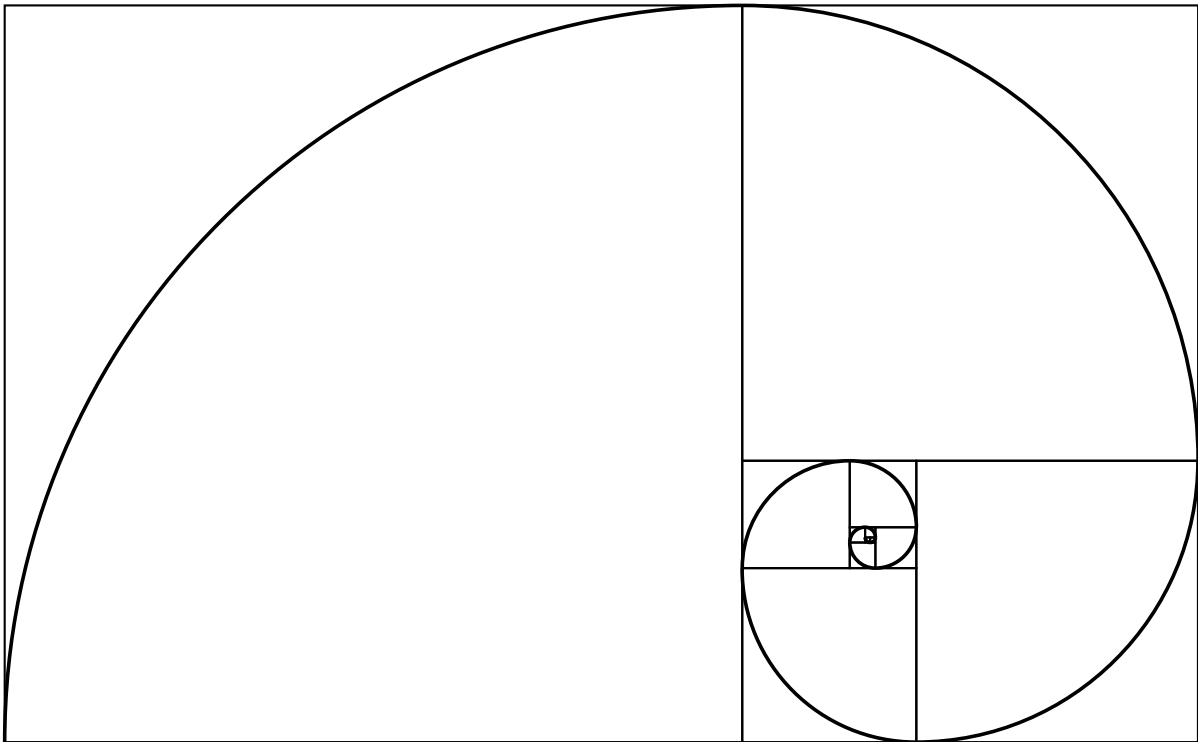


Εισαγωγή στην Στατιστική Ανάλυση με **Jamovi** και **SPSS**



Αλέξης Μπράιλας

2024

Μπράιλας Α. (2024). *Εισαγωγή στην Στατιστική Ανάλυση με Jamovi και SPSS*.
<https://doi.org/10.31219/osf.io/rbz2u>

Έκδοση 01, Απρίλιος 2024

Για ανατροφοδότηση, παροράματα και προτάσεις βελτίωσης ή διορθώσεις,
παρακαλώ στείλτε email στο: abrailas@panteion.gr

Περιεχόμενα

Πρόλογος	4
Βασική Περιγραφική Στατιστική	11
t-test για ανεξάρτητα δείγματα	12
Προετοιμασία των δεδομένων	14
t-test για επαναλαμβανόμενες μετρήσεις	15
Γραμμική συσχέτιση δύο συνεχών μεταβλητών	17
Πρακτική - <i>Dr Harpo's Statistics Lessons</i>	21
Ανάλυση Διακύμανσης (ANOVA).....	22
Cronbach's alpha	32
Ο Έλεγχος χ^2 (chi-square test).....	35
Πολλαπλή Γραμμική Παλινδρόμηση	42
Βιβλιογραφία	54

Πρόλογος

Ήταν μια ευτυχής σύμπτωση. Το συγκεκριμένο εγχειρίδιο αποτελεί έναν οδηγό διεξαγωγής βασικών στατιστικών αναλύσεων με λογισμικό SPSS και Jamovi μέσω ελεύθερα διαθέσιμων βιντεομαθημάτων. Αναπτύχθηκε στο πλαίσιο του μαθήματος «Πρακτικές Ασκήσεις Στατιστικής» του Τμήματος Ψυχολογίας του Παντείου Πανεπιστημίου. Ήταν η άνοιξη του 2020. Ήταν η πρώτη φορά που αναλάμβανα ομάδα διδασκαλίας σε αυτό το μάθημα. Ήταν η περίοδος της πρώτης καραντίνας που επιβλήθηκε λόγω της πανδημίας του COVID-19. Ήταν ο χειμώνας της απελπισίας. Ήμουν τσακισμένος από ένα ατύχημα που είχα στις αρχές του 2020 με τρία κατάγματα σπονδύλων και κάποιες άλλες ελαφρότερες κακώσεις. Τίποτα δεν είναι δεδομένο. Ήταν υπαρξιακό σοκ. Ζούσα τη χαρμολύπη ότι τουλάχιστον έχω αποφύγει τα χειρότερα, και δεν είναι κάτι που δεν διορθώνεται. Πονούσα καθιστός, πονούσα όρθιος, πονούσα ξαπλωμένος, πονούσα γενικώς. Και την ίδια στιγμή, ήταν πόνος που γνώριζα ότι θα ήταν περαστικός. Στη αρχή μέρα με τη μέρα, μετά βδομάδα με τη βδομάδα, μετά μήνα με το μήνα, γινόμεουν καλύτερα. Ήταν η άνοιξη της ελπίδας.

Υποχρεωτική κατεπείγουσα τηλεκπαίδευση (remote emergency teaching), αυτός ο όρος επινοήθηκε για να το περιγράψει. Έπρεπε να επινοηθούν τρόποι. Όλα τα μαθήματα, αλλά και οι ψυχοθεραπείες και τα ψυχοεκπαιδευτικά προγράμματα, άρχισαν να προσφέρονται σε ελάχιστο χρόνο, σχεδόν άμεσα, μέσω περιβαλλόντων σύγχρονης και ασύγχρονης εκπαίδευσης και τηλεδιάσκεψης. Η Παρασκευή ερχόταν κάθε βδομάδα και εγώ απελπιζόμουν, συχνά έκλαιγα. Ήμουν ήδη εξαντλημένος από τις πολλές ώρες εργασίας μπροστά στον υπολογιστή και χρειαζόμουν ξεκούραση. Και ήξερα ότι μόνο ξέγνοιαστο δεν θα ήταν το σαββατοκύριακο που ακολουθούσε, είχε ατέλειωτες ώρες σκληρής δουλειάς στην καρέκλα με τα κόκαλά μου να πονάνε. Προετοίμαζα το μάθημα της Δευτέρας, οργάνωνα το υλικό της ενότητας, έβρισκα το κατάλληλο dataset, οργάνωνα όλη τη διδασκαλία που θα έκανα κατά τη διάρκεια του τηλεμαθήματος. Δεν ξέρω πως το σκέφτηκα, αλλά ξεκίνησα και κατέγραφα στον υπολογιστή (recording) όλο το μάθημα που θα έκανα τις επόμενες μέρες. Με πολύ κόπο και χρόνο. Μάλλον ήταν ένας τρόπος να νιώθω και εγώ σίγουρος για το μάθημα. Και τη Δευτέρα ξαναέκανα το μάθημα ζωντανά στη τηλεδιάσκεψη, αλλά

ταυτόχρονα διέθετα στο eclass το βίντεο το μαθήματος που είχα ήδη καταγράψει το σαββατοκύριακο. Πολύ κουραστικό αλλά και ξέγνοιαστο μετά. Δεν με απασχολούσε αν κάποιοι δεν είχαν καλή σύνδεση ή αν δεν προλάβαιναν να καταλάβουν τη στατιστική ανάλυση της εβδομάδας κατά τη διάρκεια της σύγχρονης διδασκαλίας. Το βίντεο υπήρχε εκεί και μπορούσαν να το δουν ασύγχρονα όσες φορές θέλανε. Ταυτόχρονα ήταν μια παρακαταθήκη για μένα, το μάθημα στήθηκε για τις επόμενες χρονιές. Κάθε χρονιά ξαναβλέπω τα βίντεο που έφτιαξα τότε για να θυμηθώ τις αντίστοιχες ενότητες. Ευτυχής σύμπτωση; Το ρόδο και ο όμορφος ανθος, γεννιέται μέσ' το αγκάθι.

Έτσι ήρθε και η επόμενη άνοιξη του 2021, πάλι μαζί με υποχρεωτική καραντίνα. Τώρα όμως ήταν όλα στημένα για το μάθημα. Και εγώ πονούσα πολύ λιγότερο. Το μόνο σταθερό στη διάρκεια της ζωής, είναι η αποσταθεροποίηση. Η ζωή είναι αλλαγή. Και αυτή τη φορά ήταν κάποιοι φοιτητές με προβλήματα όρασης που δεν μπορούσαν να δουλέψουν με το SPSS γιατί δεν μπορούσε να υποστηριχθεί από το ειδικό λογισμικό που χρησιμοποιούσαν. Ωστόσο, το Jamonι συνεργαζόταν πολύ καλά με αυτό το εξειδικευμένο λογισμικό. Το Jamonι είναι ένα λογισμικό στατιστικής ανάλυσης ανοικτού κώδικα (open source) που σημαίνει ότι αναπτύσσεται δωρεάν από μια κοινότητα εθελοντών και θα προσφέρεται πάντα δωρεάν. Έχει επικρατήσει η έκφραση ότι *αν κάτι είναι δωρεάν, εσύ είσαι το προϊόν που πωλείται και αγοράζεται* (If something is free, you are the product). Αυτό δεν ισχύει στην περίπτωση του λογισμικού ανοικτού κώδικα, όπου η ανοικτή και ελεύθερη διάθεση υποστηρίζεται από πληθοποριστικά (crowdsourcing) συνεργατικά εγχειρήματα και οικολογίες ομότιμης παραγωγής. Ταυτόχρονα το Jamonι ανοίγει αυτόματα τα αρχεία που έχουν δημιουργηθεί με SPSS και είναι αρκετά διαισθητικό (intuitive) στη χρήση του, που σημαίνει ότι αν κάποιος ξέρει να χρησιμοποιεί το SPSS ή κάποιο άλλο λογισμικό στατιστικής ανάλυσης, πρακτικά μπορεί να χρησιμοποιήσει και το Jamonι (ή το JASP αντίστοιχα), με λίγη ως ελάχιστη προσπάθεια. Έτσι λοιπόν, ξεκίνησα και κατέγραφα τα μαθήματα και σε Jamonι, ώστε να είναι διαθέσιμα και σε αυτή την υποομάδα φοιτητών.

Άνοιξη του 2024. Ξέρουμε από τη θεωρία των πολύπλοκων ζωντανών συστημάτων ότι η διαφοροποίηση διασφαλίζει την ανθεκτικότητα (diversity reassures resilience). Αν κάποιος μπορεί να χρησιμοποιεί διαφορετικά λογισμικά για να κάνει τις στατιστικές αναλύσεις του, τότε, όταν κάποιο από αυτά τα λογισμικά δεν θα είναι πια διαθέσιμο για τον οποιοδήποτε λόγο, θα μπορέσει να προσαρμοστεί καλύτερα στις νέες συνθήκες. Η ευτυχής σύμπτωση των φοιτητών που δεν μπορούσαν να χρησιμοποιήσουν το SPSS, οδήγησε στην επέκταση του εκπαιδευτικού υλικού με βίντεο για το Jamonι. Και τώρα έχω στη διάθεση μου ένα πλουσιότερο και διαφοροποιημένο μαθησιακό υλικό για την εκπαίδευση των φοιτητών στην βασική

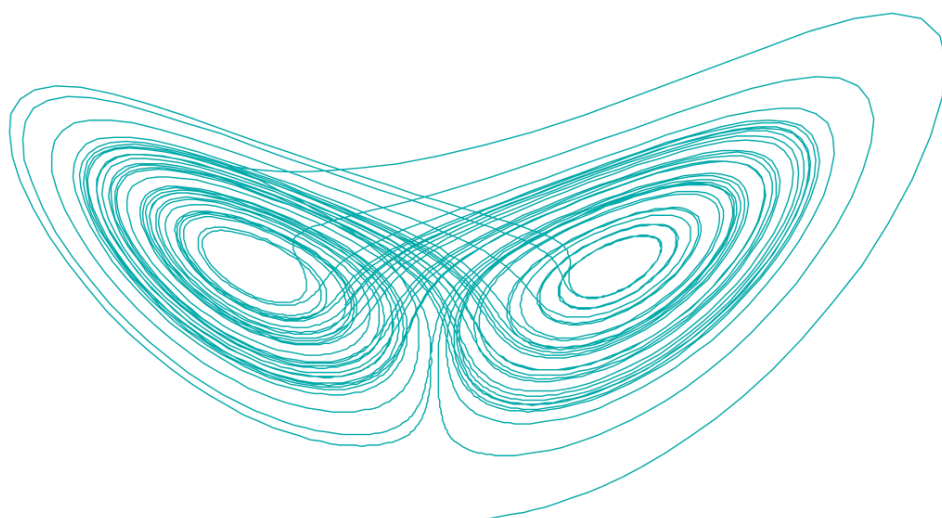
στατιστική ανάλυση. Το παρόν εγχειρίδιο περιλαμβάνει όλο αυτό το υλικό. Έχει ένα σαφή εκπαιδευτικό-παιδαγωγικό προσανατολισμό, αποσκοπώντας στην σταδιακή και εύκολη κατανόηση των σχετικών διαδικασιών από τους φοιτητές διαμέσου της μελέτης των αναλυτικών βίντεο-μαθημάτων που περιλαμβάνονται σε κάθε ενότητα.

Πολλοί παλαιότεροι φοιτητές μου ζητάνε να συνεχίσουν να έχουν πρόσβαση στο eclass του μαθήματος και στα βιντεομαθήματα των αναλύσεων. Εύχομαι αυτός ο οδηγός να φανεί χρήσιμος σε κάθε ενδιαφερόμενο. Καλές και δημιουργικές αναλύσεις!

