

ΒΑΣΙΛΗΣ ΓΙΑΛΛΑΜΑΣ
ΚΩΝΣΤΑΝΤΙΝΟΣ ΛΑΒΙΔΑΣ
ΔΙΟΝΥΣΗΣ ΜΑΝΕΣΗΣ

ΣΤΑΤΙΣΤΙΚΕΣ ΜΕΘΟΔΟΙ ΚΑΙ ΤΕΧΝΙΚΕΣ ΣΤΙΣ ΚΟΙΝΩΝΙΚΕΣ ΕΠΙΣΤΗΜΕΣ ΜΕ ΤΗ ΧΡΗΣΗ SPSS



ΒΑΣΙΛΗΣ ΓΙΑΛΛΑΜΑΣ
Καθηγητής ΤΕΑΠΗ, ΕΚΠΑ

ΚΩΝΣΤΑΝΤΙΝΟΣ ΛΑΒΙΔΑΣ
Επίκουρος Καθηγητής ΤΕΕΑΠΗ, Πανεπιστήμιο Πατρών

ΔΙΟΝΥΣΗΣ ΜΑΝΕΣΗΣ
Εργαστηριακό Διδακτικό Προσωπικό ΤΕΑΠΗ, ΕΚΠΑ

ΣΤΑΤΙΣΤΙΚΕΣ ΜΕΘΟΔΟΙ ΚΑΙ ΤΕΧΝΙΚΕΣ ΣΤΙΣ ΚΟΙΝΩΝΙΚΕΣ ΕΠΙΣΤΗΜΕΣ ΜΕ ΤΗ ΧΡΗΣΗ SPSS



Στατιστικές μέθοδοι και τεχνικές στις κοινωνικές επιστήμες με τη χρήση SPSS

Συγγραφή

Βασίλης Γιαλαμάς

Κωνσταντίνος Λαβίδας

Διονύσης Μάνεσης

Συντελεστές έκδοσης

Γλωσσική Επιμέλεια: Ανδριανή Π. Καλογερά

Γραφιστική Επιμέλεια: Κατερίνα Γαλάτη

Copyright © 2024, ΚΑΛΛΙΠΟΣ, ΑΝΟΙΚΤΕΣ ΑΚΑΔΗΜΑΪΚΕΣ ΕΚΔΟΣΕΙΣ
(ΣΕΑΒ + ΕΛΚΕ-ΕΜΠ)



Το παρόν έργο αδειοδοτείται υπό τους όρους της άδειας Creative Commons Αναφορά Δημιουργού - Μη Εμπορική Χρήση - Παρόμοια Διανομή 4.0. Για να δείτε ένα αντίγραφο της άδειας αυτής επισκεφτείτε τον ιστότοπο <https://creativecommons.org/licenses/by-nc-sa/4.0/deed.el>

ΚΑΛΛΙΠΟΣ
Εθνικό Μετσόβιο Πολυτεχνείο
Ηρώων Πολυτεχνείου 9, 15780 Ζωγράφου

www.kallipos.gr

ISBN: 978-618-228-204-5

Βιβλιογραφική Αναφορά: Γιαλαμάς, Β., Λαβίδας, Κ., & Μάνεσης, Δ. (2024). *Στατιστικές μέθοδοι και τεχνικές στις κοινωνικές επιστήμες με τη χρήση SPSS* [Μεταπτυχιακό εγχειρίδιο]. Κάλλιπος, Ανοικτές Ακαδημαϊκές Εκδόσεις. <https://dx.doi.org/10.57713/kallipos-439>

*Το βιβλίο αφιερώνεται σε όλους τους φοιτητές αυριανούς ή και
νυν εκπαιδευτικούς, που συμμετείχαν στις διδασκαλίες των
προγραμμάτων προπτυχιακών και μεταπτυχιακών σπουδών, των
τμημάτων νηπιαγωγών του Πανεπιστημίου Πατρών και του
Εθνικού και Καποδιστριακού Πανεπιστημίου Αθηνών.*

Πίνακας περιεχομένων

Πίνακας περιεχομένων	7
Πίνακας Συντομογραφιών/Ακρωνυμίων.....	11
Πρόλογος	13
Σύντομα βιογραφικά των συγγραφέων	14
Κεφάλαιο 1 Εισαγωγή στην ανάλυση δεδομένων στο υπολογιστικό περιβάλλον του SPSS	15
1.1 Εισαγωγή στη στατιστική	15
1.1.1 Βασικές στατιστικές έννοιες	16
1.1.2 Στάδια διεξαγωγής μιας ποσοτικής έρευνας	18
1.2 Εισαγωγή στο SPSS	18
1.2.1 Βασικά παράθυρα του SPSS	19
1.2.2 Μενού του SPSS.....	20
1.2.3 Καρτέλα Data View.....	21
1.2.4 Καρτέλα Variable View	21
1.2.5 Εισαγωγή δεδομένων στο SPSS.....	25
1.2.6 Μεταφορά δεδομένων στο SPSS από αρχείο Excel.....	30
1.2.7 Αποθήκευση αρχείων δεδομένων.....	31
1.3 Προβλήματα για εξάσκηση και ανατροφοδότηση	32
Βιβλιογραφία	33
Κεφάλαιο 2 Μετασχηματισμός Δεδομένων	35
2.1 Επέκταση των δεδομένων.....	35
2.1.1 Αυτόματη επανακωδικοποίηση – Automatic Recode	35
2.1.2 Επανακωδικοποίηση στις ίδιες μεταβλητές – Recode into Same Variables.....	37
2.1.3 Επανακωδικοποίηση σε διαφορετική μεταβλητή – Recode into Different Variables	39
2.1.4 Διαδικασία υπολογισμού τιμών νέας μεταβλητής – Compute Variable	43
Βιβλιογραφία	47
Κεφάλαιο 3 Χειρισμός Δεδομένων.....	49
3.1 Επεξεργασία και χειρισμός δεδομένων	49
3.1.1 Ταξινόμηση περιπτώσεων – Sort Cases	49
3.1.2 Ταξινόμηση μεταβλητών – Sort Variables.....	50
3.1.3 Διαχωρισμός αρχείου σε υποομάδες – Split File	51
3.1.4 Επιλογή περιπτώσεων – Select Cases	53
3.1.5 Στάθμιση περιπτώσεων – Weight Cases	56
3.1.6 Συγχώνευση αρχείων – Merge Files.....	58
3.1.6.1 Merge Files – Add Cases	59
3.1.6.2 Merge Files – Add Variables	60
Βιβλιογραφία	63
Κεφάλαιο 4 Περιγραφική στατιστική στο SPSS.....	65
4.1 Εισαγωγή στην περιγραφική στατιστική ανάλυση.....	65

4.1.1 Πίνακες συχνοτήτων	65
4.1.2 Περιγραφικά στατιστικά.....	67
4.1.3 Κατασκευή πίνακα κατανομής συχνοτήτων στο περιβάλλον του SPSS.....	69
4.1.3.1 Περιγραφικά στατιστικά μέσω της διαδικασίας «Frequencies».....	72
4.1.3.2 Περιγραφικά στατιστικά μέσω της διαδικασίας «Descriptives»	73
4.1.3.3 Περιγραφικά στατιστικά μέσω της διαδικασίας «Explore».....	75
4.1.4 Σύνθετη παρουσίαση δύο μεταβλητών	78
4.1.4.1 Σύνθετη παρουσίαση μιας ποιοτικής και μιας ποσοτικής μεταβλητής: «Compare means».....	78
4.1.4.2 Σύνθετη παρουσίαση μιας ποιοτικής και μιας ποσοτικής μεταβλητής: «Custom Tables».....	79
4.1.4.3 Σύνθετη παρουσίαση δύο ποιοτικών μεταβλητών: «Crosstabs»	81
4.2 Προβλήματα για εξάσκηση και ανατροφοδότηση	84
Βιβλιογραφία	87
Κεφάλαιο 5 Ανάλυση ερωτήσεων πολλαπλών απαντήσεων.....	89
5.1 Εισαγωγή στην ανάλυση ερωτήσεων πολλαπλών απαντήσεων.....	89
5.1.1 Καταμέτρηση επιλεγμένων επιλογών και εντοπισμός μη ερωτηθέντων	92
5.1.2 Δημιουργία σετ πολλαπλών απαντήσεων στο SPSS.....	94
5.1.3 Δημιουργία πίνακα συχνοτήτων για το σετ πολλαπλών απαντήσεων	95
5.2 Προβλήματα για εξάσκηση και ανατροφοδότηση	98
Βιβλιογραφία	99
Κεφάλαιο 6 Εισαγωγή στην Εκτιμητική και στον Έλεγχο Υποθέσεων.....	101
6.1 Εισαγωγή στην επαγωγική στατιστική.....	101
6.1.1 Η Εκτιμητική.....	102
6.1.1.1 Υπολογισμός του διαστήματος εμπιστοσύνης με το SPSS.....	103
6.1.2 Έλεγχος Υποθέσεων.....	104
6.1.2.1 Διαδικασία ελέγχου υποθέσεων	105
6.2 Σχεδιασμοί ανεξάρτητων και εξαρτημένων δειγμάτων.....	108
6.3 Έλεγχος δεδομένων (Data Screening) με το SPSS	109
6.3.1 Έλεγχος της κανονικότητας των τιμών μιας ποσοτικής μεταβλητής.....	112
6.4 Προβλήματα για εξάσκηση και ανατροφοδότηση	114
Βιβλιογραφία	116
Κεφάλαιο 7 Έλεγχος μέσων τιμών (T-Tests).....	117
7.1 Εισαγωγή στο t-test.....	117
7.1.1 Έλεγχος t ενός δείγματος – One-Sample T-Test.....	117
7.1.2 Έλεγχος t ανεξάρτητων δειγμάτων – Independent-Samples T-Test	123
7.1.3 Έλεγχος t δειγμάτων κατά ζεύγη – Paired-Samples T-Test.....	129
7.2 Προβλήματα για εξάσκηση και ανατροφοδότηση	133
Βιβλιογραφία	135
Κεφάλαιο 8 Έλεγχος χ^2	137
8.1 Έλεγχος χ^2	137

8.1.1 Έλεγχος καλής προσαρμογής	138
8.1.1.1 Έλεγχος χ^2 καλής προσαρμογής με χρήση λογισμικού (SPSS).....	140
8.1.2 Έλεγχος ανεξαρτησίας δύο ποιοτικών μεταβλητών.....	141
8.1.2.1 Πίνακες συνάφειας.....	141
8.1.2.2 Σχετικές συχνότητες κατηγορικών συμμεταβλητών.....	142
8.1.2.3 Περιγραφή της συνάφειας κατηγορικών μεταβλητών από πίνακα σύμπτωσης	143
8.1.2.4 Γραφική απεικόνιση συνάφειας	143
8.1.2.5 Υπολογισμός της τιμής χ^2 σε πίνακες σύμπτωσης και βαθμοί ελευθερίας.....	144
8.1.2.6 Διαδικασία ελέγχου ανεξαρτησίας.....	146
8.1.2.7 Έλεγχος ποσοστών στηλών.....	146
8.1.2.8 Προϋποθέσεις στην εκτέλεση του ελέγχου χ^2	147
8.1.2.9 Συντελεστές συνάφειας.....	147
8.1.2.10 Πίνακες συνάφειας και έλεγχος χ^2 ανεξαρτησίας στο περιβάλλον του SPSS	149
8.2 Προβλήματα για εξάσκηση και ανατροφοδότηση	152
Βιβλιογραφία	154
Κεφάλαιο 9 Ανάλυση Διακύμανσης (ANOVA)	155
9.1 Εισαγωγή στην ανάλυση διακύμανσης	155
9.1.1 Ανάλυση διακύμανσης προς έναν παράγοντα – One-Way ANOVA.....	155
9.1.2 Ανάλυση διακύμανσης επαναλαμβανόμενων μετρήσεων	167
9.1.3 Ανάλυση διακύμανσης με δύο ή περισσότερους παράγοντες.....	173
9.2 Προβλήματα για εξάσκηση και ανατροφοδότηση	183
Βιβλιογραφία	185
Κεφάλαιο 10 Μη Παραμετρικοί Έλεγχοι.....	187
10.1 Εισαγωγή στους μη παραμετρικούς ελέγχους	187
10.1.1 Έλεγχος Mann-Whitney για δύο ανεξάρτητα δείγματα	188
10.1.2 Έλεγχος Wilcoxon για δύο εξαρτημένα δείγματα.....	191
10.1.3 Έλεγχος Kruskal-Wallis για περισσότερα από δύο ανεξάρτητα δείγματα	194
10.1.4 Έλεγχος Friedman για περισσότερα από δύο εξαρτημένα δείγματα	198
10.2 Προβλήματα για εξάσκηση και ανατροφοδότηση	203
Βιβλιογραφία	204
Κεφάλαιο 11 Συσχέτιση και Απλή Γραμμική Παλινδρόμηση	205
11.1 Εισαγωγή στη συσχέτιση.....	205
11.1.1 Χαρακτηριστικά της συνάφειας	206
11.1.2 Ο Συντελεστής συσχέτισης Pearson και η αξιολόγηση της τιμής του	209
11.1.3 Έλεγχος υποθέσεων για την τιμή του συντελεστή συσχέτισης Pearson r.....	212
11.1.4 Πίνακας συσχετίσεων και έλεγχος υποθέσεων για την τιμή του συντελεστή συσχέτισης Pearson στο SPSS.....	212
11.1.5 Ερμηνεία του συντελεστή Pearson.....	214
11.1.6 Πού χρησιμοποιείται ο συντελεστής γραμμικής συσχέτισης;.....	215

11.2 Εισαγωγή στην παλινδρόμηση	216
11.2.1 Η ευθεία των ελαχίστων τετραγώνων	217
11.2.2 Η παλινδρόμηση και η ευθεία των ελαχίστων τετραγώνων με το SPSS	218
11.3 Προϋποθέσεις εγκυρότητας της παλινδρόμησης και του συντελεστή r Pearson και μη παραμετρικοί συντελεστές συσχέτισης	221
11.3.1 Ο Υπολογισμός του συντελεστή συσχέτισης Spearman	221
11.4 Προβλήματα για εξάσκηση και ανατροφοδότηση	225
Βιβλιογραφία	227
Παράρτημα	229
Κεφάλαιο 12 Πολλαπλή Γραμμική Παλινδρόμηση	231
12.1 Εισαγωγή στην Πολλαπλή παλινδρόμηση	231
12.1.1 Περιγραφή και έλεγχος υποδείγματος.....	231
12.1.2 Κωδικοποίηση κατηγορικών ανεξάρτητων μεταβλητών	232
12.1.3 Προϋποθέσεις εγκυρότητας του υποδείγματος της πολλαπλής παλινδρόμησης	232
12.1.4 Η προσαρμογή του υποδείγματος	239
12.1.5 Η σημαντικότητα και η ερμηνεία των μερικών συντελεστών	240
12.1.6 Η κατασκευή ενός υποδείγματος (παλινδρόμηση πάνω σε όλα τα υποσύνολα των ανεξάρτητων μεταβλητών και Στατιστική Παλινδρόμηση).....	243
12.1.7 Περιπτώσεις με ισχυρή επίδραση στην προσαρμογή και στους συντελεστές του υποδείγματος	246
12.1.8. Η διασταυρούμενη εγκυρότητα του υποδείγματος	250
12.2 Προβλήματα για εξάσκηση και ανατροφοδότηση	251
Βιβλιογραφία	254
Κεφάλαιο 13 Λογαριθμική Παλινδρόμηση	255
13.1 Εισαγωγή στη λογαριθμική παλινδρόμηση	255
13.1.1 Η λογαριθμική παλινδρόμηση στο SPSS	256
13.1.2 Σχετική πιθανότητα και η συνάρτηση «logit».....	260
13.1.3 Οι μερικοί συντελεστές και η ερμηνεία τους	261
13.1.4 Η εκτίμηση και η σημαντικότητα των μερικών συντελεστών και οι δείκτες καλής προσαρμογής του υποδείγματος.....	261
13.1.5 Προϋποθέσεις της ΛΠ.....	267
13.2 Προβλήματα για εξάσκηση και ανατροφοδότηση	268
Βιβλιογραφία	272
Κεφάλαιο 14 Διερευνητική Παραγοντική Ανάλυση	273
14.1 Εισαγωγή στην Παραγοντική Ανάλυση	273
14.1.1 Οι στόχοι της Διερευνητικής Παραγοντικής Ανάλυσης	275
14.1.2 Η εκτέλεση της ΔΠΑ.....	275
14.2 Παραγοντική εγκυρότητα και αξιοπιστία των παραγόντων	286
14.3 Προβλήματα για εξάσκηση και ανατροφοδότηση	291
Βιβλιογραφία	293

