرية المرتبطة بتباين بين المجموعات (SS_{Between}):

$$df_{B,G} = K - 1$$

- درجات الحرية المرتبطة بتباين بين الأفراد: $df_{B,i} = n - 1$

- B.G : Between Groups (بين المجموعات)

- B.I : Between individual (بين الأفراد)

- درجات الحرية المرتبطة بتباين الخطأ (داخل الأفراد):

$$df_w = (K - 1) (n - 1)$$

و df_w درجات الحرية الخطأ أو البواقي df_e أو df_{res}

- درجات الحرية المرتبطة بالتباين الكلي: $df_t = N - 1$

حيث N حجم العينة الكلي .

مسلمات اختبار ANOVA للقياسات المتكررة

هي نفس مسلمات ANOVA المستقلة:

1. الاعتدالية لبيانات المتغير التابع لكل مستوى أو مجموعة من مجموعات تصميم

أحادي العامل للقياسات المتكررة، ويرى (2014) Green & Salkind أن قيمة P

الاحتمالية المقبولة لحجم عينة متوسط هو 30 فردًا والقوة الإحصائية لهذا الاختبار تنقص مع نقصان حجم العينة.

2. الاستقلالية داخل المجموعات: لا تتأثر درجة فرد ما بدرجة فرد آخر ولكن توجد اعتماديه بين اداء نفس الفرد عبر كل المعالجات وتكون الثقة كبيره فى نتائج ANOVA المرتبطة إذا كانت درجات مستقلة.

3. تجانس التباينات: تباين مجتمع درجات كل مجموعة متساوى ويشار لهذه المسلمة بـ Sphericity Assumption، وإذا لم تتحقق هذه المسلمة فإنه يوجد طرق تكيفية أو تعديلية لتصحيح درجات الحرية كما هو الحالة فى اختيار Welch T. وعدم تحقق هذه المسلمة يحدث تضخم لبسط F بالتالى تضخم الخطأ من النوع الأول.

4. معاملات الارتباط فى المجتمع بين كل زوج من درجات المعالجات متساوي.

وإذا لم تتحقق المسلمة 3 ، 4 كما اوضحنا فإن قيمه الخطأ من النوع الأول تتضخم ويرتبط بهذه المسلمة تجانس مصفوفه التغاير .

وفلسفة التعامل مع ANOVA المرتبطة تتبع إحصاء النموذج المتدرج مثل MANOVA. وتنفيذها فى برنامج SPSS يكون من خلال General Linear Model Repeated Measures Procedure لأنها هى فلسفه الإحصاء المتدرج حيث إن المتغير المستقل متعدد المستويات بينما توجد عدة قياسات للمتغير التابع ويمكن اعتبار كل قياسه فى حد ذاتها متغير تابع وبالتالى إذا وجدت ثلاث معالجات فإنه يوجد ثلاث متغيرات تابعة. وعليه فهو اختبار يقع ضمن اختبارات النموذج البسيط فى المجلدات الإحصائية ولكن فى حقيقة تنفيذه الإحصائى وفلسفته فهو اختبار متدرج لأنه ملف SPSS يتضمن متغيرات تابعة عديدة حيث يدخل قياسات المتغير التابع فى المعالجة الأولى فى عمود ثم قياسات المتغير التابع فى المعالجة الثانية فى عمود ثانى وهكذا.

تطبيقات ANOVA المرتبطة فى البحوث

يستخدم هذا الأسلوب فى التصميمات أو الدراسات البحثية الآتية: الدراسات التجريبية، الدراسات شبه التجريبية، الدراسات غير التجريبية، والدراسات الطولية.

اختبارات الفروض لقضية بحثية (2014) Gravetter & Wallanu:

أراد باحث المقارنة بين أربعة استراتيجيات للإعداد للامتحان فى محتوى ما وهى القراءة، القراءة وإعادةؤها، حل أسئلة معدة مسبقاً، وإعداد أسئلة من عند الطالب والاجابة عليها ، وطبق ذلك على ستة طلاب واراد التحقق ماذا كانت توجد فروق بين درجات التحصيل بين الاستراتيجيات الاربعة ؟

خطوات اختبارات الفروض الصفرية:

1. الفروض الإحصائية:

$$H_0 : \mu_1 = \mu_2 = \mu_3 = \mu_4 ,$$

$$H_A : \mu_1 \neq \mu_2 \neq \mu_3 \neq \mu_4$$

• μ_1 ، μ_2 متوسط مجتمع العينة الأولى والثانية على التوالى.

2. الاختبار الإحصائى ومسلماته: هو اختبار ANOVA للقياسات المتكررة، حيث

$$F = \frac{MS_{\text{Between}}}{MS_{\text{Error(Res)}}}$$

3. مستوى الدلالة الإحصائية وقاعدة القرار: تبنى الباحث $\alpha = 0.01$ وللكشف عن القيمة

الحرية لـ F يتم حساب درجات الحرية كالاتى:

• درجات الحرية بين المعالجات:

$$df_{B,G} = K - 1 = 4 - 1 = 3$$

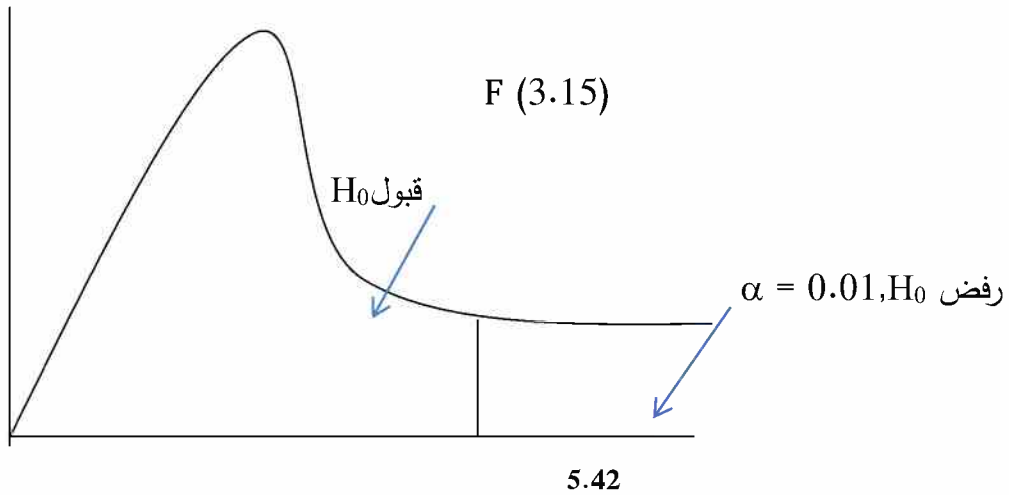
• درجات الحرية بين الأفراد أو الاشخاص:

$$df_{B,I} = n - 1 = 6 - 1 = 5$$

• درجات الحرية الخطأ أو البواقي داخل المجموعات:

$$\begin{aligned} df_{w(\text{error})} &= (k - 1)(n - 1) = df_{BG} * df_{B,I} = \\ &= 3 \times 5 = 15 \end{aligned}$$

بالبحث في جدول توزيع F ب $\alpha = 0.01$ و $df_{B,G} = 3$ البسط و $df_{error} = 15$ المقام فإن $F_{C.V} = 5.42$ ، وعليه فإن:



4. الحسابات:

CNO	قراءة	قراءة واعادتها	حل أسئلة معدة	إعداد أسئلة وحلها	مجموع الأفراد
A	3	5	8	8	P=24
B	3	3	5	9	P=20
C	4	3	8	7	P=24
D	6	7	9	10	P=32
E	6	8	8	10	P=32
F	8	8	10	10	P=36
$T = \sum X_i = 30$	$T = 36$	$T = 48$	$T = 54$		
$\bar{X}_1 = 5$	$\bar{X}_2 = 6$	$\bar{X}_3 = 8$	$\bar{X}_4 = 9$		
$SS_1 = 20$	$SS_2 = 20$	$SS_3 = 20$	$SS_4 = 8$		

وتحسب ANOVA بمرحلتين: (n عدد الاشخاص في كل مجموعة = 6 ، K عدد

المعالجات = 4 ، N عدد الاشخاص في كل المجموعات = 24:

وحيث يتم حساب الآتي:

- مجموع كل الدرجات في المعالجات الأربعة $\sum X = \sum T$:

$$\sum T = 3+3+4+6+6+8+5+3+....+10+10+10=168$$

- مجموع مربع كل الدرجات

$$\sum X^2,$$

$$= (3)^2 + (3)^2 + (4)^2 + (6)^2 + (6)^2 + (8)^2 + (5)^2 + (3)^2 + (10)^2 + (10)^2 + (10)^2 + (10)^2 = 1298$$

- مجموع درجات كل فرد في المعالجات الأربعة:

$$P_1 = 3+5+8+8+8=24$$

$$P_2 = 5+3+5+7+8+8=20$$

$$P_3 = 8+5+8+9+8+10=32$$

$$P_4 = 8+9+7+10+10+10=36$$

- مجموع درجات كل معالجة T :

$$T_1 = 3+3+4+6+6+8=30$$

$$T_2 = 5+3+5+7+8+8=36$$

$$T_3 = 8+2+8+9+8+10=48$$

$$T_4 = 8+9+7+10+10+10=54$$

- متوسط كل مجموعة: $\bar{x}_1=5, \bar{x}_2=6, \bar{x}_3=8, \bar{x}_4=9$

ولحساب نسبة F فلا بد من حساب أولاً:

أولاً: مجموع المربعات SS وهي:

- مجموع المربعات الكلى:

$$SS_T = \sum X^2 - \frac{\sum T^2}{N}$$

$$= 1298 - \frac{(168)^2}{24} = 122 - 1176 = 122$$

و درجات الحرية المرتبطة بها: $df_T = N-1 = 24-1 = 23$

• حساب مجموع المربعات بين المعالجات:

$$SS_{B.G} = \sum \left(\frac{T_1^2}{n} + \frac{T_2^2}{n} + \frac{T_3^2}{n} - \frac{(\sum T)^2}{N} \right)$$

$$= \left(\frac{(30)^2}{6} + \frac{(36)^2}{6} + \frac{(48)^2}{6} + \frac{(54)^2}{6} \right) - \frac{(168)^2}{24} = 60$$

$$df_{B.G} = K - 1 = 4 - 1 = 3$$

• حساب مجموع المربعات بين الأفراد:

$$SS_{B.I} = \sum \frac{P_i^2}{K} - \frac{(\sum T)^2}{N}$$

$$= \frac{(24)^2}{4} + \frac{(20)^2}{4} + \frac{(24)^2}{4} + \frac{(32)^2}{4} + \frac{(32)^2}{4} + \frac{(36)^2}{4} - \frac{(168)^2}{24}$$

$$= 144 + 100 + 144 + 256 + 324 + 324 - 1176 = 48$$

$$df_{B.I} = 6 - 1 = 5 \text{ و}$$

• مجموع مربعات البواقي أو الخطأ:

$$SS_{Res(error)} = SS_T - SS_{B.G} - SS_{B.I}$$

$$= 122 - 60 - 48 = 14$$

ثانيًا : حساب متوسط مجموع المربعات:

• متوسط مجموع المربعات بين المعالجات أو المجموعات:

$$MS_{B.G} = \frac{SS_{B.G}}{df_{B.G}} = \frac{60}{3} = 20$$

• متوسط مجموع المربعات بين الأفراد أو الاشخاص:

$$MS_{B.I} = \frac{SS_{B.I}}{df_{B.I}} = \frac{48}{5} = 9.6$$

• متوسط مجموع مربعات البواقي أو الخطأ:

$$MS_{e(Res)} = \frac{S_{error}}{df_{error}} = \frac{14}{15} = 0.933$$

ثالثًا: حساب نسبة F:

$$F = \frac{MS_{BG}}{MS_{Res}} = \frac{20}{0.933} = 21.43$$

ونلاحظ أن البسط (الفروق بين المعالجات) هي أكبر 22 مرة من الفروق الموجودة بفرض عدم وجود معالجات (الصدفة) وهذا يعطى وجود تأثير حقيقى قوى للمعالجة.

5. القرار والتفسير: وبما أن نسبة $F (21.43) < F_{c.v}$ الدرجة 5.42، وعليه يرفض الباحث الفرض الصفري ومن ثم توجد فروق ذات دلالة إحصائية بين متوسطات أو درجات التحصيل بين الاستراتيجيات الأربعة.

6. حجم التأثير: يحسب حجم التأثير لـ F للقياسات المتكررة كالتالى:

• مؤشر η_p^2 :

$$\eta_p^2 = \frac{SS_{B.G \text{ بين المعالجات}}}{SS_{B.G \text{ (المعالجات)}} + SS_{Res(error)}}$$

$$\eta_p^2 = \frac{60}{60+14} = \frac{60}{74} = 0.811$$

وهذا يعنى أن 81.1% من التباين فى البيانات يرجع إلى الفروق بين المعالجات أو للاستراتيجيات المستخدمة.

• مؤشر ω_p^2 (Omega-Squared Partial): مؤشر η_p^2 يعتبر تقدير

عالى Overestimate لحجم التأثير لكن مؤشر ω_p^2 أكثر تحفظاً من η_p^2

ويتم تقديره كالتالى:

$$\omega_p^2 = \frac{SS_{B.G} - df_{B.G}(MS_{error})}{(SS_{B.G} + SS_{error}) + (MS_{error})}$$

$$\omega_p^2 = \frac{60-3 (0.933)}{(60+14)+0.933} = 0.79$$

وبالتالى:

وعليه 79% من تباين تحصيل الموضوع يرجع إلى الاستراتيجيات المستخدمة.

وعرض (2012) Huck أهم مؤشرات حجم التباين ANOVA:

المؤشر	صغير	متوسط	كبير
η^2	0.01	0.06	0.14
ω^2	0.01	0.06	0.14
η_p^2	0.01	0.06	0.14
ω_p^2	0.01	0.06	0.14
Cohen F	0.10	0.25	0.40
r (η)	0.10	0.24	0.37

كتابة نتائج ANOVA المتكررة في تقرير البحث وفقاً APA

بحساب التباينات والمتوسطات لدرجات تحصيل الاستراتيجيات الاربعة كما في الجدول التالي. واطهر تحليل التباين للقياسات المتكررة وجود دلالة إحصائية بين متوسطات الاستراتيجيات الاربعة حيث $F(3,15)=21.43$, $P<0.01$, $\eta^2 = 0.811$

الاستراتيجية الأولى	الثانية	الثالثة	الرابعة	
5	6	8	9	$\bar{X}(M)$
2	2	1.67	1.26	SD

ولكن السائد في تقارير البحوث يتم عرض جدول تحليل التباين كالآتي:

المصدر	SS	df	MS	F	Fc.v
بين المعالجات	60	3	20		
بين الأفراد	48	5		21.43	5.42
داخل المعالجات	62	20			
الخطأ	14	15	.933		
الكلي	122	23			

• دالة عند 0.01

إجراء المقارنات البعدية لـ ANOVA المتكررة:

بعد أن أظهر اختبار F دلالة إحصائية فروق بين متوسطات الاستراتيجيات الأربعة وهذا يعنى وجود فروق بين متوسطين على الأقل من المتوسطات الأربعة ولتحديد أى فروق دالة إحصائياً بين المجموعات الأربعة فإن الباحث يجرى تحليل تتبعى لـ F ويسمى بالمقارنات المتعددة. ولاختبار ANOVA المتكررة يستخدم اختبار توكى (HSD) وشيفيه كما سبق استخدامه فى ANOVA للقياسات المستقلة.

فى المثال السابق: $n=6$, $df_{error}=15$, $MS_{error}=0.933$ وعليه:

$$HSD = Q \sqrt{\frac{MS_{error}}{n}}$$

حيث Q هى قيمة Q الجدولية لـ $df=15$ و $n=5$:

$$= 4.08 \sqrt{\frac{0.933}{6}} = 4.08 (0.394) = 1.61$$

وإى فرق بين أى مجموعتين أكبر من 1.61 يكون دال إحصائياً عند 0.05 وبالتالى:

$$\bar{X}_2 - \bar{X}_1 = 6 - 5 = 1 \text{ غير دال إحصائياً و } \bar{X}_3 - \bar{X}_1 = 8 - 5 = 3 \text{ دال إحصائياً}$$

$$\bar{X}_4 - \bar{X}_1 = 9 - 5 = 4 \text{ دال إحصائياً و } \bar{X}_3 - \bar{X}_2 = 5 - 6 = -2 \text{ دال إحصائياً}$$

تحليل التباين للقياسات المتكررة واختبار T المرتبطة

عند تقييم الفروق بين متوسطى عينتين يستخدم اختبار T أو ANOVA واختبار T و ANOVA يشتركا فى مظاهر عديدة منها:

- يصلوا إلى نفس الاستنتاج أو القرار عن الفرض الصفري
- العلاقة الأساسية بين T و F هى: $F = T^2$
- درجة الحرية لاختبار T مشابهة مع درجة الحرية فى مقام F

- مربع القيمة الحرجه لـ T تساوى القيمة الحرجة لـ F
وفيما يلى مثال لتوضيح كيفية حسابه باستخدام T و F المرتبطة:

الأفراد	المعالجة I	المعالجة II	D
A	3	5	2
B	4	14	-10
C	5	7	-2
E	4	6	-2

فالفروض الإحصائية لـ T:

$$H_0: \mu_D(\mu_1 - \mu_2) = 0$$

$$H_A: \mu_D \neq 0$$

وبالتالى الاختبار ذو ذيلين و $df = n - 1 = 4 - 1 = 3$ ولمستوى دلالة الإحصائية $\alpha = 0.05$ فان:
 $T_{c.v} = 3.182$ وعليه فإن T المحسوبة:

$$T = \frac{\bar{x}_D}{s_D} = \frac{4}{2} = 2$$

وعليه يتم قبول H_0

وفى حالة ANOVA للقياسات المتكررة:

الأفراد	المعالجة I	المعالجة II
A	3	5
B	4	14
C	5	7
E	4	6

وعليه فإن : $\sum T = 48$ ، $\sum X^2 = 372$ و $N = 8$ ، و تكون الفروض الإحصائية:

$$H_0: \mu_1 = \mu_2 \quad , \quad H_A: \mu_1 \neq \mu_2$$

$$\text{و } df_{\text{error}} = 4 - 1 = 3 \quad , \quad df_{B.i} = 4 - 1 = 3 \quad , \quad df_{B.G} = 2 - 1 = 1$$

وبالتالى df تكون (3,1) وبالبحت فى توزيع F نجد أن قيمتها الحرجة:

$$F_{C.V}(0.05) = 10.13$$

$$F_{C.V} = (T_{C.V})^2 = (3.182)^2 = 10.13$$

وللبينات: $SS_T = 84$ و $SS_{B.G} = 32$ و $SS_{B.i} = 28$

$$SS_{error} = SS_T - SS_{B.i} - SS_{B.G} = 84 - 32 - 28 = 24$$

$$MS_{B.G} = \frac{32}{1} = 32 \quad \text{وبالتالى:}$$

$$MS_{error} = \frac{24}{3} = 8$$

$$F = \frac{32}{8} = 4 \quad \text{إذا:}$$

وبالتالى فإن: $(2)^2 F(4) = T^2$

ويتم قبول H_0 لأن F المحسوبة (4) $>$ الجدولية (10.13) وهو نفس القرار فى حالة T المرتبطة.

مدخل ANOVA المرتبطة لـ (Howell 2013):

قضية: استخدم باحث الاسترخاء لتخفيف الصداع وكانت البيانات كالاتى: " مدة استمرار الصداع (المتغير التابع)"

الفرد	بعد التدريب					متوسطات الأفراد
	الأسبوع الأول	الثاني	الثالث	الرابع	الخامس	
1	21	22	8	6	6	12.6
2	20	19	10	4	4	11.4
3	17	15	5	4	5	9.2
4	25	30	1.3	12	17	19.4
5	30	27	18	8	6	16.8
6	19	27	8	7	4	13.0
7	26	16	5	2	5	10.0
8	17	18	8	1	5	9.8
9	26	24	14	8	9	16.2
متوسط الأسبوع	22.333	22.00	9.333	5.778	6.778	13.249

• الفروض الإحصائية: $H_0 : \mu_1 = \mu_2 = \mu_3 = \mu_4 = \mu_5$

$H_A : \mu_1 \neq \mu_2 \neq \mu_3 \neq \mu_4 \neq \mu_5$

• الحسابات :

$$\begin{aligned}
 SS_{\text{total}} &= \sum (x - \bar{x}_i)^2 \\
 &= (21 - 13.44)^2 + (20 - 13.44)^2 + \dots + (22 - 13.44)^2 + \\
 &\dots + (8 - 13.44)^2 + (6 - 13.44)^2 + \dots + (6 - 13.44)^2 + \dots + \\
 &\quad (5 - 13.44)^2 + (9 - 13.44)^2 = 3166.31
 \end{aligned}$$

$$\begin{aligned}
 SS_{\text{sub(B.I)}} &= W \sum (\bar{x}_s - \bar{x}_i)^2 = 5[(12.6 - 13.244 + \dots + (16.2 - \\
 &\quad 13.44)^2] = 486.71
 \end{aligned}$$

حيث $\bar{x}_s$: متوسط درجات الاشخاص ، $5=W$ عدد المعالجات (الأسابيع).

$$\begin{aligned}
 SS_{\text{weeks(B.G)}} &= n \sum (\bar{x}_w - \bar{x})^2 = 9(22.333 - 13.244)^2 + \\
 &\quad \dots + (6.778 - 13.244)^2 = 2449.20
 \end{aligned}$$

حيث n عدد الأفراد ، $\bar{x}_w$ متوسط درجات الأسبوع و

$$SS_{\text{Res(error)}} = SS_{\text{total}} - SS_{B.I} - SS_{B.G}$$

$$= 3166.31 - 486.71 - 2449.20 = 230.40$$

$$MS_{w(B.G)} = \frac{SS_{B.G}}{df} = \frac{2449.20}{5-1} = 612.30 \quad \text{و}$$

$$MS_{\text{error}} = \frac{SS_{\text{error}}}{(k-1)(n-1)} = \frac{230.40}{4 \times 8} = 7.20$$

$$F = \frac{MS_{B.G}}{MS_{\text{error}}} = \frac{612.30}{7.20} = 85.04 \quad \text{ومن ثم فإن:}$$

وبالبحث في جداول F بـ $df = 4$ البسط، $df = 32$ المقام و $\alpha = 0.05$ ، فإن F الجدولية أو الحرجة $F_{c.v} = 2.68$ ، وعليه فإن :

$$F \text{ المحسوبة } (85.04) < F_{c.v} (2.68) \text{ وبالتالي نرفض } H_0$$

وملخص الجدول:

F	MS	SS	df	مصدر التباين
		486.71	8(n-1)	بين الأفراد
		2679.609	36	داخل الأفراد
85.04	612.30	2449.20	4	بين المعالجات
	7.20	2350.40	32	الخطأ
		3166.31	44	الكلي

المقارنة المتعامدة Constrasts Comparison

وهذه تدخل في المقارنات المتعددة المخطط لها. وفي هذه الحالة افترض الباحث مقارنة

الأسبوعين الأولين μ_1, μ_2 مع باقى الاسبوع الثلاثة وعليه فإن:

$$\mu_1 \neq \mu_2 \neq \mu_3 + \mu_4 + \mu_5$$

إذا كان عدد معاملات $\mu_1, \mu_2 = 2$ ، وعدد المتوسطات الباقية = 3

$$\frac{1}{2}\mu_1 \neq \frac{1}{2}\mu_2 \neq \frac{1}{3}\mu_3 + \frac{1}{3}\mu_4 + \frac{1}{3}\mu_5 \quad \text{ومن ثم فإن:}$$

$$\frac{1}{2} + \frac{1}{2}\mu_2 - \frac{1}{3}\mu_3 - \frac{1}{3}\mu_4 - \frac{1}{3}\mu_5 = 0 \quad \text{وعليه فإن:}$$

الأسبوع الأول	الثاني	الثالث	الرابع	الخامس	
المعاملات	$\frac{1}{2}$	$\frac{1}{3}$	$\frac{1}{3}$	$-\frac{1}{3}$	
المتوسطات	22.333	22.000	9.333	5.778	6.778

وتصبح المقارنات :

$$\hat{\Psi} = \sum a_i \bar{x}_i$$

حيث a_i هي المعامل، $\bar{x}_i$ متوسط الأسبوع

$$\begin{aligned}
 \hat{\Psi} &= \frac{1}{2} (22.333) + \frac{1}{2} (22.00) - \frac{1}{3} (9.333) - \frac{1}{3} (5.778) - \frac{1}{3} (6.778) \\
 &= \frac{22.333 + 22.00}{2} - \frac{9.333 + 5.778 + 6.778}{3} \\
 &= \frac{44.333}{2} - \frac{21.889}{3} \\
 &= 22.160 - 7.296 = 14.870
 \end{aligned}$$

ويمكن اختبار هذه المقارنة للأسبوع الأول والثاني في مقابل الأسابيع الثلاثة الباقية من خلال اختبار F أو T ولكن يتم استخدام T:

$$\begin{aligned}
 T &= \frac{\hat{\Psi}}{\sqrt{\frac{\sum (a_i)^2 MS_{Res}}{n}}} \\
 &= \frac{14.870}{\sqrt{\frac{0.833 (7.20)}{9}}} = \frac{14.870}{\sqrt{0.667}} = 18.21
 \end{aligned}$$

وحيث قيمة T المحسوبة أكبر من قيمتها الجدولية نلاحظ أنها دالة إحصائية.

بالتالى تمت مقارنات الأسبوع الأول والثاني معاً لوحده (قبل التدريب) بالأسابيع الثلاثة اللاحقة (الثالث - الرابع - الخامس) كوحده معاً واتضح وجود دلالة إحصائية.

$$\begin{aligned} &= \frac{\hat{\Psi}}{S_{\text{error}}} \hat{d} \\ &= \frac{14.870}{\sqrt{7.20}} = \frac{14.870}{4.2} = 3.5 \end{aligned}$$

حجم التأثير:

وهذا حجم تأثير كبير جدًا وفقًا لمعايير Cohen(1988).

كتابة النتائج وفقًا لـ APA

أظهر تحليل التباين للقياسات المتكررة فروق داله بين الأسابيع الخمسة $F(4,32) = 85.04$, $P < 0.05$. وان متوسط الأسبوعين قبل المعالجة $\left(\frac{22+22.33}{2}\right)$ 22.166 وان متوسط زمن الصداغ بعد المعالجة $\left(\frac{0.77+5.778+9.33}{3}\right)$ واتضح وجود فروق بين المتوسطين معا (قبل وبعد) دالة إحصائية $P < 0.05$, $T(32) = 18.21$ وتم حساب حجم التأثير $d = 4.95$ وهذا يؤكد أهمية العلاج بالاسترخاء لمعالجة الصداغ .

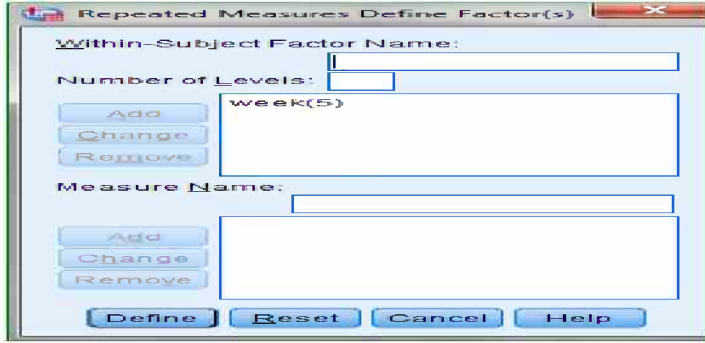
إجراء تحليل التباين الأحادي للقياسات المتكررة المرتبطة في SPSS

أولاً: مرحلة إدخال البيانات: 1. اضغط على Variable view.

2. تحت عمود Name أكتب مسمى المتغيرات كالاتى: الأسبوع الأول Weak1، الأسبوع الثانى Weak2، الأسبوع الثالث Weak3، الأسبوع الرابع Weak4، الأسبوع الخامس Weak5.

3. اضغط على Dataview يظهر الشاشة بها خمسة أعمدة لخمس متغيرات، إبدأ فى إدخال البيانات.

ثانيًا: تنفيذ الامر: 1. اضغط Analyze ثم Model Linear General ثم Repeated Measures تظهر الشاشة الاتية:



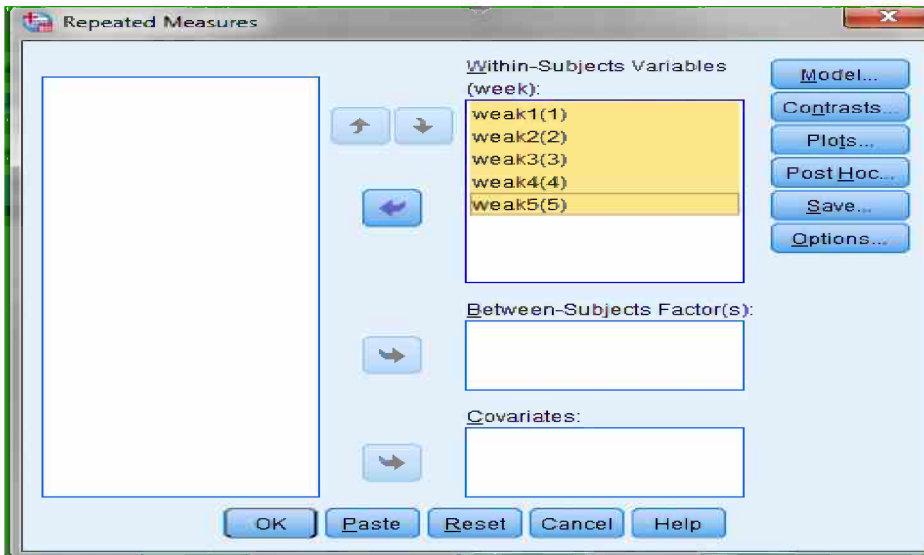
2. فى المربع أمام Within-subject، أكتب فى هذا المربع اسم العامل وهو Weak.

3. أمام Number of levels أكتب عدد المستويات = 5

4. اضغط Add.

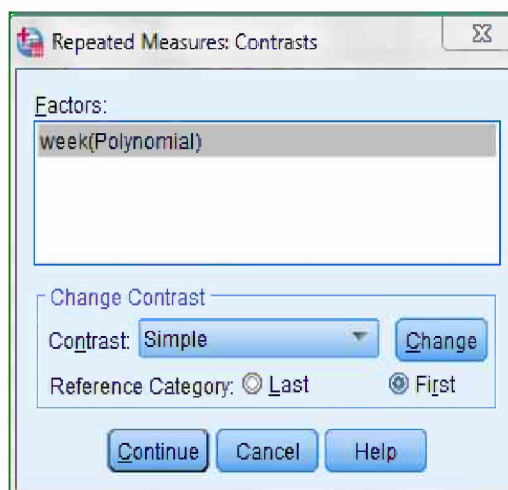
5. يمكن كتابة مسمى المتغيرات التابع فى مربع Measure Name (اختيارى).

6. اضغط Define أسفل الشاشة تظهر الشاشة الآتية:



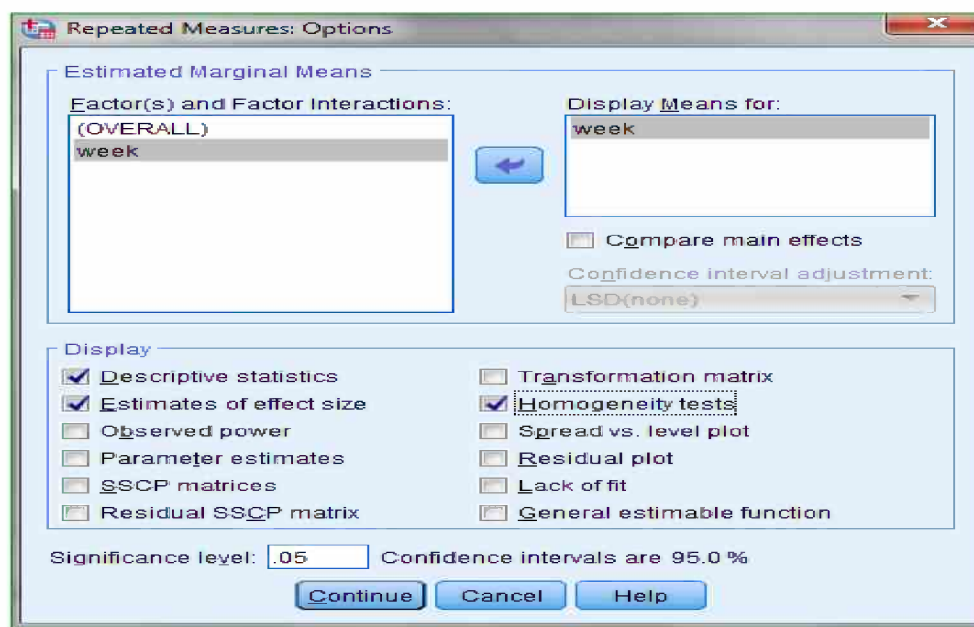
7. انقل المتغيرات الخمسة إلى مربع Within Subject variable ويمكن أن تضغط على مفتاح Ctrl وتنقلهم مرة واحدة.

8. اضغط على Contrast وحدد نوع المقارنات التي ترغبها وهذه تتضمن المقارنات البعدية المخططة لها ويمكن أن يقارن بين المجموعات الاربعة بالمجموعة الأولى على أساس أنها هي المجموعة المرجعية ولتنفيذ ذلك:



- اضغط contrast تظهر الشاشة:
- اضغط على السهم لأسفل امام Contrast تظهر عدد البدائل.
- اختار Simple.
- في Reference category اختار First.
- اضغط Change. ثم اضغط Continue.

9. اضغط اختيار Options تظهر الشاشة الآتية:

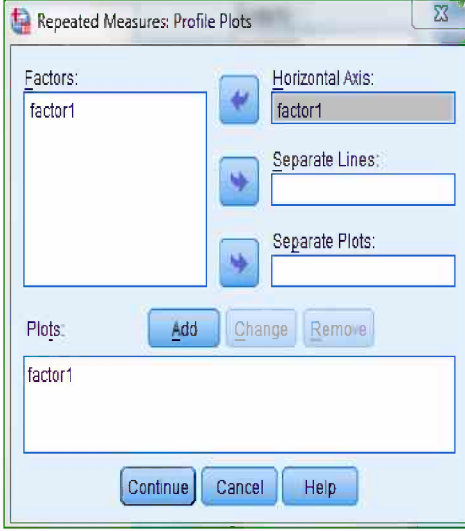


10. انقل Weak إلى مربع Display Mean.

11. تحت مربع Display اختار Descriptive statistics، Estimated of effect size، ثم Homogeneity test.

12. نشط compare main effects

13. تحت confidence interval adjustment اضبط على السهم واختار المقارنة المتعددة sidak



14. اضبط على Continue.

15. اضبط على اختبار Plots تظهر الشاشة الآتية:

16. انقل Weak إلى مربع Horizontal Axis.

17. اضبط Add.

18. اضبط Continue ثم اضبط على Ok.

ثالثا: المخرج:

Descriptive Statistics			
	Mean	Std. Deviation	N
weak1	22.3333	4.58258	9
weak2	22.0000	5.33854	9
weak3	9.3333	3.39116	9
weak4	5.7778	3.41971	9
weak5	6.7778	4.11636	9

● الإحصائيات الوصفية كالاتي:

اعطى الإحصاء الوصفى مثل المتوسط والانحراف المعياري لقياسات الصداع فى كل أسبوع. ومن الواضح أن أدنى متوسط للأسبوع الخامس وأعلامهم

الأسبوع الأول وكذلك للثاني كما أعطى الأفراد فى كل قياس وهم نفس الأفراد.

● اختبار تجانس التباين كالاتي:

وما يهمنا فى هذا الجدول وهو قيمة Sig (P) وهى 0.811 وعلى ذلك يتحقق الفرض الصفري تساوى التباينات أو بكلمات أخرى تتطلب أن التباينات بين درجات فروق القياسات متساوية وهذا يتم التحقق منه باستخدام اختبار Mauchly.

Mauchly's Test of Sphericity ^a							
Measure: MEASURE_1							
Within Subjects Effect	Mauchly's W	Approx. Chi-Square	df	Sig.	Epsilon ^b		
					Greenhouse-Geisser	Huynh-Feldt	Lower-bound
week	.282	8.114	9	.537	.684	1.000	.250
Tests the null hypothesis that the error covariance matrix of the orthonormalized transformed dependent variables is proportional to an identity matrix.							
a. Design: Intercept Within Subjects Design: week							
b. May be used to adjust the degrees of freedom for the averaged tests of significance. Corrected tests are displayed in the Tests of Within-Subjects Effects table.							

وفى بيانات المثال الحالى تحققت مسلمة استخدام تحليل التباين للقياسات المتكررة إذا لم تتحقق يوجد تصحيح لذلك من خلال الجزء الأخير من الجدول عن طريق تصحيحين هما Greenhouse-Geisser و Huynh-Feldt.

● جدول Multivariate test:

Multivariate Tests ^a							
Effect		Value	F	Hypothesis df	Error df	Sig.	Partial Eta Squared
week	Pillai's Trace	.986	86.391 ^b	4.000	5.000	.000	.986
	Wilks' Lambda	.014	86.391 ^b	4.000	5.000	.000	.986
	Hotelling's Trace	69.113	86.391 ^b	4.000	5.000	.000	.986
	Roy's Largest Root	69.113	86.391 ^b	4.000	5.000	.000	.986

a. Design: Intercept
Within Subjects Design: week
b. Exact statistic

حقيقة أن اختبار ANOVA للقياسات المتكررة يتبع النموذج المتعدد المتدرج حيث مرات القياس (الزمن) متغير مستقل والقياسات الخمسة متغيرات تابعة ولذلك فإن هذا الجدول يبين الفروق فى الاتحاد الخطى للقياسات الخمسة (تخليق متغير جديد) لمرات القياس الخمسة معًا واعطى اختبارات دلالة إحصائية كثيرة هى Pillai's trace و

Lambda Wilks وغيرها. وما يهمنا في هذا الجدول قيمة Sig لاختبارين Pillai's و Wilks وهي واحدة 0.00 وعليه $0.05 < 0.00$. وعليه توجد فروق ذات دلالة إحصائية بين الأسابيع الخمسة في الاتحاد الخطي للقياسات الخمسة معًا. ولكن هذا ليس هدف الباحث الأساسي وإنما هدفه هو دراسة الفروق بين القياسات الخمسة على حدة عبر مرات القياس ولذلك فإن البرنامج عرض الجدول الآتي:

Tests of Within-Subjects Effects							
Measure: MEASURE_1							
Source		Type III Sum of Squares	df	Mean Square	F	Sig.	Partial Eta Squared
week	Sphericity Assumed	2449.200	4	612.300	85.042	.000	.914
	Greenhouse-Geisser	2449.200	2.738	894.577	85.042	.000	.914
	Huynh-Feldt	2449.200	4.000	612.300	85.042	.000	.914
	Lower-bound	2449.200	1.000	2449.200	85.042	.000	.914
Error(week)	Sphericity Assumed	230.400	32	7.200			
	Greenhouse-Geisser	230.400	21.903	10.519			
	Huynh-Feldt	230.400	32.000	7.200			
	Lower-bound	230.400	8.000	28.800			

الجزء الأول من المخرج هو الهدف الرئيسي للباحث حيث عرض تأثير الزمن في حالة ما إذا تحققت مسلمة Sphercity وكذلك في حالة عدم تحققها عن طريق اختبارات Greenhouse-Geisser وغيرها ولكن في البيانات الحالية تحققت المسلمة.

وبالتالي فإن مجموع المربعات $2449.2 = \text{Sum of squares}$ ودرجات الحرية $df = 4 = 5 - 1$ وقيمة $F = 53.25$ وأن $\text{Sig (P)} = 0.000$ إذا $0.05 < 0.000$ وعليه نقبل الفرض البديل بوجود فروق بين القياسات الخمسة للصداع، وأعطى البرنامج حجم التأثير $= 0.914$. لاحظ نتائج Greenhouse-Geisser و Huynh Feildt واحدة ويحدث اختلاف بينهما في حالة عدم تحقق مسلمة تساوى التباينات بين فروق الدرجات.

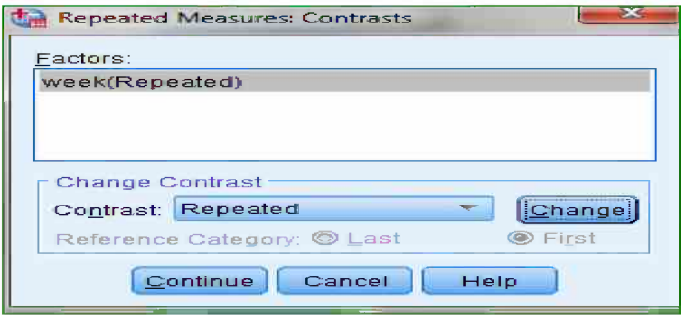
حقيقة أن الاعتماد على F MANOVA لا يجعلنا في حيرة من أمرنا إذا لم تتحقق التجانس للتباينات حيث إنه يدرس الفروق في الاتحاد الخطي للقياسات الخمسة معًا،

ثم أعطى البرنامج Contrast بين متوسط المجموعة بالمجموعة الأولى على اعتبار أنها الضابطة كالآتى:

Tests of Within-Subjects Contrasts							
Measure: MEASURE_1							
Source	week	Type III Sum of Squares	df	Mean Square	F	Sig.	Partial Eta Squared
week	Level 2 vs. Level 1	1.000	1	1.000	.038	.849	.005
	Level 3 vs. Level 1	1521.000	1	1521.000	108.643	.000	.931
	Level 4 vs. Level 1	2466.778	1	2466.778	144.868	.000	.948
	Level 5 vs. Level 1	2177.778	1	2177.778	93.556	.000	.921
Error(week)	Level 2 vs. Level 1	208.000	8	26.000			
	Level 3 vs. Level 1	112.000	8	14.000			
	Level 4 vs. Level 1	136.222	8	17.028			
	Level 5 vs. Level 1	186.222	8	23.278			

وما يهمننا هو الجزء الخاص ب Weak ويظهر أن لا توجد فروق بين القياسات فى الأسبوع الأول والقياسات فى الأسبوع الثانى Sig= 0.849 بينما توجد فروق بين قياسات الأسبوع الثالث والأول Sig= 0.000 وكذلك فروق بين الأسبوع الرابع والأول Sig= 0.000 وفروق بين الأسبوع الخامس والأول Sig= 0.000.

ويمكن دراسة الفروق بين كل قياسات كل أسبوع والأسبوع الذى يليه من خلال تحديد اختيار Repeated فى Contrast ثم اضغط Change:



وتكون نتائج المقارنات كالآتى:

Tests of Within-Subjects Contrasts							
Measure: MEASURE_1							
Source	week	Type III Sum of Squares	df	Mean Square	F	Sig.	Partial Eta Squared
week	Level 1 vs. Level 2	1.000	1	1.000	.038	.849	.005
	Level 2 vs. Level 3	1444.000	1	1444.000	115.520	.000	.935
	Level 3 vs. Level 4	113.778	1	113.778	18.876	.002	.702
	Level 4 vs. Level 5	9.000	1	9.000	1.286	.290	.138
Error(week)	Level 1 vs. Level 2	208.000	8	26.000			
	Level 2 vs. Level 3	100.000	8	12.500			
	Level 3 vs. Level 4	48.222	8	6.028			
	Level 4 vs. Level 5	56.000	8	7.000			

ثم اعطى البرنامج المقارنات البعدية باستخدام اختبار Sidak كالآتى:

Pairwise Comparisons						
Measure: MEASURE_1						
(I) week	(J) week	Mean Difference (I-J)	Std. Error	Sig. ^b	95% Confidence Interval for Difference ^b	
					Lower Bound	Upper Bound
1	2	.333	1.700	1.000	-6.153-	6.820
	3	13.000*	1.247	.000	8.240	17.760
	4	16.556*	1.375	.000	11.306	21.805
	5	15.556*	1.608	.000	9.418	21.693
2	1	-.333-	1.700	1.000	-6.820-	6.153
	3	12.667*	1.179	.000	8.169	17.164
	4	16.222*	.909	.000	12.751	19.693
	5	15.222*	1.441	.000	9.722	20.722
3	1	-13.000-*	1.247	.000	-17.760-	-8.240-
	2	-12.667-*	1.179	.000	-17.164-	-8.169-
	4	3.556*	.818	.024	.432	6.679
	5	2.556	1.156	.450	-1.856-	6.967
4	1	-16.556-*	1.375	.000	-21.805-	-11.306-
	2	-16.222-*	.909	.000	-19.693-	-12.751-
	3	-3.556-*	.818	.024	-6.679-	-.432-
	5	-1.000-	.882	.967	-4.366-	2.366
5	1	-15.556-*	1.608	.000	-21.693-	-9.418-
	2	-15.222-*	1.441	.000	-20.722-	-9.722-
	3	-2.556-	1.156	.450	-6.967-	1.856
	4	1.000	.882	.967	-2.366-	4.366

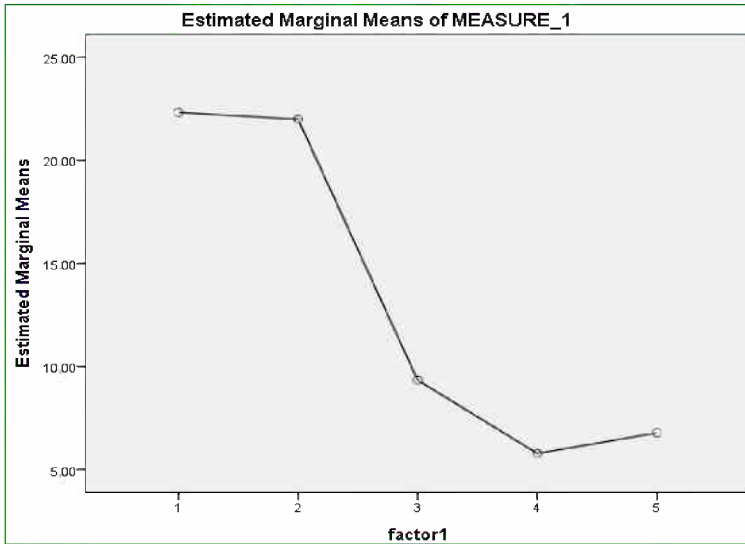
Based on estimated marginal means

*. The mean difference is significant at the .05 level.

b. Adjustment for multiple comparisons: Sidak.

ويمكن إجراء المقارنات المتعددة باستخدام طريقة Boneforoni. ويمكن إجراء المقارنات بين أزواج القياسات باستخدام اختبار T المرتبطة بين كل زوج من خلال اختيار Analyze ثم Compare means ثم Paired sample T-test.

ثم أعطى الرسم البياني لمتوسطات القياسات كالاتي:



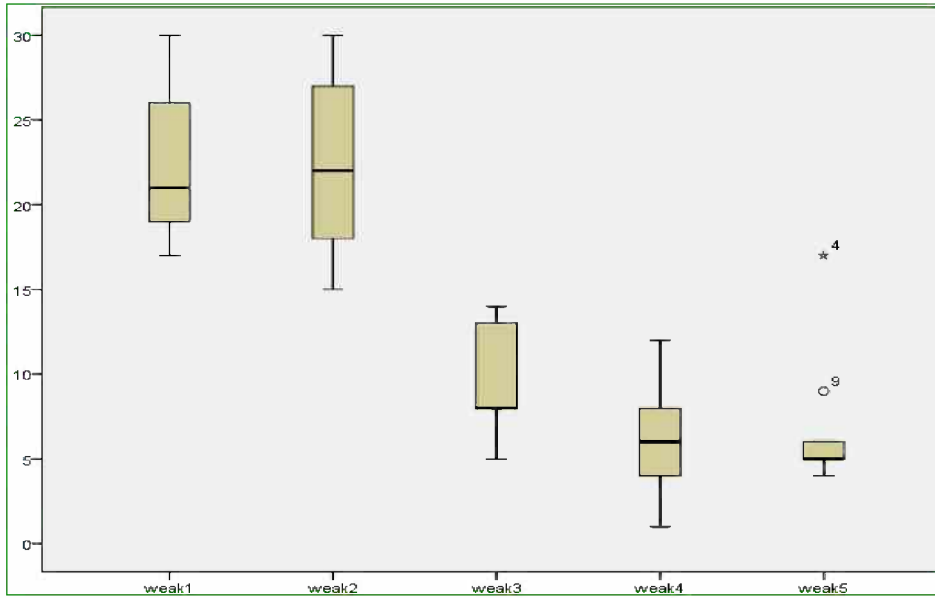
حيث يظهر أن المتوسط يقل كلما زاد الزمن حيث متوسطات قياسات الصداع تقل كلما زاد فترة الاسترخاء. عرض النتائج في صورة العرض البياني

1. اضغط على Graphs ثم Legacy dialogs ثم Boxplot.

2. اضغط Simple ثم اضغط Summaries of separate variable.

4. اضغط Define.

5. انقل المتغيرات إلى مربع Boxes represent ثم اضغط Ok.



الملاحظ أن وسيط قياسات الأسبوع الأول والثاني أعلى من باقى القياسات وأن القياسات الخمسة بها قيم متطرفة وهى للحالة 9 وهى تعادل الدرجة 9 كما تتضمن قيمة متطرفة جدًا وهى للحالة 4 وتناظر الدرجة 17.

عرض نتائج تحليل التباين الأحادى للقياسات المتكررة فى ضوء APA

عرض جدول الإحصائيات الوصفية ويتضمن المتوسط M والانحراف المعياري SD:

SD	M	
4.58	22.33	القياسة الأولى
4.11	6.77	القياسة الخمسة

واظهرت نتائج Repeated ANOVA وجود فروق بين القياسات الخمسة فى الاتحاد الخطى للمتغيرات الخمسة معًا:

Wilks Lambda (4,5) = 0.014, F = 86.391, P = 0.001, $\eta_p^2 = 0.970$

بينما أظهرت نتائج ANOVA فروق بين القياسات الخمسة على حدة:

F(4,32) = 53.252, p = 0.000, $\eta_p^2 = 0.869$

وأظهرت المقارنات المتعددة باستخدام اختبار Sidak وجود فروق بين المجموعة الأولى والخامسة $P= 0.000$ ، بين المجموعة الأولى والرابعة وهكذا.

وأخيرًا عليك المقارنة بين النتائج التي توصل إليها (2013) Howell والنتائج التي حصلت عليها من المخرج في SPSS.

الفصل الرابع

التحليل التباين العاملى أو تحليل التباين ذو اتجاهين بين المجموعات

Factorial Analysis of Variance or Two- way Between Groups ANOVA

فى هذا الفصل سوف تتسع مناقشتنا لنتضمن تحليل التباين المتعدد أو ذو اتجاهين حيث يستخدم فى الدراسات التى تتضمن متغيرين مستقلين فأكثر. ويشار اليه بتحليل التباين ثنائى الاتجاه Two way ANOVA . وعلى ذلك فهو يتضمن:

- أكثر من متغير مستقل كيفى تصنيفى.

- متغير واحد تابع فترى متصل.

وعلى ذلك فإن هذا الأسلوب يقع تحت عنوان عام وهو التحليلات الأحادية Univariate. وعند هذه النقطة فإن التعقيد للتصميم الإحصائى يتنوع فى ضوء:

1. التركيز على تغير مستويات عامل واحد: ويستخدم لهذا اختبارات لدراسة الفروق داخل مجموعة واحدة (Z, T) أو فروق بين مجموعتين (T مستقلة أو مرتبطة) أو اختبارات الفروق بين متوسطين فأكثر لمستويات العامل (ANOVA).

2. كيف تتم القياسات على المشاركين: تتم القياسات على أفراد مختلفين أو مجموعات مستقلة حيث تتضمن كل مجموعة أفراد مختلفين عن المجموعة الأخرى (تصميم بين الأفراد) أو تتم القياسات على نفس الأفراد على مستويات العامل (تصميم داخل الأفراد أو القياسات المتكررة).

وتحليل التباين أحادى الاتجاه يتعامل مع تصميم المتغير المستقل الوحيد أو التجربة أحادية العامل Single Factor experiment وعليه فإنه يوجد متغير مستقل (معالجة) واحدة فقط ويوجد لها مستويين فأكثر.

ولكن تحليل التباين ليس قاصرًا فقط على تصميم أو تجربة العامل الواحد. وفى الواقع يمكن دراسة أثر عوامل (متغيرات مستقلة) عديدة فى نفس الوقت فى التجربة الواحدة وهذه التصميمات يطلق عليه بالتجارب أو التصميمات العاملية Factorial

Experiments or Designs. وعلى ذلك فإن التجربة النموذجية تتعامل مع متغير مستقل واحد على متغير تابع واحد، ولكن في مجال العلوم الإنسانية فإن الظاهرة ليست بهذه البساطة فلا يمكن عزل تأثير متغير ما عن بقية المتغيرات الأخرى ولكن الباحث يريد التعامل أو دراسة الظاهرة كما تحدث في واقعها الطبيعي، وعلى ذلك يسعى إلى تضمين عدد أكبر من المتغيرات المستقلة التي يرى أنها أكثر تأثيراً على ناتج معين (متغير تابع). وعلى ذلك فيمكن للباحث أن يدرس أثر تغير متغير ما أو أكثر على التغير الحادث على ظاهرة ما (متغير تابع). وفي التصميمات التجريبية يطلق على متغير المعالجة بالعامل Factor . وعليه فعند التعامل مع بيانات تصميم يتضمن عاملين أو متغيرين مستقلين فأكثر باستخدام تحليل التباين يطلق عليه يسمى تحليل التباين العاملي Factorial ANOVA.

وهذا يعنى أن ANOVA تتعامل مع متغيرين مستقلين فأكثر بينما المتغير التابع وحيد. وإذا كان التصميم التجريبي يتضمن معالجتين أو متغيرين مستقلين يطلق عليه التصميم العاملي ذو اتجاهين وفي التجربة إذا حدث تزاوج بين كل مستوى للعامل مع كل مستوى للعامل الآخر فإنه يطلق عليه تصميم عاملي Factorial Design . وعلى ذلك فإن التصميم العاملي يتضمن كل الاتحادات أو التفاعلات Interactions بين كل مستويات المتغيرات المستقلة (العوامل) والتصميم عاملي يسمى حسب عدد العوامل المتضمنة فيه، فالتصميم العاملي الذي يتضمن متغيرين مستقلين أو عاملين يسمى Two-way factorial Design، وإذا كان عدد المتغيرات المستقلة (العوامل) ثلاثة يطلق عليه Three-way Factorial Design، إذا كان أحد التصميمات العاملية تتضمن معالجتين وأحد المعالجات بخمس مستويات والمتغير المستقل الآخر مستويين فإنه يطلق عليه 2×5 Factorial ANOVA. وإذا كان يتضمن ثلاث متغيرات مستقلة اثنان منهما لهما ثلاث مستويات بينما الثالث له أربعة مستويات يطلق عليه $3 \times 3 \times 4$ Factorial ANOVA وعلى ذلك فإن تحليل التباين ثنائي الاتجاه (متغيرين) أو ثلاثى الاتجاه (ثلاث متغيرات مستقلة) أو تحليل التباين عالى الرتبة High-order ANOVA أو يطلق عليه Factorial ANOVA.

ويرى (2012) Nolan & Heinzen أنه تحليل إحصائي يتعامل مع متغير تابع متصل (فترى على الأقل) ومتغيرين مستقلين اسميين على الأقل (عاملين على الأقل) ويطلق عليه أحيانا ANOVA Multifactorial حيث يشير Multi إلى تعدد المتغيرات المقاسة وFactor إلى المتغير المستقل. لاحظ أن هذا يختلف تمامًا عن Multivariate ANOVA (MANOVA) حيث تشير Multivariate إلى تعدد المتغيرات التابعة عكس تحليل التباين المتعدد Multiple ANOVA وهو تعدد في المتغيرات المستقلة بينما التابع وحيد.

أوجه التشابه بين تحليل التباين البسيط وتحليل التباين العاملي

يتشابه تحليل التباين الأحادي مع تحليل التباين المتعدد في:

- كلاهما قائم على متوسطات المجموعات.
- اختبارات الفروض قائمة على متوسطات المجتمعات سواء باستخدام مدخل اختبارات الفروض الصفرية أو مدخل فترات الثقة.
- لهما نفس المسلمات .
- يستخدموا تحليل تتبعي (لاحق) باستخدام المقارنات المتعددة سواء كانت المخطط لها (قبلية) أو البعدية.

مميزات التصميمات العاملية

للتصميمات العاملية مميزات عديدة عن التصميمات أحادية العامل يحددها (2013) Howell:

- تسمح بتعميم أكثر قابلية وموثوقية للنتائج لأنه تتضمن متغيرات أكثر تمثيلاً لواقع الظاهرة.
- يعطى قدره تفسيرية أفضل للنتائج خاصة إذا حدثت تفاعلات بين مستويات المتغيرات المستقلة وهذه هي طبيعة الظواهر السلوكية حيث إنها معقدة ومتفاعلة وهذا يعطى عمق للظاهرة المراد دراستها.
- أحد مميزات التصميم العاملي هو الاقتصادية Economy حيث تسمح بدراسة تأثيرات أحد المتغيرات عبر مستويات متغير آخر، حيث يمكن لنفس أفراد العينة

تعرضهم لمعالجتين بدلاً من معالجة واحدة. وعلى ذلك فهو اقتصادى فى حجم العينة وكذلك فى التكلفة فبدلاً من تناول كل معالجة أو متغير مستقل فى تجربة على حدة يتم تناول المتغيرين أو العاملين معاً فى نفس التجربة فهو اقتصادى فى الوقت والجهد والتكلفة.

• الكفاءة Efficiency : مع التحليل التلازمى لمتغيرين مستقلين كأن يتم دمج نتائج من إجراء بحثين لكل متغير مستقل على حدة. ويمكن اختبار ما إذا كان المستويات المختلفة للمتغيرين المستقلين تؤثر فى المتغير التابع وهو ما يطلق عليه التفاعلات بين المتغيرين المستقلين ولذلك ففى التصميم العاملى يمكن أن تدرس تأثيرات المتغيرات المستقلة على حدة وتسمى التأثيرات الرئيسية Main effect وكذلك تأثير التفاعل بين المستويات المختلفة للمتغيرات المستقلة ويسمى تأثيرات التفاعل Interaction effect.

• أحد المميزات هو ضبط Control للمتغير الثانى عن طريق تضمينه فى التصميم كمتغير مستقل، فعلى سبيل المثال يهتم الباحث بدراسة التحصيل الرياضى عبر ثلاث طرق تدريسية مختلفة فاذا اراد الباحث ضبط الذكاء IQ فإنه يتضمنه فى التصميم التجريبى كمتغير مستقل ثانى (منخفض - متوسط - مرتفع) ويصبح التصميم العاملى 3×3 .

هيكلية تحليل التباين ثنائى الاتجاه

من الطبيعى فهم بناء التصميم العاملى أو تحليل التباين العاملى. فاذا اعتبرنا تجربة أو دراسة تتضمن متغيرين مستقلين (A , B) . حيث عدد مستويات المتغير المستقل الأول A هى $a_1, a_2, \dots, a_j$ وعدد مستويات المتغير المستقل الثانى B هى $b_1, b_2, \dots, b_k$ فإن بناؤه كالاتى (Hinkle et al., 1994):

		مستويات المتغير المستقل الثاني			
		b_1	b_2	b_3	b_k
مستويات المتغير المستقل الأول	a_1	X_{111}	X_{112}	X_{113}	X_{11k}
		X_{211}	X_{212}	X_{313}	X_{21k}
		X_{n11}	X_{n12}	X_{n13}	X_{n1k}
	a_2	X_{121}	X_{122}	X_{123}	X_{12k}
		X_{221}	X_{222}	X_{223}	X_{22k}
		X_{n21}	X_{n22}	X_{n23}	X_{n2k}
	a_3	X_{131}	X_{132}	X_{133}	X_{n3k}
		X_{231}	X_{232}	X_{233}	--
		X_{n31}	X_{n32}	X_{n33}	X_{nj3}
	a_j	-----	-----	-----	-----
		----	----	----	----
		X_{1j1}	X_{1j2}	X_{1j3}	X_{1j3}
		X_{2j1}	X_{2j2}	X_{2j3}	X_{2jk}
		X_{nj1}	X_{nj2}	X_{nj3}	X_{nj3}
		$\bar{x}.1$	$\bar{x}.2$	$\bar{x}.3$	$\bar{x}.k$
					$\bar{x}$

- $\bar{x}_{jk}$ تناظر متوسط درجات الصف J والعمود K ، $\bar{x}_j$ متوسط درجات الصف j ، $\bar{x}_k$ متوسط درجات العمود K ، $\bar{x}_1$ متوسط كل الدرجات (القياسات) في الصف الأول (المستوى الأول للعامل الأول). $\bar{x}_2$ متوسط كل درجات الصف الثاني. $\bar{x}.1$ متوسط درجات العمود الأول (المستوى الأول من العامل الثاني). $\bar{x}_{11}$ متوسط درجات الخلية الأولى في المصفوفة حيث X تعبر عن درجات الأفراد للمتغير التابع في التكوينات المختلفة للمستويات.

عمومًا المتغير المستقل الأول يمثل الصفوف، والمتغير المستقل الثاني يمثل الأعمدة.

- X_{ijk} هي درجة الطالب i في الصف j في العمود K .
- X_{111} درجة الفرد الأول في المستوى الأول للمتغير المستقل الأول والمستوى الأول للمتغير المستقل الثاني.
- X_{223} درجة الطالب الثاني في الصف الثاني (المستوى الثاني من المتغير المستقل الأول) في العمود الثالث (المستوى الثالث للمتغير المستقل الثاني).
- X_{njk} درجة الطالب n في الصف (المتغير المستقل الأول) j في العمود (المتغير المستقل الثاني) k .

وإذا كان نفس العدد من الأفراد (n) في كل خلية من التصميم العامل فإن العدد الكلي من الأفراد في كل الخلايا هو (N). وإذا كان عدد مستويات المتغير المستقل الأول (الصفوف) J وعدد مستويات المتغير المستقل الثاني (الأعمدة) K فإن عدد خلايا التصميم $K = J$ ، والعدد الكلي من القياسات أو الأفراد: $N = n * J * K$. بفرض أن عدد مستويات العامل الأول 2، وعدد مستويات العامل الثاني 3 فإن عدد الخلايا ($2 \times 3 = 6$)

وبفرض أن كل خلايا تتضمن ستة أفراد فإن العدد الاجمالي للعينة:

$$= 6 \times 2 \times 3 = 36$$

ويمكن عرض التصميم بعوامله و خلايا كالاتي:

مستويات العامل الأول	مستويات العامل B				
		b_1	b_2	b_c	
	a_1	Cell ₁₁	Cell ₁₂	Cell _{1c}	r_1
	a_2	Cell ₂₁	Cell ₂₂	Cell _{2c}	r_2
	a_3	Cell ₃₁	Cell ₃₂	Cell _{3c}	r_3
	A	Cell r_1	Cell r_2	Cell r_c	R
					Σx

- C تمثل العمود (Column) وهي مثل K في العرض السابق.

• r تمثل الصف (Row) وهي مشابهه لـ J في العرض السابق.

ودعنا نطرح مثال لباحث تجريبي في مجال التربية الرياضية اراد إجراء تجربة لدراسة أثر توقيت إجراء التمارين الرياضية (صباحًا - مساءً) وشدة التمارين (قوية-خفيفة) على مدة النوم وتم عرض التصميم العاملى كما فى الشكل الآتى:

شدة التمرين (العامل B)			
وقت التمرين (A)		خفيف b_1	شديد b_2
	a_1 (صباحًا)	$b_1 \times a_1$ درجات	$a_1 \times b_2$ درجات
	a_2 (مساءً)	$a_2 \times b_1$ درجات	$a_2 \times b_2$ درجات

نلاحظ أنه يوجد متغيرين مستقلين أو عاملين كل عامل يتضمن مستويين ويسمى تصميم 2×2 (Two by Two) وإذا تضمن أحد العوامل ولتكن فترة التمرين (صباحًا-ظهرًا-مساءً) ثلاثة مستويات فإنه يسمى 2×3 . ويوجد مجموعات مستقلة ويسمى تصميم المجموعات المستقلة ويتم اختيار الأفراد عشوائيًا فى كل مجموعة أو فى كل خلية حيث الخلية هى اتحاد بين أحد مستويات العامل الأول بأحد مستويات العامل الثانى. لاحظ استخدام المصطلحات الآتية: عامل = متغير مستقل فى التجربة، مستوى = مستويات أو تصنيفات أو شروط العامل، خلية = اتحاد بين كل مستويين لعاملين مختلفين. وعندما يتم تحديد مستويات العامل أو المتغير المستقل مسبقًا بواسطة

المجرب فإن هذا تصميم التأثيرات الثابتة Fixed effects design

ويوجد ثلاثة تحليلات فى تحليل التباين العاملى وهى تجيب عن ثلاثة أسئلة:

1. هل يوجد تأثير دال للعامل الأول وهو وقت التمرين؟ بمعنى هل توجد فروق فى عدد ساعات النوم بين إجراء التمارين فى الصباح وفى المساء؟ وهذا يسمى تأثير رئيسى Main effect للعامل الأول. وعليه فإن التأثيرات الرئيسية هى فروق المتوسطات بين مستويات كل عامل على حدة، وعندما يتم التعبير عن العامل الأول

فى الصفوف والعامل الثانى فى الأعمدة وعليه فإن الفروق بين متوسطات الصفوف تصف التأثير الرئيسى الأول بينما الفروق بين متوسطات الأعمدة تصف التأثير الرئيسى الثانى. وإذا نظرنا إلى الفروق فى ساعات النوم فى ضوء توقيت التمارين نتجاهل الفروق نتيجة شدة التمرين .

وإذا تناول الباحث أثر أحد العوامل (A) على درجات مستوى واحد فقط من العامل الثانى فإنه يطلق عليه التأثير البسيط أو التأثير الشرطى Simple effect or Conditional effect. وهذا عكس التأثير الرئيسى حيث يتم تجاهل تأثير المعالجة الأخرى.

وعلى ذلك فإن هذا التحليل يهدف إلى تحديد ما إذا يوجد تأثير دال للعامل A (وقت إجراء التمرين)؟

وتكون الفروض الإحصائية: $H_0: \mu_{A1} = \mu_{A2}$

$H_A: \mu_{A1} \neq \mu_{A2}$

2. هل يوجد تأثير دال لشدة التمرين على مدة النوم ؟ وهذا أيضًا تأثير رئيسى للعامل

الثانى وتكون الفروض الإحصائية: $H_0: \mu_{B1} = \mu_{B2}$

$H_A: \mu_{B1} \neq \mu_{B2}$

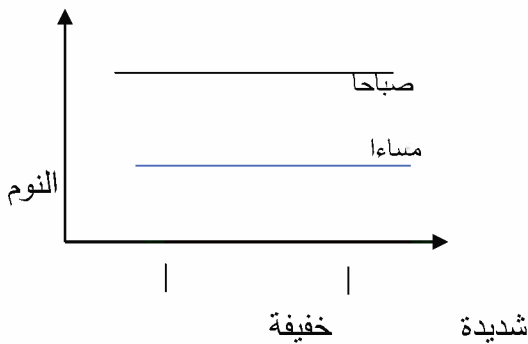
3. هل يوجد تأثير دال للتفاعل بين وقت التمرين (A) وشدة التمرين (B) على مدة النوم؟ ويشار إلى هذا التأثير بـ Interaction effect فمثلاً هل تختلف درجات الخلية التى تتضمن شدة تمرين خفيف ووقت التمرين فى الصباح عن بقية الخلايا الأخرى بالتالى تكون الفروض الإحصائية:

H_0 : لا يوجد تفاعل بين مستويات العاملين A , B أو كل الفروق بين المتوسطات فى المعالجات (الخلايا) يتم تفسيرها من خلال التأثيرات الرئيسية للعاملين.

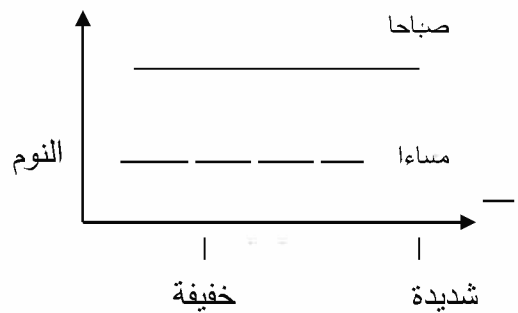
H_A : يوجد تفاعل بين العاملين. الفروق فى متوسطات الخلايا أو شروط المعالجة لا يمكن التنبؤ بها من التأثيرات الرئيسية للعاملين. والتفاعل يحدث عندما يكون أثر أحد

العوامل ليس بنفس الدرجة لكل مستويات العامل الآخر. والتفاعل هو التأثير المنتج عندما يعمل العاملين معاً، وعلى ذلك توجد اعتمادية Dependence بين العوامل. وفكرة التفاعل قائمة على ما يقوم به الصيدلي بصرف دواء ويوصى بعدم اخذ دواء اخر لأنه يعطل أو يشوه مفعول الدواء الأول. وعلى ذلك فإن تأثير العقار الأول (العامل A) يعتمد على تأثير الدواء الثانى (العامل B) وعليه فإن التفاعل يتعامل مع متوسطات الخلايا وليس متوسطات التأثيرات الرئيسية وفيها يتم تمثيل مستويات أحد العوامل على المحور الافقى (السينى) وفى هذه الحالة توضع مستويات المعالجة (B) ويوضع المتغير التابع وهو كمية النوم (مدتها) على المحور الصادى (الرأسى) ثم يتم عرض خطين، الخط الأسفل يمثل (صباحاً) بينما الخط المنقوط (----- مساءً) وبين الخطوط العلاقة بين مستويات العامل (B) ومدة النوم. وعليه يمثل فى ضوء محور سيني يمثل مستويات أحد العوامل ومحور صادى يمثل المتغير التابع وخطوط منفصلة تمثل مستويات العامل الآخر. وفيما يلى الاشكال الآتية لتوضيح التأثيرات الرئيسية والتفاعلات:

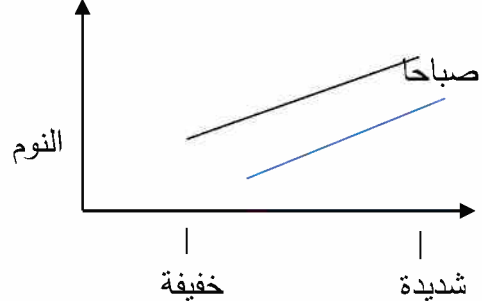
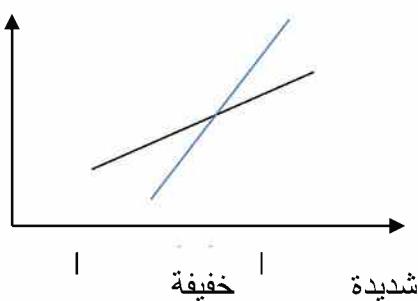
(A) لا تأثيرات دالة



(B) تأثير دال للتوقيت فقط

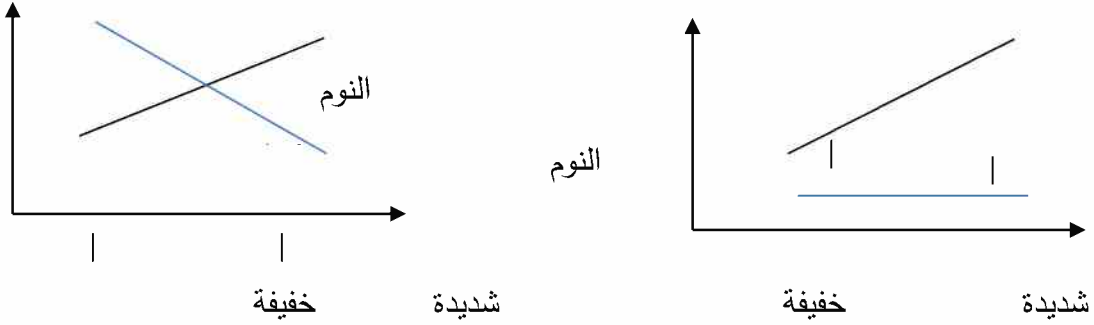


(D). تأثير دال للتوقيت وشدة التمرين ولا تفاعل (C). تأثير دال لشدة التمرين



(e). تأثير دال للتفاعل ودال للتأثيرات الرئيسية

(F). تأثير دال لوقت التمرين وكذلك لتأثير التفاعل



الشكل (1.4): بعض النواتج المحتملة لتصميم تجريبي لعاملين شدة التمرين وتوقيته.

فالشكل السابق يوضح بعض النواتج المحتملة للتجربة ففي الشكل (A) لا يوجد تأثيرات دالة إحصائية، في الجزء (B) يوجد تأثير رئيسي دال لوقت التمارين ولا يوجد تأثير دال لشدة التمارين أو للتفاعل. وعليه فالأفراد يكون أكثر نومًا لو كان أداء التمارين صباحًا وليس مساءً. في الجزء (C) توجد دلالة لشدة التمارين ولا يوجد تأثير دال لوقت التمارين أو للتفاعل، فأداء التمارين بدرجة شديدة يؤدي إلى نومًا أكثر في مقابل تمارين خفيفة. والجزء (D) يوضح تأثيرات رئيسية دالة لوقت التمارين ولشدتها ولا يوجد تأثير للتفاعل و الجزء E ، F فإن تأثيرات التفاعل دالة إحصائية. وعليه فإن تأثير التفاعل لاحتد العوامل ليس هو نفسه عبر كل مستويات العامل الآخر. ففي الجزء (E) يوجد تأثير دال للتفاعل بين شدة التمارين وتوقيتها ولكن تأثير شدة التمارين ليست هي نفسها كما في مستويات التوقيت، فإذا أجريت التمارين في المساء فإن التمارين الخفيفة تؤدي إلى نوم أكثر في مقابل التمارين الشديدة في المساء. وعلى الجانب الآخر فإن إذا كان أداء التمارين صباحًا فإن التمارين الشديدة تؤدي إلى نوم أكثر من التمارين الخفيفة. في الجزء (F) يوجد تأثير دال لوقت التمرين وكذلك تأثير دال للتفاعل، فعند عمل التمارين في الصباح فإنه يؤدي إلى نوم أكثر من أدائها في المساء بغض النظر عن شدتها سواء كانت ضعيفة أو شديدة بالإضافة إلى وجود تفاعل بين شدة التمرين ووقت التمرين . وبالتالي لا يوجد تأثير دال لشدة التمرين بمفردها. وعليه يمكن تحديد التفاعل من نمط المسارات في الشكل السابق ولاحظ في الشكل (B) ، (D) فإن الخطين متوازيين بمعنى المسافة بين الخطين ثابتة وهذا يعني عدم وجود تفاعل.

ولكن عندما يكون الخطين غير متوازيين حيث تتغير المسافة بين الخطين من على الشمال إلى اليمين كما في E ، F فهذا يدل على وجود تفاعل وعلى ذلك فكلما حدث تقاطع بين الخطين أو المسافة بينهما غير متساوية فإن هذا يشير إلى وجود تفاعل.

تصميمات أو اشكال التحليل التباين العاملي

- عندما يوجد تصميم يتضمن عاملين وتوجد أفراد مختلفين داخل كل خلية فإنه يطلق عليه تحليل التباين العاملي بين المجموعات أو بين الأفراد. حيث يوجد أفراد مختلفين في كل خلية ناتجة عن اتحاد مستويات العاملين. حيث عدد الأفراد داخل كل خلية n وعدد أفراد العينة الاجمالي في كل الخلايا N حيث $N = nrc$

- عندما يوجد تصميم عاملي يتضمن عاملين ويوجد نفس الأفراد في كل الخلايا فإنه يسمى Factorial analysis ANOVA within groups أو إذا كان عدد مستويات العامل الأول 3 وعدد مستويات العامل الثاني 2 فإنه يطلق عليه 3×2 within subject ANOVA . حيث يوجد نفس العدد من الأفراد في الخلايا الستة فإن $N = n$.

- عندما نتعامل مع تصميم عاملي يتضمن عاملين وفيه يتم ملاحظة أفراد مختلفين في كل مستوى من مستويات أحد العوامل و يتم قياس نفس الأفراد على مستويات العامل الثاني ويسمى Mixed-subject ANOVA

التباينات في تحليل التباين المتعدد

في تحليل التباين ذو اتجاهين يتحدد اربعة تباينات:

- متوسط مجموع المربعات للصفوف MS_{rows}

- متوسط مجموع المربعات للأعمدة $MS_{columns}$

- متوسط مجموع المربعات للتفاعل $MS_{interaction}$

- متوسط مجموع المربعات داخل الخلايا MS_{within}

وتقدير التباين داخل الخلايا يناظر التباين داخل المجموعات (الخطأ) في تحليل التباين أحادي الاتجاه MS_{rows} هو تقدير التباين للصفوف وهو قائم على الاختلاف أو الفروق بين متوسطات الصفوف ولذلك فهو يمثل تأثير العامل الأول

(A)، و MS columns قائم على الفروق بين متوسطات الأعمدة وهو يمثل تأثير العامل الثانى (B) و MS interactions تقدير تباين التفاعلات وهو قائم على الاختلاف بين متوسطات الخلايا وهو يمثل تأثير التفاعل.

ولاختبار الدلالة الإحصائية لهذه التأثيرات الثلاثة يتم حساب نسب F ثلاث مرات وهى:

- للعامل أو للمتغير المستقل A:

$$F = \frac{MS_{row}}{MS_{within-cell}}$$

- للعامل أو للمتغير المستقل B:

$$F = \frac{MS_{column}}{MS_{within-cell}}$$

- للتفاعل:

$$F = \frac{MS_{int}}{MS_{within-cell}}$$

تقدير التباين داخل الخلايا $MS_{within-cell}$

يتم تقديره من خلال تباين درجات كل خلية، فكل الأفراد داخل كل خلية تلقوا نفس المستوى من المتغيرات المستقلة A, B, فإن تباين درجاتهم لا ترجع إلى الفروق بين المعالجات. وتباين داخل المجموعات هو مشابه لتباين داخل المجموعات فى حالة تحليل التباين أحادى الاتجاه. واعلم أن هذا التباين لا يتأثر بأى من العوامل A, B, أو التفاعل بينهما ويطلق عليه تباين الخطأ أو البواقي. وعليه:

$$MS_{within-cell} = \frac{SS_{within-cell}}{df_{within-cell}}$$

• SS هى مجموع مربعات داخل الخلايا

• df درجات الحرية داخل الخلايا

وحيث :

$$SS_{within-cell} = SS_{11} + S_{12} + SS_{13} + \dots + S_{21} + S_{22} + \dots + SS_{rc}$$

• SS_{11} مجموع مربعات الدرجات فى الخلية المحددة بالصف الأول والعمود الأول.

- SS_{12} مجموع مربعات الدرجات فى الخلية المحددة بالصف الأول و العمود الثانى.
- SS_{rc} مجموع المربعات الدرجات فى الخلية المحددة بالصف r والعمود c (الخلية الأخيرة فى المصفوفة)

ويتم تحديد مجموع المربعات داخل الخلية بالآتى:

$$SS_{\text{within-cell}} = \sum x^2 - \left[\frac{(\sum x_{11})^2 + (\sum x_{12})^2 + \dots + (\sum x_{rc})^2}{n_{\text{cell}}} \right]$$

- $\sum x^2$ مجموع مربعات كل الدرجات فى الخلايا حيث يعبر عنها
- $\sum x^2 = T^2$

- $(\sum x_{11})^2$ مربع مجموع درجات الخلية التى تشغل الصف والعمود الأول.
- n_{cell} عدد أفراد فى الخلية.

لاحظ نتعامل مع تصميم يتعامل مع نفس العدد الأفراد فى كل خلية ولذلك $n_{\text{cell}} = n$

وتتحدد درجات الحرية لمجموع المربعات داخل الخلايا كالآتى:

$$df_{\text{within-cell}} = rc(n-1)$$

- r عدد الصفوف.

- c عدد الأعمدة .

تباين الصفوف (المتغير المستقل الأول) (MS_A) MS_{rows}

هذا التقدير قائم على الفروق بين متوسطات الصفوف ($X_1, X_2, \dots, X_r$) وهو مشابهاً للتباين بين المجموعات فى تحليل التباين أحادى الاتجاه وهو التباين الذى يعود إلى أثر المتغير المستقل حيث إذا كان متوسطات الصفوف متساوية فإن هذا يمثل الفرض الصفرى للمعالجة A :

$$H_0: \mu_{a1} = \mu_{a2} = \mu_{a3} = \dots = \mu_{ar}$$

وتقدر قيمة متوسط المربعات:

$$MS_{\text{rows}} = \frac{SS_{\text{rows}}}{df_{\text{rows}}}$$

- SS_{rows} مجموع مربعات الصفوف.

- df درجات الحرية للصفوف.

$$SS_{rows} = n_{row}[(\bar{x}_{row1} - \bar{x}_G)^2 + (\bar{x}_{row2} - \bar{x}_G)^2 + \dots + (\bar{x}_{rowr} - \bar{x}_G)^2]$$

- $\bar{x}_{row} = \frac{\sum x}{n_{r1}}$: متوسط درجات الصف الأول

- $\bar{x}_{row.r} = \frac{\sum x}{n_{rr}}$: متوسط درجات الصف r

- $\bar{x}_G = \frac{\bar{x}_{row1} + \bar{x}_{row2} + \dots + \bar{x}_{rowr}}{n_r}$: متوسط المتوسطات

- nr عدد الصفوف

ويتم حساب SS_{rows} من الدرجات الخام كالآتي:

$$SS_r = \left[\frac{(\sum x_{r1})^2 + (\sum x_{r2})^2 + (\sum x_{r3})^2 + \dots + (\sum x_{rr})^2}{n_r} \right] - \frac{(\sum x)^2}{N}$$

- $(\sum x_{r1})^2$ مجموع درجات الصف الأول.

- $\sum x_{rr}$ مجموع درجات الصف r .

- $(\sum x)^2$ مربع مجموع درجات كل الخلايا ($\sum x = T$).

- n_r عدد أفراد الصف

وتقدر درجات الحرية للصفوف كالآتي:

$$df = r - 1$$

تقدير تباين الأعمدة MS_{column}

هذا التباين قائم على الفروق بين متوسطات الأعمدة وهو تمامًا مثل MS_{rows} ولكنه يتعامل مع متوسطات الأعمدة بدلاً من متوسطات الصفوف وهو يمثل تأثير المتغير

المستقل الثاني B ويمثل الفرض الصفري : $H_0: \mu_{b1} = \mu_{b2} = \dots = \mu_{bc}$

$$MS_{column} = \frac{SS_{column}}{df_{column}} \text{ حيث:}$$

- SS_{column} مجموع المربعات بين الأعمدة.

- df_{column} درجات الحرية للأعمدة .

وتقدر SS_{column} كالآتي:

$$SS_{column} = n_{col}[(\bar{x}_{col1} - \bar{x}_G)^2 + (\bar{x}_{col2} - \bar{x}_G)^2 + \dots + (\bar{x}_{colc} - \bar{x}_G)^2]$$

$$\bar{x}_{col1} = \frac{\sum x}{n_{col1}} \quad \text{حيث:}$$

$$\bar{x}_{colc} = \frac{\sum x}{n_{colc}}$$

$$\bar{x}_G = \frac{\bar{x}_{col1} + \bar{x}_{col2} + \dots + \bar{x}_{colc}}{n_{cLoc}}$$

وتقدر من الدرجات الخام كالآتي:

$$SS_{column} = \left[\frac{(\sum X_{col1})^2 + (\sum X_{col2})^2 \dots + (\sum X_{cc})^2}{n_{col}} \right] - \frac{(\sum X)^2}{N}$$

- $\sum X_{col1}$ مجموع درجات العمود الأول.
- n_{col} حجم العينة في كل عمود.
- $(\sum X)^2 (T^2)$ مربع مجموع درجات كل الدرجات في كل الأعمدة.
- N حجم العينة في كل الأعمدة.

تقدير تباين التفاعلات ($MS_{Int.}$) $MS_{Interactions}$

يوجد التفاعل عندما يكون تأثير أحد المتغيرات على مستويات المتغير الثاني مختلفة. ويتم تقدير تباين التفاعلات بين مستويات العوامل أو المتغيرات المستقلة لتقويم التفاعل بين A ، B وهو قائم على الفروق بين متوسطات الخلايا ويتم تقديره للتحقق من الفرض الصفري:

$$H_{0int.}: \mu_{a1b1} = \mu_{a1b2} = \mu_{a2b1} = \mu_{a2b2} = \dots = \mu_{arbc}$$

ويتم تقدير $MS_{int.}$ كالآتي:

$$MS_{int.} = \frac{SS_{interaction}}{df_{interaction}}$$

• $SS_{int.}$ مجموع مربعات التفاعلات.

• $df_{int.}$ درجات الحرية للتفاعلات.

ويتم تقدير التباين للتفاعلات بين A , B كالاتى:

$$SS_{int.} = n_{cel} [(\bar{X}_{cel11} - \bar{X}_G)^2 + (\bar{X}_{cel12} - \bar{X}_G)^2 + \dots + (\bar{X}_{celrc} - \bar{X}_G)^2] - SS_{row} - SS_{coloumn}$$

$$SS_{int} = \left[\frac{(\sum^{cell11} X)^2 + (\sum^{cell12} X)^2 + \dots + (\sum^{rc} X)^2}{n_{cel}} - \frac{(\sum^{All} X)^2}{N} \right] - SS_r - SS_{col}$$

• $(\sum^{cell11} X)^2$ مربع مجموع الدرجات فى الخلية الصف الأول والعمود الأول.

• $T^2 = (\sum^{All} X)^2$ مربع مجموع درجات كل الخلايا.

• n_{cel} حجم العينة فى الخلية الواحدة.

وتتحدد درجات الحرية كالاتى:

$$df_{int.} = (r-1)(c-1)$$

ومجموع المربعات الكلى:

$$SS_{Total} = SS_{row} + SS_{coloumn} + SS_{within.cell} + SS_{int.}$$

و درجات الحرية: $df_{total} = N-1$

حساب نسب F

يتم حساب ثلاثة نسب لـ F وهى:

• للتأثيرات الرئيسية:

١- التأثير الرئيسى للمتغير A (تأثير الصف):

$$F = \frac{MS_{row}}{MS_{witin.cell}} = \frac{\text{تأثير المتغير A} + \sigma^2}{\sigma^2}$$

إذا كان المتغير A ليس له تأثير رئيسى فإن تقدير:

$$MS_{row} = \sigma^2$$

وإذا كان A له تأثير رئيسي فإن MS_{row} سوف تكون أكبر من $MS_{within.cell}$ وعليه تزيد قيمة F وبالتالي يرفض H_0

ب- التأثير الرئيسي للمتغير B (تأثير العمود):

$$F = \frac{MS_{col}}{MS_w} = \frac{\sigma^2 + B \text{ تأثير المتغير}}{\sigma^2}$$

وإذا كان المتغير B ليس له تأثير رئيسي بالتالي فهو تقدير مستقل لـ σ^2 . وإذا كان له تأثير فإن MS_{col} تزيد وعليه يكون البسط أكبر من المقام بالتالي تزيد قيمة F بالتالي تزيد الاحتمال لرفض H_0

ج- لتأثير التفاعل B X A:

$$F = \frac{MS_{int.}}{MS_{w.c}} = \frac{\sigma^2 + B, A \text{ تفاعل}}{\sigma^2}$$

وعليه فإن التوزيع العيني F كما هو الحال في حالة F في تحليل التباين أحادي الاتجاه توزيع ملتوى ناحيه اليمين.