Δρ. Αλέξης Μπράιλας
Τμήμα Ψυχολογίας
Πάντειο Πανεπιστήμιο

Ήταν τα καλύτερα χρόνια, μα συνάμα και τα χειρότερα, ήταν η εποχή της σοφίας, ήταν η εποχή της απερισκεψίας, ήταν η περίοδος της πίστης, ήταν η περίοδος της δυσπιστίας, ήταν η εποχή της Φωτιάς, ήταν η εποχή του Σκότους, ήταν η άνοιξη της ελπίδας, ήταν ο χειμώνας της απελπισίας, είχαμε τα πάντα μπροστά μας, δεν είχαμε τίποτα μπροστά μας, πηγαίναμε όλοι κατευθείαν προς τον Παράδεισο, πηγαίναμε όλοι προς την αντίθετη κατεύθυνση –κοντολογίς η περίοδος εκείνη είχε τόσες ομοιότητες με τη σημερινή, που ορισμένες από τις πιο εντυπωσιακές γραπτές μαρτυρίες επιμένουν να την περιγράφουν αποκλειστικά με επίθετα υπερθετικού βαθμού, προκειμένου να περιγράψουν είτε την καλή είτε την κακή της πλευρά¹

¹ Τσαρλς Ντίκενς, *η Ιστορία δύο πόλεων*



Εξοικείωση με το περιβάλλον εργασίας

Εξοικείωση με το περιβάλλον εργασίας του λογισμικού στατιστικής ανάλυσης SPSS

Παρακολουθήστε το [εισαγωγικό βίντεο για το περιβάλλον του SPSS](#).

Περιηγηθείτε μόνοι σας στα παράθυρα και τα μενού του SPSS ώστε να εξοικειωθείτε σταδιακά με το περιβάλλον εργασίας.

Ανοίξτε το αρχείο δεδομένων anorectic.sav από τον φάκελο Samples του SPSS ή κατεβάστε το από [εδώ](#).

Πόσα cases υπάρχουν στο αρχείο αυτό; (οριζόντιες γραμμές στο data view). Πόσες μεταβλητές υπάρχουν σε αυτό (κάθετες γραμμές στο data view).

Μεταφερθείτε στην καρτέλα variable view. Πόσες από τις μεταβλητές είναι ιεραρχικής κλίμακας μέτρησης (ordinal) και πόσες κατηγορικής (nominal)

Εξοικείωση με το περιβάλλον εργασίας του λογισμικού στατιστικής ανάλυσης JAMONI

Ανοίξτε το Jamoni και περιηγηθείτε το οριζόντιο μενού επιλογών, παρατηρώντας το περιβάλλον και τις διαθέσιμες λειτουργίες. Ο στόχος είναι να εξοικειωθείτε με την αίσθηση του περιβάλλοντος και να μην σας φαίνεται πλέον terra incognita.

Στη συνέχεια, ανοίξτε το ίδιο αρχείο δεδομένων anorectic.sav (το Jamoni ανοίγει και χειρίζεται κανονικά όλα τα αρχεία που έχουν δημιουργηθεί με το SPSS).

Πόσα cases υπάρχουν στο αρχείο αυτό; (οριζόντιες γραμμές στο data view). Πόσες μεταβλητές υπάρχουν σε αυτό (κάθετες γραμμές στο data view).

Μεταφερθείτε στην καρτέλα variable view. Πόσες από τις μεταβλητές είναι ιεραρχικής κλίμακας μέτρησης (ordinal) και πόσες κατηγορικής (nominal)

Εισαγωγή δεδομένων (data entry) στο SPSS και στο Jamovi.

Το παράθυρο Data Editor. Data view και Variable view. Δημιουργία μεταβλητών. Εισαγωγή Δεδομένων. Αποθήκευση και άνοιγμα αρχείου δεδομένων *.sav.

Παρακολουθήστε το [βίντεο για την εισαγωγή δεδομένων στο SPSS εδώ](#).

Το SPSS χρησιμοποιεί τρεις κλίμακες μέτρησης μεταβλητών (levels of measurement): Nominal, Ordinal και Scale. Nominal είναι οι ονομαστικές/κατηγορικές, πχ φύλο, θρησκεία. Ordinal είναι οι ιεραρχικές, πχ σειρά κατάταξης σε αγώνες. Scale είναι οι ποσοτικές, όπως ηλικία, θερμοκρασία, επίδοση, ύψος, ταχύτητα (ισοδιαστημικές ή αναλογικές για το SPSS είναι Scale). Αντίθετα, η κλίμακα likert (likert scale) είναι Ordinal (πχ. Διαφωνώ πολύ, διαφωνώ, ούτε διαφωνώ ούτε συμφωνώ, συμφωνώ, συμφωνώ πολύ). Δείτε και αυτό το [βίντεο στα Αγγλικά αν θέλετε περισσότερα](#).

Στη συνέχεια, δημιουργήστε το ίδιο dataset, αυτή τη φορά στο περιβάλλον του Jamovi.

Βασική Περιγραφική Στατιστική

Η ενότητα αυτή προσφέρεται μόνο για το SPSS και θα προστεθεί σε επόμενη έκδοση του οδηγού το Jamovi. Ένας χρήσιμος οδηγός για το Jamovi στα Ελληνικά από τον Πέτρο Ρούσσο βρίσκεται [ΕΔΩ](#).

Εισαγωγή δεδομένων από σχεδιασμούς ανεξαρτήτων δειγμάτων και βασική περιγραφική στατιστική

Παράδειγμα: Θερμίδες Hotdogs [Dataset από www.statstutor.ac.uk]

Ένα αμερικάνικο περιοδικό έντυπο έκανε μια έρευνα για τις θερμίδες που υπάρχουν σε διαφορετικές μάρκες hotdog. Μετρήθηκαν και καταγράφηκαν οι θερμίδες 20 μοσχαρίσιων hotdog και 17 από κρέας πουλερικών, οπότε προέκυψε το παρακάτω dataset:

Μοσχάρι

186, 181, 176, 149, 184, 190, 158, 139, 175, 148, 152, 111, 141, 153, 190, 157, 131, 149, 135, 132

Πουλερικό

129, 132, 102, 106, 94, 102, 87, 99, 170, 113, 135, 142, 86, 143, 152, 146, 144

Ερευνητικό ερώτημα: Υπάρχει διαφορά στις θερμίδες ανάμεσα στα hotdog από κρέας μοσχαριού και κοτόπουλου; (αυτό θα το ερώτημα θα απαντηθείς σε επόμενη ενότητα, ωστόσο εδώ θα κάνουμε μια αρχική διερεύνηση)

Δραστηριότητα αυτής της εβδομάδας

1. Περάστε αυτό το dataset στο SPSS χρησιμοποιώντας τις μεταβλητές meat_type (nominal) και calories (scale). Αποθηκεύστε το dataset με ένα σχετικό όνομα (και επεκταση .sav).
2. Υπολογίστε κάποια βασικά περιγραφικά στατιστικά μεγέθη (Analyze -> Descriptive statistics--> Explore). Αποθηκεύστε το output σε ένα αρχείο με ένα σχετικό όνομα (και επέκταση .spv).

Το βίντεο της ενότητας σε SPSS

Στη συνέχεια, μπορείτε για πρακτική να επιχειρήσετε να δημιουργήσετε το ίδιο dataset, αυτή τη φορά στο περιβάλλον του Jamovi, και να τις αντίστοιχες περιγραφικές αναλύσεις.

t-test για ανεξάρτητα δείγματα

Το t-test χρησιμοποιείται για να εξετάσουμε αν δύο μέσοι όροι διαφέρουν σημαντικά μεταξύ τους ή όχι.

Υπάρχουν δύο βασικές εκδοχές του t-test που θα μας απασχολήσουν, το t-test για ανεξάρτητα δείγματα και το t-test για εξαρτημένα.

Σε αυτή την ενότητα θα δούμε το t-test για ανεξάρτητα δείγματα. Θα χρησιμοποιήσουμε ένα υποθετικό dataset ([μπορείτε να το κατεβάσετε από εδώ](#)) στο οποίο έχουν καταγραφεί οι επιδόσεις στο frisbee, συγκεκριμένα η απόσταση ρίψης, για δύο ομάδες συμμετεχόντων: ιδιοκτήτες σκύλων και μη ιδιοκτήτες. Μας ενδιαφέρει να απαντήσουμε το ερώτημα αν οι ιδιοκτήτες σκύλων ρίχνουν το frisbee κατά μέσο όρο πιο μακριά από τους μη ιδιοκτήτες.

Το t-test είναι παραμετρικό, δηλαδή κάνουμε κάποιες υποθέσεις για τους πληθυσμούς από τους οποίους προέρχονται τα δείγματα.

Προϋποθέσεις του t-test

- Η ανεξάρτητη μεταβλητή είναι κατηγορική διχοτομική (με δύο επίπεδα) και κάθε συμμετέχοντας μπορεί να ανήκει μόνο σε μια από τις δύο ομάδες (πχ ή θα είναι άντρας ή θα είναι γυναίκα)
- η εξαρτημένη μεταβλητή συνεχής/ποσοτική (scale)
- Τα δείγματα να προέρχονται από πληθυσμούς που ακολουθούν την κανονική κατανομή ([ένα ενδιαφέρον άρθρο για τους ελέγχους κανονικότητας εδώ](#))
- Να υπάρχει ισοδιακύμανση, δηλαδή να έχουν την ίδια περίπου διακύμανση (equality of variances) οι δύο ομάδες (πχ η πειραματική με την ομάδα ελέγχου)

Η μη-παραμετρική ανάλυση αντίστοιχη του t-test είναι το Mann-Whitney U-Test.

Οδηγίες διεξαγωγής της ανάλυσης σε SPSS

1. Μεταφορτώστε στον υπολογιστή σας το σχετικό dataset ([μπορείτε να το κατεβάσετε από εδώ](#)).
2. κάντε τους απαραίτητους ελέγχους για την κανονικότητα όπως περιγράφονται στα σχετικά βίντεο παρακάτω.

Συγκεκριμένα για το SPSS, Analyze --> Descriptive statistics --> Explore. Επιλέγουμε στα plots Histograms και Normality plots with tests. Τα ιστογράμματα θα μας δώσουν μια πρώτη οπτική ένδειξη για το σχήμα της κατανομής, σε συνδυασμό με τα ιστογράμματα μπορούμε να δούμε και τα τεστ κανονικότητας Kolmogorov-Smirnov και Shapiro-Wilk's.

3. Διενεργούμε τον έλεγχο t-test για ανεξάρτητα δείγματα (Analyze --> Compare means --> Independent-Samples t-test)

4. Υπολογίζουμε το μέγεθος της επίδρασης (effect size) μέσω του δείκτη Cohen's d. Πλέον ο δείκτης Cohen's d υπολογίζεται από το SPSS (τώρα υπάρχει ένα σχετικό check box) άρα δεν χρειάζεται να το υπολογίζετε εσείς ξεχωριστά (όταν έφτιαξα το βίντεο της ανάλυσης, η τότε έκδοση του SPSS δεν το υπολόγιζε ακόμα, οπότε θα δείτε μια διαφοροποίηση στο σημείο αυτό). Το Jamoni επίσης υπολογίζει τον δείκτη Cohen d.

5. Ανοίγουμε ένα αρχείο κειμένου και γράφουμε την ερευνητική μας αναφορά για τον έλεγχο, συμπληρώνοντας με τα σωστά στοιχεία την παρακάτω παράγραφο.

"Οι ιδιοκτήτες σκύλων ρίχνουν μακρύτερα ένα frisbee (mean = ____ μέτρα) σε σχέση με τους μη ιδιοκτήτες (mean = ____ μέτρα). Η μέση διαφορά μεταξύ τους ήταν 11.378 και το 95% διάστημα εμπιστοσύνης για την μέση διαφορά στην εκτίμησή μας για τον πληθυσμό είναι μεταξύ ____ και _____. Το μέγεθος της επίδρασης (effect size) είναι μεγάλο ($d = ______$). Το t-test για ανεξάρτητα δείγματα έδειξε ότι η διαφορά ανάμεσα στις δύο ομάδες ήταν στατιστικά σημαντική ($t = ______, df = ______, p = ______, one-tailed$)."

Το βίντεο της ανάλυσης σε **SPSS**

Οδηγίες διεξαγωγής της ανάλυσης σε Jamoni

Δείτε το παρακάτω βίντεο στο οποίο η διαδικασία περιγράφεται σε περιβάλλον Jamoni για το ίδιο dataset.

Το βίντεο της ανάλυσης σε **JAMONI**

Προετοιμασία των δεδομένων

Η ενότητα αυτή προσφέρεται μόνο για το SPSS και θα προστεθεί σε επόμενη έκδοση του οδηγού το Jamovi

Σε αυτή την ενότητα θα δούμε τη λειτουργία δύο βασικών εντολών χειρισμού και προετοιμασίας των δεδομένων στο SPSS, των Recode και Compute. Το μάθημα βρίσκεται διαθέσιμο εδώ: <https://vimeo.com/409844809>

Για τη δραστηριότητα αυτής της ενότητας:

Κατεβάστε το παρακάτω dataset: Survey of Teachers ICT_simplified_10_countries.sav

Κάντε Recode την μεταβλητή TE24Q05 σε μια νέα μεταβλητή TE24Q05_reversed, αντιστρέφοντας τις τιμές της (1 --> 4, 2 --> 3, 3 --> 2, 4 --> 1, else --> copy).

Κάντε Recode τις μεταβλητές TE24Q05_reversed, TE24Q06, TE24Q07, TE24Q08, TE24Q09, στις μεταβλητές TE24Q05_reversed_temp, TE24Q06_temp, TE24Q07_temp, TE24Q08_temp, TE24Q09_temp, ακολουθώντας τον κανόνα (999 -> 2.5, else --> copy).

Δημιουργήστε μια νέα μεταβλητή T24Q5_Q09_mean_score με την εντολή Compute, από τον μέσο όρο των μεταβλητών TE24Q05_reversed_temp, TE24Q06_temp, TE24Q07_temp, TE24Q08_temp, TE24Q09_temp.

** Το Dataset που χρησιμοποιείται σε αυτή την ενότητα είναι μια απλοποιημένη έκδοχή του αρχικού που ανακτήθηκε από το data.europa.eu.

Το βίντεο της ενότητας σε **SPSS**

t-test για επαναλαμβανόμενες μετρήσεις

Το t-test για repeated measures (ή paired samples) χρησιμοποιείται για να εξετάσουμε αν δύο μέσοι όροι που προέρχονται από διαφορετικές μετρήσεις για τους ίδιους συμμετέχοντες διαφέρουν σημαντικά μεταξύ τους ή όχι.