Πίνακας Συντομογραφιών/Ακρωνυμίων

ANOVA	Analysis of Variance
AVE	Average Variance Extracted
CFA	Confirmatory Factor Analysis
CI	Confidence Interval
EFA	Exploratory Factor Analysis
HTMT	Heterotrait-Monotrait Ratio of Correlations
IQ	Intelligence Quotient
IQR	Interquartile Range
KMO	Kaiser–Meyer–Olkin
ML	Maximum Likelihood
OR	Odds ratio
PAF	Principal Axis Factoring
SATS	Survey of Attitudes Toward Statistics
SPSS	Statistical Package for the Social Sciences
ΔΕ	Διάστημα Εμπιστοσύνης
ΔΠΑ	Διερευνητική Παραγοντική Ανάλυση
ΕΠΑ	Επιβεβαιωτική Παραγοντική Ανάλυση
IBM	International Business Machines
ΛΠ	Λογαριθμική Παλινδρόμηση
MME	Μέσα Μαζικής Ενημέρωσης
ΤΠΕ	Τεχνολογίες της Πληροφορίας και των Επικοινωνιών

Πρόλογος

Το βιβλίο «Στατιστικές μέθοδοι και τεχνικές στις κοινωνικές επιστήμες με τη χρήση SPSS» γράφτηκε στο πλαίσιο της Δράσης «ΚΑΛΛΙΠΟΣ+» και αποτελεί μια σύντομη περιγραφή των βασικότερων διαδικασιών στατιστικής ανάλυσης ποσοτικών και ποιοτικών δεδομένων. Βασίζεται στη μεγάλη διδακτική εμπειρία των συγγραφέων στη μεθοδολογία της εκπαιδευτικής έρευνας και κυρίως στην περιγραφική και στην επαγωγική στατιστική, σε προγράμματα προπτυχιακών και μεταπτυχιακών σπουδών, κυρίως αυριανών εκπαιδευτικών.

Το βιβλίο αυτό σχεδιάστηκε κυρίως για φοιτητές ή και εκπαιδευτικούς που παρακολουθούν μαθήματα μεθοδολογίας εκπαιδευτικής έρευνας και για όλους όσους επιθυμούν να σχεδιάσουν και να υλοποιήσουν τη δική τους έρευνα. Λαμβάνοντας υπόψη ότι η στατιστική προκαλεί άγχος και δυσφορία στους φοιτητές που έχουν μικρή εμπειρία στα μαθηματικά και στη στατιστική, στο βιβλίο αυτό προσπαθούμε να παρουσιάσουμε τη στατιστική με τέτοιο τρόπο που μπορεί να γίνει εύκολα κατανοητή και αξιοποιήσιμη. Παρουσιάζεται βήμα προς βήμα πώς χρησιμοποιείται ένα εύχρηστο περιβάλλον ανάλυσης δεδομένων, καθώς και τι πρέπει να κάνει κάποιος για να προετοιμάσει και να αναλύσει τα δεδομένα του, ποιες στατιστικές τεχνικές είναι διαθέσιμες, πώς χρησιμοποιούνται και πώς ερμηνεύονται τα παραγόμενα αποτελέσματα, αλλά και πώς τα αποτελέσματα αυτά παρουσιάζονται σε μια ερευνητική αναφορά στο πλαίσιο μιας δημοσίευσης.

Υπό το πρίσμα αυτό, η παρουσίαση των στατιστικών διαδικασιών ανάλυσης δεδομένων γίνεται παράλληλα με την παρουσίαση του περιβάλλοντος επεξεργασίας και στατιστικής ανάλυσης δεδομένων SPSS (Statistical Package for the Social Sciences – Στατιστικό Πακέτο για τις Κοινωνικές Επιστήμες). Τονίζουμε ότι οι εικόνες που συνοδεύουν την παρουσίαση του περιβάλλοντος SPSS έχουν δημιουργηθεί από τους συγγραφείς κατά τη χρήση του περιβάλλοντος αυτού, το οποίο αποτελεί και το βασικό «εργαλείο» του συγγράμματος. Σε κάθε περίπτωση, γίνεται αρχικά μια εισαγωγή του θεωρητικού μέρους των στατιστικών διαδικασιών και στη συνέχεια, ακολουθώντας συγκεκριμένα παραδείγματα κυρίως της εκπαιδευτικής έρευνας, παρουσιάζονται στο περιβάλλον του SPSS λεπτομερώς η ανάλυση των δεδομένων, η ανάγνωση και η ερμηνεία των αποτελεσμάτων και η διατύπωση των συμπερασμάτων. Το βιβλίο συνοδεύεται με αρχεία δεδομένων του SPSS τα οποία υποστηρίζουν την ανάλυση των δεδομένων στα περισσότερα κεφάλαια του βιβλίου. Τα αρχεία αυτά μπορείτε να τα κατεβάσετε στον υπολογιστή σας από τον δικτυακό τόπο του βιβλίου: <http://bookfilesspss.weebly.com/>.

Το βιβλίο αποτελείται από δεκατέσσερα κεφάλαια. Πιο συγκεκριμένα, στο πρώτο κεφάλαιο γίνεται μια περιγραφή των βασικών στατιστικών εννοιών που χρειάζονται για τις αναλύσεις με τη χρήση του SPSS, καθώς και των κύριων σταδίων της διεξαγωγής μιας ποσοτικής έρευνας. Ακολουθεί μια αναλυτική περιγραφή των διαθέσιμων επιλογών των μενού του SPSS, καθώς και η περιγραφή της διαδικασίας της εισαγωγής δεδομένων στο SPSS, με τη χρήση συγκεκριμένου ερωτηματολογίου. Έπειτα, παρουσιάζονται οι ενέργειες που απαιτούνται για τη μεταφορά δεδομένων στο SPSS από αρχείο Excel. Στο δεύτερο κεφάλαιο αναφέρεται ο μετασχηματισμός των δεδομένων που είναι απαραίτητος για τη συμβατότητα των μεταβλητών με την εκάστοτε μεθοδολογία. Στο τρίτο κεφάλαιο παρουσιάζονται οι βασικότερες εντολές για τον χειρισμό και την τροποποίηση των δεδομένων, οι οποίες είναι απαραίτητες σε αρκετές αναλύσεις.

Στο τέταρτο κεφάλαιο γίνεται αρχικά μια σύντομη παρουσίαση των βασικών στατιστικών που περιγράφουν την κατανομή μιας μεταβλητής. Στη συνέχεια, υπολογίζονται και εξάγονται πίνακες συχνοτήτων, καθώς και περιγραφικά στατιστικά μιας ποσοτικής μεταβλητής. Τέλος, το κεφάλαιο περιλαμβάνει σύνθετη παρουσίαση δύο μεταβλητών και επιπροσθέτως προτάσεις για τον αρχικό έλεγχο των δεδομένων και τον έλεγχο της κανονικότητας των τιμών μιας ποσοτικής μεταβλητής. Στο πέμπτο κεφάλαιο γίνεται αρχικά μια σύντομη παρουσίαση των ερωτήσεων πολλαπλών απαντήσεων. Στη συνέχεια, υποδεικνύονται τρόποι οργάνωσης των δεδομένων των ερωτήσεων πολλαπλών απαντήσεων στο περιβάλλον του SPSS.

Στο έκτο κεφάλαιο γίνεται αρχικά η εκτίμηση σημείου και διαστήματος, με έμφαση στον υπολογισμό διαστήματος εμπιστοσύνης με το SPSS. Ακολουθεί μια εισαγωγή στον έλεγχο υποθέσεων, με έμφαση στη διάκριση μεταξύ αμφίπλευρου και μονόπλευρου ελέγχου. Τέλος, παρουσιάζονται οι ερευνητικοί σχεδιασμοί ανεξάρτητων και εξαρτημένων δειγμάτων. Στο έβδομο κεφάλαιο εξετάζονται τα τρία είδη παραμετρικών ελέγχων t-test: ο έλεγχος t-test ενός δείγματος, ο έλεγχος t-test δύο ανεξάρτητων δειγμάτων και ο έλεγχος t-test δύο εξαρτημένων δειγμάτων ή δειγμάτων κατά ζεύγη. Στο όγδοο κεφάλαιο αρχικά παρουσιάζεται μια σύνοψη της θεωρητικής κατανομής χ^2 . Στη συνέχεια, στο περιβάλλον του SPSS γίνεται η παρουσίαση της αξιοποίησης της θεωρητικής κατανομής χ^2 για τον έλεγχο α) καλής προσαρμογής, δηλαδή πόσο καλά ταιριάζει η κατανομή των παρατηρήσεων με μια γνωστή θεωρητική κατανομή και β) ανεξαρτησίας δύο ποιοτικών μεταβλητών,

δηλαδή του ελέγχου της σχέσης μεταξύ δύο ποιοτικών μεταβλητών. Στο ένατο κεφάλαιο εξετάζονται τα τρία είδη ελέγχου ανάλυσης διακύμανσης: ανάλυση διακύμανσης προς έναν παράγοντα – One-Way ANOVA, ανάλυση διακύμανσης επαναλαμβανόμενων μετρήσεων και ανάλυση διακύμανσης με δύο ή περισσότερους παράγοντες. Στο δέκατο κεφάλαιο διερευνώνται οι τέσσερις βασικοί μη παραμετρικοί έλεγχοι, Mann-Whitney για δύο ανεξάρτητα δείγματα, Wilcoxon για δύο εξαρτημένα δείγματα, Kruskal-Wallis για περισσότερα από δύο ανεξάρτητα δείγματα και Friedman για περισσότερα από δύο εξαρτημένα δείγματα.

Στο ενδέκατο κεφάλαιο παρουσιάζεται η γραμμική συσχέτιση μεταξύ ποσοτικών μεταβλητών με τη βοήθεια του διαγράμματος σκεδασμού. Δίνεται επίσης η ποσοτικοποίηση της κατεύθυνσης και της έντασης μιας γραμμικής συσχέτισης με τον συντελεστή συσχέτισης Pearson r . Επιπλέον, εξετάζεται το υπόδειγμα της απλής γραμμικής παλινδρόμησης. Στο δωδέκατο κεφάλαιο ερευνάται το υπόδειγμα της πολλαπλής γραμμικής παλινδρόμησης για μία εξαρτημένη ποσοτική μεταβλητή και μία ομάδα ανεξάρτητων που μπορεί να ανήκουν σε οποιονδήποτε τύπο μεταβλητής. Παρουσιάζονται και ελέγχονται οι προϋποθέσεις καταλληλότητας των κατανομών με τη βοήθεια των τυποποιημένων υπολοίπων, καθώς και η γραμμικότητα των σχέσεων μεταξύ της εξαρτημένης και ανεξάρτητων μεταβλητών. Παρουσιάζεται και ερμηνεύεται πίνακας με δείκτες προσαρμογής του υποδείγματος, καθώς και ο πίνακας των μερικών συντελεστών της εξίσωσης παλινδρόμησης με τη στατιστική τους σημαντικότητα. Παρουσιάζονται στατιστικές μέθοδοι ελάττωσης του πλήθους των ανεξάρτητων μεταβλητών. Δίνονται και εξετάζονται δείκτες που προσδιορίζουν περιπτώσεις με ισχυρή επίδραση στους συντελεστές και στην προσαρμογή του υποδείγματος. Στο δέκατο τρίτο κεφάλαιο δίνεται αρχικά μια εισαγωγή στη λογαριθμική παλινδρόμηση και, στη συνέχεια, παρουσιάζεται η λογαριθμική παλινδρόμηση στο περιβάλλον του SPSS.

Τέλος, στο δέκατο τέταρτο κεφάλαιο γίνεται μια εισαγωγή στη διερευνητική παραγοντική ανάλυση (Factor analysis), καθώς επίσης στα βήματα που απαιτούνται για την ανάλυση αυτή στο περιβάλλον του SPSS. Επιπλέον, παρουσιάζονται ζητήματα εγκυρότητας και αξιοπιστίας που θα πρέπει να λαμβάνονται υπόψη κατά τη διερεύνηση της παραγοντικής δομής μιας κλίμακας.

Σύντομα βιογραφικά των συγγραφέων

Ο **Βασίλης Γιαλαμάς** είναι Ομότιμος Καθηγητής του Τμήματος Εκπαίδευσης και Αγωγής στην Προσχολική Ηλικία, του Εθνικού και Καποδιστριακού Πανεπιστημίου Αθηνών. Το γνωστικό του αντικείμενο είναι «Στατιστική επεξεργασία ποσοτικών και ποιοτικών δεδομένων» και τα ερευνητικά του ενδιαφέροντα εστιάζονται κυρίως στη μεθοδολογία και στις εφαρμογές της Πολυδιάστατης Στατιστικής Ανάλυσης, με έμφαση στη Διερευνητική και Επιβεβαιωτική Παραγοντική Ανάλυση και στα Υποδείγματα Δομικών Εξισώσεων, στις Κοινωνικές Επιστήμες.

Ο **Κωνσταντίνος Λαβίδας** είναι Επίκουρος Καθηγητής στο Τμήμα Επιστημών Εκπαίδευσης και Αγωγής στην Προσχολική Ηλικία (ΤΕΕΑΠΗ) του Πανεπιστημίου Πατρών και παράλληλα Καθηγητής-Σύμβουλος στη θεματική ενότητα «Εκπαιδευτική Έρευνα στην Πράξη» των Μεταπτυχιακών Προγραμμάτων των Ανθρωπιστικών Επιστημών του Ελληνικού Ανοικτού Πανεπιστημίου. Το γνωστικό του αντικείμενο είναι Ποσοτικές μέθοδοι στην εκπαιδευτική έρευνα και μαθηματική εκπαίδευση. Διδάσκει Μεθοδολογία Εκπαιδευτικής Έρευνας, Ποσοτικές Προσεγγίσεις στο μεταπτυχιακό πρόγραμμα σπουδών και περιγραφική και επαγωγική Στατιστική στο προπτυχιακό πρόγραμμα σπουδών του ΤΕΕΑΠΗ. Τα ερευνητικά του ενδιαφέροντα εστιάζονται κυρίως σε ζητήματα μεθοδολογίας της εκπαιδευτικής έρευνας, ανάλυσης ποσοτικών δεδομένων, πολυμεταβλητής ανάλυσης αλλά και αξιοποίησης των ΤΠΕ για τη διδασκαλία των μαθηματικών.

Ο **Διονύσης Μάνεσης** είναι διδάκτωρ του Εθνικού και Καποδιστριακού Πανεπιστημίου Αθηνών και μέλος του Εργαστηριακού Διδακτικού Προσωπικού του Τμήματος Εκπαίδευσης και Προσχολικής Αγωγής του Εθνικού και Καποδιστριακού Πανεπιστημίου Αθηνών. Ασχολείται με την ποσοτική μεθοδολογία και ανάλυση δεδομένων και την αξιοποίηση των ψηφιακών μέσων με έμφαση στα ψηφιακά παιχνίδια στην εκπαιδευτική διαδικασία.