اختبارات الفروض لقضية بحثية (Pagano , 2013):

اجرى باحث فى التربية البدنية تجربة لمقارنة شدة التدريبات (خفيفة-متوسطة-شديدة) ووقتها (فى الصباح-المساء) على كمية او وقت النوم. وتم اختيار ستة طلاب عشوائياً فى كل خلية، حيث سمح لستة طلاب بممارسة رياضة شديدة (الجرى ثلاثة اميال) وستة طلاب ممارسة رياضة متوسطة (جرى ميل واحد) وستة طلاب ممارسة رياضة خفيفة (ربع ميل) وتم إجراء التمارين صباحاً 7:30 ومساءً 7:00 وأراد الباحث التحقق ما إذا كانت لشدة التمرين وتوقيتها ولتفاعلها معاً اثرعلى مده النوم؟

الخطوات البحثية

1. سؤال البحث: توجد ثلاثة أسئلة بحثية وتنقسم الى:

• التأثيرات الرئيسية:

أ- تأثير شدة التمارين (عامل A): هل يوجد تأثير لشدة التمارين على مدة النوم؟ أو هل توجد فروق بين إجراء التمارين بصورة شديدة وخفيفة ومتوسطة على مدة النوم؟

ب- تأثير توقيت التمرين (عامل B): هل يوجد تأثير لتوقيت التمارين على مدة (عدد ساعات) النوم ليلاً أو هل توجد فروق بين إجراء التمارين صباحاً ومساءً في مدة أو زمن النوم ليلاً؟

• تأثير التفاعل بين A, B: هل يوجد تأثير لتفاعل شدة التمارين وتوقيتها على عدد ساعات النوم ليلاً؟

2. فروض البحث: توجد ثلاثة فروض، فرضان للتأثيرات الرئيسية هما:

أ- يوجد أثر لشدة التمرين على ساعات النوم أو توجد فروق في عدد ساعات النوم بين إجراء التمارين بشدة وبدرجة معتدلة وبدرجة خفيفة.

ب- يوجد أثر لوقت التمارين عدد ساعات النوم أو توجد فروق بين إجراء التمارين صباحاً وإجرائها مساءً في عدد ساعات النوم

ج- فرض لتأثير التفاعل: يوجد تأثير لتفاعل شدة التمارين وتوقيتها على عدد ساعات النوم ليلاً.

3. التصميم البحثي: يستخدم هذا الاختبار في:

• الدراسات التجريبية: وهي تتضمن التصميمات العاملية حيث يوجد أكثر من عامل (متغير مستقل) ومتغير تابع واحد.

• الدراسات التجريبية.

• الدراسات غير التجريبية: مثل دراسة أثر متغيرات كيفية غير معالجة مثل الجنس (ذكر - انثى) والشعبة (علمي - أدبي) على التحصيل وفي المثال الحالي فإنه يستخدم تصميم بحث تجريبي.

4. متغيرات البحث: يطلق على المتغيرات المستقلة بالعوامل والمتغيرات هي:

- شدة التمرين: اسمي - مستقل - ثلاث مستويات (3) (B).
- وقت التمرين: مستقل - اسمي - بمستويين (2) (A).
- ساعات النوم: تابع - نسبي - متصل.

5. الإحصاء المستخدم: إحصاء النموذج المتعدد وهو يتعامل مع أكثر من متغير مستقل وتابع واحد وهو إحصاء بارامترى، الاختبار المناسب: ولان هدف الباحث دراسة فروق بين متوسطات المجتمعات وعليه فإن الاختبار المناسب هو تحليل التباين ذو اتجاهين أو تحليل التباين العاُملى (2X3).

خطوات اختبارات الفروض الصفرية

1. الفروض الإحصائية: توجد ثلاثة فروض صفرية وثلاثة بديلة وهى:

أ- فروض صفرية للتأثيرات الرئيسية وتتضمن وباعتبار أن المتغير المستقل الأول (A) هو

$$H_{0A} : \mu_{a1} = \mu_{a2}$$

وباعتبار أن المتغير المستقل (العامل) الثانى هو شدة التمرين (B):

$$H_{0B} : \mu_{b1} = \mu_{b2} = \mu_{b3}$$

ب- فرض صفرى لتأثير التفاعل:

$$H_{0AXB} : \mu_{a1b1} = \mu_{a1b2} = \mu_{a1b3} = \mu_{a2b1} = \mu_{a2b2} = \mu_{a2b3}$$

وعليه تكون الفروض البديلة:

$$H_{A(A)} : \mu_{a1} \neq \mu_{a2}$$

$$H_{A(B)} : \mu_{b1} \neq \mu_{b2} \neq \mu_{b3}$$

$$H_{A(AXB)} : \mu_{a1b1} \neq \mu_{a1b2} \neq \mu_{a1b3} \neq \mu_{a2b1} \neq \mu_{a2b2} \neq \mu_{a2b3}$$

2. الاختبار الإحصائى المناسب ومسلماته: الاختبار هو تحليل التباين ذو الاتجاهين 2X3

ومسلماته فى الآتى:

- الاعتدالية: هذه المسلمة فى غاية الأهمية خاصة للعينات الصغيرة حيث تتضمن خلايا التصميم (الاتحاد بين مستويات العاملين) مجتمعات مختلفة فى التصميم 3X2=6 خلايا ويفترض أن درجات المتغير التابع فى

الخلايا الستة اعتدالية التوزيع. يرى (Green & Salkind (2014) انه
فى بعض التطبيقات مع المجتمعات غير الاعتدالية فإن حجم عينة
15 لكل خلية كافى لإعطاء قيمة دقيقة لـ P.

- **تجانس التباينات:** يفترض أن تباينات المجتمعات فى خلايا التصميم متساوى حيث يفترض
فى تصميم 2 X 3 وجود ستة تباينات متساوية وعدم تحقق هذه المسلمة يزيد احتمالية
ارتكاب الخطأ من النوع الأول.

- **الاستقلالية:** راجع ANOVA أحادى الاتجاه.

- **المعاينة العشوائية:** راجع ANOVA أحادى الاتجاه.

حيث يتم حساب ثلاث قيم لـ F.

3. **مستوى الدلالة الإحصائية وقاعدة القرار:** تبني الباحث $\alpha=0.05$ ولحساب تأثير العامل
(A) تقدر درجات الحرية للعامل للصف:

$$df_{\text{row}} = r - 1 = 2 - 1 = 1 \quad (\text{البسط})$$

ودرجة الحرية بين الخلايا (الخطأ):

$$df_{\text{within}} = rc (n-1) = 2 \times 3 (6-1) = 30 \quad (\text{المقام})$$

بالكشف عن F الحرجة بدرجتى حرية 1 ، 30 و $\alpha = 0.05$ فإن :

$$F_{c.v} = 4.17$$

ولتأثير العامل (B) أو العمود: يتم تقدير F الحرجة من خلال حساب درجة الحرية للعامل (B)
(العمود):

$$df_{\text{col}} = C-1 = 3-1 = 2$$

ودرجة الحرية داخل الخلايا:

$$df_{\text{within}} = rc (n-1) = 3 \times 2 \times 5 = 30$$

وبالبحث فى جدول F بـ $\alpha = 0.05$ و $df = 2$ البسط و $df = 30$ المقام وعليه فان: $F_{c.v} =$

3.32

ولتأثير التفاعل بين A , B:

$$df_{int.} = (c-1) (r-1) = 2 \times 1 = 2 \text{ ، } df_{witin} = 30 \text{ ، وبالكشف في جدول F}$$

$$F_{c.v} = 3.32 \text{ فان:}$$

وتكون قاعدة القرار : إذا كانت F المحسوبة < F الحرجة يرفض H0

4.الحسابات: فيما يلي بيانات التجربة:

		شدة التمرين					
توقيته	صباحًا	خفيفة		متوسطة		شديدة	
		6.5	7.4	7.4	7.3	8.0	7.6
		7.3	7.2	6.8	7.6	7.7	6.6
	مساءً	6.6	6	6.7	7.4	7.1	7.2
		7.1	7.7	7.4	8	8.2	8.7
		7.9	7.5	8.1	7.6	8.5	9.6
		8.2	7.6	8.2	8	9.5	9.4

ويمكن التعبير عن الجدول السابق في ضوء:

		B			Σ
		خفيف	متوسط	شديد	
A	صباحًا	n = 6	n = 6	n = 6	ΣX=129.20
		$\bar{X}= 6.97$	$\bar{X}= 7.20$	$\bar{X}= 7.37$	ΣX ² = 930.50
					n = 18
					$\bar{X}= 7.18$
	مساءً	n = 6	n = 6	n = 6	ΣX=147.20
		$\bar{X}= 7.67$	$\bar{X}= 7.88$	$\bar{X}= 8.98$	ΣX ² =1212.68
					n = 18
					$\bar{X}= 8.18$
	Σ	ΣX=87.0	ΣX=90.50	ΣX=98.10	ΣX=276.40
		ΣX ² = 645.30	ΣX ² = 685.07	ΣX ² = 812.81	ΣX ² =2143.18
		n = 12	n = 12	n = 12	N = 36
		$\bar{X}= 7.32$	$\bar{X}= 7.54$	$\bar{X}= 8.18$	

والحسابات وفقاً للخطوات الآتية:

أولاً: حساب مجموع المربعات:

• حساب مجموع مربعات الصفوف (A):

$$SS_{\text{row}} = \left[\frac{(\sum^{r=1} X)^2 + (\sum^{r=2} X)^2}{n_r} \right] - \frac{(T)^2}{N}$$

حيث: (T)² هي $\left(\sum X \right)^2$ لكل الدرجات مربع مجموع درجات كل الأفراد الـ 36 ، و n_r حجم العينة في

الصف :

$$A = \frac{(129.20)^2 + (147.20)^2}{18} - \frac{(276.40)^2}{36} = 9.0$$

• حساب مجموع مربعات الأعمدة (B):

$$SS_{col} = \left[\frac{(\sum^{col1} X)^2 + (\sum^{col2} X)^2 + (\sum^{col3} X)^2}{n_{col}} \right] - \frac{(T)^2}{N}$$

$$= \frac{(87.8)^2 + (90.5)^2 + (98.1)^2}{12} - \frac{(276.4)^2}{36}$$

$$= 4.754$$

• حساب مجموع مربعات التفاعلات:

$$SS_{int} = \left[\frac{(\sum^{cel11} X)^2 + (\sum^{cel12} X)^2 + (\sum^{cel13} X)^2 + (\sum^{cel21} X)^2 + (\sum^{cel22} X)^2 + (\sum^{cel23} X)^2}{n_{cel}} \right]$$

$$- \frac{(T)^2}{N} - SS_{row} - SS_{column}$$

$$= \left[\frac{(41.8)^2 + (43.2)^2 + (44.2)^2 + (46.0)^2 + (47.3)^2 + (53.9)^2}{6} \right] - \frac{(276.4)^2}{36} - 9.00 - 4.754$$

$$= 1.712$$

• حساب مجموع المربعات داخل الخلايا :

$$SS_{within} = \sum \text{كل الدرجات} X^2 -$$

$$\left[\frac{(\sum^{cel11} X)^2 + (\sum^{cel12} X)^2 + (\sum^{cel13} X)^2 + (\sum^{cel21} X)^2 + (\sum^{cel22} X)^2 + (\sum^{cel23} X)^2}{n_{cel}} \right]$$

$$= 2143.18 - \frac{(41.8)^2 + (43.2)^2 + (44.2)^2 + (46.0)^2 + (47.3)^2 + (53.9)^2}{6}$$

$$= 5.577$$

• حساب مجموع المربعات الكلى :

$$SS_{total} = SS_{column} + SS_{row} + SS_{within} + SS_{int}$$

$$= \sum \text{لكل الدرجات} X^2 - \frac{\left(\sum \text{لكل الدرجات} X \right)^2}{N}$$

$$= \sum X^2 - \frac{T^2}{N}$$

$$= 2143.18 - \frac{(276.4)^2}{36} = 21.042$$

ثانيًا: حساب تباينات أو متوسط مجموع المربعات:

$$MS_{\text{row}} = \frac{SS_{\text{rows}}}{df_{\text{rows}}} = \frac{9.0}{1} = 9$$

$$MS_{\text{coloumn}} = \frac{SS_{\text{col}}}{df_{\text{col}}} = \frac{4.74}{2} = 2.377$$

$$MSS_{\text{Int}} = \frac{SS_{\text{int}}}{df_{\text{int}}} = \frac{1.712}{2} = 0.856$$

$$MSS_{\text{within.cell}} = \frac{SS_{\text{w.cell}}}{df_{\text{w.cell}}} = \frac{5.577}{30} = 0.186$$

ثالثاً: حساب نسبة F:

لتأثير الصف (العامل الأول):

$$F = \frac{MS_{\text{row}}}{MS_{\text{within.cell}}} = \frac{9.00}{0.186} = 48.42$$

لتأثير العمود (العامل الثانى):

$$F = \frac{MS_{\text{coloumn}}}{MS_{\text{within.cell}}} = \frac{2.377}{0.186} = 12.78$$

لتأثير التفاعل:

$$F = \frac{MS_{\text{int}}}{MS_{\text{within.cell}}} = \frac{0.856}{0.186} = 4.60$$

5. القرار والتفسير: بالنسبة لتأثير الصف (توقيت التمارين):

F المحسوبة (48.42) < F الحرجة (4.170)

إذا يرفض H_0 وعليه لتوقيت ممارسه التمارين أثر على عدد ساعات النوم أو توجد فروق فى ساعات النوم بين ممارسة التمارين صباحاً ومساءً ويمكن عزوها إلى صالح المتوسط الأعلى وهو ممارسة التمارين مساءً.

بالنسبة لتأثير العمود (العامل B) وهو شدة التمارين:

F المحسوبة (12.78) < F الحرجة (3.32)

إذا يرفض H_0 وعليه توجد فروق ذات دلالة إحصائية فى ساعات النوم بين شدة التمارين الثلاثة (خفيف- متوسط -شديد) ولتحديد اتجاه الفروق يتم إجراء المقارنات المتعددة (لاحظ أكثر من متوسطين).

بالنسبة لتأثير تفاعل الصف X العمود (AXB):

$$F \text{ المحسوبة } (4.660) < F \text{ الحرجة } (3.32)$$

وعليه يرفض H_0 وبالتالي يوجد تأثير دال إحصائيًا لتفاعل شدة التمارين وتوقيتها على ساعات النوم ولتحديد أى المجموعات (الخلايا) التى أحدثت الفروق يتم إجراء المقارنات المتعددة وكذلك عرض النتائج فى ضوء الرسومات البيانية.

6. حجم التأثير: يتم قياس حجم التأثير فى ضوء مؤشرات عديدة أهمها:

• مؤشر η^2 أو R^2 ويكون كالآتى:

$$\begin{aligned} \eta^2_{\text{rows}} (A) &= \frac{SS_{\text{rows}}}{SS_{\text{total}}} \\ &= \frac{9}{21.042} = 4.42 \end{aligned}$$

• مؤشر η^2_{partial} :

$$\eta^2_{\text{P row}} = \frac{SS_{\text{row}}}{SS_{\text{row}} - SS_{\text{within.cell}}}$$

وبالنسبة لتأثير العمود (العامل الثانى):

$$\eta^2_{\text{column}} = \frac{SS_{\text{column}}}{SS_T} = \frac{4.754}{21.04}$$

بينما $\eta^2_{\text{P.col}}$:

$$\eta^2_{\text{P.col}} = \frac{SS_{\text{column}}}{SS_{\text{column}} + SS_{\text{within.cell}}}$$

وبالنسبة للتفاعل A X B:

$$\eta^2_{AXB} = \frac{SS_{int}}{SS_T} = \frac{1.712}{21.04}$$

بينما $\eta^2_{P.int}$:

$$\eta^2_{P.int} = \frac{SS_{int}}{SS_{int} + SS_{within.cell}}$$

• مؤشر ω^2 : مؤشر ω^2 لتأثير الصف (العامل A التوقيت)

$$\begin{aligned}\omega^2_{row} &= \frac{SS_{row} - df_{row}(MS_{within.cel})}{SS_T + MS_{within.cel}} \\ &= \frac{9 - (2-1)(0.186)}{21.42 + 0.186}\end{aligned}$$

ومؤشر ω^2 لتأثير العمود (العامل B الشدة):

$$\begin{aligned}\omega^2_{col} &= \frac{SS_{col} - df_{col}(MS_{w.cel})}{SS_T + MS_{w.cel}} \\ &= \frac{4.754 - (3-1)(0.186)}{21.42 + 0.186}\end{aligned}$$

ومؤشر ω^2 لتأثير تفاعل A , B (الشدة X التوقيت):

$$\begin{aligned}\omega^2_{int} &= \frac{SS_{int} - df_{int}(MS_{within.cel})}{SS_T + MS_{within.cel}} \\ &= \frac{1.712 + (2-1)(3-1)(0.186)}{21.42 + 0.186}\end{aligned}$$

• مؤشر Cohen f: يقدر كالاتى:

$$f = \sqrt{\frac{(k-1)(F-1)}{Nk}}$$

• مؤشر f لتأثير الصف (العامل الأول):

$$f = \sqrt{\frac{(2-1)(48.42-1)}{2 \times 36}} = .658$$

• مؤشر f لتأثير العمود (العامل الثانى):

$$f = \sqrt{\frac{(3 - 1)(12.78 - 1)}{3 \times 36}} = \sqrt{\frac{23.56}{108}} = 0.46$$

- مؤشر f للتفاعل:

$$f = \sqrt{\frac{(3 - 1)(2 - 1)(4.60 - 1)}{36 \times 3 \times 2}} = \sqrt{\frac{9.2}{216}} = 0.042$$

7. تقدير القوة الإحصائية لـ ANOVA باستخدام برنامج G-Power

لتحديد القوة الإحصائية للمثال السابق باستخدام برنامج G-Power اتبع الآتي:

أولاً: لتأثير الصف (التوقيت):

1. افتح البرنامج تظهر الشاشة الافتتاحية.

2. تحت Type of power analysis اختار:

Type of power analysis
Post hoc: Compute achieved power - given α , sample size, and effect size

لاحظ أننا اعتمدنا على التحليل البعدي Post hoc

3. أسفل Test family اختار F- test

4. أسفل Statistical test اختار:

Test family	Statistical test
F tests	ANOVA: Fixed effects, special, main effects and interactions

5. ادخل المعالم Input parameters:

- حجم التأثير المقدر (f) = 0.658 (حجم تأثير كبير)

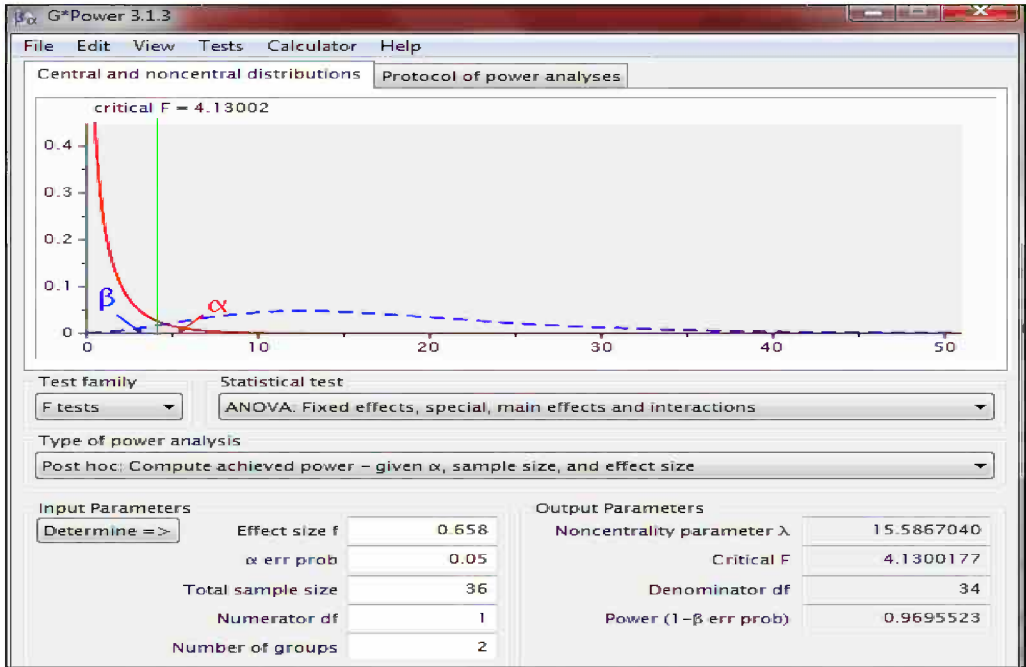
- مستوى الدلالة الإحصائية $\alpha = 0.05$

- عدد الصفوف K=2

- حجم العينة = 36

Input Parameters	
Determine =>	
Effect size f	0.658
α err prob	0.05
Total sample size	36
Numerator df	1
Number of groups	2

6.اضغط Calculated تظهر شاشة الناتج كالآتي:

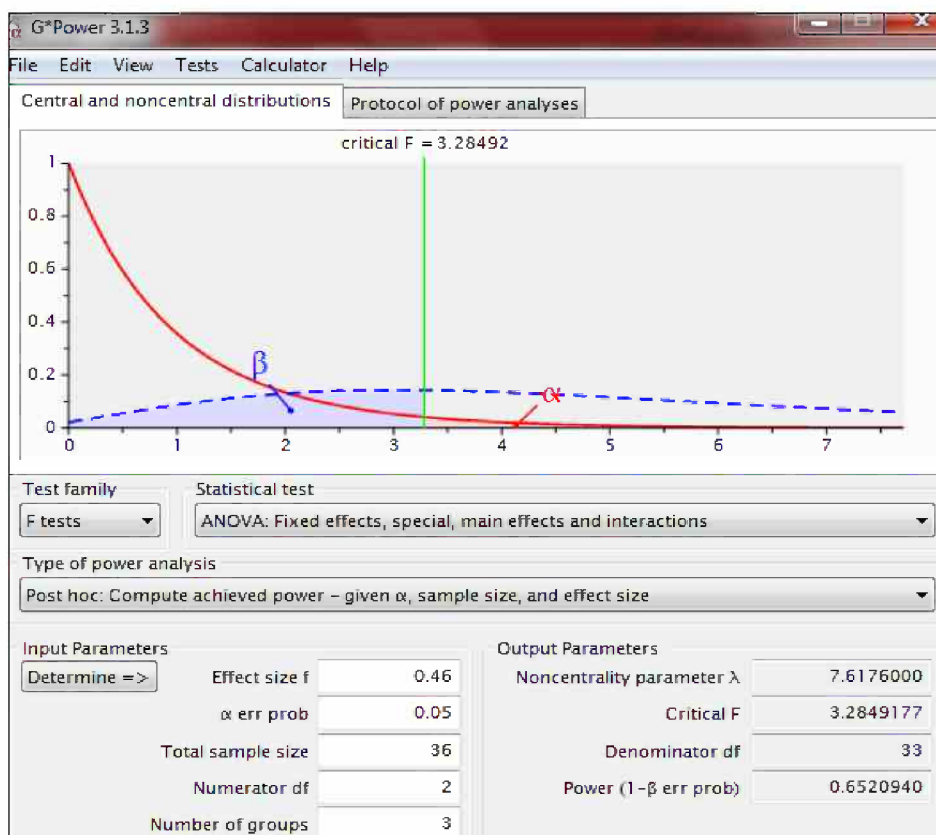


يتضح أن القوة الإحصائية = 0.969 وهو مستوى مرغوب.

ثانيًا: لتأثير العمود (شدة التمرين): اتبع الخطوات السابقة وتكون معالم المدخلات كالآتي:

Input Parameters	
Determine =>	
Effect size f	0.46
α err prob	0.05
Total sample size	36
Numerator df	2
Number of groups	3

وتكون المخرج كالآتي:

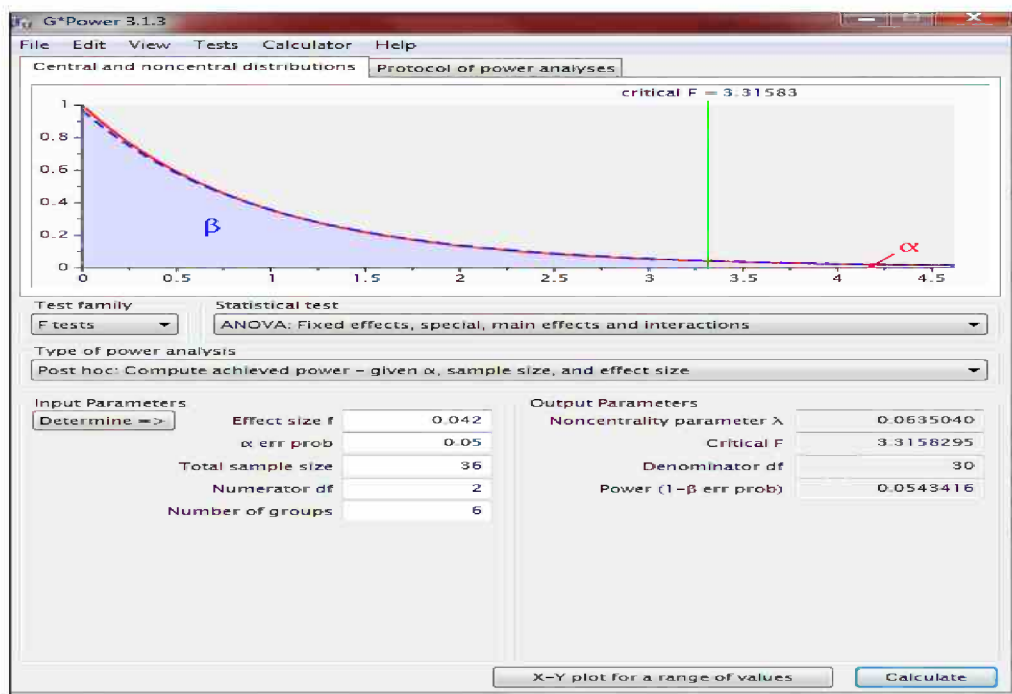


يتضح أن القوة الإحصائية = 0.652 وهذا مستوى متوسط ولكن لم يصل إلى المستوى المرغوب وهو 0.80.

ثالثاً: القوة الإحصائية للتفاعل: اتبع الخطوات السابقة بحيث تكون المعالم:

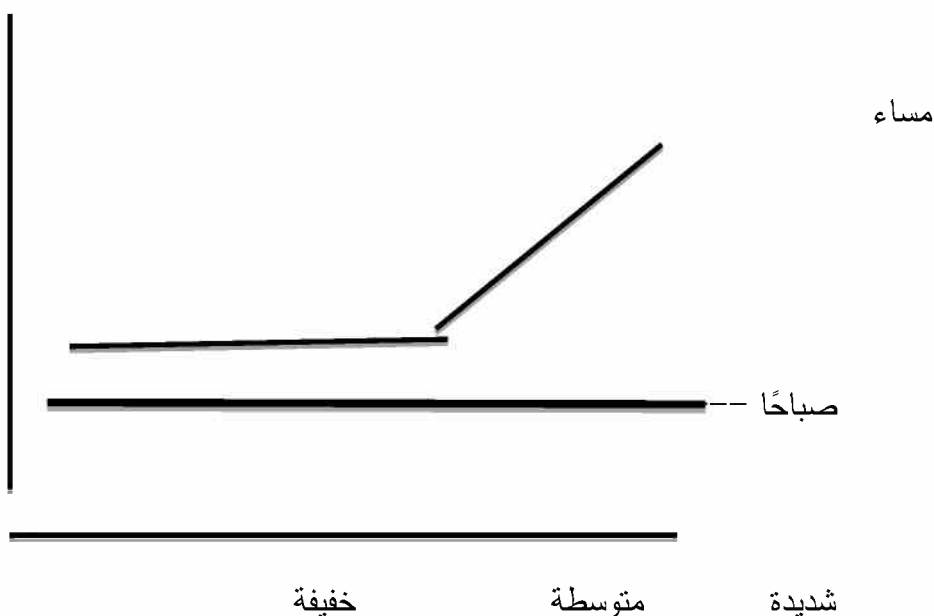
Input Parameters	
Determine =>	Effect size f
	α err prob
	Total sample size
	Numerator df
	Number of groups

ويكون المخرج كالاتى:



يتضح أن القوة الإحصائية = 0.0543 وهذا مستوى ضعيف جداً.

تفسير النتائج: يمكن تفسير الدلالة الإحصائية لتوقيت التدريب (تأثير الصف) وكذلك الدلالة الإحصائية لشدة التمرين (تأثير العمود) والدلالة الإحصائية لتفاعل الوقت وشدة التمرين (صف x عمودين) خلال الشكل الآتي:



يتضح من الشكل أن متوسط النوم مساءً أكبر من متوسط ساعات النوم جراء التدريبات صباحاً وان متوسط النوم جراء التدريبات شديداً أعلى من إجراء التدريبات متوسطة وأيضاً خفيفة. كما يتضح أن الخطئين ليس متوازيين مما يشير إلى وجود تفاعل بين توقيت التمرين وشدة حيث إن ساعات النوم تزيد من إجراء التدريبات الخفيفة إلى الشديدة عند إجراء التدريبات مساءً عن اجرائها في الصباح.

كتابة نتائج تحليل التباين ثنائي الاتجاه في تقرير البحث وفقاً لـ APA
تم إجراء تحليل التباين المتعدد (2x3) لتقييم أثر شدة التمرين وتوقيته على عدد ساعات النوم ليلاً، وكانت المتوسطات والانحرافات المعيارية لعدد ساعات النوم كدالة للعاملين في الجدول التالي المتوسطات والانحرافات المعيارية لساعات النوم:

Mean	الشدة	التوقيت
6.97	خفيف	صباحاً
7.20	متوسط	
7.37	شديد	
7.67	خفيف	مساءً
7.88	متوسط	
7.98	شديد	

واوضحت نتائج ANOVA 2 x3 وجود تأثيرات رئيسية دالة إحصائياً لـ التوقيت $F_{(2,30)}=12.78, P< 0.05, \eta^2=$ ولشدة التمرين $F_{(1,30)}=48.42, P< 0.05, \eta^2=$ ووجود تأثير لتفاعل التوقيت والشدة دالة إحصائياً $F_{(2,30)}=4.6, P< 0.05, \eta^2=$ أو تعرض النتائج في جدول كالآتي:

مصدر التباين	SS	df	MS	F	Fc.v
الصفوف (الوقت)	9.00	1	9.00	48.42*	4.17
الأعمدة (الشدة)	4.754	2	2.377	12.78*	3.32
التفاعل	1.712	2	0.856	4.60*	3.32
داخل الخلايا (الخطأ)	5.577	30	0.186		
الكلي	21.042	35			

*دالة عند مستوى الفا 0.05

قضية أخرى (حل آخر) (Hinkle et al. 1994): أهتم باحث بدراسة أثر طريقة التدريس (محاضرة - مناقشة - عصف ذهني) على التحصيل لدى الذكور والإناث فاختار عينة عشوائية من الإناث مكونة من 24 طالبة وعينة عشوائية من الذكور 24 طالب وتم تقسيم كل عينة إلى ثلاث مجموعات تلقت كل مجموعة ثلاث طرق تدريسية وكانت البيانات كالاتي:

طريقة التدريس					
عصف ذهني B ₃		مناقشة B ₂		محاضرة B ₁	
46	36	26	8	28	2
39	47	33	1	23	7
43	42	27	4	25	2
40	35	25	5	21	9
26	24	25	7	16	8
30	27	32	1	19	2
32	33	26	7	24	0
29	25	24	5	31	5

خطوات اختبارات الفروض الصفرية

لتحليل التباين تتأى الاتجاه ثلاثة فروض صفرية كالاتي:

1. الفروض الإحصائية:

● فرض صفري للصف (الجنس): $H_{01}: \mu_1 = \mu_2$

في المجتمع فإن متوسطات الذكور يساوي متوسطات تحصيل الإناث.

● فرض صفري للأعمدة (طريقة التدريس):

$$H_{02}: \mu_{.1} = \mu_{.2} = \mu_{.3}$$

$\mu_{.1}$: النقطة تشير إلى الصف

متوسطات تحصيل الطرق الثلاثة متساوية في المجتمع

● فرض التفاعل:

$$H_{03}: \text{All } (\mu_{JK} = \mu_{J.} = \mu_{.K}) \neq \mu$$

μ_{JK} متوسط الخلية في الصف J والعمود K، $\mu_{J.}$ متوسط الصفوف،

μ_K متوسط الأعمدة، μ المتوسط الكلي

أو يصاغ كالاتي:

$$H_{03}: \mu_{11} = \mu_{12} = \mu_{13} = \mu_{21} = \mu_{22} = \mu_{23}$$

وعليه يوجد ثلاثة فروض بديلة هي:

• الصف $H_{A1}: \mu_{1.} \neq \mu_{2.}$

$\mu_{1.}$ متوسط الصف الأول، $\mu_{2.}$ متوسط الصف الثاني

• الأعمدة $H_{A2}: \mu_{.1} \neq \mu_{.2} \neq \mu_{.3}$

$\mu_{.1}$ متوسط العمود الأول، $\mu_{.2}$ متوسط العمود الثاني

• للتفاعل $H_{A3}: \mu_{11} \neq \mu_{12} \neq \mu_{13} \neq \mu_{21} \neq \mu_{22} \neq \mu_{23}$

2. الاختبار الإحصائي المناسب ومسلماته:

يتم حساب ثلاث نسب لـ F وهي:

• لتأثير الصفوف (الجنس) (J): $F_J = \frac{MS_J}{MS_W}$

• لتأثير الأعمدة (K): $F_K = \frac{MS_K}{MS_W}$

• لتأثير التفاعل JK: $F_{JK} = \frac{MS_{JK}}{MS_W}$

3. مستوى الدلالة الإحصائية وقاعدة القرار: مستوى الدلالة الإحصائية $\alpha = 0.05$

وبالنسبة لقيمة F الجدولية للصف فإن: $df_J = J - 1 = 2 - 1 = 1$

$$df_W = Jk (n-1) = 3 \times 2 \times (8-1) = 42$$

وبالكشف في جدول F بدرجتى حرية 2 (البسط) و 42 (المقام) وبـ $\alpha = 0.05$ فإن:

$$F_{c.v} = 4.07$$

وبالنسبة لـ F الجدولية للعمود K: $df_K = K - 1 = 3 - 1 = 2$ ، $df_W = 42$

وبالكشف في جدول F فإن: $F_{c.v} = 3.22$

وبالنسبة لـ F الجدولية للتفاعل فإن: $df_{JK} = (J-1) (K-1) = (2-1) (3-1) = 2$

وبالكشف في جدول F بدرجتى حرية 2 (البسط) و 42 (المقام) فإن: $F_{c.v} = 3.22$

وبالتالى فإذا كانت F : المحسوبة $F < F$ الجدولية يرفض H_0

4. الحسابات:

طريقة التدريس				
عصف ذهني		مناقشة	محاضرة	
$T_{1.}=744$	$T_{13}=328$	$T_{12}=219$	$T_{11}=197$	الإناث
$\bar{X}_{1.} = 31.0$	$\bar{X}_{13} = 41$	$\bar{X}_{12} = 27.38$	$\bar{X}_{11} = 24.63$	
$\sum X^2=24622$	$\sum X^2=13580$	$\sum X_i^2=6065$	$\sum X_i^2=4977$	
$T_{2.}=618$	$T_{23}=226$	$T_{22}=217$	$T_{21}=175.0$	الذكور
$\bar{X}_{2.} = 25.75$	$\bar{X}_{23} = 28.25$	$\bar{X}_{22} = 27.13$	$\bar{X}_{21} = 21.88$	
$\sum X^2=16392$	$\sum X^2=6460$	$\sum X_i^2=5945$	$\sum X^2=3987$	
$T=1,362$	$T_{.3}=554$	$T_{.2}=436$	$T_{.1} = T_{11}+T_{21}=372$	
$\bar{X} = 28.38$	$\bar{X}_{.3} = 34.36$	$\bar{X}_{.2} = 27.25$	$\bar{X}_{.1} = \frac{197 + 175}{16} = 23.25$	
$\sum X^2=41014$	$\sum X^2=20040$	$\sum X^2=12010$	$\sum X^2=4977+3987=8.964$	

حيث مجموع درجات الصف: $\sum T_{j.}^2=(744)^2+(618)^2=935460$

مجموع درجات العمود: $\sum T_{.k}^2=(372)^2+(436)^2+(554)^2=635396$

مجموع درجات الخلايا:

$$\sum T_{j.k}^2=(197)^2+(219)^2+(328)^2+(175)^2+(217)^2+(226)^2=323144$$

مجموع كل الدرجات فى الخلايا: $T=\sum X_{ijk}=1,362$

مجموع مربع كل الدرجات فى الخلايا: $\sum X_{ijk}^2=41,014$

وعليه فان:

$$SS_j = \frac{1}{nK} \sum T_{j.}^2 - \frac{T^2}{N} = \frac{935460}{8 \times 3} - \frac{(1362)^2}{48} = 38977.50 - 38646.75 = 330.75$$

$$SS_k = \frac{1}{nJ} \sum T_{.K}^2 - \frac{T^2}{N} = \frac{635396}{8 \times 2} - \frac{(1362)^2}{48} = 39712.25 - 38646.75 = 1065.50$$

$$SS_{JK} = \frac{1}{n} \sum T_{JK}^2 - \frac{1}{nK} \sum T_j^2 - \frac{1}{nJ} \sum T_{.K}^2 + \frac{T^2}{N} = \frac{323144}{8} - \frac{935460}{8 \times 3} - \frac{635396}{8 \times 2} + \frac{(1362)^2}{48}$$

$$= 40393.00 - 38977.50 - 39712.25 + 38646.75 = 350.00$$

$$SS_w = \sum X_{ijk}^2 - \frac{1}{n} \sum T_{JK}^2 = 41014.00 - 40393.00 = 621.00 \quad \text{و}$$

$$SS_T = \sum X_{ijk}^2 - \frac{T^2}{N} = 41014.00 - \frac{(1362)^2}{48} = 41014.00 - 38646.75 = 2367.22$$

وعليه فإن:

$$F_j = \frac{MS_j}{MS_w} = \frac{330.75}{14.79} = 22.36$$

$$F_k = \frac{MS_k}{MS_w} = \frac{532.75}{14.79} = 36.07$$

$$F_{jk} = \frac{175.00}{14.79} = 11.83$$

5. القرار والتفسير: F المحسوبة $< F$ الجدولية، وعليه يتم رفض الفروض الصفرية الثلاثة وهى للتأثيرات الرئيسية (الجنس - طريقة التدريس). ويمكن عرض جدول ملخص كالاتى:

مصدر التباين	SS	df	MS	F	Fc.v
(الجنس)	330.75	1	330.75	22.36*	4.07
(طريقة التدريس)	1065.50	2	532.75	36.02*	3.22
التفاعل	350.00	2	175.00	11.83*	3.22
داخل الخلايا	621.00	42	14.79		
الكلية	2367.25	47			

*دالة عند 0.05

6. حساب حجم التأثير: تقدر من خلال مؤشر ω^2 وتقدر لتأثير الصفوف كالاتى:

$$\omega^2 = \frac{ss_J - (J - 1)MS_w}{SS_T + MS_w} = \frac{330.75 - (2 - 1)(14.79)}{2367.25 + 14.79} = 0.13$$

ومؤشر ω^2 لتأثير الأعمدة (طريقة التدريس):

$$\omega^2 = \frac{ss_K - (K - 1)MS_w}{SS_T + MS_w} = \frac{1065.50 - (3 - 1)(14.79)}{2367.25 + 14.79} = 0.44$$

ومؤشر ω^2 لتأثير التفاعل:

$$\omega^2 = \frac{ss_{JK} - (J - 1)(K - 1)MS_w}{SS_T + MS_w} = \frac{350 - (2)(14.79)}{2367.25 + 14.79} = 0.13$$

وعليه فإن البيانات اشارت إلى وجود حوالى 13% من تباين من درجات التحصيل ترجع إلى الفروق بين الذكور والإناث، و44% من تباين من درجات التحصيل ترجع إلى طريقة التدريس و13% من تباين من درجات التحصيل ترجع إلى التفاعل بين المتغيرات المستقلة.

إجراء المقارنات المتعددة للتأثيرات الرئيسية

بعد الحصول على دلالات إحصائية للتأثيرات الرئيسية وهى الجنس وطريقة التدريس ولمعرفة من المسبب لهذه الدلالة، فإذا كان المتغير بمستويين فيمكن تحديد ذلك من خلال المتوسط الأعلى، ولكن إذا كان المتغير مكوناً من ثلاث مستويات فأكثر فلا مفر من إجراء المقارنات المتعددة لتحديد أى زوج من أزواج المقارنات الذى أحدث دلالة إحصائية، وفى هذا المثال نطبق طريقة توكى اختبار فروض مرتبطة بطريقه بأزواج المقارنات لمتوسطات الأعمدة (طريقة التدريس) وعليه نختبر الفروض الصفرية الآتية:

$$H_{01}: \mu_1 = \mu_2 \quad , \quad H_{02}: \mu_1 = \mu_3 \quad , \quad H_{03}: \mu_2 = \mu_3$$

$$Q = \frac{\bar{X}_i - \bar{X}_k}{\sqrt{\frac{MS_W}{n}}} \quad \text{ولاختبار ذلك نستخدم طريقة توكى:}$$

• $\bar{X}_I$ متوسط العينة للمجموعة i ، $\bar{X}_K$ متوسط العينة للمجموعة K .

• MS_W متوسط المربعات داخل الخلايا.

• n عدد الملاحظات أو الأفراد في المجموعات التي يتم مقارنتها.

$$Q_{12} = \frac{\bar{X}_{.1} - \bar{X}_{.2}}{\sqrt{\frac{MS_W}{n}}} = \frac{23.25 - 27.25}{\sqrt{\frac{14.79}{8 \times 2}}} = 4.17: Q$$

وهكذا لباقي المقارنات وحيث درجات الحرية $df_w = 42$ و $\alpha = 0.05$ ، فإن $Q_{c.v} = 3.44$ ، ولـ

$\alpha = 0.01$ فإن $Q_{c.v} = 3.37$

ويتم تلخيص قيمة Q في الجدول الآتي:

طريقة التدريس	$\bar{X}_i$	$X_i - \bar{X}_k$	Q
المحاضرة	23.25		4.17*
المناقشة	27.25	4.00	
العصف الذهني	34.63	11.38, 7.38	11.85**, 7.69**

**دالة عند 0.01

*دالة عند 0.05

وعليه يتم رفض H_0 للمقارنات الثلاثة عند 0.05 ، بينما يتم رفض الفرضين الصفرين الثانى والثالث عند 0.01 ، وعليه يستطيع الباحث استنتاج أن المجموعة الثالثة (طريقه التدريس الثالثة) أكثر تحصيلاً من المجموعتين الأولى والثانية، وهذا يتضح من خلال قيمة المتوسط لهذه المجموعة أعلى من بقيه المجموعات وان المجموعة الثانية (طريقة المناقشة) أكثر تحصيلاً من طريقه المحاضرة.

تفسير التفاعل: كما نعلم أن التفاعل يحدث عندما يكون تأثير مستويات المتغير المستقل الأول ليست واحده غير مستويات المتغير المستقل الثانى وأحد الإجراءات لفحص التفاعل هو عرض الرسم البيانى لمتوسطات الخلايا حيث يوضع المتغير التابع على المحور الصادى وتوضع مستويات أحد المتغيرات المستقلة على المحور الافقى السينى ويعرض ذلك كالاتى:



وإذا لم يوجد تفاعل بين المتغيرين المستقلين فإن الخطين يكونان متوازيان أو اقرب للتوازي، ولكن في الشكل السابق فإن الخطين غير متوازيين وعليه يوجد تفاعل بين مستويات المتغيرين، ويلاحظ للإناث أن متوسط درجات التحصيل باستخدام الطريقة (العصف الذهني) أكبر بكثير من متوسطات تحصيل طريقة المحاضرة وأيضًا طريقة المناقشة، ولعينة الذكور نلاحظ أن متوسطات طريقة العصف الذهني والمناقشة أعلى من متوسط تحصيل طريقة المحاضرة وهذا يدل على أن تأثير طريقته التدريس ليست واحدة وليست بنفس الدرجة لعينة الذكور ولعينة الإناث، ولذلك يوجد تفاعل بين الجنس وطريقة التدريس وإذا كان التفاعل غير دال إحصائيًا نستنتج أن الرسم البياني يمثل خطين متوازيين.

اختبارات التأثيرات البسيطة Test of simple effects

في حاله الحصول على دلالة إحصائية لتأثير التفاعل فإن إجراء المقارنات البعدية في هذه الحالة يطلق عليها اختبار التأثيرات البسيطة وتستخدم مقرونة بالرسم البياني السابق لتفسير التفاعلات، وفي المثال الحالي يتضح أن الفروق بين الذكور والإناث ليست بنفس الدرجة عبر طرق التدريس الثلاثة واختبار التأثيرات البسيطة يعتبر الفروق بين متوسطات الخلايا داخل مستويات المتغيرين المستقلين ويستخدم لتحديد ما إذا كان متوسطات طرق التدريس الثلاثة تختلف للذكور وأيضًا تحديد ما إذا كان متوسطات الطرق

الثلاثة تختلف للإناث، وعلى ذلك يمكن إجراء ANOVA لمتوسطات الذكور وكذلك تحليل التباين الأحادي للإناث ويتم حساب نسبه F ويستخدم متوسطات المربعات داخل الخلايا (الخطأ) في تحليل التباين ثنائي الاتجاه في المقام:

الحسابات (في الصف الأول): يتم حساب مجموع مربعات التأثيرات البسيطة للإناث عبر الطرق الثلاثة في الصفوف:

$$SS_k \text{ (في الصف الأول)} = \sum \frac{(T_{jk})^2}{n} - \frac{T_{j.}^2}{nk}$$

$$SS_k \text{ (في الصف الأول)} = \frac{(197)^2}{8} + \frac{(219)^2}{8} + \frac{(328)^2}{8} - \frac{(744)^2}{24} \text{ وعليه:}$$

$$= 24294.25 - 23064.00 = 1230.25$$

$$SS_k \text{ (في الصف الثاني)} = \frac{(175)^2}{8} + \frac{(217)^2}{8} + \frac{(226)^2}{8} - \frac{(618)^2}{24}$$

$$= 16098.75 - 15913.50 = 185.5$$

ولحساب نسبة F للأعمدة وهو اختبار الفروق بين متوسطات الذكور والإناث في كل مستوى من مستويات المتغير المستقل طرق التدريس الثلاثة وتقدر مجموع المربعات كالاتي:

$$SS_J \text{ (في العمود K)} = \sum \frac{(T_{jk})^2}{n} - \frac{T_{.k}^2}{nJ}$$

$$SS_J \text{ (العمود الأول)} = \frac{(197)^2}{8} + \frac{(175)^2}{8} - \frac{(372)^2}{16} = 8679.25 - 8649$$

$$= 30.25$$

$$SS_J \text{ (العمود الثاني)} = \frac{(219)^2}{8} + \frac{(217)^2}{8} - \frac{(436)^2}{16} = 11881.25 - 11881.00$$

$$= 0.25$$

$$SS_J \text{ (العمود الثالث)} = \frac{(328)^2}{8} + \frac{(226)^2}{8} - \frac{(554)^2}{16} = 19832.50 - 19182.25 = 650.25$$

ويتم تلخيص الحسابات في الجدول الآتي:

المصدر	SS	df	MS	F	Fc.v
الصف (الجنس)	330.75	1	330.75	22.36	4.07
الصفوف فى العمود 1	30.25	1	30.25	2.05	6.73
الصفوف فى العمود 2	.25	1	.25	.02	6.73
الصفوف فى العمود 3	650.25	1	650.25	43.97	6.73
أعمدة (طريقة التدريس)	1065.50	2	532.75	36.02	3.22
الأعمدة فى الصف 1	1230.25	2	615.13	41.59	4.43
الأعمدة فى الصف 2	185.25	2	92.63	6.26	4.43
التفاعل	350	2	175	11.83	3.22
داخل الخلية	621	42	14.79		
الكلى	2367.25	47			

وعليه يوجد ثلاثة اختبارات لتأثيرات رئيسية بسيطة للصفوف واحد لكل صف واختبارين للتأثيرات الرئيسية البسيطة للأعمدة واحد لكل عمود وفى تحديد القيمة الحرجة لـ F لكل اختبار من الاختبارات السابقة ويتم قسمة مستوى الفا 0.05 على عدد التأثيرات البسيطة داخل التأثير الرئيسى، وفى هذا المثال يتم اختبار ثلاثة تأثيرات رئيسية للصفوف داخل كل عمود وعليه فأن:

$$\alpha = \frac{0.05}{3} = 0.017$$

وبالمثل التأثيرات الرئيسية للعمود (طريقة التدريس) فإن اختبار التأثيرات الرئيسية للعمودين الذكور والإناث هي اثنان:

$$\alpha = \frac{0.05}{2} = 0.025$$

والمشكلة الأساسية فى تحديد القيم الحرجة لـ F بهذه مستويات الدلالة حيث لا يتضمنها جداول F وللحصول على القيم الحرجة فنلجأ إلى التقريب فتصبح:

$$\alpha = 0.017 \leq 0.01$$

$$\alpha = 0.025 \leq 0.05$$

وعليه فإن قيمة F الحرجة للتأثيرات الرئيسية للصف عبر الطرق الثلاثة فأن:

$$F(1,42), \alpha = 0.05 = 4.07$$

ولتحديد قيمة F الحرجة للتأثير الرئيسى لكل عمود عبر الصفيين:

$$F(1,42), \alpha = 0.01 = 7.29$$

وبالتالى القيمة الحرجة لـ F عند $\alpha = 0.017$ هي 7.29 وعند $\alpha = 0.05$ هي 4.07 , ولتحديد هذه القيمة تكون كالاتى:

0.05	0.05	7.29
- 0.01	- 0.017	- 4.07
0.04	0.033	3.22

ولتقدير F الحرجة عند $\alpha = 0.017$ فيجب أضافة:

$$F_{c.v} = 4.07 + \left(\frac{0.033}{0.04}\right)(3.22) = 4.07 + 2.66 = 6.73$$

ولتقدير F الحرجة للتأثيرات الرئيسية لكل صف فى الأعمدة الثلاثة فأن:

$$F_{c.v}(2,42), \alpha = 0.05 = 3.22$$

$$F_{c.v}(2,42), \alpha = 0.01 = 5.16$$

وعليه فإن $F_{c.v}$ لـ $\alpha = 0.025$ هي بين 3.22 , 5.16 وبالتالى:

0.05	0.05	5.16
-0.01	-0.025	-3.22
0.04	0.025	1.94

$$F_{c.v} = 3.22 + \left(\frac{0.025}{0.04}\right)(1.94) = 3.22 + 1.21 = 4.43$$

وعليه فإن النتائج كما فى الجدول السابق لتحليل التأثيرات البسيطة للصفوف فى كل عمود تشير إلى أن طريقة التدريس المثالية (العصف الذهنى) وتوجد فروق دالة إحصائية بين متوسط درجات تحصيل الذكور ($\bar{x} = 28.25$) متوسط درجات تحصيل الإناث وأشارت نتائج التأثيرات البسيطة للأعمدة ($\bar{x} = 41$).

وأشارت نتائج تحليل التأثيرات البسيطة فى كل صف بوجود فروق فى متوسط التحصيل فى كل من طرق التدريس الثلاثة بين الذكور والإناث ولتحديد أى متوسطات احدثت الفروق تستخدم طريقه توكى للمقارنات البعدية.

تنفيذ تحليل التباين ذو اتجاهين أو العاملى فى SPSS مثال Pagano(2013)

أولاً: إدخال البيانات: 1. اضغط Variable view

2. تحت عمود name اكتب مسمى المتغيرات وهى كالاتى:

المتغير المستقل الأول وهو وقت التمرين time ومستوياته صباحاً = 1، ومساءً = 2.

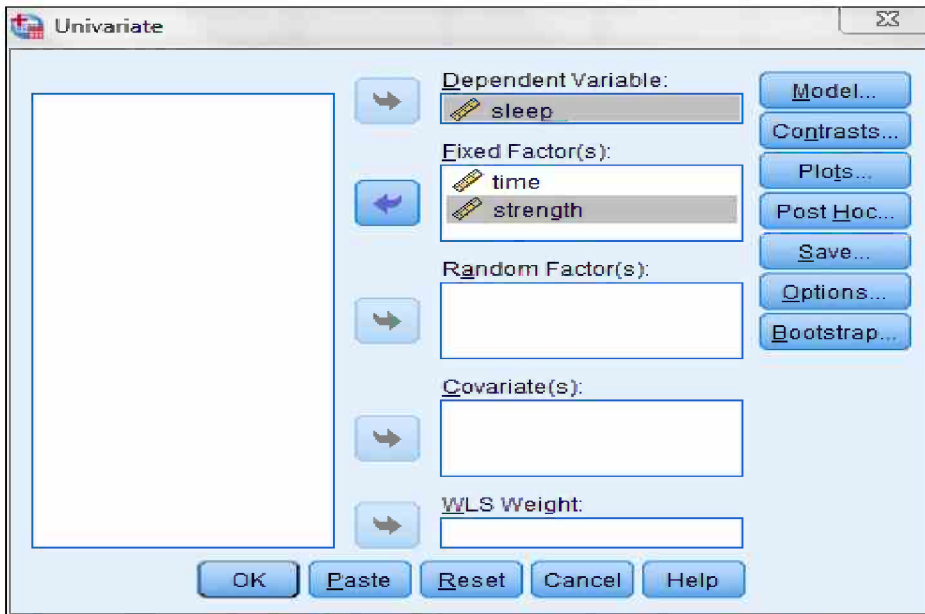
المتغير المستقل الثانى وهو شدة التمرين strength ومستوياته هى: خفيف = 1، متوسطة = 2، شديدة = 3.

المتغير التابع وهى مدة النوم sleep.

3. اضغط على عمود values وحدد مسميات الأكواب.

4. اضغط data view تظهر شاشة إدخال البيانات ابدأ فى إدخال البيانات.

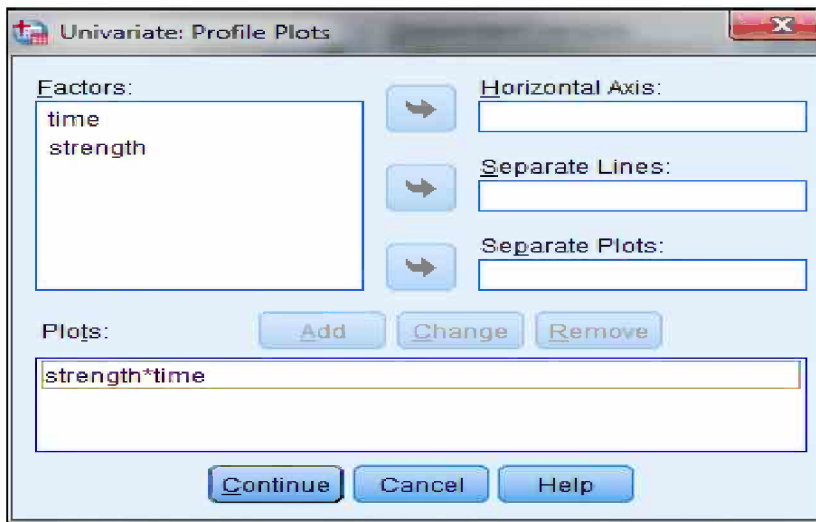
ثانياً: تنفيذ الامر: 1. اضغط على Analyze ثم General Linear Model ثم univariate تظهر الشاشة الآتية:



2. انقل المتغير التابع sleep إلى مربع Dependent variable.

3. انقل المتغيرات المستقلة strength و time إلى مربع Fixed factors.

4. اضغط اختيار Plots لتوضيح الرسم البياني للتفاعلات بين مستويات المتغيرات المستقلة. يظهر الشاشة التالية:



5. انقل المتغير المستقل ذو المستويات الثلاثة شدة التمرين strength وانقله إلى مربع

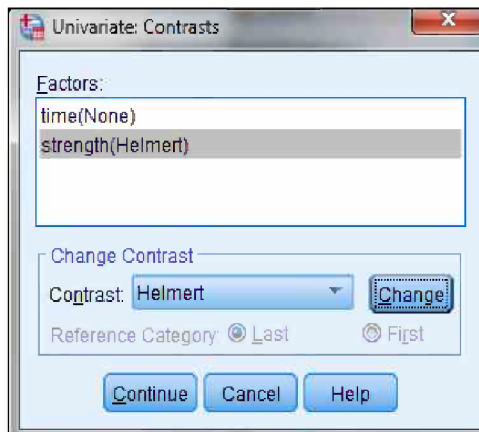
Horizontal axis.

6. اختار المتغير المستقل الثانى time وانقله إلى مربع Separate lines.

7. اضغط على add و plots وإذا كان يوجد متغير مستقل الثالث انقله إلى مربع .Separatelines

8. اضغط Continue.

9. اضغط على اختبار المقارنات :Contrasts



10. في مربع Change contrast امام

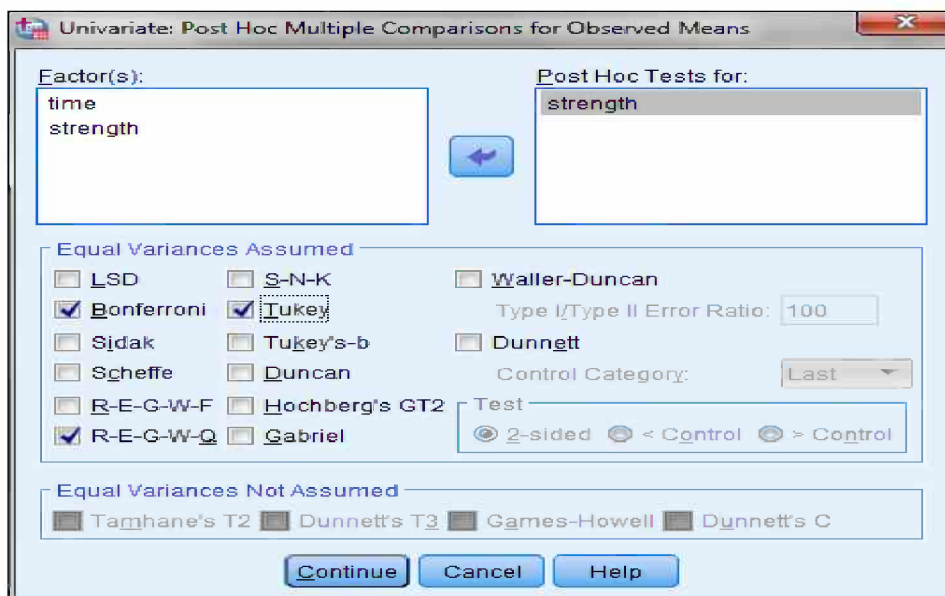
Contrast اضغط على السهم الأسفل تظهر عدة بدائل. اختار نوع المقارنات حسب

طبيعة الدراسة وسبق توضيحها وهنا يمكن اختبار نوع المقارنة helmert وفيها تحدث المقارنة بين كل مستوى (ما عدا الأخير) بمتوسط المستويات اللاحق له. والمقارنة تحدث للمتغير المستقل ذو الثلاث مستويات فاكتر وليس له داعى لإجراء مقارنة للمتغير المستقل بمستويين لأنها يمكن التعرف الدلالة لصالح المتوسط الأعلى.

11. اضغط Change.

12. اضغط Continue.

13. اضغط على المقارنات البعدية Post Hoc تظهر الشاشة التالية:

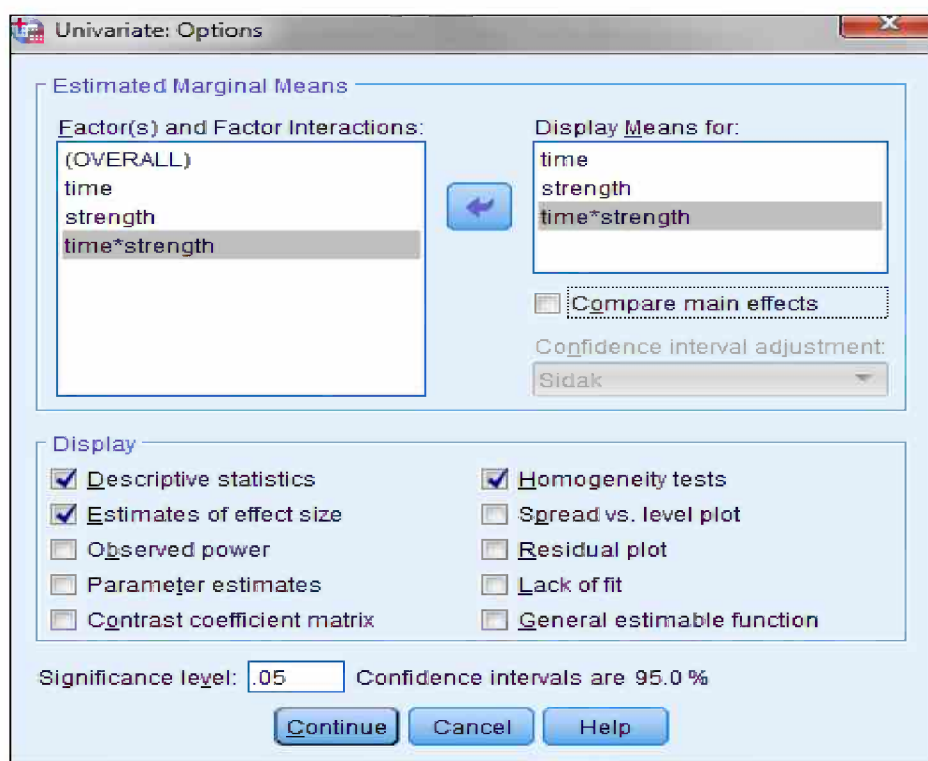


14. انقل المتغير المستقل المراد إجراء اختبارات المقارنة البعدية بين مستوياته وعليه انقل strength إلى مربع Post Hoc tests ولا داعي لنقل متغير time لأنه بمستويين فقط.

15. اضغط على Benferroni أو Tukey أو R.E.G.W.G.

16. اضغط Continue.

17. اضغط اختيار Options على يمين الشاشة تظهر الشاشة الآتية:



18. انقل متغيرات strength و time و time*strength إلى مربع Display means.

19. اضغط Descriptive statistics و Estimate of effect size و Homogeneity test. ويمكن الضغط على effectcompare mean لإجراء المقارنات المتعددة باستخدام LSD أو Bonferroni أو Sidak لكن لا داعي لذلك لأنه تم تنفيذهم من خلال اختبار PostHoc.

20. اضغط على Continue ثم اضغط على Ok.

ثالثاً: تفسير المخرج:

Between-Subjects Factors			
		Value Label	N
time	1.00	صباحاً	18
	2.00	مساءً	18
strength	1.00	خفيفة	12
	2.00	متوسطة	12
	3.00	شديدة	12

الجدول الآتى:

اعطى كل متغير مستقل بمستوياته وعدد الأفراد فى كل مستوى ووضح أن أسلوب ANOVA يعطى القيم 2X3. وعليه يوجد

سنة خلايا للمتغير التابع ساعات النوم وجدول الإحصائيات الوصفية كالاتى:

Descriptive Statistics				
Dependent Variable: sleep				
time	strength	Mean	Std. Deviation	N
صباحاً	خفيفة	6.9667	.38297	6
	متوسطة	7.2000	.36332	6
	شديدة	7.3667	.50067	6
	Total	7.1778	.42917	18
مساءً	خفيفة	7.6667	.37238	6
	متوسطة	7.8833	.31252	6
	شديدة	8.9833	.59133	6
	Total	8.1778	.72400	18
Total	خفيفة	7.3167	.51316	12
	متوسطة	7.5417	.48140	12
	شديدة	8.1750	.99282	12
	Total	7.6778	.77537	36

واعطى المتوسطات والانحرافات المعيارية للدرجات فى كل مستوى من مستويات المتغيرات المستقلة بمعنى متوسطات الخلايا الستة، فمثلاً فى خلية توقيت صباحاً وممارسة رياضة شديدة فإن المتوسط = 7.366 والانحراف المعيارى = 0.500 وحجم

العينة 6 وكذلك أعطى الإحصائيات الوصفية للعدد الكلى من العينة.

Levene's Test of Equality of Error Variances ^a			
Dependent Variable: sleep			
F	df1	df2	Sig.
1.775	5	30	.148
Tests the null hypothesis that the error variance of the dependent variable is equal across groups.			
a. Design: Intercept + time + strength + time * strength			

اختبار تجانس التباينات كالاتى:

يتضح اختبار مسلمة تجانس التباينات Leven's test للتحقق من وجود فروق بين

تباينات المجموعات عبر المستويات ويتضح أن: $F = 1.775, P = 0.148 > 0.05$

وعليه نقبل الفرض الصفري القائل بتساوي التباينات عبر المجموعات الستة وعليه تحققت مسلمة استخدام ANOVA.

جدول ANOVA الرئيسي كالاتي:

Tests of Between-Subjects Effects					
Dependent Variable: sleep					
Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	15.466 ^a	5	3.093	16.640	.000
Intercept	2122.138	1	2122.138	11416.163	.000
time	9.000	1	9.000	48.416	.000
strength	4.754	2	2.377	12.787	.000
time * strength	1.712	2	.856	4.604	.018
Error	5.577	30	.186		
Total	2143.180	36			
Corrected Total	21.042	35			

وهذا الجدول يحمل إجابات لأسئلة وفروض الدراسة ويتضح بالنسبة لتأثير الوقت يتضح أن: $F(1,30) = 48.416, P = 0.000 < 0.05$ ويرفص H_0 وعليه توجد فروق ذات دلالة إحصائية في ساعات النوم بين ممارسة التمارين صباحًا و مساءً. و 61.7% من تباين ساعات النوم ترجع إلى وقت التمرين.

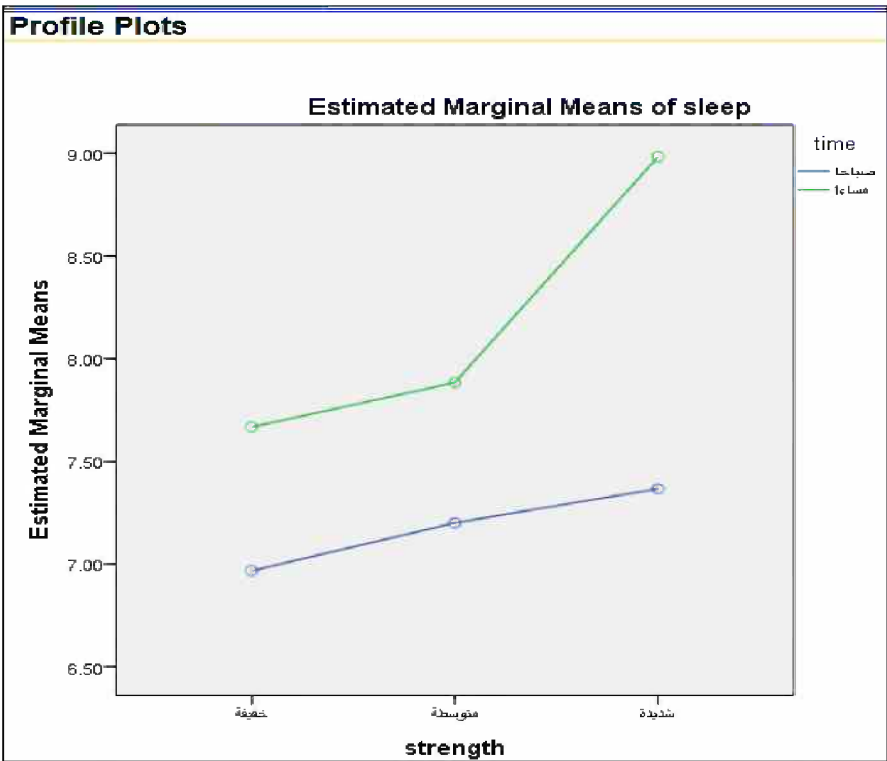
Partial Eta Squared
.735
.997
.617
.460
.235

وبالنسبة لأثر شدة التمرين يتضح أن: $F(2,30) = 12.377, P = 0.00 < 0.05$

وعليه توجد فروق ذات دلالة إحصائية في ساعات النوم بين شدة التمارين (شديدة – متوسطة – خفيفة) و 46% من تباين ساعات النوم ترجع إلى شدة التمرين.

وبالنسبة لتأثير التفاعل بين وقت التمرين وشدة على ساعات النوم $F(2,30) = 4.604, P = 0.018 < 0.05$ ، و توجد دلالة إحصائية لتفاعل وقت التمرين وشدة على ساعات النوم وهذا التفاعل فسر 23.5% من تباين ساعات النوم.

مما يؤكد وجود هذا التفاعل وهو شكل profile plot كالاتي:



حيث يظهر من الشكل أن الخطين ليس متوازيين حيث إذا كان التمارين بشدة صباحًا تختلف عن ساعات النوم عن ممارسة التمارين مساءً وبشدة، أي أن ساعات النوم تختلف باختلاف تفاعل مستويات المتغيرات المستقلة.

جدول المقارنات المخطط لها المخطط لها Contrasts كالاتي:

Contrast Results (K Matrix)		
		Dependent Variable
strength Helmert Contrast		sleep
Level 1 vs. Later	Contrast Estimate	-.542-
	Hypothesized Value	0
	Difference (Estimate - Hypothesized)	-.542-
	Std. Error	.152
	Sig.	.001
	95% Confidence Interval for Difference	Lower Bound Upper Bound -.853- -.230-
Level 2 vs. Level 3	Contrast Estimate	-.633-
	Hypothesized Value	0
	Difference (Estimate - Hypothesized)	-.633-
	Std. Error	.176
	Sig.	.001
	95% Confidence Interval for Difference	Lower Bound Upper Bound -.993- -.274-

لاحظ أننا اختارنا Helmert حيث المقارنة بين كل مجموعة والتي تليها ولاحظ أننا اجرينا المقارنات للمتغير المستقل بثلاثة مستويات.

وللمقارنة بين المستوى الأول الممارسة الخفيفة والمستوى التالى وهى الممارسة المتوسطة فى ساعات النوم يتضح وجود فروق دالة إحصائياً:

$$\text{Contrast estimate} = -0.542, \text{Sig} = 0.001 < 0.05$$

كما يتضح وجود فروق دالة إحصائياً فى ساعات النوم بين ممارسة النشاط بدرجة متوسطة وبدرجة خفيفة:

$$0.05 < \text{Contrast estimates} = -.633, \text{Sig} = .0015$$

واعطى البرنامج المقارنات المتعددة باستخدام الطرق الآتية:

Multiple Comparisons						
Dependent Variable: sleep						
(I) strength	(J) strength	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
Tukey HSD	منوسطة	-.2250-	.17602	.418	-.6589-	.2089
	شدیده	-.8583 [*]	.17602	.000	-1.2923-	-.4244-
	منوسطة	.2250	.17602	.418	-.2089-	.6589
	شدیده	-.6333 [*]	.17602	.003	-1.0673-	-.1994-
	منوسطة	.8583 [*]	.17602	.000	.4244	1.2923
	شدیده	.6333 [*]	.17602	.003	.1994	1.0673
Scheffe	منوسطة	-.2250-	.17602	.451	-.6783-	.2283
	شدیده	-.8583 [*]	.17602	.000	-1.3116-	-.4051-
	منوسطة	.2250	.17602	.451	-.2283-	.6783
	شدیده	-.6333 [*]	.17602	.005	-1.0866-	-.1801-
	منوسطة	.8583 [*]	.17602	.000	.4051	1.3116
	شدیده	.6333 [*]	.17602	.005	.1801	1.0866
Bonferroni	منوسطة	-.2250-	.17602	.633	-.6713-	.2213
	شدیده	-.8583 [*]	.17602	.000	-1.3047-	-.4120-
	منوسطة	.2250	.17602	.633	-.2213-	.6713
	شدیده	-.6333 [*]	.17602	.003	-1.0797-	-.1870-
	منوسطة	.8583 [*]	.17602	.000	.4120	1.3047
	شدیده	.6333 [*]	.17602	.003	.1870	1.0797
Based on observed means. The error term is Mean Square(Error) = .186. *. The mean difference is significant at the .05 level.						

وفى ضوء مقارنات Tukey اتضح وجود فروق دالة إحصائيا بين الممارسة الخفيفة والشديدة، وبين الممارسة المتوسطة والشديدة، فى ساعات النوم وبكل تأكيد هذه الفروق لصالح المتوسطات الأعلى وكذلك عدم وجود فروق بين الممارسة الخفيفة والممارسة المتوسطة وافرت مقارنات Sheffe و Bonferroni نفس نتائج طريقة Tukey.

الفصل الخامس

تحليل التباين

Analysis of Covariance(ANCOVA)

يعتبر تحليل التباين اتساع لتحليل التباين حيث يتم دراسة التأثيرات الرئيسية والتفاعلات لمتغيرات مستقلة على متغير تابع متصل بعد تصحيح درجات المتغير التابع من الفروق المرتبطة أو التداخلات من متغير أو أكثر دخيل أو مجموعة من المتغيرات Covariates التي تقاس قبل المتغير التابع وترتبط معه، كأن يتم دراسة أثر البرنامج (المتغير المستقل) على التحصيل بعد تصحيح درجات التحصيل من متغير تباينى أو متلازم معه ومرتبطة به مثل الذكاء. بكلمات أخرى يسمح بتقدير ما إذا كانت توجد فروق بين المجموعات على متغير تابع متصل وحيد بعد ضبط أثر متغير مستقل فأكثر متصل وتسمى تباينات.

والأسئلة التي يحاول أن يجيب عنها تحليل التباين هي نفسها مثل أسئلة تحليل التباين هل توجد فروق في متوسط التحصيل بين المجموعة الضابطة والتجريبية في القياس البعدى بعد تصحيح درجات القياس البعدى من تداخل متغير تباينى مثل الذكاء؟ ويعتبر تحليل التباين هو حلقة وصل أو منطقة عبور من الطرق البسيطة Univariate and Bivariate methods مثل تحليل التباين وتحليل الانحدار

المتعدد (متغير تابع وحيد) إلى الطرق المتدرجة المعقدة مثل تحليل التباين المتدرج (MANOVA) وتحليل الارتباط المتعدد Canonical Correlation (أكثر من متغير تابع). وعليه فإن ANCOVA يسمح بالتعامل مع متغير مستقل فأكثر تصنيفى ومتغير تابع متصل وحيد بالإضافة إلى متغير مستقل فأكثر متصل (تباين واحد فأكثر متصل) و تحليل التباين هو أسلوب يدمج تحليل التباين (ANOVA) وتحليل الانحدار المتعدد Multiple Regression (MR). وفى ANOVA يتعامل مع المتغيرات المستقلة التصنيفية وأثرها على المتغير التابع، بينما فى الانحدار المتعدد يحدد أثر المتغيرات المتصلة (التباينات) ومدى ارتباطها بقوة مع المتغير التابع.

وينظر إلى ANCOVA على أساس أنه مثل تحليل معامل الارتباط الجزئي Partial Correlation حيث يهدف إلى دراسة العلاقة بين X_1 و Y بعد حذف تباين X_2 على Y . وفي تحليل التباين يتم دراسة العلاقة بين المتغير المستقل (التصنيفى) والمتغير التابع المتصل بعد ضبط تأثير المتغير التباينى (المتصل) على المتغير التابع.

أهداف تحليل التباين

يستخدم تحليل التباين لثلاثة أهداف رئيسية (Tabachnick & Fidell, 2007):

- زيادة الدقة والحساسية لاختبار التأثيرات الرئيسية والتفاعل عن طريق تقليل قيمة الخطأ Error term وذلك من خلال تصحيحه أو تقليله عن طريق استبعاد العلاقة بين المتغير التابع والمتغيرات التباينية.

- تصحيح المتوسطات Adjusted Means للمتغير التابع بعد استبعاد المتغيرات الدخيلة.

- الاستخدام الثالث لتحليل التباين يحدث فى تحليل التباين المتدرج (MANOVA)

كاختبار تتبعى ملازم لتحليل التباين المتدرج MANOVA الدال إحصائيا حيث كل

متغير تابع فى MANOVA يستخدم لاحقاً كمتغير وحيد فى التحليل باستخدام

ANCOVA وكل المتغيرات التابعة الأخرى تستخدم كتباينات. فمثلاً إذا اتضح

وجود فروق دالة إحصائية بين المجموعات الثلاثة على متغيرين تابعين فى

MANOVA ثم لاحقاً يتم إجراء ANCOVA لأحد المتغيرين التابعين بينما

الآخر يكون كمتغير تباينى. وهذا يفيد فى اختبار الفروق بين المجموعات على

متغير تابع عبر أى فروق على المتغير التابع الآخر.

والاستخدام الأول هو الأكثر استخداماً فى المواقف التجريبية حيث يزيد أسلوب

ANCOVA القوة الإحصائية لاختبار F للتأثيرات الرئيسية وللتفاعلات عن طريق

حذف أو استبعاد التباين المتوقع المرتبط بمتغير التباينى من مجموع مربعات الخطأ

(البواقي).

ويستخدم التباين لتقدير التشويش المتداخل وهو تباين غير مرغوب فى المتغير التابع

ويتم تقديره من خلال درجات المتغير التباينى وبالتالي تظهر قوة تحليل التباين فى

تقليل تباين الخطأ وكما نعلم أن تحليل التباين في طبيعته يعبر عنه في ضوء معادلة حيث يوجد متغيرات ثنائية مستقلة تصنيفية Dummy وكما نعلم أن الانحدار المتعدد يتضمن العديد من المتغيرات المنبئة المتصلة ولذلك فليس من المستغرب أن يتم اتساع معادلة الانحدار لتحليل التباين لتشمل متغير متصل مستقل فأكثر يستخدم للتنبؤ بالمتغير التابع وهذه المتغيرات المستقلة ليست جزء من المعالجة التجريبية ولكنها تؤثر على ناتج التجربة ولذلك يعرف هذه المتغيرات بالتغايرات Covariates.

يوجد سببين لتضمين التغايرات في تحليل التباين هما:

- تقليل تباين الخطأ أو التباين داخل المجموعات حيث يمكن تفسير بعض التباينات التي لم تفسر ويطلق عليها التباين غير المفسر (التباين داخل المجموعات SS_w) وبذلك يتم تقدير دقيق لتأثير المتغير المستقل.
 - استبعاد التداخلات Confounds ففي أي تجربة فإن المتغيرات غير المتضمنة في التحليل تتداخل مع النتائج ولذلك فإن ANCOVA يساعد على استبعاد هذه التدخلات من هذه المتغيرات عن طريق قياسها وإدخالها في التحليل كتغايرات مما يساعد على إجراء ضبط إحصائي.
- وفي تحليل التغاير يتم تقدير انحدار متغير تغايري واحد فأكثر في المتغير التابع وتصحيح متوسط درجات المتغير التابع من التأثيرات الخطية للمتغير التغايري قبل التحليل، وعليه يتم الحصول على خط الانحدار يربط بين المتغير التابع DV ومتغير التغايري CV. ويمكن استخدام المتغيرات التغايرية في كل تصميمات تحليل التباين بين الأفراد وداخل الأفراد وداخل وبين الأفراد (المختلطة).

الأسئلة البحثية في تحليل التغاير

كما هو الحال في ANOVA فإن تحليل التغاير يهدف إلى تحديد ما إذا كانت توجد فروق في متوسطات المتغير التابع بين المجموعات ويهدف إلى دراسة:

1. التأثيرات الرئيسية للمتغيرات المستقلة: حيث تتم دراسة التغيرات في المتغير التابع جراء المستويات المختلفة لأي متغير مستقل تصنيفي: هل يتأثر قلق الاختبار

نتيجة المعالجة أو البرنامج بعد ضبط الفروق في العمر والحالة الاقتصادية (تغايرات) ؟ ويكون الفرض الصفري لا أثر للمتغير المستقل على المتغير التابع، ومع وجود أكثر من متغير مستقل فإنه توجد اختبارات إحصائية لكل متغير مستقل على حده كأن يضاف إلى السؤال السابق. هل يتأثر قلق الاختبار نتيجة للمعالجة والجنس بعد ضبط الفروق في العمر والحالة الاقتصادية؟

2. **التفاعلات بين المتغيرات المستقلة:** عندما يتضمن التصميم أكثر من متغير مستقل تصنيفي فالباحث يرغب في تحديد ما إذا كان التغير في السلوك عبر مستويات أحد المتغيرات المستقلة يعتمد على التغير في مستويات المتغير المستقل الآخر. فهل توجد فروق في قلق الاختبار بين المجموعة التجريبية والضابطة في القياس البعدي عبر المستويات المختلفة من المتغير المستقل الآخر (ذكور- إناث) بعد ضبط العمر والحالة الاقتصادية؟

3. **المقارنات المحددة وتحليل الاتجاه Specific comparison and trend analysis:** عندما تظهر دلالات إحصائية في التصميم لمتغيرات مستقلة تتضمن أكثر من مستويين فمن الأفضل تحديد طبيعة اتجاه الفروق كما هو الحال في المقارنات المتعددة في ANOVA . فأى لمجموعات احدثت الدلالة إحصائياً؟ فعندما يتغير مستوى القلق فلا بد من تحديد أى المجموعتين التجريبيتين (تصميم ثلاث مجموعات ، مجموعتين تتلقى معالجتين مختلفتين والمجموعة الثالثة ضابطة) أكثر كفاءة وفاعلية في تقليل القلق مقارنة بالمجموعة الضابطة بعد ضبط وتصحيح درجات قلق الاختبار من المتغيرات التغايرية وكذلك يمكن تحديد أى من المجموعتين التجريبيتين أكثر كفاءة في تخفيض قلق الاختبار وهذا يتحقق من خلال إجراء المقارنات المتعددة المخطط لها مسبقاً في ضوء الإطار النظري أو مقارنات متعددة بعدية يتم إجرائها بعد الحصول على دلالة إحصائية بغض النظر عن وجود نظرية أو عدم وجودها .

4. **تأثيرات التغايرات Effects of covariates:** تحليل التغاير قائم على الانحدار الخطي بين المتغير التغايري (CV) والمتغير التابع (DV) ولكن لا يوجد ضمان

للحصول على دلالة إحصائية لمعامل الانحدار وقيم الانحدار إحصائياً عن طريق اختيار المتغيرات التغايرية باعتبارها مصدر يسهم في تباين درجات المتغير التابع ويمكن اعتبار أن قلق الاختبار في القياس القبلي متغير تغايري لدرجات القلق البعدي، وعليه يمكن دراسة أثر الدرجات القبلية للقلق على الدرجات البعدية.

5. حجم التأثير: إذا كان التأثيرات الرئيسية والتفاعلات دالة إحصائياً وحدث تغير في المتغير التابع فالسؤال المنطقي هو تحديد حجم هذا التأثير أو التغير؟ بمعنى ما هو حجم التباين المفسر في المتغير التابع المصحح من المتغيرات التغايرية يرجع إلى المتغيرات المستقلة؟

6. تقديرات المعالم Parameter estimates: بعد الحصول على دلالة إحصائية للتأثيرات الرئيسية والتفاعلات فيجب تحديد ما هي معالم المجتمع المقدرة مثل المتوسط المصحح والانحراف المعياري وفترات الثقة لكل مستوى من المتغير المستقل. ففي قلق الاختبار إذا وجد تأثير رئيسي للمعالجة فما هو المتوسط المصحح لدرجات القلق البعدي في كل مجموعة من المجموعات في التصميم (تصميم المجموعتين أو تصميم المجموعات الثلاثة).

أنماط الارتباطات في تحليل التغير

يتطلب تحليل التغير أنماط من الارتباطات بين المتغير التابع والمتغير التغايري كالآتي:

1. المتغيرات التغايرية يجب أن ترتبط بدرجة متوسطة Moderately على الأقل مع المتغير التابع ($r > 0.30$). وإذا كان الارتباط بينهما صغير فإنه نسبة إسهامه في المتغير التابع صغيرة، وإذا كان الارتباط بين المتغير التغايري والمتغير التابع جوهرياً ولكن ليس كبيراً ودالاً إحصائياً فإن هذا يزيد الحساسية نتيجة تقليل تباين الخطأ.

2. المتغيرات التغايرية لا بد أن تكون عالية الثبات حيث معامل الثبات 0.7 فأكثر لان العلاقة بين المتغيرات غير الثابتة تكون أقل قوة وغير دقيقة وفي البحوث التجريبية إذا كان التغير غير ثابت فإن هذا يؤدي إلى فقد القوة وبالتالي اختبار إحصائي

متحفظ وفى المواقف غير التجريبية يؤدي إلى تقدير أقل أو أعلى للمتوسطات المصححة، وفى البحوث غير التجريبية يفضل أن يكون ثبات التغير أكبر من 0.80.

3. يجب أن تكون العلاقات بين التغيرات ضعيفة إذا كان تحليل التغير يتضمن متغيرات تغايرية عديدة وذلك لان وجود العلاقات الارتباطية المرتفعة يؤدي إلى ظهور إشكالية التلازمية الخطية Collinearity.

متى يستخدم تحليل التغير ANCOVA؟

يوجد ثلاث استخدامات رئيسية لتحليل التغير يحددها (Harlow 2005):

1. التصميمات التجريبية: يستخدم تحليل التغير فى التصميم التجريبى لضبط بعض المتغيرات المتصلة الدخيلة التى لم يتم ضبطها مثل درجات الاختبار القبلى للمتغير التابع أو الذكاء أو غيرها. وفى التصميم التجريبى يتم معالجة متغير تجريبى تصنيفى قد يكون بمستويين أو أكثر ويتم توزيع أفراد العينة عشوائيًا على المجموعات المختلفة. فمثلا عند دراسة فعالية برنامج للإقلاع عن التدخين لثلاث مجموعات حيث تتلقى كل مجموعة معالجة مختلفة فإن القياسات القبلية للتدخين ومتغيرات دخيلة أخرى مثل الضغوط وعدد أفراد الأسرة المدخنين هى تغيرات. وعلى ذلك فإن تحليل التغير مفضل عن تحليل التباين عند إدخال متغيرات دخيلة فى التصميم وهذا يمكن أن يساعد فى عملية الضبط الإحصائى Statistical Control للمتغيرات الدخيلة التى أحيانًا من الصعب ضبطها تجريبياً.

2. التصميمات غير أو شبه التجريبية: يمكن استخدام ANCOVA للضبط الإحصائى للفروق فى متغير أو أكثر فى التصميمات غير أو شبه التجريبية.

أحد المشاكل المرتبطة بهذا الاستخدام هو عدم وجود طريقة مناسبة لتضمين أو اعتبار كل المتغيرات الدخيلة وعليه فإن الضبط الإحصائى يكون مضللاً.

محددات تحليل التغير

قضايا نظرية: كما هو الحال فى ANOVA فإن الاختبار الإحصائى ليس ضماناً للتأكيد على أن التغيرات فى المتغير التابع حدثت نتيجة للمتغير المستقل فاستنتاج السببية هو عملية منطقية نظرية وليست قضية إحصائية وتعتمد على الأسلوب الذى يتم من خلاله توزيع الأفراد على مستويات المتغير المستقل ومعالجة المستويات المختلفة للمتغير المستقل عن طريق الباحث والية الضبط أو التحكم فى البحث.

والمنطقية فى اختيار التغيرات هو أن القاعدة العامة فى اختيار عدد محدود من CV ترتبط مع المتغير التابع ولا يوجد ارتباطات بين المتغيرات التغيرية CV. والهدف هو الحصول على أقصى توافق للمتغير التابع مع ادنى فقد لدرجات الحرية للخطأ (داخل المجموعات فى تحليل التباين)، وفى العمل التجريبي يوجد تحذير وهو أن التغيرات يجب أن يكون مستقل عن المعالجة وهذا يفرض علينا أن يتم جمع بيانات التغيرات قبل إجراء المعالجة وعدم الالتزام بهذا يؤدي إلى استبعاد جزء من تأثير المتغير المستقل على المتغير التابع. ويوجد مصادر للتحيز فى تحليل التغيرات ولا يمكن ادعاء الاستنتاج أو ادعاء السببية باستخدام ANCOVA فى التصميمات شبة التجريبية لان توزيع أفراد العينة على المجموعات غير عشوائى.

ويتم اختيار التغيرات بناءً على أسس نظرية وليس أى متغير يوضع فى التحليل كتغيرات حيث من المفترض أن يكون له دور فى تفسير المتغير التابع، وإذا لم تتوفر أسس نظرية فيكون الاختيار على أسس إحصائية عن طريق معامل الارتباط المرتفع مع المتغير التابع، وإذا توافرت قياسات عديدة للتغيرات يفضل استخدام طريقة الانحدار خطوة خطوة لانتقاء أهم التغيرات، ويفضل البساطة فى اختبار التغيرات حيث إن العدد المحدود أفضل من العدد الكبير.

قضايا عملية

1. أحجام العينات غير المتساوية Unequal sample sizes: إذا وجد بيانات مفقودة على المتغير التابع فى تحليل التغيرات بين الأفراد (مجموعات مستقلة) فإن هذا يولد مشكلة عدم تساوى أحجام العينات عبر مستويات المتغير المستقل أو عبر تفاعلات هذه المستويات وتظهر هذه الإشكالية بوضوح إذا كانت البيانات المفقودة على

المتغير التباينى وعلى المتغير التابع خاصة فى نموذج ANCOVA داخل الأفراد، ويجب أن يكون حجم العينة فى كل خلية كافى للحصول على مستوى قوة مناسب خاصة أن إضافة التباينات فى التحليل يزيد من القوة إذا كانت مرتبطة بالمتغير التابع، وتوجد آليات للتعامل مع البيانات المفقودة .