Παράδειγμα

Ένας καθηγητής δημιούργησε τρία σετ θεμάτων για το ίδιο μάθημα και θέλει να δει αν είναι της ίδιας δυσκολίας. Ζητάει από τους φοιτητές να απαντήσουν στα θέματα τα οποία χορηγεί με τυχαία σειρά στον καθένα. 19 εθελοντές συμμετέχουν στη διαδικασία. Οι σωστές απαντήσεις σε κάθε τεστ και για κάθε συμμετέχοντα καταγράφονται σε ένα dataset [compare-exams.sav](https://www.spss-tutorials.com/spss-paired-samples-t-test/). Πηγή dataset: <https://www.spss-tutorials.com/spss-paired-samples-t-test/>

Το t-test είναι παραμετρικό, δηλαδή κάνουμε κάποιες υποθέσεις για τους πληθυσμούς από τους οποίους προέρχονται τα δείγματα.

Προϋποθέσεις του t-test

- Οι διαφορετικές μετρήσεις αφορούν το ίδιο δείγμα (τους ίδιους συμμετέχοντες)
- Η εξαρτημένη μεταβλητή είναι συνεχής/ποσοτική (scale). Σημείωση: οι διαφορετικές μετρήσεις της εξαρτημένης μεταβλητής στο SPSS θα αποθηκευτούν σε ξεχωριστές εσωτερικές μεταβλητές.
- Κανονική κατανομή της διαφοράς των μετρήσεων (προσεγγιστικά)
- Να μην υπάρχουν ακραίες τιμές στις διαφορές των μετρήσεων

Για τον έλεγχο της κανονικότητας και των ακραίων τιμών, χρειάζεται να χρησιμοποιήσουμε μια μεταβλητή που να αντιπροσωπεύει τις διαφορές των μετρήσεων.

Ο Έλεγχος της κανονικότητας γίνεται όπως και στο t-test για ανεξάρτητα δείγματα αλλά για μια μεταβλητή που έχει τη διαφορά των δύο συνθηκών (κάνω compute μια νέα μεταβλητή και σε αυτή ελέγχω να πληροί τις προϋποθέσεις της κανονικότητας). Η μη-παραμετρική ανάλυση αντίστοιχη του t-test για paired samples είναι το Wilcoxon test.

Οδηγίες διεξαγωγής της ανάλυσης σε SPSS

Το βίντεο της ενότητας σε SPSS βρίσκεται ΕΔΩ.

1. Analyze --> Compare means --> Paired-Samples t-test
2. Μπορούμε να διενεργήσουμε πολλούς ταυτόχρονους ελέγχους t-test. Στην αρχή ο έλεγχος αναφέρει κάποια χρήσιμα περιγραφικά στατιστικά για κάθε ζευγάρι που ελέγχεται.
3. Αναφέρεται επίσης και ένας έλεγχος συσχέτισης Pearson. Αυτόν μπορούμε να τον αγνοήσουμε αν δεν μας ενδιαφέρει. Γενικά δίνει χρήσιμες πληροφορίες, πχ αν υπάρχει ισχυρή συσχέτιση στη μεταβολή των τιμών.
4. Αναφορά αποτελεσμάτων: Υπάρχει στατιστικά σημαντική διαφορετική επίδοση στη 3ο σετ θεμάτων σε σχέση με το 1ο ($t = 2.455$, $df = 18$, $p = 0.025$, two-tailed).
5. Ερώτηση: Τι γίνεται στα άλλα δύο ζευγάρια?
6. Effect size. Τώρα πλέον η τελευταία έκδοση του SPSS υπολογίζει αυτόματα το μέγεθος της επίδρασης (δείκτης Cohen's d) οπότε δεν χρειάζεται να το υπολογίσετε εσείς. Επιδράσεις της τάξης το 0.2 θεωρούνται μικρές. Δηλαδή η διαφορά έστω και αν είναι στατιστικά σημαντική, θεωρείται πολύ μικρή. Μια μέση επίδραση είναι περίπου 0.5, ενώ μια επίδραση της τάξης του 0.8 θεωρείται μεγάλη.

Το βίντεο της ανάλυσης σε SPSS

Οδηγίες διεξαγωγής της ανάλυσης σε Jamovi

Δείτε το παρακάτω βίντεο στο οποίο η διαδικασία περιγράφεται σε περιβάλλον Jamovi για το ίδιο dataset.

Το βίντεο της ανάλυσης σε JAMOVİ

Γραμμική συσχέτιση δύο συνεχών μεταβλητών

Γραμμική συσχέτιση δύο συνεχών μεταβλητών

Σε μια έρευνα συχνά θέλουμε να ελέγξουμε αν υπάρχει βαθμός γραμμικής συσχέτισης ανάμεσα σε δύο μεταβλητές. Για παράδειγμα αν υπάρχει συσχέτιση ανάμεσα στο κάπνισμα και στην αναπνευστική λειτουργία. Σε αυτή την περίπτωση, θα μπορούσαμε να καταγράψουμε την ημερήσια κατανάλωση τσιγάρων και την απόδοση σε μια σπιρομέτρηση και να ελέγξουμε αν υπάρχει συσχέτιση σε αυτές τις μεταβλητές.

Στον έλεγχο συσχέτισης δεν έχουμε εξαρτημένες και ανεξάρτητες μεταβλητές. Απλά έχουμε δύο μεταβλητές που ελέγχουμε αν υπάρχει συσχέτιση μεταξύ τους.

“Correlation does not imply causality”. Η ύπαρξη συσχέτισης δεν συνεπάγεται και σχέση αιτίας-αιτιατού. Μπορεί να υπάρχει μια τρίτη μεταβλητή που να ερμηνεύει τη μεταβολή και στις δύο μεταβλητές. Παράδειγμα, κατανάλωση παγωτού και πνιγμοί το καλοκαίρι.

Για παραμετρικά δεδομένα χρησιμοποιούμε τον συντελεστή Pearson r . Ο δείκτης αυτός κυμαίνεται από -1 έως $+1$. Πολύ κοντά στο -1 σημαίνει ισχυρή αρνητική συσχέτιση (δηλαδή μεγάλη αύξηση της μιας μεταβλητής προκαλεί μεγάλη μείωση στην άλλη μεταβλητή). Πολύ κοντά στο $+1$ σημαίνει ισχυρή θετική συσχέτιση (αύξηση στη μία συνεπάγεται αύξηση στην άλλη). Το 0 σημαίνει ότι δεν υπάρχει καμία συσχέτιση.

Για να τρέξουμε έναν έλεγχο συσχέτισης Pearson r είναι απαραίτητο να έχουμε ένα δείγμα επαρκώς μεγάλο, τουλάχιστον 100 cases. Διαφορετικά η συσχέτιση μπορεί να αλλοιωθεί ή να κρυφτεί λόγω ακραίων τιμών στα δεδομένα μας. Σε αυτή την περίπτωση ένα διάγραμμα σκέδασης (scatterplot) μπορεί να φανεί ιδιαίτερα χρήσιμο.

Για παράδειγμα, χορηγούμε μια δοκιμασία πριν και μετά από μια παρέμβαση ή κάνουμε μια μέτρηση επίδοσης υπό διαφορετικές συνθήκες.

Οι μεταβλητές που ελέγχουμε χρειάζεται να είναι συνεχείς και να έχουμε δύο τουλάχιστον μετρήσεις που μπορεί να αντιπροσωπεύουν διαφορετικές χρονικές στιγμές ή διαφορετικές συνθήκες (για τους ίδιους πάντα συμμετέχοντες).

Παράδειγμα διερεύνησης γραμμικής συσχέτισης με τον δείκτη Pearson r

Πραγματοποιήσαμε 5 διαφορετικές ψυχολογικές δοκιμασίες σε 128 εφήβους ηλικίας 12-14. Θέλουμε να εξετάσουμε αν υπάρχει συσχέτιση ανάμεσα στις επιδόσεις στις διαφορετικές δοκιμασίες. Οι δοκιμασίες ήταν οι εξής: (α) Wechsler IQ Test Score (β) Depression Test Score (γ) Anxiety Test Score (δ) Social Functioning Test Score, και (ε) General Well Being Test Score. Τα αποτελέσματα των μετρήσεων αποθηκεύτηκαν στο dataset `adolescents.sav` (Πηγή dataset: <https://www.spss-tutorials.com/spss-correlation-analysis/>).

Οδηγίες διεξαγωγής της ανάλυσης σε SPSS

Το βίντεο της ενότητας σε SPSS βρίσκεται ΕΔΩ.

Προϋποθέσεις του ελέγχου συσχέτισης Pearson r

- Οι διαφορετικές μετρήσεις αφορούν το ίδιο δείγμα (τους ίδιους συμμετέχοντες)
- Κλίμακα μέτρησης συνεχής/ποσοτική (scale).
- Η συσχέτιση των δύο μεταβλητών πρέπει να είναι ευθύγραμμη (γραμμική).
- Οι μεταβλητές κατανέμονται κανονικά

Αρχικός έλεγχος ορθότητας των δεδομένων μας

Αρχικά ελέγχουμε τα δεδομένα μας για να εντοπίσουμε περίεργες τιμές, λάθη κατά την εισαγωγή των δεδομένων, που μπορεί να επηρεάσουν την ανάλυσή μας. Θυμηθείτε ότι οι στατιστικοί έλεγχοι είναι ευαίσθητοι στις ακραίες τιμές. Επίσης ελέγχουμε αν τα δεδομένα κατανέμονται κανονικά.

Analyze --> Descriptive statistics --> Frequencies (Frequencies ή/και Histograms)

Έλεγχος διαγραμμάτων σκεδασμού.

Στη συνέχεια ελέγχουμε τα διαγράμματα σκεδασμού για να διαπιστώσουμε αν η σχέση μεταξύ των δύο μεταβλητών είναι ευθύγραμμη, και για να έχουμε για γενικότερη εικόνα του πως συσχετίζονται τα δεδομένα μας.

Graphs --> Chart builder (αν είχα ένα μόνο ζευγάρι μεταβλητών να ελέγξω)

Graphs --> Graphboard Template Chooser (Scatterplot matrix) (για αυτή την περίπτωση που έχω 10 ζευγάρια μεταβλητών να ελέγξω ταυτόχρονα).

Διεξαγωγή του ελέγχου

Analyze --> Correlate --> Bivariate

Επιλέγουμε τις μεταβλητές για τις οποίες θέλουμε να ελέγξουμε αν υπάρχει γραμμική συσχέτιση και στους correlation coefficients επιλέγουμε Pearson.

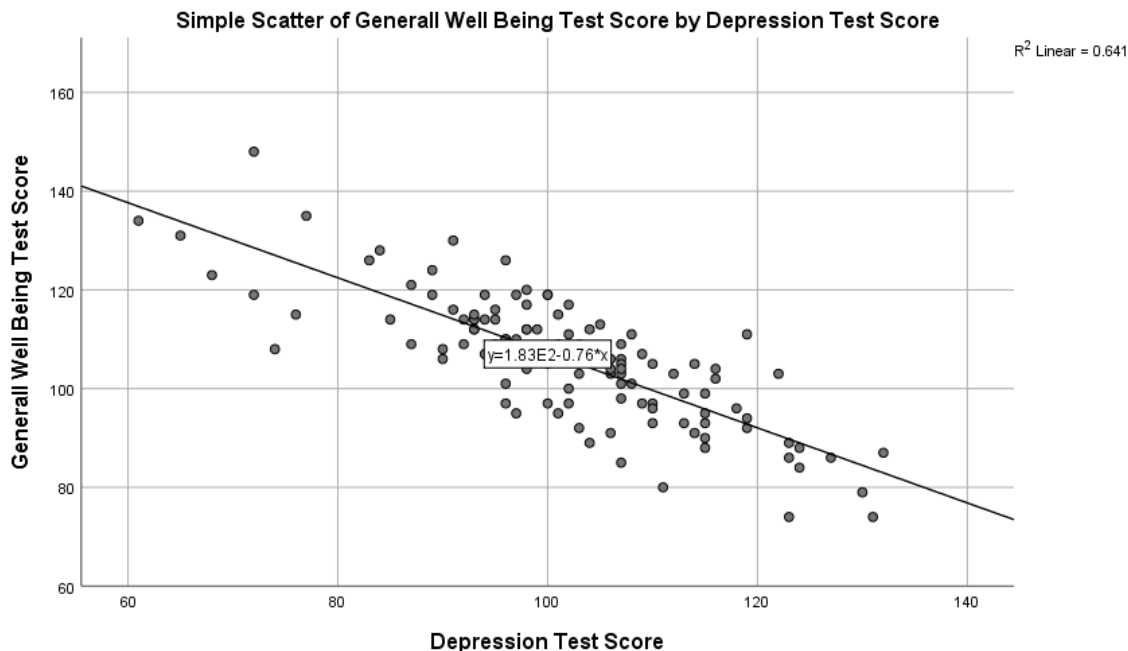
Ελέγχουμε για ποια ζευγάρια μεταβλητών υπάρχει ισχυρός συντελεστής συσχέτισης Pearson r και ιδιαίτερα για ποια ζευγάρια μεταβλητών ο συντελεστής συσχέτισης είναι στατιστικά σημαντικός (δηλαδή μπορώ να γενικεύσω στον πληθυσμό).

Ερμηνεύω τα αποτελέσματα σε σχέση με τα διαγράμματα σκεδασμού. Για να φτιάξω ένα εξειδικευμένο διάγραμμα σκεδασμού: Graphs --> Chart builder

Ο συντελεστής προσδιορισμού R^2 είναι ένα μέτρο του ποσού της μεταβλητότητας που μοιράζονται οι δυο μεταβλητές.

Αναφορά των αποτελεσμάτων

Υπάρχει στατιστικά σημαντική αρνητική συσχέτιση ανάμεσα στο Ευ Ζειν και την Κατάθλιψη ($r = -.801$, $N = 117$, $p < .0005$, two-tailed) [ή $.00025$, one-tailed αν η αρχική μου υπόθεση ήταν μονόπλευρη]. Ο συντελεστής προσδιορισμού R^2 είναι 0.641 (64.1% της μεταβολής εξηγείται από αυτή τη συσχέτιση). Το διάγραμμα σκεδασμού εμφανίζει μια ικανοποιητική κατανομή των τιμών των μεταβλητών γύρω από την γραμμή παλινδρόμησης σε μια γραμμική σχέση χωρίς ακραίες τιμές.