Κεφάλαιο 1

Εισαγωγή στην ανάλυση δεδομένων στο υπολογιστικό περιβάλλον του SPSS

Σύνοψη

Στο κεφάλαιο αυτό γίνεται περιγραφή των βασικών στατιστικών εννοιών που χρειάζονται για τις αναλύσεις με τη χρήση του SPSS, καθώς και των κύριων σταδίων διεξαγωγής μιας ποσοτικής έρευνας. Στη συνέχεια, γίνεται μια αναλυτική περιγραφή των διαθέσιμων επιλογών των μενού του SPSS. Ακολουθεί η περιγραφή της διαδικασίας εισαγωγής δεδομένων στο SPSS, με τη χρήση συγκεκριμένου ερωτηματολογίου. Έπειτα, παρουσιάζονται οι ενέργειες που απαιτούνται για τη μεταφορά δεδομένων στο SPSS από ένα αρχείο Excel.

Προαπαιτούμενη γνώση

Βασικές γνώσεις Excel.

1.1 Εισαγωγή στη στατιστική

Στις μέρες μας, η Στατιστική καταλαμβάνει μεγάλο μέρος της καθημερινότητάς μας. Τα Μέσα Μαζικής Ενημέρωσης (ΜΜΕ) χρησιμοποιούν πίνακες και σχεδιαγράμματα με μετρήσεις και αποτελέσματα αναλύσεων. Γίνεται δηλαδή προσπάθεια να συμπυκνωθεί μια μεγάλη ποσότητα πληροφοριών σε λίγα σχήματα ή αριθμούς. Η Στατιστική αποτελεί έναν κλάδο εφαρμοσμένων μαθηματικών, που ασχολείται με τη συλλογή και την ανάλυση των δεδομένων, την ερμηνεία και την εύληπτη παρουσίαση των αποτελεσμάτων, καθώς και με την εξαγωγή συμπερασμάτων για τον πληθυσμό. Στην κατεύθυνση αυτή, η Στατιστική συμβάλλει στη λήψη αποφάσεων, μετά από την επεξεργασία και την ανάλυση διάφορων αριθμητικών δεδομένων, οι οποίες δεν μπορούν να εξαχθούν με μια πρώτη ματιά, λόγω του μεγάλου αριθμού των δεδομένων.

Το πρώτο στάδιο για τη λήψη μιας οποιασδήποτε απόφασης αφορά τη συλλογή των δεδομένων και αποτελεί μια πάρα πολύ σημαντική διαδικασία, η οποία επηρεάζει την αξιοπιστία των αποτελεσμάτων. Αν η συλλογή των δεδομένων δεν πραγματοποιηθεί με τον σωστό τρόπο, τότε οποιαδήποτε απόφαση ληφθεί θα είναι βασισμένη σε λανθασμένα στοιχεία και κατά συνέπεια θα είναι και αυτή λανθασμένη. Η συλλογή των δεδομένων μπορεί να διαχωριστεί σε δύο μεγάλες κατηγορίες: α) σε μετρήσεις που πραγματοποιούνται σε ολόκληρο τον πληθυσμό και β) σε μετρήσεις που πραγματοποιούνται σε ένα μέρος μόνο του πληθυσμού, το οποίο αποκαλείται *δείγμα*.

Με τον όρο *πληθυσμός* δεν εννοούμε μόνο έναν δημογραφικό πληθυσμό, αλλά οποιοδήποτε συγκεκριμένο πεδίο, τα στοιχεία του οποίου μπορούν να καταμετρηθούν. Καθένα από αυτά τα στοιχεία ή άτομα ονομάζεται *μονάδες του πληθυσμού, υποκείμενα ή περιπτώσεις*. Συνήθως το μέγεθος του πληθυσμού, δηλαδή ο αριθμός των μελών του, συμβολίζεται με (n), ενώ το μέγεθος του δείγματος με (N). Για παράδειγμα, το σύνολο των Ελλήνων εκπαιδευτικών πρωτοβάθμιας εκπαίδευσης αποτελεί έναν πληθυσμό, ενώ οι εκπαιδευτικοί αυτοί αποτελούν μονάδες του συγκεκριμένου πληθυσμού. *Μονάδα ανάλυσης ή δειγματοληπτική μονάδα* (sampling unit) είναι το άτομο (στην περίπτωση αυτή ο/η εκπαιδευτικός) που παρέχει την πληροφορία για ανάλυση. Μπορεί να είναι μεμονωμένα άτομα, ομάδες, χώρες, πόλεις κ.λπ. Τα χαρακτηριστικά των μονάδων του πληθυσμού τα οποία μπορούν να καταμετρηθούν αποκαλούνται *μεταβλητές*. Τέτοια χαρακτηριστικά (μεταβλητές) μπορεί να είναι οτιδήποτε μετρήσιμο, όπως για παράδειγμα το βάρος, το ύψος, η ικανότητα απομνημόνευσης, η επίδοση στο μάθημα των μαθηματικών, η καταγωγή, το φύλο κ.λπ.

Όταν έχουμε στη διάθεσή μας ολόκληρο τον πληθυσμό, μπορούμε να απογράψουμε, δηλαδή να μετρήσουμε κάποια συγκεκριμένα χαρακτηριστικά (μεταβλητές) σε όλες τις μονάδες του πληθυσμού. Παρ' όλα αυτά, για να μετρηθεί ένας ολόκληρος πληθυσμός απαιτούνται συνήθως μεγάλο κόστος και μεγάλος αριθμός εξειδικευμένων ατόμων που θα πραγματοποιήσουν την απογραφή. Όταν δεν έχουμε στη διάθεσή μας ολόκληρο τον πληθυσμό, αλλά μόνο ένα μέρος του, το λεγόμενο *δείγμα*, τότε χρησιμοποιούνται *επαγωγικές μέθοδοι* ανάλυσης για την εξαγωγή συμπερασμάτων που αφορούν τον πληθυσμό. Για παράδειγμα, αν από το

σύνολο των Ελλήνων εκπαιδευτικών, που αποτελεί έναν πληθυσμό, επιλέξουμε 1000 εκπαιδευτικούς, αυτοί θα αποτελούν ένα δείγμα του πληθυσμού.

Υπάρχουν διάφορες μέθοδοι για την επιλογή ενός δείγματος από έναν πληθυσμό, οι οποίες ονομάζονται *δειγματοληπτικές μέθοδοι* και χωρίζονται σε *πιθανοτικές* και *μη πιθανοτικές* μεθόδους (Bryman, 2017). Σε αυτό το σημείο θα πρέπει να τονίσουμε ότι για να μπορούν να αναχθούν τα όποια αποτελέσματα, που θα προκύψουν από την ανάλυση των δεδομένων του δείγματος, στον αντίστοιχο πληθυσμό θα πρέπει το δείγμα να είναι *αντιπροσωπευτικό*. *Αντιπροσωπευτικό* είναι το δείγμα που επιλέχθηκε από πιθανοτική μέθοδο δειγματοληψίας. Στο αντιπροσωπευτικό δείγμα θεωρητικά θα πρέπει να περιλαμβάνονται υποκείμενα του πληθυσμού τα οποία έχουν την ίδια πιθανότητα να βρεθούν στο εν λόγω δείγμα. Για μια αναλυτική παρουσίαση των πιθανοτικών μεθόδων δειγματοληψίας μπορείτε να ανατρέξετε στο βιβλίο του Γιαλαμά (2005).

Όσον αφορά τώρα τα δεδομένα, είναι απαραίτητο να διακρίνεται η προέλευσή τους. Με λίγα λόγια, θα πρέπει να γίνεται σαφές αν προέρχονται από ολόκληρο τον πληθυσμό ή από ένα δείγμα του. Κάθε πληροφορία σχετική με κάποιο χαρακτηριστικό του πληθυσμού, όπως για παράδειγμα η μέση τιμή του, ονομάζεται *παράμετρος του πληθυσμού*. Το πληροφοριακό στοιχείο που περιγράφει το δείγμα ονομάζεται *στατιστικό* ή *στατιστικός δείκτης*. Έτσι, η μέση τιμή του δείγματος στο παράδειγμά μας αποτελεί ένα στατιστικό. Τυπικά κάθε πληθυσμιακή παράμετρος έχει το αντίστοιχο δειγματικό στατιστικό της. Στις περισσότερες περιπτώσεις, ο τρόπος υπολογισμού μιας παραμέτρου δεν διαφέρει καθόλου ή διαφέρει ελάχιστα από τον υπολογισμό του αντίστοιχου στατιστικού. Τα σύμβολα που χρησιμοποιούνται για την παράμετρο και το στατιστικό είναι διαφορετικά ούτως ώστε η διάκριση, όταν αυτά χρησιμοποιούνται, να είναι σαφής.

1.1.1 Βασικές στατιστικές έννοιες

Παρακάτω παρουσιάζονται συνοπτικά οι βασικές στατιστικές έννοιες. Ωστόσο, για μια πιο αναλυτική περιγραφή των εννοιών αυτών, ο αναγνώστης θα μπορούσε να μελετήσει το βιβλίο του Γιαλαμά (2005).

Μέτρηση ονομάζεται κάθε εμπειρική διαδικασία που καθορίζει την αντιστοίχιση χαρακτηριστικών υποκειμένων με αριθμούς ή σύμβολα μέσα στο πλαίσιο κανόνων. *Μεταβλητή* είναι κάθε χαρακτηριστικό ενός υποκειμένου της έρευνας που μας ενδιαφέρει να μετρήσουμε και το οποίο παίρνει διαφορετικές τιμές. Για παράδειγμα, το βάρος των συμμετεχόντων, η ηλικία τους, η επίδοσή τους σε μια δοκιμασία, η στάση τους απέναντι στους μετανάστες κ.λπ. αποτελούν πιθανές περιπτώσεις μεταβλητών μιας έρευνας. Για κάθε τύπο μέτρησης έχουν αναπτυχθεί διάφορες *κλίμακες* (scales), δηλαδή ένα σύνολο κανόνων στους οποίους θα βασιστούμε για να αντιστοιχίσουμε χαρακτηριστικά που μετράμε με σύμβολα ή αριθμούς. Η διάκριση μεταξύ κλιμάκων είναι σημαντική, επειδή ορισμένες στατιστικές επεξεργασίες εξαρτώνται από τον τύπο της κλίμακας με βάση τον οποίο μετρήθηκε η συγκεκριμένη μεταβλητή.

Οι κλίμακες μέτρησης που γενικότερα χρησιμοποιούνται για όλες τις μετρήσεις είναι οι εξής τέσσερις:

1. Κατηγορική ή ονομαστική κλίμακα (nominal scale). Οι τιμές (κατηγορίες) του χαρακτηριστικού που μετράμε απλώς διακρίνονται η μία από την άλλη, π.χ. το φύλο, η κατεύθυνση σπουδών, το επάγγελμα κ.λπ. Μια ερευνήτρια η οποία μελετά την αλληλεπίδραση μεταξύ μαθητή και εκπαιδευτικού κατά τη διεξαγωγή μιας δραστηριότητας μπορεί να υποδείξει ότι η αλληλεπίδραση αυτή μπορεί να είναι «ικανοποιητική» ή «μη ικανοποιητική», χρησιμοποιώντας με αυτόν τον τρόπο μια ονομαστική κλίμακα δύο τιμών.
2. Ιεραρχική κλίμακα ή τακτική ή διάταξης (ordinal scale). Οι τιμές όχι μόνο διακρίνονται αλλά και ιεραρχούνται. Παραδείγματα μεταβλητών που μετριοούνται με βάση αυτήν την κλίμακα είναι το επίπεδο εκπαίδευσης, η ακαδημαϊκή ή η στρατιωτική ιεραρχία. Σε αυτόν τον τύπο κλίμακας συγκαταλέγεται και η κλίμακα τύπου Likert, όπου οι απαντήσεις εκφράζουν το μέγεθος συμφωνίας ή διαφωνίας με μια ορισμένη δήλωση: («συμφωνώ απόλυτα», «συμφωνώ», «είμαι αβέβαιος ή δεν έχω άποψη», «διαφωνώ», «διαφωνώ απόλυτα»). Ωστόσο, θα πρέπει να σημειώσουμε ότι η κλίμακα αυτή δεν μας δίνει πληροφορίες για τη διαφορά που υπάρχει ανάμεσα στις θέσεις κατάταξης, για παράδειγμα στη μεταβλητή «σειρά κατάταξης των αθλητών σε έναν αγώνα δρόμου».
3. Ισοδιαστημική κλίμακα (interval scale). Οι τιμές διακρίνονται, ιεραρχούνται και ταυτόχρονα τα διαστήματα μεταξύ των διαδοχικών τιμών είναι ίδια. Ωστόσο δεν υπάρχει η έννοια του απόλυτου μηδέν, αφού το μηδέν είναι αυθαίρετο και δεν σημαίνει απουσία της υπό μέτρησης ιδιότητας. Επίσης, δεν ισχύουν σχέσεις αναλογίας. Για παράδειγμα, όσον αφορά την έννοια του απόλυτου μηδέν, όταν αναφερόμαστε σε θερμοκρασία 0 βαθμών Κελσίου δεν υποδεικνύουμε ότι δεν

υπάρχει θερμοκρασία. Αντίστοιχα, για τις σχέσεις αναλογίας, όταν λέμε ότι ένας μαθητής έχει πάρει σε ένα τεστ ευφυΐας 150 Intelligence Quotient (IQ) δεν σημαίνει ότι διαθέτει τριπλάσια μόνο ευφυΐα από κάποιον που έχει 50 IQ.

4. Αναλογική κλίμακα ή κλίμακα λόγου ή πηλίκου (ratio scale). Στην κλίμακα αυτή μπορούμε να έχουμε και το απόλυτο μηδέν ως πιθανή τιμή μιας μέτρησης αλλά και σχέσεις αναλογίας ($y/x=a$) με νόημα. Πιο συγκεκριμένα, ο λόγος δύο τιμών μίας μεταβλητής έχει νόημα σε αυτήν την κλίμακα, αν για παράδειγμα σκεφτούμε ότι το ύψος 160 cm είναι διπλάσιο από το ύψος 80 cm. Επίσης, θα μπορούσε μια πιθανή τιμή να είναι και το μηδέν. Για παράδειγμα, αν μετρούσαμε την απόσταση από κάποιο στόχο θα μπορούσαμε να πούμε ότι η απόσταση είναι μηδενική. Τέλος, η κλίμακα αυτή μας δίνει τις περισσότερες πληροφορίες καθώς περιλαμβάνει τις ιδιότητες όλων των προηγούμενων κλιμάκων: διαφορά, διάταξη, σταθερά διαστήματα και αναλογία.

Οι μεταβλητές ανάλογα με τον τρόπο με τον οποίο μεταβάλλονται κατανέμονται σε ποσοτικές ή ποιοτικές μεταβλητές. Ποσοτικές μεταβλητές (quantitative variables) ονομάζονται εκείνες που μεταβάλλονται από την άποψη της ποσότητας και διακρίνονται σε συνεχείς και ασυνεχείς μεταβλητές. Οι συνεχείς μεταβλητές (continuous variables) είναι δυνατό να παίρνουν οποιαδήποτε τιμή μεταξύ δύο ακραίων τιμών μιας δεδομένης κλίμακας, π.χ. ύψος, χρόνος, βάρος κ.λπ. Στις ασυνεχείς (discontinuous) ή διακριτές μεταβλητές (discrete variables), οι μετρήσεις είναι διακριτοί αριθμοί, για παράδειγμα, η αριθμητική σειρά κατάταξης, οι ακέραιες τιμές βαθμολογίας, τα άτομα στην οικογένεια κ.λπ. Κατά την οποιαδήποτε μέτρηση μιας ποσοτικής μεταβλητής αυτό που τελικά προσδιορίζουμε είναι διακριτοί αριθμοί. Η μέτρηση δηλαδή είναι μια προσέγγιση ή στρογγυλοποίηση της πραγματικής. Επίσης, κατά τη στατιστική ανάλυση δεν γίνεται συνήθως διάκριση μεταξύ των δύο αυτών τύπων ποσοτικών μεταβλητών. Από την άλλη πλευρά, ποιοτικές μεταβλητές (qualitative variables) ονομάζονται εκείνες που μεταβάλλονται από την άποψη της ποιότητας. Όλες οι ποιοτικές μεταβλητές είναι διακριτές και δεν μεταβάλλονται σε βαθμό, ποσότητα ή μέγεθος, αλλά είναι ποιοτικά διαφορετικές. Παραδείγματα ποιοτικών μεταβλητών αποτελούν το φύλο, το θρήσκευμα, το επάγγελμα αλλά και τα είδη των μεθόδων διδασκαλίας, π.χ. δασκαλοκεντρική, μαθητοκεντρική διδασκαλία κ.λπ.