2. القيم المتطرفة **Oulters** : داخل درجات كل مجموعة يجب فحص عدم وجود قيم متطرفة فى المتغير التابع أو التباين لأن وجود القيم المتطرفة المتدرجة Multivariate outliers بين المتغير التابع والتباينات معاً يؤدي إلى تحطم مسلمة تجانس الانحدار Homogeneity regression.

3. استقلالية التباين عن تأثير المعالجة وكما هو الحال فى ANOVA حيث يقسم تأثير التجربة إلى جزئين هما تأثير المعالجة (بين المجموعات) والخطأ أو التباين غير المفسر (داخل المجموعات) وهى تأثير عوامل لم يتم قياسها وتؤثر فى المتغير التابع وفى ANCOVA حيث يسهم التباين فى تفسير جزء من التباين غير المفسر بكلمات أخرى يوجد استقلال تام لهذا التباين المفسر عن طريق التباين عن التباين المفسر عن طريق المعالجة. وإذا تداخل التباين المفسر عن طريق التباين مع التباين المفسر عن طريق المعالجة وفى هذا الموقف لا يجب استخدام ANCOVA ، وفى هذا الشأن فإن تأثير التباين يقلل تأثير التجربة ويخفى تأثير التجربة ويصبح تفسير نتائج تحليل التباين فى غاية الخطورة، ويزداد هذا الموقف خطورة عندما لا يتم توزيع الأفراد عشوائياً على مجموعات التصميم، فمثلاً يرتبط القلق بالاكتئاب ارتباطاً عالياً فعند دراسة أثر برنامج على تخفيض القلق فانك تعتقد انك تضيف الاكتئاب كتباينات لضبط تأثير الاكتئاب ولكن هذا غير صحيح.

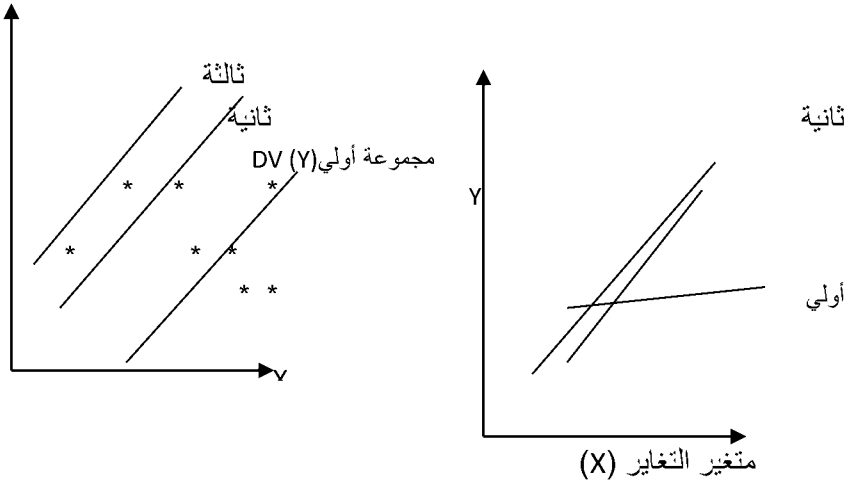
4. التلازمة الخطية المتعددة **Multicollinearity** : إذا كانت العلاقات بين التباينات كبيرة فيجب حذف أحد التباينات لأن تضمينهم معاً فى التحليل لن يحدث أى تصحيح للمتغير التابع ويحدث مشاكل حسابية أثناء تقدير نتائج ANCOVA . وعليه إذا زاد مربع معامل الارتباط المتعدد Squared multiple correlation (SMC) لـ أى متغير تباينى مع بقية المتغيرات عن 0.5 فيجب حذفه من التحليل.

5. **الاعتدالية:** يجب أن يتسم توزيع متوسطات درجات كل مجموعة على المتغير التابع بالاعتدالية وهنا يجب التأكيد أن التوزيع للمتوسطات وليس للدرجات الخام وفى حالة وجود عينات صغيرة أو غير متساوية وقيم متطرفة فمن الضروري تحويل للبيانات.

6. **تجانس التباينات:** كما هو الحال فى تحليل التباين يفترض تساوى تباينات المتغير التابع داخل كل خلية، وأيضًا أن يكون تباينات المتغير التباينى عبر المجموعات متساوية وإذا لم تتحقق هذه المسلمة للتباين فإنه يوجد بدليلين اما اختبار الدلالة الإحصائية للتأثيرات الرئيسية والتفاعلات عند مستوى ألفا صارم مثلاً 0.025 بدلا من 0.05 أو استبعاد CV من التحليل.

7. **الخطية:** تحليل التباين قائم على مسلمة أساسية وهى أن العلاقة بين المتغير التابع و المتغير وكذلك العلاقة بين كل زوج من المتغيرات التباينية هى خطية كما هو الحال فى تحليل الانحدار. وعدم توفر هذه المسلمة يقلل القوة للاختبار الإحصائى ويحدث أخطاء فى صناعة القرار الإحصائى، ووجود علاقة خطية منحنية Curvilinearity يمكن تصحيحها بتحويل بعض من المتغيرات التباينية أو حذف التباين CV الذى يحدث هذه العلاقة المنحنية.

8. **تجانس الانحدار Homogeneity of Regression:** يفترض تساوى معاملات الانحدار فى كل خلية (المتغير التباينى ← المتغير التابع) ، وهذه المسلمة تعنى أن خط الانحدار بين المتغير التابع والمتغير التباينى داخل كل خلية هو تقدير لنفس معامل الانحدار فى المجتمع وهذا يفترض أن خطوط الانحدار متساوية فى كل الخلايا، بمعنى أن معاملات الانحدار أو أوزان بيتا (β) هى نفسها فى كل خلايا التصميم، وعدم تجانس الانحدار يعنى وجود خط انحدار ($CV \rightarrow DV$) فى بعض الخلايا مختلف عن خلايا أخرى أو وجود تفاعل بين المتغير المستقل التصنيفى المعالجة والمتغير التباينى وإذا حدث هذا التفاعل فإن العلاقة بين المتغير التابع ومتغير التباين تكون مختلفة عبر مستويات المتغير المستقل. وفيما يلى شكل لثلاثة مجموعات حيث يوجد تجانس تام للانحدار (Equality of slopes):



(B): عدم تجانس خط الانحدار (A): تجانس خط الانحدار

الشكل (2.5): شكل الانتشار الميل الانحدار في المجموعات الثلاث

اختبارات الفروض لقضية بحثية: (Tabachnick & Fidell, 2013)

قام باحث بعمل برنامج علاجي لتسعة اطفال يعانون من صعوبات التعلم وقام بتقسيم العينة عشوائياً إلى ثلاث مجموعات يتلقون معالجات مختلفة (مجموعتين تجريبيتين ومجموعة ضابطة) وعدد الأفراد كل مجموعة ثلاثة ولكل فرد من التسعة درجتين هي الدرجة القبلية (المتغير التبايري) والدرجة البعدية (المتغير التابع) على نفس الاختبار (وليكن التحصيل) و عليه فإن درجات التحصيل القبلي هي التباير ولأن الهدف من ANCOVA هو تصحيح درجات التحصيل البعدي من الفروق الفردية في التحصيل قبل تأثرها بالمعالجة فيما يلي البيانات:

المجموعة					
المعالجة 1		المعالجة 2		الضابط	
قبلي	بعدي	قبلي	بعدي	قبلي	بعدي
85	100	86	92	90	95
80	98	82	99	87	80
92	105	95	108	78	82
257	303	263	299	255	257

(يمكن أن يكون التباين الذكاء مثلاً) ويصبح الهدف تصحيح التحصيل البعدى من الفروق الفردية فى الذكاء قبل تأثرها بالمعالجة ومن المتوقع أن يقلل هذا التباين من تباين درجات التطبيق البعدى داخل كل خلية من مجموعات المعالجة وهذا يؤدي إلى اختبار قوى لاختبار الفروق بين المعالجات

الخطوات البحثية: 1. سؤال البحث: هل توجد فروق فى درجات التحصيل البعدى بين المجموعات الثلاثة بعد تصحيح هذه الفروق من درجات التحصيل القبلى أو بعد استبعاد أثر التحصيل القبلى ؟

2. الفروض البحثية: توجد فروق فى درجات التحصيل البعدى بين المعالجات أو المجموعات الثلاثة بعد تصحيح هذه الفروق من أثر درجات التحصيل القبلى.

3. التصميم البحثى: تصميم تجريبى ذو ثلاث مجموعات، مجموعتين تجريبيتين ومجموعة ضابطة.

4. المتغيرات: حيث يوجد متغيرين مستقلين ومتغير تابع هما:

البرنامج أو التجربة: مستقل - تصنيفى - ثلاث مجموعات، التحصيل القبلى: مستقل - متصل (التباين)، التحصيل البعدى: تابع - متصل.

5. النموذج الإحصائى: النموذج الإحصائى المتعدد، والاختبار الإحصائى هو تحليل التباين وهو إحصاء بارامترى.

خطوات اختبارات الفروض الصفرية

1. الفروض الإحصائية:

• الفرض الصفرى (H_0): متوسطات المجتمعات الثلاثة متساوية بعد تصحيحها من التباين

$$H_0: \bar{\mu}_1 = \bar{\mu}_2 = \bar{\mu}_3$$

و كذلك يوجد فرض صفرى لمتغير التباين $H_0: \rho = 0$:

- معامل الارتباط بين المتغير التابع والتغاير في المجتمع و حيث لا يوجد علاقة ارتباطية بين التحصيل البعدي والتحصيل القبلي في المجتمع.

- $\bar{\mu}_1$ المتوسط للمجتمع بعد تصحيحه من التغاير.

• الفرض البديل (H_A):

$$H_A: \bar{\mu}_1 \neq \bar{\mu}_2 \neq \bar{\mu}_3$$

$$H_A: \bar{\mu}_1 = \bar{\mu}_2 \neq \bar{\mu}_3 \quad \text{أو}$$

الفرض البديل الخاص بمتغير التغاير:

$$H_A: \rho \neq 0$$

2. الاختبار المناسب ومسلّماته: يستخدم اختبار ANCOVA سبق عرض مسلماته.

فيما يلي ملخص لجدول ANCOVA:

df	مجموع المربعات (SS)	المصدر
C	SS _{cov}	التغاير
K-1	SS' _B	بين المجموعات
N-k-c	SS' _w	داخل المجموعات
N-1	SS' _T	الكلية

حيث C عدد التغايرات ، SS' مجموع المربعات المصححة من التغاير

و يتم حساب قيمة F للتغاير كالتالي:

$$F = \frac{MS_{cov}}{MS'_w}$$

و درجات الحرية: df = C = 1

وحساب قيمة F للمتغير المستقل أو المعالجة:

$$F = \frac{MS'_B}{MS_w}$$

ودرجات الحرية: $df = K-1 = 3-1 = 2$

ودرجات حرية الخطأ أو داخل المجموعات:

$$df = n - K - C = 9 - 3 - 1 = 5$$

3. مستوى الدلالة الإحصائية وقاعدة القرار: تبني الباحث $\alpha = 0.05$ ، وعليه بالكشف في جدول F بـ $\alpha = 0.05$ ، ودرجتى حرية 5 و 2 فكانت : $F_{c.v} = 5.79$ (الجدولية).

وإذا كانت: F (المحسوبة) للمعالجة بين المجموعات $F < F$ (الجدولية) نرفض H_0 ، وبالكشف في جدول F بدرجتى حرية 5 ، 1 ، و مستوى الدلالة الإحصائية 0.05

وإذا كانت: F المحسوبة $< F$ الحرجة

نرفض H_0 يوجد تأثير لمتغير التغيرات (التحصيل القبلى) على المتغيرات التابع.

4. الحسابات : المعادلات فى تحليل التغيرات هى اتساع لمعادلات تحليل التباين وفى تحليل التغيرات يوجد تأثير للمتغير المستقل التصنيفى و عليه فإن التباين الكلى يكون كالاتى:

$$SS_{total} = SS_{B.G} + SS_{Within}$$

$SS_{B.G}$: مجموع المربعات بين المجموعات Between groups

SS_{Within} مجموع المربعات داخل المجموعات Within groups.

وكذلك فى ANCOVA يوجد مكونين اضافيين من مجموع المربعات الأول يرجع إلى الفروق بين درجات التغيرات CV و المتوسط المشترك (GM) و يقسم إلى مجموع مربعات بين و داخل المجموعات كالاتى :

$$SS_{total} = SS_{B.G(x)} + SS_{W.G(x)}$$

وعليه فإن مجموع مربعات الفروق بين متوسط درجات المتغير التابع والمتوسط المشترك (GM) grand mean تقسم إلى مكونين مجموع مربعات الفروق بين متوسطات المجموعات (Y) و المتوسط المشترك GM، و مجموع مربعات الفروق بين الدرجات المفردة (Yij) و متوسطات المجموعات .

وبعد الحسابات وعليه فإن جدول تحليل التباين ANCOVA كالآتي:

مصدر التباين	مجموع المربعات المصححة	df	MS	F
بين المجموعات	366.202	2(k-1)	183.101	6.13
داخل المجموعات	149.439	5(n-k-c)	29.888	

$$P < 0.05$$

أما جدول تحليل التباين (ANOVA) لأثر البرنامج قبل التصحيح من درجات التحصيل القلبي كالآتي:

مصدر التباين	مجموع المربعات المصححة	df	MS	F
بين المجموعات	432.889	2(k-1)	216.444	4.52
داخل المجموعات	287.333	6(n-k-c)	47.889	

ولدرجات حرية 2 ، 6 ، ومستوى دلالة إحصائية 0.05. فإن: قيمة F الجدولية = 5.14.

لاحظ من جدول ANCOVA وجدول ANOVA أن مجموع المربعات في ANOVA أكبر بكثير من مجموع المربعات في ANCOVA خاصة للخطأ (داخل المجموعات) وكذلك من مجموع درجات الحرية الخاصة بداخل المجموعات في ANOVA تزيد بمقدار واحد عن درجات الحرية الخاصة بالخطأ أو داخل المجموعات في ANCOVA وكذلك تم رفض H_0 في ANCOVA بينما تم قبوله في ANOVA.

5. القرار والتفسير: F المحسوبة (6.13) $< F$ الحرجة (5.79) وعليه نرفض الفرض الصفري حيث توجد فروق فى التحصيل بين المجموعات الثلاثة فى القياس البعدى بعد تصحيحها من درجات التحصيل القبلى.

اما لتأثير البرنامج على التحصيل بدون حذف تأثير القياس القبلى للتحصيل (ANOVA) فان: F المحسوبة (4.52) $< F$ الحرجة (5.14) وعليه يتم قبول H_0 بمعنى لا توجد فروق فى درجات التحصيل البعدى بين المجموعات الثلاثة

6. حجم التأثير: يتم تقدير حجم التأثير باستخدام:

- مؤشر η^2_p :

$$\eta^2_p = \frac{SS'_{BG}}{SS'_{BG} + SS'_{within}} = \frac{SS_{effect}}{SS_{effect} + SS_{Residual}} = \frac{336.202}{336.202 + 149.438} = 0.71$$

وعليه فإن 71% من تباين التحصيل البعدى المصحح من التحصيل القبلى يرجع إلى

المعالجة أو البرنامج وهو حجم تأثير كبير

كما يمكن استخدام مؤشر η^2 كالاتى:

$$\eta^2 = \frac{SS'_{BG}}{SS_{Total}} = \frac{SS_{effect}}{SS_{Total}}$$

ولكن η^2_p أكثر دقة وتقدير η^2 .

- مؤشر f لـ Cohen:

$$f = \sqrt{\frac{(k-1)(F-1)}{NK}} = \sqrt{\frac{(3-1)(6.13-1)}{9 \times 3}} = 0.616$$

7. تقدير القوة الإحصائية لـ ANCOVA باستخدام برنامج G-Power

لتحديد القوة الإحصائية للمثال السابق باستخدام برنامج G-Power اتبع الاتى:

1. افتح البرنامج تظهر الشاشة الافتتاحية.

2. تحت Type of power analysis اختيار:

Type of power analysis

Post hoc: Compute achieved power - given α , sample size, and effect size

لاحظ اننا اعتمدنا على التحليل البعدي Post hoc

3. تحت Testfamily اختبار F- test

4. تحت Statistical test اختبار :

Test family

F tests

Statistical test

ANCOVA: Fixed effects, main effects and interactions

5. ادخل المعالم Input parameters :

Input Parameters

Determine ==>

Effect size f

0.616

α err prob

0.05

Total sample size

9

Numerator df

2

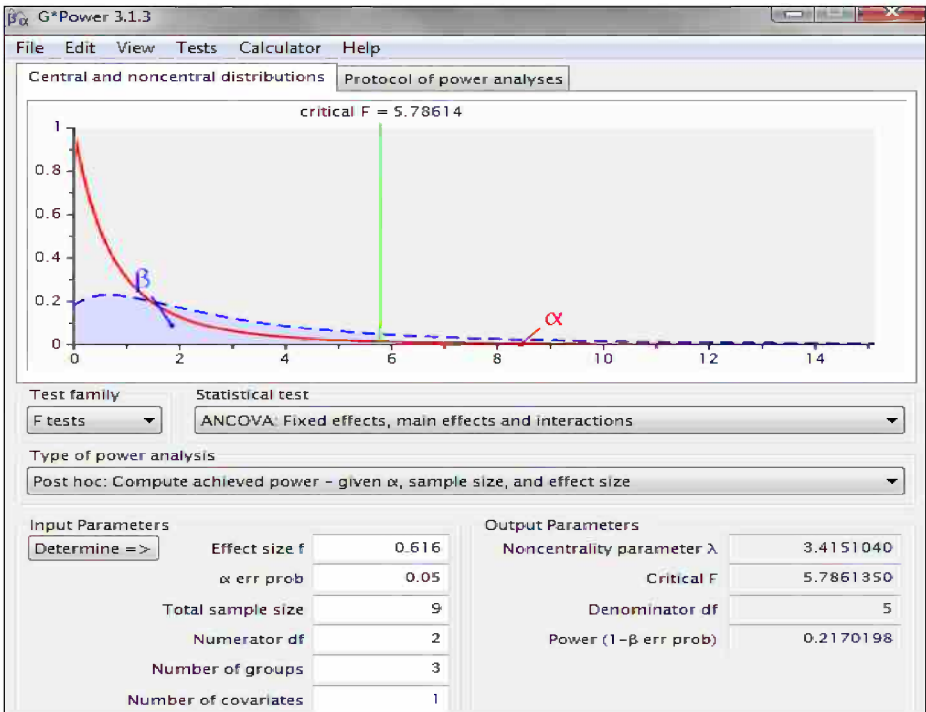
Number of groups

3

Number of covariates

1

6.اضغط Calculated تظهر شاشة الناتج كالآتى:



يتضح أن القوة الإحصائية = 0.217 وهذه قوة ضعيف وهذا يرجع إلى صغر حجم العينة.

تنفيذ ANCOVA في SPSS لبيانات مأخوذة من (Hinkle et al. 1994):

اراد الباحث دراسة أثر ثلاثة استراتيجيات للتدريس على تحصيل التلاميذ بعد ضبط متغير الذكاء وكانت درجات التحصيل البعدي والذكاء فى الاستراتيجيات التدريسية كالآتى:

المجموعة 1		المجموعة 2		المجموعة 3	
X	Y	X	Y	X	Y
98	60	104	62	102	65
102	63	109	63	117	68
104	66	104	67	108	72
103	69	117	71	117	76
112	72	120	77	105	78
113	75	113	79	116	80
118	78	117	82	111	82
120	80	126	84	120	84
115	67	113	64	107	75
106	70	109	68	104	77
112	74	125	72	128	79
122	76	118	77	119	81

أولاً : إدخال البيانات فى برنامج SPSS كالآتى:

1. اضغط على variable view أسفل الشاشة الافتتاحية للبرنامج .

2. اكتب مسمى المتغيرات تحت عمود Name حيث تم تسمية ثلاثة متغيرات:

الذكاء intelligence ، التحصيل achievement ، المجموعة group حيث المجموعة الأولى

تأخذ كود 1 ، والثانية كود 2، والثالثة كود 3.

ثانياً : التحقق من مسلمات ANCOVA:

- **استقلالية المتغير المستقل والتغاير:** كما سبق عرضة قبل إجراء تحليل التباين لا بد من فحص استقلالية المعالجة التجريبية وعلى ذلك لا بد من فحص ما إذا كان المتغير التباينى له تأثير متساوى عبر مستويات المتغير المستقل ولذلك يتم إجراء

تحليل التباين أحادى الاتجاه باعتبار أن المتغير الذكاء تابع ومتغير المعالجة (المجموعة) مستقل و يتم كالاتى:

Analyze --> Compare means- → One way ANOVA

والمخرج كالاتى:

ANOVA					
intelligence					
	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	105.056	2	52.528	.902	.415
Within Groups	1921.500	33	58.227		
Total	2026.556	35			

ويظهر من المخرج السابق أن التأثير الرئيسى للمجموعة غير دال إحصائيًا

$F(2,33)=0.902$ ($P=0.415$)

وهذا يعنى لا فروق بين المجموعات الثلاثة فى الذكاء وعليه فإنه يمكن استخدام متغير الذكاء كمتغير تباينى لأنه مستقل عن التجربة.

- اختبار العلاقة الجوهرية بين المتغير التابع والتباين: يجب أن تكون العلاقة على الأقل متوسطة 0.30 فأعلى وان لا تكون علاقة ضعيفة وللتحقق من هذا الشرط تم إجراء امر العلاقة كالاتى:

Analyze→ Correlate → Bivariate →

والمخرج كالاتى:

Correlations			
		achievement	intelligence
achievement	Pearson Correlation	1	.641**
	Sig. (2-tailed)		.000
	N	36	36
intelligence	Pearson Correlation	.641**	1
	Sig. (2-tailed)	.000	
	N	36	36

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

ويظهر من المخرج أن معامل الارتباط بين التحصيل والذكاء: $r=0.641$ ($P=0.000$) وهى علاقة قوية ودالة إحصائيًا.

- اختبار مسلمة تجانس ميل الانحدار: يفترض أن ميل الانحدار بين المتغير التباينى (الذكاء) والمتغير التابع التحصيل متساوية عبر المجموعات الثلاثة ويكون الفروض الإحصائية كالاتى:

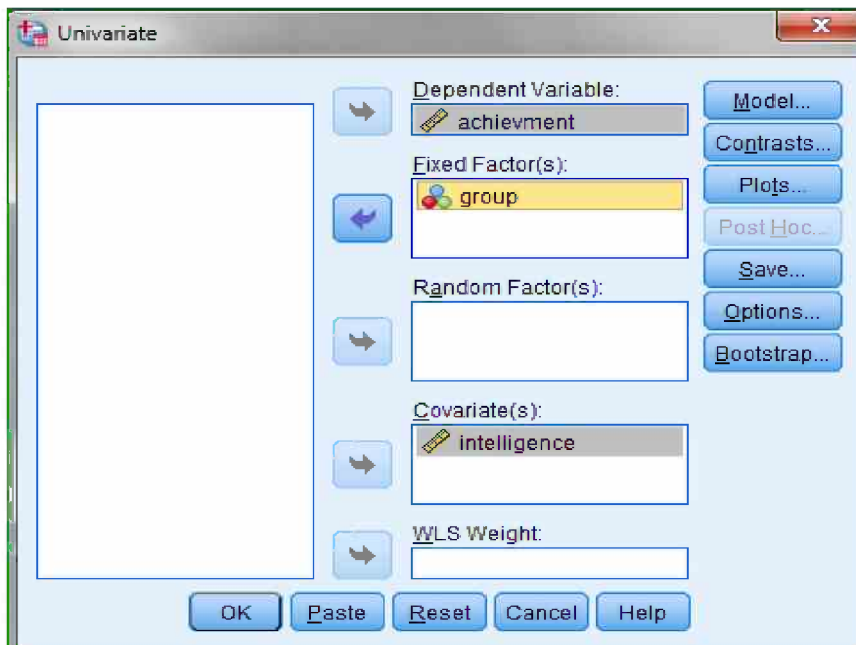
H_0 : أثر تفاعل الذكاء والمجموعة (المتغير المستقل) فى التنبؤ بالتحصيل صفر.

H_A: أثر تفاعل الذكاء والمجموعة في التنبؤ بالتحصيل لا يساوى صفر.

بمعنى وجود دلالة إحصائية للتفاعل بين متغير التباير والمتغير المستقل وهذا يعنى اختلاف خط ميل الانحدار بين المجتمعات الثلاثة نتيجة التباير وهذا يعنى عدم إجراء .ANCOVA

للتحقق من هذه المسلمة اتبع الخطوات الآتية:

1. اضغط UnivariateAnalyze → General linear model → تظهر الشاشة الآتية:

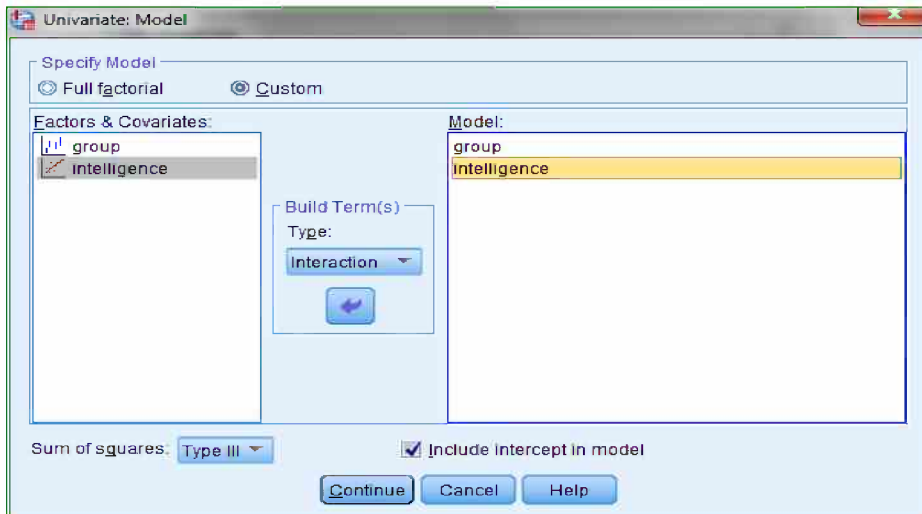


2. اضغط على achievement ثم اضغط على السهم → وانقله إلى مربع المتغير التابع .Dependent Variable.

3. اضغط على group ثم اضغط على السهم → المناظر وانقله إلى مربع Fixed Factors.

4. اضغط intelligence ثم اضغط على السهم → لتنتقله إلى مربع Covariate (S)

5. اضغط على Model ثم اضغط على Custom في النموذج:



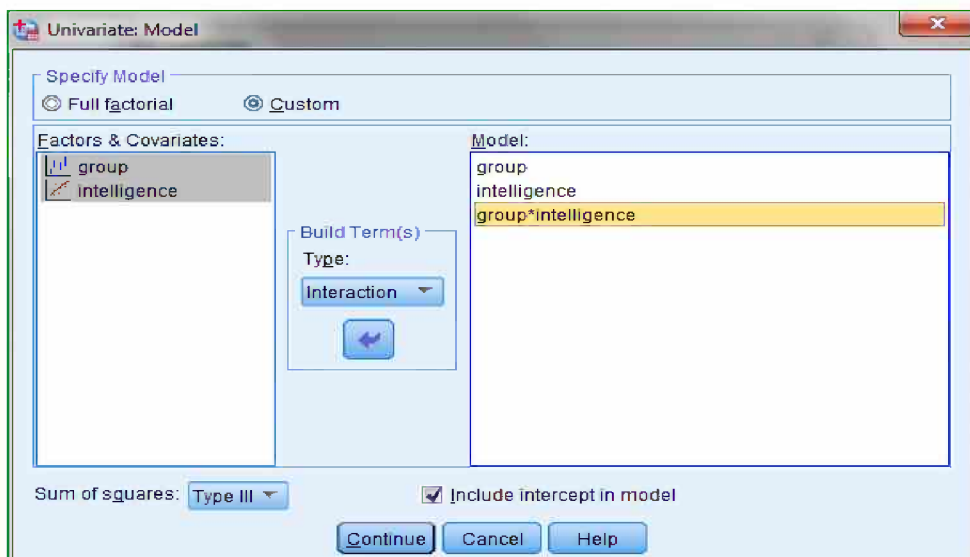
7. اضغط على متغير group في مربع Factors & Covariates ثم اضغط على → لتنتقله إلى مربع Model



8. اضغط على متغير intelligence في مربع Factors & Covariates ثم اضغط على → لتنتقله إلى مربع Model

9. في مربع (S) Build term وتحت Type اختار البديل :Interaction

10. اضغط على intelligence و group معاً أو من خلال مفتاح Ctrl (اجعلهم مظللين معاً) في مربع Factors & Covariates ثم اضغط على → لتنتقلهما إلى مربع Model كما في الشاشة الآتية:



11. اضغط Continue ثم اضغط Ok.

• المخرج كالاتى:

UNIANOVA achievement BY group WITH intelligence
/METHOD=SSTYPE(3)
/INTERCEPT=INCLUDE
/CRITERIA=ALPHA(0.05)
/DESIGN=group intelligence group*intelligence.

Univariate Analysis of Variance

Between-Subjects Factors

	N
group 1.00	12
2.00	12
3.00	12

Tests of Between-Subjects Effects

Dependent Variable: achievement

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	911.510 ^a	5	182.302	7.940	.000
Intercept	7.524	1	7.524	.328	.571
group	82.988	2	41.494	1.807	.182
intelligence	657.941	1	657.941	28.656	.000
group * intelligence	67.968	2	33.984	1.480	.244
Error	688.796	30	22.960		
Total	194175.000	36			
Corrected Total	1600.306	35			

a. R Squared = .570 (Adjusted R Squared = .498)

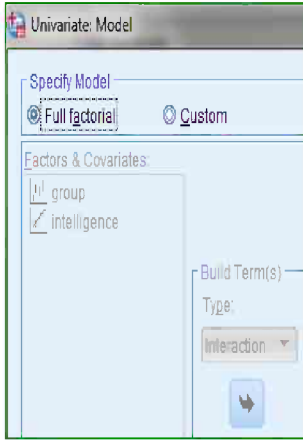
يتضح من المخرج السابق أن التفاعل بين المجموعة والذكاء غير دال إحصائياً حيث: $F(1,30)=1.480$, $P=0.244$ ، وعلى ذلك نتحقق لمسمة تساوى ميل تجانس خطوط الانحدار فى المجموعات الثلاثة بمعنى وجود علاقة مقاربة بين التحصيل و الذكاء فى المجموعات، وعلى الرغم أن حجم التأثير من النوع المتوسط وعلى ذلك يمكن إجراء ANCOVA ثم إجراء الاختبارات البعدية أو التتبعية لتقدير الفروق بين المتوسطات المصححة

ثالثاً: إجراء تحليل التباين أحادى الاتجاه One way ANCOVA

يمكن إجراء التحليل الرئيسى ANCOVA من خلال الخطوات الآتية:

1. اضغط Analyze ثم اضغط General liner Model ثم اضغط Univariate.

2. كرر الخطوات من 1 حتى 4 أثناء التحقق من مسلمة تجانس الانحدار.



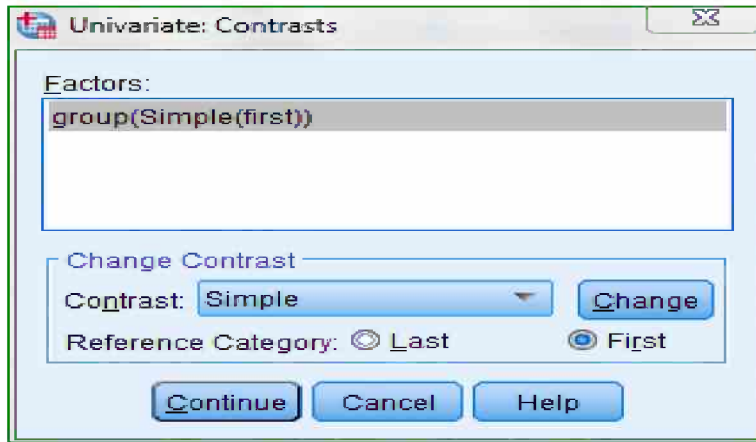
3. اضغط على اختيار Model فى شاشة الامر الرئيسية:

4. اضغط على Full Factorial.

5. اضغط على Continue (أسفل الشاشة) تعود إلى شاشة الامر الرئيسية.

ويوجد عدة مربعات حوارية منها Post Hoc وهذا المربع غير نشط ولكن كما نعلم أن المقارنات المتعددة يتم اجرائها من خلال المقارنات Contrasts حيث يتم مقارنة المجموعة الواحدة مع بقية المجموعات معاً.

6. اضغط على اختيار Contrasts تظهر الشاشة:



7. فى تحديد اى المجموعات المعيارية للمقارنة و لذلك يمكن وضع المجموعة الضابطة ترميز كمجموعة أولى ولذلك يتم مقارنة كل مجموعة تجريبية بالمجموعة الأولى (كود1) فى الشاشة الفرعية يوجد Contrast change اضغط على السهم ليعطى عدة بدائل كالاتى:

إذا اختارت None فإن البرنامج يختار المجموعة المرجعية للمقارنة ذات التصنيف أو الكود الأخير فإذا كانت ثلاث مجموعات يختار المجموعة ذات الكود3 و إذا ادخلت المجموعة الضابطة بالكود 1 تغير هذا البديل None بالبديل Simple ثم اضغط Change من القائمة فإنه يوجد اختيار لتحديد المجموعة المرجعية للمقارنة.

والطريقة الأخرى لإجراء المقارنات المتعددة البعدية:

8. اضغط على اختيار Options فى شاشة الامر الرئيسية يعطى الشاشة الفرعية الآتية:

Univariate: Options

Estimated Marginal Means

Factor(s) and Factor Interactions:

(OVERALL)
group

Display Means for:

group

☒ Compare main effects

Confidence interval adjustment:
Sidak

Display

☒ Descriptive statistics
☒ Estimates of effect size
☐ Observed power
☐ Parameter estimates
☐ Contrast coefficient matrix

☒ Homogeneity tests
☐ Spread vs. level plot
☐ Residual plot
☐ Lack of fit
☐ General estimable function

Significance level: .05 Confidence intervals are 95.0 %

Continue Cancel Help

9. اضغط على المتغير المستقل group لتنتقله إلى المربع Display Means

10. اضغط على الاختيار Compare main effect وعليه فإن الاختيار confidence interval adjustment ينشط ثم اضغط على السهم المقابل لـ LSD (None) يظهر ثلاثة اختيارات كالاتى:

- LSD لـ Tukey وهذا البديل لا يوصى باستخدامه.
- طريقة تصحيح Bonferroni يوصى باستخدامه.
- تصحيح Sidak Correction وهو مشابه لبديل بونيفيرونى ولكنه أقل تحفظاً ويفضل اختياره.

11. فى شاشة Options تحت Display اضغط على الإحصائيات الآتية: Descriptive Statistics و Estimates of effect size و Homogeneity tests

وفى نفس الشاشة يوجد عدة بدائل أخرى متاحة منها (Field, 2009):

- **Observed power** : تقدير مدى القدرة الاحتمالية لاختبار الإحصائي للكشف على فروق بين متوسطات المجموعات وهذا القياس نادر الاستخدام لأنه إذا كانت F داله إحصائية فإن هذا يعنى أن الاختبار لديه احتمالية أو قوة عالية.
- **Parameter estimates**: هذا البديل يعطى جدول معاملات الانحدار واختبارات الدلالة للمتغيرات فى نموذج الانحدار.
- **Homogeneity tests**: يعطى هذا الاختبار اختبار ليفين Leven's لاختبار مسلمة تجانس التباينات.

12. اضغط على Continue للعودة إلى الشاشة الرئيسية للأمر ثم اضغط OK

رابعًا: تفسير المخرج ANOVA:

- إجراء تحليل التباين أحادى الاتجاه مع حذف الذكاء يتم إجراءه من خلال امر تحليل التباين أحادى الاتجاه و كانت النتائج كالتالى :

ONEWAY achievement BY group /MISSING ANALYSIS.					
Oneway					
ANOVA					
achievement					
	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	204.056	2	102.028	2.411	.105
Within Groups	1396.250	33	42.311		
Total	1600.306	35			

يظهر من المخرج أن Sig = 0.105 وعليه فهى غير دالة إحصائيًا بمعنى عدم وجود فروق فى التحصيل بين المجموعات الثلاثة ويجب ملاحظة أن مجموع المربعات الكلى 1600.306 ومجموع

المربعات بين المجموعات (المعالجة) 204.056 ومجموع المربعات داخل المجموعات (الخطأ) وهو التباين غير المفسر 1396.250

تفسير مخرج ANCOVA:

الوصف الإحصائي: فيما يلى جدول الإحصاءات الوصفية وهو يعطى المتوسط والانحراف المعياري لدرجات التحصيل للمجموعات الثلاثة.

UNIANOVA achievement BY group WITH intelligence			
/CONTRAST (group)=Simple			
/METHOD=SSTYPE(3)			
/INTERCEPT=INCLUDE			
/EMMEANS=TABLES (group) WITH(intelligence=MEAN) COMPARE ADJ(SIDAK)			
/PRINT=ETASQ HOMOGENEITY DESCRIPTIVE			
/CRITERIA=ALPHA (.05)			
/DESIGN=intelligence group.			
Univariate Analysis of Variance			
Between-Subjects Factors			
group	1.00	12	
	2.00	12	
	3.00	12	
Descriptive Statistics			
Dependent Variable: achievement			
group	Mean	Std. Deviation	N
1.00	70.9333	8.11754	12
2.00	72.1667	7.56587	12
3.00	75.4167	5.68024	12
Total	73.1389	6.76188	36

Levene's Test of Equality of Error Variances^a				
Dependent Variable: achievement				
F	df1	df2	Sig.	
2.265	2	33	.120	
Tests the null hypothesis that the error variance of the dependent variable is equal across groups.				
a. Design: Intercept + intelligence + group				

ثم اعطى اختبار للتحقق من تجانس تباينات التحصيل للمجموعات الثلاثة: يتضح عدم وجود دلالة إحصائية حيث $P(0.12) > 0.05$ وبالتالي يتحقق مسلمة تجانس التباينات. ويرى Field (2013) أن اختبار Leven's للتحقق من تجانس التباينات ليس هو بالضرورة الطريقة الأمثل للحكم على ما إذا كانت التباينات متساوية والطريقة الأفضل هي فحص التباين الأدنى والتباين الأعلى.

واعطى البرنامج باقى المخرج كالاتى:

Tests of Between-Subjects Effects						
Dependent Variable: achievement						
Source	Type III Sum of Squares	df	Mean Square	F	Sig.	Partial Eta Squared
Corrected Model	843.542 ^a	3	281.181	11.890	.000	.527
Intercept	10.082	1	10.082	.426	.518	.013
intelligence	639.486	1	639.486	27.041	.000	.458
group	185.333	2	92.666	3.918	.030	.197
Error	756.764	32	23.649			
Total	194175.000	36				
Corrected Total	1600.306	35				
a. R Squared = .527 (Adjusted R Squared = .483)						

بالنظر إلى أول قيمة للدلالة الإحصائية المقابلة الذكاء Intelligence نلاحظ أن $\alpha > P(0.00)$ وعليه فإن متغير التباين (الذكاء) منبأ بالمتغير التابع

(التحصيل) على ذلك فإن ذكاء الأفراد يؤثر فى تحصيلهم: $P = 2.704$, $F(10.32)$, $\eta^2_p = 0.46$, $0.05 < (0.00)$ ، حيث إن الذكاء فسر 46% تقريباً من تباين التحصيل والشىء المثير للاهتمام أن تأثير المعالجة group دال إحصائية $P=0.03$ عند حذف أو استبعاد تأثير الذكاء ولكن مع وجود تأثير الذكاء كان التأثير غير دال إحصائياً للمجموعة كما سبق توضيحه فى تحليل التباين الأحادى حيث إن تأثير المجموعة كالاتى: $\eta^2_p = 0.33$, $P < 0.05$; $F(2.32) = 3.9185$ ، وهذا يعنى وجود علاقة قوية بين المعالجة والتحصيل، وحجم التأثير من النوع القوى حيث إنه طريقة التدريس فسرت 33% من تباين التحصيل كما ملاحظ أن مقدار التباين المفسر بواسطة النموذج (SSM) ارتفع إلى 84354 وكذلك انخفض قيمة التباين داخل المجموعات (الخطأ) من 1396 فى تحليل التباين الأحادى (بدون التغيرات) إلى 756.764 فى تحليل التغيرات مع تضمين التغيرات فى التحليل .و ذلك لان اختبار ANCOVA قدر الفروق بين المتوسطات المصححة وهى كما يلي Estimated marginal means:

Estimated Marginal Means				
group				
Estimates				
Dependent Variable: achievement				
group	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
1.00	72.099 ^a	1.425	69.197	75.001
2.00	71.029 ^a	1.421	68.135	73.923
3.00	76.288 ^a	1.404	73.429	79.148

a. Covariates appearing in the model are evaluated at the following values: intelligence = 112.6111.

وليست هى نفسها متوسطات المتغير التابع التى

group	Mean
1.00	70.8333
2.00	72.1667
3.00	76.4167
Total	73.1389

تم اعتمادها فى تحليل تباين الأحادى ANOVA

مع وجود تباين الذكاء حيث كانت كالاتى:

ولذلك تم اضافة الذكاء (التغيرات) فى

التحليل لضبط الفروق على هذا المتغير الذى لم يتضمن فى تحليل ANOVA وعلى

ذلك فإن نتائج متغير التغيرات الذى لا يتم ذكرها احياناً فى تقرير النتائج لأنها تقيم

العلاقة بين المتغير والمتغير التابع داخل المجموعات. وهذا المثال يوضح كيف لـ ANCOVA أن تساعد الباحث في إجراء ضبط تجريبي صارم اخذه في الاعتبار المتغيرات الداخلية وذلك لإعطاء صورة دقيقة عن تأثير نقى للمعالجة التجريبية بدون تشويش من متغيرات اخرى تسهم في زيادة التباين غير المفسر (البواقي). وبدون الاخذ في الاعتبار لمتغير الذكاء فيكون الاستنتاج والقرار خاطئ وهو عدم وجود تأثير لطريقة التدريس على التحصيل وهنا ينافى الحقيقة.

المقارنات Contrasts: اعطى المخرج نتيجة تحليل المقارنات:

"Contrast Results (K Matrix)			
		Dependent Variable	
group Simple Contrast ^a		achievement	
Level 2 vs. Level 1	Contrast Estimate	-1.070-	
	Hypothesized Value	0	
	Difference (Estimate - Hypothesized)	-1.070-	
	Std. Error	2.038	
	Sig.	.603	
	95% Confidence Interval for Difference		
	Lower Bound	-5.223-	
	Upper Bound	3.082	
Level 3 vs. Level 1	Contrast Estimate	4.189	
	Hypothesized Value	0	
	Difference (Estimate - Hypothesized)	4.189	
	Std. Error	2.003	
	Sig.	.045	
	95% Confidence Interval for Difference		
	Lower Bound	.109	
		Upper Bound	8.270
a. Reference category = 1			

وفى هذا الجدول تمت المقارنة بين المجموعة الثانية والمجموعة الأولى كأولى المقارنات و المقارنة بين المجموعة الثالثة والمجموعة الأولى كثنائى المقارنات وهذا حدث نتيجة تحديد المقارنات بين كل المجموعات مع المجموعة الأولى كمجموعة مرجعية وهذه المقارنة بين فروق المجموعات ثم عرضها: القيمة الفارقة Adifference value والخطأ المعياري Standard error والقيمة الدالة Significance value وفترات الثقة 95% .

وافرزت النتائج عدم جود فروق بين المجموعة الأولى والثانية $P=0.603$ ووجود فروق بين المجموعة الثالثة والأولى $P=0.045$ عند مستوى دلالة إحصائية 0.05 بالتالى الذى احدث الفروق هى بين متوسطات المجموعة الأولى والثالثة. وهذه المقارنات اخبرتنا عن وجود فروق بين المجموعات ولكن لتفسير هذه الفروق نحتاج لمعرفة المتوسطات ويتم ذلك بالنظر إلى قيم المتوسطات المصححة للمجموعات Adjusted means وليس متوسطات المجموعات الاصلية حيث إنها لم تصح من تباين التغاير ويتضح أن المتوسط المصحح للمجموعة الأولى 72.04 وللمجموعة الثالثة 76.288 وهذا الفرق أحدث الفروق بين المجموعات الثلاثة فى التحصيل وهذا الاستنتاج يمكن تأكيده والتحقق منه من خلال اختبارات المقارنات المتعددة التى سبق تحديدها بطريقة Sidak

المقارنات المتعددة البعدية: ويظهر ذلك فى المخرج الآتى:

Pairwise Comparisons						
Dependent Variable: achievment						
(I) group	(J) group	Mean			95% Confidence Interval for Differenceb	
		Difference (I-J)	Std. Error	Sig.b	Lower Bound	Upper Bound
1.00	2.00	1.070	2.038	.937	-4.065-	6.206
	3.00	-4.189-	2.003	.128	-9.236-	.858
2.00	1.00	-1.070-	2.038	.937	-6.206-	4.065
	3.00	-5.260-*	1.995	.038	-10.285-	-.234-
3.00	1.00	4.189	2.003	.128	-.858-	9.236
	2.00	5.260*	1.995	.038	.234	10.285
Based on estimated marginal means						
The mean difference is significant at the .05 level.						
b. Adjustment for multiple comparisons: Sidak.						

يتضح أن نتائج المقارنات البعدية اعطت أن الفروق بين المجموعة الأولى والمجموعة الثانية غير دال $P=0.937$ وكذلك الفروق بين المجموعة الأولى والثالثة غير دالة إحصائيًا $P=0.128$ وهذا يتناقض مع نتائج المقارنات Contasts بين المجموعة الأولى والثالثة حيث كانت دالة إحصائيًا ($P=0.045$) ولكن قيمة p كانت كبيرة تقترب من 0.05 ولكن كانت الدلالة نتيجة الفروق بين المجموعة الثانية و لمجموعة الثالثة حيث فروق المتوسطات -5.260 ، و $P=0.038$ ويمكن القول أن نتائج المقارنات البعدية المصححة Sidak أكثر دقة وتحفظًا من نتائج المقارنات Contrasts.

مقارنات Helmert: إذا لم يكن لدى الباحث توقع عن أى المعالجات والمجموعات أكثر فاعلية فإننا نستطيع اختبار بديل من مجموعات المقارنات مثل Helmert وفيها تتم المقارنة بين متوسط المجموعة الأولى (الضابطة) مع المجموعة الثانية التى تليها ومقارنة المجموعة الثانية والمجموعة التى تليها وهكذا وهذا النوع من المقارنات من المقارنات المخطط لها قبل الدراسة وفيما يلى تنفيذ هذا النوع من المقارنات:



1. اضغط على الاختيار Contrast
2. اضغط على السهم ↓ امام Contrast يعطيك مجموعة من البدائل اضغط على Helmert ثم اضغط على change تكون الشاشة الآتية:
لاحظ ظهر Helmert فى مربع العوامل

3. اضغط على Continue ثم اضغط على OK بالنظر فى المخرج:

Contrast Results (K Matrix)

		Dependent Variable
group Helmert Contrast		achievement
Level 1 vs. Later	Contrast Estimate	-1.559-
	Hypothesized Value	0
	Difference (Estimate - Hypothesized)	-1.559-
	Std. Error	1.758
	Sig.	.382
	95% Confidence Lower Bound	-5.140-
	Interval for Difference Upper Bound	2.021
Level 2 vs. Level 3	Contrast Estimate	-5.260-
	Hypothesized Value	0
	Difference (Estimate - Hypothesized)	-5.260-
	Std. Error	1.995
	Sig.	.013
	95% Confidence Lower Bound	-9.323-
	Interval for Difference Upper Bound	-1.196-

يتضح مقارنة المجموعة الأولى مع الثانية وهى غير دالة إحصائياً بينما الفروق بين متوسط المجموعة الثانية ومتوسط المجموعة الثالثة دال إحصائياً $p = 0.013$ وهذا يتفق مع المقارنات المتعددة بأسلوب Sidak

حجم التأثير لـ ANCOVA: يقدر حجم التأثير من خلال:

Partial Eta Squared
.527
.013
.458
.197

يعطيك حجم التأثير للمعالجة group وللتغاير الذكاء، فقدرت للذكاء:

$$\eta^2_p = \frac{639.486}{639.486 + 756.764} = 0.458$$

وللمجموعة:

$$\eta^2_p = \frac{185.333}{185.333 + 756.764} = 0.197$$

وهذا يعنى أن الذكاء أكثر تأثيراً على التحصيل من تأثيرا المعالجة حيث فسر الذكاء %45.8 من تباين التحصيل بينما فسرت طريقة التدريس حوال %19.7 .

كتابة تقرير نتائج ANCOVA وفقاً لـ APA

يؤثر الذكاء (التغاير) إحصائياً على التحصيل $F(1.32)=27.041$ $P(0.00) < 0.05$, $\eta^2_p = 0.458$ كما توجد فروق ذات دلالة إحصائية فى التحصيل بين طرق التدريس الثلاثة بعد ضبط تأثير الذكاء , $F(2.32)=3.198$, $P(0.030) < 0.05$, $\eta^2_p = 0.197$ كما ظهرت المقارنات المتعددة باستخدام طريقة sidak وجود فروق دالة إحصائياً بين متوسطى التحصيل للمجموعة الثالثة والثانية $P(0.038) < 0.05$.

الفصل السادس

تحليل التباين المتدرج

Multivariate analysis variance (MANOVA)

تحليل التباين المتدرج هو تعميم أو أتساع لتحليل التباين البسيط ANOVA عندما يوجد متغير مستقل واحد فأكثر تصنيفى وأكثر من متغير تابع . فإذا اراد الباحث دراسة تأثير أنواع مختلفة من المعالجات على أبعاد الابتكارية (أصالة- مرونة- طلاقة) فإنه يستخدم تحليل التباين المتدرج و كلمة متدرج تعنى تعدد فى المتغيرات التابعة، بينما إذا أراد دراسة تأثير أنواع مختلفة من معالجات على الأصالة فقط فانه يستخدم تحليل التباين البسيط Simple ANOVA ويعرف التحليل بـ Univariate حيث تعنى متغير تابع واحد فقط بينما إذا اراد دراسة أثر المعالجة والجنس (متغير مستقلين تصنيفين) على الأصالة فإنه يستخدم تحليل التباين المتعدد ذو اتجاهين Factorial ANOVA (تحليل التباين المتعدد).

ويبحث MANOVA فى ما إذا كان الاتحاد Combination بين أبعاد الابتكارية معاً تتغير كوظيفة للمعالجة وهذه الوظيفة الخطية للمتغيرات التابعة تعرف بالوظيفة التمييزية Discriminant Function، وفى الواقع فإن MANOVA يطلق عليه التحليل التمييزى Discriminant analysis. وعلى ذلك فإن ANOVA يختبر فروق المتوسطات بين المجموعات على متغير تابع واحد بينما MANOVA يختبر فروق المتوسطات بين المجموعات على اتحاد مجموعة من المتغيرات التابعة معاً وعليه فإن MANOVA يسعى إلى خلق متغير تابع جديد يمكن أن يعظم الفروق بين المجموعات نتيجة تخليق هذا المتغير من مجموعة المتغيرات التابعة وهذا المتغير التابع الجديد يكون فى اتحاد خطى مع المتغيرات التابعة المتصلة الاصلية المقاسة المكونة له و يتم إجراء ANOVA على المتغير التابع التخليقى الجديد.

وإذا تعددت المتغيرات المستقلة التصنيفية فى هذه الحالة يطلق عليه Factorial MANOVA وفيه يتم دراسة التأثيرات الرئيسية وكذلك تأثير التفاعلات كما هو الحال

فى Factorial ANOVA ، وعلى ذلك فإن MANOVA هو منهجية لدراسة الفروق بين المجموعات ويتملك MANOVA عدة مميزات عن ANOVA هى:

- بدلاً من دراسة التأثيرات على كل متغير تابع على حدة يتم دراسة التأثيرات على كل المتغيرات التابعة كلاً على حدة ومتفاعلين.
- يمتاز MANOVA عن إجراء سلسلة ANOVA لكل متغير تابع على حدة بأنه حماية أو وقاية ضد تضخم الخطأ من النوع الأول نتيجة إجراء الاختبارات المتعددة لـ ANOVA للمتغيرات التابعة المرتبطة كلا على حدة.
- كما أنه لو أجرى ANOVA لكل متغير تابع على حدة فإن أى علاقة بين المتغيرات التابعة يتم تجاهلها وعليه نفقد المعلومات عن هذه العلاقات الموجودة ولكن MANOVA يتضمن كل المتغيرات التابعة معاً فى نفس التحليل فإذا كان المفهوم يتكون من أبعاد فإن ANOVA يخبرنا عن الفروق فى كل بعد على حدة بينما MANOVA يكون لديه القوة لتحديد الفروق بين المجموعات على اتحاد الأبعاد معاً ونتيجة لذلك فإن حسابات MANOVA فى غاية التعقيد الرياضى وكذلك تتطلب عدة مسلمات هامة يجب اخذها فى الاعتبار وكذلك يوجد غموض فى التميز فى تأثير المتغير المستقل الواحد على المتغيرات التابعة معاً فماذا تعنى فقد تكون الفروق ناتجة من متغير تابع ما بعينة و ذلك ينصح بان يكون استخدام MANOVA له مبررات قوية فى اختيار المتغيرات وكذلك له تبريرات فى التفسير. وذلك لان اختيار متغيرات تابعة متوسطة العلاقة فيما بينها يقلل من قوة MANOVA.

ويمكن لـ MANOVA أن يتضمن متغيرات مستقلة منفصلة ومتغيرات مستقلة متصلة و على ذلك يكون مشابهة لتحليل التباين و يصبح هنا تحليل التباين المتدرج (Multivariate Analysis of Covariance (MANCOVA).

استخدامات لـ MANOVA

1. يستخدم فى التصميمات التجريبية وشبه التجريبية لتقدير الفروق بين المجموعات على مجموعة من المتغيرات التابعة المتصلة المفترض أن تكون مرتبطة، وأحياناً فى بداية التصميم التجريبى إذا كان البرنامج يهدف إلى تنمية القراءة فإن الباحث يختبر

ما إذا كانت توجد فروق بين المجموعة التجريبية والمجموعة الضابطة في متغيرات مثل الذكاء ودخل الأسرة ودرجات القراءة القبلية، وفي نهاية التجربة يستخدم MANOVA لدراسة الفروق بين المجموعة التجريبية و المجموعة الضابطة على أكثر من متغير تابع (الفهم- قراءة المفردات- قراءة الجمل) وإذا وجدت فروق بين المجموعات على أفضل اتحاد خطى linear Combination للمتغيرات التابعة فإن الباحث يغزو هذه الفروق إلى المعالجة.

2. يستخدم MANOVA لتحليل بيانات تصميمات القياسات المتكررة حيث يتم جمع البيانات لمتغير تابع واحد فأكثر وفي هذا المثال فإن الوقت هو متغير مستقل بمستويات تعكس فترات زمنية منفصلة ومجموعة المتغيرات التابعة عبارة عن متغير تابع وحيد يقاس عبر فترات زمنية منفصلة ويسمى هذا تصميم داخل المجموعات.

3. يستخدم MANOVA في التصميمات غير التجريبية لتقدير الفروق بين مجموعتين فأكثر على متغيرين تابعين فأكثر كأن نختبر الفروق بين الذكور والإناث على مجموعة من المتغيرات التابعة مثل التحصيل والابتكارية أو حتى أبعاد الابتكارية (الاصالة- المرونة- الطلاقة) ولو وجدت فروق دالة إحصائية بين الجنسين فإنه من الصعب أن تغزو ذلك إلى السببية لأنه لا توجد معالجة للمتغير المستقل ولم يتم ضبط مناسب للمتغيرات الداخلية.

مسلمات تحليل التباين المتدرج

- يتسم توزيع المتغيرات التابعة في المجتمع بالاعتدالية المتدرجة Multivariate Normality للمتغيرات التابعة مجتمعة وإذا كانت المتغيرات التابعة معا تتسم بالاعتدالية المتدرجة فإن كل متغير تابع على حدة يكون توزيعه اعتدالياً وإذا لم تتحقق هذه المسلمة فإن قيمه P تكون غير صادقة و تنخفض قوة الاختبار الإحصائي، ويشير (2013) Field إلى أن مسلمة الاعتدالية المتدرجة لا يمكن اختبارها في برنامج SPSS والحل العملي هو التحقق من Univariate Normality لكل متغير تابع على حدة وتحقيق الاعتدالية لكل متغير على حدة ليست ضمانة الاعتدالية المتدرجة.

- **تباينات وتغايرات المتغيرات التابعة فى المجتمع هى نفسها (متساوية) عبر كل مستويات العامل (المتغير المستقل التصنيفى) وإذا لم تتحقق هذه المسلمة فإن نتائج MANOVA تكون غير صادقة ويسمح برنامج SPSS باختبار هذه المسلمة وهى تجانس مصفوفات التباين - التغاير باستخدام إحصاء Box's M حيث الدلالة الإحصائية تشير إلى عدم تحقق المسلمة بينما عدم الدلالة تشير إلى وجود تجانس، ويمكن التحقق منها من اختبارات Univariate test لتساوى التباينات بين المجموعات لكل متغير باستخدام اختبار ليفين's leven's test، لكن قبل إجراء اختبار Box's لا بد من التحقق من توافر الاعتدالية المتدرجة قبل تفسير النتائج وكذلك يتأثر اختبار Box's بأحجام العينات الكبيرة حيث تؤدي إلى دلالة إحصائية ويفترض أيضًا تساوى أحجام العينات فى كل المجموعات.**
- **العينة مختارة عشوائيًا والدرجة على المتغير لأى فرد هى مستقلة تمامًا عن درجات الأفراد الآخرين على هذا المتغير (مسلمة الاستقلالية).**

الأسئلة البحثية فى MANOVA

- الهدف من استخدام MANOVA هو تحديد ما إذا كان سلوك المتغيرات التابعة تتغير نتيجة المعالجة أو أى حدث آخر ويهدف MANOVA إلى دراسة:
- **التأثيرات الرئيسية للمتغيرات المستقلة (قد يكون متغير مستقل واحد):** تحديد ما إذا كانت للمعالجة أو المتغير المستقل أثر على أفضل اتحاد خطى للمتغيرات التابعة: هل يوجد فروق فى درجات أبعاد الابتكارية مجتمعة بين المجموعة التجريبية والمجموعة الضابطة؟
 - وإذا تضمن البحث متغير تغاير يصبح الأسلوب MANCOVA ويكون السؤال هل توجد فروق فى درجات أبعاد الابتكارية معًا بين المجموعة الضابطة والمجموعة التجريبية بعد ضبط أثر درجات الابتكارية فى القياس القبلى؟
 - **التفاعلات بين المتغيرات المستقلة:** إذا تضمن تصميم البحث أكثر من متغير مستقل ومتغيرات تابعة فبعد دراسة التأثيرات الرئيسية لكل متغير مستقل على حدة

يحتاج الباحث إلى دراسة التفاعلات بين مستويات المتغيرات المستقلة على المتغيرات التابعة فكل تفاعل بين المتغيرات المستقلة يتم التحقق منه باستخدام اختبار خاص به.

• **اهمية المتغيرات التابعة Importance of DVS:** إذا وجدت فروق دالة إحصائية لاحد التأثيرات الرئيسية أو أكثر أو للتفاعلات فإن الباحث يتساءل عن أى المتغيرات التابعة التى تغيرت أو تأثرت وإيهما لم يتأثر بالمتغيرات المستقلة، فربما يكون أحد المتغيرات التابعة هى التى تغيرت نتيجة للمعالجة بينما المتغيرات التابعة لأخرى لم تتأثر.

• **تقديرات المعالم Parameter estimates:** تعتبر المتوسطات المصححة افضل التقديرات لمعالم المجتمع لمتوسطات التأثيرات الرئيسية والخلايا، والمتوسطات ليست متوسطات العينة بل المتوسطات المصححة Adjusted Means. وفى حالة MANOVA يتم تصحح المتوسطات من التغايرات وبالإضافة للمتوسطات يتم تقدير الانحرافات المعيارية والأخطاء المعيارية وفترات الثقة.

• **المقارنات المتعددة:** إذا كانت التأثيرات الرئيسية أو التفاعلات دالة إحصائيا فإن الخطوة التالية هى أى مستويات التأثيرات الرئيسية أو خلايا التفاعلات هى التى احدثت هذه الدلالة فاذا كانت للمعالجة ثلاث مستويات واعطت دلالة إحصائية فالباحث يتساءل أى هذه المستويات احدثت الدلالة فى الاتحاد الخطى للمتغيرات التابعة وعليه يمكن استخدام اختبارات المقارنات المتعددة سواء البعدية أو القبلية.

• **حجم التأثير:** إذا كانت للمعالجة تأثير حقيقى على السلوك (المتغيرات التابعة) فإن السؤال اللاحق هو: ما مقدار هذا التأثير أو ما حجمه أو ما هى نسبة التباين المفسر فى الاتحاد الخطى للمتغيرات التابعة الذى أحدثته المعالجة؟، وعليه فإن الباحث يلجأ إلى تقدير حجم التباين المفسر للتأثيرات الدالة إحصائياً.

• **تأثيرات التغايرات Effects of Covariates:** عندما تستخدم متغيرات التغايرات فى تحليل MANOVA فإنه الباحث يريد تقدير مدى جدواه ومدى تأثيره على المتغيرات التابعة المتحدة معاً، هل التغايرات لها دلالة إحصائية على المتغيرات التابعة بمعنى هل توجد علاقة بين التغايرات والمتغيرات التابعة.

قضايا عملية فى إجراء MANOVA (Tabachnick & Fidell, 2007)

توجد ثلاثة قضايا عملية قبل إجراء تحليل MANOVA:

1. مدى تحقق مسلمات الاختبار .
 2. اختيار اختبار الدلالة الإحصائية والجدل الموجود عن أفضل هذه الاختبارات في ضوء اعتبارات القوة الإحصائية وحجم العينة.
 3. ماذا بعد أي ما هو التحليل الذي يجب إجراءه بعد تحليل MANOVA خاصة عندما تكون MANOVA دالة إحصائية.
- فبالنسبة للمسلمات فقد تم عرضها سابقاً.
- اختبارات الدلالة الإحصائية**

يوجد أربعة اختبارات إحصائية للتحقق من الدلالة الإحصائية في MANOVA ويوجد القليل من الدراسات التي بحثت في القوة الإحصائية للاختبارات الأربعة MANOVA ، ولاحظ (1974) Olson وجود اختلافات ضئيلة في ضوء القوة الإحصائية للاختبارات الأربعة لأحجام العينات الصغيرة و المتوسطة. فإذا كانت فروق المجموعات تركز على أول وظيفة تمييزية (وهذا شائع الاستخدام في العلوم الاجتماعية) فإن إحصاء Roy's Root أكثر قوة، يليه Hotelling's λ ثم Wilk's λ وأخيراً Pillai's-trace، وعندما يتم التركيز على فروق المجموعات على أكثر من تباين فإن التفضيل يكون العكس بمعنى أن Pillai's أكثر قوة إحصائية وأقلهم قوة هو Roy's root . واوصى (1980) Stevens باستخدام أقل من عشر متغيرات تابعة ما لم تكن أحجام العينات كبيرة. وفي ضوء الضلالة الإحصائية Robustness فإن الاختبارات الأربعة تتميز بالضلالة النسبية ضد عدم تحقق الاعتدالية المتدرجة ولكن أقلهم ضلالة Roy's وأيضاً يعتبر ليس له ضلالة عندما لا تتحقق مسلمة تجانس مصفوفة التغاير .

وعندما تكون أحجام العينات متساوى اختبار فان Pillai's-trace هو أكثر ضلالة ضد عدم تحقق مسلمات MANOVA وعندما تكون أحجام العينات غير متساوية فإن هذا الاختبار أكثر تأثراً بعدم تحقق مسلمة تجانس مصفوفات التغاير ويعتبر اختبار Wilks's-lambda مقياس للتباين غير المفسر لاتحاد المتغيرات التابعة المتصلة نتيجة

المتغيرات المستقلة التصنيفية ويعتبر اختبار Wilks من أكثر الاختبارات استخدامًا ويفضل أن تكون قيمته منخفضة، وقيمة F الدالة إحصائيًا لهذا الاختبار تشير إلى وجود فروق بين على الأقل بين مجموعتين على الاتحاد الخطى للمتغيرات التابعة وعكس إحصاء Wilks نجد أن Pillai's-trace عبارة عن مجموع النسب للتباينات المفسرة على الوظائف التمييزية وهو مشابه لمؤشر $\eta^2(\frac{SSB}{SST})$ أما اختبار Hotelling's T^2 هو مجموع $(\frac{SSB}{SST})$ لكل متغير تابع على حدة.

التحليل التتبعي Follow up analysis

يوجد بعض الجدل حول ماهى انسب وسيلة لإجراء التحليل التتبعي بعد إجراء التحليل الرئيسى لـ MANOVA فالمدخل التقليدى هو إجراء تحليل التباين البسيط لكل متغير تابع على حده ولكن هذا من شأنه أن يضخم الخطأ من النوع الأول الفا ولذلك فمن الضروري استخدام تصحيح Bonferroni لتحليلات ANOVA اللاحقة. وبعض الباحثون يميلون إلى استخدام التحليل التمييزى لإيجاد الاتحاد الخطى للمتغيرات التابعة التى تفصل أو تميز المجموعات لكل متغير تابع على حده حيث يتم إجراء التحليل التمييزى مع عكس ادوار المتغيرات التابعة و المستقلة ANCOVA حيث يستخدم كل المتغيرات التابعة المتصلة MANOVA كمتغيرات مستقلة متصلة ويصبح المتغير التصنيفى هو متغير تابع والهدف هو تقدير كيف كل متغير متصل يتميز بين المجموعات والقيمة المرتفعة للأوزان المعيارية للمتغير تشير إلى فروق كبيرة على المتغير التصنيفى وفيه يظهر المتغيرات أكثر تمايزًا على المجموعات.

ويمكن إجراء تحليل التباين أحادى الاتجاه لكل متغير تابع واعتبار المتغيرات التابعة الباقية كتغايرات فى كل تحليل وإذا كشفت هذه التحليلات عن دلالة إحصائية لـ F فإنه توجد فروق دالة إحصائية بين المجموعات على المتغير التابع المتضمن فى التحليل بعد استبعاد أثر التباين المتداخل مع بقية المتغيرات التابعة الباقية المستخدمة فى MANOVA . ويتم إجراء اختبارات تتبعية على المتغيرات التابعة سواء باستخدام ANOVA أو ANCOVA . ويمكن إجراء المقارنات المتعددة البعدية

مثل اختبار توكي (HSD) للفروق بين المجموعات لكل متغير تابع على حده وكذلك يوجد طريقة Bonferroni.

مؤشرات حجم التأثير في اختبارات الإحصاء المتدرج

التعامل مع الظواهر النفسية يفترض أن تتعدد فيها المتغيرات التابعة والمستقلة ويوجد القليل من الدراسات التي تناولت مؤشرات حجم التأثير أو تحليل القوة في اختبارات الإحصاء المتدرج، فعلى سبيل المثال الفرض الصفري في تحليل التباين الأحادي يفترض تساوى المتوسطات لكل المجموعات على المتغير التابع، بينما في تحليل التباين المتدرج MANOVA يفترض تساوى متوسطات المجموعات لكل المتغيرات التابعة. ويوجد في التراث مقاييس لتقدير حجم التأثير لاختبارات الإحصاء المتعدد ففي الانحدار المتعدد يستخدم مؤشر معامل الارتباط المتعدد R^2 وتصحيحه R^2_{adj} .

ويوجد مؤشراً بديلاً لمؤشر d يستخدم في الإحصاء المتدرج وخاصة في MANOVA ويرمز له بالرمز D^2 وذلك في حالة تعدد المتغيرات التابعة وأكد (2009) Ferguson على ضرورة الاهتمام بمؤشر D^2 ويسمى distance Mahalanobis ويحسب من اختبارات الدلالة الإحصائية لاختبار MANOVA كالاتى:

$$D^2 = \frac{\bar{y}_1 - \bar{y}_2}{S(\bar{y}_1 - \bar{y}_2)}$$

$$|\bar{y}_1 - \bar{y}_2|$$

قيمة متجهة المتوسطات

$$S(\bar{y}_1 - \bar{y}_2) \text{ مصفوفة التباين.}$$

وان مؤشر D^2 هو مستقل عن حجم العينة ولتفسير D^2 وضع (2009) Stevens القيمة 0.25 ضعيف و 0.50 متوسط و 1.0 فأكثر عالى.

كما يمكن اعتبار إحصاء الدلالة الإحصائية في تحليل التباين المتدرج (MANOVA) Wilks- Lambda (Λ) مؤشر لحجم التأثير حيث يعبر عن التباين غير المفسر، وتحدد قيمته كالاتى:

$$\Lambda = \frac{|S_w|}{|S_t|}$$

$$\eta^2 = 1 - \Lambda$$

حيث S_w مجموع المربعات داخل المجموعات و S_t مجموع المربعات الكلى.

وكذلك يوجد مؤشر $\eta^2_{adjMult}$ فى اختبار MANOVA وتتحدد قيمته بالآتى:

$$\eta^2_{multi\ adj} = \eta^2_{multi} - \frac{(p^2 + q^2)(1 - \eta^2_{multi})}{3N}$$

حيث P عدد المتغيرات التابعة، q عدد درجات الحرية للفروض.

واشار (1993) Tatsuoka إلى استخدام مؤشر ω^2_{Mult} للاختبارات المتدرجة وذلك لان مؤشر η^2_{Mult} متحيز تحيز موجب ومؤشر ω^2_{Mul} قريباً من مؤشر ω^2 لـ Hayes ويمكن التعبير عنه بالصيغة الآتية:

$$\omega^2_{mult} = \frac{(N - K) - (N - 1)\Lambda}{(N - K) + \Lambda}$$

وقد عبر عنها (2000) Olejnik & Algina بالصيغة الآتية:

$$\omega^2_{mult} = 1 - \frac{N\Lambda}{df_{error} + \Lambda} \quad , \quad df_{error} = N - K$$

و df درجات الحرية وتساوى $N-K$

بالإضافة إلى مؤشرات D^2 ، η^2 ، ω^2 يوجد مؤشرات أخرى لتقدير حجم التأثير ويوصى باستخدامها وهى مؤشرات هوتلنج Hotelling و Lawley trace & Bartlett و Pillais Trace وهذه المؤشرات الثلاثة يمكن الحصول عليها من خلال SPSS.

اختبارات الفروض قضية بحثية (2009, 2013) Field:

اراد الباحث إجراء برنامج علاجى للوسواس القهرى واستخدم ثلاث مجموعات المجموعة الأولى تلقت علاج السلوكى المعرفى والمجموعة الثانية تلقت علاج سلوكى والمجموعة الثالثة لم تتلقى أى نوع من العلاج (مجموعة ضابطة) وقاس الوسواس القهرى فى

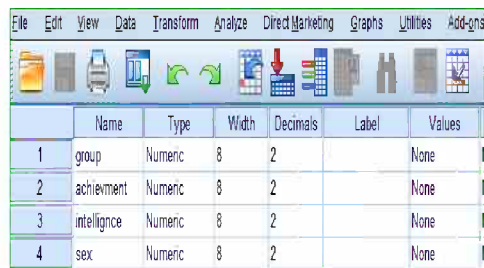
القياس البعدى فى صورة متغيرين تابعين أحدهما يتعلق بالأفعال والآخر بالأفكار وفيما يلى درجاتهم:

المتغير التابع الأول (الأفعال)						المجموعة
التابع الثانى (أفكار)						
سلوكى معرفى	سلوكى (B)	ضابطة	سلوكى معرفى	سلوكى	ضا	
5	4	4	14	14	13	
5	4	5	11	15	15	
4	1	5	16	13	14	
4	1	4	13	14	14	
5	4	6	12	15	13	
3	6	4	14	19	20	
7	5	7	12	13	13	
6	5	4	15	18	16	
6	2	6	16	14	14	
4	5	5	11	17	18	
4.90	3.70	5.0	13.40	19.20	15.	
1.20	1.77	1.05	1.90	2.10	2.3	
1.43	3.12	1.11	3.60	4.40	5.5	
10	10	10	10	10	10	
$\bar{X}_{\text{grand}} = 4.53$		$\bar{X}_{\text{grand}} = 14.53$				
$S^2_{\text{grand}} = 2.1195$		$S^2_{\text{grand}} = 4.8780$				

وتسأل الباحث عن:

1. متوسطات الفروق: هل يختلف الدرجات على متغيرات الوسواس القهرى للأفعال والوسواس القهرى الأفكار والاتحاد الخطى بينهما بين المجموعات الثلاثة ؟

إجراء MANOVA فى SPSS



	Name	Type	Width	Decimals	Label	Values
1	group	Numeric	8	2		None
2	achievement	Numeric	8	2		None
3	intelligence	Numeric	8	2		None
4	sex	Numeric	8	2		None

أولاً: إدخال البيانات: 1. اضغط على Variable

view أسفل الشاشة الافتتاحية للبرنامج تظهر الشاشة:

2. اكتب مسمى المتغيرات تحت Name وهى: متغير

المجموعة group: علاج سلوكى معرفى = 1، علاج

سلوكى = 2، مجموعة ضابطة = 3.، الوسواس القهرى (أفعال) Action، الوسواس القهرى

(أفكار) Thought

	group	actions	thoughts	var	var
1	1.00	5.00	14.00		
2	1.00	5.00	11.00		
3	1.00	4.00	16.00		
4	1.00	4.00	13.00		
5	1.00	5.00	12.00		
6	1.00	3.00	14.00		
7	1.00	7.00	12.00		
8	1.00	6.00	15.00		
9	1.00	6.00	16.00		
10	1.00	4.00	11.00		
11	2.00	4.00	14.00		
12	2.00	4.00	15.00		

3 . اضغط Data view تظهر شاشة البيانات ويمثلوا

ثلاثة أعمدة يتم إدخال البيانات لثلاثة متغيرات.

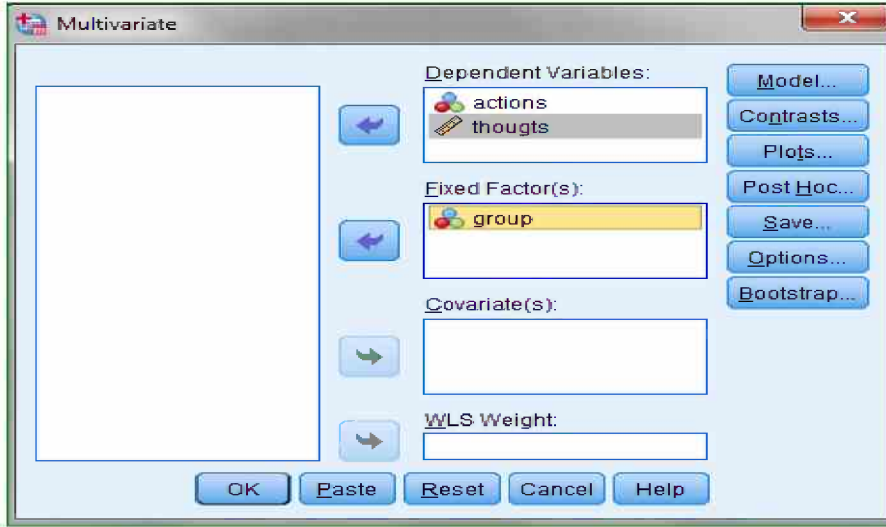
وفيما يلي ملف البيانات:

ثانياً: لإجراء MANOVA والاختبارات التتبعية

والمقارنات البعدية اتبع الآتي:

1. اضغط Analyze ثم اختار General liner model ثم أضغط على Multivariate

يظهر لك المربع أو الشاشة الآتية:



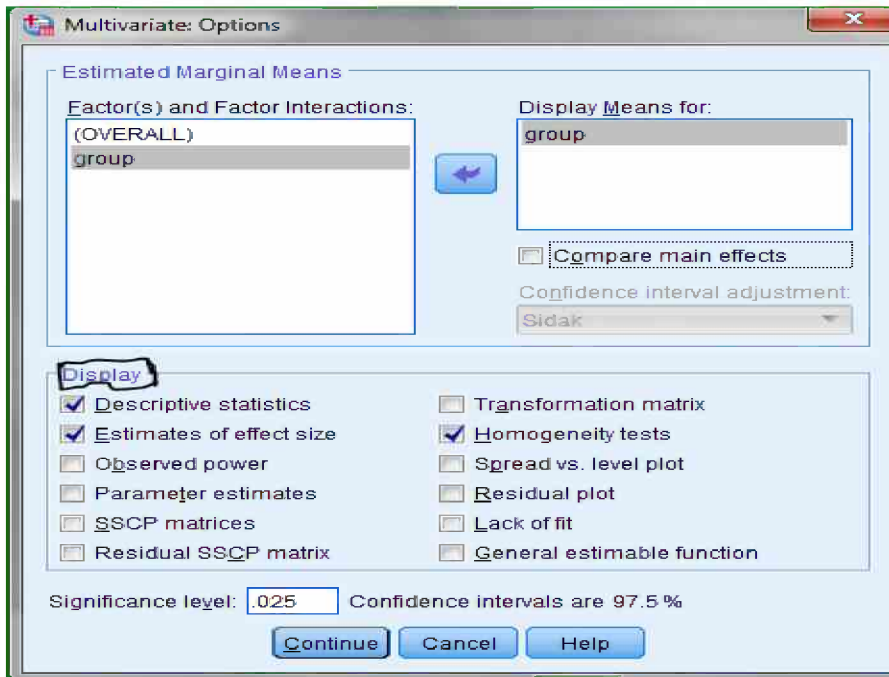
2. اضغط على المتغير Action مع الضغط على مفتاح ctrl اضغط على المتغير التابع

الآخر thought ثم اضغط على السهم → لتنتقلهم إلى مربع Dependent variables

3. اضغط على متغير Group ثم على → لتنتقله إلى مربع Fixed factor(s) لاحظ قد يوجد

أكثر من متغير مستقل تصنيفي يتم نقلهم إلى نفس المربع.

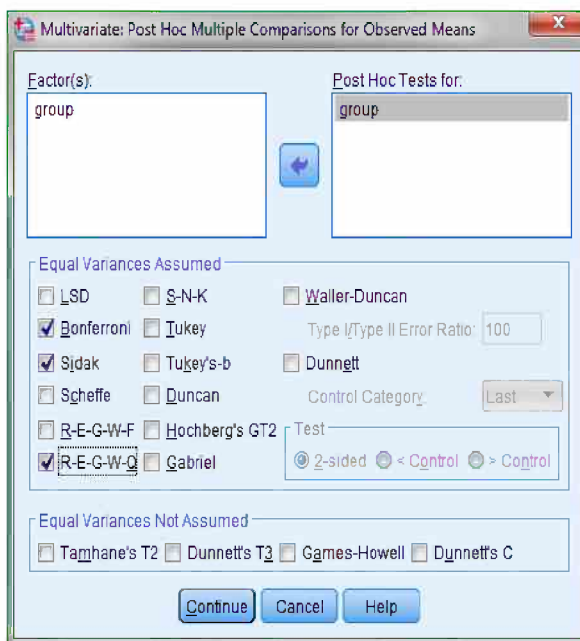
4. اضغط على اختيار Options (يمين الشاشة) تظهر الشاشة الفرعية الآتية :



5. اضغط على group في مربع Factor(s) and factors interactions ثم اضغط على → لتنتقله إلى مربع Display Means

6. في مربع Display اضغط على الاختيارات الآتية: Descriptive statistics , Homogeneity test , Estimates of effects size

7. غير مستوى الدلالة الإحصائية Significance level من 0.05 الي 0.025 و هي خارج قسمة 0.05 على عدد المتغيرات التابعة وذلك لتجنب تضخم الخطأ من النوع الأول عن إجراء المقارنات المتعددة.



8. اضغط Continue ترجع إلى الشاشة الرئيسية للأمر.

9. اضغط على اختيار Post Hoc تظهر الشاشة الآتية:

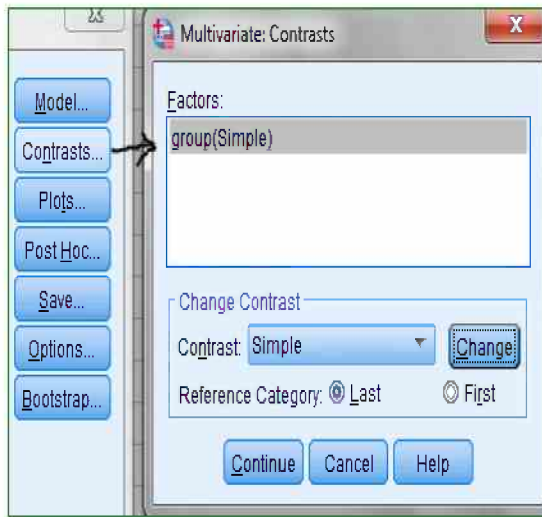
10. اضغط على group فى المربع factors ثم على السهم → لتنتقله إلى مربع Post hoc tests

11. فى مربع تساوى التباينات اضغط على اختيار Bonferroni وكذلك يمكن استخدام طريقة LSD إذا كان اقصى عدد مستويات المتغير المستقل ثلاثة وإذا زاد عن ذلك فإنه غير مناسب.

12. إذا لم يتحقق شرط تساوى التباينات اضغط تحت هذا المربع Equal variances not assumed على طريقة Dunnett's والاختبارات الأخرى يمكن استخدامها.

13. اضغط Continue ترجع إلى الشاشة

الرئيسية للأمر ثم اضغط على الاختيار Contrasts يظهر المربع الآتى:



14. اضغط على ↓ امام المربع Contrast

يعطيك بدائل عديدة اختار Simple ثم

اضغط على Change على اعتبار أن يتم

مقارنة أى من المجموعات التجريبية

بالمجموعة الضابطة لاحظ أن المجموعة

الضابطة كودت 3 (أعلى كود) و عليه يجب تغير المجموعة المرجعية Reference category

إلى last و بالتالى نشط last وإذا اخترت البديل NONE فإنه يعتبر المجموعة ذات الكود

الأخير هى مجموعة المقارنة.

15. اضغط Continue ثم اضغط OK

ثالثاً: تفسير مخرج MANOVA:

• التحليل الأولي واختبار المسلمات: اعطى البرنامج وصف إحصائى للعينة حيث كل

مجموعة تتضمن 10 أفراد:

```

GLM actions thoughts BY group
/CONTRAST(group)=Simple
/METHOD=SSTYPE(3)
/INTERCEPT=INCLUDE
/POSTHOC=group(BONFERRONI SIDAK QREGW)
/EMMEANS=TABLES(group)
/PRINT=DESCRIPTIVE ETASQ HOMOGENEITY
/CRITERIA=ALPHA(.025)
/DESIGN= group.

```

General Linear Model

Between-Subjects Factors

		N
group	1.00	10
	2.00	10
	3.00	10

Descriptive Statistics

	group	Mean	Std. Deviation	N
actions	1.00	4.90000	1.19722	10
	2.00	3.70000	1.76698	10
	3.00	5.00000	1.05409	10
	Total	4.53333	1.45586	30
thoughts	1.00	13.40000	1.89737	10
	2.00	15.20000	2.09762	10
	3.00	15.00000	2.35702	10
	Total	14.53333	2.20866	30

وعرض البرنامج الإحصائيات الوصفية لكل متغير تابع فى المجموعات الثلاثة وهى

Box's Test of Equality of Covariance Matrices^a

Box's M	9.959
F	1.482
df1	6
df2	18168.923
Sig.	.180

Tests the null hypothesis that the observed covariance matrices of the dependent variables are equal across groups.

a. Design: Intercept + group

المتوسط والانحراف المعياري وحجم العينة.

ثم اعطى البرنامج اختبار Box's

test:

ويختبر ما إذا كانت التباينات والتغايرات

للمتغيرات التابعة هى نفسها عبر كل

مستويات العامل وإذا كانت قيمة F دالة

إحصائية بالتالى فإن فرضية التجانس غير متوفرة لأنه يعنى وجود فروق فى

المصفوفات ولكن نتائج Box الدالة إحصائياً تفسر بحذر لأنها قد ترجع إلى عدم تحقق

مسلمة الاعتدالية المترتبة وكذلك النتائج غير الدالة إحصائياً قد ترجع إلى صغر حجم

العينة ونقص القوة للاختبار وقيمة $F=1.482$ $P(0.180) > 0.05$ وعليه لا توجد دلالة

إحصائية بمعنى تساوى مصفوفات التباين ولتغاير عبر المستويات الثلاثة وعليه تحقق

شرط استخدام MANOVA.

اعطى البرنامج اختبارات الدلالة فى MANOVA كالاتى:

Multivariate Tests ^a						
Effect		Value	F	Hypothesis df	Error df	Partial Eta Squared
Intercept	Pillai's Trace	.983	745.230 ^b	2.000	26.000	.000
	Wilks' Lambda	.017	745.230 ^b	2.000	26.000	.000
	Hotelling's Trace	57.325	745.230 ^b	2.000	26.000	.000
	Roy's Largest Root	57.325	745.230 ^b	2.000	26.000	.000
group	Pillai's Trace	.318	2.557	4.000	54.000	.049
	Wilks' Lambda	.699	2.555 ^b	4.000	52.000	.050
	Hotelling's Trace	.407	2.546	4.000	50.000	.051
	Roy's Largest Root	.335	4.520 ^c	2.000	27.000	.020

a. Design: Intercept + group

b. Exact statistic

c. The statistic is an upper bound on F that yields a lower bound on the significance level.

ركز على الاختبارات فى خانة group و اتضح وجود دلالة إحصائية فى اختبارات Pillai's trace (P=0.049) واختبار Wilks' lambda (p=0.05) واختبار Roy's root (P=0.02) وكلهم توصلوا إلى وجود فروق بين المجموعات الثلاثة على الاتحاد الخطى (متغير جديد) المتغيرين التابعين الافعال والأفكار معًا ولكن بالنظر الدلالة الإحصائية لـ Hotelling's trace (P=0.05) اتضح عدم وجود دلالة إحصائية ولكن كما سبق وشرنا أن اختبارات Pillai's و Roy's هى أكثر قوة إحصائية من باقى الاختبارات وعليه فإن الثقة فى نتائج الدلالة الإحصائية تكون قوية بمعنى وجود دلالة إحصائية بين المجموعات الثلاثة وعلى ذلك فإن نوعية العلاج لها أثر على الوسواس القهرى بعبدية معًا ولكن طبيعية هنا التأثير غير واضحة فى ضوء نتائج اختبارات الدلالة MANOVA لأنه لا يخبنا عن أى مجموعة التى احدثت الدلالة وكذلك لا تخبرنا عن تأثير العلاج على كلا من الأفكار والافعال على حدة ولتحديد ذلك ليمدنا البرنامج باختبار F Univariate لدراسة أثر المجموعة على كل متغير تابع عل حدة ولكن قبل إجراء ذلك الاختبارات لا بد من التحقق من مسلمة تجانس التباينات لكل متغير تابع على حده عبر المجموعات الثلاثة للتحقيق من هذه المسلمة اعطى البرنامج اختبار leven's كالتى :

Levene's Test of Equality of Error Variances ^a				
	F	df1	df2	Sig.
actions	1.828	2	27	.180
thoughts	.076	2	27	.927

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept + group

ويتضح أن قيمة الدلالة الإحصائية F للأفعال $F=1.828$, $P=0.18$ وكذلك للأفكار $F=0.076$, $P=0.927$ وهما غير دالين إحصائيًا بالتالي

يوجد تساوى التباينات ولمتغيرات التابعة وفيما يلي نتائج اختبارات Univariate F:

Tests of Between-Subjects Effects						
Source	Dependent Variable	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	Actions	10.467 ^a	2	5.233	2.771	.080
	Thoughts	19.467 ^b	2	9.733	2.154	.136
Intercept	Actions	616.533	1	616.533	326.400	.000
	Thoughts	6336.533	1	6336.533	1402.348	.000
group	Actions	10.467	2	5.233	2.771	.080
	Thoughts	19.467	2	9.733	2.154	.136
Error	Actions	51.000	27	1.889		
	Thoughts	122.000	27	4.519		
Total	Actions	678.000	30			
	Thoughts	6478.000	30			
Corrected Total	Actions	61.467	29			
	Thoughts	141.467	29			

a. R Squared = .170 (Adjusted R Squared = .109)

b. R Squared = .138 (Adjusted R Squared = .074)

والصف موضوع الاهتمام هو المعنون group ولاحظ أن النتائج امام group هي نفسها النتائج امام Corrected model وذلك لان مطابقة النموذج للبيانات تضمن متغير مستقل واحد فقط وامام group يتضمن جدول ANOVA لكل متغير تابع و يتضمن قيم مجموع المربعات للأفعال والأفكار وكذلك قيمة F ودلالاتها، بينما الصف Errors يتضمن معلومات عن مجموع المربعات داخل المجموعات (البواقي) (SSW) والصف المعنون Corrected Model يمثل مجموع المربعات الكلى (SST) الشيء المهم فى هذا الجزء هو قيم F ودلالاتها (P) والملاحظ أن قيم F المقابلة للأفعال والأفكار غير دالة إحصائيًا حيث إن قيمة P للأفعال 0.08 وقيمة P للأفكار 0.136 وهذا يقودنا إلى استنتاج وهو نوع العلاج ليس له تأثير دال إحصائيًا على أبعاد الوسواس القهرى الأفعال والأفكار على حدة، وعليك أن تتعامل مع هذا التناقض بين نتائج اختبار MANOVA ونتائج اختبار ANOVA وهذا ببساطة يرجع إلى أن MANOVA اخذ فى حسابه العلاقة بين المتغيرات التابعة وتعامل معها كمتغير واحد كاتحاد خطى بينهما وهذا جعل MANOVA أكثر قوة للكشف عن الفروق بين المجموعات وعلى ذلك فإن تفسير ANOVA عديم الجدوى وهذا يعطى دلائل على اهمية التعامل مع المتغيرات النفسية بصورة أكثر شمولية لأنها ظاهرة تتسم بتفاعل متغيراتها وهذا ما كشف عنه MANOVA بينما لم يستطع ANOVA إجراء ذلك.