Έλεγχος γραμμικής συσχέτισης με τον δείκτη *Spearman rho*

Ο συντελεστής ρ του *Spearman* χρησιμοποιείται όταν η κλίμακα μέτρησης είναι ιεραρχική (Ordinal), ή όταν υπάρχουν πολλές ακραίες τιμές στο δείγμα μας, ή όταν οι τιμές μιας μεταβλητής είναι έντονα ασύμμετρες.

Το βίντεο της ανάλυσης σε **SPSS**

Οδηγίες διεξαγωγής της ανάλυσης σε **Jamovi**

Δείτε το παρακάτω βίντεο στο οποίο η διαδικασία περιγράφεται σε περιβάλλον **Jamovi** για το ίδιο dataset.

Το βίντεο της ανάλυσης σε **JAMOV**

Πρακτική - *Dr Harpo's Statistics Lessons*

Θεωρήστε το παρακάτω σενάριο:

Suppose we have 33 students taking Dr Harpo's statistics lectures, and Dr Harpo doesn't grade to a curve. Actually, Dr Harpo's grading is a bit of a mystery, so we don't really know anything about what the average grade is for the class as a whole. There are two tutors for the class, Anastasia and Bernadette. There are $N_1 = 15$ students in Anastasia's tutorials, and $N_2 = 18$ in Bernadette's tutorials. The research question I'm interested in is whether Anastasia or Bernadette is a better tutor, or if it doesn't make much of a difference. Dr Harpo emails me the course grades, in the harpo.csv file.

Το DataSet σε μορφή csv (Comma Separated Values - μια ανοικτή μορφή Excel) βρίσκεται [ΕΔΩ](#).

Πηγή σεναρίου και dataset: [Learning Statistics with JAMOV](#)

Διενεργήστε όποια στατιστική ανάλυση κρίνετε απαραίτητη ώστε να απαντήσετε το ερευνητικό ερώτημα του σεναρίου παραπάνω. Γράψτε τη σχετική ερευνητική αναφορά.

Ανάλυση Διακύμανσης (ANOVA)

Η ANOVA είναι μια πολύ χρήσιμη στατιστική ανάλυση η οποία μας επιτρέπει να ελέγξουμε για διαφορές σε ερευνητικούς σχεδιασμούς με μια ανεξάρτητη μεταβλητή που έχει περισσότερα από δύο επίπεδα ή σε ερευνητικούς σχεδιασμούς με περισσότερες από μια ανεξάρτητες μεταβλητές.

Για την περίπτωση των ανεξάρτητων μεταβλητών με περισσότερα από δύο επίπεδα (ομάδες ή συνθήκες), η ANOVA αναφέρει την ενδεχόμενη ύπαρξη στατιστικά σημαντικών διαφορών, αλλά δεν προσδιορίζει ανάμεσα σε ποια ζευγάρια (επίπεδα και συνδυασμούς των ανεξάρτητων μεταβλητών) εντοπίζονται οι διαφορές. Για τον εστιασμένο εντοπισμό των διαφορών θα κάνουμε μια στατιστική διαδικασία που λέγεται post-hoc (ή unplanned comparisons) έλεγχος Tukey ή τον αντίστοιχο έλεγχο Bonferroni.

Η ANOVA μας επιτρέπει επίσης να διερευνήσουμε τη συνδυαστική επίδραση περισσότερων από μια ανεξαρτήτων μεταβλητών. Για παράδειγμα, μπορούμε να εξετάσουμε τη συνδυαστική επίδραση ενός φαρμάκου και μιας ψυχοθεραπευτικής παρέμβασης στην κατάθλιψη. Η ANOVA μπορεί να εξετάσει τη συνδυαστική επίδραση οποιουδήποτε αριθμού ανεξαρτήτων μεταβλητών, αλλά στην πράξη δεν μπορούμε να εξετάσουμε περισσότερες από 3-4 γιατί ο έλεγχος γίνεται εξαιρετικά περίπλοκος.

Η συνδυαστική διερεύνηση της δράσης των ανεξαρτήτων μεταβλητών είναι ένα από τα βασικά πλεονεκτήματα της ANOVA. Με ποιο τρόπο ένα φάρμακο επιδρά στη θεραπεία της κατάθλιψης, συνδυαστικά με την ψυχοθεραπεία? Θα μπορούσε ενδεχόμενα μια μορφή ψυχοθεραπείας να είναι αποτελεσματική αλλά μόνο όταν συνδυάζεται με χορήγηση αντικαταθλιπτικού ή ενδεχόμενα ανάλογα με το αντικαταθλιπτικό?

Προϋποθέσεις ελέγχου

Η ANOVA είναι παραμετρικό τεστ (δηλαδή πρέπει να ικανοποιούνται συγκεκριμένες προϋποθέσεις στον πληθυσμό από τον οποίο προήλθε το δείγμα)..

- Η κλίμακα μέτρησης της εξαρτημένης μεταβλητής είναι scale (ισοδιαστημική ή αναλογική) και των ανεξάρτητων κατηγορική.
- Ο πληθυσμός από όπου προέρχεται το δείγμα κατανέμεται κανονικά

- Υπάρχει ομοιογένεια των διασπορών (τα δείγματα προέρχονται από πληθυσμούς με την ίδια περίπου διασπορά)

Για την περίπτωση σχεδιασμών με ανεξάρτητα δείγματα, χρειάζεται να έχουμε τυχαία ανεξάρτητα δείγματα από κάθε συνθήκη (τυχαία δειγματοληψία).

Οι διαφορετικές συνθήκες δεν απαιτείται να έχουν τον ίδιο αριθμό cases.

Ορολογία

Παράγοντες (Factors): είναι οι ανεξάρτητες μεταβλητές (μπορεί να έχουμε μόνο μία).

Επίπεδα (Levels). Οι συνθήκες, οι διαφορετικές τιμές της ανεξάρτητης μεταβλητής/παράγοντα.

Between-subjects σχεδιασμοί. Κάθε συμμετέχοντας εκτίθεται σε ένα μόνο επίπεδο του κάθε παράγοντα (Σχεδιασμοί ανεξάρτητων δειγμάτων)

Within-subjects σχεδιασμοί. Κάθε συμμετέχοντας εκτίθεται σε περισσότερα από ένα επίπεδα ενός παράγοντα (Σχεδιασμοί επαναλαμβανόμενων μετρήσεων)

Μεικτοί σχεδιασμοί ANOVA. Mixed between-subjects και within-subjects. Κάθε συμμετέχοντας εκτίθεται σε ένα μόνο επίπεδο για έναν παράγοντα, αλλά σε περισσότερα επίπεδα ενός άλλου παράγοντα.

One-way ANOVA. Πειραματικοί σχεδιασμοί με έναν παράγοντα (ανεξάρτητη μεταβλητή).

Two-way ANOVA. Πειραματικοί σχεδιασμοί με δύο παράγοντες.

N-way ANOVA. Πειραματικοί σχεδιασμοί με N παράγοντες.

Παράδειγμα: ANOVA 2*3 --> Πειραματικός σχεδιασμός δύο παραγόντων (two-way ANOVA), με δύο επίπεδα ο πρώτος και τρία επίπεδα ο δεύτερος. Για παράδειγμα, άντρες/γυναίκες με επίπεδα δοσολογίας φαρμάκου 0mg, 20mg, 40mg.

Κύριες και συνδυαστικές επιδράσεις

Κύρια επίδραση είναι η επίδραση κάθε παράγοντα ξεχωριστά στην εξαρτημένη μεταβλητή. Οι συνδυαστικές επιδράσεις αναφέρονται στο πως δύο ή περισσότεροι παράγοντες επιδρούν συνδυαστικά στην εξαρτημένη μεταβλητή.

Επομένως, όταν έχουμε one-way ANOVA (δηλαδή έναν παράγοντα) έχουμε μόνο μια κύρια επίδραση.

Όταν έχουμε two-way ANOVA (δηλαδή δύο παράγοντες), τότε έχουμε δύο κύριες επιδράσεις και μια συνδυαστική επίδραση των δύο παραγόντων.

Όταν έχουμε three-way ANOVA (δηλαδή τρεις παράγοντες), τότε έχουμε τρεις κύριες επιδράσεις, τρεις συνδυαστικές επιδράσεις ανά δύο παράγοντες, και μια συνδυαστική επίδραση και των τριών. Σύνολο 7 επιδράσεις.

Όταν έχουμε four-way ANOVA (δηλαδή τέσσερις παράγοντες), τότε έχουμε τέσσερις κύριες επιδράσεις, έξι συνδυαστικές επιδράσεις ανά δύο παράγοντες, τέσσερις συνδυαστικές επιδράσεις ανά τρεις, και μια συνδυαστική επίδραση και των τεσσάρων. Σύνολο 15 επιδράσεις.

Από τα παραπάνω γίνεται κατανοητό ότι πάνω από three ή four way ANOVA, όχι τόσο η διεξαγωγή του ελέγχου αλλά η ερμηνεία των αποτελεσμάτων γίνεται πρακτικά πολύ δύσκολη.

Παράδειγμα διερεύνησης

“Suppose you’ve become involved in a clinical trial in which you are testing a new antidepressant drug called Joyzepam. In order to construct a fair test of the drug’s effectiveness, the study involves three separate drugs to be administered. One is a placebo, and the other is an existing antidepressant / anti-anxiety drug called Anxifree. A collection of 18 participants with moderate to severe depression are recruited for your initial testing. Because the drugs are sometimes administered in conjunction with psychological therapy, your study includes 9 people undergoing cognitive behavioural therapy (CBT) and 9 who are not. Participants are randomly assigned (doubly blinded, of course) a treatment, such that there are 3 CBT people and 3 no-therapy people assigned to each of the 3 drugs. A psychologist assesses the mood of each person after a 3 month run with each drug, and the overall improvement in each person’s mood is assessed on a scale ranging from -5 to +5. With that as the study design, let’s now load up the data file in clinicaltrial.csv. We

can see that this data set contains the three variables drug, therapy and mood.gain.”

Πηγή: Navarro DJ and Foxcroft DR (2019). *Learning statistics with jamovi: A tutorial for psychology students and other beginners*. Retrieved from: <http://learnstatswithjamovi.com>

Έχουμε λοιπόν έναν ερευνητικό σχεδιασμό με τρία επίπεδα φαρμάκου (Placebo, Anxifree, Joyzepam), με δύο επίπεδα θεραπείας (no therapy, CBT) και μια εξαρτημένη μεταβλητή, την θετική επίδραση στη διάθεση σε καταθλιπτικούς συμμετέχοντες.

Το dataset του παραδείγματος βρίσκεται ΕΔΩ.

Ερευνητική Υπόθεση

H_0 : Ισχύει $\mu_P = \mu_A = \mu_J$ (δηλαδή δεν υπάρχουν διαφορές στην επίδραση στην κατάθλιψη)

H_1 : Δεν ισχύει $\mu_P = \mu_A = \mu_J$ (δηλαδή υπάρχουν διαφορές στην επίδραση στην κατάθλιψη)

Οδηγίες διεξαγωγής της ανάλυσης σε SPSS

Το αναλυτικό βίντεο της ενότητας σε SPSS βρίσκεται ΕΔΩ.

Έλεγχοι των προϋποθέσεων της κανονικότητας και της ομοιογένειας των διακυμάνσεων

Ελέγχουμε την κανονική κατανομή της εξαρτημένης μεταβλητής για κάθε επίπεδο των ανεξάρτητων μεταβλητών. Προς τούτο διενεργούμε τον έλεγχο Shapiro-Wilk σε συνδυασμό με τα ιστογράμματα και τα διαγράμματα QQ.

Analyze --> Descriptive Statistics --> Explore (και μετά δηλώνουμε την εξαρτημένη μεταβλητή και τους φάκτορες).

Η ANOVA είναι σχετικά ανθεκτική μέθοδος στην παραβίαση της κανονικότητας και για αυτό απαιτείται τα δεδομένα να είναι μόνο σχετικά (δηλαδή όχι απολύτως

ιδανικά) κανονικά κατανομημένα. Για όχι πολύ μικρά δείγματα ($N \geq 25$) μπορεί και να μην χρειάζεται ο έλεγχος της κανονικότητας. Για τον έλεγχο της ομοιογένειας των διασπορών κάνουμε τον έλεγχο Levene κατά τη διενέργεια της ANOVA επιλέγοντας το σχετικό check-box. Ο έλεγχος της ισότητας των διασπορών απαιτείται κυρίως σε υποσυνθήκες διαφορετικών μεγεθών. Αν τα επίπεδα των ανεξαρτήτων μεταβλητών είναι ίσου μεγέθους, τότε δεν απαιτείται.

Η ANOVA εν γένει είναι ανθεκτική μέθοδος και ως προς το κριτήριο κανονικότητας και της ισοδιακύμανσης όταν τα μεγέθη των ομάδων είναι ίσα. Ωστόσο, όταν τα μεγέθη δεν είναι ίσα, τότε θα χρειαστεί να είμαστε ιδιαίτερα προσεκτικοί στον έλεγχο των προϋποθέσεων.

Μη παραμετρικό τεστ που θα μπορούσαμε να χρησιμοποιήσουμε αντί της ANOVA είναι το Kruskal-Wallis.

Διεξαγωγή της ANOVA

One-way between-subjects ANOVA

Για παιδαγωγικούς και μόνο λόγους, προκειμένου να δούμε πώς διενεργείται η One-way ANOVA, αρχικά θα αγνοήσουμε τον δεύτερο παράγοντα στο dataset, δηλαδή την ψυχοθεραπευτική παρέμβαση (no therapy/CBT), θα υποθέσουμε ότι δεν υπάρχει αυτή η μεταβλητή.

Α' τρόπος: Analyze --> Compare Means --> One-way ANOVA, και στα OPTIONS --> Descriptive statistics

mood.gain

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	3.453	2	1.727	18.611	.000
Within Groups	1.392	15	.093		
Total	4.845	17			

Αναφορά των αποτελεσμάτων Α' τρόπου

Υπήρχε μια στατιστικά σημαντική επίδραση του φαρμάκου στην ψυχολογική διάθεση ($F(2,15) = 18.611, p < .0005$).