Τέλος, υπάρχουν και κάποιοι άλλοι τύποι μεταβλητών ανάλογα με τον ρόλο τους στην έρευνα ή τις συνθήκες μέτρησης:

- Ανεξάρτητες μεταβλητές (independent variables): Είναι οι μεταβλητές εκείνες που επιδέχονται χειρισμό από τον ερευνητή. Οι τιμές της ανεξάρτητης μεταβλητής δεν εξαρτώνται από τις αντίστοιχες τιμές των άλλων μεταβλητών.
- Εξαρτημένες μεταβλητές (dependent variables): Είναι οι μεταβλητές εκείνες, που εξαρτώνται από τις μεταβολές των τιμών της ανεξάρτητης μεταβλητής. Στην ουσία, οι εξαρτημένες μεταβλητές είναι αυτές που ενδιαφέρουν τον ερευνητή.
- Ελεγχόμενες μεταβλητές (control variables): Ονομάζονται οι μεταβλητές τις οποίες αδρανοποιούμε (ελέγχουμε) κατά τη διάρκεια του πειράματος έτσι ώστε να μην επηρεάσουν τη σχέση που υπάρχει μεταξύ της ανεξάρτητης και της εξαρτημένης μεταβλητής. Η επίδραση της ελεγχόμενης μεταβλητής συνήθως ελέγχεται α) με την ομοιογένεια των ομάδων ως προς το χαρακτηριστικό που θέλουμε να ελέγξουμε και β) με την τυχαία δειγματοληψία.

Για παράδειγμα, οι μαθήτριες ηλικίας 10 ετών έχουν σημαντικά υψηλότερη επίδοση στη γλώσσα από ότι οι μαθητές συνομήλικοί τους. Στο παράδειγμα αυτό, στο οποίο οι συμμετέχοντες θα πρέπει να έχουν επιλεγεί τυχαία, το φύλο των μαθητών είναι η ανεξάρτητη μεταβλητή, η επίδοση στη γλώσσα είναι η εξαρτημένη μεταβλητή και η «σταθερή» ηλικία των μαθητών (10 ετών) είναι η ελεγχόμενη μεταβλητή. Άλλο παράδειγμα στο οποίο ακολουθήθηκε ένας πραγματικός πειραματικός σχεδιασμός είναι το εξής: η πειραματική ομάδα (επιλέχθηκε τυχαία) των μαθητών, που διδάχθηκε μαθηματικά με την υποστήριξη των ΤΠΕ, είχε σημαντικά υψηλότερη επίδοση στο μετα-τεστ από την ομάδα ελέγχου (επιλέχθηκε τυχαία) των μαθητών ίδιας ηλικιακής ομάδας, που διδάχθηκε μαθηματικά χωρίς την υποστήριξη των ΤΠΕ. Στο παράδειγμα αυτό η μεταβλητή μεταχείρισης, δηλαδή η ομάδα των μαθητών (ομάδα ελέγχου, ομάδα πειραματική), είναι η ανεξάρτητη μεταβλητή, η επίδοση στα μαθηματικά (όπως αυτή προσδιορίζεται από το μετα-τεστ) είναι η εξαρτημένη μεταβλητή και η ηλικία των μαθητών, το κοινό διδακτικό περιεχόμενο, ο κοινός δάσκαλος κ.λπ. είναι οι ελεγχόμενες μεταβλητές.

Αξίζει να τονίσουμε ότι ένα ζεύγος σχετιζόμενων μεταβλητών δεν σημαίνει ότι αναφέρεται κατ' ανάγκη σε μια ανεξάρτητη και μια εξαρτημένη μεταβλητή. Για να μπορέσουμε να υποστηρίξουμε ότι η μια μεταβλητή

είναι ανεξάρτητη και η άλλη εξαρτημένη, θα πρέπει ακολουθώντας αυστηρά συνήθως μια πειραματική έρευνα να μπορούμε να δείξουμε τα εξής: α) Η ανεξάρτητη να προηγείται χρονικά ή να ορίζεται αυστηρά κατά τον σχεδιασμό της έρευνας, β) να διαφαίνεται σημαντική σχέση μεταξύ των δύο μεταβλητών (ανεξάρτητης και εξαρτημένης) και γ) να μην υπάρχει άλλη μεταβλητή (ελεγχόμενη) που να συναγωνίζεται με την ανεξάρτητη στην εξήγηση της εξαρτημένης.

1.1.2 Στάδια διεξαγωγής μιας ποσοτικής έρευνας

Η διεξαγωγή μιας ποσοτικής έρευνας είναι μια πολύπλοκη διαδικασία, η οποία αποτελείται από αρκετά στάδια. Ας δούμε τι περιλαμβάνει το καθένα από αυτά τα στάδια:

- Αρχικά μελετάται η προϋπάρχουσα έρευνα στο πεδίο και ορίζονται οι στόχοι της έρευνας. Οι στόχοι αυτοί έχουν ως σκοπό τη διερεύνηση ή/και την επαλήθευση της προϋπάρχουσας θεωρίας του συγκεκριμένου επιστημονικού πεδίου. Αξίζει να τονίσουμε ότι στην ποσοτική έρευνα η διατύπωση των στόχων βασίζεται στον προσδιορισμό του ερευνητικού προβλήματος όπως αυτό καθορίζεται από τα ευρήματα της προηγούμενης έρευνας. Επομένως η συστηματική ανασκόπηση της βιβλιογραφίας θα υποστηρίξει τη διατύπωση των επιδιώξεων (ερευνητικών στόχων) της έρευνας και ταυτόχρονα τη διατύπωση επιμέρους ερευνητικών ερωτημάτων τα οποία εξειδικεύουν τους ερευνητικούς στόχους.
- Έπειτα, επιλέγεται με τυχαίο τρόπο ένα αντιπροσωπευτικό δείγμα από τον συγκεκριμένο πληθυσμό, τον οποίο θέλουμε να διερευνήσουμε.
- Καθορίζεται το μέσο συλλογής των δεδομένων της έρευνας. Συνήθως, τα δεδομένα συλλέγονται με τη συμπλήρωση ερωτηματολογίου από τους συμμετέχοντες της έρευνας, αν και μπορούν επίσης να συγκεντρωθούν με τη συμπλήρωση ερωτηματολογίου μέσω προσωπικής συνέντευξης ή ακόμα και μέσω φυσικής παρατήρησης.
- Αφού συλλεχθούν τα δεδομένα, πρέπει να καταχωριστούν στη βάση (φόρμα) δεδομένων ενός λογισμικού, π.χ. του SPSS. Στη συνέχεια, θα διεξαχθεί η στατιστική επεξεργασία των δεδομένων με τη χρήση του λογισμικού για τυχόν εύρεση σφαλμάτων κατά τη διάρκεια της καταγραφής των δεδομένων ή της καταχώρισής τους.
- Εξετάζεται η ανάγκη για τον μετασχηματισμό των δεδομένων και για την τροποποίηση ή την κατασκευή καινούριων μεταβλητών. Οι μετασχηματισμοί αυτοί αποτελούν συχνό φαινόμενο στη στατιστική ανάλυση, προκειμένου να επιτευχθεί η συμβατότητα των μεταβλητών με κάθε είδος επεξεργασίας (Γναρδέλλης, 2013).
- Σε αυτό το στάδιο αναλύονται στατιστικά τα δεδομένα αξιοποιώντας την περιγραφική και την επαγωγική στατιστική αλλά και την πολυμεταβλητή ανάλυση.
- Το τελικό στάδιο περιλαμβάνει τη συγγραφή της έκθεσης-σύνοψης των αποτελεσμάτων της έρευνας, στην οποία ερμηνεύονται κατάλληλα τα αποτελέσματα και γενικεύονται τα συμπεράσματα.

Όπως γίνεται κατανοητό, η χρήση του SPSS δεν αφορά μόνο τη στατιστική ανάλυση των δεδομένων αλλά και την εκτέλεση και άλλων σημαντικών διεργασιών, όπως αυτές που περιγράφονται παραπάνω, στα στάδια 4 και 5.

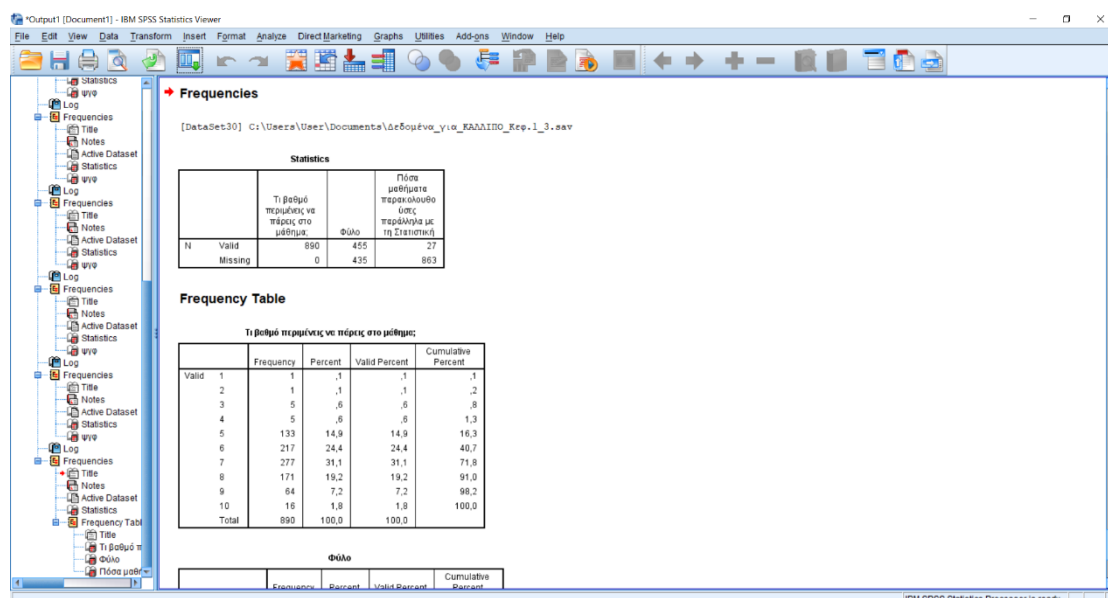
1.2 Εισαγωγή στο SPSS

Το **SPSS** (Statistical Package for the Social Sciences – Στατιστικό Πακέτο για τις Κοινωνικές Επιστήμες) αναπτύχθηκε για πρώτη φορά το 1965 από μία ομάδα φοιτητών (Norman Nie, Dale Ben & Hadlai Hull, 1965) του Πανεπιστημίου Stanford του San Francisco των ΗΠΑ. Η πρώτη έκδοση του λογισμικού (SPSS 1) κυκλοφόρησε το 1968. Το SPSS παραγόταν αρχικά από την εταιρεία SPSS Inc., αλλά αποκτήθηκε από την IBM το 2009. Οι τρέχουσες εκδόσεις (μετά το 2015) έχουν το εμπορικό σήμα «IBM SPSS Statistics». Πρόκειται για το πιο δημοφιλές στατιστικό πακέτο με ευρύτατη χρήση από όλους τους ερευνητικούς οργανισμούς και τα πανεπιστήμια, κυρίως στο πεδίο των κοινωνικών επιστημών (Norris et al., 2017).

Το SPSS αποτελεί ένα ισχυρό εργαλείο, το οποίο παρέχει πολλές δυνατότητες επεξεργασίας και ανάλυσης δεδομένων, που καλύπτουν σε μεγάλο βαθμό τις απαιτήσεις των ερευνητών (Field, 2013, Norris et al., 2017). Ο χρήστης εργάζεται σε ένα φιλικό και εύχρηστο περιβάλλον διεπαφής, οργανωμένο σε παράθυρα, που είναι πλήρως συμβατό με περιβάλλον Windows.

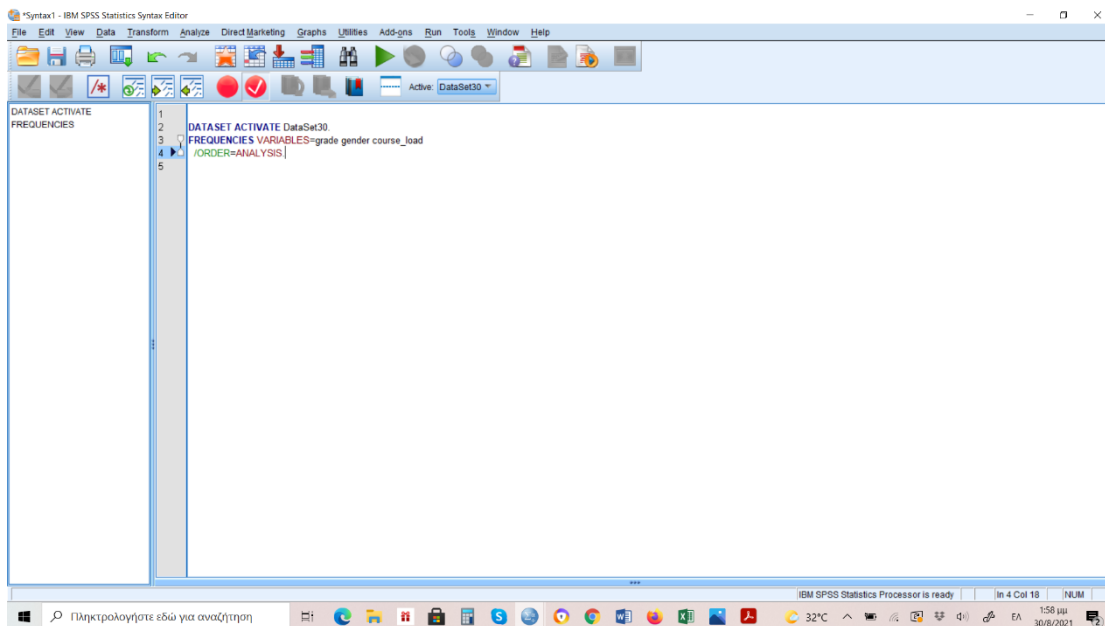
1.2.1 Βασικά παράθυρα του SPSS

Τα βασικά παράθυρα του SPSS είναι ο *Data Editor*, ο *Output Viewer* και ο *Syntax Editor*. Ο Data Editor είναι το πρώτο παράθυρο που ανοίγει όταν ξεκινάει το SPSS. Πρόκειται για έναν επεξεργαστή δεδομένων με τη χρήση του οποίου γίνεται η καταχώριση ή η διαχείριση των δεδομένων. Στον Data Editor αποθηκεύουμε ένα αρχείο δεδομένων του SPSS. Τα αρχεία δεδομένων του SPSS έχουν προέκταση (*.sav). Ένα σημαντικό πλεονέκτημα του Data Editor είναι ότι μπορεί να διαβάσει αρχεία δεδομένων άλλων λογισμικών, όπως για παράδειγμα το Excel. Στον Output Viewer παρουσιάζονται οι πίνακες με τα αποτελέσματα των αναλύσεων, καθώς και τα γραφήματα (Εικόνα 1.1). Πιο συγκεκριμένα, εκεί γίνεται η αποθήκευση των αποτελεσμάτων σε αρχεία αποτελεσμάτων του SPSS, τα οποία έχουν προέκταση (*.spv). Σημαντικό πλεονέκτημα του Output Viewer είναι ότι μπορούμε με τη διαδικασία της αντιγραφής και της επικόλλησης να μεταφέρουμε τα αποτελέσματα των αναλύσεων σε άλλα λογισμικά, όπως το Word και το PowerPoint.