المقارنات: اعطى البرنامج المقارنات كالاتى:

Contrast Results (K Matrix)

group Simple Contrast ^a		Dependent Variable	
		actions	thoughts
Level 1 vs. Level 3	Contrast Estimate	-.100-	-1.600-
	Hypothesized Value	0	0
	Difference (Estimate - Hypothesized)	-.100-	-1.600-
	Std. Error	.615	.951
	Sig.	.872	.104
	97.5% Confidence Interval for Difference	Lower Bound -1.559-	Upper Bound -3.856-
		1.359	.656
Level 2 vs. Level 3	Contrast Estimate	-1.300-	.200
	Hypothesized Value	0	0
	Difference (Estimate - Hypothesized)	-1.300-	.200
	Std. Error	.615	.951
	Sig.	.044	.835
	97.5% Confidence Interval for Difference	Lower Bound -2.759-	Upper Bound -2.056-
		.159	2.456

a. Reference category = 3

يرى البعض أن إجراء المقارنات ليس له جدوى لأن ANOVA اعطت عدم دلالة إحصائية هذا الإجراء من الضروري إجراؤه إذا كانت اختبارات Univariate دالة إحصائية ولكن تم إجراء مقارنات بين المجموعة الضابطة (3) مع كلاً من المجموعتين

التجريبيتين على حدة واعطى المخرج الجدول السابق وهو مكون من جزئين الجزء الأول Level I Vs. level 3 و Level 2 Vs. level 3 حيث يشير الرقم إلى كود المجموعات وتم إجراء هذه المقارنات على كل المتغيرات التابعة واعطى قيم هي Contrast Estimate والقيمة المفترضة Hypothesized value (دائمًا تساوى صفر) وهى القيمة التى تخبر عنها الفرض الصفرى. ويتم اختبار الدلالة الإحصائية من خلال مقارنة من خلال الفروق بين قيمة المقارنة والقيمة المفترضة فى ضوء الخطأ المعيارى وفى ضوء حدود الثقة 95% حيث إذا وقع الفرق بين القيمتين فى مدى حدود الثقة نقبل الفرض الصفرى.

ويلاحظ من الجدول السابق فى الأول level 1 Vs. level 3 عدم وجود فروق بينهما على الأفعال $P = 0.0872$ وكذلك على الأفكار $P = 0.104$ لان القيمتين أكبر من 0.05 واتضح وجود فروق دالة إحصائية بين المجموعتين الثالثة و الثانية على بعد الأفعال $p=0.044$ ، وعدم وجود فروق ذات دلالة إحصائية بين المجموعتين الثانية والثالثة فى بعد الأفكار $p = 0.835$

المقارنات المتعددة أو المزدوجة

اعطى البرنامج المقارنات المتعددة باستخدام طريقة Bonferrioni و Sidak:

Multiple Comparisons								
Dependent Variable		(I) group	(J) group	Mean Difference (I-J)	Std. Error	Sig.	97.5% Confidence Interval	
							Lower Bound	Upper Bound
actions	Bonferroni	1.00	2.00	1.2000	.61464	.184	-.5498-	2.9498
			3.00	-.1000-	.61464	1.000	-1.8498-	1.6498
		2.00	1.00	-1.2000-	.61464	.184	-2.9498-	.5498
			3.00	-1.3000-	.61464	.131	-3.0498-	.4498
		3.00	1.00	.1000	.61464	1.000	-1.6498-	1.8498
			2.00	1.3000	.61464	.131	-.4498-	3.0498
	Sidak	1.00	2.00	1.2000	.61464	.173	-.5477-	2.9477
			3.00	-.1000-	.61464	.998	-1.8477-	1.6477
		2.00	1.00	-1.2000-	.61464	.173	-2.9477-	.5477
			3.00	-1.3000-	.61464	.126	-3.0477-	.4477
		3.00	1.00	.1000	.61464	.998	-1.6477-	1.8477
			2.00	1.3000	.61464	.126	-.4477-	3.0477
thoughts	Bonferroni	1.00	2.00	-1.8000-	.95063	.207	-4.5064-	.9064
			3.00	-1.6000-	.95063	.312	-4.3064-	1.1064
		2.00	1.00	1.8000	.95063	.207	-.9064-	4.5064
			3.00	.2000	.95063	1.000	-2.5064-	2.9064
		3.00	1.00	1.6000	.95063	.312	-1.1064-	4.3064
			2.00	-.2000-	.95063	1.000	-2.9064-	2.5064
	Sidak	1.00	2.00	-1.8000-	.95063	.193	-4.5030-	.9030
			3.00	-1.6000-	.95063	.280	-4.3030-	1.1030
		2.00	1.00	1.8000	.95063	.193	-.9030-	4.5030
			3.00	.2000	.95063	.996	-2.5030-	2.9030
		3.00	1.00	1.6000	.95063	.280	-1.1030-	4.3030
			2.00	-.2000-	.95063	.996	-2.9030-	2.5030

Based on observed means.
The error term is Mean Square(Error) = 4.519.

ويتضح من الجدول بالنسبة للأفعال أن المقارنة بين المجموعة الأولى و كلاً من المجموعة الثانية والثالثة على حدة غير دالة إحصائياً وكذلك الفروق بين المجموعة الثانية والثالثة غير دالة إحصائياً، أما بالنسبة لبعد الأفكار فأتضح عدم وجود فروق بين أى من المجموعات الثلاثة ولاحظ أن المقارنات تمت على أساس مستوى دلالة إحصائية 0.025 حتى يمكن التحكم فى الخطأ من النوع الأول عكس المقارنات Contrast حيث تمت المقارنة على أساس 0.05 واتضح عدم وجود فروق فى المقارنات المزدوجة نظراً لأن قيمة F غير دالة إحصائياً لكل متغير تابع على حدة .

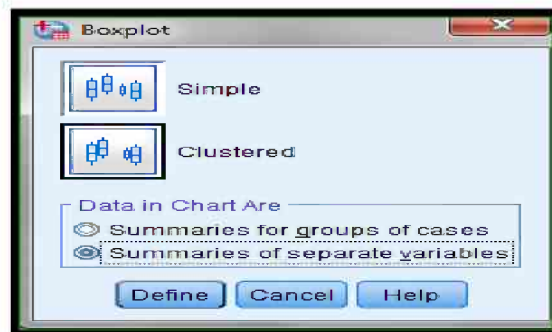
وكذلك أظهرت المقارنات المتعددة باستخدام REGWR عدم وجود فروق
 $p=0.106$ للأفعال وللفكار $P=0.160$:

Homogeneous Subsets			
actions			
	group	N	Subset 1
Ryan-Einot-Gabriel- Welsch Range ^a	2.00	10	3.7000
	1.00	10	4.9000
	3.00	10	5.0000
	Sig.		.106
Means for groups in homogeneous subsets are displayed. Based on observed means. The error term is Mean Square(Error) = 1.889. a. Alpha = .025.			
thoughts			
	group	N	Subset 1
Ryan-Einot-Gabriel- Welsch Range ^a	1.00	10	13.4000
	3.00	10	15.0000
	2.00	10	15.2000
	Sig.		.160
Means for groups in homogeneous subsets are displayed. Based on observed means. The error term is Mean Square(Error) = 4.519. a. Alpha = .025.			

العرض البياني للنتائج

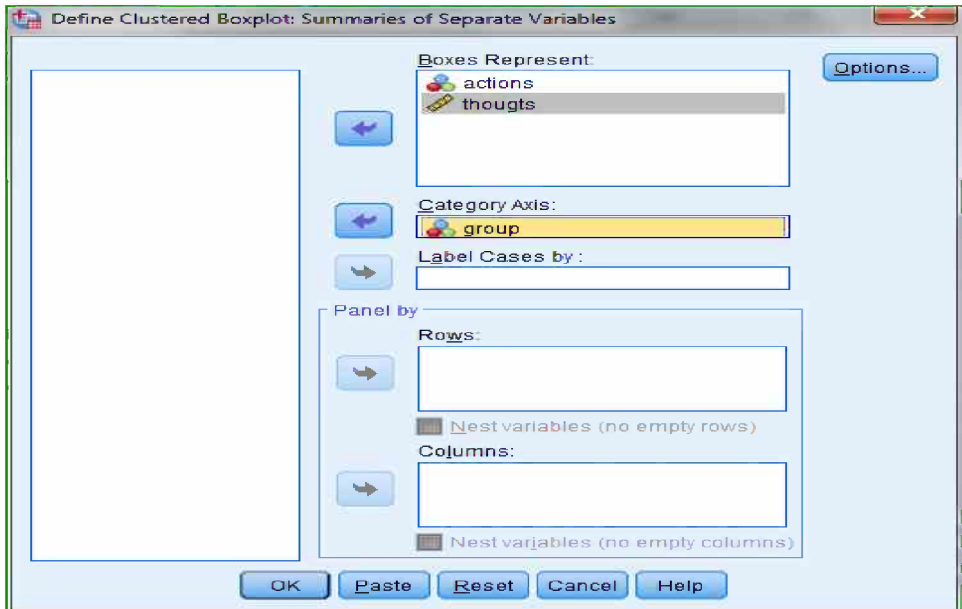
تستخدم الرسومات مثل Boxplot في عرض نتائج MANOVA ، تسمح للقارئ بتقويم الفروق بين المجموعات، وBoxplot يعرض توزيعات المتغيرات التابعة المتعددة في MANOVA. ولعمل الرسومات Boxplot لنتائج المثال السابق أتبع الخطوات الآتية:

1. اضغط على Graphs ثم اضغط على اختيار Legacy dialogs ثم اضغط على اختيار Boxplot تظهر الشاشة الآتية :



2. اضغط على Clustered ثم نشط Summaries for separate variables

3. اضغط على Define تظهر الشاشة الآتية:

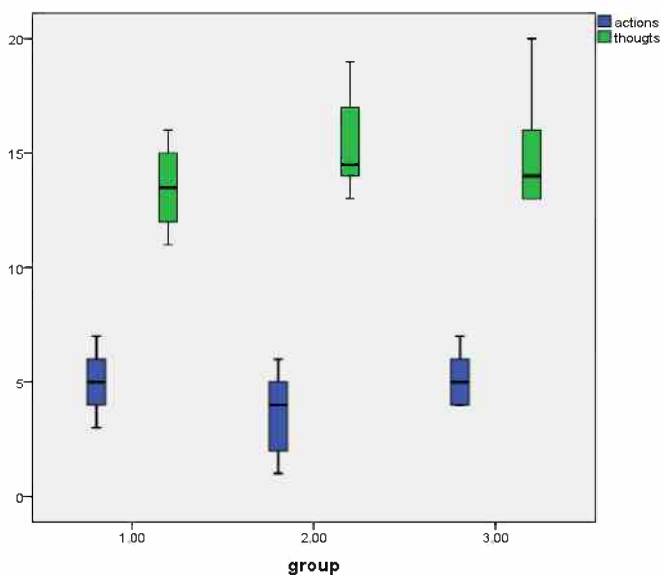


4. اضغط على group ثم على السهم → و انقله إلى مربع CategoryAxisBox

5. اضغط على مفتاح ctrl ثم اضغط على Action , thoughts معا ثم اضغط على السهم

→ و انقلهما إلى مربع Boxesrepresent

6. اضغط على OK.



المخرج:

هذا الشكل يعكس توزيعات درجات الأفكار والأفعال للوسواس القهري للمجموعات الثلاث، أعطى المجموعات على المحور الأفقى والمتغيرات التابعة على المحور الرأسى فمتغير الأفعال أعطى اللون الأزرق وأعطى Boxplot للمجموعات الثلاثة ونلاحظ أن أقل قيمة لمؤشر نزعة مركزية (الوسيط) كان للمجموعة الثانية بينما تساوى وسيط المجموعة الأولى والثالثة مما يدل أن البرنامج كان أكثر تأثيراً على المجموعة الثانية (العلاج السلوكى) حيث أقل قيمة للوسواس القهري بينما العلاج السلوكى المعرفى لم يحدث تحسناً وهكذا بالنسبة للأفكار كانت أقل قيمة للوسيط (الخط فى منتصف المستطيل) للمجموعة الأولى مما يدل على أنها هى أكثر استفادة من البرنامج السلوكى المعرفى.

كتابة نتائج MANOVA فى تقارير البحث وفقاً APA

كتابة نتائج MANOVA مشابهة لكتابة نتائج ANOVA لأن اختبارات MANOVA هى تحول إلى تقريب F ويتم عرض النتائج كالاتى:

بإجراء تحليل التباين المتدرج لتحديد أثر نوعى العلاج وهما السلوكى المعرفى والسلوكى وكذلك بالإضافة إلى المجموعة الضابطة على بعدى الوسواس القهري الأفعال والأفكار، وأتضح أن وجود تأثير ذو دلالة إحصائية لنوعية العلاج على المتغيرين التابعين معاً حيث:

$$\text{Wilks's} = 0.699, F(4,52) = 2.56, P \leq 0.05$$

ويمكن عرض أى اختبار دلالة إحصائية آخر من الاختبارات الأربعة ولكنه يفضل عرض اختبار واحد مثل Wilks's أو Pillai's. وإذا تم عرض نتائج Pillai's يكون كالاتى: وجود تأثير ذو دلالة إحصائية لنوعين العلاج على الأفكار والأفعال معاً:

$$\text{Pillai's trace}, = 0.3185, F = 2.56, P < 0.05$$

وبإجراء تحليل التباين البسيط ANOVA اتضح عدم وجود فروق بين المجموعات الثلاثة فى الوسواس القهري الأفعال حيث:

$$F(2,27) = 2.771, P > 0.05, \eta^2_p = 0.170$$

وكذلك عدم وجود فروق بين المجموعات الثلاثة فى الوسواس القهري الأفكار حيث:

$$F(2,27) = 2.154, P > 0.05, \eta^2_p = 0.138$$

الفصل السابع

الانحدار المتعدد

Multiple Regression

وتحليل الانحدار المتعدد هو مجموعة من الأساليب الإحصائية واتساع للانحدار الخطى البسيط التى تسمح بتقدير العلاقة بين متغير تابع (DV) ومتغيرات مستقلة عديدة. ويستخدم لمجموعة من البيانات حيث يوجد ارتباط بين المتغيرات المستقلة ببعضها وترتبط مع المتغير التابع بدرجات مختلفة، فمثلاً تقدير العلاقة بين المتغيرات المستقلة مثل دخل للأسرة والمستوى الاجتماعى الاقتصادى للأسرة والعمر مع التحصيل. وعلى ذلك فإن الانحدار يكون مفيد عندما ترتبط المتغيرات المستقلة ببعضها بدرجات مختلفة وهو مفيد أيضاً فى البحوث التجريبية وكذلك فى البحوث المسحية.

ويستخدم تحليل الانحدار للمتغيرات المستقلة سواء كانت متصلة أو تصنيفية، وعلى ذلك فإن ANOVA هى حالة خاصة من الانحدار حيث دراسة التأثيرات الرئيسية والتفاعلات هى سلسلة من متغيرات مستقلة تصنيفية ومشاكل ANOVA يمكن التغلب عليها فى تحليل الانحدار.

فى الانحدار البسيط معادلة الانحدار غير المعيارية:

$$\hat{Y} = bx + a$$

- X الدرجة الخام للمتغير المستقل.
- a ثابت (الخطأ).
- b معامل الانحدار غير المعيارى.
- $\hat{Y}$ القيمة المتنبأ بها للمتغير التابع.

بينما فى الانحدار المتعدد حيث متغير تابع وحيد وأكثر من متغير مستقل ومعادلة الانحدار كالتالى:

$$\hat{Y} = a + b_1X_1 + b_2X_2 + \dots + b_kX_k$$

- $b_k, \dots, b_1, b_2$ معاملات الانحدار غير المعيارية المرتبطة بالمتغيرات المستقلة $X_k, \dots, X_1, X_2$.
- e الخطأ أو البواقي.

ولذلك فإن معادلة الانحدار تتضمن معاملات الانحدار للمتغيرات المنبئة أو المستقلة وثابت الانحدار a .

والصيغة السابقة لمعادلة الانحدار هي من الدرجات الخام ولكن بفرض تحويل درجات المتغيرات المستقلة والمتغير التابع إلى درجات معيارية فإن صيغة معادلة الانحدار المعيارية تكون كالآتي:

$$Z\hat{Y} = \beta Z_1 + \beta_2 Z_2 + \dots + \beta_k Z_k$$

ويسمى معامل الانحدار بيتا β أو معاملات الانحدارية المعيارية.

وفي الانحدار البسيط تم تقدير مربع معامل الارتباط R^2 بين المتغير المستقل والتابع و نسبة التباين غير المفسر هي $(1 - R^2)$ ولكن بفرض وجود متغيرين مستقلين X_1, X_2 فإن معامل الارتباط يطلق عليه مربع معامل الارتباط المتعدد $R^2_{y.X_1, X_2}$ بين y وكلًا من X_1, X_2 ويمكن اعطائه الرمز:

$$R^2_{y.X_1, X_2} = R^2_{y.12}$$

ومربع معامل الارتباط بين X_1 و y_1 :

$$r^2_{y.X_1} = r^2_{y.1}$$

$$r^2_{y.X_2} = r^2_{y.2}$$

حيث $R^2_{y.12}$ تشير إلى التباين المفسر للمتغير التابع Y نتيجة للمتغيرين X_1, X_2

ومربع معامل الارتباط المتعدد هو حاصل جمع:

$$R^2_{y.12} = r^2_{y.1} + r^2_{y.2}$$

الأسئلة البحثية في الانحدار المتعدد (Tanachnick & Fiddel, 2007) :

يعتبر تحليل مفيد للإجابة على الأسئلة الآتية:

- **تحديد درجة العلاقة Degree of Relationship**: هل معادلة الانحدار تصلح للتنبؤ ؟ هل معامل الارتباط المتعدد (R) يختلف عن الصفر؟ ولذلك فهو يصلح لمعرفة ما إذا كان الارتباط المتعدد يختلف عن الصفر ام لا.
- **أهمية المتغيرات المستقلة Importance of IVS**: لو معادلة الانحدار تختلف عن الصفر (دالة إحصائية) فالسؤال التالي: أى من المتغيرات المستقلة على درجة كبيرة من الأهمية وأيهما ليس له أهمية فى المعادلة؟ وهذا يبين الأهمية النسبية Relative importance للمتغيرات المستقلة المختلفة فى حلول الانحدار.

● **إضافة متغيرات مستقلة:** افترض أن الباحث حسب معادلة الانحدار وأراد معرفة ما إذا كان يريد تحسين التنبؤ بالمتغير التابع عن طريق إضافة متغير مستقل أو أكثر إلى المعادلة فمثلاً أراد تحسين معادلة التنبؤ بالقدرة على القراءة للأطفال بإضافة متغيرات أخرى مثل اهتمام الآباء بالقراءة للمتغيرات المستقلة الموجودة في المعادلة (العمر والذكاء) ويرى مدى التحسن في قيمة R.

● **المقارنة بين مجموعات من المتغيرات المستقلة:** هل التنبؤ بالمتغير التابع من مجموعة من المتغيرات المستقلة أفضل من التنبؤ من مجموعة أخرى من المتغيرات المستقلة، فمثلاً هل التنبؤ بالتحصيل من المناخ الأسرى والدافعية أفضل من التنبؤ بالتحصيل من المناخ المدرسى والاتجاه نحو المدرسة؟

● **تقدير المعالم Parameter estimates:** تقدير معالم الانحدار المتعدد مثل معاملات الانحدار غير المعيارية (b) تعكس التغير في المتغير التابع المرتبط بتغير وحدة واحدة في المتغير المستقل مع جعل المتغيرات المستقلة الأخرى ثابتة فمثلاً التنبؤ بالتحصيل من مفهوم الذات:

$$\text{التحصيل} = 10 + 50 \text{ مفهوم الذات}$$

فان قيمة $b = 10$ تخبرنا بان زيادة وحدة واحدة في مفهوم الذات تؤدي إلى زيادة 10 نقطة أو وحدة في درجات التحصيل. وهذا التفسير يعتمد على مدى تحقق مسلمات الانحدار وهو قياس المتغيرات المستقلة بدون أخطاء.

محددات تحليل الانحدار

● **القضايا النظرية:** تحليل الانحدار يكشف عن العلاقات بين المتغيرات ولا يمكن أن نستنتج أن العلاقات تكون سببية فالعلاقة القوية بين المتغيرات يمكن أن تتبع من مصادر عديدة مثل تأثير متغيرات أخرى غير مقاسة وغير متضمنة في التحليل.

والقضية الأخرى أي من المتغيرات التابعة يجب قياسها وأي من المتغيرات المستقلة يجب تضمينها في النموذج والإجابة على هذا ينبع من النظرية القوية التي ينطلق منها الباحث وأحياناً من الفحص الجيد لتوزيع البواقي أو أخطاء القياس للمتغيرات (انظر مسلمات الانحدار البسيط).

ويوجد اعتبارات عديدة في اختيار المتغيرات المستقلة فتحليل الانحدار يكون مثالي إذا ارتبطت المتغير المستقل ارتباطاً قوياً مع المتغير التابع و لا يرتبط مع المتغيرات المستقلة الأخرى أو

ارتباط ضعيف جزئى، فالهدف العام من تحليل الانحدار هو تحديد أفضل عدد من المتغيرات المستقلة الضرورية للتنبؤ بالمتغير التابع. والواضح أن حلول الانحدار لها حساسية شديدة لتضمين المتغيرات المتضمنة فيه، فأهمية المتغير المستقل فى الانحدار يعتمد على المتغيرات المستقلة الأخرى، وعليه فإن أفضل متغيرات مستقلة هى الأقل ارتباطاً أو غير مرتبطة مع بعضها والأكثر ثباتاً. ويفترض تحليل الانحدار أن المتغيرات المستقلة تقاس بدون أخطاء وهذا مستحيل فى العلوم الإنسانية والاجتماعية والسلوكية وعليه فالأفضل أن تكون المتغيرات المستقلة عالية الثبات.

وعدم تضمين المتغيرات المستقلة الأكثر تأثيراً فى المتغير التابع فإنها تسهم فى زيادة الخطأ فى معادلة الانحدار خاصة إذا ارتبط هذا المتغير مع المتغيرات المستقلة المتضمنة فى نموذج الانحدار، وعليه فإن مكونات الخطأ للمتغير المستقل غير المقاس والمفترض الأكثر تأثيراً على المتغير التابع والأكثر ارتباطاً مع المتغير المستقل فى النموذج فإن هذا يؤدى إلى تحطم مسلمة استقلالية الأخطاء Independence of errors وعلى ذلك فالعلاقة بين المتغيرات المقاسة المستقلة فى التحليل والمتغير المستقل المستبعد من التحليل يغير تقدير قيمة معاملات الانحدار، فلو كانت العلاقة موجبة فإن معامل الانحدار يعتبر Overestimated أى يعطى قيمة أكبر من قيمته الحقيقية وإذا كانت العلاقة سالبة فإن تقدير معامل الانحدار Underestimated يعطى قيمة أقل من قيمته المفترضة.

القضايا العملية

يتطلب الانحدار المتعدد العديد من القضايا العملية مثل:

• نسبة الحالات إلى المتغيرات المستقلة: تحديد حجم العينة فى تحليل الانحدار يعتمد على

عدة عوامل مثل مقدار القوة المرغوبة ومستوى ألفا وعدد المنبئات وحجم التأثير المتوقع.

وأعطى (Green 1991) عدة قواعد أهمها لتحديد حجم العينة:

$$N \geq 50 + 8m$$

وذلك لاختبار معامل الارتباط المتعدد ، حيث m عدد المتغيرات المستقلة

$$\text{فإذا كان عدد المنبئات} = 6 \text{ فإن: } N \geq 50 + 8 * 6$$

$$\geq 50 + 48$$

$$\geq 98$$

وحجم العينة المطلوب لاختبار المعالم مثل معامل الانحدار والثابت:

$$N \geq 104 + m$$

وإذا كان عدد المنبئات = 6 فإن حجم العينة

$$N \geq 104 + 6 \\ \geq 110$$

والنسبة العالية للحالات بالنسبة للمتغيرات المستقلة تكون عندما يكون المتغير التابع ملتوى التوزيع إلى حدا ما وحجم تأثير منخفض متوقع وقياسات أقل ثباتًا والبيانات لا تتوزع اعتداليًا.

وطرح (Green 1991) قاعدة أخذة في الاعتبار حجم التأثير وهى:

$$N \geq (8 / f^2) + (m - 1)$$

و f حجم التأثير فى حالة تحليل التباين وهى 0.02 صغير و 0.15 متوسط و 0.35 كبير:

$$f^2 = Pr^2 / (1 - Pr^2)$$

• Pr^2 مربع معامل الارتباط الجزئى.

وكلما زاد حجم العينة فإن معامل الارتباط المتعدد يبتعد عن الصفر ولذلك نحصل على الدلالة الإحصائية وبمعنى يمكن التنبؤ بتباين محدد فى المتغير التابع ويكون دال إحصائيًا. ولا اعتبارات إحصائية وعملية يمكن اعتماد على حجم عينة ليس كبيرًا حيث يساعد فى الكشف عن العلاقة الحقيقة وكذلك دلالتها الإحصائية ولكن إذا استخدمت طريقة الانحدار خطوة خطوة Stepwise يفضل زيادة حجم العينة. وقاعدة 40 فرد لكل متغير هى منطقية ومناسبة لان الانحدار الإحصائى يعطى حلول لا يمكن تعميمها على العينة ما لم يكن حجم العينة كبيرًا (Tabachnick & Fidell, 2007)

• **القيم المتطرفة للمتغيرات المستقلة والمتغير التابع:** وجود القيم المتطرفة له تأثير كبير على حلول الانحدار ويؤثر فى دقة تقدير معاملات الانحدار وجودها (عالية جدًا - منخفضة جدًا) يؤدي إلى أن الأخطاء المعيارية لمعاملات الانحدار تكون كبيرة ولذلك يجب استبعاد القيم المتطرفة أو تحوير (تعديل) درجات المتغير أو إعادة تصحيحها. والقيم المتطرفة المتدرجة Multivariate Outlier للمتغيرات المستقلة

يمكن تشخيصها باستخدام طرق إحصائية مثل إحصاء Mahalanobis Distance

• **غياب التلازمية الخطية المتعددة Multicollinearity & Singularity :**

تعرف بالموقف الذى توجد فيه علاقات ارتباطيه قوية بين المتغيرات المنبئة (المستقلة)، وهذا يعنى وجود قياسات تقيس تقريباً نفس السمة كان يوجد بعدين فى قياس وكسلر بلفو للذكاء يقيسوا نفس الشئ ويمكن توضيحها كالآتى:

الجدول (1.7): مصفوفة ارتباط تتضمن الارتباطات العالية بين المتغيرات.

	X_5	X_4	X_3	X_2	X_1
X_1					1.00
X_2				1.00	0.60
X_3			1.00	0.32	0.70
X_4		1.00	0.21	0.21	0.70
X_5	1.00	0.22	0.35	0.13	0.95

ويتضح من الجدول السابق وجود ارتباطات عالية بين X_1 والمتغيرات الاخرى، وخاصة مع المتغير X_5 وهذا يسبب ظهور مشكلة الاعتمادية أو الازدواجية الخطية وعلى ذلك يمكن اختصار X_5 و X_1 فى متغير واحد.

وتظهر التلازمية الخطية المتعددة نتيجة عدة أسباب أهمها الآتى:

- تجميع لمجموعة من المفردات فى درجة واحدة.
- وجود القيم المتطرفة فى بيانات المتغيرات.
- حجم العينة أقل من عدد المتغيرات فى المصفوفة.

قضية بحثية (Pedhazure, 1997):

فيما يلى درجات التحصيل القرائى (Y) والاستعداد اللغوى (x_1) والدافعية للإنجاز (x_2):

	Y	X ₁	X ₂	$\hat{Y}$	Y- $\hat{Y}$ =e
	2	1	3	2.0097	
	4	2	5	3.8981	
	4	1	3	2.0097	
	1	1	4	2.6016	
	5	3	6	5.1947	
	4	4	5	6.3047	
	7	5	6	6.6040	
	9	5	7	7.1959	
	7	7	8	9.1971	
	8	6	4	6.1248	
	5	4	3	4.1236	
	2	3	4	4.0109	
	8	6	6	7.3086	
	6	6	7	7.9005	
	10	8	7	9.3098	
	9	9	6	9.4226	
	3	2	6	4.4900	
	6	6	5	6.7167	
	7	4	9	5.8993	
	10	4	6	7.6750	
مج	117	87	110		
$\bar{x}$	5.85	4.35	5.50		
SS	825	481	658		

ومعادلة الانحدار من الدرجات الخام:

$$\hat{Y}=a+b_1x_1+x_2b_2$$

ومعادلة الانحدار المعيارية: $\hat{Y} = \beta_1 Z_{X_1} + \beta_2 Z_{X_2}$

وعلى ذلك b معامل الانحدار غير المعياري، β معامل الانحدار المعياري

وللانحدار البسيط فإن معادلة الانحدار المعيارية: $Z_{\hat{Y}} = \beta Z_X$

حيث $Z_{\hat{Y}}$ الدرجة المعيارية المتنبأ بها.

ويفسر β معامل الانحدار المعياري كتغير متوقع في المتغير التابع المرتبطة بتغير وحدة في X

. وبما أن الانحراف المعياري Z 1.00 وعلى ذلك فتغير وحدة انحراف معيارية واحدة في X

يؤدي إلى تغير قيمة β في Y .

وتقدر قيمة معامل الانحدار غير المعياري b :

$$b = \frac{\sum XY}{\sum X^2}$$

وتقدر b كذلك: $b = \beta \frac{S_y}{S_x}$

ويمكن التعبير عن β كالآتي: $\beta = b \frac{S_x}{S_y}$

S_x الانحراف المعياري ل X ، S_y الانحراف المعياري ل Y

وعندما يوجد متغير مستقل واحد فإن قيمة β : $\beta = \frac{\sum Z_X Z_Y}{\sum Z_X^2}$

لاحظ عند وجود متغير مستقل واحد فإن:

(معامل ارتباط بيرسون) $r = \beta$ (معامل الانحدار المعياري).

وعند التعامل مع الدرجات المعيارية فإن $a=0$:

$$a = \bar{Y} - \beta \bar{X} = 0 - \beta 0 = 0$$

لان الدرجة المعيارية متوسطها صفر.

ولمتغيرين مستقلين فإن معادلة الانحدار المعيارية:

$$\hat{Z}_Y = \beta_1 Z_1 + \beta_2 Z_2$$

Z_2, Z_1 الدرجات المعيارية ل X_2, X_1 .

وصيغة حساب معامل β لمتغيرين مستقلين هي:

$$\beta_{x1} = \frac{r_{yx1} - r_{yx2} r_{x1x2}}{1 - r_{x1x2}^2}$$

$$\beta_{x2} = \frac{r_{yx2} - r_{yx1} r_{x1x2}}{1 - r_{x1x2}^2}$$

ولاحظ عندما لا يرتبط المتغيرين المستقلين: $r_{x1x2} = 0$ وعليه فإن:

$$\beta_{x1} = r_{yx1}, \quad \beta_{x2} = r_{yx2}$$

كما هو الحال في الانحدار البسيط للبيانات السابقة:

$$r_{x1x2} = 0.522, \quad r_{yx2} = 0.678, \quad r_{yx1} = 0.792$$

وعليه:

$$\beta_{x1} = \frac{0.792 - (0.678)(0.522)}{1 - (0.522)^2} = 0.602$$

$$\beta_{x2} = \frac{0.678 - (0.792)(0.522)}{1 - (0.522)^2} = 0.364$$

ومعادلة الانحدار المعيارية: $\hat{Z}_Y = 0.602 Z_1 + 0.364 Z_2$

وصيغة حساب معامل الانحدار غير المعيارى: $b_j = \beta_j \frac{s_y}{s_j}$

حيث $j=1,2$ ، وبما أن:

$$s_{x2} = 1.670, \quad s_{x1} = 2.323, \quad s_y = 2.720$$

$$b_1 = 0.602 \times \frac{2.720}{2.323} = 0.705 \quad \text{وعليه فإن:}$$

$$b_2 = 0.364 \times \frac{2.720}{1.670} = 0.593$$

وتقدر قيمة الثابت a كالآتي:

$$a = \hat{Y} - b_1 \bar{X}_1 - b_2 \bar{X}_2 = 5.85 - (0.705)(4.35) - (0.5919)(5.50) = -0.4705$$

وعليه فإن معادلة الانحدار غير المعيارية:

$$\hat{Y} = -0.4705 + 0.705 X_1 + 0.593 X_2$$

مربع معامل الارتباط المتعدد (R^2)

يقدر مجموع مربعات الانحدار $SS_{regression}$ كالآتي:

$$SS_{reg} = b_1 \sum X_1 Y + b_2 \sum X_2 Y$$

• $\sum X_1 Y$ مجموع حاصل ضرب قيم X_1 في Y .

• $\sum X_2 Y$ مجموع حاصل ضرب قيم X_2 في Y .

$$SS_{reg} = (0.7046)(95.05) + (0.5919)(58.05) = 101.60$$

ويقدر مربع معامل الارتباط المتعدد كالآتي:

$$R^2 = \frac{SS_{reg}}{\sum Y^2}$$

حيث $\sum Y^2$ مجموع مربعات قيم Y

$$R^2 = \frac{101.60}{140.55} = 0.723 \quad \text{إذا:}$$

وهي نسبة تباين المتغير التابع الذي تم تفسيره عن طريق المتغيرات المستقلة وعليه فإن X_1

، X_2 تفسر 72% من تباين التحصيل القرائي.

وتعتبر $R^2_{Y.12}$ مربع معامل الارتباط بين Y المقاسة و $\hat{Y}$ المتنبأ بها.

$$R^2_{Y.12} = r^2_{y\hat{y}} = \frac{(\sum y\hat{y})^2}{\sum y^2 \sum \hat{y}^2}$$

وتقدر r من خلال حساب معامل ارتباط بيرسون بين Y ، $\hat{Y}$ من المثال السابق :

$$\Sigma \hat{y}^2 = \Sigma \hat{y}^2 - \frac{(\Sigma \hat{y})^2}{N} = 786.0525 - \frac{(117)^2}{20} = 101.60$$

$$\Sigma Y \hat{Y} = \Sigma Y \hat{Y} - \frac{(\Sigma Y)(\Sigma \hat{Y})}{N} = 786.0528 - \frac{(117)(117)}{20} = 101.60$$

وعليه فإن:

$$R^2 = \frac{(101.60)^2}{(140.55)(101.60)} = 0.723$$

والجذر التربيعي يعطى مربع معامل الارتباط R :

$$R_{y.12} = \sqrt{0.723} = 0.85$$

ويمكن تقدير R^2 من الصيغة الآتية:

$$R_{y.12}^2 = \beta_{x1} r_{yx1} + \beta_{x2} r_{yx2}$$

$$R_{y.12}^2 = 0.602 \times 0.792 + 0.364 \times 0.678 = 0.724$$

ويمكن تقديرها من الصيغة الآتية:

$$R_{y.12}^2 = \frac{r_{yx1}^2 + r_{yx2}^2 + 2r_{yx1}r_{yx2}r_{x1x2}}{1 - r_{x1x2}^2}$$

$$R_{y.12}^2 = \frac{(0.792)^2 + (0.678)^2 + 2(0.792)(0.678)(0.522)}{1 - (0.522)^2} = 0.723$$

اختبار الدلالة الإحصائية والتفسيرات

اختبار الدلالة الإحصائية R^2

لاختبار دلالة R^2 من خلال اختبار F كالتى:

$$F = \frac{R^2/K}{(1-R^2)(N-K-1)}$$

• K عدد المتغيرات المستقلة.

• N حجم العينة.

ودرجة الحرية لـ F هي:

$$df_1 = K$$
$$df_2 = N - K - 1$$

وللبينات السابقة فإن:

$$F = \frac{0.723 / 2}{(1 - 0.723) / (20 - 2 - 1)} = \frac{0.3615}{0.0163} = 22.18$$

وبمقارنتها بالقيمة الجدولية فإن $P < 0.01$ وعلى ذلك فإن معادلة التنبؤ تصلح لاستخدامها للتنبؤ بالتحصيل وأن مربع معامل الارتباط بين كلا من الاستعداد والدافعية والتحصيل لا يساوى صفر في المجتمع. وبكلمات أخرى إذا كانت F دالة إحصائياً فعليه يمكن بناء معادلة تنبؤ. ويمكن تقدير F كالاتي:

$$F = \frac{SS_{reg} / df_{reg}}{SS_{res} / df_{res}}$$

• SS_{res} مجموع مربعات البواقي (الخطأ).

• res اختصار للبواقي Residuals

و درجات الحرية للانحدار:

$$df_{reg} = k = 2$$

$$df_{res} = N - K - 1 = 20 - 2 - 1 = 17$$

كما سبق تم حساب مجموع مربعات الانحدار:

$$SS_{reg} = 101.60$$

وتقدر مجموع مربعات البواقي كالاتي:

$$SS_{res} = \sum Y^2 - SS_{reg}$$
$$= 140.55 - 101.60 = 38.95$$

وعليه فإن:

$$F = \frac{101.60 / 2}{38.95 / 17} = \frac{50.80}{2.29} = 22.18$$

وبالتالى يمكن القول أن الاستعداد اللفظى والدافعية للإنجاز فسروا 72% من تباين تحصيل القراءة وهذا الناتج دال إحصائياً عند 0.01.

تقدير فترات الثقة لمعامل الانحدار

تقدر فترات الثقة للانحدار كالاتى:

$$b \pm T\left(\frac{\alpha}{2}, df\right) S_b$$

حيث T هي قيمتها الجدولية أو الدرجة التى نحصل عليها من الجداول عند: df , $\frac{\alpha}{2}$

وقيمة T عند $\frac{\alpha}{2} (0.025)$ مع درجات حرية 17 هي 2.11 ولفترتي الثقة 95%: CI95%:

$$CI95\% \rightarrow b_1 = 0.7046 \pm (2.11)(0.1752)$$

$$= (0.3349, 1.0743)$$

$$CI95\% \rightarrow b_2 = 0.5919 \pm (2.11)(0.2437)$$

$$= (0.0777, 1.1661)$$

وعليه فإن حدود الثقة لكلا من b_1 ، b_2 لم تتضمن الصفر وهى القيمة التى تختبر حولها فروض فى الفرض الصفرى ($H_0: b_1 = 0$ ، $H_0: b_2 = 0$)، وعليه فإن b_1 ، b_2 دالة إحصائياً.

الأهمية النسبية للمتغيرات

قيمة معامل الانحدار تتأثر بمستوى القياس المستخدم لقياس المتغير ويمكن أن تكون القيم الكبيرة نسبياً لـ b ليست ذو معنى جوهري أو ذو دلالة إحصائية وعلى العكس القيم الصغيرة نسبياً لـ b يمكن أن تكون ذو معنى جوهري ودالة إحصائياً، فحجم b لا يستخدم لاستنتاج الأهمية النسبية للمتغيرات ولذلك يجب على الباحث أثناء الحديث عن الأهمية النسبية المقارنة بين قيم β القائمة على الدرجات المعيارية، ففي المثال السابق $\beta_1 = 0.602$ ، $\beta_2 = 0.364$ ولذلك يقوم الباحث باستنتاج أن تأثير المتغير المستقل X_1 يعادل مرة ونصف تأثير X_2 وهذا هو التفسير الصحيح ولكنه ليس خالى من الإشكاليات لأن β تتأثر بتباين المتغير.

مسلمات الانحدار المتعدد (Cohen, Cohen, West, & Aiken, 2003; Keith, 2014)

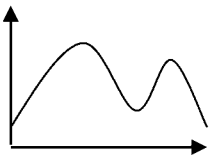
القيم المتطرفة: القيمة المتطرفة هي القيمة التي تختلف اختلافاً كبيراً عن بقية قيم المتغير، فالقيم المتطرفة لها تأثير كبير على معاملات الانحدار وكذلك على حلول الانحدار بصفة عامة. حيث وجود قيمة أو قيمتين متطرفة يؤثر في تفسير النتائج حيث يمكن أن يجعل الأخطاء المعيارية لمعاملات الانحدار صغيرة أو كبيرة ولذلك يجب التحقق من عدم وجود قيم متطرفة في الدرجات وإعادة تصحيحها أو تعديلها. ويمكن تشخيص القيم المتطرفة لكل متغير على حدة من خلال عرض شكل Boxplot، كما يمكن تشخيص القيم المتطرفة المتدرجة باستخدام إحصاء Mahalanobis distance ويتضمن تشخيص القيم المتطرفة قبل التحليل من خلال مدى بعد الحالة أو القيمة من مركز بقية الحالات Centeroid وهو دمج متوسطات كل المتغيرات معاً ويقوم هذا الإحصاء من خلال إحصاء χ^2 حيث التقدير شديد التحفظ هو عند $P < 0.001$ χ^2 تعتبر وجود قيم متطرفة متدرجة. ومؤشر التأثير Influence أو القيم المؤثرة يقدر التغير في معاملات الانحدار عند حذف الحالة، والحالات التي تزيد معامل التأثير لها عن 1.00 من المحتمل أن تكون قيم متطرفة ويتم تقدير ذلك من خلال مؤشرات Cook Distance أو DFFITs أو DBETAS (انظر مخرج الانحدار في SPSS). حيث إذا كان مؤشر $Cook(D) < 1$ فإنه يعتبر كبير ويبدل على وجود قيمة مؤثرة. وهذه المؤشرات Mahalanobis, Cook, يمكن حفظ قيمتهما في ملف البيانات وإجراء الإحصاء الوصفي عليهم. والقيم المتطرفة في المتغير هي الحالات التي لها درجات معيارية كبيرة Z-Score. فالحالات التي تزيد الدرجات المعيارية لها عن 3.29. بالإضافة إلى الدرجة المعيارية توجد الطرق البيانية مثل المدرج التكرارى و Boxplot Normal Probability plot, لتشخيص القيم المتطرفة. وأيضاً يمكن الكشف عن القيم المتطرفة عن طريق Studentized Residuals وإذا كانت أكبر من 2 تعتبر قيمة متطرفة.

اعتدالية البيانات: كل إجراءات الإحصاء المتدرج تتطلب مسلمة الاعتدالية المتدرجة Multivariate Normality بمعنى توفر الاعتدالية لكل متغير على حدة وللاتحاد الخطى للمتغيرات معاً، وعندما تتحقق هذه المسلمة فإن البواقي في التحليل تتوزع

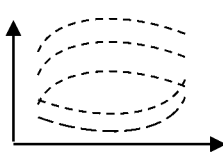
اعتداليًا وتكون مستقلة ويمكن فحص الاعتدالية من خلال طرق إحصائية ورسومات بيانية، ويوجد مؤشرين للاعتدالية هما الالتواء والتفرطح والمتغير ذو البيانات الملتوية متوسطة لا يقع في منتصف أو مركز التوزيع، والمتغير المتفرطح له قمة مدببة أو قمة سميكة متفرطحة، وفي التوزيع الاعتدالي قيمتي الالتواء والتفرطح تساوي صفر. وإذا كان حجم العينة صغيرًا أو متوسطًا فإن يمكن الحكم على التفرطح والالتواء من خلال الدلالة الإحصائية ولكن إذا كان حجم العينة كبيرًا فالحكم عليها من خلال شكل التوزيع أفضل.

اعتدالية وتجانس البواقي Homoscedasticity: وفيها يفترض الانحدار التباينات البواقي (أخطاء القياس) اعتدالية التوزيع للمتغيرات المستقلة ولها تباينات متماثلة أو متساوية عبر كل مستويات المتنبئات ولكن عندما تكون تباينات البواقي أو الأخطاء الواقعة على المتغيرات المستقلة (التباين غير المفسر) للمتغيرات المستقلة غير اعتدالية التوزيع فإنه يطلق عليها ظاهرة **Hetero-scedasticity**، وتحدث نتيجة وجود القيم المتطرفة وعدم تحقق هذه المسلمة يزيد من احتمالية الخطأ من النوع الأول وتكون نتائج اختبار F غير موثوق فيها و تؤدي إلى استنتاجات خاطئة. ويمكن فحص هذه المسلمة من خلال العرض البياني للبواقي للمتغيرات المستقلة:

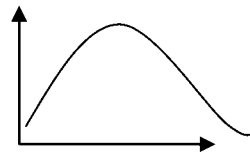
اعتدالية غير متجانسة غير خطية



C



B



A

الشكل (1.7): الشكل البياني للبواقي أو الأخطاء على المتغيرات المستقلة.

وفحص شكل الانتشار للبواقي (الفرق بين $\hat{Y}$ ، Y) يعطى اختبار لمسلمة الاعتدالية والخطية للبواقي حيث يجب أن تكون البواقي علاقة خط أفقي مستقيم مع القيم المتنبأ

بها للمتغير التابع، وتباين البواقي حول المتغير التابع المتنبأ به هي نفسها لكل الدرجات المتنبأ بها وشكل الانتشار للبواقي يتم تمثيلة بوضع الدرجات المتنبأ بها على أحد المحاور وأخطاء التنبؤ (البواقي) على المحور الآخر حيث الدرجات المتنبأ بها $\hat{Y}$ والبواقي (e) درجات معيارية. وتتحقق مسلمة الاعتدالية للبواقي تعنى أن أخطاء التنبؤ اعتدالية التوزيع وعدم تحقق هذه المسلمة يؤدي إلى تحيز لمعاملات الانحدار ولكن خطورتها تكون اشد على اختبارات الدلالة (T,F) وفترات الثقة.

ففي العينات الكبيرة عدم اعتدالية البواقي لا يؤدي إلى مشاكل خطيرة سواء كانت في اختبارات الدلالة أو حدود الثقة. ولكنها تعتبر إشارة على درجة كبيرة من الأهمية إلى وجود مشاكل في نموذج الانحدار مثل سوء تحديد النموذج بمتغيراته.

وعلى ذلك فإن توزيع أخطاء البواقي والدرجات المتنبأ بها في شكل (C) يعنى عدم تحقق مسلمة الخطية للبواقي والتوزيع الاعتدالي وعلية تكون نتائج تحليل الانحدار غير صادقة. ومسلمة **Homoscedastiaty** تعنى أن الانحرافات المعيارية والتباين لأخطاء التنبؤ (البواقي) تقريبا متساوية في مسافتها من كل القيم المتنبأ بها بمعنى أن التباينات للبواقي حول خط الانحدار تكون ثابتة بغض النظر عن قيم X ولكن إذا زادت أخطاء التباين كلما زاد حجم التنبؤ تنشأ قضية عدم تجانس تباينات الأخطاء.

استقلالية البواقي أو أخطاء القياس: هذه المسلمة تختبر من خلال تحليل البواقي وهي أن أخطاء التنبؤ $((Y-\hat{Y})=e)$ مستقلة عن بعضها البعض، بمعنى لا توجد علاقة بين بواقي أى مجموعة من الحالات في التحليل وهذه المسلمة تتحقق إذا كان اختيار العينة عشوائياً وإذا حدث تكوين تجمعات للبيانات فإن البواقي لا تكون مستقلة، وعدم استقلالية البواقي شائعة في تصميمات القياسات المتكررة وعدم تحقق هذه المسلمة لا يؤثر في تقديرات معاملات الانحدار ولكنه يؤثر في الأخطاء المعيارية وهذه المشكلة تؤدي إلى اختبارات دالة إحصائية وأخطاء معيارية غير صحيحة. ويوجد إحصاء **Durbin-Watson** لقياس الارتباط الذاتي Autocorrelation للأخطاء عبر سلسلة من الحالات ولو أنه دال إحصائياً دل على عدم استقلالية للأخطاء، والقيم الموجبة للارتباط الذاتي يجعل تقديرات تباين الخطأ صغيرة ويؤدي إلى تضخم الخطأ

من النوع الأول بينما القيم السالبة تجعل تقديرات تباين الخطأ كبيرة ويؤدي إلى انخفاض القوة الإحصائية.

التحديد الصحيح لشكل العلاقة بين المتغيرات المستقلة والمتغيرات التابعة: من أهم المسلمات للانحدار المتعدد هو التحديد المناسب للشكل الرياضي للعلاقة بين Y وأى من المتغيرات المستقلة فى المجتمع . ويفترض أن تكون كل العلاقات خطية ويتم اختبارها بيانياً من خلال شكل الانتشار Scatter plot .

التحديد الصحيح للمتغيرات المستقلة فى نموذج الانحدار: لا بد أن توجد نظرية مفترضة فى ضوءها يتم تحديد المتغيرات المستقلة الأكثر تأثيراً فى المتغير التابع. إذا تبنى الباحث نظرية ما واختبار صحيح وتم قياس للمتغيرات بدون أخطاء وشكل العلاقة بين المتغيرات المستقلة والتابعة خطية فإن البواقي تكون مستقلة عن بعضها وتقديرات معاملات الانحدار تكون غير متحيزة.

المتغيرات المستقلة بدون أخطاء قياس (تامة الثبات): أحد تعريفات الثبات بأنه العلاقة بين المتغير كما يقاس هو انعكاس للدرجة الحقيقية. وعندما لا يوجد أخطاء قياس فى المتغير المقاس X فإن $r_{xx}=1.00$. ولكن فى العلوم السلوكية فإن الثبات لا يكون واحد صحيح إنما يأخذ قيم متفاوتة من صفر إلى الواحد الصحيح ولكن لم تتحقق هذه المسلمة فإن تقدير r_{xy} يكون متحيز ويحدث تقلص لمعاملات الانحدار b وكما أن R^2 يكون تقديرها منخفض عن قيمتها الحقيقية. ويتم التحقق من الثبات من خلال مقاييس الاتساق الداخلى مثل معامل ألفا كرونباخ وهو يعرض متوسط الارتباطات بين كل نصفين محتملين وكذلك يقدر الثبات من خلال الاختبار وإعادة.

التلازمة الخطية المتعددة Multicollnearity and Singularity: حساب معاملات الانحدار يتطلب معكوس مصفوفة الارتباطات بين المتغيرات المستقلة وهذا المعكوس مستحيل الحصول عليه ما لم يكن محدد المصفوفة موجب. وتحدث قضية التلازمة الخطية نتيجة الارتباطات العالية بين المتغيرات المستقلة أو وجود تفاعلات

بين المتغيرات المستقلة، ومع التلازمية الخطية فإن محدد المصفوفة لا يكون صفرًا تمامًا.

ويتم تشخيص التلازمية الخطية عن طريق:

- **مربعات معامل الارتباط المتعدد (SMC):** القيمة المرتفعة لمربع معامل الارتباط المتعدد بين أحد المتغيرات المستقلة وبقية المتغيرات الأخرى مستقلة وإذا كانت قيمة SMC مرتفعة فإن المتغير يرتبط ارتباطًا عاليًا مع بقية المتغيرات المستقلة وتوجد قضية التلازمية وإذا كان $SMC=1$ فإن المتغير يرتبط ارتباطًا تامًا مع بقية المتغيرات المستقلة.

- **القيمة المنخفضة لمعامل توليرانس: $Tolerance = 1 - SMC$** وإذا كانت قيمة المعامل $Tolerance$ أعلى من 0.5 أو 0.6 وهذا يؤدي إلى صعوبات في اختبار وتفسير معاملات الانحدار.

- **فحص الارتباطات البسيطة بين المنبئات في مصفوفة الارتباط** حيث إذا زاد قيمة معامل الارتباط بين أي متغيرين مستقلين عن 0.70 أو 0.80 فإنه توجد التلازمية الخطية بين المتغيرين وعليه يجب حذف أحد المتغيرين أو دمجها معًا.

- **فحص مؤشر Variance Inflation Factors (VIF):** وهو يشير إلى وجود علاقة خطية قوية بين أحد المتغيرات المنبئة وبقية المتغيرات المنبئة وإذا زادت قيمته عن 1.00 فإنه يوجد قضية التلازمية الخطية للمتغير المنبئ ويجب أن يحذف من تحليل الانحدار.

والتلازمية الخطية تسبب مشاكل خطيرة في تحليل الانحدار أهمها:

- **تقلص حجم قيمة R**
- **تجعل تحديد أهمية المنبئات صعبة لأن تأثير المنبئات يكون متداخل نتيجة الارتباطات العالية بينها.**
- **يزيد تباينات معاملات الانحدار وهذا يجعل معادلة التنبؤ غير صادقة.**
- **تزيد من الأخطاء المعيارية لمعاملات الانحدار وعندما تكون قيمة r بين متغيرين مستقلين $r=0.9$ فإن الأخطاء المعيارية لمعاملات الانحدار تتضاعف وعليه فإن**

أى من معاملات الانحدار لا تكون دالة إحصائيًا نتيجة الحجم الكبير للأخطاء المعيارية.

توجد العديد من الطرق للتغلب على التلازمة الخطية منها:

• دمج المتغيرات المنبئة عالية الارتباط معًا، فمثلاً 0.9، $r_{x_1x_2}=0.8$ وعالية دمج المتغيرين x_1 ، x_2 معًا.

• إذا وجد مجموعة كبيرة من المنبئات فيتم إجراء التحليل العاملى أو تحليل المكونات الرئيسية لتقليل العدد الكبير من المنبئات عالية الارتباط إلى مجموعة أقل منها.

• استخدام طريقة انحدار ريدج Ridge regression

طرق بناء نموذج الانحدار المتعدد

توجد ثلاث استراتيجيات تحليلية فى الانحدار المتعدد وهى الانحدار المتعدد المعيارى Standard multiple regression وطريقة الانحدار التسلسلى الهرمى Sequential Hierarchical regression والانحدار الإحصائى ذو الخطوة التدريجية Statistical stepwise regression والاختلاف بينهم يتضمن ماذا يحدث اثناء تداخل التباين نتيجة المتغيرات المستقلة المرتبطة ومن يحدد ترتيب إدخال المتغيرات المستقلة فى المعادلة.

أولاً: الانحدار المتعدد المعيارى: أحد الطرق المستخدمة لقواعد البيانات الصغيرة وفى الطريقة التلازمية أو المعيارية كل المتغيرات المستقلة يتم إدخالها فى معادلة الانحدار فى نفس الوقت أو اللحظة ولذلك تسمى الانحدار فى نفس الوقت Simultaneous أيضاً يشار اليه بالانحدار الإدخال القصرى Forced entry regression لان كل المتغيرات تجبر على الإدخال فى المعادلة فى نفس الوقت حيث كل متغير مستقل يتم تقديره كم لو أنه ادخل فى الانحدار بعد إدخال كل المتغيرات المستقلة، ويتم تقييم كل متغير مستقل فى ضوء ماذا اضاف فى التنبؤ بالمتغير التابع. فى هذه الطريقة يمكن للمتغير المستقل أن يرتبط ارتباطاً مرتفعاً مع المتغير التابع ولكن يظهر غير ذات اهمية فى إسهامه فى التنبؤ بالمتغير التابع وتسمى فى برنامج SPSS بالطريقة Enter.

وهذه الطريقة مفيدة فى البحث الاستكشافى لتحديد الدرجة التى يؤثر فيها متغير واحد أو أكثر على ناتج معين وأيضاً مفيد فى تحديد التأثير النسبى لأحد المتغيرات المستقلة وأيضاً فى تقدير التأثيرات المباشرة لأحد المتغيرات المستقلة على متغير تابع.

ثانياً: الانحدار المتعدد التسلسلى (الهرمى) : وهى مشابهة بالطريقة السابقة وتستخدم لأغراض الاستكشاف حيث يتم إدخال المتغيرات المستقلة إلى المعادلة فى ترتيب محدد من قبل الباحث وكل متغير مستقل يتم تقييمه فى ضوء ماذا أضاف إلى المعادلة عند إدخاله فيها، وعليه يتم تحديد ترتيب إدخال المتغيرات وفقاً إلى اعتبارات نظرية أو منطقية ويمكن أن يتم إدخال المتغيرات المستقلة فى نفس الوقت أو كمجموعة Blocks ثم يقوم البرنامج بإجرائه فى خطوات وعلى ذلك فإن إدخال المتغيرات المستقلة يتوقف على اعتبارات نظرية وفى بعض الاحيان يكون لاختيار تسلسل المتغيرات المستقلة غير واضح فى ضوء النظرية.

ثالثاً: الانحدار لإحصائى ذو الخطوات التدريجية: وفيها يتم إدخال المتغيرات المستقلة فى ضوء ترتيب قائم على معايير إحصائية. فالقرار على أى المتغيرات يتم تضمينها وأى المتغيرات يتم استبعادها قائم على حسابات رياضية. فمثلاً كلاً من المتغير المستقل الأول والمتغير المستقل الثانى يرتبطوا ارتباطاً جوهرياً مع المتغير التابع بينما المتغير المستقل الثالث أقل ارتباطاً، فالمفاضلة بين المتغير المستقل الأول والثانى لإدخالهم أولاً إلى المعادلة قائم على أى منهما أكثر ارتباطاً مع المتغير التابع فإذا كان المتغير المستقل الأول أكثر ارتباطاً يتم إدخاله أولاً فى نموذج الانحدار ثم يتم إدخال المتغير الثانى لنرى مدى إسهامه فى تحسين قيمة R^2 ويستمر إدخال المتغيرات حتى نصل إلى مرحلة فيها إضافة أى من المتغيرات المستقلة لا يضيف شيئاً فى تباين المتغير التابع R^2 وعليه يتم استبعادها من معادلة الانحدار وتوجد ثلاثة صور للانحدار الإحصائى:

1. الانحدار الامامى Forward method : فيها يبدأ البرنامج بتحديد النموذج المبدئى فى ضوء الثابت فقط (a) ثم يبدأ البرنامج فى البحث عن افضل المنبئات بالمتغير التابع (اعلى معامل ارتباط) ولو اضاف هذا المتغير تحسن جوهري فى قدرة النموذج للتنبؤ بالناتج فانه يبقى فى المعادلة ثم يبحث البرنامج على منبئ ثانى

والقاعدة فى اختبار المنبئ الثانى هى اعلى معامل ارتباط جزئى Partial correlation مع الناتج أو المتغير التابع وهكذا فمثلا إذا كان الذكاء يسهم بنسبة 40% من تباين التحصيل فيبقى 60% لم تفسر فيقوم البرنامج بالبحث عن متغير مستقل يستطيع أن يفسر أكبر جزء من تباين الـ 60% غير المفسرة. ومعامل الارتباط الجزئى هو مقياس لكمية التباين الجديد فى المتغير التابع التى يمكن أن تفسر عن طريق بقية المنبئات مع استبعاد المتغير الأول. على ذلك تبدأ المعادلة بالمتغير الأكثر إسهاماً ثم تضاف إليها متغير تلو الآخر ويتم دراسة الدلالة الإحصائية لمدى إسهامه فى المعادلة حتى نصل إلى مرحلة اضافته أى متغير لا يسهم فى تباين المتغير التابع.

2. الانحدار الخلفى Backward method : وهى طريقة عكس طريقة الانحدار الامامى حيث يبدأ البرنامج بوضع كل المنبئات فى النموذج أو المعادلة ويقدر اسهام كل منبئ عن طريق فحص الدلالة الإحصائية لـ T لكل منبئ وإذا كان معامل الانحدار دالة إحصائياً يبقى فى المعادلة وإذا كان غير دالة إحصائياً يستبعد من المعادلة ثم يعاد تقدير النموذج بالمتغيرات التى تسهم اسهاماً فعالاً فى المتغير التابع.

وتمر هذه الطريقة بالخطوات الآتية:

أ. حساب معادلة التنبؤ لكل المنبئات.

ب. حساب F لكل متغير منبئ مع وجود المتغير الأكثر اسهاماً فى المعادلة.

ج. القيمة الصغيرة لـ F يتم التحقق من دلالتها وإذا كانت غير دالة يتم حذف المتغير المنبئ من المعادلة ويعاد تقدير المعادلة مع المتغيرات الباقى.

3. طريقة الانحدار التدريجى (خطوة خطوة) method Stepwise : هى خليط بين

الطريقتين السابقتين ما عدا توقيت اضافة المتغير المنبئ إلى المعادلة. حيث تبدأ المعادلة بالثابت وتضاف المتغيرات المستقلة فى نفس الوقت إذا توافر لها المعايير الإحصائية وهكذا ولكن يتم اختبار الدلالة الإحصائية للمتغير المنبئ الأقل فائدة أو إسهاماً فى المعادلة أولاً وعليه فإن المعادلة تبدأ بالمتغيرات المستقلة كلها وليس واحد

تلقوا الآخر، في نفس الوقت وإذا استوفت المعايير الإحصائية من حيث دلالتها تبقى ويمكن حذفها في أي خطوة عندما لا تسهم بفاعلية في الانحدار.

وعندما يستخدم الباحث طريقة الانحدار التدريجي فإن البرنامج ينفذ سلسلة من الخطوات، في الخطوة الأولى يحدد البرنامج المتغير المستقل الأكثر اسهاماً في تفسير معظم التباين في المتغير التابع R^2 وتظل هذه الخطوات بمعنى يختار المتغير الثاني المسئول عن تفسير قدر كبير من R^2 بالتزامن مع المتغير الأول إلى أن يصل إلى الخطوة التي يكون فيها إضافة المتغير المستقل لا يضيف أي شيء في تباين المتغير التابع وتتوقف هذه العملية. يتم الاختبار والإدخال في المعادلة على أساس اسهامها في R^2 . ولكن توجد خطوة في أن نترك للبرنامج الحرية في وضع تسلسل للمتغيرات في المعادلة، وطريقة Stepwise صممت للاختيار من مجموعة من المتغيرات المستقلة وتحديد أحد المتغيرات يتم إدخاله في كل مرحلة بحيث يكون له اسهاماً عالياً في R^2 . والبرنامج يقف في اضافته أي من المتغيرات المستقلة إلى المعادلة عندما لا يكون له اسهاماً أو دلالة إحصائية.

وعندما يكون لدى الباحث مجموعة كبيرة من المتغيرات وليس لديه أسس نظرية قوية ترشده في اختيار أي من هذه المتغيرات فإن البرنامج يختار هذه المتغيرات ويضع القرار فيما يخص الأسبقية أو الأولوية المنطقية أو السببية التي يصعب الوصول إليها قبل التحليل. وعليه فإن تفسير النتائج في الانحدار الإحصائي ليس سهلاً. وعلى ذلك فإن لانحدار الإحصائي التدريجي يكون مفيد في البحوث الاستكشافية وفي استخدام الانحدار الهرمي فإن الباحث لديه رؤية واضحة عن ترتيب المتغيرات واسبقيتها في حدوث السببية قبل متغيرات أخرى.

وتوجد مشاكل خطيرة من استخدام الانحدار التدريجي أو التسلسلي عندما يوجد عدد كبير نسبياً من المتغيرات المستقلة لأن اختبارات الدلالة لإسهام كل متغير مستقل في R^2 وفترات الثقة المرتبطة به تتم في تجاهل المتغيرات المستقلة المتنافسة معه وهذا يزيد من فرصة أن تكون فترات الثقة تحت التقدير (أقل من قيمتها الحقيقية)، وتظهر مشكلة أخرى وهي أن تحليل الانحدار التدريجي حيث يعطى معالم في العينة قد تختلف عن معالم و معادلة الانحدار في عينة أخرى لنفس المجتمع.

وعندما يوجد عدد من المتغيرات المستقلة التى تتنافس فى الدخول إلى المعادلة ففى خطوة معينة توجد فروق تافهة بين الارتباطات الجزئية مع Y وهنا يكون البرنامج فى حيرة فى اختيار معامل الارتباط الجزئى الأكبر فى هذه الخطوة. وعلى الرغم أن طريقة الانحدار التدريجى صممت لتعظيم R^2 مع عدد محدد من المتغيرات المستقلة للينة ولكنها لم تتجح جيداً فى الممارسة الواقعية. ويمكن زيادة فاعلية طريقة Stepwise بالآتى:

• هدف البحث التنبؤ وليس التفسير لأن التفسير الفعال لنتائج الانحدار التدريجى يشوبها التحذيرات السابقة.

• حجم العينة كبير جداً وعدد المتغيرات المستقلة ليس كبيراً حيث نسبة المتغيرات المستقلة (K) إلى العينة (n) هى 1 إلى 40 كافي.

• إذا كانت النتائج لها قابلية للتفسير فيفضل إجراء Cross-validation ويمكن إجراء ذلك بتقسيم العينة الأساسية عشوائياً إلى عینتين ويتم معالجتها بحثياً وإحصائياً بنفس الطريقة.

والطرق الإحصائية الثلاثة الانحدار الخلفى والانحدار الامامى والانحدار التدريجى يقعوا تحت عنوان عام وهى الطريقة التدريجية Stepwise Methods حيث يعتمدوا على انتقاء البرنامج للمتغيرات فى ضوء اسس إحصائية، وعليه فإن هذه الطرق تستخدم للبحوث الاستكشافية ولكن عندما يوجد اطر نظرية قوية فى البحث فيفضل استخدام الانحدار التسلسلى أو الهرمى. وعليه فأن:

1. لتبسط العلاقات بين المتغيرات وتجيب عن السؤال الرئيسى للبحث فإن الطريقة المعيارية الاجبارية مفضلة.

2. يفضل استخدام الانحدار التسلسلى فى وجود اطار نظرى قوى لاختبار الفروض وفيها تتحدد اهمية المتغيرات المستقلة فى معادلة التنبؤ عن طريق الباحث وفقاً للنظرية والمنطق.

3. فى الانحدار الإحصائى يتم إدخال وترتيب المتغيرات المستقلة فى المعادلة فى ضوء إحصائيات مقدرة من بيانات العينة ولذلك فهى استراتيجية لبناء النموذج

Model Building وليس لاختبار النموذج Model Testing

4. نتائج الانحدار الإحصائى يمكن أن تكون مضللة ما لم تكون أحجام العينات كبيرة وممثلة تمثيلاً جيداً للمجتمع.

5. إذا استخدم الانحدار المعيارى الاجبارى فيوجد سؤالين رئيسيين:

أ- ما هو حجم العلاقة بين المتغير التابع والمتغير المستقل R ؟

ب- ما مقدار العلاقة أو التباين التى أسهم بها كل متغير مستقل؟

6. إذا استخدم الانحدار الهرمى التسلسلى فإن السؤال هو:

هل اضافته متغير الذكاء اضاف إحصائياً إلى التنبؤ بالتحصيل بعد استبعاد

متغيرات الميل والدافعية مثلاً؟

7. إذا استخدم الانحدار الإحصائى فإن السؤال هو:

ما هو أفضل اتحاد خطى للمتغيرات المستقلة للتنبؤ بالمتغير التابع فى العينة؟

تنفيذ الانحدار المتعدد بمسلماته فى SPSS

جمع باحث بيانات لـ 175 طالب لأربعة متغيرات مستقلة هى:

الدافعية: يتضمن بعدين داخلية mot1 وخارجية mot2

الاتجاه: يتضمن بعدين هما الاتجاه نحو المادة att1 ونحو المدرسة att2

والمتغير التابع هو التحصيل achievement

أولاً: إدخال البيانات: 1. اضغط variable view وتحت عمود Name اكتب مسمى

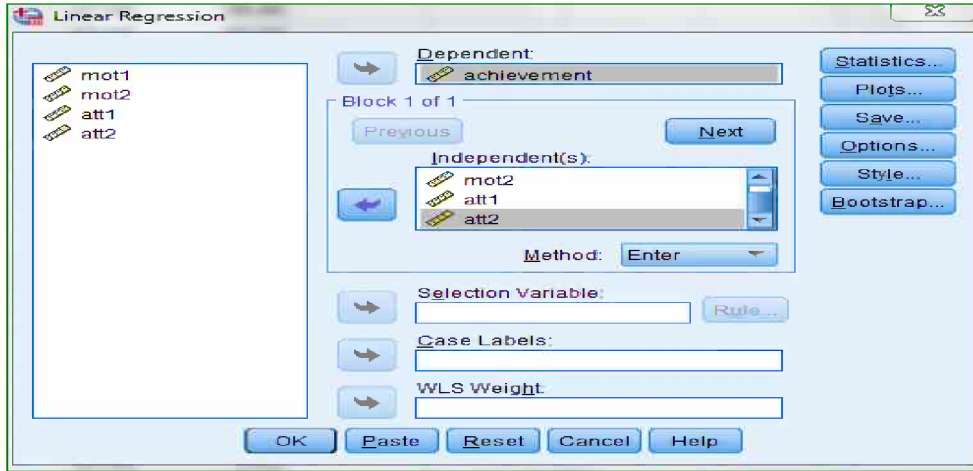
المتغيرات وهى المستقلة mot1 , mot2 , att1 , att2 والمتغير التابع achievement

2. اضغط Data view تظهر شاشة بها خمس أعمدة كل عمود يمثل متغير

3. ابدأ فى إدخال البيانات

ثانياً: تنفيذ الأمر بمسلماته:

1. اضغط Analyze → Regression → Linear يعطى شاشة مماثلة للانحدار البسيط:

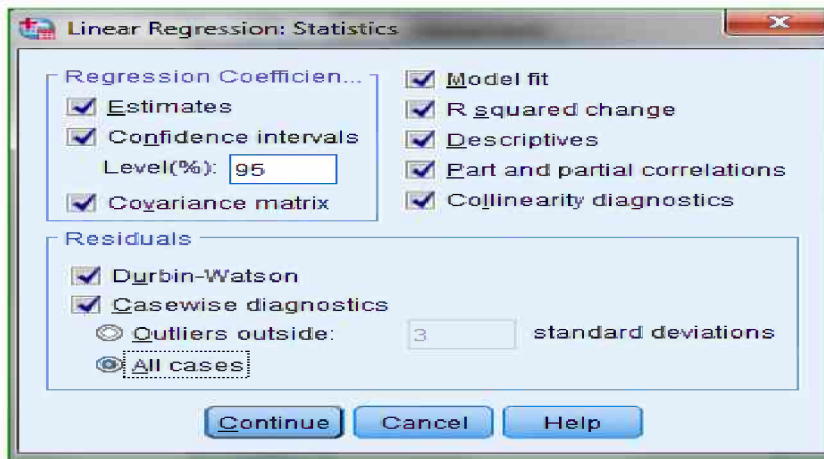


2. انقل المتغير التابع achievement إلى مربع Dependent

3. انقل المتغيرات المستقلة mot1 , mot2 , att1 , att2 إلى مربع Independent

4. في مربع Block 1 of 1 يوجد الطريقة Method ويظهر في المربع الذى امامها طريقة Enter وهى default للبرنامج واذا ضغطت على السهم تظهر طرق الانحدار المتعدد والباحث اختار Enter

5. اضغط على اختيار Statistics تظهر الشاشة الآتية:



6. على الجانب الأيمن من الشاشة يوجد عدة اختيارات وهى:

- اختيار Model fit : (هى منشطة) وهذا يفيد فى صياغة إحصائيات المطابقة وتتضمن R^2 و $adjustedR^2$ والخطأ المعياري للتقدير وجدول ANOVA للنموذج
- Descriptive : يعطى المتوسط والانحراف المعياري لكل متغير فى التحليل وعدد الأفراد لكل المتغيرات ويمكن أن يعطى مصفوفة الارتباط بين كل المتغيرات فى التحليل وكذلك اختبار الدلالة الإحصائية لكل معامل ارتباط لاختبار ذو ذيل واحد وهى مفيدة فى فحص قضية التلازمية الخطية عن طريق فحص معاملات الارتباط.

- Part and Partial correlations: وهذا الاختبار يعطى معامل ارتباط بيرسون r ومعامل الارتباط الجزئى Partial وهى تفيد فى حساب معامل الارتباط بين المتغير التابع ومتغير مستقل محدد بعد ضبط تأثيرات بقية المتغيرات المستقلة الاخرى فى النموذج وهذا المعامل يقوم عليه تحليل الانحدار. بينما معامل الارتباط شبه الجزئى Semipartial يتناول العلاقة بين كل متغير مستقل (منبئ) وجزء من تباين المتغير التابع الذى لم يفسر عن طريق بقية المتغيرات الاخرى فى النموذج.

- Collinearity diagnostics: يعطى مؤشرى Tolerance وقيم VIF لكل متغير مستقل فى التحليل وهذا يفيد فى تشخيص التلازمية الخطية بين المتغيرات المستقلة.

اما فى مربع Regression Coefficient توجد عدة اختيارات منها:

- R square change: يعرض التغير فى R^2 جراء تضمين متغير منبئ جديد (او مجموعة من المتغيرات) وهذا المقياس مفيد لتقدير اسهام المنبئات الجديدة فى تفسير تباين المتغير التابع.

- Estimates يعطى معاملات الانحدار اللامعيارية والمعيارية (b , β) واختبارات الدلالة T لمعاملات الانحدار والخطأ المعيارى.

- Confidence intervals: يعطى فترات الثقة لمعاملات الانحدار غير المعيارية وهى مفيدة لتقدير القيمة الاحتمالية لمعاملات الانحدار فى المجتمع.

- Covariance Matrix: يعرض مصفوفة التغايرات ومعاملات الارتباطات والتباينات بين معاملات الانحدار لكل متغير فى النموذج.

اما فى مربع Residuals توجد عدة اختيارات:

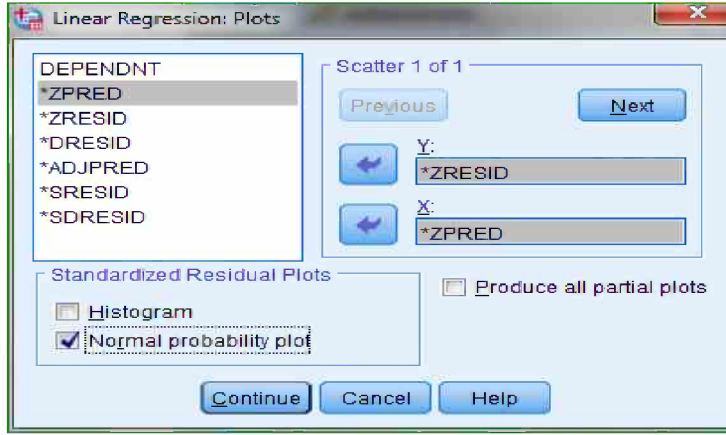
- Durbin - Watson: إحصاء للتحقق من استقلالية الأخطاء أو البواقي وبرنامج SPSS لا يعطى اختبار دلالة إحصائية لهذا الإحصاء وإذا كانت قيمته تختلف بدرجة كافية عن 2 فانه لا توجد استقلالية بين أخطاء القياس.

- Casewise Diagnostic: يعطى القيمة المقاسة للمتغير التابع Y والقيمة المتنبأ بها $\hat{Y}$ والفرق بينهما (الأخطاء أو البواقي المعيارية) و يعطى هذه القيم لكل الحالات (اختار All Cases) أو للحالات التى تزيد لها البواقي المعيارية عن 3 ويمكن تغير 3

إلى 2 من خلال Outliers outside

7. اضغط Continue

8. اضغط على اختيار Plot على يمين الشاشة تظهر الشاشة الآتية:



وهذا الاختيار لفحص مسلمة الاعتدالية وتجانس التباينات للبواقي، فعلى يمين الشاشة تظهر مسمى المتغيرات وهي:

- درجات المتغير التابع الخام Dependent

- ZPRED : قيم المتغير التابع المعيارية المنتبأ بها فى ضوء معادلة الانحدار التى تم بنائها.

- ZRESID : البواقي أو الأخطاء المعيارية وهى الفروق المعيارية بين الدرجة المقاسة Y والدرجة المنتبأ بها عن طريق النموذج وكما نعلم أن البواقي غير المعيارية هى الفرق بين $\hat{Y}$ و Y ولا يمكن تحديد نقطة قطع لتحديد البواقي الكبيرة. وللتغلب على هذه المشكلة نستخدم البواقي المعيارية وهى خارج قسمة البواقي اللامعيارية على الانحراف المعياري ويمكن أن تقارن البواقي المعيارية Z للنماذج المختلفة فاذا كان:

- القيمة المطلقة للبواقي أكبر من 3.29 (تقريباً 3) لا بد أن تؤخذ فى الاعتبار.
- إذا وجد 1% من الحالات لها بواقي معيارية أكبر من 2.58 فهذا دليل على أن مستوى الخطأ داخل النموذج غير مقبول بمعنى أن النموذج ضعيف المطابقة مع البيانات.

- إذا وجد 5% من الحالات لها بواقي معيارية 1.96 فأكثر (تقريباً 2) فإنه دليل على أن النموذج لا يمثل البيانات تمثيلاً دقيقاً.

اما الشكل الآخر لبواقي هو Studentized Residual وهو بواقي غير معيارية مقسومة على الانحراف المعياري ولها نفس خصائص البواقي المعيارية ولكنها عادة تعطينا تقدير دقيق لتباين الخطأ.

- DRESID : البواقي المحذوفة

- ADJPRED : القيم المتنبأ بها المصححة

- SRESID : Studentized Residual البواقي المعدلة

- SDRESID : البواقي المعدلة المحذوفة

9. ولتقييم الاعتدالية وتجانس التباينات للبواقي Homoscedasticity اتبع الآتي:

أ- اضغط على متغير ZRESID وانقله إلى مربع Y

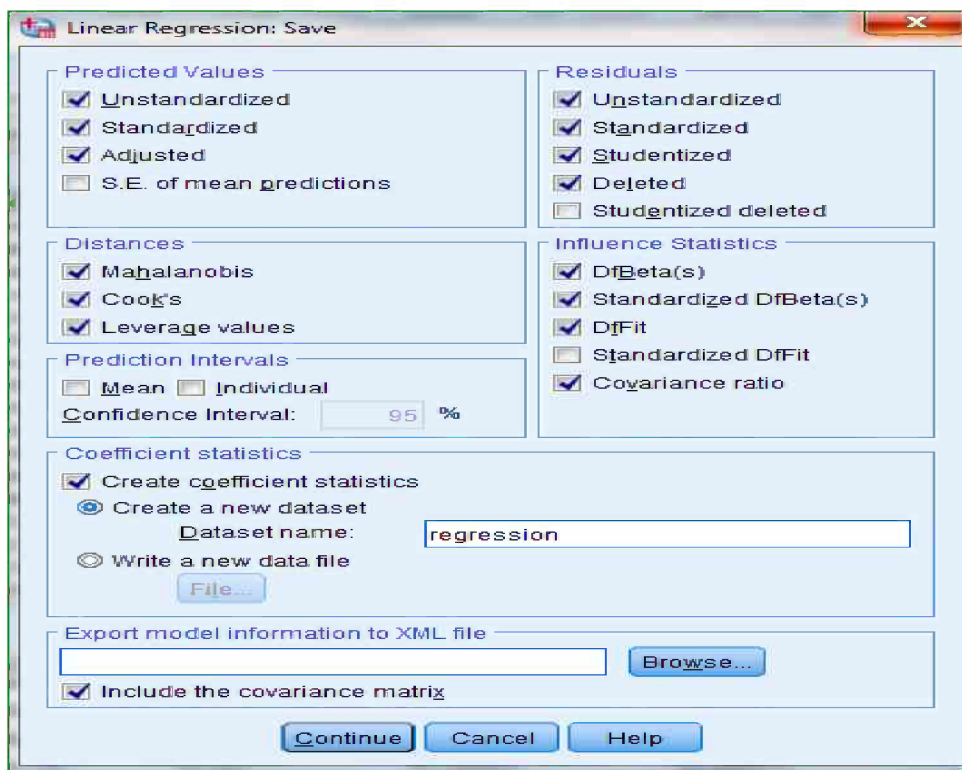
ب- اضغط على متغير ZPRED وانقله إلى مربع X

ت- تحت جزء Studentized Residual plots اضغط على مسلمة Normal

histogram و probability plot لتقويم اعتدالية البواقي.

10. اضغط Continue

11. اضغط على اختيار Save تظهر الشاشة الآتية:



حيث يمكن حفظ المتغيرات التشخيصية الناتجة من الخطوات السابقة فى ملف البيانات. حيث يقوم البرنامج بحساب هذه الإحصائيات وخلق أعمدة جديدة فى ملف البيانات

12. فى مربع Predicted value اضغط على:

أ- الدرجة المتنبأ بها Unstandardized

ب- الدرجة المتنبأ بها المعيارية (Z) Standardized

ت- Adjusted Predicted Value

13. فى مربع Residuals اضغط على:

أ- البواقي الخام Unstandardized

ب- البواقي المعيارية Standardized

ت- Studentized residual

ج- Deleted residual هو البواقي عندما يستبعد الفرد من التحليل وهى الفرق بين المتغير

التابع والقيمة المتنبأ بها المصححة عند استبعاد الحالة.

14. فى مربع Distance اضغط:

- Mahalanobis distance: وهو مؤشر لمدى ابتعاد قيمة المتغير المستقل عن

المتوسط العام لتشخيص القيم المتطرفة.

- Cook distance: إلى أى درجة البواقي لكل الأفراد تتغير إذا استبعد الفرد من التحليل

- Leverage : مقياس لمدى تأثير قيم المتغير المستقل على مطابقة نموذج الانحدار.

15. إحصائيات التأثير: تتضمن إحصائيات تقيس ماذا يحدث لو استبعدنا الحالة من التحليل وهى:

- DfBeta: الفرق فى معاملات الانحدار المعيارية (Beta) لو استبعدت الحالة من التحليل

- DfFit: التغير فى القيمة المتنبأ بها لو استبعدت الحالة أو الفرد من التحليل

- Standardized DfBeta: الدرجة المعيارية ل DfFit

- Covarince ratio : نسبة محدد مصفوفة التغاير فى حالة استبعاد حالة أو فرد محدد من التحليل إلى محدد المصفوفة التعاير مع وجود الفرد من التحليل والنسب القريبة من 1.0 تشير إلى خطأ قليل.

16. اضغط Create coefficient Statistics لحفظ معاملات الانحدار فى قاعدة بيانات جديدة واكتب مسمى قاعدة البيانات الجديدة كالاتى:

- Create anew dataset

- اكتب اسم الملف Datasetname واكتب مسماة وليكن Regression

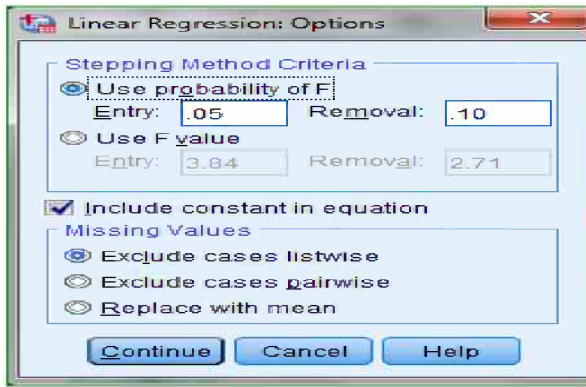
او - اضغط data

- write anew file

وحدد موقع الملف الجديد فمثلا تم تحديده فى مسار D مثلا واكتب اسمه Regression واضغط Save.

17. اضغط Continue

18. اضغط اختيار Options تظهر الشاشة الآتية :



- يمكن تغيير محركات إدخال المتغيرات فى الانحدار التدريجى Stepwise ومن الافضل استبقاء محك 0.05 ويمكن تغييره إلى معيار أكثر صرامة وهو 0.01.

- يمكن تحديد طريقة التعامل مع

البيانات الغائبة سواء استبعاد الحالة من التحليل التى بها قيمة غائبة على أحد المتغيرات اختار Listwise أو الاستبقاء على الحالة التى بها بيانات غائبة على أحد المتغيرات للاستفادة منها فى متغيرات اخرى Pairwise أو استبدال القيم الغائبة بمتوسط المتغير.

19. اضغط Continue

20. اضغط OK لتنفيذ الامر

ثالثا : المخرج :

```
REGRESSION
/DESCRIPTIVES MEAN STDDEV CORR SIG N
/MISSING LISTWISE
/STATISTICS COEFF OUTS CI(95) BCOV R ANOVA COLLIN TOL CHANGE ZPP
/CRITERIA=PIN(.05) POUT(.10)
/NOORIGIN
/DEPENDENT achievement
/METHOD=ENTER mot1 mot2 att1 att2
/SCATTERPLOT=(*ZRESID ,*ZPRED)
/RESIDUALS DURBIN NORMPROB(ZRESID)
/CASEWISE PLOT(ZRESID) ALL
/SAVE PRED ZPRED ADJPRED MAHAL COOK LEVER RESID ZRESID SRESID DRESID DFBETA SDBETA DFFIT COVRATIO
/OUTFILE=COVB(regression).
```

• اعطى المخرج الإحصاء الوصفى للمتغيرات فى التحليل:
حيث اعطى المتوسط Mean والانحراف المعياري وعدد افراد أو القياسات فى المتغيرات الاربعة.

Descriptive Statistics			
	Mean	Std. Deviation	N
achievement	17.5714	9.97205	175
mot1	25.5314	4.43832	175
mot2	29.0514	3.94687	175
att1	40.7257	7.15382	175
att2	45.9829	9.24412	175

• مصفوفة الارتباط:

Correlations						
		achievement	mot1	mot2	att1	att2
Pearson Correlation	achievement	1.000	.398	.226	.308	.348
	mot1	.398	1.000	.545	.138	.326
	mot2	.226	.545	1.000	.077	.281
	att1	.308	.138	.077	1.000	.679
	att2	.348	.326	.281	.679	1.000
Sig. (1-tailed)	achievement	.	.000	.001	.000	.000
	mot1	.000	.	.000	.034	.000
	mot2	.001	.000	.	.157	.000
	att1	.000	.034	.157	.	.000
	att2	.000	.000	.000	.000	.
N	achievement	175	175	175	175	175
	mot1	175	175	175	175	175
	mot2	175	175	175	175	175
	att1	175	175	175	175	175
	att2	175	175	175	175	175

حيث فى الجزء الأول اعطى معاملات الارتباط لبيرسون وفى الجزء الثانى دلالتها الإحصائية عند 01. لاختيار ذو ذيلين والجزء الثالث عدد الأفراد التى دخلت قياساتها فى حساب معامل الارتباط. والملاحظة فى مصفوفة معاملات الارتباط أن العلاقة بين المتغيرات المستقلة والتابعة منخفضة إلى حدًا ما. كما أن العلاقة بين mot1 و mot2 مرتفعة إلى حدًا ما 545. كما أن العلاقة بين المتغيرين المستقلين att1 و att2 مرتفعة 6790. وهذا غير مرغوب فى تحليل الانحدار حيث يفترض أن تكون العلاقة

بين المتغيرات المستقلة ببعضها البعض منخفضة وعلى ذلك من خلال فحص مصفوفة الارتباط نتوقع ظهور قضية التلازمية الخطية إلى حدًا ما خاصة بين att1 و att2 . ويتضح أن كل معاملات الارتباط دالة إحصائية عند 0.05. ما عدا معامل الارتباط بين att1 و mot2، وكذلك دالة عند 0.01. ما عدا معامل الارتباط بين att1 و mot2 وكذلك معامل الارتباط بين att1 و mot1 .

• اعطى البرنامج الجدول الآتى:

Variables Entered/Removed ^a			
Model	Variables Entered	Variables Removed	Method
1	att2, mot2, mot1, att1 ^b	.	Enter
a. Dependent Variable: achievement			
b. All requested variables entered.			