B' τρόπος: Analyze --> General Linear Model --> Univariate, και στα OPTIONS --> Descriptive statistics και Estimates of effect size

Tests of Between-Subjects Effects

Dependent Variable: mood.gain

Source	Type III Sum of Squares	df	Mean Square	F	Sig.	Partial Eta Squared
Corrected Model	3.453 ^a	2	1.727	18.611	.000	.713
Intercept	14.045	1	14.045	151.383	.000	.910
drug	3.453	2	1.727	18.611	.000	.713
Error	1.392	15	.093			
Total	18.890	18				
Corrected Total	4.845	17				

a. R Squared = .713 (Adjusted R Squared = .674)

Αναφορά των αποτελεσμάτων B' τρόπου

Υπήρχε μια στατιστικά σημαντική επίδραση του φαρμάκου στην ψυχολογική διάθεση ($F(2,15) = 18.611$, $p < .0005$, μερικό $\eta^2 = .71$).

Graphs --> Chart Builder --> Bars --> Simple Error Bar

Unplanned ή post-hoc συγκρίσεις (Bonferroni)

Στο πλαίσιο διαλόγου για την ANOVA (compare means --> one-way ANOVA ή General Linear Model --> Univariate) επιλέγουμε Post Hoc και μετά Bonferroni (ή Tukey).

Multiple Comparisons

Dependent Variable: mood.gain

Bonferroni

(I) drug	(J) drug	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
Placebo	Anxifree	-.267	.1759	.451	-.740	.207
	Joyzepam	-1.033*	.1759	.000	-1.507	-.560
Anxifree	Placebo	.267	.1759	.451	-.207	.740
	Joyzepam	-.767*	.1759	.002	-1.240	-.293
Joyzepam	Placebo	1.033*	.1759	.000	.560	1.507
	Anxifree	.767*	.1759	.002	.293	1.240

Based on observed means.

The error term is Mean Square(Error) = .093.

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

Στην αναφορά μας προσθέτουμε:

Με τον post-hoc έλεγχο Bonferroni ανιχνεύθηκαν σημαντικές διαφορές ανάμεσα στην ομάδα που έλαβε Joyzepam και την ομάδα έλαβε Placebo ($p < .0005$), όπως και ανάμεσα στην ομάδα που έλαβε Joyzepam και την ομάδα που έλαβε Anxifree ($p = .002$). Δεν βρέθηκε σημαντική διαφορά ανάμεσα στην ομάδα που έλαβε Placebo και την ομάδα που έλαβε Anxifree ($p = .451$).

Two-way between-subjects ANOVA (3*2)

Τώρα που είδαμε πως διενεργούμε one way ANOVA, λαμβάνουμε κανονικά υπόψιν και τους δύο παράγοντες, οπότε θα χρειαστεί να κάνουμε κανονικά two-way ANOVA.

Analyze --> General Linear Model --> Univariate, και στα OPTIONS --> Descriptive statistics και Estimates of effect size

Tests of Between-Subjects Effects

Dependent Variable: mood.gain

Source	Type III Sum of Squares	df	Mean Square	F	Sig.	Partial Eta Squared
Corrected Model	4.192 ^a	5	.838	15.398	.000	.865
Intercept	14.045	1	14.045	257.969	.000	.956
drug	3.453	2	1.727	31.714	.000	.841
therapy	.467	1	.467	8.582	.013	.417
drug * therapy	.271	2	.136	2.490	.125	.293
Error	.653	12	.054			
Total	18.890	18				
Corrected Total	4.845	17				

a. R Squared = .865 (Adjusted R Squared = .809)

Αναφορά των αποτελεσμάτων

Υπάρχει μια στατιστικά σημαντική κύρια επίδραση του φαρμάκου στην ψυχολογική διάθεση ($F(2,12) = 31.714$, $p < .0005$, μερικό $\eta^2 = .841$).

Υπάρχει μια στατιστικά σημαντική κύρια επίδραση της ψυχολογικής θεραπείας στην ψυχολογική διάθεση ($F(1,12) = 8.582$, $p = .013$, μερικό $\eta^2 = .417$).

Δεν υπάρχει στατιστικά σημαντική συνδυαστική επίδραση του παράγοντα του φαρμάκου και του παράγοντα της ψυχολογικής θεραπείας στην ψυχολογική διάθεση ($F(2,12) = 2.490$, $p = .125$, μερικό $\eta^2 = .293$).

Graphs --> Chart Builder --> Bars --> Simple Error Bar

Unplanned ή post-hoc συγκρίσεις (Bonferroni)

Στο πλαίσιο διαλόγου για την ANOVA (compare means --> one-way ANOVA ή General Linear Model --> Univariate) επιλέγουμε Post Hoc και μετά Bonferroni.

Warnings

Post hoc tests are not performed for therapy because there are fewer than three groups.

Multiple Comparisons

Dependent Variable: mood.gain

Bonferroni

(I) drug	(J) drug	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
Placebo	Anxifree	-.267	.1347	.214	-.641	.108
	Joyzepam	-1.033*	.1347	.000	-1.408	-.659
Anxifree	Placebo	.267	.1347	.214	-.108	.641
	Joyzepam	-.767*	.1347	.000	-1.141	-.392
Joyzepam	Placebo	1.033*	.1347	.000	.659	1.408
	Anxifree	.767*	.1347	.000	.392	1.141

Based on observed means.

The error term is Mean Square(Error) = .054.

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

Στην αναφορά μας προσθέτουμε:

Με τον post-hoc έλεγχο Bonferroni ανιχνεύθηκαν σημαντικές διαφορές ανάμεσα στην ομάδα που έλαβε Joyzepam και την ομάδα έλαβε Placebo ($p < .0005$), όπως και ανάμεσα στην ομάδα που έλαβε Joyzepam και την ομάδα που έλαβε Anxifree ($p < .0005$). Δεν βρέθηκε σημαντική διαφορά ανάμεσα στην ομάδα που έλαβε Placebo και την ομάδα που έλαβε Anxifree ($p = .214$).

Για την πρακτική σας άσκηση στην ANOVA

1. Κατεβάστε τον αντίστοιχο dataset από [εδώ](#).
2. Κάντε τους ελέγχους που περιγράφονται στο σχετικό Βίντεο (Προσοχή μόνο για την two-way ANOVA)

Το βίντεο της ανάλυσης σε **SPSS**

Οδηγίες διεξαγωγής της ανάλυσης σε Jamovi

Δείτε το παρακάτω βίντεο στο οποίο η διαδικασία περιγράφεται σε περιβάλλον Jamovi για το ίδιο dataset.

Το βίντεο της ανάλυσης σε **JAMOV**

Cronbach's alpha

Ο δείκτης Cronbach alpha είναι ένας συντελεστής αξιοπιστίας μιας κλίμακας. Για την ακρίβεια είναι ένα μέτρο της εσωτερικής συνεκτικότητας/συνέπειας (consistency) μιας ομάδας μετρήσεων, συνήθως σε κλίμακα Likert. Αν για παράδειγμα ένα σύνολο ερωτήσεων εξετάζει ένα συγκεκριμένο χαρακτηριστικό της προσωπικότητας, συμπεριφορά ή εμπειρία, τότε ένας υψηλός δείκτης Cronbach alpha (>0.7) αποτελεί μια καλή ένδειξη για την συνεκτικότητα αυτής της (υπό)κλίμακας. Από την άλλη, μια υψηλή τιμή του δείκτη Cronbach alpha δεν διασφαλίζει ότι μια κλίμακα μετράει μόνο ένα πράγμα, δηλαδή μόνο μια εννοιολογική κατασκευή (unidimensional), έναν μόνο παράγοντα. Μια κλίμακα θα μπορούσε να μετράει περισσότερους από έναν σχετιζόμενους παράγοντες και για αυτό να εμφανίζει υψηλό alpha.

Παράδειγμα διερεύνησης

“Twenty-five personality self-report items (see Figure 15.2) taken from the International Personality Item Pool (<http://ipip.ori.org>) were included as part of the Synthetic Aperture Personality Assessment (SAPA) web-based personality assessment (SAPA: <http://sapa-project.org>) project. The 25 items are organized by five putative factors: Agreeableness, Conscientiousness, Extraversion, Neuroticism, and Openness. The item data were collected using a 6-point response scale:

1. Very Inaccurate
2. Moderately Inaccurate
3. Slightly Inaccurate
4. Slightly Accurate
5. Moderately Accurate
6. Very Accurate.

A sample of $N=250$ responses is contained in the dataset bfi sample.csv.”

Variable name	Question / Item (short phrases that you should respond to by indicating how accurately the statement describes your typical behaviour or attitudes)	Coding (R: reverse)
A1	<u>Am</u> indifferent to the feelings of others.	R
A2	Inquire about others' well-being.	
A3	Know how to comfort others.	
A4	Love children.	
A5	Make people feel at ease.	
C1	<u>Am</u> exacting in my work.	
C2	Continue until everything is perfect.	
C3	Do things according to a plan.	
C4	Do things in a half-way manner.	R
C5	Waste my time.	R
E1	Don't talk a lot.	R
E2	Find it difficult to approach others.	R
E3	Know how to captivate people.	
E4	Make friends easily.	
E5	Take charge.	
N1	Get angry easily.	
N2	Get irritated easily.	
N3	Have frequent mood swings.	
N4	Often feel blue.	
N5	Panic easily.	
O1	Am full of ideas.	
O2	Avoid difficult reading material.	R
O3	Carry the conversation to a higher level.	
O4	Spend time reflecting on things.	
O5	Will not probe deeply into a subject.	R

Πηγή: Navarro DJ and Foxcroft DR (2019). *Learning statistics with jamovi: A tutorial for psychology students and other beginners*. Retrieved from: <http://learnstatswithjamovi.com>

Επιθυμούμε να εξετάσουμε την εσωτερική συνεκτικότητα των 5 υποκλιμάκων A, C, E, N, O (Agreeableness, Conscientiousness, Extraversion, Neuroticism, and Openness).

Το dataset για την ανάλυση αυτή βρίσκεται εδώ.

Οδηγίες διεξαγωγής της ανάλυσης σε SPSS

Το βίντεο της ενότητας αυτής σε SPSS βρίσκεται ΕΔΩ.

Analyze --> Scale --> Reliability Analysis

Στη συνέχεια στο Model, επιλέγουμε Alpha για να υπολογίσουμε μόνο τον δείκτη Cronbach Alpha. Επιλέγουμε τις μεταβλητές της υποκλίμακας που θέλουμε να ελέγξουμε. Στα Statistics, επιλέγουμε για περισσότερες χρήσιμες πληροφορίες ανά μεταβλητή της κλίμακας: item, scale, scale if item deleted, Correlations.

Στην αναφορά των αποτελεσμάτων μας γράφουμε:

Ο δείκτης Cronbach alpha για την κλίμακα χχχχ στο συγκεκριμένο δείγμα είναι .xx. Τα περισσότερα στοιχεία της κλίμακας χρειάζεται να παραμείνουν, με ενδεχόμενη απομάκρυνσή τους να μειώνει τον δείκτη Cronbach alpha. Εξαίρεση αποτελεί η ερώτηση χχχ η οποία αν αφαιρεθεί αυξάνει την εσωτερική συνοχή της κλίμακας (Cronbach α = ...). Για αυτό η συμπερίληψή της στο ερωτηματολόγιο χρειάζεται να επανεξεταστεί.

Το βίντεο της ανάλυσης σε **SPSS**

Οδηγίες διεξαγωγής της ανάλυσης σε Jamovi

Δείτε το παρακάτω βίντεο στο οποίο η διαδικασία περιγράφεται σε περιβάλλον Jamovi για το ίδιο dataset.

Το βίντεο της ανάλυσης σε **JAMOV**

Ο Έλεγχος χ^2 (chi-square test)

Ο έλεγχος χ^2 χρησιμοποιείται για να ελέγξουμε την ύπαρξη σχέσης ανάμεσα σε δύο κατηγορικές (nominal or ordinal) μεταβλητές. Για συσχέτιση ανάμεσα σε συνεχείς/ποσοτικές (scale) μεταβλητές χρησιμοποιούμε τον Pearson r. Συνήθης εφαρμογή είναι σε ερωτηματολόγια για τον έλεγχο συνάφειας ανάμεσα σε κατηγορικές μεταβλητές.

Ο έλεγχος χ^2 αποτελεί μόνο έναν έλεγχο συνάφειας, όχι συσχέτισης (association not correlation). Ένα θετικό αποτέλεσμα υποδηλώνει ότι υπάρχει ένα είδος σχέσεις ανάμεσα σε δύο κατηγορικές μεταβλητές, αλλά δεν προσδιορίζει το είδος αυτής της σχέσης.

Τα κατηγορικά δεδομένα αποθηκεύονται εσωτερικά στο SPSS με έναν αριθμό, αλλά αυτός ο αριθμός είναι αυθαίρετος και χρησιμοποιείται απλά για να αντιπροσωπεύει την κατηγορία. Για παράδειγμα 0 = άντρας και 1 = γυναίκα. Ή 1 = γυναίκα και 2 = άντρας.

Στους ερευνητικούς σχεδιασμούς με ανεξάρτητα δείγματα, η ανεξάρτητη μεταβλητή (πχ dog owner, not dog owner, ή ομάδα ελέγχου, πειραματική ομάδα) είναι κατηγορική μεταβλητή (nominal). Ωστόσο, χρειάζεται και η εξαρτημένη μεταβλητή να είναι κατηγορική για να κάνουμε τον έλεγχο χ^2 .

Ο έλεγχος χ^2 χρησιμοποιείται για ανεξάρτητα δείγματα. Για επαναλαμβανόμενες μετρήσεις χρησιμοποιούμε τον έλεγχο McNemar.

Επισημάναμε ότι οι μεταβλητές θα πρέπει να είναι κατηγορικές, ονομαστικές ή ιεραρχικές (nominal ή ordinal). Ωστόσο, μπορούμε να μετατρέψουμε και συνεχείς μεταβλητές σε κατηγορικές, πχ. Χωρίζοντας τις επιδόσεις στο IQ σε κατηγορίες (low, low middle, middle, middle high, high).

Παράδειγμα διερεύνησης συσχέτισης με το χ^2

(Πηγή dataset: Mollie Gilchrist and Peter Samuels, www.statstutor.ac.uk)

Θα ελέγξουμε αν υπάρχει σχέση ανάμεσα στον τύπο της προσωπικότητας (εσωστρεφής ή εξωστρεφής) και στην προτίμηση σε χρώμα (κόκκινο, πράσινο, κίτρινο, μπλε).