Εικόνα 1.1 Ενδεικτικό παράθυρο του «Output Viewer».

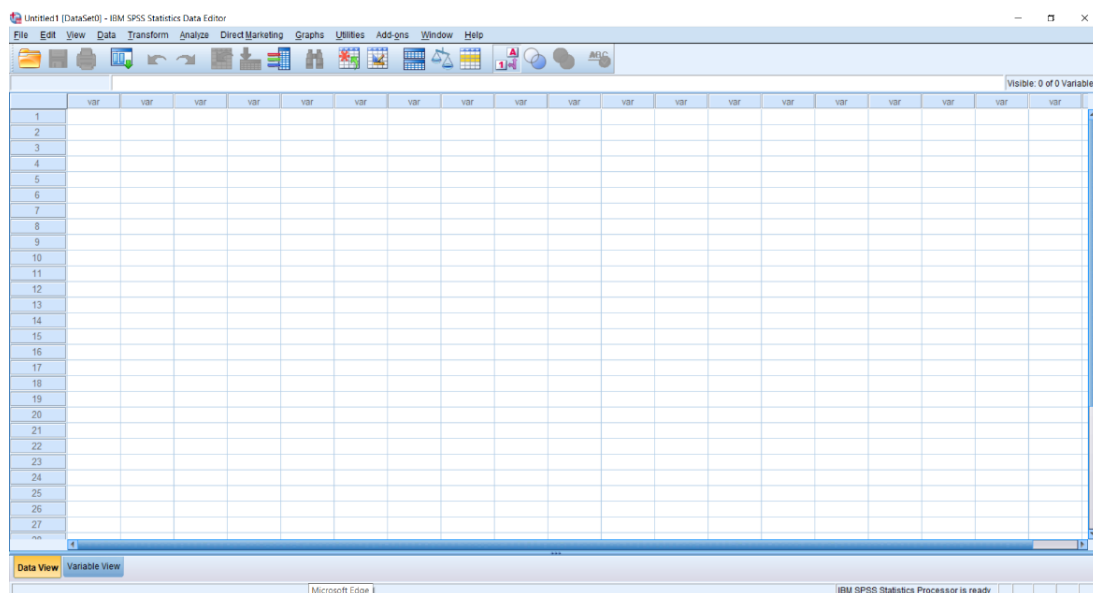
Στο παράθυρο του Syntax Editor μπορούμε να συντάξουμε εντολές στη γλώσσα εντολών που παρέχει το SPSS, η οποία ονομάζεται *command syntax* (Εικόνα 1.2). Στο συγκεκριμένο παράθυρο έχουμε τη δυνατότητα να επικολλήσουμε σε μορφή εντολών τις διεργασίες που εκτελέσαμε κατά τη διάρκεια χρήσης των μενού του SPSS, πατώντας την εντολή **Paste**. Οι εντολές του Syntax Editor αποθηκεύονται σε αρχεία του SPSS που έχουν προέκταση (*.sps). Τα αρχεία αυτά είναι πολύ χρήσιμα, ιδιαίτερα σε διαδικασίες του SPSS που επαναλαμβάνονται αρκετές φορές, γιατί οι εντολές τους μπορούν να εκτελεστούν άμεσα με το πάτημα του κουμπιού **Run Selection** στο παράθυρο του Syntax Editor.



Εικόνα 1.2 Ενδεικτικό παράθυρο του «Syntax Editor».

1.2.2 Μενού του SPSS

Αφού εκκινήσουμε το IBM SPSS Statistics 26 από το μενού **Έναρξη** ή από την Επιφάνεια Εργασίας, κάνουμε κλικ στο κουμπί **Cancel** στο πρώτο παράθυρο και στη συνέχεια εμφανίζεται ο Data Editor του SPSS, ο οποίος είναι κενός, όπως φαίνεται στην Εικόνα 1.3.



Εικόνα 1.3 Παράθυρο του «Data Editor».

Η γραμμή των μενού του SPSS περιλαμβάνει τα εξής μενού: **File**, **Edit**, **View**, **Data**, **Transform**, **Analyze**, **Direct Marketing**, **Graphs**, **Utilities**, **Add-ons**, **Window** και **Help**.

Από το μενού **File** διαχειριζόμαστε διάφορους τύπους αρχείων του SPSS (άνοιγμα, κλείσιμο, αποθήκευση). Από το μενού **Edit** διεξάγουμε την αντιγραφή, την αποκοπή και την επικόλληση των δεδομένων καθώς και την εισαγωγή μεταβλητών και περιπτώσεων. Μέσω του μενού **View** μπορούμε να ρυθμίσουμε την εμφάνιση του Data Editor. Στο μενού **Data** διαχειριζόμαστε τα δεδομένα ενός ή και περισσότερων αρχείων μας. Μπορούμε, για παράδειγμα, να ταξινομήσουμε τα δεδομένα μας ή να διαχωρίσουμε το αρχείο σε ξεχωριστές υποομάδες δεδομένων. Από το μενού **Transform** μετασχηματίζουμε τις τιμές των μεταβλητών του αρχείου

δεδομένων, υπολογίζουμε τις τιμές νέων μεταβλητών, αντικαθιστούμε τις ελλείπουσες τιμές με άλλες μέσω κάποιας συγκεκριμένης μεθόδου κ.τ.λ. Το μενού [Analyze](#) είναι το σημαντικότερο μενού του SPSS. Περιέχει συγκεντρωτικά όλες τις στατιστικές αναλύσεις που μπορούμε να πραγματοποιήσουμε, από έναν απλό υπολογισμό συχνοτήτων έως την εκτέλεση πολύπλοκων πολυμεταβλητών αναλύσεων. Το μενού [Direct Marketing](#) περιλαμβάνει μια σειρά τεχνικών που επιτρέπουν την πραγματοποίηση αναλύσεων των πελατών ή των επαφών του ερευνητή ώστε να βελτιωθούν οι εκστρατείες διαφήμισης. Το μενού [Graphs](#) παρέχει τη δυνατότητα κατασκευής μιας μεγάλης ποικιλίας γραφημάτων, τα οποία μπορούν να τροποποιηθούν ώστε να βελτιωθεί η παρουσίασή τους. Το μενού [Utilities](#) περιέχει συμπληρωματικές επιλογές για την περιγραφή των μεταβλητών, την ενσωμάτωση σχολίων όταν αποθηκεύουμε ένα αρχείο δεδομένων κ.ά. Το μενού [Add-ons](#) παρέχει πρόσβαση σε άλλες εφαρμογές και υπηρεσίες της IBM SPSS Inc., που δεν συμπεριλαμβάνονται στην τρέχουσα έκδοση του λογισμικού. Στο μενού [Window](#) εμφανίζονται τα ανοικτά παράθυρα του SPSS. Παρέχει επίσης τη δυνατότητα της διαίρεσης του τρέχοντος παραθύρου σε πολλαπλά παράθυρα, με σκοπό να βλέπουμε ταυτόχρονα διαφορετικά τμήματα ενός μεγάλου όγκου δεδομένων σε ένα αρχείο. Τέλος, το μενού [Help](#) παρέχει on-line πρόσβαση στο βασικό σύστημα βοήθειας του SPSS για όλα τα θέματα για τα οποία χρειάζεται αναλυτικές οδηγίες ο χρήστης.

1.2.3 Καρτέλα Data View

Η εισαγωγή και η επεξεργασία των δεδομένων γίνονται στον Data Editor. Ο τρόπος με τον οποίο εισάγουμε τα δεδομένα είναι αντίστοιχος με αυτόν ενός υπολογιστικού φύλλου του Excel. Πρόκειται δηλαδή για τη δημιουργία ενός πίνακα, η κάθε γραμμή του οποίου αντιστοιχεί σε μια περίπτωση ενός συνόλου δεδομένων. Για παράδειγμα, ένας φοιτητής που συμπληρώνει ένα ερωτηματολόγιο μιας έρευνας αποτελεί μια περίπτωση. Οι γραμμές λοιπόν του Data Editor ονομάζονται περιπτώσεις (cases). Κάθε στήλη του πίνακα αντιστοιχεί σε μία μεταβλητή.

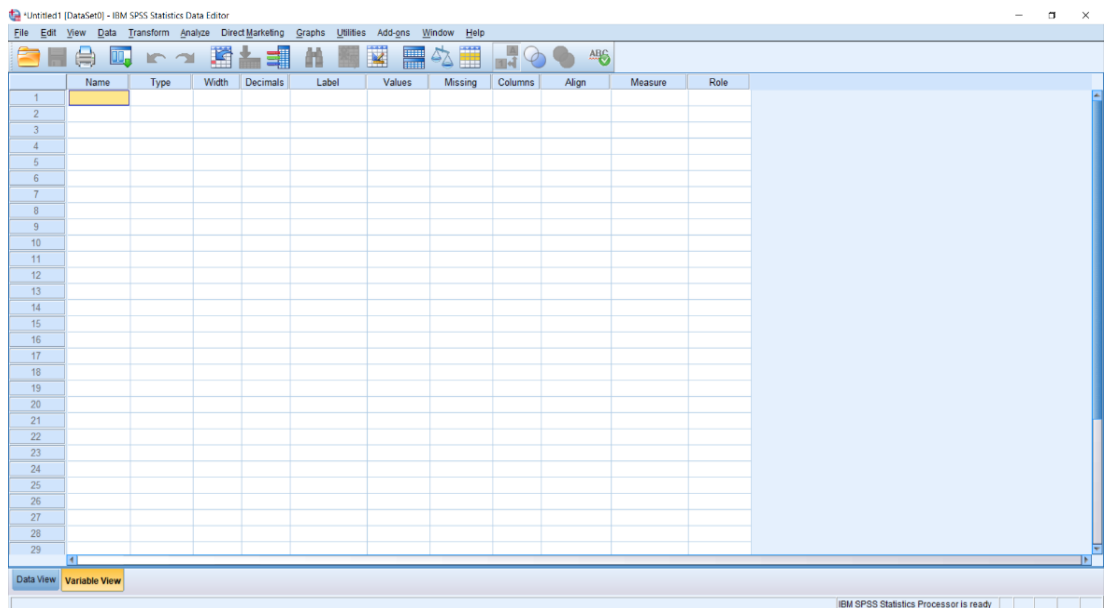
Όπως μπορούμε να δούμε παραπάνω, στην Εικόνα 1.3, στην κάτω αριστερή γωνία του παραθύρου του Data Editor υπάρχουν δύο καρτέλες: η καρτέλα Data View και η καρτέλα Variable View. Η επιφάνεια εργασίας της καρτέλας Data View που είναι επιλεγμένη, όπως βλέπουμε στην παραπάνω εικόνα, είναι η καρτέλα στην οποία πληκτρολογούμε τις τιμές των δεδομένων μας. Επομένως, ο πίνακας της καρτέλας Data View αποτελείται από τόσες γραμμές όσες είναι οι παρατηρήσεις/συμμετέχοντες της έρευνας και από τόσες στήλες όσες είναι οι μεταβλητές.

1.2.4 Καρτέλα Variable View

Όταν επιλέξουμε από τον Data Editor την καρτέλα Variable View, εμφανίζεται το παράθυρο της Εικόνας 1.4. Εδώ, αναγράφονται οι ιδιότητες των μεταβλητών με την εξής μορφή: οι στήλες του Data Editor αποτελούν τις ιδιότητες των μεταβλητών και οι γραμμές αποτελούν τις μεταβλητές. Στην καρτέλα Variable View διαχειριζόμαστε τις μεταβλητές του αρχείου δεδομένων μας. Μπορούμε να ορίσουμε ή να τροποποιήσουμε τις ιδιότητες των μεταβλητών, καθώς επίσης και να προσθέσουμε νέες μεταβλητές ή να διαγράψουμε τις ήδη υπάρχουσες.

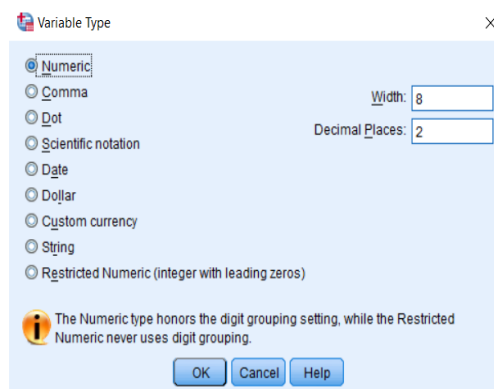
Οι ιδιότητες των μεταβλητών που μπορούμε να ορίσουμε ή να τροποποιήσουμε είναι οι εξής:

- **Name** (Όνομα): Μπορούμε να ορίσουμε οποιοδήποτε όνομα για μια μεταβλητή είτε με ελληνικούς είτε με αγγλικούς χαρακτήρες. Υπάρχουν όμως κάποιοι περιορισμοί στην ονοματολογία των μεταβλητών: το όνομα κάθε μεταβλητής πρέπει να ξεκινά με ένα μικρό ή με ένα κεφαλαίο γράμμα ελληνικού ή αγγλικού αλφαβήτου, να μην περιέχει κενά, να περιλαμβάνει μέχρι 64 χαρακτήρες και να μην περιέχει ειδικούς χαρακτήρες, όπως για παράδειγμα (!, *, / κ.λπ.).



Εικόνα 1.4 Καρτέλα «Variable View» του Data Editor.

- **Type** (Τύπος): Κάνοντας κλικ στο κελί **Type**, εμφανίζεται το αντίστοιχο πλαίσιο διαλόγου (Εικόνα 1.5). Το SPSS επιλέγει εξ' ορισμού οι μεταβλητές να είναι αριθμητικές (**Numeric**), δηλαδή μεταβλητές που περιέχουν μόνο αριθμούς, με ακρίβεια 2 δεκαδικών ψηφίων (**Decimal Places**) και με συνολικό πλάτος (**Width**) 8 θέσεων .



Εικόνα 1.5 Πλαίσιο διαλόγου «Variable Type».

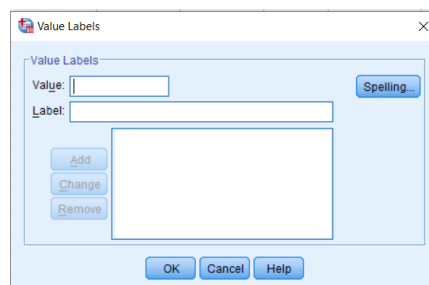
Οι υπόλοιποι τύποι μεταβλητών που εμφανίζονται στο πλαίσιο διαλόγου Variable Type εκτός από τις Numeric που είδαμε παραπάνω, όπως φαίνονται και στην Εικόνα 1.5, είναι οι εξής:

- **Comma**: Είναι οι αριθμητικές μεταβλητές που χρησιμοποιούν το κόμμα ανά τρεις θέσεις και την τελεία ως υποδιαστολή. Για παράδειγμα, η μεταβλητή 1,972,144.48.
- **Dot**: Είναι οι αριθμητικές μεταβλητές που χρησιμοποιούν την τελεία ανά τρεις θέσεις και το κόμμα ως υποδιαστολή. Για παράδειγμα, η μεταβλητή 1.972.144,48.
- **Scientific notation**: Οι τιμές των μεταβλητών αυτών γράφονται με εκθετική μορφή δυνάμεων του 10. Για παράδειγμα, η έκφραση με τη μορφή 1,44E +2 δηλώνει τον αριθμό 1,44 X 10², δηλαδή τον αριθμό 144.
- **Date**: Οι μεταβλητές τύπου ημερομηνίας είναι αριθμητικές μεταβλητές, οι οποίες εμφανίζονται με ημερολογιακή ή και ωρολογιακή μορφή.
- **Dollar**: Οι μεταβλητές αυτές περιέχουν τιμές με ή χωρίς το σύμβολο του δολαρίου (\$). Χρησιμοποιούν το κόμμα ανά τρεις θέσεις και την τελεία ως υποδιαστολή. Για παράδειγμα, η μεταβλητή 2,500.23.