ويتضح أن المتغيرات المستقلة جميعها ادخلت المعادلة ولا توجد متغيرات مستبعدة وذلك لان الطريقة المستخدمة هي ENTER ولكن إذا استخدم الباحث أى

من طرق الانحدار الأخرى فيوجد متغيرات مستبعدة وأوضح الجدول أن المتغير التابع هو achievement

• ملخص النموذج Summary of Model كالاتى :

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	.480 ^a	.230	.212	8.85232
a. Predictors: (Constant), att2, mot2, mot1, att1				
b. Dependent Variable: achievement				

فى هذا الجزء يصف النموذج ككل حيث اعطى معامل الارتباط المتعدد = 0.480 R هى معامل الارتباط المتعدد بين المتغير التابع والمتغيرات المنبئة معًا، ومربع معامل الارتباط المتعدد $R^2 = 0.230$ ، ومربع معامل الارتباط المصحح $Adjusted\ R\ Square = 0.212$.

وان المتغيرات الاربعة المستقلة فسرت 23% من تباين التحصيل ويتم حساب R^2 :

$$R^2 = \frac{SS_{Reg}}{SS_{total}}$$

انظر فى جدول ANOVA نلاحظ ان:

$$= \frac{3981.055}{17302.857} = 0.230$$

وتعتبر R^2_{adj} المصححة أكثر دقة من R^2 لأنها تصحح من صغر حجم العينة و درجات الحرية المفقودة اثناء الانحدار ويتم حسابها من خلال تصحيح stein و احيانًا تستخدم صيغة Wherry فى بعض البرامج :

$$\text{Adjusted } R^2 = 1 - \left[\left(\frac{175-1}{175-4-1} \right) \left(\frac{175-1}{175-4-2} \right) \left(\frac{175+1}{175} \right) \right] (1 - 0.48) = 0.212$$

وهذه القيمة مشابهة تقريبًا لقيمة $R^2 (0.230)$ مما يدل على أن الصدق التعميمى للنموذج جيد. واعطى قيمة الخطأ المعيارى لقيمة Std.Error of the Estimate وهى تدل على الدقة العامة لنموذج الانحدار وهى الجذر التربيعى لمتوسط مجموع مربعات البواقي فى جدول ANOVA:

$$= \sqrt{MS_R} = \sqrt{78.364}$$

● إحصائيات التغير Change Statistics كالآتى:

Model Summary ^b					
Change Statistics					Durbin-Watson
R Square Change	F Change	df1	df2	Sig. F Change	
.230	12.701	4	170	.000	1.208

وهى ماذا يحدث لو اضيف متغير جديد إلى معادلة الانحدار هذا الجزء غير ذات اهمية فى طريقة ENTER لان كل المتغيرات ادخلت إلى المعادلة حيث إن إحصائيات هذا الجزء هى قيم R^2 و F و $df2$ و $df1$ و Sig وهى مثل قيمة R^2 وجدول تحليل التباين ولا تتغير قيمتها ويمكن أن يحدث تغير إذا استخدم الباحث طرق الانحدار الإحصائى والهرمى. وفى نهاية الجدول يوجد إحصاء Durbin-Watson وهو يخبرنا ما إذا كانت مسلمة استقلال البواقي أو الأخطاء تحققت ام لا. والقاعدة الشديدة التحفظ تقول أن القيم الأقل من الواحد الصحيح 1.0 أو الأكبر من 3.0 تشير إلى وجود استقلالية إلى حدًا ما بينما القيمة القريبة من 2 هى افضل وتشير إلى تحقق المسلمة بدرجة كبيرة، وفى المثال هى 1.208 وعليه فالقيمة ليست قريبة من 2

ولا هي أقل من 1.0 وبالتالي يمكن القول إلى توافر المسلمة بدرجة متوسطة وليست جيدة.

• جدول ANOVA كالاتى:

ANOVA ^a						
Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	3981.055	4	995.264	12.701	.000 ^b
	Residual	13321.802	170	78.364		
	Total	17302.857	174			

a. Dependent Variable: achievement

b. Predictors: (Constant), att2, mot2, mot1, att1

وهو يتضمن معلومات عن اختبارات فروض لنموذج الانحدار أو التنبؤ ككل بمعنى أن المتغيرات الاربعة تصلح أو لا تصلح للتنبؤ بالتحصيل. وهذا الجدول يتضمن صف Regression وصف البواقي أو الأخطاء Residual وهو يتضمن معلومات عن الأخطاء وهى تمثل الفرق الكلى بين النموذج Y والبيانات المقاسة (Y) ويتضمن نسبة F وهى تعرض نسبة التحسن فى التنبؤ نتيجة لمطابقة النموذج بالنسبة إلى عدم الدقة (الأخطاء) أو البواقي الموجودة فى النموذج وتقدر كالاتى:

$$F = \frac{SS_{\text{regression}} / df_{\text{reg}}}{SS_{\text{Residual}} / df_{\text{residual}}} = \frac{MS_{\text{Reg}}}{MS_{\text{Res}}} = \frac{3981055/4}{13321.802/17} = \frac{995.264}{78.364} = 12.701$$

أو تقدر من الصيغة الآتية :

$$F = \frac{R^2 / K}{(1 - R^2)(N - K - 1)} = \frac{0.23/4}{(1 - 0.23)(175 - 4 - 1)} = 12.701$$

وما يهمنا فى الجدول هو التحقق من الفروض الإحصائية الآتية:

$$H_0 : R \text{ or } \rho = 0$$

$$H_A : R \text{ or } \rho \neq 0$$

حيث ρ معامل الارتباط المتعدد بين التحصيل والمتغيرات المستقلة فى المجتمع وبما أن $0.05 < \text{Sig} (0.000)$ وعليه نقبل الفرض البديل القائل يمكن للمتغيرات

الاربعة المستقلة أن تستخدم للتنبؤ بالتحصيل بمعنى يمكن بناء معادلة انحدار. لاحظ أن دلالة F لاختبار نموذج الانحدار ككل، وعليه $F(4,170) = 12.701$ ، $P < 0.05$ ، لاحظ أن المتغيرات الاربعة المستقلة ليس شرط أن تستخدم كلها فى بناء معادلة الانحدار ويمكن القول أن الفرض البديل على الأقل أحد المتغيرات المستقلة يمكن أن تتنبأ تنبؤ دال إحصائيًا بالتحصيل بفرض توافر مسلمات الانحدار.

• معالم النموذج أو تفاصيل النموذج أو معادلات الانحدار:

Model	Unstandardized Coefficients		Standardized Coefficients	t	Sig.
	B	Std. Error	Beta		
1 (Constant)	-17.608	6.182		-2.848	.005
mot1	.760	.185	.338	4.114	.000
mot2	-.011	.206	-.005	-.055	.956
att1	.256	.130	.184	1.979	.049
att2	.123	.106	.114	1.161	.247

a. Dependent Variable: achievement

كما سبق أن تناولنا ما إذا كان النموذج ككل له قدرة على التنبؤ بالمتغير التابع. والجزء الآتى يهتم بمعالم ومعاملات الانحدار للنموذج ويتضمن معلومات ضرورية لبناء معادلة الانحدار وتحديد أى المتغيرات المستقلة أكثر اهمية أو اسهامًا نسبيًا وكذلك مؤشر VIF لتشخيص التلازمة الخطية المتعددة.

ففى جزء المعاملات غير المعيارية Unstandardized: العمود B وهو عبارة عن الثابت a ومعاملات الانحدار غير المعيارية للمتغيرات المستقلة (ميل خط الانحدار) . وعليه فإن معادلة الانحدار غير المعيارية هي:

$$\hat{Y} = -17.608 + 0.760X_1 - 0.011X_2 + 0.256X_3 + 0.106X_4$$

وبصيغة اخرى:

$$\text{Achievement} = -17.608 + 0.76 \text{ mot}_1 - 0.011 \text{ mot}_2 + 0.256 \text{ att}_1 + 0.106 \text{ att}_2$$

وهذه الصيغة تستخدم للتنبؤ بتحصيل أى طالب من درجاته على المتغيرات المستقلة. ولاحظ أن قيم b تشير إلى اسهام المتغير فى التنبؤ بالتحصيل بالنموذج وأيضًا تخبرنا

عن العلاقة بين التحصيل وكل متغير منبئ وإذا كانت موجبة فيوجد ارتباط موجب بين المنبئ أو التحصيل إذا كانت سالبة فالارتباط بين المنبئ والتحصيل سالب. وعليه إذا زادت الدافعية الداخلية يزيد فالتحصيل $b=0.76$ وهذه القيمة تشير إلى إذا زادت الدافعية الداخلية بمقدار وحدة واحدة فالتحصيل يزيد بمقدار 0.76 وحدة وكلا من المتغيرين يقاسوا بدرجات خام فزيادة درجة واحدة في الدافعية الداخلية تسهم في زيادة 0.76 من الدرجة في التحصيل.

حقيقة المقارنة بين الأهمية النسبية للمتغيرات في ضوء b غير صحيح لاختلاف وحدات القياس وعلية فلا بد من الاعتماد على معامل الانحدار المعياري Beta ويفسر مثل معامل ارتباط بيرسون r والقيمة المطلقة ل Beta تستخدم لتحديد أي المتغيرات المستقلة أكثر اسهاما في التنبؤ. والمدقق يلاحظ أن معامل الانحدار المعياري Beta للدافعية الداخلية 0.338 بينما للدافعية الخارجية -0.005 على ذلك فإن تأثير الدافعية الداخلية على التحصيل أكثر من تأثير الدافعية الخارجية، حيث زيادة وحدة انحراف معياري في الدافعية الداخلية يسبب زيادة مقدارها 0.338 وحدة انحراف معياري في التحصيل. وعلى ذلك فالأهمية النسبية للمتغيرات الأربعة بالترتيب هي الدافعية الداخلية والاتجاه نحو المادة الدراسية والاتجاه نحو المدرسة وأخيراً الدافعية الخارجية على الترتيب.

وكذلك لمعامل الانحدار المعياري دلالة إحصائية حيث يختبر باستخدام اختبار T حيث الفروض الإحصائية:

$$H_0 : \text{Beta} = 0$$

$$H_A : \text{Beta} \neq 0$$

ويتضح أن معامل Beta ل Mot_1 دال إحصائياً حيث $T = 4.114, P < 0.05$ وكذلك معامل Beta ل Att_1 دال إحصائياً بينما معاملات Beta للدافعية الخارجية والاتجاه نحو المدرسة غير دالة إحصائياً وعلية يتضح أن معادلة الانحدار المعيارية: Z

$$\text{achievement} = 0.338 Z_{\text{mot}_1} + 0.184 Z_{\text{att}_1}$$

والجدول الآتي:

95.0% Confidence Interval for B		Correlations			Collinearity Statistics	
Lower Bound	Upper Bound	Zero-order	Partial	Part	Tolerance	VIF
-29.811-	-5.405-					
.396	1.125	.398	.301	.277	.669	1.494
-.418-	.395	.226	-.004-	-.004-	.681	1.468
.001	.512	.308	.150	.133	.524	1.908
-.086-	.332	.348	.089	.078	.469	2.130

عبارة عن معاملات ارتباط بيرسون وكذلك معاملات الارتباط الجزئية وشبه الجزئية لاحظ أن مربع معامل الارتباط شبه الجزئي Part يعطى مقدار التباين الخاص فى المتغير التابع المفسر عن طريق كل متغير مستقل ومعامل الارتباط الجزئي Partial هو معامل الارتباط بين المتغير التابع والمتغير المستقل بعد ضبط تأثير المتغيرات المستقلة الاخرى، فمثلاً معامل الارتباط الجزئي بين mot_1 والتحصيل 0.301 بعد استبعاد تأثير المتغيرات الثلاثة الباقية على كلاً من mot_1 والتحصيل بينما معامل الارتباط شبه الجزئي هو معامل الارتباط بين المتغير المنبئ والتحصيل بعد ضبط تأثير المتغيرات الثلاثة الباقية على التحصيل فقط. فى حالة استخدام الانحدار التدريجي يتم إدخال اول متغير الاكثر معامل ارتباط مع المتغير التابع ثم بعد ذلك فى ضوء أعلى معامل ارتباط شبه الجزئي.

● **التحقق من مسلمة التلازمية الخطية المتعددة:** كما نعلم أن تشخيص هذه المسلمة من خلال مؤشرين هما Tolerance وإذا كانت قيمته أقل من 0.10 ومؤشر VIP إذا كانت قيمته أكبر من 10.0 تشير إلى وجود التلازمية الخطية (Cohen et al., 2003)، وقيمة VIP الأكبر من 10 تعنى أن الأخطاء المعيارية لـ b أكبر ثلاث مرات مما هو إذا كانت المتغيرات المستقلة غير مرتبطة.

قدم (Field (2009) إرشادات لـ VIP و Tolerance وهى كالاتى:

- $VIF > 10$ وجود التلازمية الخطية وتسبب مشاكل وتوضع فى الاعتبار.
- إذا كان متوسط VIF للمنبئات أكبر من 1.00 فإن نتائج الانحدار تكون متحيزة.
- $Tolerance > 0.10$ تشير إلى وجود التلازمية الخطية وتسبب مشاكل خطيرة.
- $Tolerance > 0.20$ تشير إلى مشاكل محتملة.

وللمثال السابق نجد أن قيم Tolerance مناسبة حيث إنها كانت أكبر من 0.10 أو 0.20 بالتالى لا توجد تلازمية خطية. اما قيم VIP فكانت قيمة أقل من 10 وعليه لا توجد تلازمية خطية متعددة ويمكن حساب متوسط VIP:

$$VIF = \frac{1.494 + 1.468 + 1.908 + 2.130}{4} = \frac{7}{4} = 1.75$$

ولكن متوسط VIP زاد عن الـ واحد الصحيح بما يدل على وجود تحيز فى نتائج الانحدار .

● اعطى جدول تشخيص التلازمية الخطية مرة اخرى كالآتى:

Collinearity Diagnostics ^a								
Model	Dimension	Eigenvalue	Condition Index	Variance Proportions				
				(Constant)	mot1	mot2	att1	att2
1	1	4.931	1.000	.00	.00	.00	.00	.00
	2	.036	11.763	.01	.13	.06	.13	.13
	3	.016	17.526	.26	.34	.02	.08	.28
	4	.011	21.654	.00	.52	.36	.33	.34
	5	.007	27.343	.73	.01	.57	.46	.24

a. Dependent Variable: achievement

وهى مؤشرات اضافية لتشخيص التلازمية الخطية مثل:

- القيم الكامنة ومؤشر الشرطية Condition Index ويقدر من القيم الكامنة. والاراء فيما يخص الحدود لتحديد أو لتشخيص التلازمية الخطية المتعددة باستخدام مؤشر الشرطية هو أن القيم اعلى من 15 أو 30 تشير إلى وجود التلازمية الخطية. وعليه لم تزيد اى من القيم عن 30.

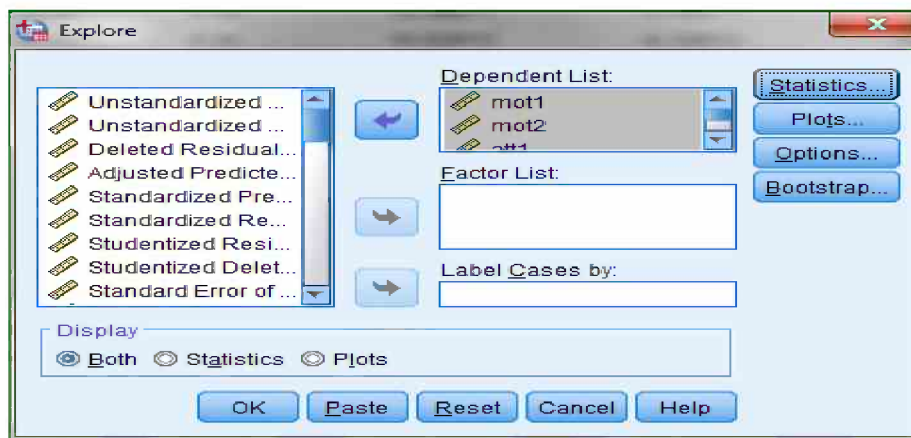
● **Casewise diagnostic** اعطى البرنامج ملخص لإحصائيات البواقى وفيما يلى جزء من المخرج للخمسة عشر حالة الأولى :

Casewise Diagnostics ^a				
Case Number	Std. Residual	achievement	Predicted Value	Residual
1	.416	19.00	15.3148	3.68517
2	.917	28.00	19.8781	8.12190
3	.179	25.00	23.4116	1.58845
4	-.789-	8.00	14.9808	-6.98077-
5	-.466-	6.50	10.6281	-4.12813-
6	.899	22.00	14.0408	7.95925
7	-.359-	7.50	10.6817	-3.18173-
8	.688	27.00	20.9122	6.08780
9	-.172-	12.00	13.5218	-1.52180-
10	.795	22.50	15.4668	7.03323
11	-1.703-	1.50	16.5781	-15.07805-
12	1.226	32.50	21.6435	10.85647
13	.277	16.00	13.5488	2.45115
14	-.786-	13.50	20.4574	-6.95742-
15	-.927-	7.00	15.2068	-8.20681-

ويستخدم لفحص القيم أو الحالات المتطرفة حيث اى حالات لها بواقى معيارية أقل من 2- أو أكبر من 2 تعتبر قيم متطرفة بالنظر إلى قيم البواقى المعيارية نلاحظ أن 100% من القيم أو الحالات تقع فى المدى من 2+.

حقيقة تشخيص الحالات المتطرفة يتم من خلال الإحصاء الوصفى من خلال الآتى :

1.اضغطAnalysis→Descriptivestatistics→Explore



2. انقل المتغيرات المستقلة والتابعة إلى مربع Dependentlist

3. اضغط Continue ثم ضغط OK

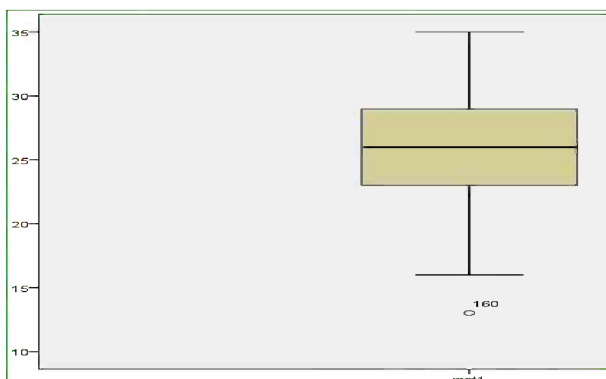
المخرج: تشخيص الحالات المتطرفة من

خل الشكل Boxplot الآتى:

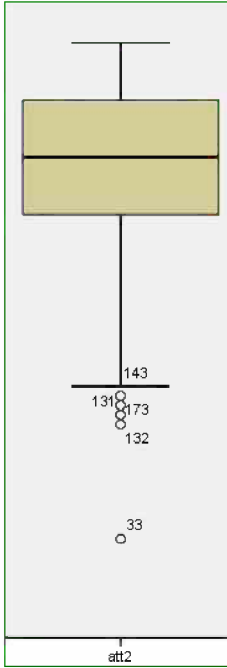
فمثلا ل mot₁ أوضح أن الحالة 160هى

قيمة متطرفة وهى تتاظر القيمة 13 وتم

تأكيد ذلك من شكل Stemleaf حيث



اوضح أن القيم أقل أو تساوى 13 تعتبر قيم متطرفة:



Frequency	Stem &	Leaf
1.00	Extremes	(=<13.0)
2.00	16 .	00
6.00	17 .	000000
4.00	18 .	0000
4.00	19 .	0000
8.00	20 .	00000000
10.00	21 .	0000000000
7.00	22 .	0000000
15.00	23 .	00000000000000
13.00	24 .	00000000000000
14.00	25 .	00000000000000
11.00	26 .	000000000000
18.00	27 .	0000000000000000
14.00	28 .	00000000000000
14.00	29 .	00000000000000
11.00	30 .	000000000000
7.00	31 .	00000000
8.00	32 .	00000000
4.00	33 .	0000
3.00	34 .	000
1.00	35 .	0

ولكن يبدو أن المتغير المستقل att2 به حالات كثيرة متطرفة:
وهذه الحالات يجب استبعادها من التحليل.

● **إحصائيات التأثير Influence:** لتحديد أى الملاحظات أو الحالات أكثر تأثيرًا فى

تقديرات الانحدار مقارنة بالحالات الأخرى فى قاعدة البيانات ولتحديد ذلك توجد عدة إحصائيات هى:

1. مؤشر : Leverage : وهو دالة لدرجات المتغيرات المستقلة وهو مقياس للقيم المتطرفة للمتغيرات المستقلة معًا وتقدر كالاتى:

$$h_i = \frac{1}{N} + \frac{(X - \bar{X})^2}{\sum X^2}$$

N حجم العينة ، X الدرجة ، $\bar{X}$ متوسط المتغير ، $\sum X^2$ مجموع مربعات X

ويعتبر (Pedhazure 1997) الحالة ذات تأثير عالى إذا كان:

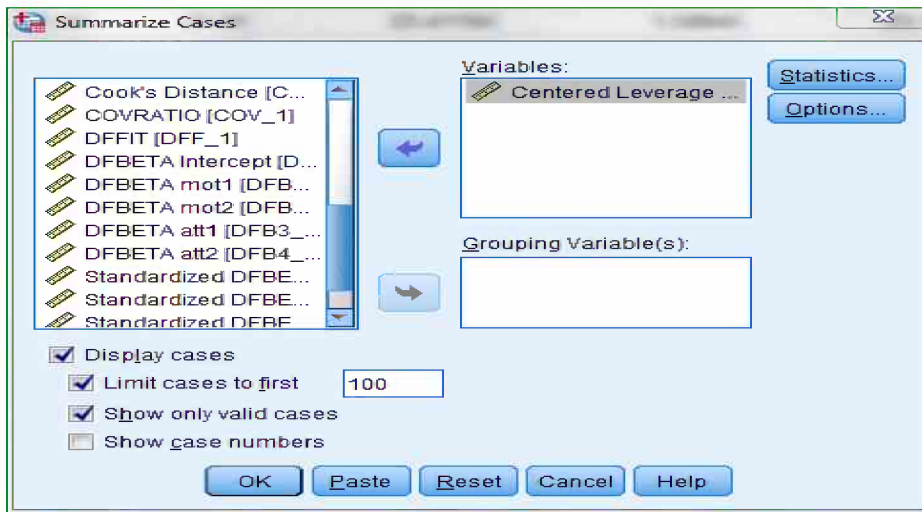
$$\begin{aligned} h_i &> 2(K+1)/N \\ &> 2(4+1)/175 \\ &> 0.06 \end{aligned}$$

وتسهم الحالة فى معادلة الانحدار بتأثير متوسط إذا كان:

$$\begin{aligned} &= (K+1)/N \\ &= \frac{5}{175} = 0.03 \end{aligned}$$

واقصى قيمة لـ h_i الواحد الصحيح.

وبإجراء امر **Case Summaries** لمؤشر **Leverge** (موجود في ملف البيانات) كالاتي:
اضغط **Analyze→Reports→Case Summarize** يعطى الشاشة الآتية:



2. انقل متغير **CenteredLeverage** إلى مربع **Variables**

3. اضغط **OK**

المخرج: فيما يلي عرض لبعض الحالات:

Case Summaries	
	Centered Leverage Value
1	.01717
2	.00247
3	.01315
4	.01519
5	.03491
6	.02712
7	.01396
8	.02206
9	.01625
10	.01113

23	.03220
24	.00668
25	.01728
26	.03291
27	.06403

ش

نلاحظ أن الحالة 5 لها تأثير متوسط 0.0349 وكذلك الحالة 3 ، اما الحالة 27 لها تأثير عالى في معادلة الانحدار 0.064

مؤشر Cook D

القيمة الصغيرة نسبياً لهذا المؤشر تشير إلى أن

للحالة تأثير كبير في معادلة الانحدار. وبإجراء الامر السابق لقيم مؤشر **Cook Distance** في ملف البيانات:

Analyze → reports→Summarize Cases

المخرج:

33	.00000
34	.00622
35	.00242
36	.00517
37	.00449
38	.00088
39	.00009
40	.00399
41	.00102
42	.00377
43	.00376

وإذا تأملنا إلى قيمة مؤشر D للحالة 33 أن لها $D = 0.00$ ، لاحظ أن القيمة العالية لـ h_i والمنخفضة لـ D تدل على أن الحالة أو الفرد له تأثير في معادلة الانحدار. فالحالات 28 , 33 , 39 , 128 تعتبر أكثر الحالات تأثيرًا.

لاحظ ما افرضه مؤشر h_i leverage لا يتفق مع مؤشر Cook D

إحصاء DfBeta

Case Summaries	
	DFBETA mot1
1	-.00816
2	.00476
3	.00346
4	.00669
5	.00843

وإذا كانت قيمته تقع في المدى من ± 1 فإن الحالات ليس لها تأثير في معادلة الانحدار. ففي ضوء هذا المؤشر لا توجد أي حالة لها تأثير على معالم معادلة الانحدار وذلك لمتغير $motx1$ ، وهكذا للمتغيرات المستقلة الأخرى.

حيث يتم حساب معادلة الانحدار بعد استبعاد كل حالة على حدة وملاحظة التغير في معالم المعادلة.

مؤشر Mahalanobis distance

وهذا المؤشر مثل χ^2 حيث القيمة الحرجة عند $\alpha = 0.01$ و $df=4$ هي 18.46 وعليه إذا قيمة إحصاء Mahalanobis أكبر من هذه القيمة تدل على قيمة متطرفة متدرجة وبإجراء امر

امر **Case Summaries** لقيم هذا المؤشر (موجود في ملف البيانات):

Reports→ Case Summaries

المخرج

28	3.56771
29	7.55685
30	3.21367
31	4.85081
32	8.03547
33	18.87951
34	2.45916
35	2.13726
36	4.23953

عرض بعض الحالات لاحظ أن البرنامج اعطى قيمة المؤشر للحالات الـ (175)، يتضح من المخرج أن قيمة المؤشر لكل الحالات أقل من 18.46 ماعدا الحالة 33 فقيمتها زادت عن هذه القيمة 18.879 وعلية يجب استبعادها من التحليل لاحظ أن مؤشر Mahalanobis يشخص القمة المتطرفة المتدرجة

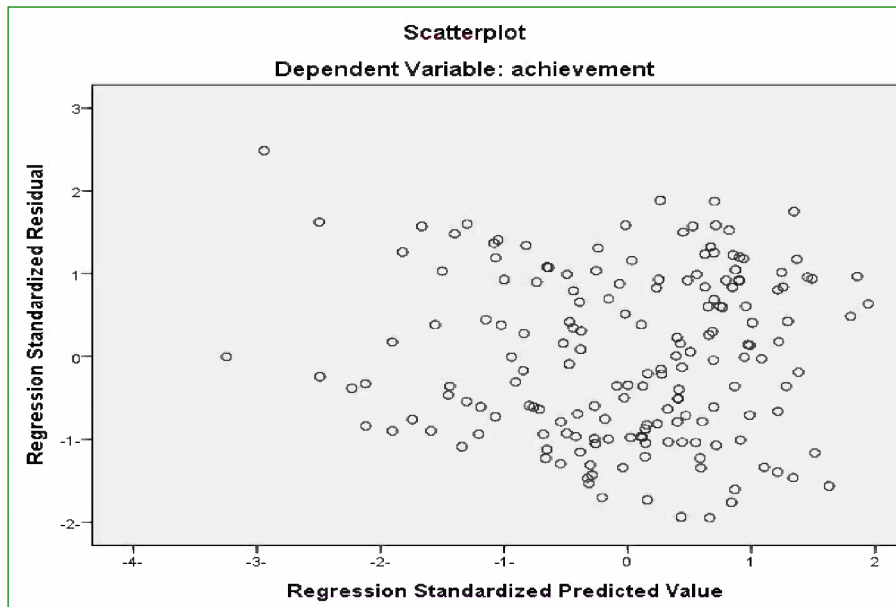
للحالة فى علاقتها مع كل الحالات.
واعطى البرنامج ملخص لمعظم الإحصائيات كالآتى:

Residuals Statistics ^a					
	Minimum	Maximum	Mean	Std. Deviation	N
Predicted Value	2.0333	26.8750	17.5714	4.78327	175
Std. Predicted Value	-3.248-	1.945	.000	1.000	175
Standard Error of Predicted Value	.736	2.992	1.444	.394	175
Adjusted Predicted Value	1.4911	26.6478	17.5608	4.83460	175
Residual	-17.23361-	22.00326	.00000	8.74998	175
Std. Residual	-1.947-	2.486	.000	.988	175
Stud. Residual	-2.007-	2.596	.001	1.004	175
Deleted Residual	-18.30855-	24.00890	.01062	9.02799	175
Stud. Deleted Residual	-2.025-	2.642	.001	1.007	175
Mahal. Distance	.209	18.880	3.977	2.863	175
Cook's Distance	.000	.123	.006	.012	175
Centered Leverage Value	.001	.109	.023	.016	175

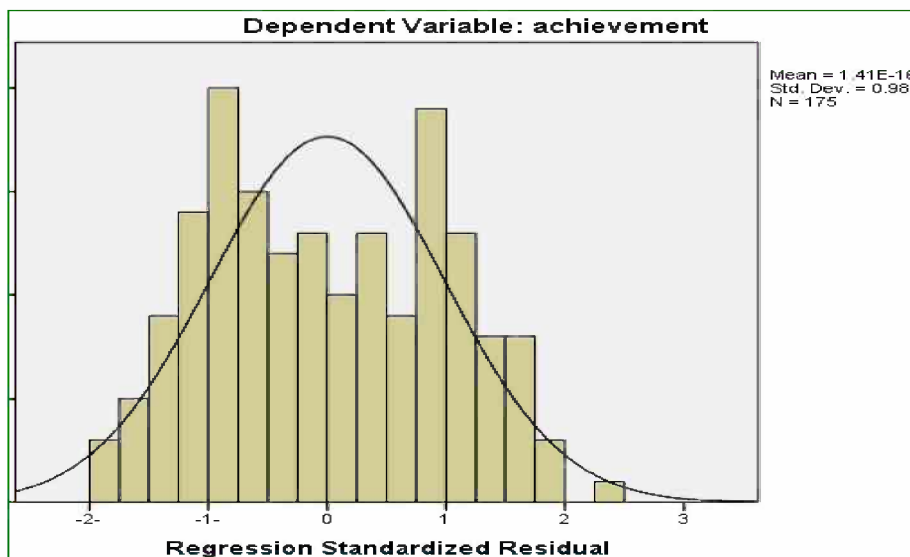
a. Dependent Variable: achievement

التمثيل البياني للبواقي

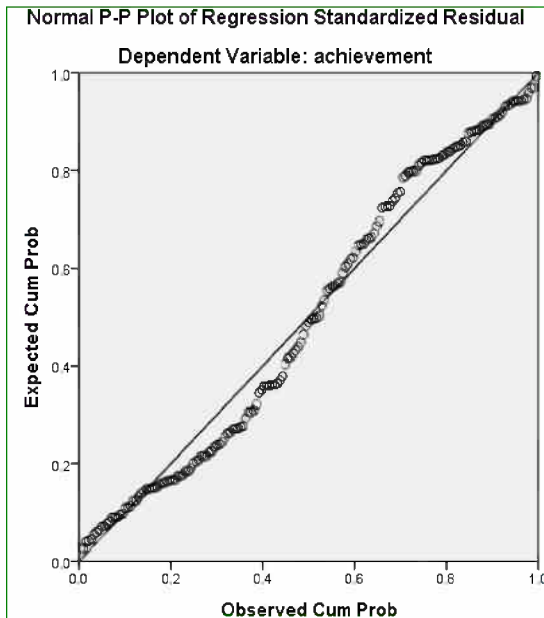
المرحلة الأخيرة من التحليل فحص مسلمات النموذج وسبق وان عرضنا كثيرًا منها . وفى امر الانحدار يوجد الرسم البيانى لـ ZRESID و ZPRED وكذلك المدرج و Normal Probability للبواقي. وإذا كان توزيع البواقي يشبه القمع فإنه لا تتوفر مسلمة تجانس تباينات البواقي homoscedasticity فى البيانات والشكل الآتى يوضح الشكل البيانى للبواقي المعيارية فى مقابل القيم المعيارية المتنبأ بها:



لاحظ التوزيع العشوائى للنقاط وانتشارها فى الرسم البيانى وهذا يشير إلى توفر مسلمة تجانس تباينات البواقي إلى حدًا ما. ولكن إذا كان الشكل يشبه القمع فيدل على عدم توافر مسلمة التجانس للبواقي حقيقة أن الخطية لم تتوفر بدرجة تامة. اما اعتدالية البواقي تم الحصول عليها من خلال عرض Normal probability plot, Histogram وكان كالاتى:



ويبدو أن المنحنى ليس اعتدالى بدرجة كبيرة بل يوجد به نوعًا من التفرطح حيث $kurtosis = -1.01$ ولكنه فى المجمل العام يتسم بالاعتدالية. كما عرض p.plot كالاتى:



حيث الخط المستقيم يعكس التوزيع الاعتدالي والنقاط حوله تعكس البواقي المعيارية فاذا كانت النقاط تقع تماماً على الخط المستقيم فإن التوزيع تام الاعتدالية ولكن يتضح أن النقاط انتشرت قريباً على جانبي الخط المستقيم وهذا يدل على أن الاعتدالية ليست تامة وهذا يعطى شك على مدى موثوقية نتائج معالم نموذج الانحدار. ولكن القول الفاصل في التحقق من الاعتدالية

من عدمها إجراء اختبار كولوموجوروف -سميرنوف للبواقي المعيارية وكانت النتائج كالآتي:

Tests of Normality						
	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
Standardized Residual	.086	175	.003	.970	175	.001

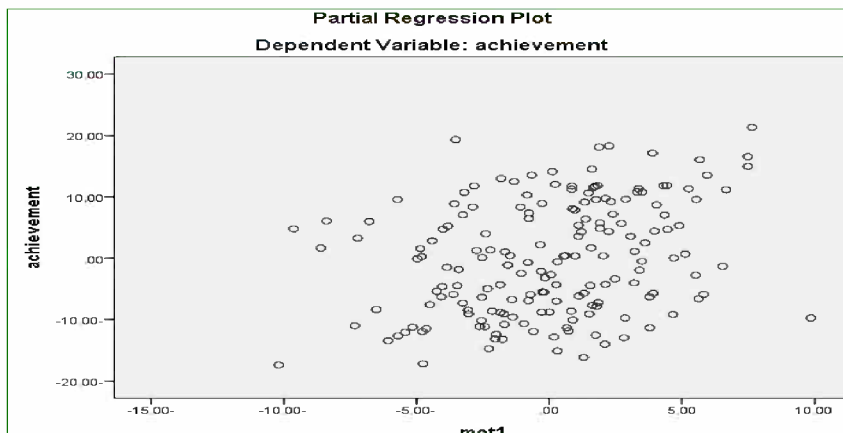
a. Lilliefors Significance Correction

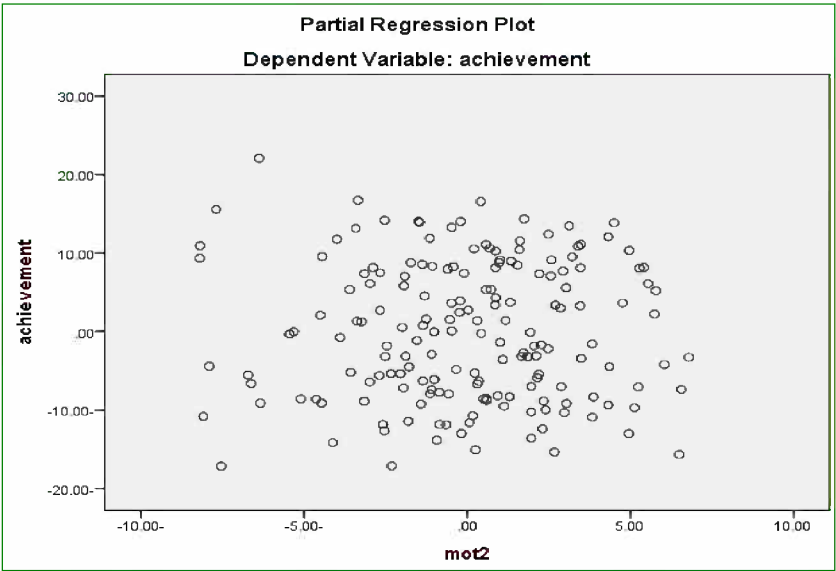
حيث $K.S(175) = 0.086$, $P=0.003$,وعليه فإن توجد دلالة إحصائية وعليه فإن توزيع البواقي المعيارية غير اعتدالية التوزيع، وعليه يجب الحذر في تعميم نتائج تحليل

الانحدار ولا بد من إجراء صدق النتائج على عينات أخرى Cross- validation.

الاشكال البيانية للبواقي للمتغير التابع وكل متغير منبئ على حدة:

ويمكن تشخيص العلاقات غير الخطية وعدم تساوى التباين البواقي من هذه الرسومات وهى كالآتي:





فالعلاقة بين mot_1 والتحصيل تبدو خطية، غير أن العلاقة بين mot_2 والتحصيل بها غير خطية والحال نفسة في العلاقة بين التحصيل وكلا من att_1 , att_2 .

مخرج الانحدار باستخدام Stepwise:

كل خطوات المخرج السابق ما عدا تغيير الطريقة Stepwise –Method:

Variables Entered/Removed ^a			
Model	Variables Entered	Variables Removed	Method
1	mot1		Stepwise (Criteria: Probability-of-F-to-enter <= .050, Probability-of-F-to-remove >= .100).
2	att1		Stepwise (Criteria: Probability-of-F-to-enter <= .050, Probability-of-F-to-remove >= .100).

a. Dependent Variable: achievement

حيث اعطى معادلة التنبؤ في ضوء متغيرين فقط هما mot_1 و att_1 وجدول ANOVA كالآتي:

ANOVA

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	2747.680	1	2747.680	32.658	.000 ^b
	Residual	14555.177	173	84.134		
	Total	17302.857	174			
2	Regression	3873.685	2	1936.843	24.807	.000 ^c
	Residual	13429.172	172	78.077		
	Total	17302.857	174			

a. Dependent Variable: achievement

b. Predictors: (Constant), mot1

c. Predictors: (Constant), mot1, att1

Model Summary ^c				
Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	.398 ^a	.159	.154	9.17246
2	.473 ^b	.224	.215	8.83610

a. Predictors: (Constant), mot1
b. Predictors: (Constant), mot1, att1
c. Dependent Variable: achievement

حيث إن اسهام المتغير الأول mot₁ هو $R^2 = 0.159$, Adjusted $R^2 = 0.154$ بينما اسهام المتغيرين معًا mot₁ و att₁ هو $R^2 = .224$ و Adjusted $R^2 = .215$ والملاحظ في طريقة Enter أن قيمتي $R^2 = .230$ و $R^2 = .221$ Adjusted أكبر من قيمتهما في طريقة Stepwise لان في طريقة Enter تم اضافته اسهام att₂ , mot₂ إلى التباين المفسر R^2 بينما في Stepwise تم استبعادهما.

ومعالم النموذج كالاتى:

Coefficients ^a					
Model		Unstandardized Coefficients		Standardized Coefficients	Sig.
		B	Std. Error	Beta	
1	(Constant)	-5.288	4.060		.194
	mot1	.895	.157	.398	.000
2	(Constant)	-17.871	5.126		.001
	mot1	.815	.152	.363	.000
	att1	.359	.095	.258	.000

a. Dependent Variable: achievement

القوة الإحصائية للانحدار المتعدد باستخدام برنامج G-power
لحساب القوة الإحصائية باستخدام برنامج G-Power اتبع الاتى:

1. افتح البرنامج تظهر الشاشة الافتتاحية.

2. تحت Type of power analysis اختيار:

Type of power analysis
Post hoc: Compute achieved power - given α , sample size, and effect size

3. تحت Test family اختيار F- tests

4. تحت Statistical test اختيار:

Test family	Statistical test
F tests	Linear multiple regression: Fixed model, R^2 deviation from zero

5. ادخل المعالم الآتية تحت Input parameters:

- حجم التأثير f^2 كالآتي:

$$f^2 = \frac{R^2}{1 - R^2} = \frac{0.230}{1 - 0.230} = \frac{0.723}{0.77} = 0.939$$

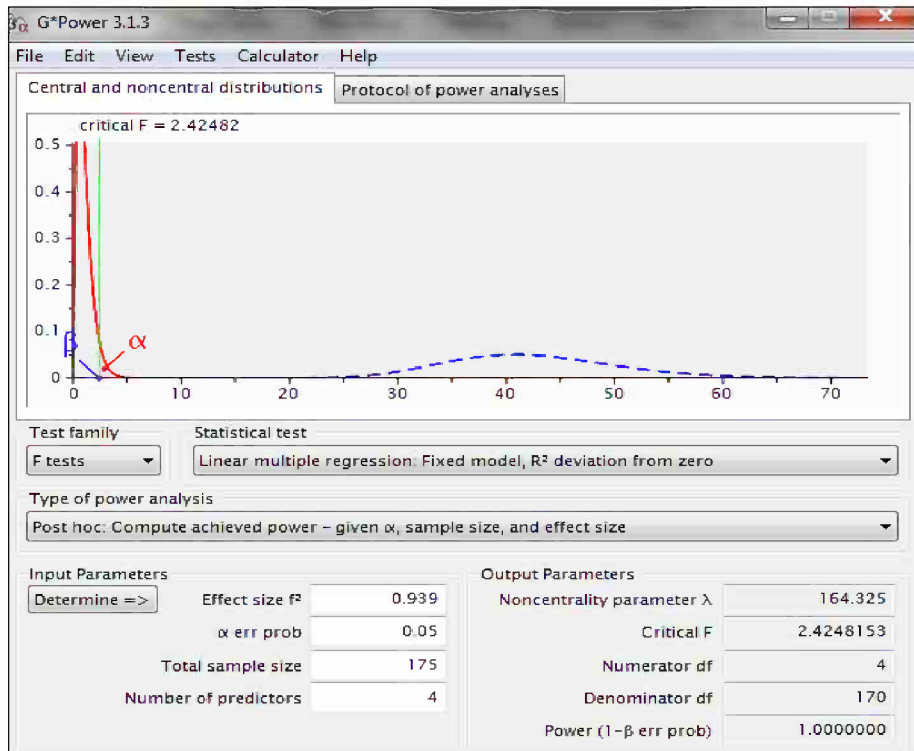
- مستوى الدلالة الإحصائية: 0.05

- حجم العينة الإجمالي = 175

- عدد المتنبئات = 4

Input Parameters	
Determine =>	Effect size f^2 0.939
	α err prob 0.05
	Total sample size 175
	Number of predictors 4

6. اضغط Calculated تظهر المخرجات الآتية:



القوة الإحصائية=1.000 وهذا مستوى قوة تام وربما يرجع هذا إلى أن حجم العينة كبيرًا كبرًا كافيًا.

كتابة نتائج الانحدار المتعدد في تقرير البحث وفقا لـ APA

بعد التحقق من مسلمتي التلازمية الخطية المتعددة وتجانس التباين للبواقي والخطية للبواقي ولم تتحقق مسلمة اعتدالية البواقي واتضح أن نتائج تحليل الانحدار للمتغيرات المستقلة الأربعة الدافعية الداخلية والخارجية والاتجاه نحو المادة الدراسية ونحو المدرسة حيث: $F(4,170) = 12.701$, $\text{Adjusted } R^2 = 0.212$, $P < 0.05$

فيما يلي معالم النموذج:

Model	b	SE.b	Beta
a	-17.60	6.182	8
mot ₁	0.760	0.185	0.338*
mot ₂	-0.011	0.206	-0.005
att ₁	0.256	0.130	0.184**
att ₂	0.123	0.106	0.114

*P < 0.01 , **P < 0.05

الفصل الثامن

التحليل العاملي الاستكشافي

Exploratory factor analysis

التحليل العاملي هو أسلوب يستخدم لتحديد العوامل التي تفسر إحصائيًا التباينات والتغايرات بين مجموعة من القياسات (المفردات)، وعمومًا عدد العوامل أصغر من عدد القياسات أو المفردات وبالتالي هو أسلوب لتقليل أو خفض البيانات - Data Reduction technique حيث يقلل عدد كبير من القياسات المتداخلة (المرتبطة) إلى مجموعة صغيرة من العوامل، فلو أن دراسة تضمنت مجموعة كبيرة من القياسات (مفردات) تعكس أبعاد مختلفة لمفهوم عام فإن التحليل العاملي يعطي عوامل تمثل هذه الأبعاد. لذلك فإن العوامل تناظر ابنيه Constructs (متغيرات كامنة غير ملاحظه) للنظرية التي تساعدنا على فهم سلوك ظاهرة ما. وأمثلة لهذه الابنية كمفهوم الذات وتقدير الذات والدافعية والابتكارية وغيرها.

وعلى ذلك فإن التحليل العاملي أسلوب لتحديد مجموعات أو تجمعات من المتغيرات.

استخدامات التحليل العاملي

يستخدم التحليل العاملي في الآتي:

1. فهم طبيعة البناء لمجموعة من المتغيرات و تحديد أبعادها.
2. تقليل عدد المتغيرات إلى عدد أقل من العوامل بمعنى تحديد أبعاد مقياس موجود مما يساعدنا في استخدامها في تحليلات لاحقة مثل تحليل التباين و الانحدار و التحليل التمييزي.
3. تحديد الطبيعة الأحادية البعد Unidimensionality لبناء نظري مفترض
4. تقييم الصدق البنائي للمقياس أو الاختبار .
5. يستخدم لبناء أبنية نظرية.
6. أحيانًا يستخدم لإثبات أو دحض بعض النظريات المفترضة.

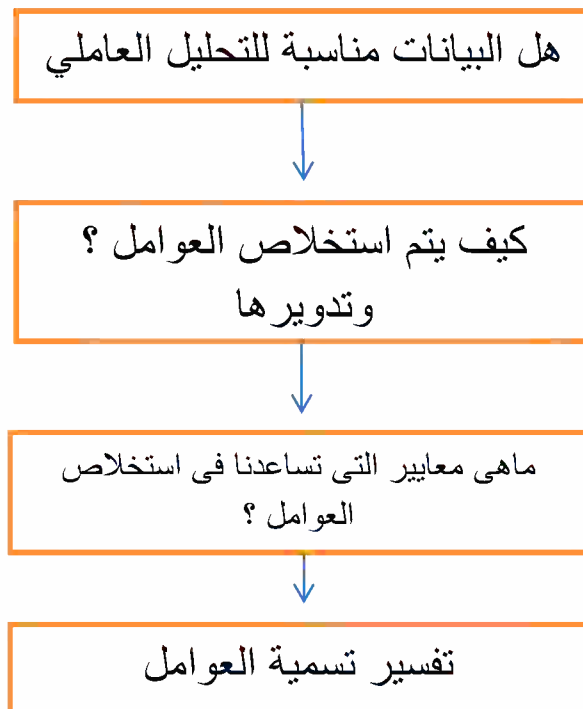
7. حل لإشكالية التلازمية الخطية Collinearity (معاملات الارتباط المرتفعة بين المتغيرات) حيث هذه المتغيرات ذات العلاقة المرتفعة يمكن التعبير عنها في ضوء عامل.

8. تحديد مؤشرات الابنية.

والشئ المهم يجب التأكيد عليه هو ارتباط التحليل العاملى ببناء ادوات القياس و التأكد من مصداقية بنائها ولذلك يعتبر هو قلب و روح قياس المفاهيم النفسية.

مراحل إجراء التحليل العاملى

لإجراء التحليل العاملى نتبع الخطوات الخمسة الآتية:



أولاً: متطلبات التحليل العاملى (هل البيانات مناسبة لتحليل): وفى هذا الشأن تظهر عدة قضايا أهمها:

• **حجم العينة:** توجد آراء عديدة فى هذا الشأن فيشير (Comry & lee, 1992) أن المتفق عليه أن حجم العينة 100 يكون ضعيف و 200 مناسب و 300 جيد و 500 جيد جداً و 1000 فأكثر ممتاز. حقيقة أن تحديد حجم العينة يتوقف على عوامل عديدة منها التشبعات بالعامل حيث إذا زادت عن 0.6 وتم تمثيل العامل بمفردات عديدة (ثلاثة على الأقل) ومعاملات الارتباط بين المفردات أكبر من 0.8 فإنه يمكن إجراء التحليل

العاملى لعينات صغيرة مثلاً 50 حالة ويكون التحليل مناسب والنتائج ذات موثوقية ودقة.

وترى (Tabachnick & Fidell, 2013) أن حجم عينة 300 على الأقل يكون كافى، ولكن لا توجد قواعد صارمة بالضبط لتحديد حجم العينة لأن هذا يتحدد بطبيعة البيانات (Fabrigar, Wegener, MacCullum, & Strahan, 1999) وعمومًا لبيانات قوية فإن حجم العينة الصغير يعطى تحليل ذو موثوقية وثبات. والبيانات القوية هي عالية الشيع مع عدم وجود تشبعات مزدوجة Cross – loadings (تشبع المفردات على عاملين) بالإضافة إلى وجود متغيرات عديدة تتشبع بدرجة قوية بالعوامل (Costello & Osborne, 2005) ولكن هذه الشروط غير متوفرة فى الواقع العملى .

• **نسبة حجم العينة إلى المتغيرات:** يوجد توصيات أخرى موجودة فى التراث لتحديد حجم العينة فى ضوء عدد المتغيرات المتضمنة فى التحليل حيث يمكن تمثيل المتغيرات بـ 3 حالات (1:3) أو المتغير لخمس حالات (1:5) أو عشر حالات (1:10) أو خمسة عشر حالة (1:15) أو عشرون حالة (1:20) ، فإذا كان التحليل يتضمن عشرين مفردة فى مقياس إذا يمكن إجراء التحليل لـ (20X30) 600 أو $5 \times 20 = 100$ أو 200 أو 400 وهكذا.

• **البيانات المفقودة:** إذا وجدت بيانات مفقودة على بعض الحالات فاستخدم طريقة Pairwise حيث يتم استبعاد هذه الحالات من التحليل، ولكن استخدام أحد الطرق التعويضية مثل احلالها بالمتوسط أو استخدام الانحدار من شأنه أن يسبب معاملات ارتباطات عالية وهذا قد يسبب فى ظهور عوامل احيانًا لا تمثل الظاهرة.

• **الاعتدالية:** يستخدم التحليل العاملى لتلخيص العلاقات بين مجموعة كبيرة من المتغيرات ويفترض أن تكون توزيعات درجات أو بيانات هذه المتغيرات تتسم بالاعتدالية بقدر الإمكان وإذا لم تتوافر الاعتدالية فتكون الحلول لا معنى لها وهنا يجب التأكيد من الاعتدالية المتدرجة عند استخدام الأسلوب الإحصائى لتحديد عدد العوامل. والاعتدالية المتدرجة هي الدمج والاتحاد الخطى لكل المتغيرات الداخلة فى

التحليل ويمكن تقدير الاعتدالية لكل متغير على حده من خلال مؤشرات الالتواء والتفرطح.

- **الخطية:** الاعتدالية المتدرجة تتطلب أن تكون العلاقة بين كل زوج من المتغيرات خطية وإذا لم تكن العلاقة خطية فيجب عدم إجراء التحليل و يتم التحقق من الخطية من خلال فحص شكل Scatter plots وإذا وجدت عدم الخطية يمكن اللجوء إلى عملية التحويل للبيانات.

- **القيم المتطرفة:** وجود قيمة متطرفة وأكثر في كل أساليب الإحصاء المتدرجة له تأثير سلبي على حلول التحليل العاملي.

- **التلازمية الخطية المتعددة و Singularity:** في تحليل المكونات الرئيسية PCA فان التلازمية لا تمثل اى مشكلة لأنه لا حاجة لتحويل أو لقلب المصفوفة ولكن لمعظم طرق التحليل العاملي الاستكشافي فإن قضيتي التلازمية الخطية و Singularity تمثل اشكالية في التحليل وإذا كان محدد المصفوفة أو القيم الكامنة المرتبطة ببعض العوامل صفر فانه يدل على وجود التلازمية الخطية المتعددة. ويتم فحص Singularity من خلال فحص مربع معامل الارتباط المتعدد (SMC) لكل متغير على حده حيث يتعامل المتغير كتابع وبقية المتغيرات كمستقلة فاذا كانت اى قيمة لـ SMC هي واحد فإنه توجد قضية Singularity وكذلك التلازمية الخطية المتعددة ويجب حذف هذا المتغير من التحليل.

- **مصفوفة الارتباط R :** يجب أن تتضمن مصفوفة الارتباط الداخلة في التحليل أحجام مختلفة حيث إن أحجام العينات الكبيرة تتجه نحو انتاج معاملات ارتباطات منخفضة إذا لم تزيد معاملات الارتباطات عن 0.3 فإن استخدام هذه المصفوفة لإجراء التحليل العاملي يكون موضوع تساؤل لأنه لا تعطى فرصة أو احتمال لتحليل العاملي و انتاج عوامل و صنف (1995) Hair et al. بأن معاملات الارتباط كحد أدنى 0.30 و 0.40 مهم و 0.50 دالة عملياً و على ذلك يفضل أن تكون معاملات الارتباط في المصفوفة متوسطة الحجم فأعلى، ومعاملات الارتباط المرتفعة في

المصفوفة ليس دليل على أن تولد هذه المصفوفة عوامل (لاحظ أن العوامل تتولد من الارتباطات المرتفعة بين مجموعة معينة من المتغيرات).

يوجد اختبار Bartlett's (1954) لتشخيص Sphericity وهو للتحقق من فرضية أن معاملات الارتباطات في مصفوفة الارتباط تساوى صفر بمعنى أن الفرق الصفرى $H_0: p=0$ وهذا الاختبار يعطى دلالة لأحجام العينات الكبيرة حتى لو كانت الارتباطات منخفضة جداً واستخدام هذا الإحصاء مفيد إذا كانت أحجام العينات صغيرة ويمثل المتغير بـ 5 حالات، والدلالة الإحصائية لهذا الإحصاء يشير إلى الاستخدام الفعال لمصفوفة الارتباط.

وتوجد اختبارات عديدة لاختبار مدى قابلية المصفوفة للتحليل العاملي Factorability حيث يعطوا اختبارات الدلالة الإحصائية للارتباطات منها مصفوفة Anti-image correlation matrix و Kaiser test لقياس مناسبة المعاينة sampling adequacy فالدلالة الإحصائية لهذه الاختبارات تشير إلى صلاحية العلاقات بين كل زوج من المتغيرات للتحليل العاملي و لذلك إذا كانت المصفوفة قابلة للتحليل العاملي نتوقع وجود دلالات إحصائية عديدة بين كل زوج من المتغيرات و مؤشر Kaiser وهو يعرف في SPSS باختبار Kaiser - Meyer - Olkin (KMO)

$$= \frac{\text{مجموع معاملات مربعات الارتباط}}{\text{مجموع معاملات مربعات الارتباط} + \text{مجموع مربعات الارتباط معاملات}}$$

والقيمة 0.60 فأعلى تشير إلى صلاحية جيدة لمصفوفة معاملات الارتباطات لاستخدامها في التحليل العاملي ويرى (Hair et al 1995) أن القيمة 0.5 - KMO تشير إلى مناسبة المصفوفة لاستخدامها في التحليل العاملي و يرى Kaiser (1974) أن الحد الأدنى 0.5 والقيم بين 0.5 إلى 0.7 مناسبة والقيم من 0.7 إلى 0.8 جيدة والقيم 0.8 إلى 0.9 جيد جداً وأعلى من 0.9 ممتاز ويوصى باستخدامه إذا كانت نسبة المتغيرات إلى حجم العينة هي 1 إلى 5.

ثانياً: طرق استخلاص العوامل: طريقة المكونات الرئيسية Principal Component Analysis (PCA) هي تعتبر الطريقة Default لاستخلاص

العوامل فى معظم البرامج الإحصائية مثل SAS ، SPSS وهذا اكسبها انتشارًا واستخدامًا. وحقيقة أن PCA ليست طريقة من طرق التحليل العاملى ويوجد خلاف بين المتخصصين عن متى تستخدم وبعض الباحثين ينصحون بالاستخدام المحدود والمقيد لها (Mulaik , 1990 , Widman , 1990) والبعض يختلف معهم حيث يروا لا فروق تقريبًا بين طريقة PCA و طرق التحليل العاملى الأخرى ويفضلون استخدام(Guadagnolia & velicer , 1988 , Stieger , 1990).

ويرى (Castello & Osbone (2005 أن طرق التحليل العاملى مفضلة عن طريقة تحليل المكونات الرئيسية لأنها فقط طريقة لتقليل البيانات وأنها كانت أكثر استخدامًا عندما لم تتوفر إمكانيات برامج كمبيوترية عالية الجودة لأنها سهلة الحسابات حيث يتم حسابها بدون الحاجة لأى بناء تحتى مسبق اثناء استخلاص العوامل. ولكن طرق التحليل العاملى (FA) تحلل التباين المشترك ويعطى نفس السلوك مع تجنب تضخم تقديرات التباين المفسر.

وتوجد طرق عديدة لاستخلاص العوامل

1. طريقة المكونات الرئيسية (PCA)

ب. طرق التحليل العاملى وتتضمن:

1. طريقة عوامل المحاور الرئيسية (Principal axis factoring (PAF)

2. طريقة الاحتمال الأقصى (Maximum likelihood (ML)

3. طريقة المربعات الدينا غير الموزونة (Unwiegthed least squares(ULS)

4. طريقة المربعات الدينا المعممة (Generalized least squares(GLS)

5. عوامل الفا (Alpha factoring(AF)

6. العوامل التخيلية (Image factoring(IF)

ومهما يكن فإن طريقتى PAF ، PCA أكثر استخدامًا فى البحوث و الدراسات المنشورة والمفاضلة بين هاتين الطريقتين اخذ مجال للحوار و المناقشة بين الباحثين،

ويرى (Thompson 2004) أن الفروق العملية بينهما غالبًا ليست لها دلالة خاصة إذا كانت المتغيرات عالية الثبات ويوجد في التحليل 30 متغير فأكثر، ويشير Steven (2007) مع 30 متغير فأكثر وقيم شيوخ أعلى من 0.7 لكل المتغيرات فإن الفروق قليلة جدًا ، ومع عدد متغيرات أقل من 20 وقيم شيوخ منخفضة أقل من 0.4 فالفرق تكون واضحة وتعتبر طريقة PCA هي الطريقة المستخدمة في البرنامج الإحصائي إذا لم تحدد طريقة معينة وينصح باستخدامها عندما لا توجد نظرية أو نموذج لطبيعة البناء، ويمكن اعتبار حلول PCA حلول مبدئية أو قيم مبدئية لـ EFA وعليه فالباحثين أحيانًا يستخدموا PCA لتقليل البيانات ثم على هذه المكونات يجرى تحليل عاملي باستخدام أحد أساليبه.

ثالثًا: معايير تحديد عدد العوامل: بعد استخلاص العوامل باستخدام أى الطرق السابقة فإن الباحث يواجه قضية وهي ما عدد العوامل التى سيبقيها فى التحليل لتدخل عملية التدوير. فعندما تتم عاملية المتغيرات فإن العدد الكلى من العوامل يساوى عدد المتغيرات الداخلة فى التحليل (Thompson , 2004). ولأن هذه العوامل كلها لا تسهم بصورة جوهرية مع الحلول النهائية فى عملية التفسير و هذه العوامل تعتبر خطأ أو Noise ولأن الهدف من التحليل العاملي الاستكشافى هو الابقاء على عدد محدود من العوامل التى تفسر معظم تباينات المتغيرات المقاسة فمن المهم للباحث استخلاص العدد الصحيح من العوامل لأن هذا سيؤثر على تفسير النتائج.

وحيث إدخال كل العوامل فى مرحلة التدوير يمكن أن يكون له تأثيرات غير مرغوبة على النتائج وإذا لم يحدد الباحث عدد العوامل فإن البرنامج يقوم باستبقاء العوامل التى يزيد لها القيم الكامنة Eigenvalues عن الواحد الصحيح وهذا اتفاق أو تقليد فى التراث ولكن هذه أقل الطرق دقة فى تحديد عدد العوامل المستبقة فى التحليل (Velicer & Jakson , 1990) لأنهم اختبروا هذه القاعدة و توصلوا أن 36% من العينات استبقت على عوامل كثيرة ليس لها داعى فى التحليل وذلك باستراتيجية المحاكاة مونت كارلو.

Comunalities الشيوع

قيمة الشيوع للمتغير هي التباين المفسر في المتغير عن طريق كل العوامل وهو مجموع مربعات تشبعات المتغير بالعوامل، ونسبة التباين المفسر لمجموعة من المتغيرات المقدرة من خلال العامل هي عبارة عن مجموع مربعات التشبعات للمتغيرات عبر العامل مقسومًا على عدد المتغيرات

فاذا كان تشبعات أربعة متغيرات على عامل ما هي:

$$-0.86 , 0.071 , 0.994 , 0.997$$

فان التباين المفسر في هذه المتغيرات جراء العامل هي:

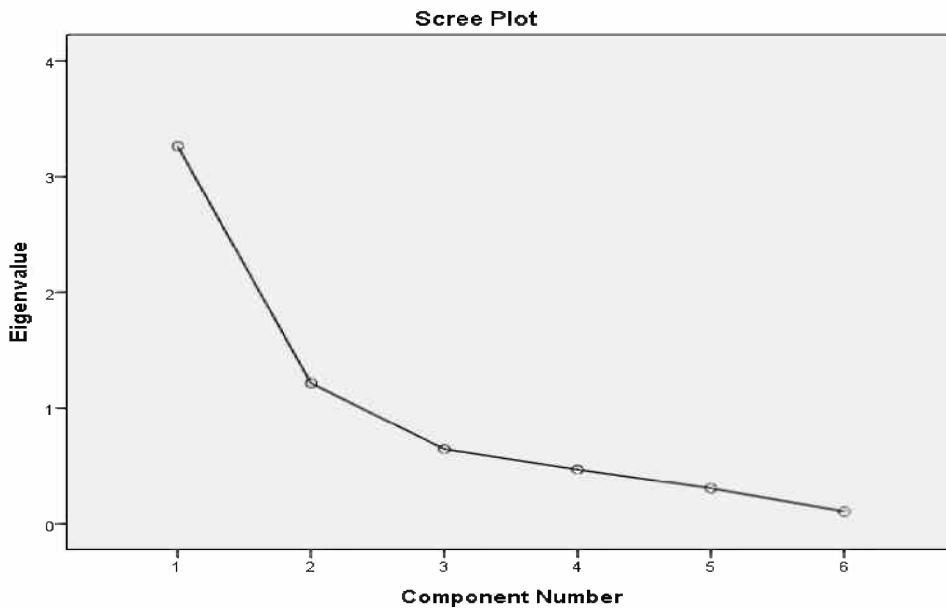
$$= \frac{(-.86)^2 + (-0.071)^2 + (0.994)^2 + (0.997)^2}{4} = \frac{1.994}{4} = 0.50$$

وهذا مفاده أن العامل الأول فسره 50% من تباين المتغيرات الاربعة و كما نعلم أن التباين الكلى لمتغير معين له مكونين جزء منه يتشارك مع متغيرات اخرى و يسمى التباين العام Common variance الذى يفسره العامل وجزء منه خاص به و يسمى التباين الفريد Unique variance ويستخدم التباين الفريد ليشير إلى التباين الخاص بالمتغير و يوجد نوع اخر من التباين وهو خاص بالمتغير و لكنه ثابت و يسمى التباين الخطأ أو العشوائى Error variance ونسبة التباين المشترك أو العام فى المتغير تسمى الشيوع واذا كان التباين الخاص أو الفريد = صفر فإن قيمة الشيوع تساوى الواحد الصحيح واذا لم يتشارك لمتغير مع تباين متغير اخر فإن قيمة الشيوع تساوى صفر ولذلك فإن التحليل العاملى يهتم فقط بالتباين العام.

ويعتبر استخلاص عدد كبير من العوامل أفضل حيث يزيد من نسبة التباين المفسر فى البيانات ولكن اخذ كل العوامل يجعل البناء أقل بساطة Less Parsimonious واختيار عدد العوامل قضية اعقد من اختيار طريقة التدوير ولكن فى طرق التحليل العاملى التوكيدى مثل طريقة ML يمكن تحديد عدد العوامل فى ضوء الاطر النظرية للبحث أو للمفهوم المراد دراسته.

ويوجد عدة طرق لتحديد العدد المناسب من العوامل التي يجب أن نستبقها في التحليل، وتعتبر أحد الطريقة السريعة لتقدير عدد العوامل هو أحجام القيم الكامنة وحيث إذا كان الجذر الكامن للعامل 1 فأكثر فإنه يدخل في التحليل و التفسير.

والمحك الآخر اقترحه (Cattel 1966) وهو اختبار Scree test للقيم الكامنة والعوامل حيث يتم عرضة في شاكل تصاعدي ، فالقيمة الأعلى للقيم الكامنة هي تمثل اول عامل وهكذا حتى نصل إلى نقطة معينة وهي نقطة التواء وانكسار واضحة وتكون القيم الكامنة أقل من الواحد الصحيح ، وهو يتضمن عرض بياني لرسم القيم الكامنة وبالبحت في نقطة الانكسار في البيانات أو النقطة التي يحدث عندها تحول بدرجة ملحوظة لمسار الخط المنحنى وعدد نقاط البيانات (الاحداثيات) فوق نقطة الانكسار هي تعتبر عدد العوامل التي نستخدمها في التحليل والتفسير كما في الشكل الاتي:



كما توجد طريقتي محك (velicer'sMAP) Minimum average partial correlation و التحليل الموازي Parallel analysis. وعلى الرغم أن هاتين الطريقتين أكثر دقة وسهلة الاستخدام الا أنهما غير متاحيتين في معظم البرامج الإحصائية ويرى (Costello & Osborne 2005) أن افضل طريق لتحديد عدد العوامل هي Scree test وهي موجودة في البرامج الإحصائية.

وقام Zwick & velicer (1986) بإجراء دراسة محاكاة للمقارنة بين Scree test Parallel, Velicer's map, Parallel test, MAP, Scree test ، بينما يعتبر محك $EV > 0.1$ Overestimate لعدد العوامل المستبقة بمعنى يتبقى عدد عوامل أكبر من المفترض و لكن هذا يتناقض مع ما توصل اليه Mote (1970) إلى أن محك القيمة الكامنة < 1 يتبقى عدد عوامل أقل من المفروض الابقاء عليها و عليه فإن Scree test طريقة مفضلة لتحديد عدد العوامل ولكن طريقة Parallel, MAP أكثر دقة وموضوعية.

وطرح Thompson(1988) استخدم طريقة Bootstrap لتحديد العوامل،بينما توصل Gorsuch (1983) إلى أن نتائج Scree test تكون أكثر ثباتًا عندما يكون حجم العينة كبيرًا و قيم الشيوخ عالية وتشبعات المتغيرات عالية.

وينصح Henson & Roberts (2006) باستخدام محكات متعددة لتحديد عدد العوامل. ويوصي Kaiser (1960) بإبقاء كل العوامل التي تزيد القيم الكامنة لها عن الواحد الصحيح و القيم الكامنة تمثل مقدار التباين المفسر عن طريق العامل و اشار Jolliffe (1972) إلى أن محك كايزر شديد الصرامة وتقتصر استبقاء العوامل التي لا قيم كامنة أكبر من 0.70 و توجد فروق كبيرة بين محك Kaiser ومحك Jolliffe حيث يعتبر محك Jolliffe تقدير مبالغ لعدد العوامل. ولكن هذا المحك يكون دقيق عندما يكون عدد المتغيرات أقل من 20 وقيم الشيوخ بعد الاستخلاص تزيد عن 0.70. للعوامل وحجم العينة عن 250 فاكتر (Field, 2009). ويرى Steven (2009) أن استخدام Scree test يكون دقيق لحجم عينة أكثر من 200. وبرنامج SPSS يستخدم محك Kaiser لاستخلاص عدد العوامل إذا لم تطلب من البرنامج أى محك اخر ويرى Field (2013) أن تحديد عدد العوامل يتوقف على الهدف من التحليل فإذا كان الباحث يهدف للتغلب على قضية التلازمة الخطية المتعددة فى تحليل الانحدار فمن الافضل استخلاص عدد كبير من العوامل. وبعد تحديد عدد العوامل فمن المهم النظر فى التشبعات بعد التدوير لتحديد عدد المتغيرات التى تشبع على كل عامل وإذا تشبع متغير واحد على عامل تشبعًا عاليًا فإن هذا العامل يعتبر غير محدد تحديدًا جيدًا ، وإذا تشبع متغيرين على العامل يمكن اخذه فى الاعتبار فى التفسير وهذا يتوقف على

نمط العلاقات بين المتغيرين وكذلك بين المتغيرين وبقية المتغيرات فى مصفوفة الارتباط فإذا كانت الارتباط بينهما عالى (أكثر من 70. مثلاً) وغير مرتبطتين مع عينة المتغيرات فى التحليل ففى هذه الحالة يعتبر هذا العامل على درجة معقولة من الأهمية ويعتمد عليه فى التفسير وعموماً فى تفسير العوامل التى تتضمن متغيراً أو متغيرين يمثل درجة كبيرة من الخطورة Hazardous .

وتظهر قضية أخرى وهى من الأفضل الإبقاء على عدد محدود ام عدد كبير من العوامل إذا كان عددها غير محدد مسبقاً؟ فالباحث دائماً يستبقى العوامل التى لها قيم كامنّة أكبر من الواحد الصحيح ولكن فى بعض الأحيان توجد عوامل جذرها الكامن أقل من الواحد وتعتبر فى غاية الأهمية فى تفسير الظاهرة وهذا قد يكون سبب مقنع لتضمينه فى عملية التحليل والتفسير.

رابعاً: طرائق تدوير العوامل Rotation: بعد استخلاص العوامل يحدث أن تتشعب معظم المتغيرات على العامل الأول بتشعبات عالية، ثم تتوزع على بقية العوامل بتشعبات أقل، وهذا يضع الباحث فى حيرة؛ حيث يكون هذا مخالف لطبيعة البناء المراد التحقق من بنيته العاملية، وعليه تكون عملية التفسير صعبة جداً للبناء؛ ولذلك يلجأ الباحثون إلى إجراء عملية تدوير العوامل Factor rotation؛ ليحدث تمايز للعوامل بالمتغيرات التى تتشعب عليها. وهى عملية رياضية وفيها يجرى تدوير محاور العوامل؛ ولذلك فالتدوير يمدنا ببنية عاملية مناسبة Factorial suitability.

وإستراتيجيات التدوير متعددة وتصنف إلى تصنيفين عريضين، هما: التدوير المتعامد Orthogonal والتدوير المائل Oblique، ومعظم الباحثين يلجؤون إلى تدوير نتائج التحليل العاُملى لتسهيل تفسير هذه العوامل. وفى التحليل العاُملى الاستكشافى إسهام المتغير فى عامل معين يشار إليه بمعاملات نمط العوامل المستهدفة Factor pattern coefficient وأيضاً يشار إليه بالمعاملات البنائية للعوامل Factor

structure coefficient وهذه تمثل العلاقات بين المتغيرات المقاسة والعوامل. وفي التحليل العاملي، فإن المصفوفة البنائية للعوامل تعطى العلاقات بين كل المتغيرات المقاسة وكل العوامل (المتغيرات الكامنة) المستخلصة، وعندما تكون العوامل متعامدة، فإنها تبقى غير مرتبطة، والمصفوفة البنائية للعوامل تناظر مصفوفة العوامل المستهدفة، وبناء عليه فإن المصفوفة البنائية هي نتاج حاصل ضرب المصفوفة المستهدفة (النمطية للعوامل) في مصفوفة ارتباطات العوامل. وعملية التدوير تعظم التشعبات العالية وتقلل التشعبات المنخفضة للمفردات؛ ولذلك فهي تعطى نتائج أكثر بساطة وقابلة للتفسير.

والقرار في تحديد نوعية التدوير سواء كان متعامداً أو مائلاً؛ ففي التدوير المتعامد فإن العوامل غير مرتبطة، والتدوير المتعامد يسهل في تفسير النتائج، أما إذا اعتقد الباحث أن العمليات أو الأبنية التحتية مرتبطة، فيجب أن يستخدم التدوير المائل، وعلى الرغم من أنها تتناسب مع طبيعة البناء على مستوى النظرية، ولكنها لها محددات عملية في التفسير والوصف وتقرير النتائج.

• **التدوير المتعامد:** توجد العديد من الطرائق للتدوير المتعامد منها Varimax وEquimax وQuartimax، وهي متاحة في برامج SPSS وSAS. ولكن تعد طريقة Varimax أسهلهم استخداماً عن كل طرائق التدوير. والهدف من هذه الطريقة هو تبسيط العوامل عن طريق تعظيم تباين التشعبات بالعوامل عبر المتغيرات، فالتشعبات العالية قبل التدوير تكون أعلى بعد التدوير، والتشعبات المنخفضة قبل التدوير تزداد انخفاضاً بعد التدوير. وطريقة الفاريماكس تتجه إلى توزيع التباين داخل العوامل،

وتصبح متساوية تقريباً في أهميتها. فالتباين المستخلص عن طريق العامل الأول قبل التدوير يجرى توزيعه على العوامل الأخرى بعد التدوير، فهي تحاول أن يكون التشبع بدرجة عالية لعدد محدود من المتغيرات لكل عامل معين، وهو يؤدي إلى الحصول على تجمعات أو عوامل قابلة للتفسير.

وطريقة Quartimax تتبع نفس سلوك طريقة Varimax ولكنها ليست شائعة الاستخدام مثل: Varimax. وطريقة Equimax هي خليط أو مزيج بين الطريقتين السابقتين، وتكون عديمة الجدوى ما لم يحدد الباحث مسبقاً عدد العوامل. وتعد طريقة Varimax هي default للبرامج SPSS و SAS ما لم يجرى تحديد طريقة معينة. التدوير المائل Oblique rotation: يعطى التدوير المائل ارتباطات بين العوامل، ومقدار الارتباطات المفترض به بين العوامل يتحدد في ضوء البنية النظرية للمفهوم.

وتوجد طرائق عديدة للتدوير المائل، مثل: Direct oblimin و Promax. والتدوير المائل أكثر تعقيداً من التدوير المتعامد، وينصح (2009) Field باستخدام طريقة Promax؛ لأنها إجراء سريع وتصلح لقواعد البيانات الكبيرة، وعلى الرغم من تعدد طرائق الاستخلاص وطرائق التدوير، ولكن في الممارسة الفروق بينهما ضئيلة (Fava 1992, Velicer &، ويرى (1999) Fabirgar et al. أن طرائق التدوير المائل تعطي نتائج متماثلة لطرائق التدوير المتعامد، ويرى (2009) Field أن بعض الباحثين يرى أن استخدام التدوير المتعامد في البحوث النفسية غير ملائم؛ لأن المفاهيم النفسية في طبيعتها ذات بناء تحتى مرتبط ومتفاعل.

الجدول (1.8): مقارنة بين أساليب التدوير المختلفة.

أسلوب التدوير	البرنامج	نوعه	ملاحظة
Varimax	SPSS, SAS	متعامد	أكثر استخدامًا في البحوث.
Quartimax	SPSS, SAS	متعامد	يكون العامل الأول عامًّا مع عوامل أخرى للمتغيرات.
Equimax	SPSS, SAS	متعامد	ربما يكون غير مناسب.
Direct oblimin	SPSS	مائل	يسمح بإعطاء عدد كبير من العوامل.
Direct Quartimin	SPSS	مائل	يسمح بارتباطات عالية بين العوامل.
Oblique	SPSS, SAS	مائل ومتعامد معًا	

سريع وسهل الحسابات.	مائل	SPSS, SAS	Promax
مفيد فى التحليل العاملى التوكيدى.	مائل	SAS	Procrustes

حدود القطع لقبول تشبعات العوامل

من الممكن قبول قيمة تشبع المتغير بالعامل فى ضوء دلالتها الإحصائية لأنها عبارة عن معامل ارتباط أو معامل انحدار ولكن هذا المعيار نادر الاستخدام فى البحوث لان الدلالة الإحصائية تعتمد على حجم العينة و يختار الباحث دلالة التشبع حيث إذا كانت قيمة 0.30 فاكتر فانة يقبل كمتغير على العامل، ولكن (Steven 2009) يضع قيمة حرجة لقبول التشبع و يرى لحجم العينة 50 فإن قيمة التشبع 0.772 تعتبرها دلالة و لحجم العينة 100 فإن قيمة التشبع 0.512 فأكتر، ولحجم عينة 200 فإن التشبع يجب أن يكون أكبر من 0.364 لقبوله كمتغير يتشبع بالعامل، ولحجم عينة 300 فالقيمة 0.298 أو اكتر، ولحجم عينة 600 يجب أن يكون التشبع 0.21 فاكتر، ولحجم عينة 1000 يجب أن يكون قيمة التشبع 0.162 وهذه القيم على أساس قيمة الفا 0.01 لاختبار ذو ذيلين.

اقترح (Comrey & lee 1992) قيمة التشبع 0.71 فاكتر تعتبر ممتاز (50% من التباين)، و 0.63 جيد جدا، و 0.55 (30% من التباين) جيد، و 0.45 (20% من التباين) يعتبر متوسط و 0.32 (10% من التباين) ضعيف.

وعلى ذلك فمع أحجام العينات الكبيرة فإن التشبعات صغيرة الحجم يجب وضعها فى الاعتبار و لها معنى ذو دلالة إحصائية وبرنامج SPSS لا يقدم اختبارات دلالة لتشبعات العوامل ولكنه يطبق ارشادات (Steven 2009) ومربع قيمة التشبع يعطى معامل التحديد R² وهى نفس نسبة التباين المفسر فى المتغير جراء العامل.

خامسًا: تفسير العوامل Interpretation of factors: يحاول الباحث تسمية العامل و هذه يتطلب نوع من الحساسية لدى الباحث، ومدى قابليته للتفسير ترجع إلى الباحث فى ضوء الاطر النظرية، وتختلف تفسير العوامل باختلاف طريقة التدوير التى

استخدمها الباحث و من الافضل كما يرى (2013) Tabachink & Fidell أن يتم تفسير تشبعات المفردات بالعوامل في ضوء مصفوفة Patten matrix.

ولكن يرى (2005) Castelo & Osbtern إذا استخدم الباحث التدوير المتعامد تفسر مصفوفة العوامل المدورة Rotated factor matrix وعند استخدام التدوير المائل تفسر النتائج في ضوء المصفوفة المستهدفة Pattern matrix. وكذلك يتم تفسير وعرض مصفوفة ارتباطات العوامل Factor correlation matrix وتسمية العوامل تعتمد على التعريفات النظرية للمفاهيم وأبعادها.

وعلى الباحث الاخذ في الاعتبار عدة امور اهمها:

1. تعتبر شيوخ المفردات عالية إذا كانت قيمتها 0.80 فأعلى ولكن هذا نادر الحدوث في البيانات، فمعظم القيم السائدة في العلوم الاجتماعية هي من منخفضة 0.40 إلى متوسطة 0.70 ولو أن قيم الشيوخ للمتغير أو المفردة أقل من 0.40 ربما لا يرتبط مع مفردات اخرى أو من المفترض وجود عامل اضافي في البيانات يجب اكتشافه (Castelo& Osbtern, 2005).

2. يشير (2013) Tabachnik & Fidell أن قيمة 0.32 لقبول تشبع المفردة هي الحد الأدنى وهي تساوى 10% من التباين المتداخل مع مفردات اخرى في العامل، بينما التشبعات المزدوجة عند 0.32 أو على عاملين فأكثر تكون مقبولة وتدخل في تفسير البنية العاملية ، وعلى الباحث أن يقرر أن يستبعد المفردات ذات التشبعات المزدوجة من التحليل وربما يكون هو اختيار مفضل إذا وجدت مفردات ذات تشبعات قوية على العامل (0.5 فأعلى) وإذا تشبع المتغير على أكثر من عامل فهذا ربما يدل على ضعف صياغة أو كتابة المفردة .

3. العامل الذى يتشبع عليه أقل من ثلاث مفردات يعتبر ضعيف أو غير مستقر أو غير ثابت فخمسة مفردات أو أكثر ذات تشبعات عالية بالعامل (0.5 فأكثر) شئ مرغوب لتمثيل العامل وهذا عامل مستقر Solid factor.

تنفيذ قضية بحثية في SPSS:

قام باحث بتطبيق استبيان لقياس الاكتئاب والقلق معًا وتكون من ستة مفردات ثلاثة تقيس القلق وهي: A₁ أنا عصبى، A₂ أنا قلق، A₃ أنا هادئ وثلاثة مفردات تقيس الاكتئاب: D₁ أنا مكتئب، D₂ اشعر انى ليس لى قيمة فى الحياة، D₃ أنا سعيد. وكانت بدائل الاجابة هى خمسة وهى ليس على الإطلاق (1) أحيانًا (2) وغالبًا (3) ومعظم الوقت (4) وطول الوقت (5) وتكونت العينة من 9 أفراد و كانت بياناتهم كالآتى:

الحالات	A ₁	A ₂	A ₃	D ₁	D ₂	D ₃
1	2	1	3	1	2	5
2	1	2	3	4	3	3
3	3	3	4	2	1	4
4	4	4	3	3	2	3
5	5	5	2	3	4	4
6	4	5	2	4	3	1
7	4	3	2	5	4	1
8	3	3	4	4	4	3
9	3	5	3	3	4	1

واراد الباحث التحقق من الصدق العاملى لهذا الاستبيان وتحديد الأبعاد التى تتجمع حولها هذه المفردات بالتالى:

هل يوجد بعد واحد ام بعدين للمفردات الستة؟ أو ما هو البناء العاملى لمفردات مقياس القلق – الاكتئاب

تنفيذ التحليل العاملى في SPSS

أولاً: إدخال البيانات: 1. بعد فتح البرنامج اضغط Variabl view ثم اكتب مسمى التغيرات تحت Name وهى A₁ و A₂ , A₃ و D₁ و D₂ و D₃

2. اضغط data view يظهر ستة أعمدة كالآتي:

	A1	A2	A3	D1	D2	D3
1	2.00	1.00	3.00	1.00	2.00	5.00
2	1.00	2.00	3.00	4.00	3.00	3.00
3	3.00	3.00	4.00	2.00	1.00	4.00
4	4.00	4.00	3.00	3.00	2.00	3.00
5	5.00	5.00	2.00	3.00	4.00	4.00
6	4.00	5.00	2.00	4.00	3.00	1.00
7	4.00	3.00	2.00	5.00	4.00	1.00
8	3.00	3.00	4.00	4.00	4.00	3.00
9	3.00	5.00	3.00	3.00	4.00	1.00

ثانيًا: تنفيذ التحليل العاملي: لتنفيذ الامر اتبع الخطوات الآتية: ولكن قبل إجراء التحليل العاملي فالممارسة الجيدة تقول لا بد من تقدير مؤشر الالتواء والتفرطح للمتغيرات خاصة إذا كان المتغيرات هي مفردات لها استجابات محدودة، ويتم ذلك من خلال: Analyze اختار Frequencies اختار Descriptive وتحت مربع Distribution اضغط على Kurtosis , Skewness والنتائج كالآتي:

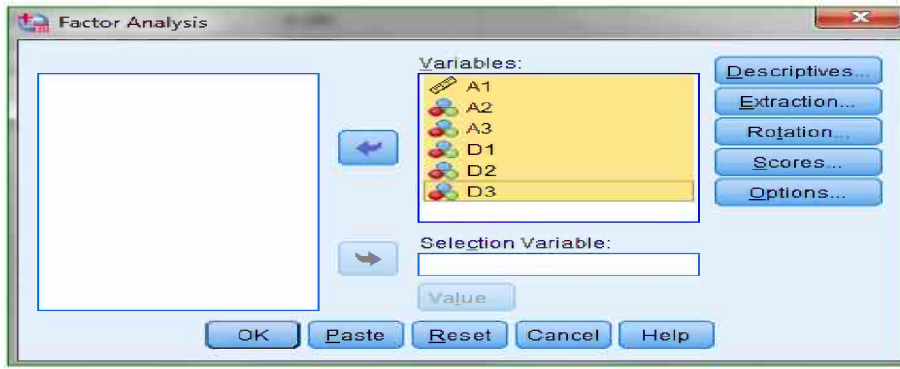
DESCRIPTIVES VARIABLES=A1 A2 A3 D1 D2 D3									
/STATISTICS=MEAN STDDEV MIN MAX KURTOSIS SKEWNESS.									
Descriptive Statistics									
	N	Minimum	Maximum	Mean	Std. Deviation	Skewness		Kurtosis	
	Statistic	Statistic	Statistic	Statistic	Statistic	Statistic	Std. Error	Statistic	Std. Error
A1	9	1.00	5.00	3.2222	1.20185	-.537-	.717	.270	1.400
A2	9	1.00	5.00	3.4444	1.42400	-.357-	.717	-.804-	1.400
A3	9	2.00	4.00	2.8889	.78174	.216	.717	-1.041-	1.400
D1	9	1.00	5.00	3.2222	1.20185	-.537-	.717	.270	1.400
D2	9	1.00	4.00	3.0000	1.11803	-.690-	.717	-.800-	1.400
D3	9	1.00	5.00	2.7778	1.48137	-.109-	.717	-1.300-	1.400
Valid N (listwise)	9								

ويتم إجراء التحليل العاملي في مرحلتين استخلاص العوامل ومرحلة التدوير كالآتي:

• مرحلة الاستخلاص Factor extraction:

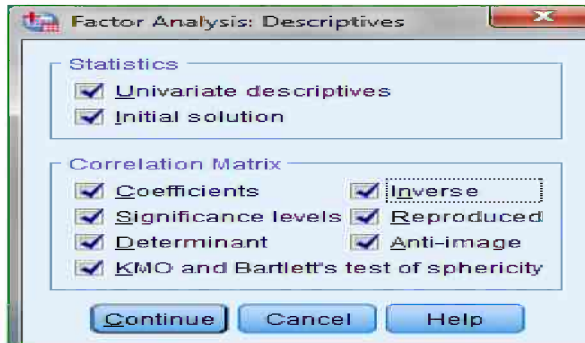
1. اضغط Analyze ثم اضغط Dimension Reduction أو Data Reduction ثم اختار

Factor تظهر الشاشة الآتية :



2. انقل المتغيرات الموجودة في المربع الأول إلى مربع Variables

3. اضغط على اختيار Descriptives تظهر الشاشة الآتية:



4. في مربع Statistics نشط أو

اضغط على univariate

descriptive ليعطى المتوسط و

الانحراف المعياري والانحراف

المعياري لكل متغير.

5. في مربع Correlation matrix اضغط على:

- Coefficients ليعطى مصفوفة الارتباط بين المتغيرات (R)

- Significance ليعطى الدلالة الإحصائية لكل معامل الارتباط

- محدد المصفوفة Determinant: وهذا الاختيار في غاية الأهمية لاختبار

multicollinearity.

- KMO and Bartlett's test: يعطى مقياس Kaiser- Meyer – Olkin لاختبار مناسبة

مصفوفة الارتباط للتحليل في التحليل العاملي ويجب أن تزيد قيمتها عن 0.50 أو 0.60.

- اختيار Reproduced يمدنا بمصفوفة الفروق بين مصفوفة البيانات ومصفوفة

الارتباط المشتقة جراء النموذج وهو ما يطلق عليه بمصفوفة البواقي Residual

وكما نعلم أن قيم هذه المصفوفة يكون صغير جداً ويجب ألا تزيد قيمتها لمعظم

معاملات الارتباط عن 0.05.

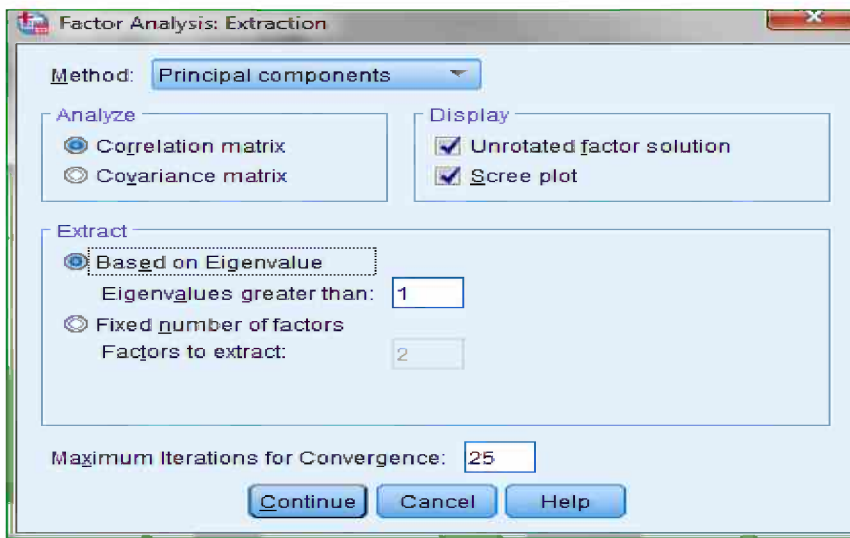
- اضغط على Inverse للحصول على مقلوب المصفوفة.