Τα αποτελέσματα από μια έρευνα σε 400 φοιτητές φαίνονται στο παρακάτω 2x4 πίνακα διασταύρωσης ή πίνακα συνάφειας (cross-tabulation or contingency table):

	Μπλε	Πράσινο	Κόκκινο	Κίτρινο	Σύνολο
Εσωστρεφής	44	30	20	6	100
Εξωστρεφής	36	50	180	34	300
Σύνολα	80	80	200	40	400

Σημείωση: Για τις κατηγορικές μεταβλητές δεν έχουν νόημα τα μέτρα κεντρικής τάσης (μέσος όρος, διάμεσος) ή τα μέτρα διασποράς (διακύμανση, εύρος, τυπική απόκλιση). Τα μόνα περιγραφικά μεγέθη που έχουν νόημα είναι το πλήθος, η συχνότητα και το ποσοστό.

Οι υποθέσεις μας

H_0 : Η προτίμηση χρώματος δεν σχετίζεται με τον τύπο προσωπικότητας

H_1 : Η προτίμηση χρώματος σχετίζεται με τον τύπο προσωπικότητας

Πώς λειτουργεί το chi-square? Αν δεν υπάρχει διαφορά στην προτίμηση χρώματος ανάμεσα σε εσωστρεφείς και εξωστρεφείς, τότε οι μετρήσεις των εσωστρεφών αναμένονται να είναι ίδιες (σχεδόν) με τις μετρήσεις των εξωστρεφών. Πρακτικά, ο έλεγχος χ^2 υπολογίζει τις αναμενόμενες τιμές (expected values) στην μια κατηγορία, αν δεν υπήρχαν διαφορές με την άλλη κατηγορία, και τις συγκρίνει με τις παρατηρούμενες τιμές. Αν βρει ότι αυτές διαφέρουν μεταξύ τους σημαντικά, αποφαινεται ότι υπάρχουν στατιστικά σημαντικές διαφορές.

Κατεβάστε το dataset για τον έλεγχο χ^2 από εδώ

Οδηγίες διεξαγωγής της ανάλυσης σε SPSS

Το βίντεο της ενότητας σε SPSS βρίσκεται ΕΔΩ.

Analyze --> Descriptive Statistics --> Crosstabs

- Επιλέγουμε την μία μεταβλητή ως row variable και την άλλη ως column variable

- Στο Statistics επιλέγουμε Chi-square
- Στο cells επιλέγουμε Observed, Expected counts και τα
- Επιλέγουμε Display chartered bar charts ή το κάνουμε μετά από το κεντρικό μενού Graphs που δίνει πιο πολλές επιλογές

Αποτελέσματα

Τον παρακάτω πίνακα των παίρνουμε αν επιλέξουμε στο Cells μόνο Observed counts. Αν επιλέξουμε και Expected counts ή και percentages, τότε προστίθενται τα αντίστοιχα στοιχεία.

Personality * ColourPreference Crosstabulation

Count

		ColourPreference				Total
		Red	Yellow	Green	Blue	
Personality	Introvert	20	6	30	44	100
	Extrovert	180	34	50	36	300
Total		200	40	80	80	400

Ο παρακάτω πίνακας είναι τα αποτελέσματα του ελέγχου Chi-Square:

Chi-Square Tests

	Value	df	Asymptotic Significance (2- sided)
Pearson Chi-Square	71.200 ^a	3	.000
Likelihood Ratio	70.066	3	.000
Linear-by-Linear Association	69.124	1	.000
N of Valid Cases	400		

a. 0 cells (0.0%) have expected count less than 5. The minimum expected count is 10.00.

Εγκυρότητα του ελέγχου

Παρατηρήστε την υποσημείωση: a. 0 cells (0.0%) have expected count less than 5. The minimum expected count is 10.00.

Αυτή ή υποσημείωση μας υποδηλώνει ότι ο έλεγχος είναι έγκυρος. Και αυτό γιατί απαιτείται κάποιος ελάχιστος αριθμός περιπτώσεων σε κάθε συνθήκη (κελί) του πίνακα διπλής εισόδου. Συγκεκριμένα, χρειάζεται λιγότερο από το 20% να έχουν expected count λιγότερο από 5, και σε όλες τις περιπτώσεις το expected count να είναι τουλάχιστον 1. Στην περίπτωση που ο έλεγχος chi-square αναφέρει ότι υπάρχουν ένα ή περισσότερα κελιά με expected count λιγότερο από 5, τότε δεν μπορείτε να προχωρήσετε και να αναφέρετε τα αποτελέσματα, χρειάζεται να κάνετε κάποιες παραπάνω ενέργειες, που δεν θα τις περιγράψουμε εδώ. Και σε ονομαστικές μεταβλητές μπορούμε να κάνουμε ομαδοποιήσεις, αρκεί οι ομαδοποιήσεις αυτές να έχουν νόημα.

Ωστόσο, αν υπάρχει πρόβλημα με την υποσημείωση αυτή, δηλαδή υπάρχουν κελιά με πολύ λίγες μετρήσεις, τότε μπορούμε να προσπαθήσουμε να παρακάμψουμε εμμέσως αυτό το πρόβλημα. Για παράδειγμα αν η μία μεταβλητή είναι ιεραρχική με 5 επίπεδα που αντιπροσωπεύουν απαντήσεις σε μια κλίμακα Likert (1= Συμφωνώ απόλυτα, 2 = Συμφωνώ, 3 = Ουδέτερα, 4 = Διαφωνώ, 5 = Διαφωνώ απόλυτα), τότε θα μπορούσαμε να μετατρέψουμε την μεταβλητή (Transform --> Recode) με ενοποιήσεις τιμών, ώστε να έχει μόνο 3 επίπεδα (1=Συμφωνώ, 2=Ουδέτερα, 3=Διαφωνώ) (και επομένως περισσότερες μετρήσεις σε κάθε επίπεδο) και να ξανακάνουμε τον έλεγχο.

Άλλος τρόπος να αντιπαρέλθουμε το πρόβλημα αυτό είναι να φροντίζουμε να έχουμε εξ αρχής αρκετά μεγάλο δείγμα ώστε να μην υπάρχουν κελιά με μικρό αριθμό παρατηρήσεων.

Αναφορά των αποτελεσμάτων

Το SPSS αναφέρει πολλά διαφορετικά μέτρα της πιθανότητας να υπάρχει στατιστικά σημαντική διαφορά ανάμεσα στις μεταβλητές. Από αυτά εστιάζομαι στην πρώτη γραμμή, στο Pearson Chi-Square (Ο Karl Pearson ήταν και αυτός που πρώτη φορά δημιούργησε των έλεγχο χ^2). Παρατηρούμε ότι οι βαθμοί ελευθερίας

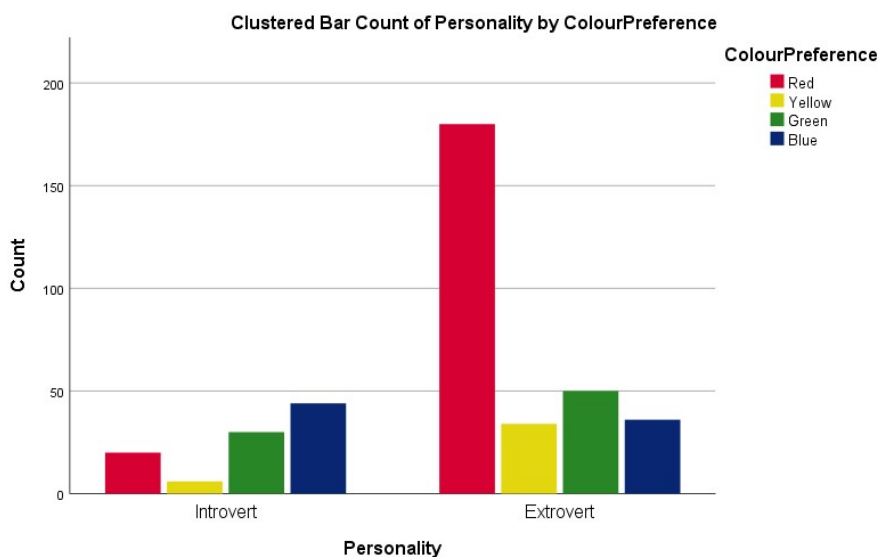
df στον έλεγχο χ^2 δεν σχετίζονται με το μέγεθος N του δείγματος (είναι τα επίπεδα στην μια μεταβλητή μείον ένα, επί τα επίπεδα στην άλλη μεταβλητή μείον ένα).

Στην αναφορά μας γράφουμε:

Υπάρχει συσχέτιση ανάμεσα στην προσωπικότητα και στην προτίμηση χρώματος, $\chi^2 (3, N=400) = 71.2, p < .0005$.

Είναι σημαντικό να επισημάνουμε ότι ο έλεγχος chi-square μας δίνει μόνο την πληροφορία ότι υπάρχει μια σχέση ανάμεσα στις 2 κατηγορικές μεταβλητές, αλλά χωρίς περισσότερες πληροφορίες για το είδος της σχέσης, το μοτίβο που συνδέει τις δύο μεταβλητές. Για αυτό χρειάζεται να κοιτάξουμε τον πίνακα διασταύρωσης σε συνδυασμό με το ραβδωτό διάγραμμα, για να γίνει περισσότερο σαφές ποιο είναι το μοτίβο της σχέσης. Για να φτιάξουμε το bar-chart, επιλέγουμε:

Graphs --> Chart Builder --> Bar --> Clustered Bar



Με το bar-chart φαίνεται να υπάρχει μια προτίμηση των εξωστρεφών προς το κόκκινο και μια σχετική προτίμηση των εσωστρεφών προς το μπλε. Από τον πίνακα διασταύρωσης μπορούμε να πάρουμε περισσότερες πληροφορίες για τα ποσοστά μας και να αναφέρουμε στην αναφορά μας:

Ανάμεσα στους εξωστρεφείς φαίνεται να υπάρχει μια πλειοψηφία (60%) που προτιμάει το κόκκινο χρώμα, τη στιγμή που στους εσωστρεφείς η προτίμηση στο κόκκινο αποτελεί μια σχετική μειοψηφία (20%). Αντίθετα στους εσωστρεφείς

φαίνεται να υπάρχει μια σχετική πλειοψηφική προτίμηση για το μπλε (44%) έναντι των εξωστρεφών όπου η προτίμηση στο μπλε είναι πολύ μικρότερη (12%).

Στη συνέχεια παραθέτουμε και το ραβδωτό διάγραμμα.

Παρατήρηση: Να επισημαίνουμε ότι τις κατηγορικές μεταβλητές δεν έχουν νόημα τα μέτρα κεντρικής τάσης (μέσος όρος, διάμεσος) ή τα μέτρα διασποράς (διακύμανση, εύρος, τυπική απόκλιση). Τα μόνα περιγραφικά μεγέθη που έχουν νόημα είναι το πλήθος, η συχνότητα και το ποσοστό. Για την γραφική τους αναπαράσταση μπορούμε να χρησιμοποιήσουμε bar-charts αλλά όχι ιστογράμματα. Ιστογράμματα μπορούμε να χρησιμοποιήσουμε μόνο για ordinal και scale μεταβλητές.

Υποσημείωση: μπορώ να πάρω ένα μέτρο του μεγέθους της συνάφειας με την επιλογή Phi and Crammer's V από το Statistics στο παράθυρο που εμφανίζεται με το Crosstabs.

Symmetric Measures

		Value	Approximate Significance
Nominal by Nominal	Phi	.422	.000
	Cramer's V	.422	.000
N of Valid Cases		400	

Ο δείκτης Φ υποδηλώνει το μέγεθος της συνάφειας και ερμηνεύεται όπως ο δείκτης Pearson r . Επομένως εδώ έχουμε ένα σχετικά μικρό μέγεθος συνάφειας, $\Phi = 0.422$. Επίσης μπορούμε να τον τετραγωνίσουμε για να πάρουμε το αντίστοιχο της συνδιακύμανσης (R^2). Σε αυτή την περίπτωση το $\Phi^2 = 0.18$ και μπορούμε να ισχυριστούμε ότι το 18% της προτίμησης στο χρώμα συνολικά μπορεί να αποδοθεί στον τύπο προσωπικότητας. Ωστόσο, από τον πίνακα συνάφειας και από το διάγραμμα μπορούμε να πάρουμε περισσότερες πληροφορίες για συγκεκριμένα χρώματα, όπως είδαμε παραπάνω.

Κατεβάστε το dataset για τον έλεγχο χ^2 από εδώ

Το βίντεο της ανάλυσης σε **SPSS**

Οδηγίες διεξαγωγής της ανάλυσης σε Jamovi

Δείτε το παρακάτω βίντεο στο οποίο η διαδικασία περιγράφεται σε περιβάλλον Jamovi για το ίδιο dataset.

Το βίντεο της ανάλυσης σε **JAMOVİ**

Πολλαπλή Γραμμική Παλινδρόμηση

Η πολλαπλή παλινδρόμηση είναι μια στατιστική τεχνική που μας επιτρέπει να διερευνήσουμε την σχέση ανάμεσα σε περισσότερες από δύο μεταβλητές. Στην ανάλυση παλινδρόμησης έχουμε μια εξαρτημένη μεταβλητή (criterion variable) και μια (απλή παλινδρόμηση) ή περισσότερες (πολλαπλή παλινδρόμηση) ανεξάρτητες μεταβλητές (ή αλλιώς προβλεπτικές μεταβλητές – predictor variables). Με την ανάλυση παλινδρόμησης δημιουργούμε ένα μοντέλο πρόβλεψης: Ποια αναμένουμε να είναι η τιμή της εξαρτημένης μεταβλητής ανάλογα με τις τιμές των προβλεπτικών μεταβλητών (ανεξάρτητες μεταβλητές).

Η ανθρώπινη συμπεριφορά είναι εγγενώς απρόβλεπτη και για αυτό όταν μιλάμε για μοντέλο πρόβλεψης εννοούμε απλά μια καλή εκτίμηση για την εξαρτημένη μεταβλητή. Όπως και με τη διμεταβλητή συσχέτιση (Pearson r) η συσχέτιση δεν σημαίνει αιτιότητα, εκτός και αν έχουμε εφαρμόσει έναν πειραματικό σχεδιασμό.

Παράδειγμα διερεύνησης

Ένας καθηγητής ψυχολογίας ενδιαφέρεται να διερευνήσει τη σχέση ανάμεσα στην επανάληψη της ύλης που κάνει ένας φοιτητής (ο αριθμός των ωρών που αφιέρωσε), το ενδιαφέρον του για το αντικείμενο του μαθήματος (μέτρηση σε κλίμακα αυτοαναφοράς 0 έως 100) και την επίδοσή του στις εξετάσεις. Στο dataset έχουμε τις μετρήσεις από 29 φοιτητές.