- **Custom currency:** Οι μεταβλητές αυτές εκφράζονται με τον συμβολισμό του νομίσματος που έχει ορίσει ο ερευνητής.
- **String:** Πρόκειται για μη αριθμητικές μεταβλητές που λαμβάνουν χαρακτήρες οποιασδήποτε μορφής. Για παράδειγμα, μια διεύθυνση ή ένα ονοματεπώνυμο. Οι συγκεκριμένες μεταβλητές δεν επιδέχονται πράξεις.
- **Restricted Numeric:** Οι μεταβλητές αυτές είναι κατάλληλες όταν θέλουμε να δώσουμε μη αρνητικές ακέραιες τιμές.

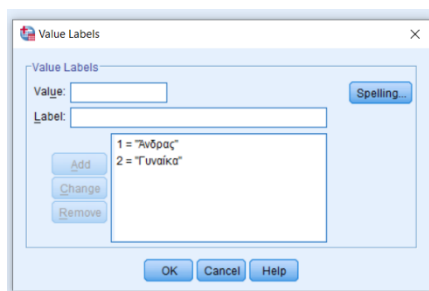
Έχοντας δει τους τύπους μεταβλητών όπως αυτοί εμφανίζονται στο πλαίσιο διαλόγου του Variable Type (Εικόνα 1.5), ας επιστρέψουμε στις ιδιότητες των μεταβλητών που μπορούμε να ορίσουμε ή να τροποποιήσουμε στην καρτέλα Variable View του Data Editor:

- **Width (Πλάτος):** Ο μέγιστος αριθμός χαρακτήρων που μπορεί να λάβει μια μεταβλητή.
- **Decimals (Δεκαδικά ψηφία):** Το μέγιστο πλήθος δεκαδικών ψηφίων που μπορεί να λάβει η μεταβλητή.
- **Label (Ετικέτα):** Λεκτικές περιγραφές σχετικές με το περιεχόμενο των μεταβλητών. Λαμβάνουν μέχρι 256 χαρακτήρες.
- **Values (Ετικέτες τιμών):** Στη στήλη Values κωδικοποιούμε τις κατηγορίες των κατηγορικών μεταβλητών σε τιμές, δηλαδή αντιστοιχούμε έναν αριθμό σε καθεμία κατηγορία της κατηγορικής μεταβλητής. Οι αριθμοί αυτοί είναι απλώς κωδικοί και έχουν μαθηματική αξία. Για παράδειγμα, η μεταβλητή «φύλο» κωδικοποιείται δίνοντας 1 για τους άνδρες και 2 για τις γυναίκες. Η διαδικασία αυτή γίνεται ως εξής: Επιλέγοντας ένα κελί στη στήλη **Values** εμφανίζεται ένα κουμπί με τρεις τελείες. Κάνοντας κλικ σε αυτό το κουμπί εμφανίζεται το πλαίσιο διαλόγου **Value Labels** (Εικόνα 1.6).



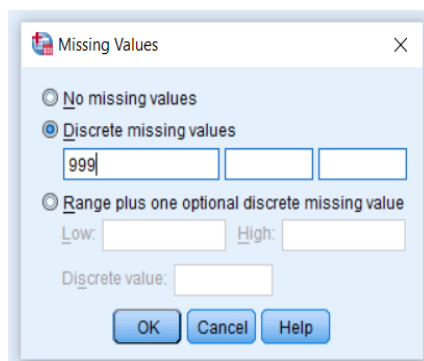
Εικόνα 1.6 Πλαίσιο διαλόγου «Value Labels».

Στο πεδίο **Value** πληκτρολογούμε την τιμή που αντιστοιχεί στην κατηγορία. Στο πεδίο **Label** πληκτρολογούμε την ετικέτα της κατηγορίας. Μόλις ολοκληρωθεί αυτή η διαδικασία, πατάμε **Add** και εισάγεται η τιμή με την κατηγορία που αντιστοιχεί σε αυτή. Με τον ίδιο ακριβώς τρόπο πληκτρολογούμε όλες τις κατηγορίες, μαζί με τις αντίστοιχες τιμές τους. Ας δούμε ένα παράδειγμα της διαδικασίας για τη μεταβλητή «φύλο». Στο πεδίο **Value** πληκτρολογούμε την τιμή 1 και στο πεδίο **Label** πληκτρολογούμε Άνδρας και έπειτα πατάμε **Add**. Στη συνέχεια, στο πεδίο **Value** πληκτρολογούμε την τιμή 2 και στο πεδίο **Label** πληκτρολογούμε Γυναίκα και έπειτα πατάμε **Add**. Η εικόνα που θα έχουμε φαίνεται παρακάτω στην Εικόνα 1.7. Όταν ολοκληρωθεί η διαδικασία, πατάμε το **OK**. Μπορούμε να αλλάξουμε την τιμή (Value) ή την ετικέτα (Label) ή και να τις διαγράψουμε επιλέγοντάς τες, μία κάθε φορά και πατώντας τα κουμπιά **Change** και **Remove** αντίστοιχα.



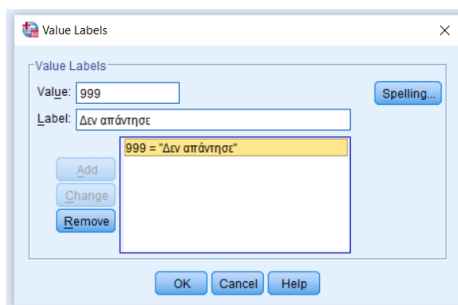
Εικόνα 1.7 Πλαίσιο διαλόγου «Value Labels» της μεταβλητής «Φύλο».

- **Missing** (Ελλείπουσες τιμές): Στη στήλη Missing ορίζουμε τον τρόπο με τον οποίο δηλώνονται στο SPSS οι ελλείπουσες τιμές. Αν ο ερευνητής ξέχασε να καταχωρίσει κάποιες τιμές μιας ποσοτικής μεταβλητής, το SPSS θεωρεί τις τιμές αυτές ως ελλείπουσες και εισάγει αυτόματα στα αντίστοιχα κενά κελιά μια τελεία (.). Οι τιμές αυτές θεωρούνται *system missing*. Επιπλέον, κατά την εισαγωγή δεδομένων στο SPSS ο ερευνητής έχει τη δυνατότητα να ορίσει δικές του τιμές ως ελλείπουσες. Συνήθως ορίζει ως ελλείπουσες τιμές αυτές που δεν θέλει να συμπεριληφθούν στην ανάλυσή του. Τέτοιες τιμές είναι αυτές που θεωρεί λανθασμένες ή πολύ ακραίες. Ο ερευνητής συνήθως υποδεικνύει αυτές τις τιμές με τον αριθμό (999). Η τιμή αυτή (999 για παράδειγμα) θα δηλωθεί στη στήλη Missing (Εικόνα 1.8) και θα καταχωριστεί από τον ερευνητή ώστε να αντικαταστήσει την άλλη. Το SPSS σε μια ανάλυση, όπως για παράδειγμα σε έναν πίνακα συχνοτήτων, θα εμφανίζει ξεχωριστά στην ένδειξη Missing τη συχνότητα που αντιστοιχεί στη δηλωμένη από τον χρήστη ελλείπουσα τιμή. Στην ένδειξη **Missing System** θα εμφανιστεί η συχνότητα των κενών κελιών. Μπορούμε να προσθέσουμε μια ετικέτα που θα αντιστοιχεί στην τιμή της ελλείπουσας (999), στο πλαίσιο διαλόγου Value Labels. Στο πεδίο **Value** πληκτρολογούμε την τιμή 999 και στο πεδίο **Label** πληκτρολογούμε «Δεν απάντησε» (Εικόνα 1.9). Μόλις ολοκληρωθεί αυτή η διαδικασία, επιλέγουμε **Add**.



Εικόνα 1.8 Ορισμός τιμής «999» ως ελλείπουσας τιμής.

- **Columns** (Στήλες): Στη στήλη αυτή ορίζουμε το πλάτος της στήλης κάθε μεταβλητής στην καρτέλα Data View.
- **Align** (Στοίχιση): Στη στήλη αυτή ορίζουμε τον τρόπο στοίχισης των δεδομένων κάθε μεταβλητής: (Left-Αριστερά) - (Right-Δεξιά) - (Center-στο κέντρο).



Εικόνα 1.9 Πρόσθεση ετικέτας που αντιστοιχεί στην τιμή της ελλείπουσας.

- **Measure** (Κλίμακα μέτρησης): Σε αυτή τη στήλη ορίζεται η κλίμακα μέτρησης κάθε μεταβλητής. Υπάρχουν τρεις επιλογές: **Scale** για τις ποσοτικές μεταβλητές, **Ordinal** για τις κατηγορικές με διάταξη μεταβλητές και **Nominal** για τις κατηγορικές μεταβλητές.
- **Role** (Ρόλος): Ορισμένα πλαίσια διαλόγου του SPSS υποστηρίζουν την επιλογή μεταβλητών σύμφωνα με τον ρόλο τους και τις θεωρούν ως προεπιλεγμένες. Κάνοντας κλικ σε ένα κελί της στήλης Role, εμφανίζονται οι εξής επιλογές: **Input** (Μεταβλητή – Είσοδος): Η επιλογή αυτή χρησιμοποιείται κυρίως για ανεξάρτητες μεταβλητές. Εδώ είναι άξιο να σημειωθεί ότι οι μεταβλητές διαδραματίζουν εξ ορισμού τον ρόλο Input, **Target** (Μεταβλητή – Στόχος): Η επιλογή αυτή χρησιμοποιείται κυρίως για εξαρτημένες μεταβλητές, **Both** (Μεταβλητή ως είσοδος και στόχος), **None** (Κανένας ρόλος), **Partition**: Η μεταβλητή αυτή χρησιμοποιείται για να διαχωρίσει τα δεδομένα σε ξεχωριστά δείγματα και **Split**: Η μεταβλητή αυτή χρησιμοποιείται για την κατασκευή μοντέλων για κάθε πιθανή τιμή της μεταβλητής.

Πρέπει να επισημανθεί ότι το SPSS παρέχει τη δυνατότητα τροποποίησης όλων των ιδιοτήτων των μεταβλητών. Κάνοντας δεξί κλικ στη γραμμή αρίθμησης των μεταβλητών ή από το μενού **Edit**, μπορούμε να αντιγράψουμε (**Copy**), να επικολλήσουμε (**Paste**), να διαγράψουμε (**Clear**) και να εισάγουμε μια νέα μεταβλητή (**Insert Variable**). Στην τελευταία περίπτωση, η ήδη υπάρχουσα μεταβλητή μετατοπίζεται μία θέση προς τα κάτω και στη θέση της εισάγεται η νέα μεταβλητή. Η καινούρια αυτή μεταβλητή θα λάβει τις εξ ορισμού ιδιότητες των μεταβλητών.

Επιπλέον, μπορούμε να αντιγράψουμε μία προς μία ή όλες μαζί τις ιδιότητες των μεταβλητών που έχουμε ήδη ορίσει, σε άλλες μεταβλητές. Κάνουμε δεξί κλικ στα κελιά ιδιοτήτων της μεταβλητής, αν θέλουμε να επιλέξουμε ξεχωριστά κάποια χαρακτηριστικά ή κάνουμε δεξί κλικ στη γραμμή αρίθμησης των μεταβλητών, αν θέλουμε να επιλέξουμε όλες μαζί τις ιδιότητες. Στη συνέχεια, πατάμε **Copy**, επιλέγουμε τα κελιά ή τις γραμμές των μεταβλητών στις οποίες θέλουμε να αντιγράψουμε τις ιδιότητες και έπειτα πατάμε **Paste**.

1.2.5 Εισαγωγή δεδομένων στο SPSS

Όπως αναφέρθηκε και προηγουμένως, η εισαγωγή και η επεξεργασία των δεδομένων στο SPSS γίνονται στην καρτέλα Data View του Data Editor. Η εισαγωγή των δεδομένων μπορεί να γίνει είτε πληκτρολογώντας κατά γραμμές (περιπτώσεις) είτε πληκτρολογώντας κατά στήλες (μεταβλητές). Είναι προτιμότερο να ορίσουμε πρώτα τις ιδιότητες των μεταβλητών της έρευνάς μας στην καρτέλα Variable View και έπειτα να καταχωρίσουμε τα δεδομένα στην καρτέλα Data View. Με αυτόν τον τρόπο καταχωρίζουμε πιο σωστά τα δεδομένα μας, αφού βλέπουμε απευθείας στο Data View τις κωδικοποιήσεις των μεταβλητών μας.

Υπάρχουν κάποιες βασικές αρχές για την εισαγωγή των δεδομένων:

- Για να καταχωριστεί μια τιμή πρέπει να πατηθεί το πλήκτρο **Enter** ή να επιλεγθεί κάποιο άλλο κελί.
- Η καταχώριση αριθμητικών δεδομένων πραγματοποιείται με πληκτρολόγηση των τιμών τους στα κελιά της καρτέλας Data View.
- Η καταχώριση δεδομένων στις κατηγορικές και στις διατάξιμες μεταβλητές γίνεται ως εξής: Όταν το κουμπί **Value Labels** είναι πατημένο (Εικόνα 1.10), τότε, κάνοντας κλικ στο κελί που θέλουμε να καταχωρίσουμε την κατηγορία της επιλογής μας εμφανίζεται ένα βελάκι προς τα κάτω. Πατώντας το βελάκι αυτό, εμφανίζονται σε μία λίστα οι κατηγορίες της μεταβλητής που

έχουμε ορίσει. Από εκεί επιλέγουμε την επιθυμητή κατηγορία χωρίς να την πληκτρολογήσουμε. Εναλλακτικά, αν το κουμπί **Value Labels** δεν είναι πατημένο, τότε μπορούμε να πληκτρολογήσουμε την τιμή που αντιστοιχεί στην επιθυμητή κατηγορία.



Εικόνα 1.10 Πατημένο κουμπί «Value Labels».

- Με τις διαδικασίες **Cut-Paste** και **Copy-Paste** μπορούμε να μεταφέρουμε και να αντιγράψουμε αντίστοιχα τις τιμές των αριθμητικών δεδομένων, καθώς και τις κατηγορίες των κατηγορικών/διατάξιμων μεταβλητών σε άλλα κελιά. Σε αυτή τη στήλη ορίζεται η κλίμακα μέτρησης κάθε μεταβλητής.

Οι διαδικασίες του ορισμού των ιδιοτήτων των μεταβλητών και της εισαγωγής δεδομένων θα περιγραφούν με τη χρήση του ερωτηματολογίου **SATS** (Survey of Attitudes Toward Statistics), με ένα μικρό μέρος δεδομένων της έρευνας των Λαβίδα, Μπαρκάτσα, Μάνεση και Γιαλαμά (2020). Το ερωτηματολόγιο SATS είναι ένα διεθνές εργαλείο μέτρησης των στάσεων απέναντι στη Στατιστική. Ένα απόσπασμα του ερωτηματολογίου αυτού παρουσιάζεται παρακάτω.