• اضغط على Anti - image correlation حيث تفيد فى التحقيق من مدى صلاحية المتغيرات فى التحليل العاىلى وهى قيمة KMO للخلايا القطرية لكل متغير على حده.

• اضغط على Initial solutions ليعطى الحلول الأولية قبل التدوير

6. اضغط على Continue لى شاشة الامر الرئيسية.

7. اضغط على الاستخلاص Extraction يعطى الشاشة الآتية:



8. اضغط على السهم امام Method تظهر عدم بدائل أو طرق للاستخلاص و لكن تظهر طريقة المكونات الرئيسية (PCA) كالبطريقة التى يقوم بها البرنامج ما لم تحدد له طريقة اخرى على الرغم أنها ليست طريقة من طرق التحليل العاىلى ولكنها تعطى نتائج مثل نتائج التحليل العاىلى ويمكن أن تختار اى من الطرق سواء ML , UIS , PAF أو غيرها وهنا تختار طريقة PCA

9. فى مربع Anlayze حدد نوع المصفوفة المراد تحليلها ويوجد بدلين مصفوفة الارتباط ومصفوفة التغاير وللمصفوفتين صيغ مختلفة لنفس الشئ. فمصفوفة الارتباط هى الصورة المعيارية لمصفوفة التغاير ويفضل فى التحليل العاىلى تحليل مصفوفة الارتباط لا إذا كانت القياسات تأخذ مقاييس أو وحدات مختلفة فإن هذا لا يؤثر فى التحليل وإذا اعتمد الباحث على مصفوفة الارتباط وفى هذا المثال المتغيرات

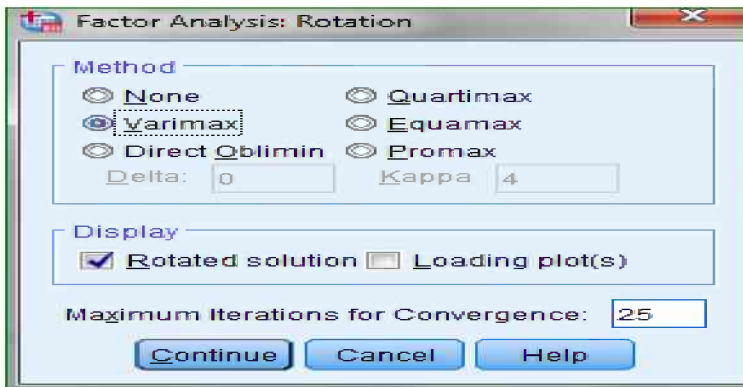
تقاس بنفس وحده القياس (مقياس ليكرت الخماسي) وحتى إذا كانت المتغيرات لها نفس الوحدة فإن لها تباينات مختلفة وهذا يسبب مشاكل لطريقة المكونات الرئيسية ((Field, 2009)، فاستخدم مصفوفة الارتباط تجنباً هذه المشاكل وعليه اضغط على .Correlation

10 . فى مربع Display يوجد Unrotated factor solution اختاره وهى الحلول الأولية قبل التدوير وضغط على اختيار Scree لتحديد عدد العوامل .

11 . فى مربع Extract يمكن تحديد عدد العوامل فى ضوء محك القيم الكامنة أكثر من 1 وهى default للبرنامج و هو محك kaiser ولكن تمكن تغييرها إلى محك Jolliffe's وهى 0.70 أو أى قيمة أخرى و لكن من المفضل إجراء التحليل عند قيم كامنة وعدد العوامل المحدد فى ضوء الاطار النظرى Fixed number of factor ويفضل إذا كان البحث استكشافى لا يتم تثبيت أو تحديد عدد العوامل مسبقاً. وإذا اختلف عدد العوامل فى صورة محك Screeplot محك kaiser فعليك فحص قيم الشيوخ للعوامل.

12. اضغط Continue لترجع إلى شاشة الامر الرئيسية.

13. اضغط على التدوير Rotation تظهر الشاشة الآتية:



14. فى صندوق Method اختار طريقة التدوير سواء المتعامدة أو المائل و بما أن المثال الحالى يتكون من مكونين إذا كان غير مرتبطين اختار تدوير متعامد (يفصل Varmix) وإذا كانوا مرتبطين اختار طريقة تدوير مائل سواء Direct oblinmim أو Promax فاضغط على Varmix

15. فى مستطيل Display يوجد Rotated solution و يعطى الحلول بعد التدوير و هى تعتبر default للبرنامج حيث يمدنا بها البرنامج سواء حددتها أو لم تحددها بينما اختيار Loading plot هو تمثيل بياني لكل متغير فى ضوء العوامل المستخلصة (أقصى عدد عوامل ثلاثة) وإذا كان عدد العوامل أربعة فإن البرنامج لا يمدنا بها ويسهل تفسير هذا الرسم إذا وجد عاملين اضغط عليها.

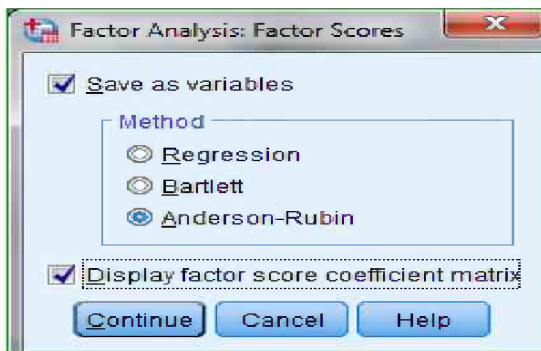
16. اختيار Maximum Iterations for convergence حيث يقوم البرنامج بأخذ عدد من المحاولات لإعطاء حلول عاملية وعليه فهى عدد المرات أو المحاولات التى يحاول البرنامج إجرائها للحصول على حلول لمصفوفة الارتباط ويحددها البرنامج بـ 25 مرة و لكن مع قواعد البيانات الكبيرة فإن الكمبيوتر يجد صعوبة فى الوصول إلى الحلول ويعطى رسالة تفيد ذلك وعليه يجب تغيير هذه القيمة إلى 40 أو 50 أو إلى 100 مثلاً و عليه اضغط عليها و غيرها إلى 40 مثلاً.

17. اضغط على Continue لتعود إلى شاشة الامر الرئيسية.

18. اضغط على اختيار Scores وهى

الدرجات العاملية ويعطى الشاشة الفرعية الآتية:

وهذا يسمح بحفظ الدرجات العاملية لكل فرد فى ملف البيانات و برنامج SPSS يقوم بتكوين عمود جديد لكل



عامل مستخلص و يضع الدرجة العاملية لكل فرد فى الأعمدة ويمكن أن تستخدم هذه الدرجات العاملية فى تحليلات لاحقة أو تستخدم لتحديد مجموعة من الأفراد الذين يحصلون على درجات عاملية على العوامل و لذلك:

- اضغط على اختيار save As variable
- حدد الطريقة المستخدمة فى الحصول على الدرجات العاملية و هى

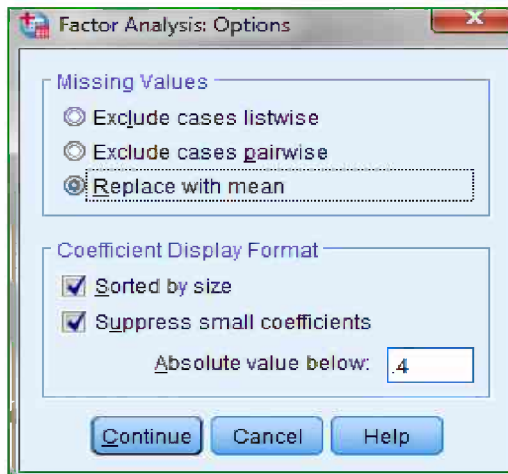
طريقة الانحدار Regression method وفى هذه الطريقة تصحح تشبعات العوامل من الارتباطات الأولية بين المتغيرات وللحصول على مصفوفة معاملات الدرجات

العاملية يتم ضرب مصفوفة تشبعات العوامل بمقلوب مصفوفة الارتباط (R^{-1}) بمعنى قسمة مصفوفة التشبعات على مصفوفة الارتباطات وهذا يعطى مصفوفة الدرجات العاملية وهذه المصفوفة تعطى مؤشر نقي للعلاقات الفريدة بين المتغيرات والعوامل.

وتوجد طرق أخرى مثل طريقة Bartlett وطريقة Rubin - Anderson ويفضل استخدامهما للحصول على درجات عاملية غير مرتبطة وإذا وجدت ارتباطات بين درجات عاملية يفضل استخدام طريقة الانحدار ولذلك اضغط على Anderson - Rubin لأنها تتضمن معلومات لحساب الدرجات العاملية.

19. اضغط Continue لتعود إلى شاشة الامر الرئيسية.

20. اضغط على اختيار Options تظهر الشاشة الفرعية:



أ. التعامل مع البيانات المفقودة Missing data : حيث يعطى البرنامج بدائل للتعامل مع البيانات المفقودة سواء استبعاد الحالة أو تقدير قيمة لها. فإذا كان توزيع البيانات الغائبة غير اعتدالي وحجم العينة صغير فعليك تقدير قيمة البيانات المفقودة ويقوم برنامج SPSS

بإستبدالها بمتوسط درجات المتغير و هذا البديل يؤدي إلى الحصول على تباين منخفض للمتغيرات التي بها بيانات مفقودة و يؤدي إلى نتائج غير دالة إحصائيًا وعليه اضغط على Replace with Mean ولكن إذا كانت البيانات المفقودة عشوائية فيجب استبعادها من التحليل وبرنامج SPSS أما أن يستبعد الحالة كلها التي لها بيانات مفقودة على أي متغير من التحليل وعليه اضغط على Exclude Case list wise أو استبعاد الحالة التي لا بيانات غائبة على المتغير الواحد فقط و إبقائها في التحليل للمتغيرات كاملة البيانات وعليه اضغط على Pair-Wise وهذا من شأنه الحصول على تحليلات إحصائية بأحجام عينات مختلفة و هذا غير مرغوب و لكن

يمكن استخدامها إذا كان حجم العينة صغير ولذلك اضغط على اختبار List-wise هذا يمثل البديل المحافظ للتعامل مع البيانات الغائبة.

ب. مربع Coefficient display format وتتضمن بديلين هما:

- Sorted by size وهى تعنى أن يتم عرض تشبعات العوامل فى ترتيب تنازلى من الأكبر إلى الأصغر على العامل اضغط على هذا البديل.

- والبديل الثانى Suppress absolute value less وبرنامج SPSS يحدد 0.1 لظهور قيمة التشبع مطبوعة فى المخرج ولكن يمكن اعطاء البرنامج رسالة وهى عرض التشبعات التى لها قيمة 0.3 فأكبر أو 0.40 فأكبر وينصح Field (2009) بوضع القيمة 0.4 لقبول التشبع على العامل اضغط على هذا الاختيار وغيرها الى 0.40.Absolute value below

21. اضغط Continue لترجع إلى شاشة الامر الرئيسية ثم اضغط OK لتنفيذ الامر.

ثالثاً: تفسير المخرج:

- الإحصاء الوصفى اعطى البرنامج الآتى:

```
FACTOR
/VARIABLES A1 A2 A3 D1 D2 D3
/ANALYSIS A1 A2 A3 D1 D2 D3
/PRINT UNIVARIATE INITIAL CORRELATION SIG DET KMO INV REPR
AIC EXTRACTION ROTATION FSCORE
/FORMAT SORT BLANK(4)
/PLOT EIGEN
/CRITERIA MINEIGEN(1) ITERATE(25)
/EXTRACTION PC
/CRITERIA ITERATE(25)
/ROTATION VARIMAX
```

Descriptive Statistics

	Mean	Std. Deviation ^a	Analysis N ^a	Missing N
A1	3.2222	1.20185	9	0
A2	3.4444	1.42400	9	0
A3	2.8889	.78174	9	0
D1	3.2222	1.20185	9	0
D2	3.0000	1.11803	9	0
D3	2.7778	1.48137	9	0

اعطى المتوسطات والانحرافات المعيارية وحجم العينة لكل متغير وهذا الجدول لا يعطى حرية كاملة للباحث لتفسيره.

• واعطى البرنامج مصفوفة الارتباط

a. For each variable, missing values are replaced with the variable mean.

كالآتي:

Correlation Matrix^a

	A1	A2	A3	D1	D2	D3
Correlation A1	1.000	.738	-.503-	.221	.279	-.250-
A2	.738	1.000	-.399-	.300	.393	-.540-
A3	-.503-	-.399-	1.000	-.370-	-.429-	.408
D1	.221	.300	-.370-	1.000	.651	-.741-
D2	.279	.393	-.429-	.651	1.000	-.528-
D3	-.250-	-.540-	.408	-.741-	-.528-	1.000
Sig. (1-tailed) A1		.012	.084	.284	.234	.259
A2	.012		.144	.216	.148	.067
A3	.084	.144		.164	.125	.138
D1	.284	.216	.164		.029	.011
D2	.234	.148	.125	.029		.072
D3	.259	.067	.138	.011	.072	

a. Determinant = .039

وهى مصفوفة قطرها الواحد الصحيح والجزء الأعلى يمثل قيم الارتباطات بين المتغيرات الستة فمعامل الارتباط بين A_1 و A_2 0.738 فى النصف الأسفل اعطى الدلالة الإحصائية للارتباط حيث $P = 0.012$ وهى دالة إحصائية عند 0.05 لاختبار ذو ذيل واحد، اما معامل الارتباط بين A_1 و A_3 هى 0.503 وقيمة $p=0.084$ و بالتالى فإنها غير دالة إحصائية على الرغم أنه معامل ارتباط كبيرالا أن عدم الدلالة الإحصائية ترجع إلى صغر حجم العينة ولاحظ أن الارتباط بينهما سالب أى عكسى على الرغم أن المتغيرين يمثلوا بعد القلق وهذا يعود إلى أن صياغة A_3 ايجابية بينما صياغة A_2 سالبة وهذا يشير إلى أهمية إعادة ترميز استجابات المتغير A_3 و الارتباطات المعقولة أو المتوسطة هى مناسبة للتحليل العاملى وليست العالية

والارتباطات المنخفضة التى تقترب من الصفر 0.07 ، 0.01 لا يفضل استخدامها فى أسلوب التحليل العاملى ويفضل استبعادها وبنظرة سريعة على قيم معاملات الارتباطات نلاحظ أنها فى معظمها 30. فأكثر واعطى البرنامج محدد المصفوفة 0.039 والاهم أن لا يكون محدد المصفوفة سالب لأنه يعطى نتائج غير موثوق فيها ويفضل أن يزيد محدد المصفوفة عن (Field, 2009) 0.00001، ولم تزيد أى ارتباطات فى المصفوفة عن 0.80 وهذا يدل على عدم وجود قضية التلازمة الخطية المتعددة و عليه فالمصفوفة مرضية لعملية التحليل

• اعطى البرنامج مقلوب المصفوفة R^{-1} كالآتى:

Inverse of Correlation Matrix						
	A1	A2	A3	D1	D2	D3
A1	3.138	-2.608-	.967	-.889-	.354	-1.491-
A2	-2.608-	3.802	-.499-	1.364	-.673-	2.261
A3	.967	-.499-	1.653	-.192-	.436	-.615-
D1	-.889-	1.364	-.192-	3.291	-1.264-	2.364
D2	.354	-.673-	.436	-1.264-	1.999	-.333-
D3	-1.491-	2.261	-.615-	2.364	-.333-	3.675

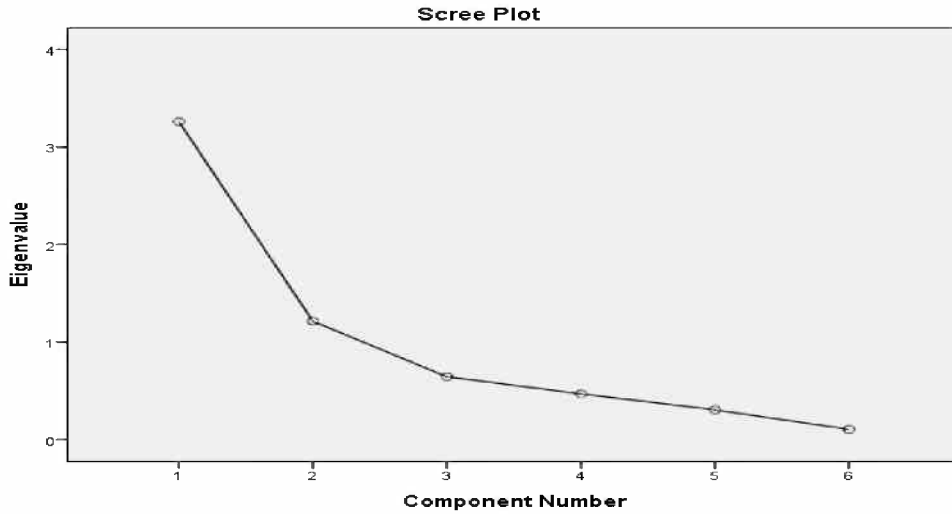
- اعطى البرنامج اختبار KMO , Bartlett's للتحقق من مسلمات التحليل العاىلى وهى كالاتى :

KMO and Bartlett's Test

Kaiser-Meyer-Olkin Measure of Sampling Adequacy.	.578	ويتضح أن قيمة
Bartlett's Test of Sphericity	Approx. Chi-Square	KMO هى
	Df	0.578 اى أنها زادت
	Sig.	عن 0.50 كما اوصى
		بها Hair et al.

(1995) وتدل على مناسبة مصفوفة الارتباط لإجراء التحليل العاىلى وعلى الرغم أنه من الافضل أن تكون 0.60 فأعلى وكلما زادت حجم العينة زادت قيمة KMO. كما أعطى اختبار Bartlett's وعدم الدلالة الإحصائية يعنى أن معاملات الارتباط فى المصفوفة فى المجتمع تساوى صفر وهذا يؤكد على عدم استخدام مصفوفة الارتباط فى التحليل ولكن عدم الدلالة الإحصائية ترجع إلى صغر حجم العينة وكذلك يكون اختبار هام إذا تم تمثيل المتغير بخمس حالات وهذا لم يتحقق فى البيانات المحللة

- اعطى البرنامج Scree plot كالاتى:



ويتضح من الشكل أنه يوجد عاملين زادت القيم الكامنة لهم عن الواحد الصحيح حيث تم وضع المكونات (العوامل) على المحور الأفقي والقيم الكامنة على المحور الرأسى وعليه فإن الباحث يقدم تفسيراته للمقياس فى ضوء عاملين فقط فى ضوء محك Scree plot

- اعطى البرنامج Anti- image- metrices سواء للارتباطات أو للتغايرات كالاتى:

		Anti-image Matrices					
		A1	A2	A3	D1	D2	D3
Anti-image Covariance	A1	.319	-.219-	.186	-.086-	.057	-.129-
	A2	-.219-	.263	-.079-	.109	-.089-	.162
	A3	.186	-.079-	.605	-.035-	.132	-.101-
	D1	-.086-	.109	-.035-	.304	-.192-	.195
	D2	.057	-.089-	.132	-.192-	.500	-.045-
	D3	-.129-	.162	-.101-	.195	-.045-	.272
Anti-image Correlation	A1	.487 ^a	-.755-	.425	-.277-	.142	-.439-
	A2	-.755-	.512 ^a	-.199-	.386	-.244-	.605
	A3	.425	-.199-	.722 ^a	-.082-	.240	-.249-
	D1	-.277-	.386	-.082-	.571 ^a	-.493-	.680

D2	.142	-.244-	.240	-.493-	.739 ^a	-.123-
D3	-.439-	.605	-.249-	.680	-.123-	.551 ^a

a. Measures of Sampling Adequacy(MSA)

وهذه المصفوفة تمدنا بمعلومات مشابهة كما تمدنا به KMO حيث يمكن تقدير KMO لمتغيرات متعددة معا و قيمة KMO لكل متغير على حدة يتم تقديرها من خلال الخلايا القطرية لمصفوفة الارتباط Anti-image matrices حيث كل القيم القطرية يجب أن تزيد عن 0.50 لتحديد صلاحية المتغير لتضمينه في التحليل العاملي والقيم أقل من 0.50 تلقى بظلالها على التحليل حيث يفضل استبعاد هذا المتغير من التحليل وإجراء التحليل بدونة و ملاحظة الفروق حيث استبعاده يؤثر على قيمة KMO وكذلك على قيمة مصفوفة الارتباط التخيلية anti-imag والخلايا القطرية هي معامل الارتباط الجزئي Partial correlation بين المتغيرات وبفحص القيم القطرية في المصفوفة السابقة نلاحظ أن قيم الخلايا القطرية زادت عن 0.50 ما عدا المتغير A₁ حيث كانت 0.487. من المفترض استبعاده من التحليل ولكن قد يعود ذلك إلى صغر حجم العينة وعليه فسوف نبقى عليه في التحليل كمثال توضيحي فقط .

- اعطى البرنامج المصفوفة المشتقة Reproduced correlations من نموذج التحليل العاملي وكذلك مصفوفة البواقي وهي خارج قسمة مصفوفة الارتباط لبيانات العينة والمصفوفة المشتقة من قبل نموذج التحليل العاملي وهي كالآتي:

		Reproduced Correlations					
		A1	A2	A3	D1	D2	D3
Reproduced Correlation	A1	.882 ^a	.808	-.586-	.141	.266	-.297-
	A2	.808	.789 ^a	-.615-	.330	.413	-.450-
	A3	-.586-	-.615-	.516 ^a	-.418-	-.451-	.486
	D1	.141	.330	-.418-	.853 ^a	.745	-.788-
	D2	.266	.413	-.451-	.745	.675 ^a	-.716-
	D3	-.297-	-.450-	.486	-.788-	-.716-	.760 ^a
Residual ^b	A1		-.070-	.083	.080	.014	.047
	A2	-.070-		.216	-.029-	-.020-	-.090-
	A3	.083	.216		.049	.022	-.078-
	D1	.080	-.029-	.049		-.094-	.046
	D2	.014	-.020-	.022	-.094-		.188
	D3	.047	-.090-	-.078-	.046	.188	
Extraction Method: Principal Component Analysis.							
a. Reproduced communalities							
b. Residuals are computed between observed and reproduced correlations. There are 8 (53.0%) nonredundant residuals with absolute values greater than 0.05.							

الملاحظ أن النصف العلوى يعكس المصفوفة التى استهلكها نموذج التحليل العاملى والنصف الأسفل يعكس مصفوفة البواقيResidual إذا نظرنا إلى مصفوفة الارتباطات بين المتغيرات الستة ونلاحظ أن معامل الارتباط التى استهلكها النموذج بين A₁ و A₂.808 ومعامل الارتباط بيرسون بين A₁ و A₂0.738 وعليه فإن مصفوفة البواقي بين A₂, A₁ هى:

$$=0.738 - 0.808= - 0.070$$

وهى قيمة معامل الارتباط فى مصفوفة البواقي و نلاحظ أن التحليل العاىلى استهلك كل معامل الارتباط بين A1, A2 بل واكثر. ويفضل أن تكون قيم معاملات الارتباطات فى مصفوفة البواقي صغيرة جدًا بحيث لا تزيد القيم عن 0.10 ويرى البعض أن لا تزيد فيها عن 0.05 (Field, 2009).

واعطى البرنامج رسالة مفادها أن 53% من قيم الارتباطات فى مصفوفة البواقي زادت عن 0.05 وهذا مؤشر غير جيد لمطابقة النموذج وهذا يدل على أن النموذج لم يستهلك كل معاملات الارتباط فى المصفوفة وهذا مؤشر على عدم ملائمة ومطابقة نموذج التحليل العاىلى.

Communalities		
	Initial	Extraction
A1	1.000	.882
A2	1.000	.789
A3	1.000	.516
D1	1.000	.853
D2	1.000	.675
D3	1.000	.760

Extraction Method: Principal Component Analysis.

• اعطى البرنامج الشيوخ Communalities

كالآتى:

الشيوخ هو نسبة التباين العام للمتغير بمعنى أن العوامل فسرت كم فى المائة من تباين المتغير وقبل إجراء الاستخلاص فإن الشيوخ لكل المتغيرات (100%) تحت عمود Initial بينما

قيم الشيوخ تحت عمود Extraction تعكس التباين المرتبط بـ A1 وهو تباين مشترك مع كل العوامل و على ذلك فإن العوامل فسر 88.2% من تباين المتغير A1 و 51.6% من تباين المتغير A3 وهكذا ولأحجام عينات صغيرة يفضل أن تكون قيم الشيوخ عالية حيث فى حالة أقل من 30 متغير فى التحليل فيجب أن يزيد الشيوخ 0.70 و إذا كان حجم العينة 250 فإن متوسط قيم الشيوخ تكون أكبر من 0.60 و إذا تم تقدير متوسط قيم الشيوخ و هى كالآتى:

$$\frac{0.882 + 0.789 + 0.516 + 0.853 + 0.675 + 0.760}{6}$$

فإنها تزيد عن 0.60 وهذا مناسب للتحليل.

الملاحظ أن قيم الشيوخ لـ A1 قبل الاستخلاص 1.0 لان كل العوامل الممكنة فسرتة بينما بعد الاستخلاص تم استبعاد بعض العوامل التى لم يتحقق فيها محك

kiaser وكذلك ينقص قيمة تباين المتغير المفسر ولذلك فإن العوامل المتبقية عليها فى التحليل لا تستطيع تفسير كل التباين الموجود فى البيانات.

- اعطى البرنامج جدول Total variance explained كالاتى:

Total Variance Explained									
Component	Initial Eigenvalues			Extraction Sums of Squared Loadings			Rotation Sums of Squared Loadings		
	Total	% of Variance	Cumulative %	Total	% of Variance	Cumulative %	Total	% of Variance	Cumulative %
1	3.260	54.337	54.337	3.260	54.337	54.337	2.374	39.563	39.563
2	1.214	20.228	74.565	1.214	20.228	74.565	2.100	35.002	74.565
3	.645	10.748	85.313						
4	.470	7.828	93.141						
5	.306	5.101	98.242						
6	.105	1.758	100.000						

Extraction Method: Principal Component Analysis.

وهذا الجدول مكون من ثلاثة اجزاء كالآتى:

١. الجزء الأول Initial Eigen values وهى تمثل القيم الكامنة للعوامل المبدئية قبل الاستخلاص والملاحظ أنه يوجد ستة عوامل هى تساوى عدد المتغيرات حيث القيم الكامنة للعامل الأول 3.260 وأقلهم للعامل السادس 0.105 وهذه العوامل فسرت 100% من تباين المتغيرات لذلك كانت الشيوخ للمتغيرات واحد والملاحظ أن العامل أو المكون الأول فسر 54.337 من تباين المتغيرات الستة بينما فسر الثانى 20.22% وهكذا و كلها فسرت 100% من تباين المتغيرات

ب . الجزء الثانى مجموع مربعات التشبعات والملاحظ أن البرنامج استبقى على العاملين الأول والثانى فقط بعد أن طبق قاعدة كاي $Ev > 1$ ولذلك فإن القيم الكامنة للعامل الأول والثانى زادت عن الواحد الصحيح وبقية العوامل انخفضت عن الواحد الصحيح، وحيث القيمة الكامنة للعامل الأول 3.26 وفسرت 54.337 من تباين المتغيرات والملاحظ أن العاملين فسر 7456% من تباين المتغيرات

ج . الجزء الثالث القيم الكامنة بعد التدوير ويتضح وجود اختلاف فى القيم الكامنة (total) حيث يتم توزيعها بالتساوى تقريباً على العاملين وكذلك فإن العامل الأول فسر 54.33% من التباين قبل التدوير بينما فسر 39.56% بعد التدوير وهذا النقصان تم وضعه فى تباين العامل الثانى وعلى ذلك لا اختلاف فى البناء العاقل قبل التدوير وبعده ولكن الاختلاف فى القيم الكامنة والتباين المفسر .

• ثم اعطى البرنامج مصفوفة المكونات Component matrix وهى تمثل تشبعات المتغيرات بالعوامل وهى كالتالى:

Component Matrix ^a		
	Component	
	1	2
D3	-.802-	
A2	.759	.462
D1	.757	-.529-
D2	.749	
A3	-.691-	
A1	.657	.672

Extraction Method: Principal Component Analysis.
a. 2 components extracted.

الملاحظ أن المتغيرات الستة تشبع على العامل الأول حيث زادت قيمة التشبع عن 0.40 وهى التى تم تحديدها مسبقاً وعليه فالعامل الأول فسر معظم تباينات المتغيرات وتشبع على العامل الثانى متغيرات A₁, A₂, D₁ وهذه تشبعات مزدوجة والمتأمل لهذه النتائج يمكن القول إن المقياس أحادى البعد حيث تولد العامل العام.

ولكن اثناء بناء المقياس ربما تبنى

الباحث بعدين هما بعد القلق وبعد الاكتئاب وعليه فإن نتائج التحليل العاقل غير متسقة مع الاطار النظرى و لكن على الباحث أن لا يلجأ إلى التفسير الا بعد إجراء التدوير حيث يساعد فى عملية التفسير ويكسب النتائج معنى سيكولوجى أو نظرى ولكن علينا أن لا نتجاهل نتائج التحليل قبل التدوير لان على الباحث أن يكون على علم بحقيقة مهمة هوانه يجرى تحليل عاملى استكشافى فليترك للبرنامج أن يقوم بهذه المهمة.

• اعطى البرنامج مصفوفة المكونات (العوامل) بعد التدوير هى كالتالى:

Rotated Component Matrix ^a		
	Component	
	1	2
D1	.918	
D3	-.829	
D2	.786	
A1		.938
A2		.847
A3		-.603
Extraction Method: Principal Component Analysis. Rotation Method: Varimax with Kaiser Normalization. a. Rotation converged in 3 iterations.		

كما ترى حدث تغير لتشبع المفردات على العاملين فتشبع المتغيرات D_1 , D_2 , D_3 على العامل الأول وهذا يمثل بعد الاكتئاب ولاحظ أن تشبع D_3 سالب لأنه ذو صياغة سالبة بينما تشبعت المتغيرات أو المفردات A_1 , A_2 , A_3 وهذا يمثل بعد القلق على العامل الثاني وعلى ذلك فإن نتائج التحليل العاملى تدعم البناء

النظري للباحث وكما هو ملاحظ قام البرنامج بعرض تشبعت المفردات ورتبها تنازلياً من الأكبر إلى الأصغر عليك أن تقارن بين مصفوفة العوامل قبل التدوير وبعد التدوير والملاحظ أن الحلول بعد التدوير أعطيت معنى سيكولوجياً للبناء.

Factor Analysis: Rotation

Method

☐ None ☐ Quartimax

☐ Varimax ☐ Equamax

☐ Direct Oblimin ☒ Promax

Delta: 0 Kappa: 4

Display

☒ Rotated solution ☐ Loading plot(s)

Maximum Iterations for Convergence: 25

Continue Cancel Help

وإذا تم إجراء تدوير مائل هل تتغير النتائج: وبإجراء تغير فى الأمر السابق بجعل طريقة التدوير من النوع المائل تغير فى الأمر السابق بجعل طريقة التدوير من النوع المائل مثل Promax:

فالاختلاف فى النتائج يكون لمصفوفة العوامل بعد التدوير حيث يعطى مصفوفتين هي Pattern Matrix كالآتى :

Pattern Matrix ^a		
	Component	
	1	2
D1	.970	
D3	-.832	
D2	.793	
A1		1.000
A2		.848
A3		-.553
Extraction Method: Principal Component Analysis. Rotation Method: Promax with Kaiser Normalization. a. Rotation converged in 3 iterations.		

وتمثل قيم التشبعات معامل الانحدار المعياري بينما المصفوفة البنائية تأخذ في حسابها العلاقات بين العوامل و تتضمن معاملات الارتباط بين العوامل والمفردات:

Structure Matrix		
	Component	
	1	2
D1	.917	
D3	-.868-	-.443-
D2	.820	.404
A1		.927
A2	.452	.885
A3	-.517-	-.673-
Extraction Method: Principal Component Analysis.		
Rotation Method: Promax with Kaiser Normalization.		

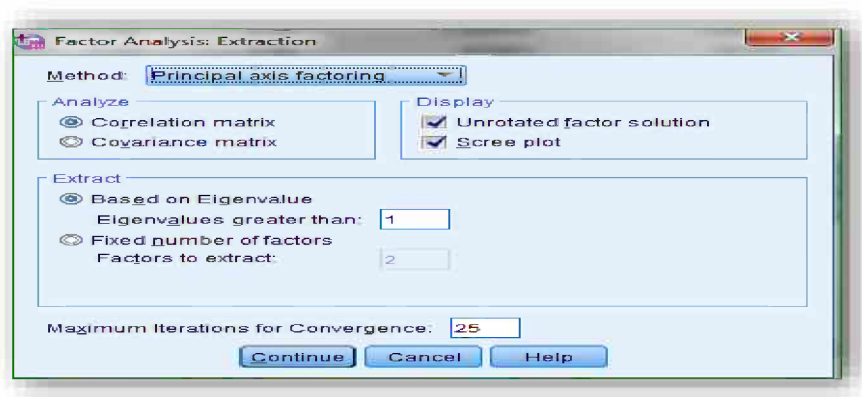
ففى التدوير المتعامد فإن هاتين المصفوفتين لهم نفس القيم والمصفوفة المستهدفة فى التدوير المائل تناظر مصفوفة العوامل فى التدوير المتعامد. ومعظم الباحثون يفسرون المصفوفة المستهدفة لأنها أكثر بساطة ويوصى بعرض المصفوفتين ففى المصفوفة المستهدفة يعطى تقريبا نفس تشبعات المفردات على العوامل كما فى مصفوفة العوامل فى التدوير المتعامد بينما فى المصفوفة البنائية تختلف حيث يتشبع على العامل الأول متغيرات D1, D2 ,D3 وكذلك A3 , A2 , A1 و هكذا بالنسبة للعامل الثانى تتشبع عليه متغيرات A3 , A2 , A1 بالإضافة إلى D3 ,D1, D2 وعلى ذلك يصبح البناء أكثر تعقيداً وتوجد صعوبة فى التفسير وهذا حدث لوجود علاقة ارتباطية بين العامل الأول والعامل الثانى وهى كالآتى:

Component Correlation Matrix		
Component	1	2
1	1.000	.432
2	.432	1.000
Extraction Method: Principal Component Analysis.		
Rotation Method: Promax with Kaiser Normalization.		

وهذه العلاقة تساوى 0.430 مما يشير إلى وجود تداخل بين العاملين وهذا يسبب وجود تشبع لمفردات A وعلى العامل الاكتئاب وتشبع لمفردات D على القلق وهذا يدل على وجود ارتباطات داخلية بين بعدى القلق والاكتئاب وليس مستقلين ولذلك من الانسب الاعتماد على التدوير المائل وليس المتعامد

وعلى ذلك نتائج التدوير المتعامد لا يمكن الوثوق فيها مثل نتائج التدوير المائل فى حالة وجود ارتباط بين العوامل.

- وإذا تم إجراء التحليل العاملى باستخدام طريقة اخرى و هل تتغير النتائج؟
- بإجراء الامر باستخدام طريقة المحاور الأساسية وذلك باتباع الخطوات السابقة ماعدا تحديد طريقة الاستخلاص Principal –Axis factoring وطريقة التدوير Varmix:



بكل تأكيد لم تتغير النتائج فيما يخص البنية العاملية قبل وبعد التدوير وهى كالاتى:

Factor Matrix ^a		
	Factor	
	1	2
D3	-.759-	
D1	.755	-.509-
A2	.717	
A1	.689	.686
D2	.669	
A3	-.586-	

Extraction Method: Principal Axis Factoring.
a. Attempted to extract 2 factors. More than 25 iterations required. (Convergence=.003). Extraction was terminated.

Rotated Factor Matrix ^a		
	Factor	
	1	2
D1	.904	
D3	-.776-	
D2	.666	
A1		.970
A2		.732
A3		-.458-

Extraction Method: Principal Axis Factoring.
Rotation Method: Varimax with Kaiser Normalization.
a. Rotation converged in 3 iterations.

ولكن إذا تم استخدام التدوير المائل Promax مع طريقة PAF نلاحظ اختلاف النتائج بعد التدوير حيث تم اعطاء مصفوفة عوامل لم يتشبع فيها المتغير A3 مع اى من العوامل :

Pattern Matrix ^a		
	Factor 1	Factor 2
D1	.955	
D3	-.780	
D2	.663	
A1		1.027
A2		.707
A3		
Extraction Method: Principal Axis Factoring. Rotation Method: Promax with Kaiser Normalization. a. Rotation converged in 3 iterations.		
Structure Matrix		
	Factor 1	Factor 2
D1	.904	
D3	-.815	-.419
D2	.706	
A1		.964
A2	.478	.782
A3	-.476	-.527
Extraction Method: Principal Axis Factoring. Rotation Method: Promax with Kaiser Normalization.		

كتابة نتائج التحليل العاىلى وفقاً لـ APA

تم إجراء تحليل المكونات الرئيسية لـ 6 مفردات لمقياس تقدير الذات للحالة المزاجية مع التدوير المتعامد باستخدام فارمكس وتم حساب مقياس Kaiser Meyer - olkin (KMO) للتحقق من مناسبة المصفوفة الارتباط للتحليل و كانت قيمته = 0.578 وقيمته لكل المتغيرات زادت عن 0.50 ماعدا المتغير A1 وهذا يدل على مناسبة المصفوفة للتحليل و المتغيرات قابلة للتحليل و كانت قيمة اختبار Bartlett's test of Sphercity هي (P>0.05) $\chi^2(15)=16.803$ وهى غير دالة إحصائياً مما تشير إلى أن العلاقات أو الارتباطات بين المتغيرات أو المفردات غير كبيرة بدرجة كافية لإجراء التحليل باستخدام PCA ولكن هذا يرجع إلى صغر حجم العينة و تم استخلاص عاملين زادت القيم الكامنة لهم عن الواحد الصحيح و فسروا معاً 74.65% من تباين المتغيرات و اوضح Scree plot قبول عاملين لتفسير تباينات المفردات فيما يلى الجدول الذى تبين تشبع المفردات على العوامل بعد التدوير:

المفردات	القلق	الاكتئاب
A ₁	.928	
A ₂	.847	
A ₃	-.603	
D ₁		.819
D ₂		-.829
D ₃		.786
القيم الكامنة	2.374	2.100
التباين المفسر	39.536%	35.002

الفصل التاسع

التحليل العااملى التوكىدى

Confirmatory Factor Analysis

يعتبر نموذج التحليل العااملى التوكىدى (CFA) Confirmatory Factor analysis أحد صور أو تطبيقات نمذجة المعادلة البنائية، وهو يهتم بدراسة العلاقات بين المتغيرات المقاسة أو المؤشرات (مفردات أو درجات الاختبار) والمتغيرات الكامنة أو العوامل. والمبدأ الأساس لـ CFA قائم على التحقق أو التأكد من الفروض (بناء معروف أبعاده أو عوامله بمفرداته مسبقاً)، عكس التحليل العااملى الاستكشافى Exploratory Factor analysis (EFA) القائم على اكتشاف طبيعة البناء العااملى (تحديد عوامله المفردات التى تنشعب عليه). وعليه فإن تحليل CFA يتم فى ضوء وجود تصور مسبق لطبيعة البناء بعدد عوامله المفترضة وبتشبعاته وذلك فى ضوء نظرية أو إطار نظرى متماسك أو دراسات سابقة. وعموماً التحليل العااملى إستراتيجية تحليلية كمحاولة الربط بين مجموعة من القياسات بعدد أقل من الأبنية أو المتغيرات التحتية (العوامل).

الفروق بين CFA و EFA

وفيما يلى أهم الفروق بين CFA و EFA:

الجدول (1.9): أهم الفروق بين CFA و EFA.

CFA	EFA	مظهر المقارنة
نعم	لا	اختبار فروض
نعم	لا	قيود على التشبعات
نعم	لا	حلول غير معيارية
نعم	نعم	حلول معيارية
لا	نعم	تدوير عوامل
نعم	لا (يمكن إعطاء مؤشر χ^2)	مؤشرات مطابقة
أكثر بساطة وأقل معالم	أكثر تعقيداً ومعالم	البساطة
لا	نعم	درجات عاملية
MPLUS , , AMOS , LISREL وغيرها	التقليدية مثل SAS ، SPSS و Minitab وغيرها	البرامج
نعم	نعم / لا	تحديد عدد عوامل
نعم	نعم	استقلال عوامل
نعم	نعم	عوامل مرتبطة
نعم	لا	نمذجة أخطاء القياس
نعم	لا	المقارنة بين نماذج
نعم	لا	أساس نظري قوي
توكيدي	وصفي استكشافي	الهدف
قبل التحليل	بعد التحليل	تسميته العوامل
نعم	لا	يدرس تشابه البناء عبر المجموعات المختلفة

استخدام وأهداف التحليل العاُملى التوكيدى:

يمكن أن تستخدم CAF فى الأغراض الآتية:

1. تطوير أو بناء مقاييس جديدة Scale development : يستخدم لاختبار البناء العاُملى لأدوات القياس مثل الاستبيانات وذلك من خلال التحقق من عدد العوامل

المفترضة ، وكذلك تشبعات المفردات بالعوامل وهذا يسهل فى تصحيح المقياس وذلك فى ضوء أبعاده الفرعية (عوامله). وكذلك يحدد طبيعة العلاقات بين العوامل أو الأبعاد. ويرى (2006) Brown أن CAF يسهم فى تقدير ثبات المقياس وذلك لتجنب مشاكل التقدير للطرق التقليدية مثل ألفاكرونباخ.

2. صدق البناء أو المفهوم Construct Validity: المفاهيم فى العلوم النفسية متعددة المظاهر أو الأبعاد. ويمدنا CFA بأدلة عن الصدق التقارى والتمييزى، فالصدق التميزى يشير إلى أن قياسات المفاهيم المختلفة متميزة (يوجد ارتباطات منخفضة بينهما). وأشار (2006) Brown إلى أن الارتباطات العالية بين المفاهيم المختلفة مثلاً (0.85) فأكثر يشير إلى صدق تميزى ضعيف. أما الصدق التقارى يشير إلى أن القياسات المختلفة لنفس المفهوم ترتبط ارتباطاً عالياً. **ولفحص الصدق العاى أو البنائى** لتحديد ما إذا كان المفهوم أحادى البعد أو متعدد الأبعاد، وكيف ترتبط الأبعاد الفرعية للمفهوم؟، وكيف ترتبط المفردات بالمفهوم أو العامل؟ فإذا كان المقياس يتكون من 40 مفردة وتتوزع على أربعة عوامل، يتشبع بكل عامل 10 مفردات ويستخدم CAF لاختبار ما إذا كانت المفردات مرتبطة بالعوامل المفترضة عليه. ولكن الكثير من الخبراء يشيرون إلى أن تشبع المفردات بالعامل هو صدق تقارى الذى يعتبر جزء من الصدق البنائى.

3. اختبار تأثيرات الطريقة Testing Method effect: تشير طبيعة العلاقات بين المتغيرات أو المفردات التى تولدت من طرق استجابة مختلفة (مفردات موجبة و مفردات سالبة) لنفس المفهوم حيث تتضمن كيفية تقديم صيغة السؤال ونوع الاستجابة المتاح. وبصفة عامة فإن الطريقة تشير إلى تأثير تحيز الاستجابة نتيجة المرغوبة الاجتماعية، فطرق الاستجابة المختلفة مثل الملاحظة، وتقدير الذات أو المفردات الموجبة فى مقابل السالبة ربما تؤدى إلى ارتباطات منخفضة بين الصور المختلفة لقياس المفهوم. فعلى سبيل المثال عندما يوجد عبارات موجبة وعبارات سالبة فإن تحليل البيانات يفرز عاملين فى حين يوجد عامل واحد متوقع فى ضوء النظرية. يعتبر مقياس تقدير الذات لـ (1965) Rosenberg مثلاً جيداً لهذا فالمقياس يتضمن عبارات سالبة وعبارات موجبة ونتائج التحليل العاى الاستكشافى أفرزت عاملين أحدهما تقدير الذات الموجب والآخر تقدير الذات السالب على الرغم أن النظرية لا تفترض عاملين، ولكن نتائج التحليل العاى التوكيدى اثبتت أن نموذج العامل العام

مع وجود عوامل الطريقة أكثر مطابقة للبيانات من نموذج العاملين (Brown, 2006).

4. تكافؤ أو ثبات القياس عبر مجموعات مختلفة **Measurement invariance**: ثبات القياس يشير إلى اتساق وتشابه أو تكافؤ بنية المقياس عبر مجموعات من الأفراد أو عبر الزمن (Brown, 2006). وعلى ذلك يهدف CFA من التأكد من تكافؤ المقياس **Equivalence** عبر مجموعات فرعية في المجتمع. وإذا لم يثبت تكافؤ أو استقرار بناء المقياس عبر المجموعات الفرعية من المجتمع فيقال أن الاختيار متحيز ويمكن التأكد من ذلك من خلال تحليل **Multi-group CFA**.

تفسير تقديرات نموذج CFA:

يمكن تفسير تقديرات معالم CFA كما أوردها (Kline, 2016) كالاتي:

1. **تشبعات العوامل**: هي تقدير التأثيرات المباشرة للعوامل على المؤشرات، وتفسر كمعاملات انحدار فلو كانت معامل التشبع غير المعياري = 5.0 فإن فروق خمس نقاط في المؤشر تعطي فرق نقطة واحدة في العامل. ووضع التشبع = 1 وذلك لمقايضة العامل.

2. **التباين المفسر للمؤشر**: تشبعات العامل المعيارية تقدر من خلال قيمة العلاقات بين المؤشر وعامله، ومربع قيمة التشبع هي تقدر نسبة التباين المفسر R^2_{smc} نتيجة للعامل، فلو أن التشبع = 0.70 فإن العامل يفسر 49% (0.49) من تباين المؤشر. والوضع المثالي أن العامل يفسر معظم تباين المؤشر ($R^2_{smc} > 0.5$). والمؤشر الذي يتشبع على عوامل متعددة فإن التشبع المعياري يفسر كأوزان بيتا B. لأنها ليس ارتباطات ولا يمكن تربيع قيمتها لاشتقاق التباين المفسر.

أما نسبة تباين خطأ القياس غير المعياري إلى التباين الملاحظ للمؤشر في العينة هي نسبة التباين غير المفسر وواحد ناقص هذه النسبة $\frac{S_e^2}{S_x^2}$ هي نسبة التباين المفسر. افترض أن تباين X_1 هو $S_x^2 = 25$ وتباين خطأ القياس الواقع عليها، $S_e^2 = 9.0$ ، إذا التباين غير المفسر $= \frac{9.00}{25.0} = 0.36$ وعلى ذلك فإن:

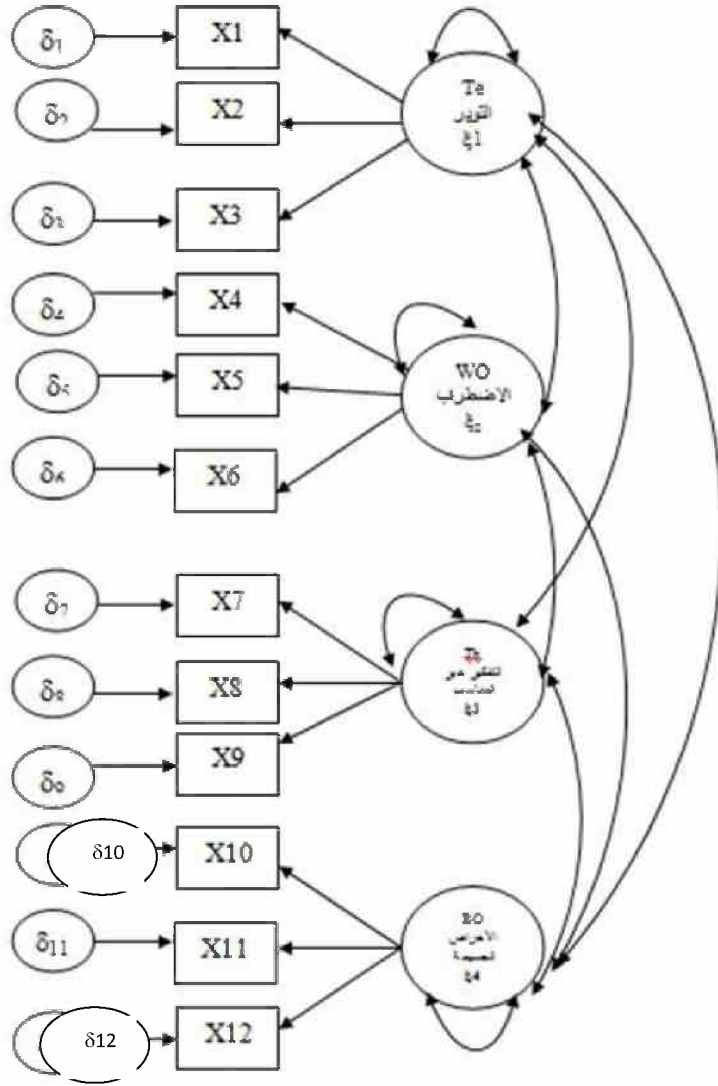
$$R^2_{smc} = 1 - 0.36 = 0.64$$

3. تقدير العلاقات بين العوامل أو أخطاء القياس: وهى تغايرات بمعنى حلول غير معيارية ولو تشبع المتغير على عامل وحيد فإن هذا التشبع يسمى معامل البناء A structure coefficient وهو تشبع معيارى.

مثال تطبيقي لتحليل نموذج CFA باستخدام LISREL (فى: Steven, 2009):

هذا المثال لبيانات مأخوذة من (1992) Benson & Bandalos وهو يتناول مقياس قلق الاختبار ويعرف بـ اختبار رود الفعل للاختبار Reactions to test scale (RTT) ويفترض أنه يقيس أربعة أبعاد وهى الاضطراب والتوتر والتفكير غير العقلانى والأعراض الجسمية.

أولاً: تخصيص النموذج: تم افتراض أن مقياس قلق الاختبار يقيس أربعة أبعاد وذلك فى ضوء تراث بحثى وذلك ما افترضه (1984) Sarason، وكما تم بناء المقياس على اعتبار أن كل بعد يمثل بـ 10 مفردات ولكن للتبسيط فى هذا العرض تم تمثيل كل متغير كامن أو عامل بثلاثة مفردات مع وجود ارتباطات بين العوامل. وفيما يلى شكل المسار عرض النموذج المقترح.



الشكل (1.9): نموذج CFA لقلق الاختبار بأبعاده الأربعة.

كما هو واضح وجود ارتباطات بين العوامل الأربعة (علاقة غير محللة)، وكذلك عدم وجود ارتباطات بين أخطاء القياس δ وهذا يؤكد على أن تباين المتغير المقاس X يفسر فقط عن طريق المتغير الكامن ξ ، ومعنى وجود علاقة بين العوامل فإن هذا يعنى وجود علاقة بين مؤشرات هذه العوامل مثل العلاقة بين X_5 و X_2 مثلاً.

ثانياً: تحديد النموذج: وهذا يرتبط بالمرحلة السابقة، وللتعرف على ماهية تحديد للنموذج، فإن عدد العناصر في مصفوفة التباين المدخلة لبيانات العينة كالتالى:

$$P = [0.5 V(V+1)] = 0.5(12)(13) = 78$$

ويتضمن عدد المعالم الحرة كالآتي = 12 تشبع عامل $\lambda + 12$ تباينات أخطاء (8) + 4 تباينات العوامل + 6 علاقات بين العوامل = 34 معلم حر.

وعلى ذلك فإن عدد المعالم الحرة أقل من عدد العناصر في المصفوفة ولذلك فإن $df = 78 - 34 = 44$ وعندئذ يقال أن النموذج فوق التحديد. ولا يتوقع حدوث اشكالية أثناء تحليله. ولكن بوضع وحدة قياس لكل متغير كامن (عامل) من خلال تثبيت تشبع X_1 على ξ_1 ، X_4 على ξ_2 ، X_7 على ξ_3 ، X_{10} على ξ_4 بالواحد الصحيح بالتالي تصبح عدد المعالم الحرة = 30 وبالتالي $df = 48$.

ثالثاً: مسح البيانات وطريقة التقدير: لم يشير الباحثان إلى كيفية التحقق من اعتدالية القياسات من X_1 إلى X_{12} ، وكذلك إلى التعامل مع البيانات الغائبة وتم تقدير الثبات باستخدام ألفا للأبعاد وتراوحته من 0.85 إلى 0.92، والثبات للمقياس ككل 0.95. وتم تطبيق المقياس على 318 طالب في مرحلة البكالوريوس والدراسات العليا، وتم استخدام طريقة ML، ويشترط حجم عينة كافية (أكبر من 200)، وبيانات من متغيرات من مستوى فترى (متصلة)، واعتدالية التوزيع، ويمكن فحص الاعتدالية البسيطة والمتدرجة والقيم المتطرفة من خلال برامج SEM مثل LISREL، EQS، AMOS ومعظم البرامج الإحصائية مثل SPSS، SAS لديها هذه الإمكانيات. وتم التعامل مع مصفوفة التباين كمدخل للبرنامج ولو تم استخدام مصفوفة الارتباط فلا بد أن تكون مقرونة بالانحرافات المعيارية حتى يستطيع البرنامج تحويل هذه المصفوفة إلى مصفوفة تباين.

إعداد ملف المدخلات للبرنامج

بعد تخصيص النموذج وفحص قضية التحديد للنموذج ومسح البيانات يتم إعداد ملف المدخلات.

إعداد ملف المدخلات في ملحق برنامج الليزرال وهو SLMPLIS كالآتي:

Title: four factor Model of RTT

Observed variables : $X_1 X_2 X_3 X_4 X_5 X_6 X_7 X_8 X_9 X_{10} X_{11} X_{12}$ or $X_1 - X_{12}$

Covariance Matrix: كما في المدخل السابق

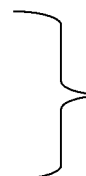
SAMPLE Size: 318

Latent variables: te wo th bo

Relationships

$X_1 X_2 X_3 (X_1 - X_3) = te$

$X_4 - X_6 = wo$



المعاري CFA

$$x_7 - x_9 = th$$

$$x_{10} - x_{12} = bo$$

$$x_1 = 1 * te$$

$$x_2 x_3 = te$$

$$x_4 = 1 * wo$$

$$x_5 x_6 = wo$$

$$x_7 = 1 * th$$

$$x_8 x_9 = th$$

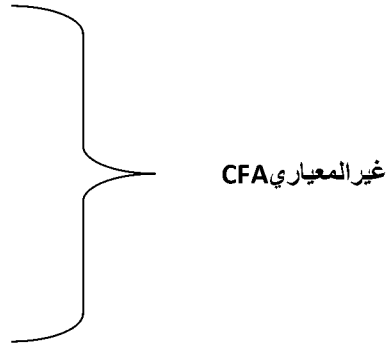
$$x_{10} = 1 * bo$$

$$x_{11} x_{12} = bo$$

$$ou : me = ml rs mi sc$$

Path Diagram

End of problem



غير المعيارى CFA

رابعًا: **تقييم النموذج**: أحد المظاهر الهامة لتقويم النموذج يحدث قبل التحليل الإحصائي من خلال بناء النموذج على أسس نظرية قوية، وعلى ذلك يكون للنموذج مقبولية نظرية ولكن بعد التحليل يتم تقييم مدى مقبولية الحلول لـ CFA فى ضوء ثلاثة مظاهر أساسية كما حددها Brown (2006) وهى كالاتى:

• **حسن المطابقة الكلية Overall goodness of fit**: كما سبق توضيحه لمؤشرات حسن المطابقة ويوصى باستخدام مؤشر χ^2 بدلالته الإحصائية بجانب مؤشر من مؤشرات المطابقة خاصة **RMSEA** أو **SRMR**، ومؤشر من مؤشرات المطابقة المقارنة خاصة **CFI** أو **NNFI**، ومؤشر من مؤشرات البساطة خاصة **PNFI** أو **AIC** أو **ECVI** لأن كلاً منهم يمدنا بمعلومات مختلفة عن مطابقة حلول نموذج CFA.

ولو أن نتائج هذه المؤشرات متسقة فيما يخص وجود مطابقة جيدة للنموذج فإن هذا يعطى تدعيم مبدئى لفكرة قبول النموذج فى الواقع الحقيقى، ولو أشارت المؤشرات إلى مطابقة سيئة فلا بد من تشخيص مصادر سوء المطابقة، وذلك من خلال إعادة تخصيص مرة أخرى. وإذا كان النموذج يعانى من سوء مطابقة فلا يجب النظر إلى تفسير تقديرات معالم النموذج مثل حجم تشبعات العوامل أو العلاقة بين العوامل وغيرها لان سوء تحديد النموذج يؤدى إلى وجود تحيز وعدم دقة فى تقديرات المعالم. ويحدث أحيانا أن يوجد عدم اتساق بين أداء مؤشرات حسن المطابقة فيمكن أن يعطى مؤشرى **CFI**، **RMSEA** مطابقة جيدة للنموذج فى حين يعطى مؤشر **AGFI** مطابقة غير جيدة للنموذج. وهنا يمكن الفصل فى ضوء مؤشر χ^2 .

أ. المطابقة لنموذج CFA المعياري (عدم اعطاء وحدة لكل متغير كامن):

فيما يلي مخرج هذا النموذج :

Four factor model of RTT
Number of Iterations = 7
LISREL Estimates (Maximum Likelihood)

LAMBDA-X				
	te	wo	th	bo
x1	0.69 (0.04) 15.59	- قيمة التّشبع - الخطأ المعياري		--
x2	0.76 (0.05) 16.01	--	--	--
x3	0.84 (0.05) 17.70	--	--	--
x4	--	0.64 (0.04) 16.18	--	--
x5	--	0.66 (0.05) 14.51	--	--
x6	--	0.67 (0.04) 16.30	--	--
x7	--	--	0.64 (0.04) 15.47	--
x8	--	--	0.67 (0.04) 16.09	--

x9	--	--	0.67 (0.04) 17.69	--
x10	--	--	--	0.38 (0.04) 10.51
x11	--	--	--	0.54 (0.05) 11.52
x12	--	--	--	0.56 (0.04) 13.29

نلاحظ أن التشبعات دالة إحصائيًا حيث زادت قيمة T عن 1.96

مصفوفات الارتباطات بين العواملφ:

PHI				
	te	wo	th	bo
te	1.00			
wo	0.55 (0.05) 11.01	1.00		
th	0.11 (0.06) 1.76	0.49 (0.05) 9.28	1.00	
bo	0.78 (0.04) 18.73	0.59 (0.05) 10.89	0.29 (0.07) 4.25	1.00

تباين أخطاء القياس (التباين غير المفسر):

THETA-DELTA

x1
x2
x3
x4
x5
x6

0.31	0.34	0.27	0.22	0.35	0.23
(0.03)	(0.04)	(0.04)	(0.03)	(0.04)	(0.03)
9.60	9.26	7.46	8.28	9.78	8.15

THETA-DELTA

x7	x8	x9	x10	x11	x12
0.27	0.25	0.16	0.26	0.41	0.27
(0.03)	(0.03)	(0.02)	(0.02)	(0.04)	(0.03)
9.31	8.66	6.55	10.68	10.10	8.49

التباين المفسر فى المتغير المقاس جراء العامل (ثبات المتغيرات):

Squared Multiple Correlations for X - Variables

x1	x2	x3	x4	x5	x6
0.61	0.63	0.72	0.65	0.56	0.66

Squared Multiple Correlations for X - Variables

x7	x8	x9	x10	x11	x12
0.61	0.64	0.74	0.36	0.42	0.54

Goodness of Fit Statistics

Degrees of Freedom = 48

Minimum Fit Function Chi-Square = 88.40 (P = .00034)

Normal Theory Weighted Least Squares Chi-Square = 86.44 (P = 0.00056)

Estimated Non-centrality Parameter (NCP) = 38.44

90 Percent Confidence Interval for NCP = (16.29 ; 68.44)

Minimum Fit Function Value = 0.28

Population Discrepancy Function Value (F0) = 0.12

90 Percent Confidence Interval for F0 = (0.051 ; 0.22)

Root Mean Square Error of Approximation (RMSEA) =0.050

90 Percent Confidence Interval for RMSEA = (0.033 ; 0.067)

P-Value for Test of Close Fit (RMSEA < 0.05) = 0.47

Expected Cross-Validation Index (ECVI) = 0.46

90 Percent Confidence Interval for ECVI = (0.39 ; 0.56)

ECVI for Saturated Model = 0.49

ECVI for Independence Model = 5.65

Chi-Square for Independence Model with 66 Degrees of Freedom = 1766.05

Independence AIC = 1790.05
 Model AIC = 146.44
 Saturated AIC = 156.00
 Independence CAIC = 1847.20
 Model CAIC = 289.31
 Saturated CAIC = 527.44
 Normed Fit Index (NFI) = 0.95
 Non-Normed Fit Index (NNFI) = 0.97
 Parsimony Normed Fit Index (PNFI) = 0.69
 Comparative Fit Index (CFI) = 0.98
 Incremental Fit Index (IFI) = 0.98
 Relative Fit Index (RFI) = 0.93

Critical N (CN) = 265.24
 Root Mean Square Residual (RMR) = .026
 Standardized RMR = 0.036
 Goodness of Fit Index (GFI) = 0.96
 Adjusted Goodness of Fit Index (AGFI) = 0.93
 Parsimony Goodness of Fit Index (PGFI) = 0.59

وكانت مؤشرات حسن المطابقة $\chi^2 = 87.811$ ($P = 0.0003$) وهي دالة إحصائية عند 0.05 مما يشير إلى رفض مطابقة النموذج (انظر محددات مؤشر χ^2) و χ^2 / df و $RMR = 0.02$, $SRMR = 0.036$ (أقل من 2) إذا النموذج جيد المطابقة و $RMSEA = 0.049$ [(90%) CI - 0.032 وكذلك 0.07 أي أنهم أقل من 0.036 و - 0.66] $CFI = 0.98$, $AGFI = 0.93$, $GFI = 0.95$, $RFI = 0.95$, $IFI = 0.98$, $NNFI = 0.98$ كل هذه المؤشرات تشير إلى مطابقة جيدة للنموذج.

Residual analysis وبجانب مؤشرات حسن المطابقة يجب فحص تحليل البواقي كما سبق الإشارة إليه يوجد ثلاثة مصفوفات مرتبطة لمخرج التحليل العاملى التوكيدى مصفوفة التباين للمقاسات S ، ومصفوفة التباين المشتقة من النموذج Σ ، ومصفوفة البواقي تعكس الفرق بين S و Σ ويمكن عرضه كالاتى:

واعطى ملخص لمصفوفة البواقي كما يلى :

Summary Statistics for Fitted Residuals

Smallest Fitted Residual = -0.06

Median Fitted Residual = 0.00

Largest Fitted Residual = 0.09

ونلاحظ أن القيمة الصغرى فى مصفوفة البواقي تساوى 0.06- والقيمة الوسطية 0.000 والقيمة العظمى 0.086 ، و إذا زادت القيمة العظمى عن 0.10 فإن هذا مؤشر على سوء مطابقة النموذج للبيانات. وعليه نلاحظ فى ضوء مصفوفة البواقي القيمة العظمى أقل من 0.10 وهذا يؤكد على المطابقة الملائمة أو المناسبة للنموذج. وعلى الرغم أن مؤشر SRMR هو يمدنا بملخص أو بمرآة عامة للفروق بين مصفوفات $(S - \sum)$ ، إلا أن مصفوفة البواقي تمدنا بمعلومات تفصيلية عن كيف كل معلم يتم إنتاجه أو اشتقاقه من مصفوفة البيانات بالتالى فإن SRMR هو مطابقة كلية بينما مصفوفة التباين للبواقي هى مطابقة لكل معلم فى النموذج وهذا يفيد فى الكشف عن مصادر سوء المطابقة فى النموذج أن وجدت. وعلى الرغم أن مصفوفة البواقي من الصعب تفسيرها وذلك لتأثيرها بالوحدة المعيارية وتباين كل متغير مقاس وبالتالى من الصعب تحديد ما إذا كانت مصفوفة البواقي مرتفعة أو منخفضة وذلك لاختلاف وحدات القياس للمؤشرات وهذه المشكلة تم حلها من خلال تقدير مصفوفة البواقي المعيارية Standardized residuals التى تقدر من خلال قسمة البواقي على الأخطاء المعيارية المقابلة لها. والبواقي المعيارية مشابهة للدرجات المعيارية للتوزيعات العينية ويمكن تفسيرها مثل الدرجة المعيارية Z.

وفيما يلى العرض البياني لمنحنى Steam plot وهى كالاتى :

Stem leaf Plot

```

- 6|2
- 4|0774
- 2|754410993300
- 0|9531099987631000000000000000
  0|123445557890667888
  2|0346992456888
  4|88
6|
  8|5

```

فى هذا الشكل يبدو أن المنحنى ذات توزيع اعتدالى مما يؤكد المطابقة الجيدة.

كما تم عرض شكل Q-plots للبواقي:

```

.
. xxx
m . . . *
a . . . * x
l . . . **
. . . x*x
Q . . . xx
u . . . *xx
a . . . *xx
n . . . x *.x
t . . . xxx
i . . . xxxxx
l . . . xxx
e . . . x
s . . . x x
. . . xx
. . . x
. . . x
. . . x
-
3.5.....
.....

```

ويتضح أن البواقى لا تتوزع اعتداليا بدرجة تامة وكل تغاير فى مصفوفة البواقى يشار إليها بالعلامة (x) والبواقى المتعددة التى لها نفس القيمة يشار إليها بالعلامة نجمة (*)، فلو كان التوزيع للبواقى قريبا من الخط المنقوط (...) فإنها تتوزع توزيعاً اعتدالياً وهذا مؤشر للمطابقة التامة، بينما فى المثال الحالى فإن المطابقة ليست جيدة وعلى ذلك توجد أخطاء فى تخصيص النموذج ولذلك يقترح البرنامج مؤشرات تعديل فى النموذج.

وتقدر مؤشرات التعديل من المعالم المثبتة مثل تشبعات العوامل المزدوجة أو تغايرات الأخطاء (الارتباطات بين الأخطاء) والمعالم المقيدة فى النموذج، أى أن مؤشر التعديل يعكس مقدار النقصان لقيمة χ^2 لو أن أحد المعالم المثبتة أو المقيدة أصبحت حرة، وتتأثر مؤشرات التعديل بحجم العينة فكلما زاد حجم العينة تزيد قيمة المؤشر ويمكن أن

لا يمثل التعديل إضافة أو تحسن جوهري في تفسير النموذج. ولكن تبدو أن التعديلات
تضيف تحسن في النموذج على الرغم من مدى محدوديتها أو تفاهتها وفيما يلي
مؤشرات التعديل للنموذج.

Modification Indices and Expected Change
Modification Indices for LAMBDA-X

	te	wo	th	bo
	-----	-----	-----	-----
x1	--	2.29	2.60	5.76
x2	--	2.41	0.06	4.45
x3	--	0.01	1.47	18.33
x4	9.58	--	0.66	8.71
x5	0.17	--	0.68	0.39
x6	11.94	--	0.00	12.36
x7	1.19	1.96	--	0.92
x8	0.50	0.00	--	0.01
x9	0.09	1.54	--	0.58
x10	0.40	0.96	0.22	--
x11	0.07	0.03	1.77	--
x12	0.12	1.03	2.62	--

Modification Indices for THETA-DELTA

	x1	x2	x3	x4	x5	x6
	-----	-----	-----	-----	-----	-----
x1	--					
x2	10.58	--				
x3	1.58	4.38	--			
x4	0.03	0.08	6.55	--		
x5	0.02	0.41	0.43	12.99	--	
x6	0.55	5.00	0.39	0.03	14.16	--
x7	1.02	0.18	0.00	0.15	0.47	3.51
x8	1.13	3.27	1.48	0.01	0.08	0.25
x9	0.33	3.18	0.10	0.64	2.15	2.45
x10	5.14	2.83	7.70	0.28	1.34	0.17
x11	1.13	0.31	4.16	1.68	0.07	0.84
x12	0.68	1.99	0.27	0.79	2.13	0.05

Modification Indices for THETA-DELTA

	x7	x8	x9	x10	x11	x12
	-----	-----	-----	-----	-----	-----
x7	--					
x8	2.61	--				
x9	0.45	1.20	--			

x10	2.79	0.33	3.49	-	-	
x11	0.01	0.24	0.36	0.90	-	-
x12	1.06	1.33	7.98	0.94	0.00	-

وفى المخرج السابق يتبين أنه بإضافة المؤشر x_3 إلى العامل الرابع bo يحدث نقصان لقيمة x^2 للنموذج بمقدار 18.3، وإضافة x_6 إلى العامل الأول te يحدث نقصان لقيمة x^2 بمقدار 11.95. وفيما يلي إعادة تخصيص النموذج فى ضوء مؤشرات التعديل حيث يتم إضافة المؤشر x_3 إلى bo و x_4 إلى bo و x_6 إلى العامل bo و x_6 إلى te وكذلك تمدنا مؤشرات التعديل بكيفية إجراء تعديل على الأخطاء الواقعة على المتغيرات المقاسة δ ومن الواضح فإن مؤشرات التعديل تقترح وجود ارتباطات بين δ_1 و δ_2 ، وبين δ_4 و δ_5 ، وبين δ_5 و δ_6 ، وبين δ_9 و δ_{12} ولإجراء هذه التعديلات يجب إعادة تخصص النموذج المفترض ويتم ذلك فى خط

Relationships

$X_1 - X_4 X_6 = te$

$X_4 - X_6 = wo$

$X_7 - X_9 = th$

$X_3 X_4 X_6 X_{10} X_{12} = bo$

Let errors $X_1 X_2$ Correlate

Let errors $X_4 X_5$ Correlate

Let errors $X_5 X_6$ Correlate

Let errors $X_9 X_{12}$ Correlate

وبإجراء التحليل بعد التعديل ظهرت تحسن أداء مؤشرات المطابقة كالاتى:

المؤشر	قبل التعديل	بعد التعديل
$\chi^2 (P)$	88.40	42.56 (0.32)
RMSEA	0.05	0.0108
90% CI for RMSEA	0.033-0.067	(0.0-0.041)
P (Clos fit)	0.0005	0.99
Ecv	0.46	0.373
χ^2 / df	88.40/48	42.56/39
AIC	146.44	118.45

الفصل العاشر

نمذجة المعادلة البنائية

Structural Equation Modeling (SEM)

أصبحت نمذجة المعادلة البنائية (SEM) Structural Equation Modeling (SEM) مكوّنًا رئيسيًا في التحليلات الإحصائية المتدرجة. وتستخدم بصورة متزايدة في العلوم التربوية والنفسية والاجتماعية والسلوكية وغيرها. وأخذت مسميات عديدة في الأدبيات البحثية منها نمذجة المعادلة البنائية (SEM)، نمذجة بنية أو بناء التغاير (Covariance Structure Modeling (CSM)، تحليل بنية التغاير (CSA)، النمذجة السببية بين المتغيرات الكامنة، تحليل المتغيرات الكامنة، نمذجة المعادلة التلازمية Simultaneous Equation Modeling، والنمذجة البنائية Structural Modeling. وأحيانًا تأخذ أسماء قريبة مثل تحليل النموذج السببي، تحليل المعادلات البنائية، ونمذجة المسارات.

وإستراتيجية SEM تعكس اتساع لمجموعة من الأساليب الإحصائية المرتبطة بالنموذج الخطى العام (GLM) General Linear Model متضمنة تحليل المسار، التحليل العاملى، التحليل التمييزى، تحليل التباين المتدرج (متعدد المتغيرات التابعة)، تحليل الانحدار المتعدد، تحليل التغاير، تحليل السلاسل الزمنية، والنمذجة النمائية الكامنة Latent Growth Modeling. وبناء عليه فهى كالمظلة التى تضم تحتها مجموعة من الأساليب الإحصائية المتدرجة التقليدية وكذلك الأساليب الأكثر تطورًا. وعلى ذلك فهى مدخل إحصائى شامل يجمع بين تحليل المسار والتحليل العاملى التوكيدى (Schumacher & Lomax, 2010). وهذه تطبيقات خاصة لنمذجة المعادلة البنائية.

مفهوم نمذجة المعادلة البنائية

تعددت تعريفات SEM فيعرفها (1995) Hoyle بأنها مدخل إحصائي شامل لاختبار فروض حول علاقات بين متغيرات مقاسه ومتغيرات كامنة. ويعرفها Hair, (1998) Anderson, Taham, & Black بأنها تكتيك إحصائي يسمح بتحليل مجموعة من المعادلات البنائية في نفس الوقت حيث يكون المتغير مستقل في معادلة وتابع في معادلة أخرى. ويعرفها (2013) Ullman & Bentler بأنها مجموعة من الأساليب الإحصائية التي تسمح بدراسة مجموعة من العلاقات بين متغير مستقل فأكثر متصل أو منفصل ومتغير تابع فأكثر متصل أو منفصل، وكلاً من المستقل والتابع متغيرات مقاسة أو كامنة. ويراهها (2012) Crockett أنها أسلوب تحليلي متدرج من الجيل الثاني يحدد إلى أي درجة بيانات العينة متطابقة مع النموذج النظري، أو تكتيك إحصائي متقدم يسمح باختبار النظريات والنماذج والبناء الكامن أو التحتى لمفهوم أو ظاهرة نظرية مجردة مثل الاتجاهات والمعارف والانفعالات تقاس بطريق غير مباشر عن طريق مجموعة من المفردات (المقاييس). ويراهها Schumacker & Lomax (2010) بأنها إستراتيجية تتضمن أنواع مختلفة من الأساليب الإحصائية لشرح العلاقات بين المتغيرات المقاسة والكامنة (المفاهيم) من ناحية، وكذلك مدى ارتباط هذه المفاهيم (الأبنية التحتية) بعضها ببعض لتحديد مدى مقبوليتها ومنطقيتها إمبيريقياً.

أهداف نمذجة المعادلة البنائية

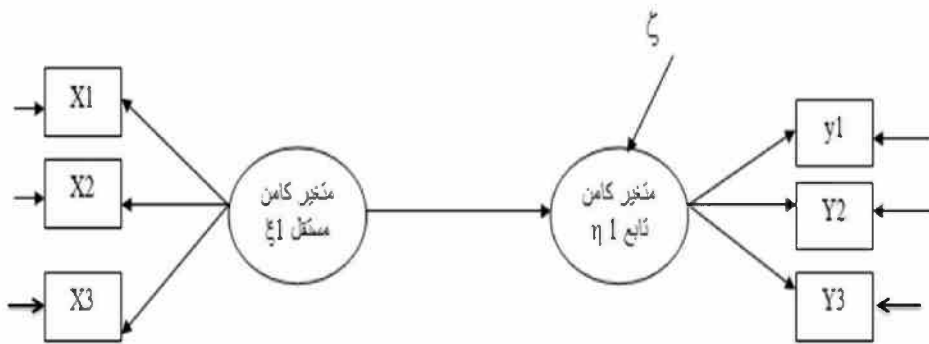
تمدنا البيانات بمصفوفة التغاير أو الارتباط (S) ثم تدخل التحليل ويعطى التحليل مصفوفة التغاير المشتقة (المحللة) (Σ) من النموذج وتسمى مصفوفة التغاير للمجتمع المقدرة. وعلى ذلك فإن السؤال الهام لـ SEM هل مصفوفة التغاير المشتقة من النموذج (المجتمع) تتقارب أو تتسق مع مصفوفة تغاير العينة (الملاحظة). وبعد الحصول على هذا التناسق أو التقارب للنموذج يتم تناول العديد من التساؤلات عن المظاهر العامة والفرعية للنموذج وهي كالآتي:

- مناسبة أو مطابقة النموذج Adequacy or fit of the model : وفيه يتم الجواب عن مدى مطابقة النموذج لبيانات العينة المتولدة من تصميمات بحثية وصفية غير تجريبية وكذلك من تصميمات تجريبية . إلى أى درجة أن النموذج النظرى يتم تدعيمه عن طريق بيانات العينة ؟
- بناء النظريات والتحقق من مقبوليتها ومطابقتها فى الواقع فى ضوء بيانات العينة وكذلك تطويرها (Hoyle & panter, 1995; Quintana & Maxwell, 1999).
- اختبار مصداقية الأبنية النظرية من خلال تأكيد البنية العاملية لأدوات قياس جديدة أو التأكد من بناء موجود فى مجتمعات جديدة، (Raykov & Marcoulides, 2006).
- المقارنة بين نماذج نظرية متنافسة أو بديلة لتحديد أيهم أكثر مطابقة مع البيانات الإمبريقية (Kline, 2016).
- اختبار أو التحقق من الظواهر النفسية المعقدة متعددة الأبعاد فى تحليل تلازمى واحد وهذا يتناسب مع طبيعة الظاهرة الإنسانية (Ullman, 2006).
- تقدير معالم النموذج مثل التشبعات والتأثيرات المباشرة وغير المباشرة وتباينات العوامل وكذلك الأخطاء المعيارية والدلالة الإحصائية لتقويم تفاصيل النموذج.
- تقدير حجم التأثير لتحديد نسبة التباين المفسر فى المتغير التابع الكامن (الداخلى) جراء المتغيرات الكامنة المستقلة (الخارجية)، أى تسمح بتقدير حجم التأثير لكل معادلة بنائية.
- دراسة التأثيرات الوسيطة من خلال تقدير التأثيرات غير المباشرة لمتغير ما على آخر من خلال متغيرات أخرى بين المتغيرين وتسمى Mediation analysis.
- دراسة الفروق بين مجموعات مختلفة فى مصفوفة التباين، وتقديرات المعالم والمطابقة. ويمكن استخدام نموذجة المعادلة البنائية لتحليل مصفوفات التباين

لمجموعات مختلفة في تحليل واحد ويطلق عليها تحليل SEM متعدد المجموعات Multi-group.

- دراسة تأثيرات التفاعلات interaction effects ، وهذا شائع في العلوم الإنسانية حيث يحدث تفاعل بين متغيرين لتوليد متغير ثالث يؤثر على الظاهرة، فتسهم منهجية SEM في كيفية توليد هذه التفاعلات ودراسة تأثيرها على المتغيرات التابعة.

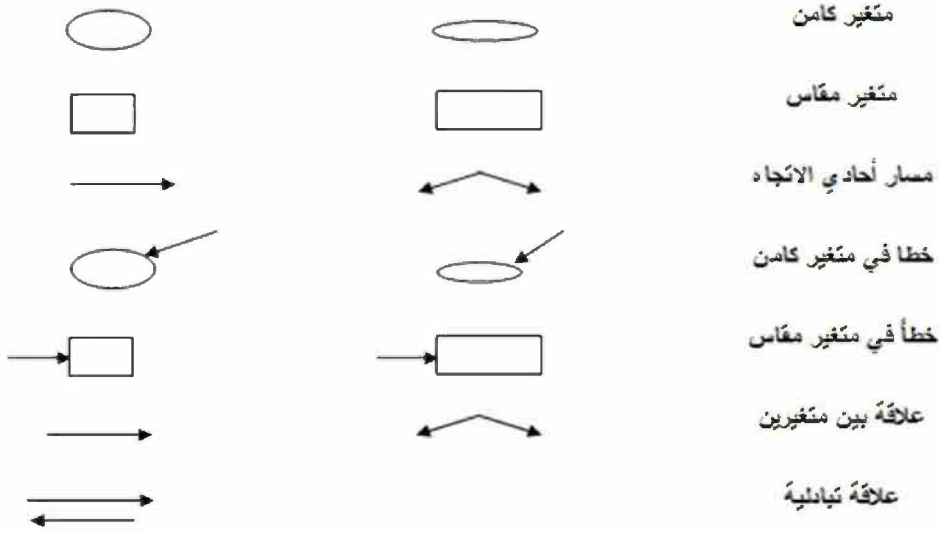
نموذج الانحدار البنائي (نموذج المعادلة البنائية) أو تحليل المسار بين المتغيرات الكامنة **Structural regression model**: يشبه التحليل العاملي التوكيدي ولكنه يفترض وجود تأثيرات سببية بين المتغيرات الكامنة ويستخدم في عدة أغراض أهمها اختبار علاقات تفسيرية (سببية) بين مجموعة من الأبنية التحتية (المتغيرات الكامنة).



الشكل (1.10): شكل المسار لنموذج المعادلة البنائية.

مكونات شكل المسار

لفهم نمذجة المعادلة البنائية لا بد من التعبير عنها في شكل المسارات. وشكل المسار هو شكل من التمثيل (تصميم) البياني للنموذج ومن خلال شكل المسار يتم ترجمة مساراته إلى رسومات ثم إلى مجموعة من المعادلات التي تحدد النموذج وعلى ذلك فهو عرض النموذج بصرياً. فشكل المسار يساعد على التواصل مع الباحثين الذين لديهم خلفيات مختلفة وفيما يلي أجزاء شكل المسارات البيانية لوصف SEM:

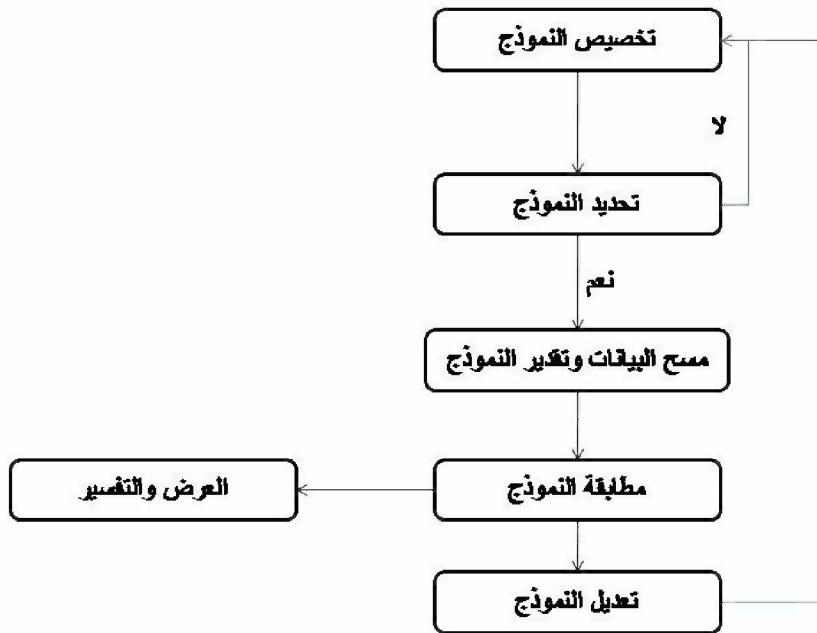


الشكل (2.10): أهم الأشكال المستخدمة لوصف نموذج SEM في شكل المسارات.

مراحل بناء نموذج المعادلة البنائية

وضع خبراء SEM تصورات عديدة فيما يخص مراحل بناء نموذج SEM، فيرى (Ullman & Bentler, 2013) أنها أربع مراحل وهي التخصيص، والتقدير، والتقويم، والتعديل، في حين يرى معظم الخبراء (Schumacker & Lomax, 2010; Kline, 2016; Bollen, 1989; Crokett, 2012) أنه لبناء نموذج المعادلة البنائية يستلزم خمسة مراحل وهي: التخصيص، والتحديد، والتقدير، واختبار المطابقة والتعديل (إعادة التخصيص)، في حين يراها Hoyle (1995) تتكون من: مراحل التخصيص، والتقدير، والتقويم، والتعديل، والتفسير.

يمكن عرضها بالآتي:



الشكل (3.10): مراحل بناء نموذج SEM .

مرحلة تخصيص النموذج Specification stage:

وفيها يتم تحديد النموذج النظرى المبدئى فى ضوء النظرية أو الأدبيات البحثية لنتائج الدراسات السابقة، وتحدد فيه المتغيرات المستقلة والتابعة، وطبيعة العلاقات بين المتغيرات المقاسة والمتغيرات الكامنة من ناحية وطبيعة التأثيرات أو العلاقات السببية بين المتغيرات الكامنة من ناحية أخرى، ويفضل أن يتم بناء النموذج المفترض فى ضوء نظرية متماسكة توضع طبيعة ديناميات العلاقات والتأثيرات بين متغيراتها. ولتخصيص أو تعيين النموذج لا بد من عرض النموذج بيانياً فى شكل مسار **Path diagram** وهو ترجمة بيانية أو عرض بصرى للنموذج النظرى يوضح العلاقات المقترحة بين المتغيرات الكامنة والمتغيرات المقاسة وكذلك التأثيرات السببية بين المتغيرات الكامنة. وقبل التعرض لشكل المسارات لا بد من عرض بعض الرموز والأشكال وهى كالتالى:

فى ضوء وظيفتها فى النموذج: المتغيرات الكامنة فى النموذج هى:

أ. متغيرات المصدر أو المستقلة أو الخارجية **Source, Independent, or Exogenous Variables**: هى المتغيرات التى يخرج منها مسارات أو أسهم أى أنها

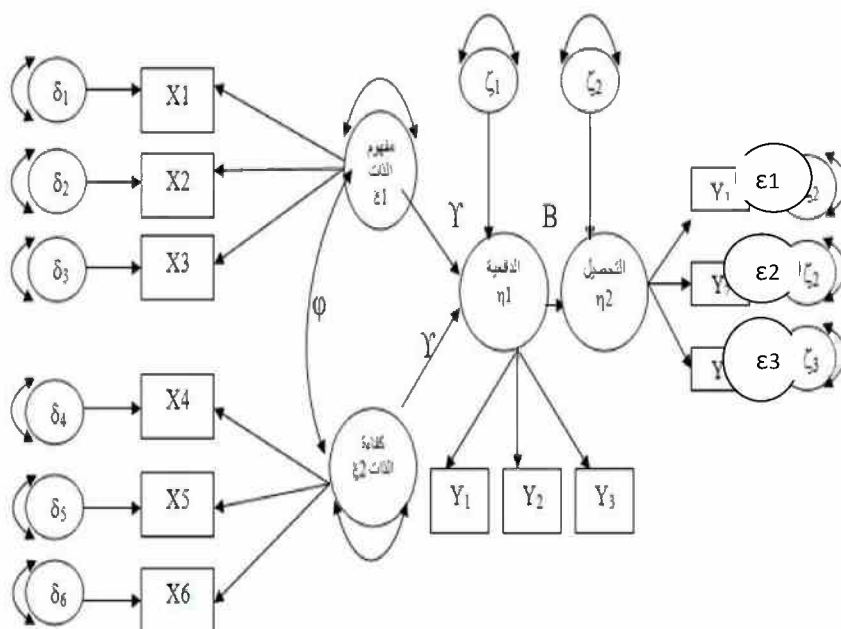
متغيرات السبب ولا يؤثر فيها متغيرات أخرى وهى تسبق فى الحدوث المتغيرات التابعة، وربما ترتبط مع بعضها البعض وعلى ذلك فهى متغيرات لا تستلم مدخلات سببية من أى متغير آخر. وأطلق عليها خارجية لأن مصادر تفسير المتغيرات المتضمنة فى تباينها يقع خارج شكل المسارات أو النموذج. وهذه المتغيرات تعكس مصادر السببية فى شكل المسار أو النموذج. فى شكل (4.10) فإن المتغيرات الخارجية أو متغيرات المصدر هى متغيرات مفهوم الذات ξ_1 (اكساي) وكفاءة الذات ξ_2 .

ب. المتغيرات التابعة أو الداخلية **Dependent or Endogenous variables** : هى

المتغيرات التى تعتمد على متغير كامن مستقل فأكثر فى النموذج أى يأتى لها مسارات، مثل متغير التحصيل يأتى له مسار من الدافعية. و فى شكل () تعتبر η_1 و η_2 متغيرات كامنة تابعة.

ج- متغيرات وسيطة **Mediator**: حيث يلعب المتغير دورين هما المستقل والتابع فى نفس

الوقت. فمتغير الدافعية هو متغير وسيط يمكن أن ينقل أثر المتغيرات الكامنة المستقلة (مفهوم الذات وكفاءة الذات) إلى التحصيل ويسمى تأثير غير مباشر. والمتغيرات الوسيطة هى متغيرات داخلية. فيما يلى نموذج معادلة بنائية مفترض بين متغيرات مفهوم الذات الأكاديمية وكفاءة الذات والتحصيل:



الشكل (4.10): نموذج SEM بين مفهوم الذات وكفاءة الذات والدافعية والتحصيل.