Πηγή παραδείγματος και dataset:

<http://www.open.ac.uk/socialsciences/spsstutorial/advanced/multiple-regression/revision>

Πριν ξεκινήσουμε την πολλαπλή παλινδρόμηση χρειάζεται να ελέγξουμε ότι ισχύουν κάποιες προϋποθέσεις. Για διδακτικούς λόγους (επειδή θα μας αποπροσανατολίσει από την κατανόηση της κύριας διαδικασίας πολλαπλής παλινδρόμησης) αυτό θα το κάνουμε στο τέλος.

Το dataset για αυτή την ενότητα βρίσκεται [εδώ](#).

Διεξαγωγή της ανάλυσης σε SPSS

Το βίντεο του της ενότητας (Regression), χωρίς τον έλεγχο των προϋποθέσεων, βρίσκεται εδώ: <https://vimeo.com/424029487>. Για τον έλεγχο των προϋποθέσεων δείτε το επόμενο (ολοκληρωμένο) βίντεο.

Analyze --> Regression --> Linear

Στη συνέχεια δηλώνουμε στα σχετικά πεδία την εξαρτημένη και τις προβλεπτικές (ανεξάρτητες) μεταβλητές.

Στο Statistics, είναι ήδη επιλεγμένα τα Estimates και Model fit, και επιλέγουμε και το Descriptives.

orrelations

		Exam score	Hours spent revising	Enjoyment of subject
Pearson Correlation	Exam score	1.000	.544	.580
	Hours spent revising	.544	1.000	.514
	Enjoyment of subject	.580	.514	1.000
Sig. (1-tailed)	Exam score	.	.001	.000
	Hours spent revising	.001	.	.002
	Enjoyment of subject	.000	.002	.
N	Exam score	29	29	29
	Hours spent revising	29	29	29
	Enjoyment of subject	29	29	29

Στην ανάλυση παλινδρόμησης θέλουμε οι προβλεπτικές (ανεξάρτητες) μεταβλητές να συσχετίζονται με την εξαρτημένη μεταβλητή, αλλιώς δεν θα υπήρχε λόγος να τις συμπεριλάβουμε στο μοντέλο που χτίζουμε. Από την άλλη, οι προβλεπτικές μεταβλητές δεν θέλουμε να συσχετίζονται ισχυρά μεταξύ τους. Αν η σχέση των ανεξάρτητων μεταβλητών είναι $r \geq 0.8$ τότε δημιουργείται μια κατάσταση που λέγεται collinearity (συγγραμικότητα) και απειλεί την αξιοπιστία του μοντέλου μας. Σε αυτή την περίπτωση μπορεί να χρειαστεί να αφαιρέσουμε μια προβλεπτική μεταβλητή του μοντέλου μας.

Model Summary

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	.647^a	.418	.373	17.97120

a. Predictors: (Constant), Enjoyment of subject, Hours spent revising

$R = 0.647$ είναι ο συντελεστής συσχέτισης ανάμεσα στην εξαρτημένη μεταβλητή με όλες τις ανεξάρτητες μαζί. Ο συντελεστής αυτός είναι σχετικά υψηλός που ερμηνεύεται ότι πρόκειται για ένα καλό μοντέλο για την πρόβλεψη της εξαρτημένης μεταβλητής. Το R^2 εκφράζει το ποσοστό της διακύμανσης στην εξαρτημένη μεταβλητή που μπορεί να εξηγηθεί από το μοντέλο μας. Εν προκειμένω, το 42% της διακύμανσης στην εξαρτημένη μεταβλητή μπορεί να εξηγηθεί από τις προβλεπτικές μεταβλητές του μοντέλου μας.

Η ANOVA που ακολουθεί ελέγχει αν η πρόβλεψη του μοντέλου μας είναι στατιστικά σημαντική.

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
	Regression	6033.628	2	3016.814	9.341	.001^b
1	Residual	8397.062	26	322.964		
	Total	14430.690	28			

a. Dependent Variable: Exam score

b. Predictors: (Constant), Enjoyment of subject, Hours spent revising

Αναφορά αποτελέσματος:

Η ανάλυση έδειξε ότι το μοντέλο πρόβλεψης της επίδοσης στις εξετάσεις είναι στατιστικά σημαντικό, $F(2,26) = 9.35$, $p = .001$.

Ενώ ο προηγούμενος πίνακας μας δίνει την πληροφορία ότι συνολικά το μοντέλο μας είναι στατιστικά σημαντικό, ο παρακάτω πίνακας προσδιορίζει με ακρίβεια το μοντέλο, δηλαδή μας προσδιορίζει τη συμμετοχή κάθε μιας ανεξάρτητης μεταβλητής στην εξαρτημένη μεταβλητή.

Coefficients^a

Model	Unstandardized Coefficients		Standardized Coefficients	t	Sig.
	B	Std. Error	Beta		
1	(Constant)	20.647		2.167	.040
	Hours spent revising	.295	.333	1.911	.067
	Enjoyment of subject	.668	.408	2.342	.027

a. Dependent Variable: Exam score

Από τον παραπάνω πίνακα γίνεται εμφανές ότι το ενδιαφέρον για το αντικείμενο (enjoyment) συνεισφέρει σημαντικά στην πρόβλεψη της επίδοσης ($p = .027$), ενώ οι ώρες επανάληψης (revision) όχι ($p = .067$).

Το μοντέλο σε μια γραμμική ανάλυση παλινδρόμησης έχει την παρακάτω γενική μορφή:

$$Y = B_0 + B_1X_1 + B_2X_2 + \dots$$

Και στο παράδειγμά μας:

$$\text{Exam score} = 20.657 + .295 * \text{Revision} + .668 * \text{Enjoyment}$$

Συνολική αναφορά των αποτελεσμάτων

Διεξάγαμε μια πολλαπλή ανάλυση παλινδρόμησης για να διερευνήσουμε αν η ένταση της επανάληψης και η ευχαρίστηση που αντλεί ο φοιτητής από το αντικείμενο αποτελούν στατιστικά σημαντικούς παράγοντες για την επίδοση στις εξετάσεις. Η ανάλυση έδειξε ότι το μοντέλο πρόβλεψης της επίδοσης στις εξετάσεις είναι στατιστικά σημαντικό, $F(2,26) = 9.35$, $p = .001$, με το 42% της διακύμανσης στην εξαρτημένη μεταβλητή να μπορεί να εξηγηθεί από τις προβλεπτικές μεταβλητές του μοντέλου. Η ευχαρίστηση από το αντικείμενο (enjoyment) συνεισφέρει σημαντικά στην πρόβλεψη της επίδοσης ($p = .027$), ενώ οι ώρες επανάληψης (revision) όχι ($p = .067$). Το τελικό προβλεπτικό μοντέλο (εξίσωση παλινδρόμησης) είναι:

$$\text{Exam score} = 20.657 + .295 * \text{Revision} + .668 * \text{Enjoyment}$$

Έλεγχος Προϋποθέσεων ανάλυσης παλινδρόμησης

“Προϋπόθεση #1 • Linearity. A pretty fundamental assumption of the linear regression model is that the relationship between X and Y actually is linear! Regardless of whether it's a simple regression or a multiple regression, we assume that the relationships involved are linear.

Προϋπόθεση #2 • Uncorrelated predictors. The idea here is that, in a multiple regression model, you don't want your predictors to be too strongly correlated with each other. This isn't “technically” an assumption of the regression model, but in practice it's required. Predictors that are too strongly correlated with each other (referred to as “collinearity”) can cause problems when evaluating the model.

Προϋπόθεση #3: • Residuals are independent of each other. This is really just a “catch all” assumption, to the effect that “there's nothing else funny going on in the residuals”. If there is something weird (e.g., the residuals all depend heavily on some other unmeasured variable) going on, it might screw things up.

Προϋπόθεση #4 • Homogeneity of variance. Strictly speaking, the regression model assumes that each residual is generated from a normal distribution with mean 0, and (more importantly for the current purposes) with a standard deviation σ that is the same for every single residual. In practice, it's impossible to test the assumption that every residual is identically distributed. Instead, what we care about is that the standard deviation of the residual is the same for all values of Y, and (if we're being especially paranoid) all values of every predictor X in the model.

Προϋπόθεση #5 • Normality. Like many of the models in statistics, basic simple or multiple linear regression relies on an assumption of normality. Specifically, it assumes that the residuals are normally distributed. It's actually okay if the predictors X and the outcome Y are nonnormal, so long as the residuals are normal.

Προϋπόθεση #6 • No ‘bad’ outliers. Again, not actually a technical assumption of the model (or rather, it's sort of implied by all the others), but there is an implicit assumption that your regression model isn't being too strongly influenced by one or two anomalous data points because this raises questions about the adequacy of the model and the trustworthiness of the data in some cases.” (Navarro & Foxcroft, 2019, pp. 309-310)

Πηγή: Navarro DJ and Foxcroft DR (2019). *Learning statistics with jamovi: A tutorial for psychology students and other beginners*. Retrieved from: <http://learnstatswithjamovi.com>

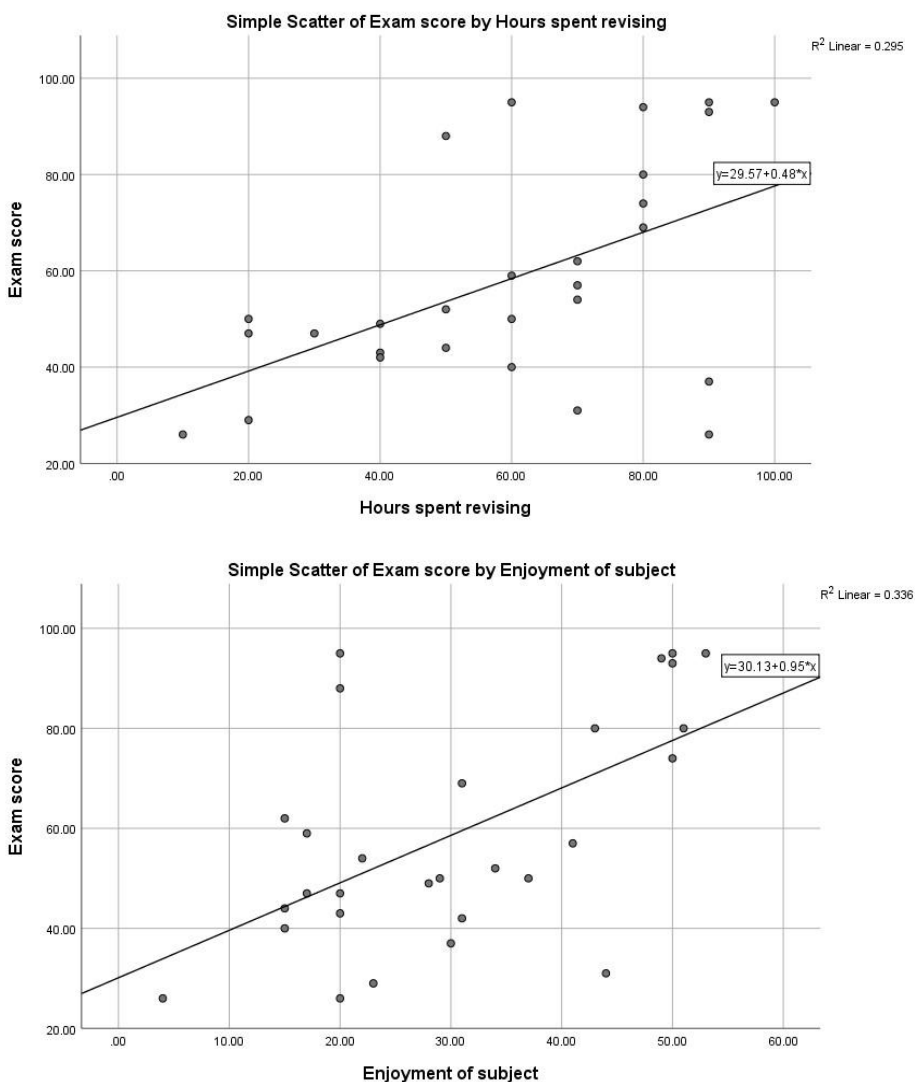
Έλεγχος προϋποθέσεων γραμμικής παλινδρόμησης

Το ολοκληρωμένο βίντεο του μαθήματος με τον έλεγχο των προϋποθέσεων (Regression Assumptions) βρίσκεται εδώ: <https://vimeo.com/424034089>

Προϋπόθεση #1: Η συσχέτιση ανάμεσα στις ανεξάρτητες μεταβλητές και την εξαρτημένη μεταβλητή είναι ευθύγραμμη.

Για αυτόν το έλεγχο δημιουργούμε ένα διάγραμμα σκεδασμού (scatter plot) για κάθε μια ανεξάρτητη μεταβλητή με την εξαρτημένη.

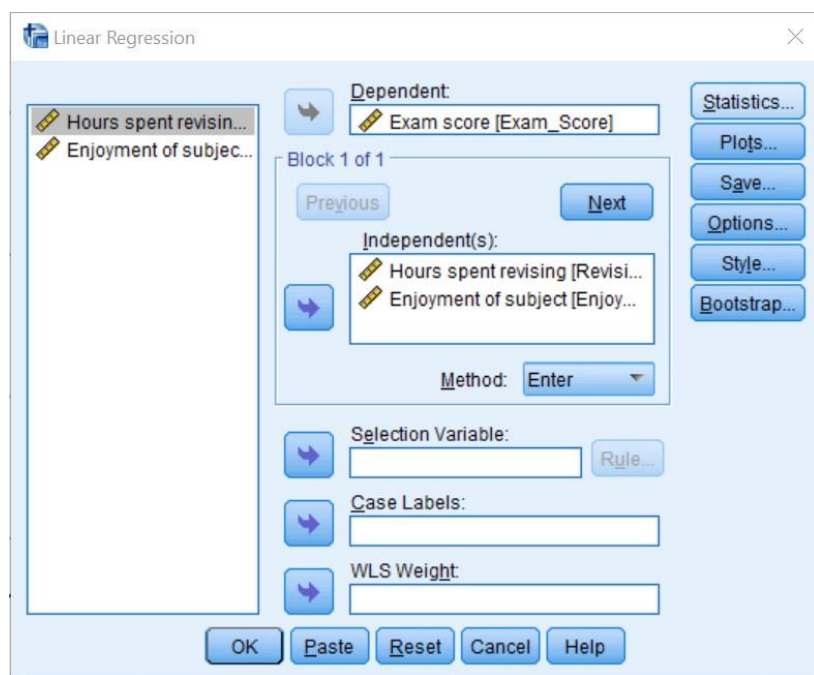
Graphs --> Chart Builder --> Scatter/dot (Simple Scatter)



Παρατηρήστε στα διαγράμματα τα κατάλοιπα/υπόλοιπα (residuals), δηλαδή την κάθετη απόσταση/απόκλιση κάθε παρατηρούμενης τιμής (κουκίδα) από την 'ιδανική' τιμή του γραμμικού μοντέλου πάνω στην γραμμή.