Ερωτηματολόγιο SATS

Έρευνα για τη Διερεύνηση των στάσεων απέναντι στη Στατιστική
Δημογραφικά χαρακτηριστικά

✓ Τι βαθμό περιμένεις να πάρεις στο μάθημα; _____

✓ Φύλο 1. Άνδρας 2. Γυναίκα

✓ Πόσα μαθήματα παρακολουθούσες παράλληλα με τη Στατιστική; _____

✓ Πόσο τακτικά παρακολουθούσες τις διαλέξεις του μαθήματος της Στατιστικής;

Σχεδόν Πάντα

Συχνά

Μερικές Φορές

Καθόλου

✓ Κρατούσες επιλεκτικές σημειώσεις από τις διαλέξεις του μαθήματος;

Ναι

Όχι

✓ Μελετούσες προσεκτικά τις σημειώσεις ή τα βιβλία του μαθήματος;

Ναι

Όχι

Δηλώσεις στάσεων

1. Διαφωνώ έντονα, ..., 4. Ούτε διαφωνώ ούτε συμφωνώ, ..., 7. Συμφωνώ έντονα							
Μου αρέσει η στατιστική.	1	2	3	4	5	6	7
Αισθάνομαι ανασφάλεια όταν βρίσκομαι αντιμέτωπος με προβλήματα στατιστικής.	1	2	3	4	5	6	7
Έχω προβλήματα στην κατανόηση της στατιστικής εξαιτίας του τρόπου σκέψης μου.	1	2	3	4	5	6	7
Οι τύποι της στατιστικής είναι εύκολο να κατανοηθούν.	1	2	3	4	5	6	7
Η στατιστική είναι άνευ αξίας.	1	2	3	4	5	6	7
Η στατιστική είναι σύνθετο αντικείμενο.	1	2	3	4	5	6	7
Η στατιστική θα έπρεπε να είναι αναπόσπαστο μέρος της επαγγελματικής μου εκπαίδευσης.	1	2	3	4	5	6	7
Οι στατιστικές ικανότητες θα με καταστήσουν πιο περιζήτητο στον χώρο εργασίας.	1	2	3	4	5	6	7
Δεν γνωρίζω τι συμβαίνει με τη στατιστική.	1	2	3	4	5	6	7
Η στατιστική δεν είναι χρήσιμη σε ένα συνηθισμένο επάγγελμα.	1	2	3	4	5	6	7

Βρίσκομαι σε απόγνωση με τη διαδικασία των τεστ στατιστικής στην τάξη.	1	2	3	4	5	6	7
Ο στατιστικός τρόπος σκέψης δεν μου είναι χρήσιμος στην υπόλοιπη ζωή μου πέραν της εργασίας.	1	2	3	4	5	6	7
Χρησιμοποιώ τη στατιστική στην καθημερινή μου ζωή.	1	2	3	4	5	6	7
Νιώθω άγχος κατά τη διάρκεια του μαθήματος της στατιστικής.	1	2	3	4	5	6	7
Με ευχαριστεί να παρακολουθώ μαθήματα στατιστικής.	1	2	3	4	5	6	7

Κάθε δήλωση ή ερώτηση ενός ερωτηματολογίου συνιστά στοιχείο παρατήρησης (ή μέτρησης) και επομένως μία μεταβλητή. Επομένως, πρέπει να οριστούν στο SPSS οι ιδιότητες των μεταβλητών και στη συνέχεια να καταχωριστούν τα δεδομένα-απαντήσεις των συμμετεχόντων της έρευνας. Σύμφωνα με τα είδη των μεταβλητών που έχουν αναφερθεί στο κεφάλαιο, η πρώτη και η τρίτη ερώτηση του ερωτηματολογίου «τι βαθμό περιμένεις να πάρεις στο μάθημα;» και «πόσα μαθήματα παρακολουθούσες παράλληλα με τη Στατιστική;» αντίστοιχα, αποτελούν δύο ποσοτικές μεταβλητές. Επομένως, απλώς θα πληκτρολογήσουμε τις απαντήσεις των συμμετεχόντων της έρευνας στα αντίστοιχα κελιά στην καρτέλα Data View. Οι υπόλοιπες ερωτήσεις των δημογραφικών χαρακτηριστικών αποτελούν κατηγορικές μεταβλητές. Αυτό σημαίνει ότι καθεμία κατηγορία κάθε κατηγορικής μεταβλητής πρέπει να κωδικοποιηθεί με μία αριθμητική τιμή. Οι αριθμητικές τιμές αυτές είναι τυχαίες και είναι απλώς κωδικοί, δηλαδή δεν έχουν μαθηματική αξία. Για παράδειγμα, η μεταβλητή φύλο έχει δύο κατηγορίες, μπορεί να κωδικοποιηθεί δίνοντας την τιμή 1 για τους άνδρες και την τιμή 2 για τις γυναίκες. Αυτό δεν σημαίνει ότι οι γυναίκες είναι διπλάσιες από τους άνδρες. Το ίδιο ισχύει και για τις μεταβλητές που παίρνουν δύο πιθανές τιμές-απαντήσεις, όπως για παράδειγμα οι μεταβλητές του τύπου Ναι/Όχι.

Τις κατηγορίες της μεταβλητής που αντιστοιχεί στην ερώτηση «πόσο τακτικά παρακολουθούσες τις διαλέξεις του μαθήματος της Στατιστικής;» μπορούμε να τις κωδικοποιήσουμε με δύο τρόπους:

- 1^{ος} τρόπος. Θα ορίσουμε τις κωδικές ονομασίες των κατηγοριών της μεταβλητής με τη σειρά που τις έχουμε ορίσει στο ερωτηματολόγιο. Δηλαδή, στην κατηγορία «Σχεδόν Πάντα» θα αντιστοιχίσουμε την τιμή 1. Στην κατηγορία «Συχνά» θα αντιστοιχίσουμε την τιμή 2. Στην κατηγορία «Μερικές φορές» θα αντιστοιχίσουμε την τιμή 3 και στην κατηγορία «Καθόλου» θα αντιστοιχίσουμε την τιμή 4. Όπως παρατηρούμε, οι χαμηλές τιμές αντιστοιχούν σε συχνή παρακολούθηση των διαλέξεων του μαθήματος.
- 2^{ος} τρόπος. Θα ορίσουμε τις κωδικές ονομασίες των κατηγοριών της μεταβλητής με αντίστροφη σειρά από αυτήν που έχουμε ορίσει στο ερωτηματολόγιο. Με αυτόν τον τρόπο οι υψηλές τιμές αντιστοιχούν σε συχνή παρακολούθηση των διαλέξεων του μαθήματος. Ακολουθώντας αυτήν την κωδικοποίηση, είναι πιο δύσκολο να κάνουμε λάθος στην ερμηνεία των αποτελεσμάτων των αναλύσεων που θα πραγματοποιήσουμε. Επομένως, στην κατηγορία «Σχεδόν Πάντα» θα αντιστοιχίσουμε την τιμή 4. Στην κατηγορία «Συχνά» θα αντιστοιχίσουμε την τιμή 3. Στην κατηγορία «Μερικές φορές» θα αντιστοιχίσουμε την τιμή 2 και τέλος στην κατηγορία «Καθόλου» θα αντιστοιχίσουμε την τιμή 1.

Οι μεταβλητές που αντιστοιχούν στις δηλώσεις των στάσεων απέναντι στη Στατιστική είναι διατάξιμες. Πρόκειται για ερωτήσεις διαφωνίας/συμφωνίας μιας επτά βαθμών κλίμακας τύπου Likert, όπου το 1 αντιστοιχεί στο «διαφωνώ έντονα» και το 7 αντιστοιχεί στο «συμφωνώ έντονα». Το κέντρο της κλίμακας έχει την τιμή 4 και αντιστοιχεί στην ουδέτερη στάση, δηλαδή στο «ούτε διαφωνώ ούτε συμφωνώ». Επομένως, θα κωδικοποιήσουμε τις κατηγορίες αυτών των μεταβλητών ακολουθώντας τη συγκεκριμένη σειρά.

Σημαντικό είναι να προσέξουμε την κατεύθυνση της στάσης απέναντι στις μεταβλητές τις οποίες μετράμε. Για παράδειγμα, στην πρώτη δήλωση «μου αρέσει η στατιστική» η τιμή 1 («διαφωνώ έντονα») εκφράζει αρνητική στάση, ενώ η τιμή 7 («συμφωνώ έντονα») εκφράζει θετική στάση. Όμως, στη δήλωση «η στατιστική είναι άνευ αξίας» η τιμή 1 («διαφωνώ έντονα») εκφράζει θετική στάση, ενώ η τιμή 7 («συμφωνώ

έντονα») εκφράζει αρνητική στάση. Αυτό δεν επηρεάζει την κωδικοποίηση των κατηγοριών, όμως θα πρέπει αργότερα να επανακωδικοποιήσουμε κάποιες μεταβλητές ώστε σε όλες οι ελάχιστες τιμές να αντιστοιχούν σε αρνητικές στάσεις και οι μέγιστες τιμές να αντιστοιχούν σε θετικές στάσεις ή το αντίστροφο. Λεπτομέρειες για αυτή τη διαδικασία θα αναφερθούν σε επόμενο κεφάλαιο.

Τις ιδιότητες για όλες τις μεταβλητές τις ορίζουμε σύμφωνα με τις διαδικασίες που περιγράφηκαν παραπάνω. Η μορφή της καρτέλας Variable View, όταν ολοκληρωθεί η διεξαγωγή των ενεργειών, είναι η παρακάτω (Εικόνα 1.11).

Name	Type	Width	Decimals	Label	Values	Missing	Columns	Align	Measure	Role
1 grade	Numeric	8	0	Τι βαθμό περιμένεις να πάρεις στο μάθημα	None	None	12	Right	Nominal	Input
2 gender	Numeric	8	0	Φύλο	(1, Άνδρας)	None	8	Right	Nominal	Input
3 course_load	Numeric	8	0	Πόσα μαθήματα παρακολουθείς παράλληλα με τη Στατιστική	None	None	8	Right	Scale	Input
4 attendance	Numeric	8	0	Πόσα τακτικά παρακολουθείς τις διαλέξεις του μαθήματος της Στατιστικής	(1, Σχεδόν Πάντα)	None	11	Right	Scale	Input
5 keep_notes	Numeric	8	0	Κρατάεις επιλεκτικές σημειώσεις από τις διαλέξεις του μαθήματος	(1, Ναι)	None	8	Right	Scale	Input
6 read_notes	Numeric	8	0	Μελετάεις προσεκτικά τις σημειώσεις ή τα βιβλία του μαθήματος	(1, Ναι)	None	8	Right	Scale	Input
7 SATS1	Numeric	8	0	Μου αρέσει η στατιστική	(1, Διαφωνώ έντονα)	None	6	Right	Nominal	Input
8 SATS2	Numeric	8	0	Αισθάνομαι ανασφάλεια όταν βρίσκομαι αντιμέτωπος με προβλήματα στατιστικής	(1, Διαφωνώ έντονα)	None	6	Right	Nominal	Input
9 SATS3	Numeric	8	0	Έχω προβλήματα στην κατανόηση της στατιστικής εξάφης του τρόπου σκέψης μου	(1, Διαφωνώ έντονα)	None	7	Right	Nominal	Input
10 SATS4	Numeric	8	0	Οι τύποι της στατιστικής είναι εύκολο να κατανοηθούν	(1, Διαφωνώ έντονα)	None	7	Right	Nominal	Input
11 SATS5	Numeric	8	0	Η στατιστική είναι άδεια αξίας	(1, Διαφωνώ έντονα)	None	7	Right	Nominal	Input
12 SATS6	Numeric	8	0	Η στατιστική είναι σύνθετο αντικείμενο	(1, Διαφωνώ έντονα)	None	7	Right	Nominal	Input
13 SATS7	Numeric	8	0	Η στατιστική θα έπρεπε να είναι αναπόσπαστο μέρος της επαγγελματικής μου εκπαίδευσης	(1, Διαφωνώ έντονα)	None	7	Right	Nominal	Input
14 SATS8	Numeric	8	0	Οι στατιστικές κανόνες θα με καταστήσουν πιο περιζήτητο στο χώρο εργασίας	(1, Διαφωνώ έντονα)	None	7	Right	Nominal	Input
15 SATS9	Numeric	8	0	Δεν γνωρίζω τι συμβαίνει με τη στατιστική	(1, Διαφωνώ έντονα)	None	8	Right	Nominal	Input
16 SATS10	Numeric	8	0	Η στατιστική δεν είναι χρήσιμη σε ένα συντηγμένο επάγγελμα	(1, Διαφωνώ έντονα)	None	9	Right	Nominal	Input
17 SATS11	Numeric	8	0	Βρίσκω σε απόγνωση με τη διαδικασία των τεστ στατιστικής στη τάξη	(1, Διαφωνώ έντονα)	None	12	Right	Nominal	Input
18 SATS12	Numeric	8	0	Ο στατιστικός τρόπος σκέψης δεν μου είναι χρήσιμος στην υπόλοιπη ζωή μου πέρα της εργασίας	(1, Διαφωνώ έντονα)	None	12	Right	Nominal	Input
19 SATS13	Numeric	8	0	Χρησιμοποιώ τη στατιστική στην καθημερινή μου ζωή	(1, Διαφωνώ έντονα)	None	12	Right	Nominal	Input
20 SATS14	Numeric	8	0	Νοιάω άσχετα κατά τη διάρκεια του μαθήματος της στατιστικής	(1, Διαφωνώ έντονα)	None	12	Right	Nominal	Input
21 SATS15	Numeric	8	0	Με ευχαριστεί να παρακολουθώ μαθήματα στατιστικής	(1, Διαφωνώ έντονα)	None	12	Right	Nominal	Input

Εικόνα 1.11 Μορφή της καρτέλας «Variable View».

Στη συνέχεια, κάνουμε κλικ στην καρτέλα Data View ώστε να καταχωρίσουμε τα δεδομένα. Όπως αναφέρθηκε παραπάνω, όταν θέλουμε να εισαγάγουμε τις κατηγορίες μιας κατηγορικής/διατάξιμης μεταβλητής, επιλέγουμε από την αναδυόμενη λίστα την επιθυμητή κατηγορία χωρίς να την πληκτρολογήσουμε. Για παράδειγμα, η μορφή της αναδυόμενης λίστας για την πρώτη μεταβλητή των στάσεων (SATS1) απ' όπου θα επιλέξουμε την κατηγορία που αντιστοιχεί στην απάντηση του ερωτηματολογίου, είναι η παρακάτω:

1: SATS1	grade	gender	course_load	attendance	keep_notes	read_notes	SATS1	SATS2	SATS3	SATS4	SATS5	SATS6	SATS7	SATS8	SATS9	SATS10	SATS11
1	9	Γυναίκα	6	Καθόλου	Όχι	Όχι	Ευφυνιστικά αρκεί	Συμφωνώ...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Συμφωνώ...	Διαφωνώ...
2	7	Γυναίκα	5	Μερικές φορές	Ναι	Ναι	Διαφωνώ έντονα	Διαφωνώ...	Όχι δια...	Διαφωνώ...	Διαφωνώ...	Συμφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
3	7	Γυναίκα	7	Συχνά	Όχι	Όχι	Διαφωνώ αρκεί	Διαφωνώ...	Όχι δια...	Διαφωνώ...	Διαφωνώ...	Συμφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
4	9	Γυναίκα	7	Σχεδόν Πάντα	Όχι	Ναι	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Διαφωνώ...	Συμφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
5	8	Γυναίκα	7	Σχεδόν Πάντα	Ναι	Όχι	Όχι διαφωνώ ούτε α...	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
6	9	Γυναίκα	7	Συχνά	Όχι	Ναι	Ευφυνιστικά αρκεί	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
7	9	Γυναίκα	6	Σχεδόν Πάντα	Ναι	Ναι	Ευφυνιστικά αρκεί	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
8	6	Γυναίκα	7	Σχεδόν Πάντα	Ναι	Ναι	Ευφυνιστικά αρκεί	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
9	6	Γυναίκα	4	Συχνά	Ναι	Ναι	Συμφωνώ...	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
10	7	Γυναίκα	7	Συχνά	Ναι	Όχι	Συμφωνώ έντονα	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
11	8	Γυναίκα	7	Συχνά	Ναι	Όχι	Συμφωνώ έντονα	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
12	7	Γυναίκα	5	Καθόλου	Όχι	Όχι	Συμφωνώ έντονα	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
13	6	Γυναίκα	7	Μερικές φορές	Όχι	Όχι	Συμφωνώ έντονα	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
14	10	Γυναίκα	6	Καθόλου	Όχι	Όχι	Συμφωνώ έντονα	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
15	8	Γυναίκα	6	Συχνά	Όχι	Όχι	Συμφωνώ αρκεί	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
16	7	Γυναίκα	6	Συχνά	Ναι	Ναι	Συμφωνώ αρκεί	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
17	5	Γυναίκα	7	Συχνά	Όχι	Όχι	Όχι διαφωνώ ούτε α...	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
18	7	Γυναίκα	6	Συχνά	Ναι	Όχι	Συμφωνώ αρκεί	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
19	9	Γυναίκα	7	Σχεδόν Πάντα	Ναι	Ναι	Συμφωνώ έντονα	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
20	7	Γυναίκα	5	Συχνά	Ναι	Ναι	Συμφωνώ έντονα	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
21	8	Γυναίκα	7	Συχνά	Ναι	Ναι	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
22	9	Γυναίκα	7	Μερικές φορές	Ναι	Ναι	Συμφωνώ έντονα	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
23	6	Γυναίκα	6	Σχεδόν Πάντα	Ναι	Όχι	Συμφωνώ έντονα	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
24	8	Γυναίκα	7	Μερικές φορές	Όχι	Όχι	Συμφωνώ έντονα	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
25	7	Γυναίκα	6	Συχνά	Ναι	Ναι	Συμφωνώ αρκεί	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
26	8	Γυναίκα	8	Συχνά	Ναι	Όχι	Συμφωνώ αρκεί	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...
27	6	Γυναίκα	6	Σχεδόν Πάντα	Ναι	Ναι	Συμφωνώ...	Συμφωνώ...	Όχι δια...	Όχι δια...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Όχι δια...	Διαφωνώ...	Συμφωνώ...	Διαφωνώ...