الشكل السابق عبارة عن مستطيلات ودوائر وأسهم بين المتغيرات، وهذا يمثل شكل المسارات ونلاحظ وجود سهم $\rightarrow$ بين المتغير الكامن المستقل الأول مفهوم الذات (ξ_1) والمتغير الكامن المستقل الثانى كفاءة الذات (ξ_2) وهو يمثل علاقة أو تغاير ويطلق عليه **علاقة غير محللة Unanalyzed Association** نتيجة حسابها عن طريق البرنامج ولا يمكن تقديرها يدوياً لعدم وجود درجات لها. وفى حالات معينة يمكن حذف هذا السهم وهذا يعنى عدم وجود علاقة بين المتغيرين (استقلالية) إذا تم افتراضه فى ضوء الإطار النظرى للظاهرة. بينما المتغيرات الداخلية التابعة التحصيل والدافعية لا ترتبط عن طريق هذا السهم. والسهم $\rightarrow$ على نفس المتغير المستقل الكامن يعكس تباين المتغير الكامن المستقل لأن مسببات المتغير الكامن المستقل (ξ) غير محددة فى النموذج ولذلك فإن ارتباطه (تغايره) أو تباينه يعتبر معلم حر، السهم $\rightarrow$ ارتباط متغير كامن مستقل مع كامن مستقل آخر $\xi_2 \xi_2$ بينما السهم $\rightarrow$

ارتباط كل متغير مقياس أو كامن مستقل مع نفسه وهذا السهم لا يوضع على المتغير الكامن التابع لأن مسبباته موجودة فى النموذج.

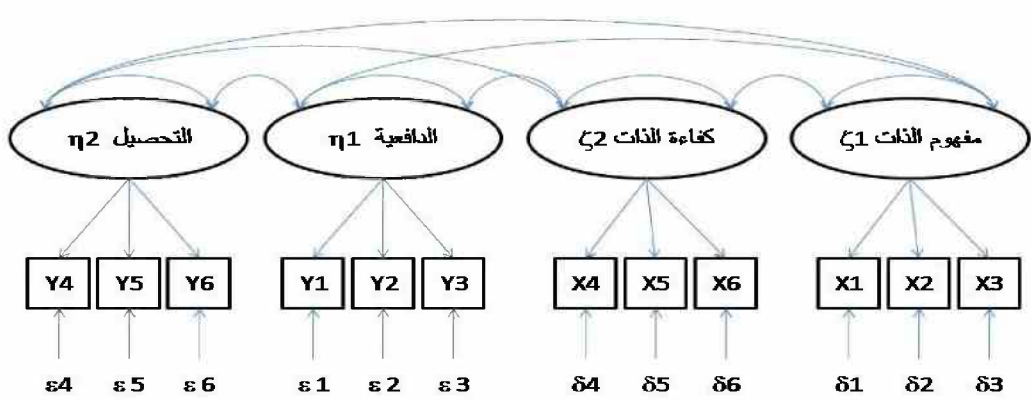
أما السهم ($\rightarrow$) يسمى تأثير أحادى الاتجاه أو مباشر والمسلمة الأساسية لهذا السهم أن التغير الذى يحدث فى ذيل السهم ينتج عنه تغير فى المتغير الذى فى رأس السهم، ويشير إلى وجود علاقة خطية بين المتغيرات المستقلة والتابعة من ناحية وبين المتغيرات الكامنة وبعضها البعض من ناحية أخرى وتسمى **علاقة سببية مباشرة Direct Causal** وهو يفترض أن يكون موجود فى نموذج SEM بين المتغيرات الكامنة والمقاسة من ناحية وبين الكامنة ببعضها البعض.

مكونات نموذج SEM

ويتكون نموذج المعادلة البنائية من جزئين كما وضحا (Hoyle, 1995; Joreskog & Sorbom, 1993) كالآتى:

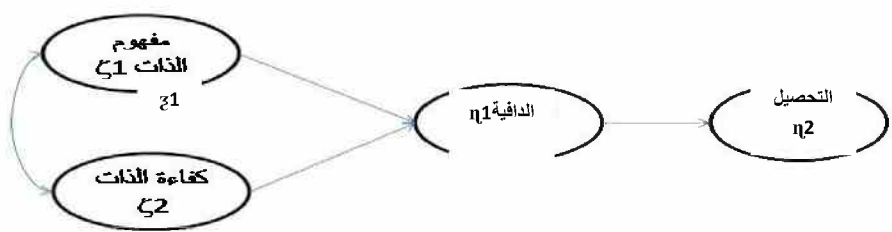
اولاً: **نموذج القياس Measurement Model**: يهتم بالعلاقات بين المتغيرات المقاسة والمتغيرات الكامنة، أو كيف يتم التعبير عن المتغيرات الكامنة فى ضوء

المتغيرات المقاسة. وهذا النموذج يعكس أسلوب التحليل العاملى التوكيدى الذى يستخدم فى التحقق من مصداقيته بنية المقاييس (Ullman, 2006)،. والسهم () يعنى الارتباط أو التغاير بين المتغيرات الكامنة المستقلة وهذا الإجراء يتشابه مع التحليل العاملى الاستكشافى عندما يستخدم طرق التدوير المائل التى تسمح بالعلاقات بين العوامل، ويفترض أن العلاقة بين المتغيرات المقاسة والكامنة خطية. ويمكن عرض النموذج المقاسة كالاتى:



شكل (5.10): النموذج المقاسة المشتقة من نموذج SEM السابق.

ثانيًا: النموذج البنائى **Structural Model**: وهو يعكس المظهر الأساسى لـ SEM وفيه توضح العلاقات أو التأثيرات السببية بين المتغيرات الخارجية الكامنة والمتغيرات الداخلية الكامنة، وهذا النموذج يعكس تحليل المسار بين المتغيرات الكامنة كما فى شكل (4.10). وعليه فإن نموذج SEM هو تكامل بين نموذج مقاس (تحليل عاملى توكيدى) ونموذج بنائى (تحليل مسار بين متغيرات كامنة). يمكن عرض النموذج البنائى كالاتى:



الشكل (6.10): النموذج البنائي المشتق من نموذج SEM السابق.

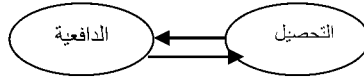
وفى نمذجة المعادلة البنائية وتحليل المسار الكلاسيكى يسمى المتغير المستقل بالمتغير الخارجى Exogenous وهو المتغير الذى يخرج منه مسارات ولا يؤثر فيه أى متغير آخر، بينما يطلق على المتغير التابع بـ Endogenous variable حيث يمتلك مسار على الأقل يودى إليه ويمكن أن تخرج منه تأثيرات أو مسارات إلى متغيرات داخلية أخرى (Bollen, 1989)، بينما فى تحليل الانحدار يطلق على المتغيرات بالمستقلة Independent والتابعة Dependent وذلك لأن تحليل الانحدار لا يتعامل مع المتغيرات الوسيطة .

عدد المتغيرات المقاسة لكل متغير كامن: يرى (Nunnly & Bernstein 1994) بأنه يجب تمثيل المتغير الكامن بثلاثة مفردات أو مؤشرات لكى يعطى حلول وتقديرات متسقة. وأشار العديد من الخبراء بأنه يجب تمثيل المتغير الكامن بثلاثة مؤشرات على الأقل أو أربعة وذلك لإعطاء نتائج ذات معنى (Bentler & Chou, 1987; Bollen, 1989) وذلك لان تمثيل المتغير الكامن بمؤشرين أو أقل يودى إلى أن النموذج يكون غير محدد (Bentler & Chou, 1987). ويرى البعض أنه يجب ألا يزيد عدد المؤشرات عن أربع (Quintana & Maxwell, 1999). فى حين يرى (Mulaik & Millsap 2000) أنه لضمان مصداقية أى بناء كامن فإنه يجب تمثيل البناء أو المتغير الكامن بأربعة مؤشرات على الأقل.

تخصيص النموذج البنائي: الأساس فى تخصيص النموذج البنائي مشابيه لتحليل المسار وهو أن العلاقات بين المتغيرات الكامنة يجب تحديد طبيعتها فيمكن أن تكون ارتباطيه كما فى نموذج التحليل العاظمى التوكيدى أو سببيه بين المتغيرات الكامنة الخارجية والكامنة الداخلية فى نموذج المعادلة البنائية، وتكون هذه العلاقة فى ضوء تأصيل واعتبارات نظرية أو منهجية. والاعتبار الأساسى فى تخصيص النموذج البنائي هو تحديد طبيعة العلاقات بين المتغيرات الكامنة وفى ضوء ذلك فإن طبيعة التأثيرات فى النموذج البنائي تكون كالاتى:

1. نماذج أحادية التأثير بين المتغيرات الكامنة **Unidirectional Causal Model or Recursive models**: وهى نماذج تحتوى على علاقات أو تأثيرات سببية أحادية الاتجاه (→) ومعظم النماذج فى الدراسات النفسية أو غيرها فى هذا الاتجاه نماذج أحادية التأثير فى اتجاه واحد.

2. نماذج تبادلية التأثير بين المتغيرات الكامنة **Non-recursive or bidirectional models**: هى النماذج التى تتضمن علاقات تأثيرية تبادلية Reciprocal أو Feedback loop بمعنى وجود تأثيرات فى اتجاهين ((. وفى هذه الحالة تكون النماذج أكثر تعقيداً وهذا يقود فى كثير من الأحيان إلى نماذج غير متطابقة مع البيانات.



الشكل (7.10): العلاقة التبادلية بين التحصيل والدافعية

فى شكل السابق يوجد تأثير سببى تبادلى بين التحصيل والدافعية، كلا منهما يؤثر فى الآخر. ويوصى (Baumgartner & Homburg 1996) بأنه فى حالة النماذج التى تتضمن علاقات تبادلية يجب تحديدها تحديداً دقيقاً لأن هذه النماذج تعاني من قضية عدم التحديد Non-identification.

النماذج البديلة أو المكافئة **Alternative Model**

يبدأ الباحث بصياغة نموذج مستهدف فى ضوء أسس نظرية ويعتقد أنه يشرح ويفسر العلاقات أو التأثيرات بين مجموعة من المتغيرات، ويمكن القول إن النموذج متولد Nested عندما يكون أحد النماذج المتولدة لها نفس المعالم الحرة لنموذج آخر بالإضافة لمعالم حرة أخرى غير موجودة فى النموذج الأخير. بكلمات أخرى فإن النموذجين متكافئين ما عدا مجموعة فرعية من المعالم الحرة الموجودة فى أحد النماذج ومثبتة أو مقيدة فى الآخر.

ويجب صياغة نماذج بديلة قبل التحليل أفضل وأكثر أماناً من إجراء تعديل فى النموذج خاصة إذا كان التعديل على أسس إحصائية إمبريقية.

ويرى (2005) Martines أن صياغة النماذج البديلة لا بد أن تكون على أسس نظرية ودراسات سابقة، وذلك لان صياغتها بدون منطقية نظرية تؤدي إلى صعوبة فى التفسير وهى مضیعة للوقت. وتوجد ميزة فى صياغة نماذج بديلة قبل التحليل للنموذج المستهدف وهو تجنب التحيزات التوكيدية **Confirmation biases** حيث يبدى الباحث درجة من التعصب للقبول بالنموذج المستهدف فى ضوء النظرية دون النظر إلى وجود تفسيرات أخرى للمتغيرات التى تحكم الظاهرة موضع الدراسة (MacCallum & Austin, 2000; Martines, 2005).

وكذلك تفيد صياغة النماذج البديلة للأطر النظرية التى تتضمن رؤى مختلفة لتفسير العلاقات بين المتغيرات، وهذا سائد فى العلوم الاجتماعية حيث يحدث غالباً تعارض بين نتائج الدراسات السابقة. وهذا يعتبر مدعاة إلى صياغة العديد من النماذج البديلة لاختبارها فى ضوء نتائج هذه الدراسات. وتتم صياغة نماذج بديلة فى ضوء أسس إحصائية وذلك من خلال ما تمدنا به مؤشرات التعديل **Modification indices** من إضافة مسارات ويحدث ذلك بعد إجراء تقدير النموذج. ولكن صياغة نماذج بديلة قبل التحليل أفضل وأكثر أماناً من صياغتها بعد التحليل حيث صياغتها قبل تكون لها معقولة نظرية وتفسيرية ولكن صياغتها بعد التحليل تكون على أسس إحصائية يمدنا بها البرنامج، ولكن يمكن أن تكون النماذج البديلة بعد التحليل مقبولة بدرجة كبيرة مثل صياغتها قبل التحليل إذا توفر تفسير وتبرير نظرى قوى للمسارات التى اقترح البرنامج بإضافتها.

مرحلة تحديد النموذج Identification

من أهم القضايا التى تواجهه تحليل نموذج SEM هى قضية تحديد النموذج، والنموذج يكون محدداً إذا كان البرنامج قادر على اشتقاق تقديرات كل معالم النموذج من مصفوفة التغاير.

ومفهوم التحديد يشير إلى فكره وجود على الأقل حل فريد وحيد لكل معلم غير معروف في النموذج باستخدام البيانات المجمعة والقيود التي توضع على بعض المعالم، وحتى إذا لم يوجد قيم مؤكده للمعالم الحرة فإن القيم يجب الحصول عليها تعاد انتاج مصفوفه التباين Σ للنموذج المفترض حتى يمكن مقارنتها لمصفوفه التباين للعينه S.

يوجد اعتبارين أساسين لتحديد نموذج المعادلة البنائية هما:

1. درجات الحرية للنموذج يجب أن تكون على الأقل صفر ($df \geq 0$).

2. أى متغير كامن لا بد أن تحدد له وحدة قياسية مترية.

والشرط الضروري لتحديد النموذج هو ما إذا كانت عدد المعالم الحرة المراد تقديرها للنموذج (تشبعات العوامل، تباينات الأخطاء، العلاقات بين المتغيرات الكامنة، معاملات الانحدار، وغيرها) لا تزيد عن عدد معاملات الارتباطات (المعادلات) فى مصفوفة التباين أو مصفوفة الارتباط (Baumgartner & Homburg, 1996; Bollen, 1989; Ullman, 2006). وتهتم مرحلة تحديد النموذج بتقدير قيم المعالم الحرة للنموذج المفترض من مصفوفة التباين، وفى بعض الحالات لا يتم التحليل أو يصل لحلول بعد 100 محاولة تدوير Iteration مثلاً لان النموذج يعانى من سوء تحديد (Schumacker & Lomax, 2010).

وعندما تكون عدد المعالم الحرة المراد تقديرها فى النموذج مساوية لعدد المعاملات فى مصفوفة التباين أو الارتباط فإن درجات الحرية تساوى صفراً ($df = 0$)، ولذلك يقال أن النموذج محدد تماماً أو متشبع Just identification، وهذا يمدنا بحلول وحيدة وفريدة للمعالم (Hoyle, 1995; Schumacker & Lomax, 2010).

وعندما تكون عدد المعالم المراد تقديرها أكبر من عدد المعادلات المشتقة من مصفوفة التباين أو الارتباط ($df < 0$) يكون النموذج تحت التحديد Under-identification وتكون درجات الحرية سالبة (Hoyle, 1995, Shah & Goldstein, 2006)، فى هذه الحالة فإن معالم النموذج لا يمكن تقديرها إلا إذا تم تثبيت بعض المعالم، أو وضع قيود معينة فى النموذج. وتظهر مشاكل عند تقدير هذا

النموذج فلا تستطيع برامج SEM انجاز تحليل المصفوفة، وتظهر رسائل تحذيرية عديدة منها محدد المصفوفة سالب أو وجود خطأ ما وحتى إذا تم التحليل فإن تقديرات المعالم والمطابقة غير متسقة لا نستطيع تفسيرها (Chou & Bentler, 1995).

وتظهر قضية عدم التحديد للنموذج ليس لعيوب في تخصيص النموذج فقط، ولكن نتيجة عيوب في البيانات أو في المصفوفة المحللة (Hoyle, 1995)، أو نتيجة إعداد ملف المدخلات للبرنامج حيث يحدث خطأ في وصف ملف المدخلات. وإذا استمر البرنامج بإعطاء إفادة بعدم القدرة على تقدير معالم النموذج فيجب إعادة تخصيص أو وصف النموذج وصفاً له معنى (MacCalum, 1995). وغالباً يحدث عدم التحديد للنموذج في النماذج التي تتناول دراسة تأثيرات تبادلية بين المتغيرات الكامنة (Kline, 2011).

وعندما تكون عدد المعالم المراد تقديرها للنموذج أقل من عدد المعاملات في مصفوفة التغاير أو الارتباط يقال أن النموذج فوق التحديد **Over-identification** (درجات الحرية تكون واحد فأكثر)، وهذا النموذج مفضل ومرغوب لوجود أكثر من معادلة مستخدمة لتقدير المعالم، وهذا يؤدي إلى الحصول على تقديرات دقيقة ومتسقة (Bollen, 1989; Shah & Goldstein, 2006)، وعلى ذلك فإن تقدير درجات الحرية شيء أساسي لفهم النموذج وتحديده ومطابقته. وينصح Shah & Goldstein (2006) بأهمية تقدير درجات الحرية لكل نموذج مختبر.

من أهم الإجراءات لتجنب قضية عدم التحديد تثبيت أحد التشعبات للمتغيرات المقاسة على المتغير الكامن بالواحد الصحيح ويسمى **مؤشر معياري** (Bollen, 1989). ونلاحظ أن البواقي (الأخطاء) في نموذج المعادلة البنائية يتم عرضها في شكل المسارات كمتغيرات كامنة ويجب أن يتم وضع مقياسيه لكل متغير كامن في النموذج.

مرحلة التقدير للنموذج Estimation

بعد مرحلة بناء النموذج، وجمع البيانات، وإعدادها للتحليل من خلال فحص الاعتدالية والبيانات الغائبة، التلازمية الخطية، وكذلك التحقق من قضية التحديد بعدها يبدأ الباحث في مرحلة تقدير النموذج. وتوجد عدة طرق لتقدير معالم نموذج SEM أهمها:

• طريقة تقدير الاحتمال أو الترجيح الأقصى Maximum likelihood Estimation

هي الطريقة الحرة Default لمعظم برامج SEM بمعنى يقوم البرنامج باستخدامها إذا لم تحدد طريقة تقدير أخرى. ومعظم تحليلات نماذج SEM التي استخدمت في التراث البحثي اعتمدت على طريقة ML، وذلك لأن استخدام أى طريقة أخرى غيرها يتطلب مبررات. وخصائص طريقة ML تتبع من النظرية الاعتدالية Normality Theory، وذلك لأن هذه الطريقة تتطلب الاعتدالية المتدرجة لتوزيعات المتغيرات المقاسة في المجتمع وتفترض أن يكون محدد المصفوفة موجب (McDonald & Ho, 2002; Shah & Goldstein, 1996). وإذا كانت المتغيرات التابعة غير متصلة (Kline, 2016)، وتعاني عدم الاعتدالية بدرجة شديدة فإنه يستخدم طرق تقدير بديلة لـ ML. وتستخدم بفاعلية عندما يكون النموذج محدد تحديدًا جيدًا ولكن إذا كان النموذج يعاني من سوء تحديد فإن النتائج المتحصل عليها ربما تكون غير صادقة وفي هذه الحالة تستخدم طريقة TSLS كمكمل لـ ML. وتعطى هذه الطريقة نتائج أقل قوة إحصائية عند استخدام أحجام عينات أقل من 200 فردًا وحلول غير مستقرة (Quintana & Maxwell, 1999). وتوصل West et al. (1995) إلى أن طريقة ML لديها مناعة أو ضلالة إذا لم تتوفر الاعتدالية (توفرت بدرجة ضعيفة أو متوسطة) وليس لديها مناعة ضد الاعتدالية الشديدة ($Ku > 7$, $SK > 3$). وتوجد صيغة خاصة لطريقة ML للتعامل مع البيانات غير كاملة البيانات (المفقودة) وهي متاحة في برامج SEM مثل AMOS, MPLUS, LISREL.

• المربعات الدنيا غير الموزونة Un-weighted least squares (ULS): وهو

نوع من إستراتيجية تقدير المربعات الصغرى الترتيبية OLS التي تقلل فروق مجموع المربعات بين مصفوفات العينة والمصفوفات المشتقة من النموذج (المجتمع) ،

وتعطى تقديرات غير متحيزة للمعالم عبر عينات عشوائية ولكنها ليست بكفاءة تقديرات ML (Kaplan, 2009). تمتاز هذه الطريقة عن ML فى أنها لا تتطلب محدد مصفوفة موجب، ويمكن استخدامها لتوليد القيم المبدئية للتحليل التالى لنفس النموذج والبيانات، ولذلك يمكن استخدامها إذا كان محدد المصفوفة سالب.

• طريقة المربعات الدنيا التعميمية (GLS) Generalized least squares:

تنتمى هذه الطريقة إلى عائلة طرق التقدير المربعات الدنيا الموزونة وهى Weighted Least Squares (WLS) (طريقة المربعات الدنيا الموزونة)، ولها نفس مسلمات طريقة ML، ولكنها أقل تشددًا فيما يخص شرط الاعتدالية فيمكن استخدامها لبيانات غير اعتدالية، غير أنه توجد دلائل قليلة على أن أداء GLS أكثر مناعة أو ضلالة Robust من طريقة ML فى حالة عدم توافر مسلمة الاعتدالية المتدرجة (West et al., 1995). وهذه الطريقة تقوم على أساس التناقض بين النموذج والبيانات أى بين مصفوفة التباين للعينة (S)، والمصفوفة التباين المشتقة من النموذج (Σ). وتعطى طريقة GLS تقديرات المعالم والأخطاء المعيارية وكذلك تعطى إمكانية لحساب مؤشرات حسن المطابقة. وعكس طريقة ULS فإن طريقة GLS تتطلب أن تكون المتغيرات المقاسة حرة المقياس وأيضًا غير حرة. وتتميز طريقة GLS على ML فى أنها تحتاج عمليات حسابية وذاكرة كمبيوتر أقل، وهذا ليس له معنى هذه الأيام فى ظل التطورات المتلاحقة لبرامج SEM. وعلى ذلك فإن طريقة ML مفضلة على ULS و GLS. وعندما يكون النموذج محدد تحديدًا جيدًا فإن طريقتى ML و GLS تعطى نفس القيم لتقديرات المعالم لمؤشرات المطابقة.

• للتعامل مع البيانات غير الاعتدالية طور (1984) Brown طريقة (ADF) Asymptotically distribution free وتسمى فى برنامج LISREL بطريقة المربعات الدنيا الموزونة Weighted least squares، والحسابات المطلوبة انجاز طريقة ADF معقدة، وأن عدد الصفوف وكذلك عدد الأعمدة يقدر من الصيغة $0.5V(V+1)$ حيث V عدد المتغيرات المقاسة. فإذا كان لدينا نموذج به 10 متغيرات مقاسة فإن المصفوفة اللازمة لتقدير ADF هى 55 (صف) 55X (عمود) (Kline,

(2016). وهذه الطريقة غير عملية في العلوم الاجتماعية لأنها تتطلب أحجام عينات كبيرة (Baumartner & Homburg, 1996). وأشارت دراسات المحاكاة إلى أن الحد الأدنى المطلوب لاستخدامها يتراوح من 1000 إلى 5000 وذلك للحصول على نتائج مرضية (Boomsma & Hoogland, 2001; Chou & Bentler, 1995).

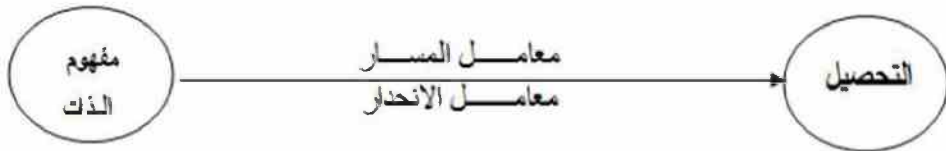
ويرى الباحثون والخبراء في مجال SEM أنه من الأفضل لبيانات تصنيفية (1,0) استخدام مصفوفة الارتباط الرباعي Tetra-choric مع طريقة WLS أو ADF، وليبيانات تصنيفية ترتيبية (3, 2, 1) استخدام مصفوفة Poly-choric مع طريقة التقدير WLS، ولكن متطلبات طريقة WLS يجعلها غير عملية للباحث النفسى والسلوكى لأنها تتطلب حجم نموذج صغير وأحجام عينات كبيرة جداً تقترب من 1000 كحد أدنى أو بأقل تقدير 500، ويزداد هذا الحجم مع زيادة درجة تعقيد النموذج.

تقديرات معالم SEM : دائماً ينشغل الباحثون بتقدير المطابقة للنموذج دون إعطاء الأهمية الكافية لتفسير معالم النموذج. تتضمن معالم نموذج SEM التنبؤات للمتغيرات المقاسة بالعوامل أو معاملات المسار (معاملات الانحدار بين المتغيرات للنموذج البنائى) أو وتباينات الأخطاء الواقعة على المتغيرات، وتكون هذه المعاملات إما معيارية وغير معيارية. وهذه المسارات تسمى تأثيرات وهذه التأثيرات إما تكون مباشرة أو غير مباشرة أو كلية. ويتضمن المخرج **الأخطاء المعيارية Standard errors** وقيمة الدلالة الإحصائية لكل معلم، والأخطاء المعيارية هي ضرورية لتحديد النسبة **الدرجة Critical Ratio** وهي نسبة إحصاء المعلم (التأثير مثلاً) مقسوماً على الخطأ المعيارى وهو الانحراف المعيارى لتوزيع المعاينة، ويقدر خطأ المعاينة من خلال الفروق بين إحصائيات العينة ونظيرتها لمعالم المجتمع.

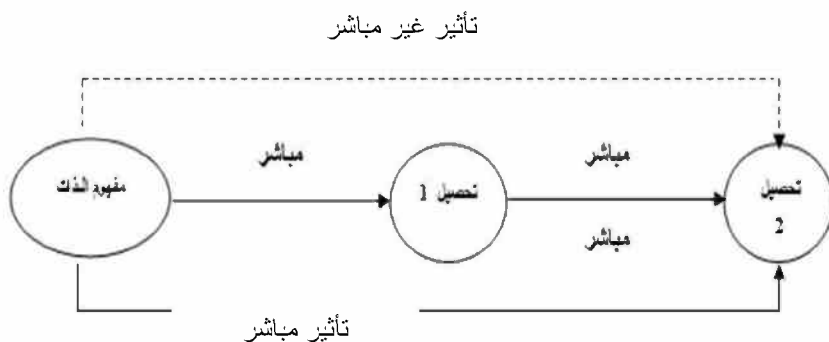
وفى ضوء ذلك فهو ضرورى لتحديد القيمة الحرجة للاختبار التى فى ضوءها يتم قبول أو رفض الدلالة الإحصائية لمعالم نمذجة المعادلة البنائية مثل التأثيرات وتباينات الأخطاء والتغايرات وغيرها. إذا زادت خارج قسمة قيمة المعلم على الخطأ المعيارى عن 1.96 (هى قيمة اختبار T ذو ذيلين عند 0.05) أو 2.58 (هى قيمة اختبار T ذو

ذيلين عند 0.01) فإن المعلم دال إحصائياً وهى أيضاً قيمة اختبار Z وذلك فى العينات الكبيرة.

التأثيرات أو المسارات السببية المباشرة وغير المباشرة والكلية: يمكن تعريف التأثير المباشر بين متغيرين كامنين بينهما تأثير سببى بأنه تأثير مباشر من أحد المتغيرين إلى الآخر بمعنى وجود سهم ذو اتجاه واحد يربط بينهما، ويتم التعبير عنه كمياً فى ضوء المعامل البنائى (معامل الانحدار المعيارى)، وهو معامل المسار أو معامل الانحدار. على ذلك فإن أى تغير لمفهوم الذات نتوقع حدوث تغير فى التحصيل مباشرة وعلى ذلك فإن لمفهوم الذات تأثير مباشر إلى التحصيل.



فإذا كان التأثير المباشر يساوى 0.42 فإن هذا يعنى أن زيادة وحدة واحد من مفهوم الذات تنتبأ بزيادة 0.42 وحدة من التحصيل. أما التأثير غير المباشر **Indirect effect** عندما يوجد تأثيرات بين متغيرين ويتوسطهم متغير أو أكثر من المتغيرات الوسيطة **Mediating variables** بمعنى لا يوجد خط مستقيم أو أسهم مباشر يربط بين المتغيرين وفى الشكل (2.8) فإن المتغير الكامن مفهوم الذات يصل تأثيره إلى المتغير الكامن التحصيل (2) من خلال التحصيل (1) أو أكثر حسب المتغيرات الوسيطة فى النموذج ، إذا تغير مفهوم الذات فإن تحصيل (1) يتغير وبدوره يؤثر فى تحصيل (2) ويمكن القول أن مفهوم الذات له تأثير سببى على تحصيل (2) ولكنه غير مباشر ويقاس من خلال معامل المسار أو الانحدار ويمكن التعبير عنه بصورة معيارية أو بصورة غير معيارية ويطلق عليها التأثيرات الوسيطة **Mediation effect**. وهى لا تضاف أو ترسم فى شكل المسارات ويمكن أن تكون إشارة التأثير المباشر وغير المباشر سالبة أو موجبة. أما التأثيرات الكلية **Total effect** فهو مجموع كلاً من التأثير المباشر والتأثير غير المباشر.



الشكل (8.10): مثال للتأثيرات غير المباشرة بين متغيرين.

ويشير (2007) Tabachnick & Fidell إلى أن التشبعات العالية أفضل، فالتشبع أقل من 0.3 غير مقبول وضعيف، والتشبع من 0.45 فأعلى مقبول. كما أن الدلالة الإحصائية تقدر من خلال قيمة T المقابلة لكل تأثير (يعطيها الليزرال) أو قيمة Z (يعطيها EQS) وهي قيمة التأثير مقسوما على الخطأ المعياري فإذا كانت $T > 1.96$ لاختبار ذو ذيلين فأكثر فإن التأثير ضروري ودال إحصائياً.

مرحلة مطابقة النموذج Model Fit

يتم تقييم النموذج في ضوء تفاصيل النموذج مثل التأثيرات، والتشبعات بالعوامل، وتقدير التباين المفسر لكل معادلة بنائية، وأيضاً في ضوء تقدير مطابقة النموذج ككل overall Model fit Assessment من خلال مؤشرات حسن المطابقة goodness of fit indexes. وتعرف المطابقة إجرائياً بمدى التعارض أو الفروق بين المصفوفة المشتقة من النموذج (Σ) ومصفوفة التغاير أو الارتباط لبيانات العينة (S)، ويقال أن النموذج تام المطابقة إذا كان الفرق بين Σ و S مساوياً للصفر، أي يتم قبول الفرض الصفرى $H_0 : \Sigma = S$ ، وكلما زادت الفروق بين المصفوفتين قلت مطابقة النموذج للبيانات ، ويطلق على الفرق بينهما بالتناقض الإمبريقي وأرجعه (1988) Marsh, Balla & McDonald إلى أخطاء في البيانات أو في النموذج أو في العينة. ويرى (2000) MacCullum & Austin أن الباحثين واعين لحقيقة أساسية هو عدم وجود نموذج حقيقي وحيد يعكس

الظاهرة، وان كل النماذج بها درجة من الخطأ وان أفضل نموذج لا بد أن يتميز بالبساطة وله معنى مقبول (جوهرى).

ومؤشرات حسن المطابقة وسيلة لتكميم التباين المفسر فى النموذج وهو يشبه مؤشر R^2 فى الانحدار المتعدد (Hu & Beutler, 1995).

فيما يلي عرض هذه المؤشرات:

أولاً: مؤشرات المطابقة المطلقة: وهى تجيب على تساؤل: هل البواقي أو التباين غير المفسر فى مصفوفة البواقي بعد مطابقة النموذج يؤخذ فى الاعتبار أى أنها مؤشرات تعتمد على الفروق بين مصفوفة التباين للعينة ومصفوفة التباين المشتقة أو المستهلكة من النموذج وهى كالاتى:

• إحصاء (χ^2) : هو المؤشر التقليدى لتقدير المطابقة فى نمذجة المعادلة البنائية، ويقدر مقدار التعارض أو الاختلاف بين مصفوفة التباين للعينة S ومصفوفة التباين المشتقة من النموذج وإذا كانت قيمة χ^2 للنموذج تساوى صفر فإن النموذج يعتبر محدد تماماً ويتطابق تماماً مع البيانات، بمعنى كل مصفوفة تباين مقاسة تساوى كل مصفوفة تباين مستخلصة من النموذج، ولو كان النموذج فوق التحديد فإن قيمة χ^2 تزيد عن الصفر. وعلى ذلك فإن زيادة χ^2 هى مقياس لسوء مطابقة النموذج **Badness of fit**، أى كلما زاد الفرق بين S و Σ كلما زادت احتمالية رفض النموذج (قبول الفرض البديل) بمعنى وجود دلالة إحصائية (Barrett, 2007; Engle, Mossbrugger & Muller 2003; Hu & Bentler, 1995; Kline, 2016).

أهم الانتقادات الموجهة لمؤشر χ^2 :

1. تتأثر قيمته بعدم توافر الاعتدالية المتدرجة، وهذا يؤدى إلى تضخم (زيادة) قيمة χ^2 وبالتالي الحصول على دلالة إحصائية. ويظهر النموذج سوء المطابقة مع البيانات على الرغم أنه محدد تحديداً دقيقاً وحقيقياً (West et al., 1995).

2. تتأثر قيمة مؤشر χ^2 بحجم العينة ، فزيادة حجم العينة يؤدى إلى الحصول على دلالة إحصائية مما يترتب عليه رفض النموذج على الرغم من بناءه فى ضوء نظرية

متماسكة، وهذا الرفض نتيجة لزيادة حجم العينة حيث يحدث تضخم للخطأ من النوع الأول ومن الأفضل استخدام مؤشر χ^2 لعينة تتراوح من 100 حتى 200 (Hair et al., 1993; Joreskog & Sorbom, 1998). وعند استخدام مؤشر χ^2 للحكم على مطابقة النماذج مع لحجم عينة تتراوح من 200 حتى 300 يحدث دائماً رفض للنموذج، ولذلك أشار (Engel et al., 2003)، عدم إعطاء اهتمام كبير لدلالة مؤشر χ^2 للحكم على مطابقة النموذج.

3. تعقيد النموذج: أحد عيوب مؤشر χ^2 هو أن قيمته تقل كلما أضيفت معالم جديدة إلى النموذج وهذا نتيجة نقصان درجات الحرية.

5. الثبات المنخفض للقياسات لتحليل متغيرات بثبات منخفض يؤدي إلى قوة إحصائية منخفضة وعليه قبول الفرض الصفري أى وجود مطابقة للنموذج، وفي هذه الحالة فالمطابقة ليس نتيجة للتحديد الجيد للنموذج بل إلى القوة الإحصائية المنخفضة وهى فشل الاختبار فى الحصول على دلالة إحصائية.

- مؤشر مربع كا ٢ المعيارية Normed chi-Square (NC): وفى محاولة للتغلب على الحساسية الشديدة Hypersensitivity لمؤشر χ^2 لحجم العينة، اقترح البعض أن يكون مؤشر χ^2 مقروئاً بـ درجات الحرية χ^2/df ، ويوجد ثلاثة محددات لهذا المؤشر كما حددها (Kline 2016) وهو أنه حساس بدرجة متوسطة لحجم العينة، وإن درجات الحرية ليست لها علاقة بحجم العينة، وليس له حدود قطع واضحة لتحديد مطابقة النموذج، وذلك لعدم وجود أساس منطقي أو إحصائي لهذا المؤشر. ويجب أن لا يكون له دور فى تقدير مطابقة النموذج ولكن الشائع عند بعض الباحثين أنه إذا كانت $\chi^2/df \geq 2$ فإن النموذج يعكس مطابقة جيدة، وأحياناً يقترح البعض $\chi^2/df \geq 3$

- مؤشر جذر متوسط مربع الخطأ للتقريب Root Mean Square Error of Approximation (RMSEA): هو مقياس لسوء مطابقة النموذج. ونتيجة لمحددات χ^2 فإن مؤشر RMSEA يعتبر المدخل الأكثر ملائمة لتقدير مطابقة النموذج ولذلك لأن اختبارات الفروض للمطابقة التامة exact fit (نعم، لا) مثل X^2

تستبدل باختبارات الفروض للمطابقة المتقاربة **Close fit** (Browne & Cudeck, 1993).

وحدد (1993) Browne & Cudeck أن المطابقة القريبة **close fit** أو الجيدة عندما تكون قيمته $RMSEA \geq 0.05$ والقيم بين 0.05 و 0.08 أو 0.10 تعنى مطابقة متوسطة (عادية) mediocre بينما القيم أكبر من 1.0 تعنى مطابقة غير مقبولة. ولكن توصل (1999) Hu & Bentler إلى أن المطابقة الجيدة فى ضوء هذا المؤشر تكون 0.06 فأقل ويمكن أن تقبل المطابقة حتى 0.08.

• مؤشرى جذر متوسطات مربعات البواقي وجذر متوسطات مربعات البواقي

المعيارية **Standardized, Root Mean Square Residual**

(RMR, SRMR): مؤشر RMR ابتدعه (1981) Joreskog & Sorbom

هو مؤشر لسوء المطابقة وقائم على أساس البواقي المتطابقة وهى الفروق $S - \sum(\Theta)$ لمعالم النموذج، والواضح أنه إذا كانت $RMR = 0$ فإنه يعنى مطابقة تامة ولكن هذا المؤشر يعتمد على أحجام تباينات وتغايرات المتغيرات المقاسة، وبدون الأخذ فى الاعتبار وحدات قياس أو مقياسية للمتغيرات فى الحساب فإنه من غير الممكن القول بأن قيمة RMR تشير إلى مطابقة جيدة أو سيئة. وعلى ذلك أحد القضايا المتعلقة بهذا المؤشر حسابه للمتغيرات غير المعيارية، ومداه يعتمد على وحدة القياس للمتغيرات المقاسة. فلو كانت وحدات قياس المتغيرات فى المصفوفة مختلفة فمن الصعب تفسير قيمته. ويوجد مؤشر SRMR هو متاح فى برنامج الليزرال و EQS وغيرها. والقيمة صفر تشير إلى مطابقة تامة. ويرى (2005) Martnes أن مؤشر SRMR يسلك مسلك مؤشر RMSEA من حيث فلسفته وحدود القطع الخاصة به ولكنه أقل تأثراً لمحددات χ^2 . وتوصل (1999) Hu & Bentler أن القيمة أقل من 0.10 حتى 0.07 تدل على مطابقة مقبولة والقيمة 0.07 فأقل تدل على مطابقة جيدة.

• مؤشر حسن المطابقة **Goodness of fit index (GFI)**: هذا المؤشر

اقترحه (1989) Joreskog & Sorbom لطريقة التقدير ML، وهو مقياس لمقدار

التباين أو التغيرات في مصفوفة بيانات العينة (S) عن طريق النموذج وهذا المؤشر مشابهة لمعامل التحديد في الانحدار المتعدد (Mulaik, Jams, Alstine, Bonnett, lind & Stilwell 1989) وقيمته تنحصر من الصفر إلى الواحد الصحيح حيث تشير القيم العليا إلى مطابقة جيدة ، والقاعدة العامة هي القيمة 0.95 فاعلى تشير إلى مطابقة جيدة بينما القيم أكبر من 0.90 حتى 0.94 تشير إلى مطابقة مناسبة وأقل من 0.90 تشير إلى مطابقة ضعيفة (HU & Bentler, 1999) ومن محددات هذا المؤشر هو تأثيره الواضح بحجم العينة أى أن قيمته تزيد مع زيادة حجم العينة (HU & Bentler 1995, 1999)، ويتأثر بطريقة التقدير المستخدمة ولذلك فإنه توجد صيغ مختلفة لهذا المؤشر باختلاف طريقة التقدير.

• **مؤشر حسن المطابقة المصحح Adjusted Goodness of fit index (AGFI):** هذا المؤشر طوره (Joreskog & Sorbom 1989) وذلك لعلاج تحيز مؤشر GFI الناتج عن تعقيد النموذج تتراوح قيمته من الصفر حتى الواحد الصحيح، والقيم العالية تشير إلى مطابقة أفضل، ويمكن أن تكون قيمته سالبة، والقيمة 0.90 فأكبر تشير إلى مطابقة جيدة، والقيم أعلى من 0.85 تشير إلى مطابقة مقبولة (Engle, Mossbrugger & Muller, 2003). يتأثر هذا المؤشر بحجم العينة، ودرجة تعقيد النموذج، ولكن بدرجة أقل من مؤشر GFI وكذلك يتأثر بسوء التحديد للنموذج (Muailk et al., 1989).

وتعتبر المؤشرات المطلقة أكثر اعتماداً على حجم العينة ماعدا مؤشري RMSEA و SRMR اعتماداً بدرجة قليلة على حجم العينة.

ثانيًا: مؤشرات المطابقة المتزايدة أو المقارنة أو النسبية **Relative, Comparative, or Incremental fit Indexes**: وهى تقيس أو تقييم مطابقة النموذج المستهدف Target model فى ضوء علاقته بالنموذج القاعدى أو الصفرى Basline or Null Model وهو نموذج سىء المطابقة مع البيانات، أى أنها تقيس التحسن النسبى فى المطابقة للنموذج عن طريق مقارنته بنموذج أكثر قيوداً، وهو النموذج الذى يفترض فيه كل المتغيرات المقاسة بدون تباينات الأخطاء مثبتة عند

الصفر أو كل تشبعات العوامل مثبتة عند الواحد الصحيح، وكل المتغيرات غير مرتبطة (كل التغيرات بين العوامل = صفر) (Joreskog & Sorbom, 1993) وعلى ذلك لا يمكن تقدير معالم له وعادة تكون مطابقة النموذج الأساس أو القاعدى أو الصفرى سيئة، وتكون أساس للمقارنة بالنموذج المستهدف وقد قسم HU & Bentler (1995) وتتضمن المؤشرات الآتية:

أ- **مؤشر المطابقة المعيارى (Normed fit index (NFI)**: وهذا المؤشر اقترحه Bentler & Bonnett (1980) وهذا المؤشر يأخذ قيم من 0 إلى 1.0. فالقيمة العالية تشير إلى مطابقة أفضل، والقيمة 0.95 فاعلى تشير إلى مطابقة جيدة بالنسبة للنموذج القاعدى بينما القيمة 0.90 تدل على مطابقة مقبولة (Hu & Bentler, 1999; Kaplan, 2000; Schumaker & Lomax, 2010). ولكن من محددات هذا المؤشر هو تأثيره بحجم العينة ، وكذلك بعدم توافر الاعتدالية للمتغيرات وغير حساس لسوء تحديد النموذج (Hu & Bentler 1999; Mulaik et al., 1989).

ب - **مؤشر المطابقة غير المعيارى (Non-normed fit index (NNFI)**: وللتخلص من محددات مؤشر NFI وخاصة تأثيره بحجم العينة وتتراوح قيم هذا المؤشر فى المدى من الصفر إلى الواحد الصحيح، فالقيمة 0.95 فأكبر تشير إلى مطابقة جيدة بينما القيمة 0.90 تشير إلى مطابقة مقبولة (Hu & Bentler, 1999). و يرى Engle et al. (2003) أن القيمة 0.97 تبدو أكثر منطقية كدلالة للمطابقة الجيدة بدلا من 0.95.

ومن أهم مميزات هذا المؤشر هو أقل تأثراً بحجم العينة وخاصة لطريقة التقدير ML (Hu & Bentler 1995, 1998; Marsh, Balla & McDonald 1988)، وكذلك يعتبر هذا المؤشر حساس إلى حدًا ما لسوء تحديد النموذج، ويميل إلى رفض النماذج فى حالة العينات الصغير (Hu & Bentler 1999).

ج- **مؤشر المطابقة التزايدى (Incremental fit index (IFI)** — Bollen (1989): Bollen's non-normed index (Delta) وهو تحسين لمؤشر NFI وأشار Bollen (1989) إلى وجود علاقة ارتباط ضعيفة بين أداء هذا المؤشر وحجم

العينة، ولكن (Hu & Bentler (1993 أشار إلى أن هذا المؤشر أكثر تأثراً بحجم العينة في حالة استخدام طريقة GLS مقارنة بطريقة ML.

د. مؤشر المطابقة المقارن **(CFI) Comparative fit index** : هذا المؤشر طوره (Bentler (1990، وهو صيغة منقحة لمؤشر **Relative Non Centrality fit index** (RNI) لـ (McDonald & Marsh (1999، ويمكن لقيمته أن تقع خارج المدى 0-1.00. ووجد (McDonald & Marsh (1990 أن هذا المؤشر غير حساس لحجم العينة، ويوصى باستخدامه عند المقارنة بين النماذج البديلة وذلك لتجنب قضية التحيز للعينات الصغيرة لمؤشر NFI.

وتتراوح قيمته من الصفر إلى الواحد الصحيح، والقيمة المرتفعة تشير إلى مطابقة مناسبة والقيمة 0.95 فأكثر تدل على مطابقة جيدة، بينما القيمة 0.90 تدل على مطابقة مناسبة أو مقبولة (Hu & Bentler, 1995)، بينما أكد Engel et al. (2003) أن القيمة 0.97 تشير إلى مطابقة جيدة، بينما القيمة أقل من 0.95 تشير إلى مطابقة مقبولة.

و. **المؤشر اللامركزي النسبي (RNI) Relative non centrality index**: هذا المؤشر من أقل مؤشرات المطابقة المتلازمة لـ (McDonald & Marsh (1990 ويستخدم للحكم على مطابقة النموذج وتتراوح قيمته بين 0 إلى 1 ويعتبر النموذج متطابق إذا كانت قيمته 0.90 فأعلى، وتوصى (Hu & Bentler (1990 بأن المطابقة الجيدة في ضوء هذا المؤشر هي 0.95 فأكثر. ويعتبر هذا المؤشر أفضل المؤشرات حيث إنه لا يتأثر بحجم العينة وأكثر حساسية لسوء تحديد النموذج، ويفضل استخدامه عند استخدام نماذج تتضمن مؤشرات عديدة، وتشبعات العوامل تكون قيمتها 0.50 فأكثر (Sharma, Mukherjee, Kumar & Dillon, 2005). وتعتبر مؤشرات NNFI و CFI و RNI و IFI من أفضل المؤشرات المتلازمة للحكم على مطابقة النموذج بأنها أقل تأثراً بحجم العينة وأكثر حساسية لسوء تحديد النموذج.

ثالثاً: **مؤشرات البساطة Adjusted or parsimony indexes**: وهي تقيس كيف يكون للنموذج مطابقة مع بساطته معاً، وتستخدم للمقارنة بين النماذج. وتشير البساطة إلى عدد المعالم المقدرة المطلوبة لتحقيق مستوى محدد من المطابقة. وتعتبر البساطة عامل هام أثناء تقدير

مطابقة النموذج، وتلعب دورًا أساسيًا للاختبار بين نماذج بديلة. ويرى (Crockett 2012) أن هذه المؤشرات تستخدم لتحديد هل أثر إضافة معالم إضافية إلى النموذج تؤثر على مطابقته ومن أهم مؤشرات البساطة الآتى:

أ - مؤشر حسن المطابقة للبساطة (PGFI) Parsimony Goodness of fit Index :

وهذا المؤشر لـ (Mulaik et al. 1989) وهو تعديل لمؤشر GFI ويأخذ قيم تقع فى المدى من الصفر حتى الواحد الصحيح، والقيم المرتفعة تشير إلى نموذج أكثر بساطة وهذا المؤشر يصحح مؤشر GFI من درجة تعقيد النموذج.

ب- مؤشر المطابقة المعيارى للبساطة Parsimony Normed fit index (PNFI):

ويستخدم هذا المؤشر فى حالة المقارنة بين نماذج بنائية بديلة، وهذا المؤشر هو تعديل لمؤشر NFI وقد اقترحه (James, Mulaik, & Brett 1982) وهو يأخذ فى الاعتبار درجات الحرية المطلوبة للحصول على مستوى مطابقة معين ويأخذ قيم محصورة من الصفر حتى الواحد الصحيح. وكلما ارتفعت قيمته يدل على أن النموذج أكثر بساطة.

ج - مؤشر محك المعلومات الايكى Akaike information Criterion (AIC): هذا

المؤشر اقترحه (Akiake 1974) وهو تصحيح لمؤشر X^2 من عدد المعالم المقدرة ، ويستخدم للمقارنة بين عدة نماذج متنافسة.

والنموذج الذى له أقل قيمة لـ AIC هو أكثر مطابقة للبيانات وأكثر بساطة. ويرى (Kaplan 2000) أن مؤشر AIC هو مؤشر يعبر عن سوء المطابقة.

و- مؤشر الصدق التعميمى المتوقع Expected Cross-validation Index (ECVI):

وابتدع هذا المؤشر (Browne & Cudeck 1989, 1993) وهو من عائلة المؤشرات التى تستخدم للمفاضلة أو المقارنة بين النماذج البديلة من خلال الحكم على بساطة النموذج، ومدى قابليته للتعميم عبر عينات أخرى من نفس المجتمع وهو مؤشر استدلالى لمعالم المجتمع

وهو مقياس للتناقض بين مصفوفة التباين المشتقة من النموذج للعينة المحللة ومصفوفة التباين المشتقة من النموذج فى عينة أخرى بنفس الحجم (Joreskog & Sorbom, 1993)، أى أنه يقيم كيف يكون أداء النموذج المتطابق لعينة التحليل Calibration Sample فى عينات مصداقية النتائج Cross (Kaplan, 2000). validation sample وعند المفاضلة بين نماذج عديدة فإن القيمة الصغرى لـ ECVI تشير إلى مطابقة أفضل، وتقوم بعض البرامج SEM بطباعة حدود الثقة لهذا المؤشر مما يسمح بتقدير دقة التقديرات مثل برنامج LISREL.

مرحلة تعديل النموذج SEM

تعتبر مرحلة تعديل النموذج هى الخطوة الأخيرة فى تحليل نموذج المعادلة البنائية فإذا اتضح أن النموذج المفترض غير متطابق مع البيانات فماذا بعد؟ هو أن يتم تفسير النتائج أو يتم إجراء تعديل فى النموذج بمعنى إجراء تعديلات فى العلاقات أو المسارات فى النموذج المبدئى حتى يتم الحصول على مطابقة أفضل ويتم عادة فى ضوء محكات إمبريقية إحصائية (MacCallum & Austin, 2000). ويكلمات أخرى نريد إجراء تحسن جوهري فى القياس أو النظرية وليس فقط للحصول على مطابقة بغض النظر عن مقبوليتها للتفسير. وكثيراً يحدث إجراءات تعديل النموذج فى العلوم الاجتماعية لأنه يحدث دائماً سوء مطابقة للنموذج المفترض (Crockett, 2000). ويحدث التعديل فى النموذج بإضافة معالم أو مسارات أو بحذفها، وهذه العملية يشار إليها بإجراءات أو بحث Specification search (MacCullum, 1986). وتحدث إستراتيجية التعديل غالباً إذا كانت إستراتيجية تحليل نموذج المعادلة البنائية توليد النموذج، وهنا يؤكد (MacCullum & Austin, 2000) أن هذه الإستراتيجية تقود إلى استنتاجات خاطئة وسوء استخدام لـ SEM، وأن النموذج المعدل المشتق من البيانات ينقصه المصداقية (MacCallum, 1986)، وقائم على درجة عالية من الشك نتيجة الصدفة (MacCallum, Roznowski & Necowit, 1992)، وأى استخدام لهذا النموذج يخضع لثلاثة شروط حددها MacCallum & Austin (2000) بالآتى:

• اشتقاقه يكون معروف في ضوء أسس إمبريقية إحصائية.

• ان يكون له معنى جوهري وتفسير نظري

• يجب أن يقيم النموذج المعدل على عينات أخرى مستقلة أى تختبر مصداقيته.

وتوجد إشكالية في استخدام التعديل البعدى للنموذج وهى سوء تطبيق لنمذجة المعادلة البنائية حيث تتحول هدفها من إجراء توكيدى إلى إجراء استكشافى (MacCallum, Wagener, Uchino & Fabrigar, 1993; Quintana & Maxwell, 1999; Ullman, 2006)، وذلك لان إضافة معالم جديدة للنموذج (مسار مثلاً) يكون على أسس إحصائية وليس على أسس نظرية وفى هذه الحالة يكون استخدام التحليل العاىلى الاستكشافى أفضل من التوكيدى.

وأوصى الخبراء بأنه لا يجب تعديل النموذج فى ضوء أسس إحصائية فقط بل أن يتم تدعيمها بتبريرات نظرية قوية، وأن تعميم النموذج المعدل ليس مضموناً وبه درجة كبيرة من عدم المصادقية ويجب أن يؤخذ بحذر شديد، ولتعميم النموذج المعدل يجب اختبار مصداقيته على عينات جديدة (Anderson & Gerbing, 1988; MacCallum et al., 1993; Austin, 2000; MacCallum et al., 1993). ويعارض كثيرًا من الخبراء استخدام إستراتيجية التعديل للنموذج المستهدف وذلك لان تحديد نماذج بديلة قبل التحليل أكثرًا أمانًا من إجراء تعديل بعدى **Post hoc modification**

(Boomsma, 2000; Hoyle & Panter, 1995; McDonald & Ho, 2002)، وأشاروا إلى أنه إذا وجدت ضرورة ملحة لإجراء التعديل فلا بد أن يكون فى أضيق الحدود وعلى أسس وتبريرات نظرية قوية. وكذلك أوصى (Hoyle & Panter 1995) إذا كانت الحاجة ملحة لإجراء التعديل فلا يجب إضافة الارتباطات بين الأخطاء لتحسين المطابقة، بل ينادى (Bramnik 1995) بعدم إجراء تعديل النموذج على الإطلاق وهذا التحذير أكد عليه (MacCallum 1986) فقد حذر من إستراتيجية تحسين النموذج الأوتوماتيكي **Automatic model improvement strategy**، وقال لو أن نموذج لم يتطابق مع البيانات فيجب أجرى التعديلات الأكثر جوهرياً للحصول على مطابقة مناسبة حتى تصل إلى X^2 غير دالة إحصائياً. وترى Ullman

(2006) بأنه لو أن تعديل النموذج أحدث تحسن في المطابقة فإنه يجب إجراؤه في أضيق الحدود بمعنى إجراء تعديلات قليلة جدًا خاصة إذا لم تتوفر عينات مصداقية النتائج.

مداخل تحليل نموذج المعادلة البنائية

يتم إجراء تحليل نموذج المعادلة البنائية في ضوء عدة مداخل أهمها:

أولاً: مدخل الخطوة الواحدة One step approach: فيه يتم تحليل نموذج التحليل العاملي التوكيدي (النموذج المقاس) والنموذج البنائي (النظرية) معًا في تحليل واحد متلازم، أي مرة واحدة. حيث تقوم البرامج مثل LISREL و EQS وغيرها، بتحليل النموذج مرة واحدة دون الفصل بين تحليل النموذج المقاس والنموذج البنائي. وهذا الإجراء متبع في اختبار النظرية وتطوير أدوات القياس وتم تداوله في المجلات في مجال المنهجية والعلوم الاجتماعية، (Fornell & Yi 1992). ويشدد Kline (2016) إلى أن إذا كانت النتائج تشير إلى سوء مطابقة فما مصدر سوء مطابقة، هل هو النموذج المقاس أم البنائي أم كلاهما معًا، وعلى ذلك فباستخدام مدخل الخطوة الواحدة من الصعب تحديد مصدر سوء التخصيص في النموذج.

ثانيًا: مدخل الخطوتين Two-Step Approach: نادى (Anderson & Gerbing, 1988; Burt, 1976; Herting & Costner, 1985) باستخدام مدخل الخطوتين في التحقق من نموذج المعادلة البنائية حيث يبدأ الباحث بتحليل نموذج التحليل العاملي التوكيدي (النموذج المقاس) أولاً وعندما يتم تهيئته وإعادة تخصيصه ونتحقق من مطابقته مع البيانات، يبدأ الباحث في التحقق من نموذج المعادلة البنائية الكامل (تحليل توكيدي وبنائي معًا)، ويُقيم المطابقة الكلية للنموذج. ولو أن مطابقة النموذج المقاس سيئة فإن مطابقة النموذج البنائي تكون سيئة أيضًا حتى لو أن النموذج محددًا تحديدًا حقيقيًا. على ذلك فالمبدأ الأساسي لمدخل الخطوتين هو الفصل بين النظرية والقياس.

وعموماً توجد عدة أسباب لفصل القياس عن النظرية أهمها:

1. **التداخل التفسيري interpretational confounding**: اقترح هذا المصطلح Burt (1976) واستخدمه Anderson & Gerbing (1988) وتحدث عندما تتم تغيرات جوهرية في تشبعات العوامل في حالة تقديرها في نموذج التحليل العااملى التوكيدى عن تقديرها في نموذج المعادلة البنائية، وعليه فإن النموذج المقاس اختلف معالمة (تشبعاته) وهذا يعنى أن التقديرات الإحصائية للمفاهيم (مثل تشبعات العوامل) تتغير اعتماداً على النموذج المحلل سواء مقاس أو بنائى.

2. **سوء التحديد**: وكما سبق فى استخدام مدخل الخطوة الواحدة غير قادر على كشف مصادر الخطأ أو سوء التخصيص فى النموذج.

وانتقد Fornell & Yi (1992) مدخل الخطوتين فى أنها تفترض الآتى:

1. النظرية والقياس مستقلين عن بعضها البعض.

2. تقديرات معالم مدخل الخطوتين غير متسقة وتختلف من خطوة إلى أخرى.

3. الاختبار الإحصائى فى الخطوة الأولى (CFA) مختلف عن الاختبار الإحصائى فى الخطوة الثانية (SEM).

ولكن Anderson & Gerbing (1988) يؤكد أن هذا فهم خاطئ من جانب Fornell & Yi (1992) حيث لا يوجد فصل بين النظرية والتطبيق ولم يدرك أن إجراء التحليل العااملى التوكيدى هو قائم على نظرية وافتراضات مسبقة لطبيعة العلاقات بين المتغيرات المقاسة وعواملها، وأن مدخل الخطوتين يعطى أولوية لتشخيص عيوب النموذج المقاس أفضل من مدخل الخطوة الواحدة. وأقر Anderson & Gerbing (1992) أن هذه المسلمات ليس مسلمات على الإطلاق.

وتوجد وجهات نظر متباينة لتنفيذ مدخل الخطوتين فيرى Anderson & Gerbing (1988) أن الخطوة الأولى تبدأ بإجراء التحليل العااملى التوكيدى لكل المتغيرات المقاسة على عواملها مع السماح بوجود ارتباطات بين العوامل أى يتم إجراء نموذج التحليل العااملى التوكيدى فى تحليل واحد لكل المتغيرات وعواملها. ثم بعد ذلك إجراء

تحليل نموذج المعادلة البنائية. ولكن رؤية (1976) Burt تختلف عن الرؤية السابقة فيرى الخطوة الأولى في تحليل عاملى توكيدى لكل متغير كامن أو عامل بمؤشرات على حدة وليست لكل العوامل فى البناء معًا، ثم الخطوة الثانية يتم تثبيت المعالم التى حصلنا عليها من الخطوة الأولى (تشبعات العوامل) وتقدير معالم النموذج البنائى (التأثيرات بين المتغيرات الكامنة) فقط، لكن القضية هو كيفية تقييم أو دراسة الصدق البنائى (تحليل عاملى توكيدى) للمتغيرات بدون التحقق من كيفية ارتباطها معًا.

ثالثًا: مدخل الخطوات الأربعة لـ (Mulaik & Millsap (2000): هى أكثر اتساعًا لمدخل الخطوتين وأكثر قدرة على تشخيص سوء التخصيص للنموذج المقاس. وفى هذه الطريقة يختبر الباحث سلسلة من أربعة نماذج هرمية ويرى (Mulaik (2000, (2009 لى تكون هذه النماذج المتولدة محددة فلا بد أن يكون العامل تشبع عليه أربعة مؤشرات على الأقل. ولخص (Mulaik & Millsap (2000 هذه الخطوات الأربعة بالتالى:

الخطوة الأولى: النموذج غير المقيد Unrestricted Model: وهو نموذج تحليل عاملى عام استكشافى (بطريقة التحليل العاملى) وليس بطريقة تحليل المكونات الأساسية، بنفس العدد من المتغيرات الكامنة أو العوامل فى نموذج المعادلة البنائية ، وعلى ذلك يكون كل متغير مقاس فى النموذج حر التشبع على كل العوامل وليس مقيد التشبع على عامل محدد (كما فى التحليل العاملى التوكيدى). ويجب أن تستخدم نفس طريقة التقدير التى سوف تستخدم فى تحليل نموذج المعادلة البنائية مثل طريقة ML مع تحديد عدد العوامل، ويمكن تنفيذ هذه الخطوة فى برنامج SPSS، ويمكن الحصول على المؤشر X^2 لحسن المطابقة، وعدم المطابقة لهذا النموذج يعطى دلالة لعدم مطابقة نموذج المعادلة البنائية. وهذا النموذج غير المقيد هو متولد من داخل النموذج المقاس وهذه الخطوة تختبر الدقة أو المصادقية حول عدد العوامل وكذلك تسمح بتحديد المشاكل المحتملة فى النموذج المقاس.

الخطوة الثانية: اختبار النموذج المقاس المتولد من التحليل العاملى الاستكشافى وهذه هى الخطوة الأولى من مداخل الخطوتين وهو التحليل العاملى التوكيدى للقياسات

الذى يختبر العلاقات بين المتغيرات المقاسة والكامنة، وفيها تحدد تشبعات المؤشرات بالعوامل المحددة لها وهذا يتم إجراؤه لكل متغير كامن وإذا ثبت النموذج المقاس مطابقة مع البيانات ينتقل الباحث إلى الخطوة الثالثة.

الخطوة الثالثة: اختبار نموذج المعادلة البنائية: يتم اختبار نموذج المعادلة البنائية وذلك بالأخذ بالاعتبار الخطوة الثانية، وعلى ذلك فإن هذا النموذج تكون فيه التأثيرات بين المتغيرات الكامنة المحددة سلفاً فى الإطار النظرى، وكذلك يعاد تقدير معالم النموذج المقاس (التحليل العاقل التوكيدى) وهذا مشابه للخطوة الثانية فى مدخل الخطوتين.

الخطوة الرابعة: اختبار الفروض حول المعالم الحرة المحددة من الباحث: ويمكن إجراؤها من خلال تثبيت أحد المعالم الحرة فى النموذج عند الصفر أو حذف أحد المسارات من النموذج وذلك على أساس الدراسات السابقة التى اعتمدت على بيانات من عينات مختلفة لنفس متغيرات النموذج وهذا يتطلب وجود نماذج بديلة.

ويتسأل (2016) Kline أى مدخل أفضل لتحليل نموذج المعادلة البنائية خطوة ام خطوتين أم أربع خطوات؟ يرى أن مدخل الخطوتين أبسط ولا يتطلب أربعة مؤشرات على الأقل للعامل ولكن كلاهما أفضل من مدخل الخطوة الواحدة، ويؤكد Bentler (2000) أن كلاً من المدخلين ليس هما المعيار أو المحك الذهبى لاختبار SEM ولا يوجد هذا المعيار الذهبى فى تحليل SEM. والخبراء فى مجال SEM ينادون بتطبيق مدخل الخطوتين عند تحليل نموذج SEM (Anderson & Gerbing, 1988; MacCallum & Austin, 2000; McDonald & Ho, 2002).

مثال لإجراء تحليل نموذج المعادلة البنائية من خطوتين

أراد باحث دراسة العلاقات السببية بين مفهوم الذات والدافعية على التحصيل وباستخدام SEM ولتحقيق ذلك يستلزم الخطوات الآتية:

أولاً: تخصيص النموذج: ويتضمن هذا تحديد طبيعة النموذج المقاس (التحليل التوكيدى) والنموذج البنائى كالاتى:

1. **تخصيص النموذج المقاس:** وتم قياس مفهوم الذات (sl) بثلاثة مؤشرات (أبعاد) s_1, s_2, s_3 حيث تمثل s مجموع المفردات الممثلة لكل بعد، وقياس الدافعية (mo) بثلاثة مؤشرات أو أبعاد هي m_1, m_2, m_3 وقياس التحصيل (ach) بثلاثة مؤشرات (أبعاد) هي a_1, a_2, a_3 ، وتم وضع وحدة قياس لكل عامل من العوامل الثلاثة وأصبح العلاقات في النموذج المقاس كالآتي في ضوء (LISREL – SIMPLIS):

Relationships

$$a_1 = 1 * ach$$

a_1 متغير مرجعي،

$$a_2, a_3 = ach$$

a_2, a_3 على الشمال (متغيرات مقاسة) تابعة

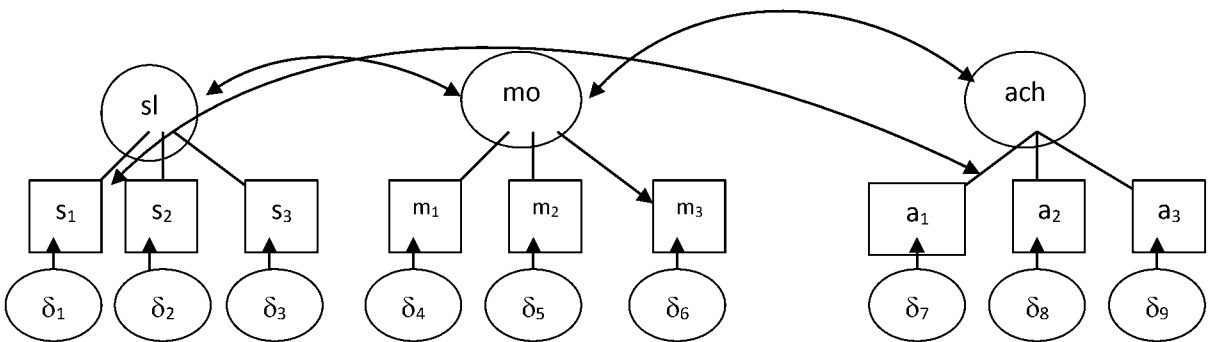
$$s_1 = 1 * sl$$

$$s_2, s_3 = sl$$

$$m_1 = 1 * mo$$

$$m_1, m_2 = mo$$

ويظهر أن المتغيرات a_1, s_1, m_1 متغيرات مرجعية وذلك تثبت تشعبها بالعوامل المحددة لها بالواحد الصحيح لتجنب قضية التحديد. وهذا التخصيص للنموذج المقاس هو لنموذج التحليل العاملى التوكيدى وفيما يلى شكل المسار لهذا النموذج:



الشكل (9.10) : النموذج المقاس لمفهوم الذات والدافعية والتحصيل.

$$s_1 = 1 * sl + \delta_1$$

المؤشر = التشعب بالعامل * العامل + خطأ القياس

$$s_2 = \lambda_2 s_1 + \delta_2$$

$$s_3 = \lambda_3 s_1 + \delta_3$$

$$m_1 = 1 * m_0 + \delta_4$$

$$m_2 = \lambda_5 m_0 + \delta_5$$

$$m_3 = \lambda_6 m_0 + \delta_6$$

$$a_1 = 1 * a_{ch} + \delta_7$$

$$a_2 = \lambda_8 a_{ch} + \delta_8$$

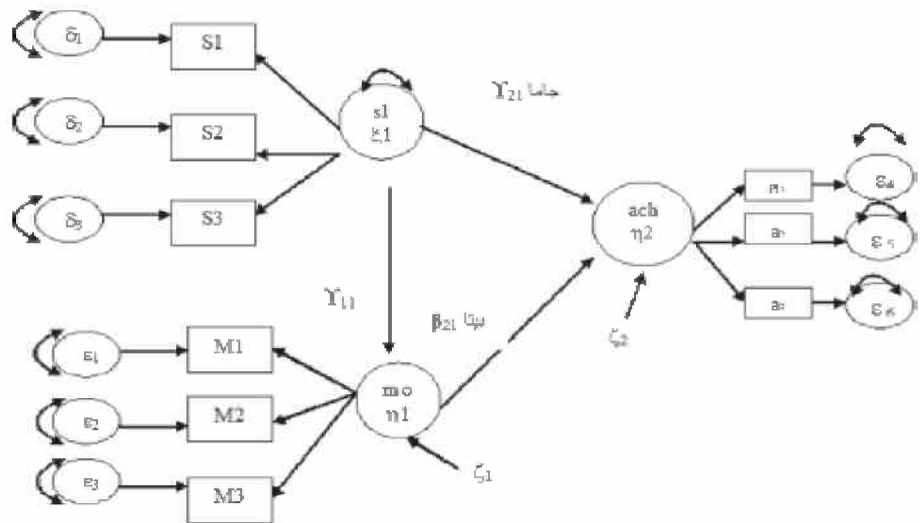
$$a_3 = \lambda_4 a_{ch} + \delta_9$$

ونلاحظ في تخصيص النموذج المقاس تم تمثيل العامل بثلاثة مؤشرات من النوع النموذج التجميعي الجزئي (انظر الفصل الرابع)، وهو وضع مثالي ولاحظ أنه تم تمثيل المتغيرات الكامنة بالحزم للمفردات وليست مفردات كل مقياس، فتخيل إذا تم تمثيل مفهوم الذات بمفردات مقياسه هي عشرون مفردة والدافعية بثلاثين مفردة والتحصيل بـ 15 مفردة فإن عرض النموذج يتميز بالتعقيد وهذا يؤثر على تقدير معالمه.

2. تخصيص النموذج البنائي (نموذج SEM الكامل): وفيه افترض الباحث وجود علاقات أو تأثيرات سببية بين المتغيرات في ضوء الدراسات السابقة والتصورات النظرية وافترض الفروض الآتية:

- يوجد تأثير مباشر (مسار) من مفهوم الذات إلى كلاً من الدافعية والتحصيل.
- يوجد تأثير مباشر (مسار) من الدافعية إلى التحصيل.

وعلى ذلك يمكن عرض شكل المسار لنموذج SEM:



الشكل (10.10) : نموذج SEM لمفهوم الذات والدافعية والتحصيل

وعلى ذلك فإن مفهوم الذات متغير خارجي (مستقل) والدافعية متغير تابع (داخلي) لمفهوم الذات ومستقل للتحصيل ولذا فهو متغير وسيط والتحصيل متغير داخلي تابع و δ_1 , δ_3 تمثل أخطاء القياس (التباين غير المفسر) الواقعة على مؤشرات المتغير الكامن المستقل مفهوم الذات بينما ϵ_1 , ϵ_6 (إبيسلون) تمثل أخطاء القياس الواقعة على مؤشرات المتغيرات الكامنة التابعة (التحصيل والدافعية)، ζ_1 (زيتا) الأخطاء الواقعة على المتغيرات الكامنة التابعة.

وفيما يلي معادلات النموذج البنائي:

$$ach = \gamma \text{ sel} + \beta \text{ mo} + \zeta_2 \text{ (خطا التنبؤ)}$$

$$mo = \gamma \text{ sel} + \zeta_1$$

γ و β معاملات الانحدار أو المسار أو المعامل البنائي.

وعلى ذلك تحدد في برنامج LISREL بإضافة معادلتين على معادلات النموذج

Relationships

المقاس:

$$ach = sel \text{ mo}$$

$$mo = sel$$

على ذلك فإن معالم هذا النموذج SEM (المقاس + البنائي) كالآتي:

1. تشعبات العوامل (λ) تشعب المتغيرات المقاسة بالعوامل المحددة لها وعددها تسعة (9).

2. معاملات المسار أو التأثيرات (γ جاما و β): وهى التأثيرات أو معاملات الانحدارية بين المتغيرات الكامنة وعددها ثلاثة (3).

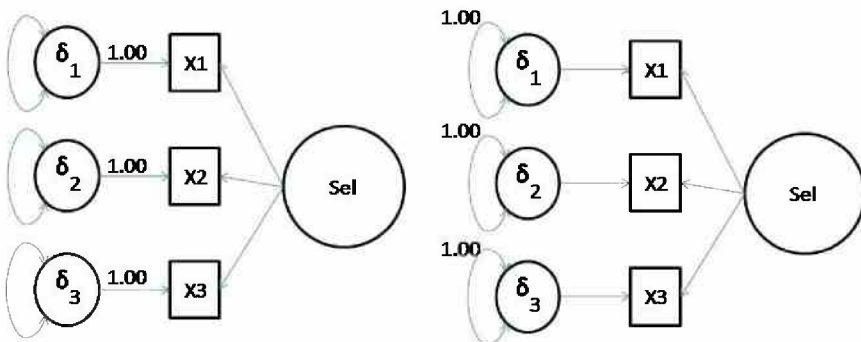
3. تباينات الخطأ: وهى على المتغيرات المقاسة الممثلة للمتغيرات الكامنة المستقلة (δ) = 3، والأخطاء على المتغيرات المقاسة الممثلة للمتغيرات الكامنة التابعة (ϵ) = 6 والأخطاء الواقعة على المتغيرات الكامنة التابعة $\zeta = 2$.

وتوجد طريقتين لنمذجة الأخطاء المرتبطة بالمتغيرات المقاسة فى النموذج ولاحظ أن أخطاء القياس يخرج منها السهم وعليه فهى متغيرات مستقلة كامنة، بينما المتغيرات المقاسة ($a_1, \dots, m_1, \dots, s_1$) يدخل إليها سهم وعليه فهى متغيرات تابعة والطريقتين هما :-

أ. تقدير تباين الأخطاء تسمى Disturbance: تثبت تشعبات الأخطاء على متغيراتها المقاسة بالواحد الصحيح $c_1 \rightarrow S_1$ حيث يشير السهم بالتباين ، وعلى ذلك فإن تباينات الأخطاء = 12 تباين خطأ

ب. تثبت تباينات الأخطاء بالواحد الصحيح ، وتقدر تشعبات الأخطاء على متغيراتها المقاسة $c_1 \rightarrow e_1$

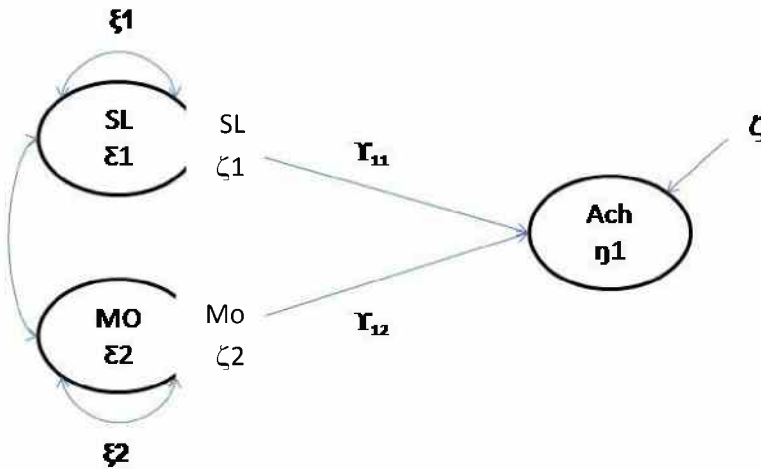
ويمكن توضيح ذلك بالآتى:



الشكل (11.10): مداخل نمذجة أخطاء القياس للمتغيرات المقاس.

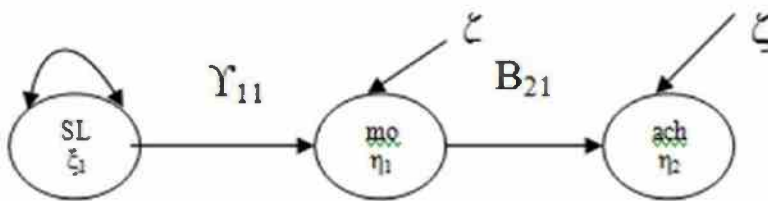
4. التباينات أو تباينات العوامل: هي العلاقة أو التباين بين المتغيرات الكامنة المستقلة (عكس) وفي النموذج لا يوجد سوى متغير مستقل واحد وعلى ذلك فإن عدد التباينات بين العوامل المستقلة = صفر والتباينات للعوامل المستقلة هي واحد للمتغير الكامن المستقل Sel

ويعتبر نموذج SEM ذو تأثيرات أحادية الاتجاه كما يوجد اعتبار آخر في مرحلة تخصيص النموذج وهو صياغة نماذج بديلة وذلك لتجنب التحيز التوكيدي للنموذج المفترض ومحاولة الادعاء بوجود سببية. ولذلك تم فرض عدة نماذج بديلة كالآتي:



الشكل (12.10): نموذج SEM البديل أو المكافئ الأول.

وفي هذا النموذج تم اعتبار sl, mo, متغيرات كامنة مستقلة بينما ach متغير كامن داخلي (تابع) وتم صياغة نموذج بديل ثانى وهو:



الشكل (13.10): نموذج ثانى بديل للنموذج المفترض.

وتم صياغة هذه النماذج البديلة في ضوء دراسات سابقة وحيث يوجد تعارض في رؤيتها للتأثيرات بين المتغيرات الثلاثة، ولا تصاغ بطريقة اعتباطية. يرى البعض لا مانع من اختبار نموذج بديل عشوائي وعلى ذلك فإن هدف الباحث من تحليل نموذج SEM هو المقارنة بين النماذج وتحديد أي هذه النماذج أكثر اتفاقاً مع بيانات العينة.

البرنامج المستخدم: تم تحليل النموذج باستخدام برنامج (8.8) LISREL والتعامل مع ملف أوامر ملف المدخلات بلغة SIMPIS.

إجراء تقدير نموذج تحليل SEM في ضوء مدخل الخطوتين لـ Anderson & Gerbing (1988).

ولذلك فإن الباحث أعد ملف المدخلات لنموذج التحليل العاملي التوكيدي أولاً كالاتي:

Title: Confirmatory model for achievement

Observed variables : $a_1 - a_3$ $s_1 - s_3$ $m_1 - m_3$

Covariance Matrix from file ex.spl

Sample Size: 220

Latent Variables: ach sl mo

Relationships

$$a_1 = 1 * ach$$

$$a_2 a_3 = ach$$

$$s_1 = 1 * sl$$

$$s_2 s_3 = sl$$

$$m_1 = 1 * mo$$

$$m_2 m_3 = mo$$

Lisrel output : ss sc rs mi

Path diagram

End of problem

ويتم ذلك بفتح البرنامج يعطى الشاشة الافتتاحية ثم اضغط على file ثم New ثم على
SIMPLES project وأحفظ الملف باسم، ولاحظ تم اعطاء النواتج فى ضوء لغة برنامج
LISREL وبإجراء التحليل أعطى البرنامج الآتى:

observed variables: a1-a3 s1-s3 m1-m3
covariance matrix
56.22
31.55 75.55
23.27 28.30 44.45
24.48 32.24 22.56 84.64
22.51 29.54 20.61 57.61 78.93
22.65 27.56 15.33 53.57 49.27 73.76
33.24 46.49 31.44 67.81 54.76 54.58 141.77
32.56 40.37 25.58 55.82 52.33 47.74 98.62 117.33
30.32 40.44 27.69 54.78 53.44 59.52 96.95 84.87 106.35
sample size : 220
latent variables: ach sf mo
relationships:
a1=1*ach
a2 a3= ach
s1=1*sf
s2 s3 = sf
m1=1*mo
m2 m3=mo
lisrel output: ss sc ef rs va pc mi
path diagram
end of problem

اتضح أن البرنامج استغرق 5 محاولات تدوير لإعطاء الحلول:

Number of Iterations = 5

وكانت التشبعات والأخطاء المعيارية والدلالة لها كالاتى:

LISREL Estimates (Maximum Likelihood)

LAMBDA-X

ach	sf	mo
-----	-----	-----

a1	1.00	--	--
a2	1.27 (0.16) 8.08	--	--
a3	0.89 (0.12) 7.70	--	--
s1	--	1.00	--
s2	--	0.92 (0.07) 13.76	--
s3	--	0.88 (0.06) 13.54	--
m1	--	--	1.00
m2	--	--	0.88 (0.05) 16.79
m3	--	--	0.88 (0.05) 18.39
PHI			
	ach	sf	mo
	-----	-----	-----
ach	1.00		
sf	0.62	1.00	
mo	0.67	0.79	1.00

THETA-DELTA

a1	a2	a3	s1	s2	s3
-----	-----	-----	-----	-----	-----
0.55	0.46	0.55	0.27	0.34	0.36

THETA-DELTA

m1	m2	m3
-----------	-----------	-----------

----- ----- -----
0.22 0.27 0.19

الجدول (1.10): الحلول المعيارية وغير المعيارية لمعالم نموذج التحليل العاملى التوكيدى.

الحلول المعيارية	الحلول غير المعيارية				
	R ² الثبات	الدلالة Ti	الخطأ المعياري	التشبع	
العامل الأول					
0.67	0.45	-	-	1.00	a ₁
0.73	0.54	8.08	0.16	1.27	a ₂
0.67	0.45	7.70	0.12	0.89	a ₂
العامل الثاني					
0.85	0.73		-	1.00	s ₁
0.81	0.66	13.76	0.07	0.92	s ₂
0.78	0.64	13.54	0.06	0.88	s ₃
العامل الثالث					
0.89	0.78	-	-	1.00	m ₁
0.86	0.73	16.79	0.05	0.88	m ₂
0.89	0.81	18.39	0.05	0.88	m ₃

ويتضح أن كل التشبعات دالة إحصائيًا ، حيث زادت قيمة T المقابلة لكل تشبع عن 1.96.

وفيما يلى مؤشرات حسن المطابقة:

Goodness of Fit Statistics

Degrees of Freedom = 24

Minimum Fit Function Chi-Square = 52.10 (P = 0.00076)

Normal Theory Weighted Least Squares Chi-Square = 48.28 (P = 0.0023)

Estimated Non-centrality Parameter (NCP) = 24.28

90 Percent Confidence Interval for NCP = (8.23 ; 48.09)

Minimum Fit Function Value = 0.24

Population Discrepancy Function Value (F0) = 0.11

90 Percent Confidence Interval for F0 = (0.038 ; 0.22)

Root Mean Square Error of Approximation (RMSEA) = 0.068

90 Percent Confidence Interval for RMSEA = (0.040 ; 0.096)

P-Value for Test of Close Fit (RMSEA < 0.05) = 0.14

Expected Cross-Validation Index (ECVI) = 0.41
90 Percent Confidence Interval for ECVI = (0.34 ; 0.52)
ECVI for Saturated Model = 0.41
ECVI for Independence Model = 5.46

Chi-Square for Independence Model with 36 Degrees of Freedom = 1177.34
Independence AIC = 1195.34
Model AIC = 90.28
Saturated AIC = 90.00
Independence CAIC = 1234.89
Model CAIC = 182.54
Saturated CAIC = 287.71
Normed Fit Index (NFI) = 0.96
Non-Normed Fit Index (NNFI) = 0.96
Parsimony Normed Fit Index (PNFI) = 0.64
Comparative Fit Index (CFI) = 0.98
Incremental Fit Index (IFI) = 0.98
Relative Fit Index (RFI) = 0.93
Critical N (CN) = 181.68
Root Mean Square Residual (RMR) = 1.99
Standardized RMR = 0.023
Goodness of Fit Index (GFI) = 0.95
Adjusted Goodness of Fit Index (AGFI) = 0.91
Parsimony Goodness of Fit Index (PGFI) = 0.51

يتضح أن χ^2 دالة إحصائية عن (p = 0.00)، ومؤشر RMSEA = 0.068 ،NFI = 0.98 ،NNFI = 0.97 ،CFI = 0.99 ،GFI = 0.95 ،AGFI = 0.91 .وعليه فإن نموذج التحليل العاُملى التوكيدى يتمتع بمطابقة جيدة فى ضوء مؤشرات NFI ،NNFI ،GFI ،GFI وبمطابقة مناسبة فى ضوء مؤشر RMSEA = 0.068 ،حيث زادت قيمته عن 0.06 يمكن فى هذه الحالة أن نبقى على هذا النموذج بدون إجراء تعديل لتحسين مطابقة النموذج.

وأمدتتا مؤشرات التعديل بضرورة جعل العلاقة بين خطأ القياس الواقع على (δ_3) وخطأ القياس الواقع على m_3 (ε_3) حرة، وتم إجراء هذا التعديل بإضافة هذا في خط العلاقات:

Covariance between S_3 and m_3 free set the error

وبإجراء التحليل نلاحظ حدث تحسن جوهري في مؤشر RMSEA فأصبحت قيمته 0.00 بدلاً من 0.068، وكذلك أصبحت قيمة χ^2 غير دالة إحصائياً.

الخطوة الثانية التحقق من نموذج SEM الكامل: وفي هذه الخطوة يتم إدخال العلاقات البنائية بين المتغيرات الكامنة:

أولاً: تحديد النموذج: لابد من تشخيص قضية التحديد قبل اختبار النموذج إحصائياً وذلك من خلال الآتى:

1. لنموذج التحليل العاُملى التوكيدى: فى شكل ()

أ. تحديد عدد التغيرات أو الارتباطات فى المصفوفة: وهى تقدر كالاتى:

$$P = \frac{V(V+1)}{2} = \frac{9 \times 10}{2} = 45$$

ب. تحديد عدد معالم النموذج: وهى كالاتى (q) = 9 تشبعات + 9 تباينات خطأ + 3 ارتباطات أو تغيرات بين العوامل + 3 تباينات عوامل مستقلة = 24 معلم.

إذا درجات الحرية كالاتى: $df = P - q = 45 - 24 = 21$

وبما أنه يوجد ثلاثة تشبعات متغيرات مقاسة مثبتة عند الواحد الصحيح إذا عدد المعالم الحرة لنموذج التحليل العاُملى التوكيدى هى $21 = 24 - 3$ وبالتالي درجات الحرية هى: $24 = 45 - 21$ ، أى أن $df > 0$ وعليه فإن النموذج فوق التحديد وعلى ذلك فإن مصفوفة التغيرات بين المتغيرات المقاسة التسعة قادرة على إعطاء حلول متسقة لمعالم النموذج.

2. التحديد للنموذج الكامل لـ SEM فى (شكل (2.19)): وعدد المعالم لهذا النموذج كالاتى = 9 تشبعات عوامل + 11 تباينات خطأ (9 لمتغيرات مقاسة + 2 لمتغيرات كامنة تابعة) + 3 معاملات مسار + 1 تباين متغير كامن مستقل = 24 معلم.

إذا درجات الحرية: $df = 45 - 24 = 21$ ، وبما أنه يوجد ثلاثة تشبعات لمتغيرات مقاسة مثبتة عند الواحد الصحيح إذا عدد المعالم الحرة فإن: $df = 45 - 21 = 24$ ، إذا نموذج المعادلة البنائية (مقاس + بنائى) فوق التحديد ولا يعانى من قضية النموذج غير المحدد.

ثانيًا: مسح البيانات وتقدير النموذج: تتم فى ضوء الآتى:

- مناسبة حجم العينة: بلغ حجم العينة الإجمالى 220 وعدد المتغيرات 9 إذا تمثل المتغير الواحد بحجم عينة 24.7، ويوجد 21 معلم حر إذا يمثل المعلم بـ 15 فردًا تقريبًا وهى تزيد عن النسبة المتفق عليها وهى عشرة.
- البيانات الغائبة: كان حجم العينة الكلى 233 ولكن بعض الأفراد كان لديهم بيانات مفقودة على استجابات مفردات المتغيرات ولذلك تم استخدام مدخل List-wise عند تقدير مصفوفة التغاير أو الارتباط ولذلك تم حذف 13 حالة.
- الاعتدالية: تم فحصها من خلال حساب معامل الالتواء والتفرطح للمتغيرات المقاس، ولم يزيد قيمتها عن 1.00 وهذا يشير إلى توافر الاعتدالية.
- التلازمة الخطية: بفحص مصفوفة الارتباطات بين المتغيرات المقاسة، اتضح أن معاملات الارتباط لم تزيد عن 0.85 وعلى ذلك فهذا مؤشر على وعدم جود هذه الظاهرة ولكن إذا زادت أحد المعاملات عن 0.90 فيجب حذف أحد المتغيرين.
- طريقة التقدير: تم استخدام طريقة ML حيث تحقق مسلمات استخدامها، وهى حجم العينة > 200 وتوافر الاعتدالية.