Για τον έλεγχο των υπόλοιπων προϋποθέσεων ξανατρέχουμε την ανάλυση παλινδρόμησης.

Analyze --> Regression --> Linear



Προϋπόθεση #2: Δεν υπάρχει collinearity ανάμεσα στις προβλεπτικές (ανεξάρτητες) μεταβλητές, δηλαδή οι προβλεπτικές μεταβλητές δεν έχουν μεγάλη συσχέτιση μεταξύ τους. (Uncorrelated predictors).

Στο Statistics επιλέγουμε Collinearity diagnostics.

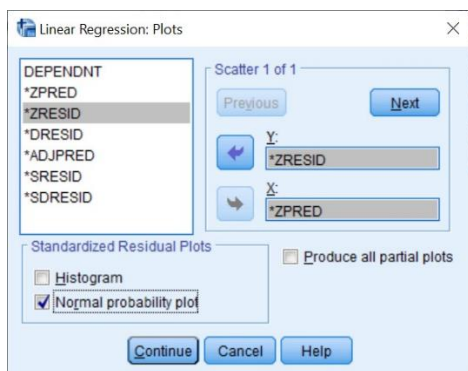
Προϋπόθεση #3: Οι τιμές των residuals είναι ανεξάρτητες μεταξύ τους. (Residuals are independent of each other)

Στο Statistics επιλέγουμε Durbin-Watson.

Προϋπόθεση #4: Η ομοσκεδαστικότητα των residuals. (Homogeneity of variance). Δηλαδή η διασπορά των residuals γύρω από τη γραμμή παλινδρόμησης πρέπει να είναι ομοιόμορφη σε όλο το μήκος της γραμμής. Για να το ελέγξουμε αυτό θα πρέπει να φτιάξουμε ένα διάγραμμα σκεδασμού που θα απεικονίζει ταυτόχρονα και τις δυο προβλεπτικές μεταβλητές.

Προϋπόθεση #5: Κανονική κατανομή των residuals. (Normality).

Τις τελευταίες δύο προϋποθέσεις τις ελέγχουμε από το μενού PLOT.

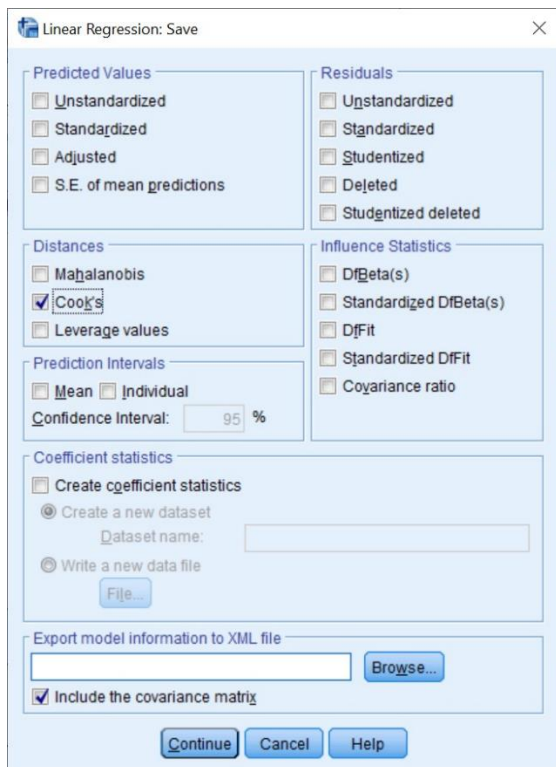


ZRESID: Οι (τυπικές) τιμές των residuals (δηλαδή τα residuals σε τιμές Z)

ZPRED: Οι (τυπικές) τιμές πρόβλεψης της εξαρτημένης μεταβλητής βάση του μοντέλου.

Προϋπόθεση #6: Δεν υπάρχουν ακραίες τιμές που να επηρεάζουν την αντιπροσωπευτικότητα του μοντέλου μας. (No 'bad' outliers).

Μπαίνουμε στο Save (κάτω από το Statistics) και μετά επιλέγουμε Cook's distance.



Στη συνέχεια το SPSS κάνει την ανάλυση παλινδρόμησης μαζί με τον έλεγχο των προϋποθέσεων.

Αποτελέσματα

Προϋπόθεση #2 (collinearity - συγγραμικότητα): Στον παρακάτω πίνακα δεν βλέπουμε κάποια ιδιαίτερα ισχυρή συσχέτιση ανάμεσα στις μεταβλητές. Αν κάποια μεταβλητή έχει συσχέτιση πάνω από 0.8, μπορεί να χρειαστεί να την αφαιρέσουμε από το μοντέλο μας.

Correlations

		Exam score	Hours spent revising	Enjoyment of subject
Pearson Correlation	Exam score	1.000	.544	.580
	Hours spent revising	.544	1.000	.514
	Enjoyment of subject	.580	.514	1.000
Sig. (1-tailed)	Exam score	.	.001	.000
	Hours spent revising	.001	.	.002
	Enjoyment of subject	.000	.002	.
N	Exam score	29	29	29
	Hours spent revising	29	29	29
	Enjoyment of subject	29	29	29

Προϋπόθεση #3 (Residuals are independent of each other):

Στον παρακάτω πίνακα, η τιμή του Durbin-Watson κυμαίνεται από 0 έως 4. Θέλουμε να είναι κοντά στο δύο. Για τιμές πάνω από 3 ή κάτω από 1, υπάρχει πρόβλημα με την ανεξαρτησία των residuals.

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Durbin-Watson
1	.647 ^a	.418	.373	17.97120	1.931

a. Predictors: (Constant), Enjoyment of subject, Hours spent revising

b. Dependent Variable: Exam score

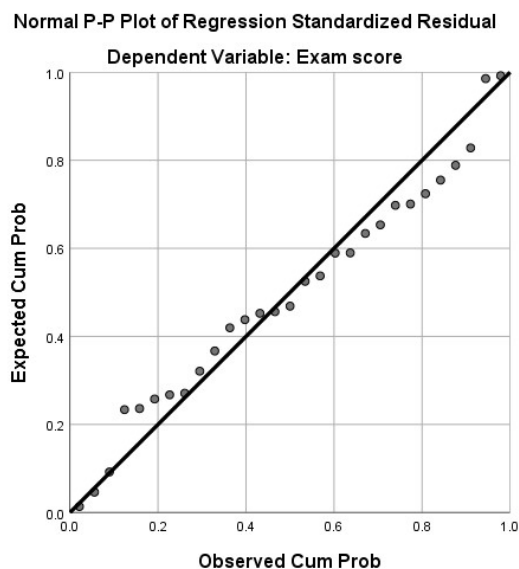
Προϋπόθεση #2 (collinearity): Στον παρακάτω πίνακα εξετάζουμε τα Collinearity Statistics ώστε να διαπιστώσουμε ότι οι προβλεπτικές μεταβλητές δεν είναι όντως ισχυρά σχετιζόμενες. Χρειάζεται το Tolerance να είναι πάνω από 0.2 και το VIF πολύ κάτω από 10.

Coefficients^a

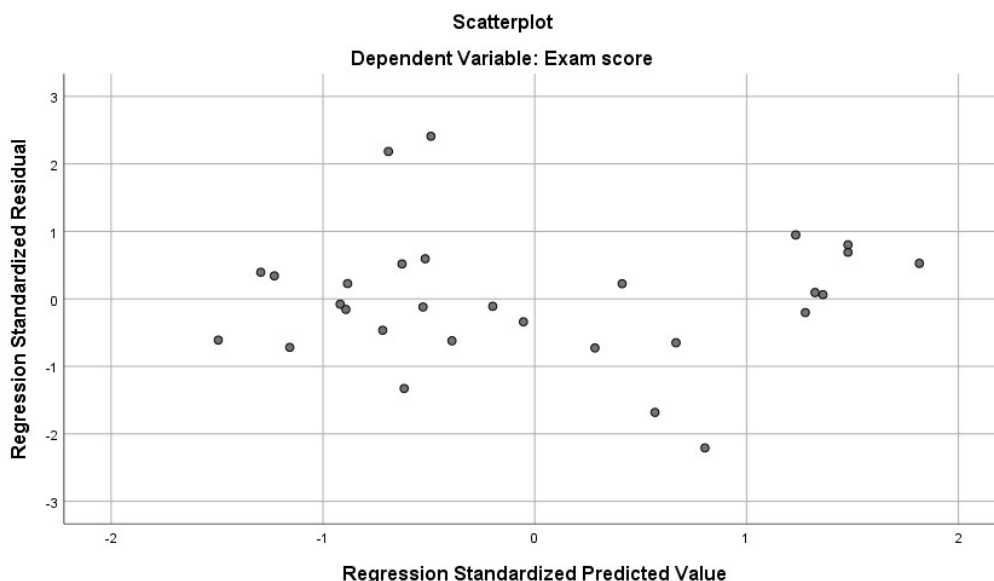
Model		Unstandardized Coefficients		Standardized Coefficients		t	Sig.	Collinearity Statistics	
		B	Std. Error	Beta				Tolerance	VIF
	(Constant)	20.647	9.530			2.167	.040		
1	Hours spent revising	.295	.154	.333		1.911	.067	.735	1.360
	Enjoyment of subject	.668	.285	.408		2.342	.027	.735	1.360

a. Dependent Variable: Exam score

Προϋπόθεση #5 (Normality - Κανονική κατανομή των residuals): Στο παρακάτω P-P διάγραμμα, αν οι τελίτσες είναι κοντά στην γραμμή, τότε έχουμε κανονική κατανομή. Σε αυτή την περίπτωση ελάχιστες τελίτσες είναι κοντά στη γραμμή, οπότε το κριτήριο της κανονικότητας δεν πρέπει να θεωρήσουμε ότι καλύπτεται ικανοποιητικά και αυτό χρειάζεται να αναφερθεί στα αποτελέσματα.



Προϋπόθεση #4 (Homogeneity of variance - Η ομοσκεδαστικότητα των residuals): Στο παρακάτω διάγραμμα απεικονίζεται οι προβλεπόμενες τιμές του μοντέλου σε σχέση με τις κανονικές τιμές των residuals των δεδομένων. Εξετάζουμε κατά πόσον στο διάγραμμα σκεδασμού καθώς αυξάνεται το X η διακύμανση των τιμών παραμένει σταθερή.



Προϋπόθεση #6 (Δεν υπάρχουν ακραίες τιμές που να επηρεάζουν την αντιπροσωπευτικότητα του μοντέλου μας No 'bad' outliers): Επιστρέφουμε στο Data View και εξετάζουμε την στήλη που έχει δημιουργηθεί με το Cook's Distance. Κάθε τιμή πάνω από 1 ενδεχόμενα αντιπροσωπεύει κάποια ακραία τιμή που μπορεί να επηρεάσει την ανάλυσή μας και ενδεχόμενα να πρέπει να απομακρυνθεί.

Στην αναφορά των αποτελεσμάτων γράφουμε πιθανές παραβιάσεις των προϋποθέσεων.

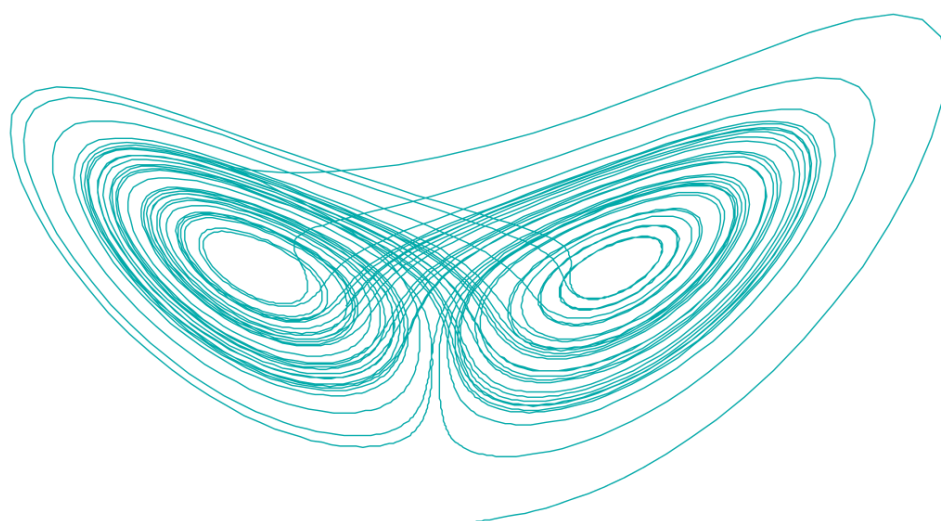
Κατεβάστε το dataset για την πολλαπλή γραμμική παλινδρόμηση από [εδώ](#)

Το βίντεο της ανάλυσης σε **SPSS**

Οδηγίες διεξαγωγής της ανάλυσης σε Jamovi

Δείτε το παρακάτω βίντεο στο οποίο η διαδικασία περιγράφεται σε περιβάλλον Jamovi για το ίδιο dataset.

Το βίντεο της ανάλυσης σε **JAMOV**



Βιβλιογραφία

- Brace, N., Kemp, R., & Snelgar, R. (2009). *SPSS for psychologists (4th ed)*. Routledge.
- Dancey, C. P., & Reidy, J. (2011). *Statistics without maths for psychology (5th ed)*. Prentice Hall/Pearson.
- Field, A. P. (2020). *Discovering statistics using IBM SPSS statistics (Fourth edition)*. SAGE Publications.
- Navarro, D. J., & Foxcroft, D. R. (2018). *Learning statistics with jamovi: A tutorial for psychology students and other beginners*. Danielle J. Navarro and David R. Foxcroft. <https://doi.org/10.24384/HGC3-7P15>
- Ρούσσος, Π., & Τσαούσης, Γ. (2011). *Στατιστική στις επιστήμες της συμπεριφοράς με τη χρήση του SPSS*. Εκδόσεις τόπος.
- Ρούσσος, Π. (2019). *Εγχειρίδιο Jamovi*. Ανακτήθηκε από: [\[link\]](#)