وأصبح ملف المدخلات لبرنامج LISREL مثل ملف مدخلات التحليل العاملي التوكيدي في الخطوة الأولى مع إضافة العلاقات البنائية من المتغيرات الكامنة في خط العلاقات:

Relationships

$$a_1 = 1 * ach$$

$$a_2 a_3 = ach$$

$$s_1 = 1 * sl$$

$$s_2 s_3 = sl$$

$$m_1 = 1 * mo$$

$$m_2 m_3 = mo$$

$$ach_2 = sl mo$$

$$mo = sl$$

Set the error covariance between S_3 and m_3 free

LisRel output: All

Path Diagram

End of problem

وفيما يلي شرح مفصل لمكونات المخرج:

observed variables: a1-a3 s1-s3 m1-m3

covariance matrix

56.22

31.55 75.55

23.27 28.30 44.45

24.48 32.24 22.56 84.64

22.51 29.54 20.61 57.61 78.93

22.65 27.56 15.33 53.57 49.27 73.76

33.24 46.49 31.44 67.81 54.76 54.58 141.77

32.56 40.37 25.58 55.82 52.33 47.74 98.62 117.33

30.32 40.44 27.69 54.78 53.44 59.52 96.95 84.87 6.35

sample size : 220

latent variables: ach sl mo

```
relationships:
a1=1*ach
a2 a3= ach
s1=1*sl
s2 s3 = sl
m1=1*mo
m2 m3=mo
ach=sl mo
mo=sl
set the error between s3 and m3 free
lisrel output: ss sc ef rs va pc mi
path diagram
end of problem
```

تم عرض تشبعات المتغيرات المقاسة بالعوامل المحدد لها وهو نفسها كما في مخرج التحليل
العاملى التوكيدى السابق، ولكن زادت بعض التشبعات وكذلك انخفض بعضها فمثلاً تشبع s2
على sl فى الخطوة الأولى التحليل العاملى التوكيدى هو 0.92 بينما تشبعها فى تحليل
النموذج ككل فى الخطوة الثانية 0.89 وهكذا بالنسبة لبقية التشبعات.

LISREL Estimates (Maximum Likelihood)

LAMBDA-Y		
	ach	mo
	-----	-----
a1	1.00	- -
a2	1.27	- -
	(0.16)	
	8.07	
a3	0.89	- -
	(0.12)	
	7.71	
m1	- -	1.00
m2	- -	0.87
		(0.05)
		17.20
m3	- -	0.86

	(0.05)
	18.39
LAMBDA-X	
	sl

s1	1.00
s2	0.89
	(0.06)
	13.89
s3	0.83
	(0.06)
	13.46

ثم أعطى المخرج قيمة التأثيرات GAMMA, Beta (حلول لا معيارية) وهذا تم عرضه فى
الجدول التالى:

BETA		
	ach	mo
	-----	-----
ach	- -	0.22
		(0.06)
		3.82
GAMMA		
	sl	

ach	0.16	
	(0.08)	
	2.13	
mo	1.01	
	(0.09)	
	11.56	

جدول (2.10): قيمة التأثيرات غير المعيارية والمعيارية لمفهوم الذات والدافعية على التحصيل وتباينات الأخطاء وقيم T.

المعلم	الحلول غير المعيارية	SE	T	الحلول المعيارية
مفهوم الذات ← التحصيل	0.16	0.08	2.13	0.26
الدافعية ← التحصيل	0.22	0.06	3.82	0.47
مفهوم الذات ← الدافعية	1.01	0.09	11.56	0.77
تباينات الأخطاء				
التحصيل	13.13	3.05	4.31	0.52
الدافعية	46.84	6.84	6.77	0.41

ويتضح أن التأثير المباشر غير المعيارى من الدافعية إلى التحصيل 0.22 وهذا يعنى أن زيادة 1 نقطة فى الدافعية تنتبأ بـ 0.22 نقطة زيادة فى التحصيل مع ضبط مفهوم الذات، وان الخطأ المعيارى المقابل لهذا التأثير $\frac{0.22}{0.06}$ - إلى ذلك فإن قيمة Z أو T = 3.82 هذا يزيد عن القيمة الحرجة لاختبار $\frac{0.05}{0.06}$ ذيلين عند 0.05 وهو 1.96 وكذلك يزيد عن قيمة اختبار T ذو ذيلين عند 0.01 .وهى 2.56 وعليه فإن التأثير دال إحصائياً. وان قيمة التأثير المباشر من مفهوم الذات إلى التحصيل بمعنى زيادة نقطة واحدة فى مفهوم الذات قادرة على التنبؤ بـ 0.16 نقطة للتحصيل، وان قيمة Z أو t = $\frac{0.16}{0.022}$ = 7.27 عند 0.05 ولكنها غير دالة عند 0.01 وهذه التأثيرات مشابهة لمعاملات الانحدار $\frac{0.16}{0.022}$ معيارية فى الانحدار المتعدد وبما أن هذه المتغيرات تقاس بوحدات قياس مختلفة فإن معاملات المسار (الانحدار) غير المعيارية لكلاً من مفهوم الذات والدافعية على التحصيل لا يمكن مقارنتها . وفى هذه الحالة يجب الاستعانة بمعاملات المسار المعيارية (مشابهة لمعاملات بيتا فى الانحدار المتعدد حيث تعتمد على الدرجات المعيارية Z).

وبفحص التأثيرات المباشرة المعيارية يتضح وجود تأثيرات من الدافعية ومفهوم الذات على التحصيل وبلغ قيمة التأثير للدافعية على التحصيل 0.47 بينما بلغت من مفهوم الذات على التحصيل 0.26 وعلى ذلك يمكن القول أن قيمة تأثير الدافعية يزيد عن تأثير مفهوم الذات على التحصيل بمقدار الضعف إلا قليلاً ولاحظ أن الحلول المعيارية في طريقة ML لا تعطى الأخطاء المعيارية ولا الدلالة الإحصائية لها.

ثم بعد ذلك اعطى المخرج مصفوفة التباين بين المتغيرات الكامنة التابعة (ETA) والمتغيرات الكامنة المستقلة (KSI) ومعامل التباين بين mo, ach 35.70 ومعامل التباين بين 25.17 .ach , ach

Covariance Matrix of ETA and KSI

	ach	mo	sl
ach	25.17		
mo	35.70	112.37	
sl	25.13	65.51	64.97

بعد ذلك مصفوفة التباين بين المتغيرات الكامنة المستقلة (PHI) وهي لـ sl وهي دالة إحصائيًا وهذه تمثل تباين المتغير المستقل الكامن.

PHI

sl
64.97
(8.34)
7.79

و اعطى المخرج الأخطاء الواقعة على المتغيرات الكامنة التابعة مصفوفة PSI:

PSI

Note: This matrix is diagonal.

ach	mo
-----	-----

13.13	46.32
(3.05)	(6.84)
4.31	6.77

ويمكن تقدير قيمة التباين غير المفسر الأخطاء الواقعة على المتغيرات الكامنة التابعة من قسمة قيم الخطأ التباين (PSI) مقسوماً على تباين (تغاير) المتغير الكامن وعلى ذلك:

$$\zeta_1(\text{mo}) = 46.32/1.37 = 0.41$$

$$\zeta_2(\text{ach}) = 13.13/25.17 = 0.52$$

وبمعنى أن 52% من تباين التحصيل لم يفسر ويمكن تفسيره من خلال متغيرات أخرى لم تتضمن في النموذج.

$S^2_{\text{Total}} = \text{explained variance} + \text{unexplained variance}$: إذا

$$S^2t = 52+48 = 100\%$$

و مربع معامل الارتباط المتعدد للمعادلات البنائية وهى تشكل نسبة التباين المفسر فى المتغير الكامن وهى $ach = 0.48$ لـ $mo = 0.59$ أى أن متغيرات أو عوامل مفهوم الذات والدافعية فسرت 48% من تباين التحصيل، ومفهوم الذات فسرت 59% من تباين الدافعية.

Squared Multiple Correlations for Structural Equations

ach	mo
-----	-----
0.48	0.59

ثم مربع معامل الارتباط المتعدد للمتغير المستقل الكامن (مفهوم الذات) على التحصيل والدافعية ، فمفهوم الذات فسرت 39% من تباين التحصيل وفسرت 59% من تباين الدافعية.

Reduced Form

sl

ach	0.39
	(0.06)
	6.82
mo	1.01
	(0.09)
	11.56

لاحظ أن تأثير مفهوم الذات على التحصيل زاد من 0.16 إلى 0.39 عند حذف تأثير الدافعية على التحصيل وبقي تأثيره على الدافعية كما هو.

وفيما يلي قيمة الأخطاء الواقعة على المتغيرات المقاسة THETA:

THETA-EPS

a1	a2	a3	m1	m2	m3
-----	-----	-----	-----	-----	-----
31.05	34.91	24.32	29.40	31.34	22.50
(3.88)	(5.05)	(3.06)	(4.28)	(3.95)	(3.28)
8.00	6.91	7.95	6.87	7.93	6.86

THETA-DELTA

s1	s2	s3
-----	-----	-----
19.67	27.71	28.54
(3.35)	(3.53)	(3.46)
5.87	7.84	8.24

ويتضح أن كل هذه الأخطاء دالة إحصائيًا وهذا يعطى انطباع بأن قياس هذه المؤشرات لم يكن بالدرجة المرضية بمعنى ثباتها منخفض بدليل أن تبايننا أخطاء القياس دالة إحصائيًا ولكن يمكن إعطاء معيارية على هذه الأخطاء من خلال قسمة هذا الخطأ الواقع على المتغير مقاس على قيمة معامل التباين للمتغير بنفسه كما في مصفوفة التباين، فتباين أو تباين $a1 = 56.22$ ، بينما قيمة الخطأ غير المعياري له $\delta = 31.05$ إذاً:

$$\delta_{\text{المعيارية}} = \frac{31.05}{56.22} = 0.55$$

وهذا يمثل التباين غير المفسر فى المؤشر a_1 (خطأ القياس) بينما التباين المفسر كما أمدنا به البرنامج من خلال المؤشر R^2 وهى 0.45 وبالنسبة لمؤشرات S1 التباين غير المفسر المعيارى له :

$$= \frac{\text{الخطأ غير المعيارى}}{\text{تباين المتغير (تغاير)}} = \frac{19.67}{84.64} = 0.23$$

بينما التباين المفسر $R^2 = 0.77$

وفيما يلى مؤشر مربع معامل الارتباط المتعدد:

Squared Multiple Correlations for Y - Variables

a1	a2	a3	m1	m2	m3
0.45	0.54	0.45	0.79	0.73	0.79

Squared Multiple Correlations for X - Variables

s1	s2	s3
0.77	0.65	0.61

وهو يمثل قدر التباين المفسر فى المؤشرات أو المتغيرات المقاسة جراء العوامل بكلمات اشمل يمثل ثبات المتغير ونلاحظ أن R^2 لمعظم المفردات زادت عن 0.50 وهو الحد الأدنى المقبول، بينما ثبات المفردتين a_1 و a_3 أقل من $0.50(0.45)$.
وأعطى البرنامج مؤشرات حسن المطابقة كالاتى:

Goodness of Fit Statistics

Degrees of Freedom = 23

Minimum Fit Function Chi-Square = 20.55 (P = 0.61)

Normal Theory Weighted Least Squares Chi-Square = 20.01 (P = 0.64)

Estimated Non-centrality Parameter (NCP) = 0.0

90 Percent Confidence Interval for NCP = (0.0 ; 11.11)

Minimum Fit Function Value = 0.094

Population Discrepancy Function Value (F0) = 0.0

90 Percent Confidence Interval for F0 = (0.0 ; 0.051)

Root Mean Square Error of Approximation (RMSEA) = 0.0
90 Percent Confidence Interval for RMSEA = (0.0 ; 0.047)
P-Value for Test of Close Fit (RMSEA < 0.05) = 0.96

Expected Cross-Validation Index (ECVI) = 0.31
90 Percent Confidence Interval for ECVI = (0.31 ; 0.36)
ECVI for Saturated Model = 0.41
ECVI for Independence Model = 5.46

Chi-Square for Independence Model with 36 Degrees of Freedom = 1177.34
Independence AIC = 1195.34
Model AIC = 64.01
Saturated AIC = 90.00
Independence CAIC = 1234.89
Model CAIC = 160.67
Saturated CAIC = 287.71

Normed Fit Index (NFI) = 0.98
Non-Normed Fit Index (NNFI) = 1.00
Parsimony Normed Fit Index (PNFI) = 0.63
Comparative Fit Index (CFI) = 1.00
Incremental Fit Index (IFI) = 1.00
Relative Fit Index (RFI) = 0.97
Critical N (CN) = 444.71
Root Mean Square Residual (RMR) = 1.27
Standardized RMR = 0.016
Goodness of Fit Index (GFI) = 0.98
Adjusted Goodness of Fit Index (AGFI) = 0.96
Parsimony Goodness of Fit Index (PGFI) = 0.50

ويتضح من المؤشرات أن قيمة χ^2 غير دالة إحصائياً، و انخفضت قيمة RMSEA عن 0.06 ولها حدود ثقة (0.051, 0.0) بمعنى أن قيمة المؤشر وقعت في هذا المدى والمؤشرات المتزايدة NFI, NNFI, CFI و IFI و RFI زادت عن 0.95، المؤشرات المطلقة GFI, AGFI, زادت عن 0.95 و SRMR انخفضت عن 0.08 وعلى ذلك فإن النموذج يتطابق مع بيانات العينة (مصفوفة التباين) بدرجة جيدة.

ثم أعطى البرنامج مصفوفة fitted Residual كالآتي:

Fitted Covariance Matrix

	a1	a2	a3	m1	m2	m3
	---	-----	-----	-----	-----	-----
a1	56.22					
a2	31.98	75.55				
a3	22.51	28.60	44.45			
m1	35.70	45.36	31.93	141.77		
m2	31.23	39.68	27.93	98.30	117.33	
m3	30.84	39.19	27.58	97.06	84.91	106.34
s1	25.13	31.93	22.47	65.51	57.31	56.59
s2	22.31	28.35	19.95	58.17	50.88	50.24
s3	20.94	26.61	18.73	54.59	47.75	59.42

Matrix Fitted Covariance

	s1	s2	s3
	-----	-----	-----
s1	84.64		
s2	57.69	78.93	
s3	54.14	48.07	73.66

Fitted Residuals

	a1	a2	a3	m1	m2	m3
	-----	-----	-----	-----	-----	-----
a1	0.00					
a2	-0.43	0.00				
a3	0.76	-0.30	0.00			
m1	-2.46	1.13	-0.49	0.00		
m2	1.33	0.69	-2.35	0.32	0.00	
m3	-0.52	1.25	0.11	-0.11	-0.04	0.01
s1	-0.65	0.31	0.09	2.30	-1.49	-1.81
s2	0.20	1.19	0.66	-3.41	1.45	3.20
s3	1.71	0.95	-3.40	-0.01	-0.01	0.10

Fitted Residuals

s1	s2	s3
-----	-----	-----

s1	0.00		
s2	-0.08	0.00	
s3	-0.57	1.20	0.10

وهى الفرق بين مصفوفة التباين المشتقة من قبل النموذج ومصفوفة التباين أو الارتباط المقاسة (المدخلة) ويفضل ألا يزيد الفرق عن 3.00 (قيمة مطلقة) ونلاحظ أن الفرق بين S_3 ، a_3 -3.40 بينما كان معامل التباين بينهما 18.73 وهذا يعنى أن النموذج لم يستطيع تفسير العلاقة بين المتغيرين وهكذا بالنسبة للتباين بين S_2 ، m_1 وعمومًا أعطى البرنامج وسيط هذه الفروق وكان صفر وهذا يشير أن النموذج استهلك معظم معاملات الارتباطات المدخلة.

وأعطى البرنامج ملخص للبواقي المعيارية حيث كان وسيطها:

Summary Statistics for Fitted Residuals

Smallest Fitted Residual = -3.41

Median Fitted Residual = 0.00

Largest Fitted Residual = 3.20

ثم أعطى البرنامج شكل Stem leaf plot لتوزيع الفروق:

Stemleaf Plot

```
- 3|44
- 2|53
- 1|85
- 0|66554311000000000000
0|1111233778
1|01223347
2|3
3|2
```

وفى هذا الشكل يقترب من الاعتدالية مما يشير إلى أن النموذج متطابق مع البيانات. وأعطى البرنامج البواقي المعيارية وهى تعكس نسبة البواقي التباين إلى الخطأ المعيارى المقابل لكل فرق تباين وهذه النسبة تفسر فى ضوء اختبار Z أو T ويلاحظ من هذه البواقي المعيارية لم تزيد T لأى منهم عن 1.96 إذا هذه البواقي غير دالة إحصائيًا وهذا متوقع لان النموذج يتطابق بدرجة جيدة مع النموذج وهو محددًا تمامًا.

ثم أعطى البرنامج مؤشرات التعديل والتغيرات المتوقعة:

Modification Indices for LAMBDA-Y

	ach	mo
	-----	-----
a1	--	0.21
a2	--	0.37
a3	--	0.03
m1	0.14	--
m2	0.02	--
m3	0.22	--

Standardized Expected Change for LAMBDA-Y

	ach	mo	
	-----	-----	
a1	--	-0.39	
a2	--	0.64	
a3	--	-0.12	النقصان فى قيمة X2 جراء اضافة هذا المسار
m1	-0.31	--	
m2	-0.10	--	
m3	0.33	--	

ومؤشرات التعديل لأخطاء القياس حيث يمكن اضافة ارتباطات بينها:

Modification Indices for THETA-EPS

	a1	a2	a3	m1	m2	m3
	-----	-----	-----	-----	-----	-----
a1	--					
a2	0.13	--				
a3	0.44	0.09	--			
m1	0.76	0.11	0.00	--		
m2	1.65	0.00	1.86	0.11	--	
m3	0.45	0.01	1.65	0.03	0.02	--

Modification Indices for THETA-DELTA-EPS

a1	a2	a3	m1	m2	m3
----	----	----	----	----	----

	-----	-----	-----	-----	-----	-----
s1	0.36	0.05	0.77	6.59	0.77	3.75
s2	0.05	0.00	0.45	8.16	0.58	4.60
s3	1.67	0.15	4.85	0.00	0.02	- -

Modification Indices for THETA-DELTA

	s1	s2	s3
	-----	-----	-----
s1	- -		
s2	0.01	- -	
s3	0.18	0.09	- -

Maximum Modification Index is 8.16 for Element (2, 4) of THETA DELTA-EPSILON

وأمدنا البرنامج بالتأثيرات غير المباشرة والكلية:

وأشار البرنامج أن أقصى مؤشر تعديل هو 8.16 للمفردات 24 ، 2 فى المصفوفة، ويمكن اضافة الارتباطات بين أخطاء القياس الواقعة على mps2 (نقصان X^2 بمقدار 8.16)، والتأثيرات الكلية من المتغير الكامن المستقل SL إلى المتغيرات التابعة ach

Total and Indirect Effects

Total Effects of KSI on ETA

	sl

ach	0.39
	(0.06)
	6.82
mo	1.01
	(0.09)
	11.56

التأثير غير المباشر من SL إلى ACH من خلال Mo وهو دال إحصائيًا:

Indirect Effects of KSI on ETA

	sl

ach	0.23

(0.06)

3.72

mo - -

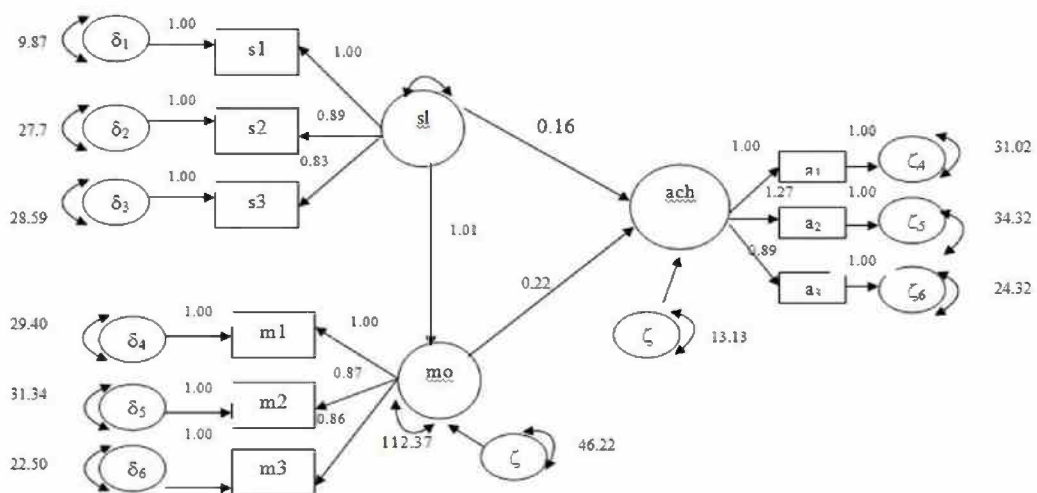
Total Effects of ETA on ETA

	ach	mo
ach	- -	0.22
		(0.06)
		3.82

ونلاحظ أن التأثير المباشر غير المعياري من SL إلى Ach هو 0.16، بينما التأثير الكلى من sl إلى ach هو 0.39 وعلى ذلك فإن قيمة التأثير غير المباشر من sl إلى ach من خلال mo هو: $0.39 - 0.16 = 0.23$

ويلاحظ وجود تأثيرات غير مباشرة من المؤشرات a1 كلا من mo و sl وكلها دالة إحصائياً ولكن هذه التأثيرات ليس لها دلالة فى التفسير النظرى بينما التأثيرات الكلية وغير المباشرة المعيارية هي كالاتى:

وفيما يلى شكل المسار فى ضوء الحلول غير المعيارية:



الشكل (14.10): الحلول غير المعيارية لنموذج SEM.

Completely Standardized Solution

LAMBDA-Y

	ach	mo
	-----	-----
a1	0.67	--
a2	0.73	--
a3	0.67	--
m1	--	0.89
m2	--	0.86
m3	--	0.89

LAMBDA-X

	sl

s1	0.88
s2	0.81
s3	0.78

BETA

	ach	mo
	-----	-----
ach	--	0.47
mo	--	--

GAMMA

	sl

ach	0.26
mo	0.77

Correlation Matrix of ETA and KSI

	ach	mo	sl
	-----	-----	-----
ach	1.00		

mo	0.67	1.00	
sl	0.62	0.77	1.00

PSI
 Note: This matrix is diagonal.

ach	mo
-----	-----
0.52	0.41

THETA-EPS

a1	a2	a3	m1	m2	m3
-----	-----	-----	-----	-----	-----
0.55	0.46	0.55	0.21	0.27	0.21

THETA-DELTA-EPS

	a1	a2	a3	m1	m2	m3
	-----	-----	-----	-----	-----	-----
s1	--	--	--	--	--	--
s2	--	--	--	--	--	--
s3	--	--	--	--	--	0.14

THETA-DELTA

s1	s2	s3
-----	-----	-----
0.23	0.35	0.39

Regression Matrix ETA on KSI (Standardized)

	sl

ach	0.62
mo	0.77

بينما التأثيرات الكلية وغير المباشرة المعيارية هي كالآتى:

Standardized Total and Indirect Effects

Standardized Total Effects of KSI on ETA

	sl

ach	0.62
mo	0.77

Standardized Indirect Effects of KSI on ETA

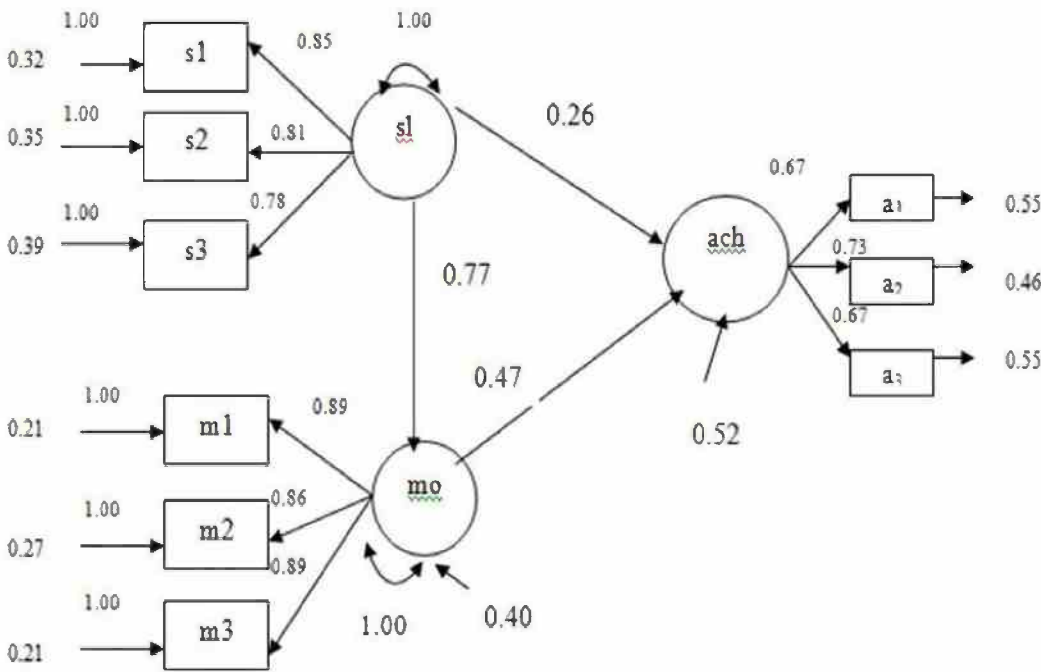
	sl

ach	0.36

Standardized Total Effects of ETA on ETA

	ach	mo
-----	-----	
ach	- -	0.47

وفيما يلي شكل المسار في ضوء الحلول المعيارية :



الشكل (15.10): الحلول المعيارية لنموذج SEM.

وفى الحلول المعيارية تعتمد على الدرجات المعيارية حيث لا يوجد وحدة قياسية للمؤشرات على متغيراتها الكامنة.

والأفضل من حيث المقارنة هو الاعتماد على الحلول المعيارية.

إذا المعادلات البنائية فى حالة الحلول اللامعيارية كالآتى:

$$\text{mo} = 1.01 \text{ SL} + 46.22 \quad \text{خطا التنبؤ}$$

$$\text{ach} = 0.26 \text{ sl} + 0.22 \text{ mo} + 13.13$$

$$\text{mo} = 0.77 \text{ SL} + 0.40 \quad \text{وفى الحلول المعيارية:}$$

$$\text{ach} = 0.26 \text{ SL} + 0.47 \text{ Mo} + 0.52$$

إرشادات وتوجيهات للاستخدام الأمثل لـ SEM:

فى ضوء الدراسات السابقة لمبادئ وأسس نمذجة المعادلة البنائية وكذلك الدراسات التقييمية لهذه المنهجية فى مختلف التخصصات يمكن عرض أهم الارشادات والتوجيهات التى يجب أن يلتزم بها الباحثين بقدر الإمكان وهى كالآتى (Breckler, 1990; Hu & Bentler, 1999; Kline, 2011; MacCallun & Austin, 2000; McDonald & Ho, 2002; Schumacker & Lomax, 2010) عامر، 2016 :

أولاً: مرحلة تخصيص النموذج:

- بناء النموذج فى ضوء أسس نظرية قوية أو نتائج دراسات سابقة وأن كان يفضل بناؤه فى وجود نظرية.
- اعرض أو ارسم النموذج النظرى فى شكل مسار يظهر فيه طبيعة التأثيرات ويفضل وضع إشارة المسار وأيضا توضيح ما إذا وجدت قيم معالم مثبتة fixed أو مقيدة constrained وأن يلتزم الباحث بالقواعد المتعارف عليها فى عرض الشكل. واعطى تفسير منطقى للقيود والفروض على معالم محددة ومدى ارتباطها بمنهجية تحليله أو مدى تفسيرها فى ضوء نظرية أو دراسات

- تمثيل المتغيرات الكامنة بعدد مناسب من المؤشرات وهنا يجب أن يمثل المتغير الكامن بثلاثة أربعة مؤشرات وهذا هو العدد الأمثل وذلك لأن تمثيل المتغير الكامن بعدد كبير من المؤشرات يؤدي إلى تعقيد النموذج.

- صغ نماذج بديلة على أسس نظرية خاصة إذا كان هدف الباحث هو المقارنة بين نماذج وكذلك صياغتها في حالة اختبار نموذج توكيدي فقط وذلك لأن صياغة نماذج بديلة يجنب الباحث من مشكلة التحيز التوكيدي في حالة اختبار نموذج واحد فقط بالتالي يحد من القدرة التعميمية للنموذج في المجتمع.

ثانيًا: قضايا مسح البيانات وإجراءات تقدير النموذج:

- لابد أن يتم تحليل نموذج SEM في ضوء حجم عينة مناسب ولا بد من تحديد عدد المعالم المقدرة حتى يتم تحديد حجم العينة في ضوءها والمتفق عليه بين الباحثين أن حجم عينة أكبر من 100 فرد مقبول ويفضل أن يكون 200 فرد فأكثر ويتوقف تحديد حجم العينة حسب حجم النموذج بمتغيراته المقاسة والكامنة وعوامل أخرى.

- بقدر الإمكان حاول تحديد حجم العينة في ضوء مستوى قوة إحصائية مرغوب يحاول الباحث تحقيقه.

- التأكد من توافر خاصية الاعتدالية للبيانات الداخلة في تحليل نموذج المعادلة البنائية وذلك لأن طرق التقدير مثل ML و GLS يتطلب توافر مسلمة الاعتدالية.

- إذا توافرت عدم الاعتدالية بدرجة خفيفة أو متوسطة (تحديدها من خلال التفرطح والالتواء) فإنه يمكن استخدام طرق التقدير السابقة لأنها تتميز بالضلاعة أو المناعة وأما إذا كانت عدم الاعتدالية شديدة فإنه أما أن يلجأ الباحث إلى إجراء تحويل للبيانات أو استخدام طريقة التقدير ADF.

- حدد كيف تحققت من وجود التلازمية الخطية بين المتغيرات وإن وجدت فيجب معالجتها.

- بقدر الإمكان اعتمد على مصفوفة التغاير كمدخل لتحليل نموذج المعادلة البنائية وذلك لأن استخدام مصفوفة الارتباط يؤدي إلى تقديرات للمعالم غير متسقة ولا بد

من عرض مصفوفة التغاير أو مصفوفة الارتباط مقرونة بالانحراف المعياري للحصول على تقديرات صحيحة للأخطاء المعيارية.

- تأكد أن محدد المصفوفة المراد تحليلها موجب لأنه إذا كان سالب فلا داعى لإجراء التحليل ويجب إعادة النظر فى تحديد النموذج مرة أخرى أو فحص البيانات.
- يجب إجراء تحليل نموذج المعادلة البنائية على خطوتين هما تحليل النموذج المقاس أولاً ثم تحليل النموذج البنائى وخبراء SEM يفضلون هذا المدخل عن مدخل الخطوة الواحدة.

ثالثاً: تقييم مطابقة النموذج:

- ضرورة عرض مؤشر χ^2 مقروناً بدرجات الحرية ودلالاتها الإحصائية خاصة إذا كانت طريقة التقدير ML.
- أكد خبراء نمذجة المعادلة البنائية على أهمية الاعتماد على مؤشر χ^2 بقدر الإمكان مقرونة بمؤشر يعتمد على تحليل البواقي مثل RMSEA أو SRMR أو كليهما بالإضافة إلى المؤشرات المتزايدة CFI و NNFI و RNI وذلك لأنهم أقل تأثراً بحجم العينة وأكثر حساسية لسوء تحديد النموذج.
- القبول بقاعدة 0.90 غير مناسبة لبعض المؤشرات ويجب أن تكون 0.95 فأعلى على المؤشرات NFI و NNFI و GFI و AGFI و CFI و RNI. ولكن يجب التأكيد على أن هذه القاعدة تكون للنماذج التى لا تتسم بالتعقيد وثبات مؤشراتها على.
- لا تعطى توصيات بأن النموذج المفترض هو الوحيد الذى يتطابق مع البيانات ولذلك فمن الأفضل صياغة نماذج بديلة قبل تحليل مطابقة النموذج المفترض لتحديد أيها أكثر مقبولة فى تفسير بيانات العينة.
- عند المقارنة بين نماذج بديلة اعتمد على مؤشرات البساطة PGFI و PNFI وكذلك مؤشر AIC.
- بالإضافة إلى تقييم مطابقة النموذج ككل يجب تقييم مطابقة المعالم أو المعادلات البنائية المختلفة للنموذج وذلك من خلال عرض مؤشرات R2 لكل معادلة بنائية.

- قيم النموذج ليس فى ضوء مؤشرات المطابقة ولكن فى ضوء النظرية وتفسير معالمه الفرعية فى ضوء المقبولية الواقعية.

رابعًا : تعديل النموذج:

- اعرض مؤشرات التعديل التى يعطيها البرنامج للنموذج المفترض لتحسين المطابقة.
- عدل النموذج المبدئى فى أضيق الحدود كلما أمكن.
- قبل التعديل هل تمت المقارنة بين مطابقة النماذج البديلة أو المتكافئة.
- وأن عدلت النموذج فلا بد أن تعتمد على الأسس الإحصائية ولا بد أن تكون مستندة إلى التفسيرات النظرية القوية لإضافة مسار أو حذفه.
- تجنب تعميم النتائج للنموذج المعدل إلا إذا أجرى له مصداقية عبر عينات أخرى Cross- Valadition.

خامسًا: التقرير والنتائج:

- اعرض الوصف الإحصائى للمتغيرات المقاسة المكونة فى مصفوفة التغير.
- اذكر تقديرات المعالم للنموذج النهائى مثل التقديرات غير المعيارية والأخطاء المعيارية والتقديرات المعيارية.
- إذا قبلت مطابقة النموذج البنائى فلا تدعى السببية بصورة مطلقة خاصة إذا كانت البيانات تولدت من تصميمات غير تجريبية.
- إذا قبلت مطابقة النموذج قدم تطبيقات مرتبطة بالمتغيرات المرتبطة بالنموذج فقط، فهل هذا أفاد فى قبول النظرية أم رفضها.
- اذكر كل معالم التقدير للنموذج مثل التشعبات ومعامل المسارات والأخطاء المعيارية والدلالة الإحصائية وتباينات الأخطاء وكذلك الحلول المعيارية.
- اعرض التأثيرات المباشرة وغير المباشرة والكلية بين المتغيرات الكامنة خاصة فى وجود متغيرات وسيطة.
- اعرض مؤشر حجم التأثير R^2 لكل معادلة بنائية فى النموذج البنائى.

- اعرض ما إذا وجدت مشاكل أثناء تقدير النموذج أو نتائج غير منطقية سواء للتأثيرات أو التباينات لان ذلك من شأنه أن يحد من تفسير النتائج وبالتالي تعميمها.
- وأخيرًا قدم Thompson (2000) عدة ارشادات وتوصيات عند تطبيق نمذجة المعادلة البنائية وهي كالآتي:
- لا تعتقد أن النموذج المفترض هو الوحيد المطابقة للبيانات فيجب صياغة نماذج أخرى بديلة يمكن أن تفسر العلاقات بين متغيرات النموذج.
- استخدم مدخل الخطوتين لـ Anderson & Gerbing (1988) عند التحقق من النموذج المقاس أولاً (التحليل العاملي التوكيدي) ثم تحقق من النموذج المقاس البنائي (نموذج العلاقات بين المتغيرات الكامنة) ثانيًا.
- قيم النماذج المفترضة في ضوء مؤشرات المطابقة والنظرية والاعتبارات العملية الأخرى.
- اعرض مؤشرات مطابقة من تصنيفات مختلفة.
- تحقق من المسلمات الواجب توافرها لتطبيق نمذجة المعادلة البنائية مثل الاعتدالية المتدرجة للبيانات.
- فضل وتحيز للنموذج الأكثر بساطة كلما أمكن.
- خذ في اعتبارك طبيعة مستوى القياس للمتغيرات (متصلة، منفصلة) وتوزيعها.
- لا تستخدم أحجام عينات صغيرة.
- التحقق من الصدق التعميمي Cross-validation من أي نموذج معدل بمعنى تم إدخال عليه تعديلات في ضوء أسس إحصائية أو نظرية.

الفصل الحادي عشر

النماذج والأساليب الإحصائية المستخدمة في البحث النفسي المصري والعربي

دراسة لعامر(2012)

المستخلص

هف الدراسة لتقييم استخدام الماذج والأساليب الإحصائية المستخدمة في البحث النفسى المصرى والعربى. وتضمنت الدراسة (39) دراسة منشورة فى مجلات تزوية فى مصر، تضمنت (199) اختبار احصائي (36) دراسة فى مجلات تزوية فى العالم العربى تضمنت (182) اختبار؛ وذلك فى الفترة الزمنية من (2000) حتى (2011). وتوصلت النتائج بالنسبة للبحث النفسى التربوى المصرى؛ أكثر الاختبارات الاحصائية استخداماً هى اختبارات النموذج البسيط (77.4%)، يليها اختبارات النموذج المتعدد (14%). وبالنسبة لبحث النفسى الزوى العربى يتضح أن اختبارات النموذج البسيط أكثر استخداماً (73.1%)، يليه اختبارات النموذج المتعدد (16.48%). وأكثر الاختبارات الاحصائية استخداماً فى البحث المصرى هو اختبار T -Test (38.2%)، يليه معامل ارتباط بيرسون (26.6%) وبالنسبة لبحث النفسى الزوى العربى فإن أكثر الاختبارات الاحصائية استخداماً معامل هو ارتباط بيرسون (35.2%) يليه اختبار T -Test (24.7%)، وبلغت نسبة استخدام الاختبارات المترية فى البحث النفسى المصرى (5%)، وفى البحث النفسى العربى (1.1%)، واتضح عدم تطور استخدام لاساليب احصائية عبر الزمن بالرغم من التزم الهائى فى الاحصاء وبرمجها الكمبيوترية.

مقدمة

أصبحت الاحصاء عامل أساسى فى إجراء البحوث النفسية والفهم العميق لاستخدام وتوظيف الإحصاء فى البحوث من المتطلبات الرئيسية لإعداد الباحثين النفسيين وذلك لمساعدتهم على تنفيذ وتصميم المنهجية البحثية (Mee, Lan, & chin, 2009). وأصبحت الاحصاء تالزم الباحث فى مشروع بحثه منذ مراجعة الدراسات السابقة وذلك من خلال توظيف منهجية ما وراء التحليل Meta-analysis لمراجعة التراث البحثى مراجعة كمية موضوعية بعيداً عن التحيز، وتدخلت فى صياغة الأسئلة من خلال تضمين متغيرات الدراسة وتصنيفها متصلة أو متقطعة ومستقلة أو تابعة، وهذا يعطى مؤشراً للباحثين عن نوع النموذج الإحصائي الذي بصدد ترجمة

الظاهرة النفسية إلى متغيرات يجمعها نموذج يشرح طبيعة العلاقة بين المتغيرات المتضمنه في دراسته.

وتتنوع الأساليب الإحصائية بتنوع النماذج الإحصائية التي تتبع لها الظاهرة موضوع الدراسة وتتضمن النموذج الإحصائي الصفرى (الصياد، 1985، 2002)، وهو النموذج الذى يتعامل مع متغير واحد فقط، والنموذج الإحصائي البسيط الذى يتعامل مع متغير مستقل وتابع واحد، والنموذج الإحصائي المتعدد الذى يتعامل مع أكثر من متغير مستقل وتابع واحد، والنموذج الإحصائي المتعدد المتدرج الذى يتضمن متغير مستقل واحد فأكثر وأكثر من تابع، والنموذج الإحصائي المتدرج الذى يتعامل مع عدد من المتغيرات بحيث يصعب تحديد المتغيرات المستقلة والتابعة (الصياد، 1985؛ 2002) ولكل نموذج إحصائي حالات ولكل حالة اختبار إحصائي يتعامل معها. وتتعدد الاختبارات الإحصائية من الاحصاء البارامترى الي الاحصاء اللابارامترى حيث يستخدم الإحصائي اللابارامترى عندما تكون المتغيرات المستقلة والتابعة من مستويات القياس الاسمية والرتبية.

ومن خلال استقراء استخدامات الاختبارات الإحصائية والنماذج الإحصائية فى البحث النفسى التربوى العربى يلاحظ الآتي:

1. يتعامل الباحثون مع الظاهرة النفسية من منظور ذرى بسيط وما يترتب على ذلك من استخدام النموذج الإحصائي البسيط وما يتضمنه من اختبارات احصائية منذ بداية الثمانينات وحتى الآن.
2. على الرغم من التقدم الهائل الذى حدث فى استخدامات الاحصاء فى البحث النفسى، إلا أن البحث العربى مازال يركز على استخدام اختبارات احصائية بسيطة سواء أكانت مقاييس علاقة أو اختبارات فارقة.
3. ما زال توظيف النموذج الإحصائي المتعدد المتدرج والمتدرج فى البحث النفسى التربوي العربى شبه منعدم.
4. يفتقد البحث النفسى التربوي العربى التعامل مع الظاهرة النفسية كما تحدث في واقعيتها من حيث تعقيدها وتشابكها وتداخلها؛ حيث أن الاختبارات الإحصائية

المثلى للتعامل معها هو تحليل التباين المتدرج (MANCOVA)، وتحليل التباين المتدرج (MANOVA) وتحليل المسار باستخدام الليزرال والتحليل العاملي والمعادلة البنائية الخطية وغيرها واستخدام هذه الأساليب ضئيل جداً في البحث النفسي التربوي العربي.

5. غالباً ما يستخدم الباحثون في العالم العربي الاختبار الاحصائي دون التحقق من شروط استخدامه وهذا يحدث أيضاً على مستوى البحث في العالم الغربي.

والدراسة تهدف إلى تقييم استخدام كلاً من الاختبارات الاحصائية والنماذج الاحصائية في البحث النفسي التربوي في كل من مصر والعالم العربي وذلك في الفترة من 2000 إلى 2011 وذلك في الدراسات المنشورة في المجالات النفسية والتربوية في مصر والعالم العربي ما بعد الدكتوراة.

ويعانى التراث البحثي ندرة شديدة في مثل هذه الدراسات التقييمية ولكن يوجد في التراث دراسات تقييم الاختبارات الاحصائية في المجالات التربوية منذ فترة زمنية ليست ببعيدة مثل Bangret & Bambreger (2005); Goodwin & Goodwin (2002); Onwuegbuzie (2002); Mee et al., (2001); (1985)، وهذا الاتجاه يهدف لتوعية الباحثين للإجراءات الاحصائية والتصميم الاحصائي الأمثل في البحوث المنشورة أو التوجيه لأنواع الاختبارات الاحصائية المناسبة للتصميم البحثي، ولم يعثر الباحث على دراسة تقييمية للاختبارات الاحصائية والنماذج الاحصائية في تخصص ما مثل علم النفس والارشاد النفسي وبحوث الشخصية وصعوبات التعلم وغيرها.

وقام Goodwin & Goodwin (1985) بتحليل (150) دراسة منشورة في الفترة الزمنية من 1979 حتى 1983 في مجلة علم النفس التربوي وصنف الأساليب الإحصائية إلى ثلاثة تصنيفات هم المستوى الأساسى ويتضمن اختبارات الاحصاء الوصفى ومعامل ارتباط بيرسون واختبار كاي² واختبار T وتحليل التباين الأحادى (ANOVA)، والمستوى المتوسط ويتضمن تحليل التباين المتعدد والمقارنات البعدية المتعددة وتحليل التباين (ANCOVA) والانحدار المتعدد، والمستوى المتقدم ويتضمن أسلوب التحليل التمييزى وتحليل التجمعات وتحليل المسار ومعامل الارتباط المتعدد

والتحليل العاملي وتحليل التباين المتدرج (MANOVA) وتحليل التباين المتدرج (MANCOVA). وتوصلا إلى أن الأساليب الإحصائية من المستوى المتوسط أكثر توظيفاً واستخدماً (43%)، يليها الاختبارات الإحصائية من المستوى الأساسي (35%)، يليها الأساليب الإحصائية من المستوى المتقدم (12%). وكان أكثر الاختبارات الإحصائية استخداماً تحليل التباين الأحادي (F-Test) (16%)، يليه معامل ارتباط بيرسون (15%)، بينما الاحصاء الوصفي هو أقل الأساليب الإحصائية استخداماً. وفي تقييمه للأساليب الإحصائية المستخدمة في رسائل الماجستير في كلية التربية بجامعة أم القرى بمكة المكرمة وكلية التربية بجامعة الملك سعود بالرياض بالسعودية توصل النجار (1999) إلى أن غالبية الأساليب الإحصائية من المستوي المتوسط وأن أكثر الأساليب الإحصائية استخداماً هو اختبار كا² والاستخدام المناسب للأساليب الإحصائية ضعيف في الكليتين. وتوصل (2002) Onwuegbuzie في فحصه لمجلد مجلة علم النفس التربوي البريطانية لعام 1998 وجد أن اختبار تحليل التباين الأحادي أكثر استخداماً (38.9%)، يليه اختبار (MANOVA) (22.2%)، ثم التحليل العاملي الاستكشافي (19.4%)، واستخدم نفس التصنيف للاختبارات الإحصائية الذي تبناه (Goodwin & Goodwin 1985).

وتوصل (2004) Hutchinson & Lovell في تقييمه للخصائص المنهجية لـ (129) في الفترة الزمنية من 1996 إلى عام 2000 دراسة منشورة في مجلة Higher Education، و(104) دراسة منشورة في Review of Higher Education، و(169) دراسة منشورة في Research of Higher Education، وتضمنت هذه الدراسات (٢٥٢) دراسة استخدمت التحليل الكمي واشتملت على (777) اختبار وجدوا أن أكثر الأساليب الإحصائية استخداماً هي تحليل الانحدار (اللوجاريتمي والمتعدد) (56.7%)، يليه معامل الارتباط (39.3%)، يليه التحليل العاملي (26.2%)، وتوزعت إلى (22.2%) تحليل عاملي استكشافي و(4%) تحليل عاملي توكيدي وتحليل تباين أحادي (21.8%) واختبار T لعينات مستقلة بنسبة (13.5%)، تحليل التباين المتعدد المتدرج (MANOVA) بنسبة (6.3%).

وأجرى (2005) Bangret & Baumberger تقييم الأساليب الإحصائية الموظفة في الدراسات المنشورة في مجلة الارشاد النفسى والنمو لـ (745) أسلوب احصائي منشور في هذه المجلة من عام 1990 حتى 2001 ووجد أن الاختبارات الإحصائية الأساسية هي أكثر استخداماً (62%)، يليها الاختبارات في المستوى المتوسط (22%)، ثم الاختبارات الإحصائية من المستوى المتقدم (16%)، وتوصلا إلى أن أكثر أن الاختبارات الإحصائية استخداماً هي الاحصاء الوصفى (31%)، يليه معامل ارتباط بيرسون (12%)، ثم الانحدار المتعدد (8%) ثم تحليل التباين الأحادى الاتجاه (7%)، و MANOVA و MANCOVA (7%).

وتوصل (2009) Mee et al. في تقييمه لـ (250) أسلوب احصائي مستخدمة في (77) أطروحة ماجستير ودكتوراة في الفترة الزمنية من 2001 إلى 2006 وذلك في مجال التربية في جامعات ماليزيا واعتمد على نفس التصنيف الذي اتبعه (1985) Goodwin & Goodwin للاختبارات الإحصائية وتوصلا إلى أن الاحصاء من المستوى الأساسى أكثر استخداماً (70.4%)، يليه الاحصاء من المستوى المتوسط (18.4%) ثم الإحصاء من المستوى المتقدم (11.2%)، وأكثر الاختبارات الإحصائية استخداماً هي الاحصاء الوصفى (19.6%)، يليه معامل ارتباط بيرسون (16%) ثم اختبار T-Test (16%) ثم تحليل التباين الأحادى الاتجاه (13.2%)، ثم الانحدار المتعدد (12.4%) ثم الاختبارات اللابارمترية (5%)، على الجانبين وجد أن توظيف الأساليب الإحصائية المتقدمة مثل تحليل المسار (0.4%) والمعادلة البنائية الخطية (0.8%) بينما يندم استخدام التحليل التمييزى وتحليل التجمعات والانحدار اللوغاريتمى.

وتوصل (2009) Kashy, Donnellan, Ackerman, & Russell لتقييمهم لـ (155) بحث منشور في مجلد علم النفس الاجتماعى والشخصية للنصف الاول من عام 2007 الي أن أكثر الاساليب الإحصائية استخداماً هي اختبار تحليل التباين الاحادي (52%)، يليه الانحدار المتعدد (41%)، و (5%) من الدراسات استخدمت التحليل العاملي الاستكشافي، و (11%) استخدمت المعادلة البنائية الخطية أو التحليل

العاملية التوكيدي، و(11%) استخدمت تكنيكات احصائية اخري مثل الانحدار اللوغاريتمي وما وراء التحليل.

وفي ضوء العرض السابق يمكن صياغة مشكلة البحث في الاسئلة الآتية:

1. ما واقع استخدام النماذج الاحصائية في البحث النفسى التربوى المصرى والعربى؟
2. ما واقع استخدام الاختبارات الاحصائية في البحث النفسى التربوى المصرى والعربى؟
3. ما أهم الأساليب الاحصائية التى كان يجب استخدامها فى البحوث النفسية والتربوية؟
4. هل تغير استخدام الأساليب الاحصائية فى ضوء النموذج الاحصائى عبر الفترة الزمنية من 2000 حتى 2011؟

الأهمية

تتبع الأهمية الدراسة فى اطار منهجها الوصفى التحليلي للدراسات إلى:

1. تقييم استخدام الاحصاء فى البحوث النفسية فى مصر والعالم العربى.
2. اقتراح أهم الأساليب الاحصائية المثلى التى يجب أن تستخدم فى البحوث النفسية التى تتضمنها الدراسة.
3. توجيه الاهتمام الي زيادة مقررات المنهجية البحثية والاحصاء وتصميم التجارب لطلاب الماجستير والدكتوراة في كليات التربية لتحسين الجودة البحثية والاحصائية.

الطريقة والمنهجية

العينة: بلغ عدد الدراسات النفسية المنشورة فى بعض المجالات النفسية فى مصر فى كليات التربية بالاسماعيلية والزقازيق والمنوفية وبنها والمنيا وأسيوط والمنصورة وقنا حوالى (39) دراسة تضمنت (199) اختبار احصائي والدراسات النفسية المنشورة فى المجالات العربية فى قطر والبحرين وسوريا وفلسطين والسعودية والمغرب حوالى (36) دراسة تضمنت (182) اختبار احصائي وذلك فى الفترة الزمنية من 200 حتى 2011 ، وهذه الدراسات فى مجال علم النفس التربوى والصحة النفسية والتربية الخاصة.

استمارة جمع البيانات: تم اعداد استمارة لجمع البيانات تضمنت عنوان الدراسة، والتوثيق، أسئلة البحث، النموذج الاحصائى، الاختبارات الاحصائية المستخدمة فى الدراسة، عددها، التحقق من

مسلمات استخدام الاختبار الاحصائي (نعم - لا)، ونوع الاحصاء (بارامترى - لا بارامترى). وتم تقدير هذه البيانات عن طريق مقدرين احدهما متخصص في مجال القياس والاحصاء التربوي والآخر باحث في مرحلة الماجستير.

تصنيف الاختبارات الاحصائية

تم تصنيف الاختبارات الاحصائية فى ضوء النموذج الاحصائي وتضمنت النموذج الصفرى (المتغير الواحد) وهو الاحصاء الوصفى واختبار T لعينة واحدة واختبارات النموذج البسيط وهى (معامل ارتباط بيرسون واختبار T لعينتين مستقلتين ومرتبطين، اختبار F ، الانحدار البسيط) والاحصاء اللابارمترى (مان ويتنى، ويلكسون، كروسكال وللاس،)، واختبارات النموذج المتعدد وهى (تحليل التباين المتعدد، الانحدار المتعدد، تحليل التباين الأحادى (ANCOVA)، واختبارات النموذج المتعدد المتدرج وهى (اختبار T هوتننج، MANOVA، MANCOVA)، واختبارات النموذج المتدرج وهى (تحليل عاملى توكيدى واستكشافى وتحليل مسار وتحليل تمييزى وتحليل التجمعات ونموذج المعادلة البنائية).

الاجراءات :

1. تم الاطلاع على المجالات التى تصدر من كلية التربية أو أي مؤسسة تربوية أخرى.
2. تم انتقاء الدراسات النفسية التى تقع فى تخصص الصحة النفسية أو علم النفس التربوى أو التربية الخاصة والتى تعتمد على المنهج الوصفى و الارتباطى والمنهج التجريبي وشبه التجريبي.
3. يمكن للدراسة أن تتضمن على عدد من الاختبارات الاحصائية للتحقق من فروض البحث والفروض الاحصائية.
4. تم تقدير البيانات الخاصة بكل دراسة علي حدة عن طريق المقدر الاول ومراجعة المقدر الثاني.

النتائج:

بعد مسح الدراسات النفسية والتربوية فى عدد من المجالات النفسية فى مصر والعالم العربى وفيما يلى توزيع هذه الاختبارات فى البحث المصرى والبحث النفسى العربى فى ضوء النماذج الاحصائية التى تم توظيفها فى هذه الدراسات.

الجدول (1): التكرارات والنسب المئوية للنماذج و الأساليب الاحصائية المستخدمة فى البحث النفسى والتربوى المصرى والعربى فى ضوء النموذج الاحصائي.

النموذج الاحصائي		البحث المصرى النفسى		البحث العربى النفسى	
الاسلوب الاحصائي		التكرار	%	التكرار	%
نموذج المتغير الواحد (الصفى)		7	3.5	15	8.2
الاحصاء الوصفى		0	0	2	1.1
اختبار "ت" لعينة واحدة					
المجموع		7	3.5	17	9.3
النموذج البسيط					
معامل ارتباط بيرسون أو سبيرمان		53	26.6	64	35.2
اختبار T – Test مستقلة ومرتبطة		76	38.2	45	24.7
اختبار F– Test		9	4.5	21	11.5
الانحدار البسيط		2		0	0
الاختبارات اللابارتية		14	70.	3	1.6

73.1	133	77.4	154	
				النموذج المتعدد
9.3	17	5.5	11	تحليل التباين المتعدد
3.9	7	8.5	17	تحليل الانحدار المتعدد
3.3	6	0	0	تحليل التغيرات المتعدد
16.48	30	14	28	المجموع
				النموذج المتعدد المترج
0	0	0	0	اختبار T هوتلنج
0	0	0	0	تحليل التباين المترج MANOVA
0	0	0.5	1	تحليل التغيرات المترج MANCOVA
0	0	0.5	1	المجموع
				النموذج المترج
1.1	2	2	4	تحليل عاملي استكشافي
0	0	1	2	تحليل عاملي توكيدي
0	0	0.5	1	تحليل مسار
0	0	0	1	معادلة بنائية خطية
0	0	0	1	تحليل تميزي
0	0	0	0	تحليل تجمعات

1.1	2	4.5	9	المجموع
	182		199	الاجمالي

كما هو واضح في الجدول (1) بالنسبة للبحث النفسي التربوي المصري يظهر أن أكثر اختبارات النماذج الاحصائية توظيفاً هي اختبارات النموذج البسيط (77.4%)، يليها اختبار النموذج المتعدد (14%)، ثم اختبار النموذج المتدرج (4.5%) وأقلهم استخداماً هي اختبارات النموذج المتعدد المتدرج (0.5%)، وإذا كانت اختبارات النموذج الصفري والبسيط تقابل المستوى الأساسي واختبارات النموذج المتعدد تقابل المستوى المتوسط واختبارات النموذج المتعدد المتدرج والنموذج المتدرج تقابل المستوى المتقدم ذلك التصنيف الذي تبنته الدراسات السابقة فإن هذا يخالف ما توصل إليه Goodwin (1985) & Goodwin حيث وجد أن اختبارات المستوى المتوسط أكثر شيوعاً (43%) في حين أن الدراسة الحالية توصلت الي أن اختبارات المستوى الأساسي أكثر استخداماً (77.4%) وهذا يتفق مع ما توصل إليه Bangret & Baumberger (2005) حيث توصلوا إلى أن اختبارات المستوى الأساسي أكثر استخداماً (62%) توصل مع Mee et al. (2009) من أن اختبارات المستوى الأساسي أكثر استخداماً (70.4%).

وبالنسبة للاختبارات الاحصائية المستخدمة في البحث النفسي التربوي المصري يظهر من الجدول السابق أن أكثر الاختبارات استخداماً اختبار T (38.2%) ، يليه معامل ارتباط بيرسون (26.6%) ، ثم تحليل الانحدار المتعدد (8.5%) ، أما أقل الاختبارات استخداماً هي اختبار تحليل التباين المتعدد المتدرج MANOVA وتحليل التجمعات (0.0%) وتحليل التغاير المتعدد (MANCOVA) (0.5%)، وتحليل المسار (0.5%)، والمعادلة البنائية الخطية (0.5%). وهذا يتفق جزئياً مع Bangret

& Baumberger 2005; Goodwin & Goodwin, 1985; Mee et al., (2009)، ويتناقض مع Kashy et al., (2004; Huttchinson & Lovell, 2009; Onwuegbuzie, 2002)، حيث وجدوا أن أكثر الاختبارات استخداماً تحليل التباين الأحادي يليه MANOVA ثم التحليل العاملي الاستكشافي.

وكما يتضح من الجدول (1) بالنسبة للبحث النفسى التربوى العربى أن اختبار النموذج البسيط أكثر استخداماً 73.1% يليه اختبار النموذج المتعدد 16.48، ثم اختبارات النموذج الصفري 8.3% ثم النموذج المتدرج 1.1% ثم اختبارات النموذج المتعدد المتدرج 0.0%. وإن أكثر الاختبارات الاحصائية استخداماً هو معامل ارتباط بيرسون (35.2%) يليه اختبار T (24.7%) ثم تحليل التباين المتعدد (9.3%) وأقل الاختبارات استخداماً هو اختبارات النموذج المتعدد المتدرج (MANOVA - Hotteling T² - MANCOVA) واختبارات النموذج المتدرج (تحليل عاملي توكيدى - تحليل مسار - معادلة بنائية خطية وغيرها) وهذه النتائج تتعارض جزئياً مع (Goodwin & Goodwin, 1985) فى أن النموذج الاحصائى الأكثر استخداماً وتتفق معهم فى الاختبارات الاحصائية الأكثر استخداماً ولكن النتائج تتفق مع (Bangret & Baumberger, 2005; Mee et al., 2009) فى أن المستوى الأساسى أكثر استخداماً يليه المتوسط ثم المتقدم ولكن النتائج تتناقض مع ما توصل إليه (Huttchinson & Lovell, 2004; Kashy et al., 2009; Onwuegbuzie, 2002).

أما بالنسبة لمقارنة النتائج فى حالة البحث التربوى النفسى العربى بنظيره المصرى يلاحظ أن كلاهما استخدام اختبارات النموذج البسيط (77.4%) فى البحث المصرى و(73.1%) فى البحث العربى، ويفتقد إلى استخدام اختبارات النموذج المتعدد المتدرج (0.5%) فى البحث المصرى و(0.0%) فى البحث العربى ولكن البحث المصرى تفوق فى استخدام اختبارات النموذج المتدرج (الاحصاء المتقدم) على نظيره العربى فبلغت فى البحث المصرى (4.5%) وفى البحث العربى (1.1%).

ويلاحظ أن أكثر الاختبارات الاحصائية استخداماً في البحث التربوى النفسى المصرى هو اختبار T (38.2%) وفى البحث العربى اختبار معامل ارتباط بيرسون (35.2%). أما بالنسبة لاستخدام الاختبارات اللابارمترية يتضح أن البحث التربوى النفسى المصرى أكثر استخداماً لها حيث بلغ نسبة استخدامها فى البحث التربوى النفسى المصرى (7.4%) وفى البحث العربى النفسى (1.6%).

أما بالنسبة للأساليب الاحصائية التى يجب على الباحثين استخدامها فى البحث التربوى النفسى المصرى والعربى. وفحص المتغيرات المتضمنة فى الدراسات وتصنيف هذه المتغيرات حسب اتصالياتها يلاحظ أن أغلب الدراسات تهدف إلى دراسة فروق أو علاقة أو تنبؤ، حيث يمكن للدراسة الواحدة أن تتضمن الأهداف الثلاثة ويتناولها الباحثون من وجهة نظر اختبارات فارقة (T و F وغيرها)، ومن وجهة نظر اختبارات ارتباطية (بيرسون) أو تنبؤية (انحدار)، ويتفحص أسئلة الباحثين وفروضهم فكان من الأمثل استخدام تحليل التباين (النموذج المتعدد)، وتحليل التباين المتدرج (MANOVA) وتحليل التباين المتعدد المتدرج (MANCOVA) واختبار Hotelling T² (النموذج المتعدد المتدرج).

ولاختبار مدى اختلاف توظيف الاختبارات الاحصائية عبر الفترتين الزمنية من 2000 حتى 2006، ومن 2007 حتى 2011 لاختبار مدى اختلاف توظيف الاختبارات الاحصائية عبر الفترتين السابقتين حتى نتأكد أن استخدام أساليب احصائية متقدمة ينمو ويتطور استخدامه بمرور الزمن، وفيما يلى جدول التكرارات للاختبارات الاحصائية عبر الفترتين الزمنية.

الجدول (2): التكرارات للاختبارات الاحصائية عبر الفترتين الزمنيتين.

المجموع	٢٠١١-٢٠٠٧	٢٠٠٦-٢٠٠٠	
٣٠٣	١٦١	١٤٢	اختبارات النموذج البسيط وما قبله
٧٨	٤١	٣٧	اختبارات النموذج المتعدد وما بعده
٣٨١	٢٠٢	١٧٩	

ولكى نتأكد من مسلمات الاختبار نلاحظ أنه لا توجد خلية بها تكرار أقل من "٥" وباجراء اختبار X^2 لاستقلالية متغيرين تصنيفين . يلاحظ أن قيمة X^2 غير دالة احصائياً ($X^2=0.013$, $df=1$, $p=0.91$) وهذا يعطى مؤشر على عدم تطور استخدام الأساليب الاحصائية عبر الزمن، وهذا مفاده ان توظيف الاساليب الاحصائية في بداية الالفية الثالثة لم يتغير عن استخدامه في العقد الماضي وإن جاز لي التعبير ان الباحثين يتعاملون مع الاحصاء في دراساتهم بمنطق وعقلية الباحث في الثمانينيات رغم التقدم المذهل الحادث في التكنيكات الاحصائية وتعاملها مع الظواهر الانسانية بصورة أكثر تعقيداً وليس بصورة ذرية بسيطة.

المناقشة والتعليق

تشير نتائج الدراسة إلى أن 81% من الاختبارات الاحصائية المستخدمة في البحث المصرى النفسى تقع فى النموذجين الصفرى والبسيط و يقابلها (82.4%) في البحث العربى النفسى و هذا يتفق مع (Mee et al., (2009 والنجار (١٩٩١) والصياد (١٩٨٥)، (Bangret & Baumberger (2005) ولكنه يتناقض مع (Goodwin & Goodwin (1985) فى تقييمه لأبحاث تخصصه فى مجال علم النفس التربوى. وهذا مفاده أن الباحثين النفسيين العرب والمصريين ما بعد الدكتوراة يتعاملون مع الظاهرة النفسية غير واعين بالنموذج الاحصائى الأمثل، وكذلك غير

مدركين بكيفية تحليل البيانات المرتبطة بدراساتهم وهذا ينطبق على الوضع المصرى والعربى على حد سواء. كما يتضح أن أكثر الاختبارات الاحصائية استخداماً فى الواقع المصرى اختبار "ت" (38.2%)، يليه معامل ارتباط بيرسون (26.2%)، وفى الواقع العربى معامل ارتباط بيرسون (35.2%)، يليه اختبار T (24.7%)، وهذا يتناقض مع (Huttchinson & Lovell, 2004; Kashy et al., 2009; Onwuegbuzie 2002). وهذا يدل على أن البحث التربوى النفسى المصرى والعربى ما زال يعتمد على اختبارات إحصائية بسيطة سواء اختبار T أو معامل ارتباط بيرسون على الرغم من التقدم الهائل الذى حدث فى الأساليب الاحصائية للتعامل مع الظواهر النفسية، وهذا يدعو إلى إجراء دورات تدريبية فى مجال الاحصاء النفسى المتقدم للباحثين على كافة مستوياتهم، وعلى الرغم من التقدم الهائل الحادث فى البرامج الاحصائية الكمبيوترية إلا أنه ما زال استخدام الاحصاء يعتمد على أساليب إحصائية بسيطة حيث ما يتقنه الباحثين النفسيين المصريين والعرب لا يتخطى 20% مما تتضمنه هذه البرامج الاحصائية مثل حزمة SPSS و LISREL و SAS وغيرها وهذا يدعو إلى إجراء دورات تدريبية لاستخدام البرامج الاحصائية فى العلوم الانسانية حتى يتسنى للباحثين توظيف مشتملات هذه البرامج من أساليب واختبارات احصائية.

ما زال البحث التربوى النفسى المصرى والعربى متدني لتوظيف لأساليب الاحصاء المتدرج Multivariate Statistics، حيث أن استخدمه البحث المصرى (5%) والبحث العربى (1.1%) من أساليب احصائية متدرجة وربما يعود هذا إلى عدم وعى الباحثين بأهمية استخدام هذه الأساليب حيث أنها الأنسب لدراسة الظاهرة النفسية من حيث تعقيدها وتشابكها وتفاعلها. وقد يعود هذا إلى أنها تحتاج إلى عمليات حسابية معقدة ولكن فى ظل وجود البرامج الاحصائية فإنها ميسرة لاستخدامها وتحتاج الي تفسير فقط. واتضح أن استخدام الإحصاء البسيط والمتدرج لا يختلف عبر المرحلة الزمنية من 2001 حتى 2011 وهذا يدل على أن الباحثون يتعاملون مع أساليب احصائية بسيطة على الرغم من التقدم الهائل الحادث في البرامج الاحصائية الكمبيوترية وهذا ينطبق على الباحثين على المستوى المصرى والعربى.

وأخيراً ما زال الباحث التربوي النفسى المصرى والعربى يتعامل مع الظاهرة النفسية من وجهة نظر ضيقة بسيطة غير مستفيد بالتغيرات الحادثة فى المنهجية الإحصائية والبرامج الكمبيوترية وهذا يدعونا إلى تدعيم المقررات الإحصائية فى برامج الدكتوراة بأساليب منهجية متطورة واجراء مزيداً من الدورات التدريبية للتوعية بالمنهجية الإحصائية المتدرجة وغيرها، وكذلك التدريب على البرامج الكمبيوترية فى هذا الشأن حتى يتسنى لهم انجاز دراساتهم بما يحقق مبدأ محاكاة الظاهرة النفسية كما تحدث فى واقعها وهذا يزيد من الصديق الخارجى والتعميمى لنتائج دراستنا فى مجال العلوم النفسية و التربوية. وعلي متخذي القرار التعامل مع توصيات الدراسات والبحوث النفسية والتربوية يشئ الحذر الشديد في تطبيق هذه التوصيات في الواقع الحقيقي حتي لا يتحول الانسان العربي الي فأر تجارب.

مراجع الدراسة

الصياد، عبد العاطى أحمد. (1985). النماذج الاحصائية في البحث التربوي والنفسي العربي بين ماهو قائم وما يجب أن يكون. الرياض: رسالة الخليج العربي، 16، 211-251.

الصياد ، عبد العاطى أحمد (2002، سبتمبر). نحو تصميم انسيابى مقترح لكيفية استخدام الطريقة الاحصائية فى إجراء الدراسات والبحوث الأمنية وموقع الاحصاء فى خطوات التفكير العلمى المتعارف عليها . الرياض: جامعة نايف العربية للعلوم الأمنية.

النجار، عبد الله عمر عبد الرحمن. (1991). دراسة تقويمية مقارنة لاساليب الاحصائية التي استخدمت في تحليل البيانات في رسائل الماجستير في كل من كلية التربية بجامعة أم القري بمكة المكرمة وكلية التربية بجامعة الملك سعود بالرياض. رسالة ماجستير غير منشورة. كلية التربية. جامعة أم القري بمكة المكرمة، السعودية.

- Bangert, A. W., & Baumberger, J. P. (2005). Research and Statistical Techniques used in The Journal of Counseling & Development: 1990 - 2001. *Journal of Counseling & Development*, 83, 480 – 488.
- Goodwin, L. D., & Goodwin, W. L. (1985). An analysis of Statistical techniques used in The Journal of Educational Psychology: 1979 – 1983. *Journal of Educational Psychologist*, 20, 13 –21.
- Hutchinson, S.R., & Lovell, C.D. (2004). A Review of methodological characteristics of Research published in key Journals in Higher Education: Implications for Graduate Research Training. *Research in Higher Education*, 45, 383-403.
- Kashy, D.A., Donnellan, M.B., Ackerman, R.A., & Russell, D.W. (2009). Reporting and interpreting research in PSPB: practices, principles, and pragmatics. *Personality and Social Psychology Bulletin*, 35, 1131-1142.
- Mee, T. M., Lan, O. S., & Chin, L. H. (2009). Statistical Techniques Employed in Educational Theses in Malaysia. *European Journal of Social Sciences*, 12, 269 – 276.
- Onwuegbuzie, A. J. (2002). Common analytical and interpretational errors in Educational Research: An analysis of the 1998 volume of the British Journal of Educational Psychology. *Educational Research Quarterly*, 26, 11 – 22.

المراجع

عامر، عبد الناصر السيد. (2007). حجم العينة في تحليل الانحدار المتعدد. *المجلة المصرية للدراسات النفسية*، 17.

عامر، عبد الناصر السيد. (2012). النماذج والإختبارات الإحصائية المستخدمة في البحث النفسي المصري والعربي. *المجلة المصرية للدراسات النفسية*، 22، 74، 355-371.

عامر، عبد الناصر السيد. (2014). تقييم استخدام تطبيقات نمذجة المعادلة البنائية في البحث النفسي. *مجلة دراسات عربية في علم النفس (رانم)*، 13، 2.

عامر، عبد الناصر السيد (2016). نمذجة المعادلة البنائية: بعض القضايا المنهجية والتوصيات. *المجلة المصرية للدراسات النفسية*، 26، 37-58.

عامر، عبد الناصر السيد. (2018). نمذجة المعادلة البنائية للعلوم النفسية والاجتماعية: الاسس و التطبيقات والقضايا (الجزء الأول). السعودية: دار جامعة نايف للنشر.

عامر، عبد الناصر السيد. (2018). نمذجة المعادلة البنائية للعلوم النفسية والاجتماعية: الاسس و التطبيقات والقضايا (الجزء الثاني). السعودية: دار جامعة نايف للنشر.

عامر، عبد الناصر السيد. (2019). الإحصاء التربوي والنفسى. مذكرة غير منشورة، كلية التربية، جامعة قناة السويس.

Akaike, H. (1987). Factor analysis and AIC. *Psychometrika*, 52, 317-332.

Anderson, J. C., & Gerbing, D. V. (1988). Structural equation modeling in practice: A review and recommended two – step approach. *Psychological Bulletin*, 103, 411-423.

- Anderson, J. C., & Gerbing, D. W. (1992). Assumptions and comparative strength of the two – step approach: Comment on Fornell and Yi. *Sociological Methods & Research*, 20, 321 – 333.
- Aron, A., Coups, E. J., & Aron, E. N. (2013). *Statistics for psychology* (6th.ed). Boston: Pearson Education, Inc.
- Bagozzi , R . P., & Yi, Y. (2012). Specifications, evaluation, and interpretation of structural equation models. *Journal of Academy of Marketing Science*, 40, 8-34.
- Bagozzi, R. P., & Heatherton, T. F. (1994). A General approach to representing multifaceted personality constructs: Application to state self-esteem. *Structural Equation Modeling*, 1, 35-67.
- Bagozzi, R. P., & Yi, Y. (1988). On the evaluation of structural equation models. *Journal of Academy of Marketing Science*, 16, 74 – 94.
- Bartlett, M. S. (1954). A note on the multiplying factors for various chi square approximations. *Journal of the Royal Statistical Society*, 16, 296-298.
- Bartlett, M. S. (1954). A note on the multiplying factors for various chi square approximations. *Journal of the Royal Statistical Society*, 16, 296-298.
- Baumgartner, H., & Homburg, C. (1996). Applications of structural equation modeling in marketing and consumer Research: A review. *International Journal of Research in Marketing*, 13, 139-161.
- Bentler, P . M., & Bonnett, D.G.(1980) . Significance tests and goodness – of - fit in the analysis of covariance structures. *Psychological Bulletin*, 112, 400-404.
- Bentler, P. M. (1990). Comparative Fit Indices in Structural Models. *Psychological Bulletin*, 107, 238-246.

- Bentler, P. M., & Chou, C. P. (1987). Practical issues in structural modeling. *Sociological Methods and Research*, 16, 78-117.
- Bentler, P. M., & Chou, C. P. (1987). Practical issues in structural modeling. *Sociological Methods and Research*, 16, 78-117.
- Bollen, K. A., & Arminger, G. (1991). Observational residuals in factor analysis and structural equation models. In P. V. Marsden (Eds.), *Sociological methodology* (PP. 235-262). Cambridge, MA: Blackwell.
- Bollen, K. A. (1989). *Structural equations with Latent variables*. New York: Wiley.
- Bollen, K. A., & Long, J. S. (1993). Introduction. In K. A. Bollen & J. S. Long (eds.), *Testing structural equation modeling* (pp. 1-9). Newburg park, CA: Sage.
- Boomsma, A. (2000). Reporting analysis covariance structure. *Structural Equation Modeling*, 7, 461- 483.
- Boomsma, A., & Hoogland, J. J. (2001). The robustness of lisRel modeling revisited . In R. Cudeck , S . dutoit , & D. Sorbom (Eds.), *Structural equation models : Present and future* (PP. 139-188). Chicago: Scientific Software International.
- Brannick, M. T. (1995). Critical comments on applying covariance structure modeling. *Journal of Organizational Behavior*, 16, 201-213.
- Breckler, S. T. (1990). Application of covariance structure modeling in psychology: Cause for concern? *Psychological Bulletin*, 107, 260 – 273.
- Browne, M. W. (1984). Asymptotic distribution free methods in the analysis of covariance structures. *British Journal of Mathematical and Statistical Psychology*, 37, 62-83 .

- Brown, T. A. (2006). *Confirmatory factor analysis for applied research*. New York: Guilford Press.
- Browne, M. W. (1984). Asymptotic distribution free methods in the analysis of covariance structures. *British Journal of Mathematical and Statistical Psychology*, 37, 62-83 .
- Browne, R. L. (1994). Efficacy of the indirect approach for estimating structural equation models with missing data: A comparison of five methods. *Structural Equation Modeling: A Multidisciplinary Journal*, 1, 287-316
- Browne, M. W., & Cudek, R. (1993). Alternative ways of assessing model fit. In K .A. Bollen & J. S. Long (Eds.), *Testing structural equation models* (PP. 136-162) . Newbury Park, CA: Sage.
- Burt, R. S. (1976). Interpretational confounding of unobserved variables in structural equation models. *Sociological methods & Research*, 5, 3-52.
- Cattell, R. B. (1966). The scree test for the number of factors. *Multivariate Research*, 1, 245-276.
- Chou, C. P., & Bentler, P. M. (1995). Estimates and tests in structural equation modeling. In. R. H. Hoyle (Eds.), *Structural equation modeling : concepts, issues, and applications* (PP. 37- 59). Thousand Oaks, CA: Sage.
- Cohen, J. (1988). *Statistical power analysis for the behavioral sciences* (2nd ed). Hillsdale: Lawrence Erlbaum Associates, Inc.
- Cohen, J., Cohen, P., West, S. G., & Aiken, L. S. (2003). *Applied multiple regression / correlation analysis for behavioral sciences* (3th.ed). Mahwah: Lawrence Erlbaum Associates, Publishers.
- Comery, A. L., & Lee, H. B. (1992). *A First course in factor analysis* (2nd .ed). Hillsadale, NJ: Erbaum.

- Cook, T. D., & Campbell, D. T. (1979). *Quasi-experimentation: Design and analysis issues for field settings*. Boston: Houghton Mifflin.
- Costello, A. B., & Osborne, J. W. (2005). Best practices in exploratory factor analysis: Recommendations for getting the Most from your analysis. *Practical Assessment, Research, and Evaluation*, 10, 1-9.
- Crockett, S. A. (2012). A five – step guide to conducting SEM analysis in counseling research. *Counseling Outcome Research and Evaluation*, 3, 30-47.
- Dattalo, P. (2008). *Determining sample size: Balancing power, precision, and practicality*. Oxford: University Press.
- Dunn, D.S. (2001). *Statistics and data analysis for behavioral sciences*. New York: McGraw-Hill Higher education, Inc.
- Engle, K .S., Mossbrugger, H., & Muller, R. O. (2003). Evaluating the fit of structural equation models test of significance and descriptive goodness of – fit measures. *Methods of Psychological Research Online*, 8, 23 – 74.
- Fabrigar, L. R., Wegener, D. T., MacCullum, R. C., & Strahan, E. J. (1999). Evaluating the use of factor analysis in psychological research. *Psychological Methods*, 4, 272-299.
- Fava, J. L., & Velicer, W. F. (1992). The effects of over extraction on factor and component analysis. *Multivariate Behavioral Research*, 27, 387-415.
- Field, A. (2009). *Discovering statistics using SPSS(3th.ed)*. London: Sage publications, LTD.
- Field, A. (2013). *Discovering statistics using SPSS (4th.ed)*. Sage Publications, Ltd.
- Fornell. C., & Yi, Y. (1992). Assumptions of the two – step approach to latent variable modeling. *Sociological Methods & Research*, 20, 291-320.

- Good, P. (1999). *Resampling Methods: A practical guide to data analysis*. Boston: Birkhauser.
- Gravetter, F. J., & Wallnau, L. B. (2014). *Essentials of statistics for the behavioral sciences* (8th.ed). Belmont: Wadsworth, Cengage Learning.
- Green, S. B. (1991). How many subjects does it take to do regression analysis? *Multivariate Behavioral Research*, 27, 499-510.
- Green, S. B., & Salkind, N. J. (2014). *Using SPSS for windows and macintosh: Analysing and understanding data* (7th.ed). Boston: Pearson Education, Inc.
- Grouch, R. L. (1983). *Factor analysis* (2nd.ed). Hillsdale, NJ: Lawrence Erlbaum.
- Guadagnoli, E., & Velicer, W. F. (1988). Relation of sample size to the stability of component patterns. *Psychological Bulletin*, 103, 265-275.
- Hair Jr. J. H., Anderson, R. E., Tatham, R. L., & Black, W. C. (1998). *Multivariate data analysis*. New Jersey: Prentice-Hall.
- Harlow, L. L. (2005). *The essence of multivariate thinking: Basic themes and methods*. New Jersey: Lawrence Erlbaum Associates, Publishers.
- Harris, R. J. (1985). *A primer of multivariate statistics*. New York: Academic.
- Henson, R. K., & Roberts, J. K. (2006). Use of exploratory factor analysis in published research: Common errors and comment on improved. *Educational and Psychological Measurement*, 66, 393-416.
- Hinkle, D. E., & Wiersma, W., & Jurs, S. G. (1994). *Applied statistics for the behavioral sciences* (3rd.ed). Boston: Houghton.

- Howell, D. C. (2013). *Statistics methods for psychology*(8th.ed). Belmont: Wadsworth, Cengage Learning.
- Hoyle, R. H. (1995). Structural equation modeling: Basic concepts and fundamental issues .*In R . H. Hoyle (Eds.), Structural equation modeling: Concepts, issues, and applications*(PP 1-15). Thousand Oaks, CA: Sage.
- Hoyle, R. H, & Panter, A. T. (1995). Writing about structural equation modeling. *In R. H. Hoyle (Eds.), Structural Equation Modeling: concepts, issues, and application* (PP.158-175) .thousand Oaks: Sage.
- Hu, L., & Bentler, P. M. (1995). Evaluating model fit. In R. H. Hoyle (Eds.), *Structural equation modeling: concepts, issues, and applications* (PP. 76- 99). Thousand Oaks, CA: Sage.
- Hu, L., & Bentler, P. M. (1998). Fit indices in covariance structure analysis: Sensitivity under parameterized model misspecification. *Psychological Methods* , 3, 424- 453.
- Hu, L., & Bentler, P. M. (1999). Cutoff criteria for fit indexes in covariance structure analysis. Conventional criteria versus new alternatives. *Structural Equation Modeling*, 6, 1- 55.
- Huck, S. W. (2012). *Reading statistics and research* (6th.ed). Boston: Pearson.
- Joreskog, K. G., & Sorbom, D. (1988). *LISREL 7: A guide to the program and applications* (2 nd. ed). Chicago: Illinois.
- Joreskog, K. G., & Sorbom, D. (1993). *LISREL 8: user's reference guide*. Mooresville, In: scientific software.
- Kaiser, H. F. (1960). An index of factorial simplicity. *Psychometrika*, 35, 401-415.
- Kaplan, D. (2000). *Structural equation modeling: Foundations and extensions*. Thousand oaks, CA: Sage.

- Keith, T. Z. (2014). *Multiple regression and beyond: An introduction to multiple regression and structural equation modeling* (2nd ed). New York: Routeledge.
- Kline, R. K. (2016). *Principles and practice of structural equation modeling* (4thed). New York: Guilford publications, Inc.
- MacCallum, R. C. (1986). Specification searches in covariance structure modeling. *Psychological Bulletin*, 100, 107-120.
- MacCallum, R. C, & Austin, J. T. (2000). Application of structural equation modeling in psychological research. *Annual Review of Psychology*, 51, 201-226.
- MacCallum, R. C., Roznowski, M., & Necowit, L. B. (1992). Model modifications in covariance structure analysis: The problem of capitalization on chance. *Psychological Bulletin*, 111, 490 – 504 .
- Marsh, H . W., Hau, T., & Grayson . D. (2005). Goodness of fit evaluation in structural equation modeling . In A. Mayday – Olivares & J . McCardle (eds.), *Psychometrics : A festschrift to Roderick P. McDonald* (PP. 275-340). Mahwah, NJ: Lawrence Erlbaum Associates, INC.
- Marsh, H. W., Balla, J. R., & McDonald, R. P. (1988). Goodness of fit indexes in confirmatory analysis: The effect of sample size. *Psychological Bulletin*, 103, 391-411.
- Martines, M. P. (2005). The use of structural equation modeling in counseling psychology research. *The Counseling Psychologist*, 33, 269- 298.
- Maruyama, G. M. (1998). *Basics of structural equation modeling*. Thousand Oaks: Sage publications, Inc.
- McDonald, R. P., & Ho, M. R. (2000). Principles and practice of reporting structural equation modeling. *Psychological Methods*, 7, 64- 82.

- McDonald, R. P., & Marsh, H. W. (1990). Choosing a multivariate model: No centrality and goodness of fit. *Psychological bulletin*, 107, 247.
- Mote, T. A. (1970). An artifact of the rotation of too few factors: Study orientation VS. trait anxiety. *Revista Interamericana De Psicologia*, 37, 267-305.
- Mulaik, S. A. (1990). Blurring the distinctions between component analysis and common factor analysis. *Multivariate Behavioral Research*, 25, 53-59.
- Mulaik, S. A. (2009). *Linear causal modeling with structural equations*. Boca Raton: Chapman & Hall-CRC.
- Mulaik, S. A., & Millsap, R. E. (2000). Doing the four step right. *Structural Equation Modeling*, 7, 36- 73.
- Mulaik, S. A., Jams, L. R., Alstine, J. V ., Bonnett, N., lind, S., & Stilwell, D. C. (1989) . Evaluation of goodness of fit indices for structural equation models. *Psychological Bulletin*, 105, 430-445.
- Muthén, L. K., & Muthén, B. O. (2002). How to use a Monte Carlo study to decide on sample size and determine power. *Structural Equation Modeling*, 4, 599-620.
- Nolan, S. A., & Henzien, T. E. (2012). *Statistics for the behavioral sciences*(2nd.ed). New York: Worth Publishers.
- Nunnally, J. C., & Bernstein, I. H. (1994). *Psychometric theory* (3rd.ed). New York: McGraw –Hill.
- Olejnik. S., & Algina, J. (2000). Measures of effect Size for comparative studies: Applications, interpretations, and limitations. *Contemporary Educational Psychology*, 25, 241 – 286

- Ormee, J. G. & Hudson, W. W. (1995). The problem of sample size estimation: Confidence intervals. *Social Work Research*, 19, 121-127.
- Pagano, R. P. (2013). *Understanding statistics in the behavioral sciences*(10th.ed). Belmont: Wadsworth, Cengage Learning.
- Pedhazur, E. J. (1997). *Multiple regression in behavioral research* (3rd.ed). Australia: Wadsworth, Thomson Learning.
- Privitera, G. J. (2015). *Statistics for the behavioral sciences*(2nd.ed). Canada: Sage Publications, Inc.
- Quintana, S. M., & Maxwell, S. E. (1999). Implications of recent development in structural equation modeling for counseling psychology. *The Counseling Psychologist*, 27, 485- 537.
- Raykov, T., & Marcoulides, G. A. (2006). *A first course in structural equation modeling* (2nd. ed). New Jersey: Erlbaum.
- Rosenberg, M. (1965). *Society and adolescents self image*. Princeton, N J: Princeton university.
- Schumacker, R. E., & Lomax, R. G. (2010). *Beginner's guide to structural equation modeling* (3rd ed). Mahwah, NJ: Lawrence Erlbaum.
- Shadish, W. R., & Cook, T. D. (2002). *Experimental and quasi-experimental designs for generalized causal inference*. Boston: Houghton Mifflin Company.
- Shah, R., & Goldstein, S. M. (2006). Use of structural equation modeling in operations management research: Looking back and forward. *Journal of Operation Management*, 24, 148-169.
- Sharama, S., Mukherjee, S., Kumar, A., & Dillon, W. R. (2005). A simulation study to investigate the use of cut off values for

assessing model fit in covariance structural models. *Journal of Business Research*, 58, 935- 943.

Steiger, J . H., & lind, J. M. (1980). *Statistical based tests for the number of common factors*. Paper presented the Meeting of Psychometric society, Iowa City, IA.

Steiger, J. H. (1990b). Some additional thoughts on component , and factor indeterminacy. *Multivariate Behavioral Research*, 25,41-45.

Steven, J. P. (1980). Power of the multivariate analysis of variance tests. *Psychological Bulletin*, 88, 728- 737.

Steven, J. P. (2009). *Applied multivariate statistics for the social sciences*. New York: Routledge.

Tabachnick, B. G., & Fidell, L. S. (2007). *Using multivariate statistics (4 th.ed)*. Boston: Allyn & Bacon.

Tabachnick, B. G., & Fidell, L. S. (2013). *Using multivariate statistics (6 th.ed)*. New Jersey: Pearson. Education, Inc

Tatsuoka, M. (1993). Effect Size. In G. Keren & C. lewis (Eds.), *A Handbook for Data analyses in the behavioral sciences: Methodological Issues*. New Jersey: Lawrence Erlbaum Associates, publishers

Thompson, B. (2000). Ten commandments of structural equation modeling. In L. G. Grimm & P. R. Yarnold(Eds.), *Reading and understanding more multivariate statistics(pp. 261-283)*. Washington, DC: Amrrican Psychological Association.

Thompson, B. (2004). *Exploratory and confirmatory factor analysis: Understanding concepts and applications*. Washington: American psychological Association.

Ullman, J. B. (2006). Structural equation modeling: Reviewing the basics and Moving forward. *Journal of Personality Assessment*, 87, 35- 50.

- Ullman, J. B., & Bentler, P. M. (2013). Structural equation modeling .In. I. B. Weiner (Eds.), *Handbook of Psychology* (2nd .ed.), (pp. 661 – 690) . John Wiley & Sons, Inc.
- Velicer, W. F., & Jackson, D. N. (1990). Component analysis versus common factor analysis: Some further observations. *Multivariate Behavioral Research*, 25, 97-114.
- West, S. G., Finch, J. F., & Curran, P. J. (1995). Structural equation modeling with non normal variables: problems and remedies. In R. H. Hoyle (Eds.), *Structural equation modeling: concepts, issues, and applications* (PP. 65- 75). Thousand Oaks, CA: Sage.
- Weston, R., & Gore, P. A. (2006). A brief guide to structural equation modeling. *The Counseling Psychologist*, 34, 719 - 751.
- Widman, K. F. (1990). Bias in pattern loadings represented by common factor analysis and component analysis. *Multivariate Behavioral Research*, 25, 89-95.
- Wilsson Van Voorhis, C.R. & Morgan, B. L. (2007). Understanding power and rules of thumb for determining sample sizes. *Tutorials in Quantative Methods for Psychology*, 3, 43-50.
- Zwick, W. R., & Velicer, W. F. (1986). Comparison of five rules for determining the number of components to retain. *Psychological Bulletin*, 99, 432-442.