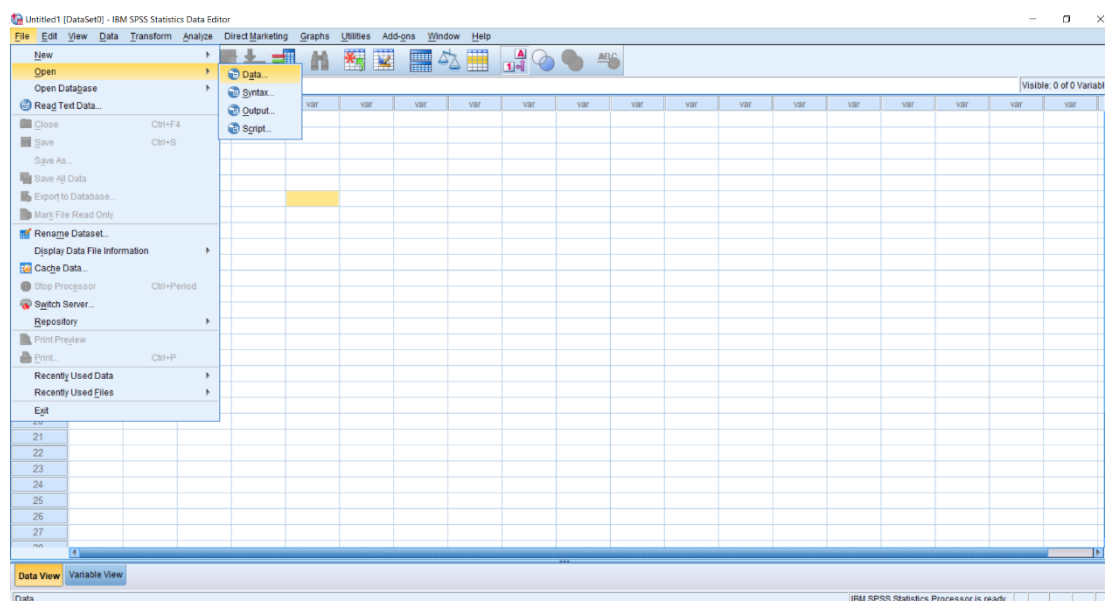
Εικόνα 1.12 Αναδυόμενη λίστα επιλογής κατηγοριών της διατάξιμης μεταβλητής «SATS1».

Στην Εικόνα 1.12 φαίνεται επίσης η μορφή της καρτέλας Data View, όταν ολοκληρωθεί η καταχώριση όλων των δεδομένων της έρευνας.

1.2.6 Μεταφορά δεδομένων στο SPSS από αρχείο Excel

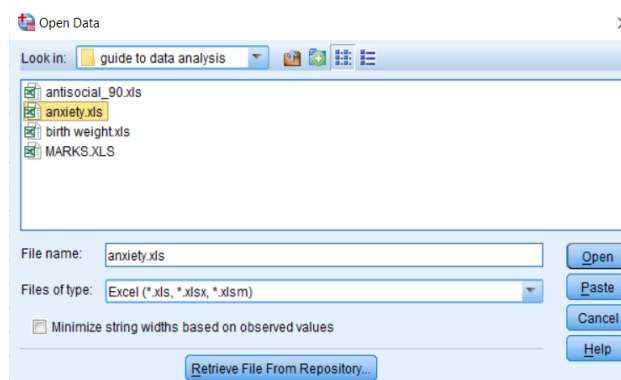
Όπως αναφέρθηκε προηγουμένως, μπορούμε να μεταφέρουμε δεδομένα από ένα αρχείο Excel στο SPSS. Αυτή η διαδικασία είναι πολύ χρήσιμη γιατί αρκετοί ερευνητές περνούν τα δεδομένα της έρευνάς τους σε ένα αρχείο Excel, ειδικά στην περίπτωση που το ερωτηματολόγιο τους συμπληρώνεται σε ηλεκτρονική μορφή, για παράδειγμα μέσω διαδικτυακών εργαλείων συγκέντρωσης δεδομένων, όπως είναι το Google Form.

Για να μεταφέρουμε τα δεδομένα ενός αρχείου Excel στο SPSS, επιλέγουμε από τη γραμμή των μενού τις εντολές **File => Open => Data** (οι επιλογές των μενού στο υπόλοιπο σύγγραμμα θα δίνονται συνοπτικά με τον παραπάνω τρόπο και θα ισοδυναμούν με τις επιλογές που εκτελεί ο χρήστης στο SPSS, οι οποίες για τη συγκεκριμένη περίπτωση δίνονται στο παρακάτω παράθυρο (Εικόνα 1.13)):



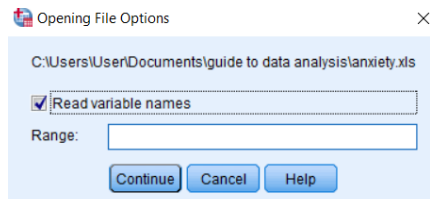
Εικόνα 1.13 Επιλογή **Open => Data** του πλαισίου διαλόγου «File».

Στη συνέχεια, στο πλαίσιο διαλόγου **Look in** του μενού **Open Data** αναζητάμε τη θέση αποθήκευσης του αρχείου δεδομένων. Έπειτα επιλέγουμε τον κατάλληλο τύπο αρχείου Excel, στο πλαίσιο διαλόγου **Files of type** (*.xls, *.xlsx, *.xlsm), όπως φαίνεται στην Εικόνα 1.14. Επιλέγοντας το αρχείο Excel που θέλουμε να ανοίξουμε, εμφανίζεται το όνομά του στο πλαίσιο διαλόγου **File name**. Τέλος, πατάμε το κουμπί **Open**.



Εικόνα 1.14 Πλαίσιο διαλόγου «Open Data» για την επιλογή αρχείου Excel.

Στη συγκεκριμένη περίπτωση, αφού επιλέξουμε το αρχείο (anxiety.xls), εμφανίζεται το πλαίσιο διαλόγου **Opening File Options** (Εικόνα 1.15).



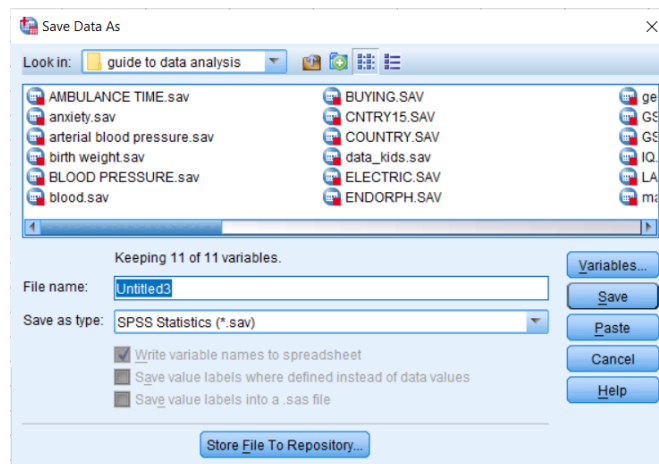
Εικόνα 1.15 Πλαίσιο διαλόγου «Opening File Options».

Στο συγκεκριμένο πλαίσιο, κάνουμε κλικ στο τετραγωνάκι της επιλογής **Read variable names**, για να δηλώσουμε αν η πρώτη γραμμή του αρχείου Excel θα διαβάσει τα ονόματα των μεταβλητών. Αν δεν είναι τσεκαρισμένη η επιλογή αυτή, τότε το SPSS δίνει στις μεταβλητές, με τη σειρά, τα ονόματα V1, V2 κ.ο.κ. Επιπλέον, μπορούμε να ορίσουμε ποια περιοχή κελιών θα διαβαστεί από το SPSS, στο πεδίο **Range**.

Στο SPSS κάθε στήλη στην καρτέλα Data View αντιστοιχεί σε μία μεταβλητή. Όταν ανοίξει στο SPSS ένα αρχείο Excel, δημιουργείται μία μεταβλητή για κάθε στήλη του αρχείου Excel.

1.2.7 Αποθήκευση αρχείων δεδομένων

Όταν ολοκληρωθεί η καταχώριση των δεδομένων, για να αποθηκεύσουμε για πρώτη φορά ένα αρχείο δεδομένων στο SPSS, επιλέγουμε από τη γραμμή των μενού τις εντολές: **File => Save As**. Στη συνέχεια, θα ανοίξει το πλαίσιο διαλόγου **Save Data As** (Εικόνα 1.16). Εδώ, στο πλαίσιο διαλόγου **Look in** θα μας ζητηθεί να ορίσουμε τη θέση αποθήκευσης του αρχείου δεδομένων. Στο πλαίσιο διαλόγου **File name** θα πληκτρολογήσουμε το όνομα του αρχείου και στη λίστα **Save as type** επιλέγουμε τον τύπο του αρχείου. Ο προεπιλεγμένος τύπος αρχείου στο SPSS Statistics είναι ο (*.sav). Αν θέλουμε να αλλάξουμε τον τύπο του προς αποθήκευση αρχείου μας, επιλέγουμε τον επιθυμητό τύπο από τον κατάλογο των διαθέσιμων επιλογών. Τέλος, πατάμε το κουμπί **Save**.



Εικόνα 1.16 Πλαίσιο διαλόγου «Save Data As».

Αν το αρχείο δεδομένων μας είναι ήδη αποθηκευμένο και κάνουμε απλώς αλλαγές στο περιεχόμενό του, τότε επιλέγουμε από τη γραμμή των μενού τις εντολές: **File => Save**. Εναλλακτικά, μπορούμε να πατήσουμε το εικονίδιο της δισκέτας από τη γραμμή εργαλείων της καρτέλας Data View ή της καρτέλας Variable View.

Συστήνεται να γίνονται και ενδιάμεσες αποθηκεύσεις όταν τροποποιούμε το περιεχόμενο του αρχείου δεδομένων μας. Οι ενδιάμεσες αποθηκεύσεις μπορούν να γίνονται είτε πατώντας απλώς το εικονίδιο της δισκέτας είτε αλλάζοντας το όνομα του αρχείου στην περίπτωση που θέλουμε να διατηρήσουμε το περιεχόμενο του αρχικού μας αρχείου.

1.3 Προβλήματα για εξάσκηση και ανατροφοδότηση

Άσκηση 1

Να κατασκευαστεί βάση δεδομένων (μεταβλητές και χαρακτηριστικά αυτών) που συλλέχθηκαν με το επόμενο ερωτηματολόγιο. Οι επόμενες ερωτήσεις είναι κάποιες μόνο από τις ερωτήσεις που περιλαμβάνονται στο ερωτηματολόγιο.

Ερωτηματολόγιο:

1. Ποιο είναι το έτος γέννησής σας;
2. Ποια ήταν η κατεύθυνση των σπουδών σας στο Λύκειο; (1. Θετική, 2. Τεχνολογική, 3. Θεωρητική)
3. Ποιο είναι το αντικείμενο των βασικών σας σπουδών; (προπτυχιακό επίπεδο)
4. Ποιες από τις επόμενες κατηγορίες υπολογιστικών εργαλείων (λογισμικών) χρησιμοποιείτε;

	Καθόλου	Λίγο	Αρκετά	Πολύ	Πάρα Πολύ
Προγράμματα επεξεργασίας κειμένου (όπως το Word)					
Προγράμματα στατιστικής επεξεργασίας δεδομένων (όπως το SPSS)					

5. Ποια από τα επόμενα έχετε διδαχθεί; (μπορείτε να επιλέξετε περισσότερα από ένα):

- Περιγραφική Στατιστική
- Επαγωγική Στατιστική
- Πολυμεταβλητή Ανάλυση

6. Έχετε ποτέ εμπλακεί σε ερευνητική διαδικασία και, αν ναι, ποιο ήταν το αντικείμενο της έρευνας;

Βιβλιογραφία

- Γιαλαμάς, Β. (2005). *Στατιστικές Τεχνικές και Εφαρμογές στις Επιστήμες της Αγωγής*. Αθήνα: Πατάκης.
- Γναρδέλλης, Χ. (2013). *Ανάλυση δεδομένων με το IBM SPSS Statistics 21*. Αθήνα: Εκδ. Παπαζήση.
- Field, A. (2013). *Discovering statistics using IBM SPSS statistics* (5th ed.). SAGE: London.
- Lavidas, K., Barkatsas, T., Manesis, D., & Gialamas, V. (2020). A Structural Equation Model Investigating The Impact Of Tertiary Students' Attitudes Toward Statistics, Perceived Competence At Mathematics, And Engagement On Statistics Performance. *Statistics Education Research Journal*, 19(2).
- Norris, G., Qureshi, F., Howitt, D., & Cramer, D. (2017). *Εισαγωγή στη Στατιστική με το SPSS για τις Κοινωνικές Επιστήμες*. Επιμ. Εμμανουηλίδης, Χ. Αθήνα: Εκδ. Κλειδάριθμος.

Κεφάλαιο 2

Μετασχηματισμός Δεδομένων

Σύνοψη

Το SPSS περιέχει πολλές εντολές για τον μετασχηματισμό των δεδομένων, οι οποίες είναι απαραίτητες για τη συμβατότητα των μεταβλητών με την εκάστοτε μεθοδολογία. Επίσης, διαθέτει επιλογές για τη δημιουργία νέων μεταβλητών από τις ήδη υπάρχουσες, καθώς και για την επανακωδικοποίηση των τιμών μιας μεταβλητής. Στο κεφάλαιο αυτό, περιγράφονται οι εντολές *Compute Variable* (δημιουργία νέων μεταβλητών με την εφαρμογή αριθμητικών εκφράσεων), *Recode into Same Variables* (επανακωδικοποίηση στις ίδιες μεταβλητές), *Recode into Different Variables* (επανακωδικοποίηση σε διαφορετικές μεταβλητές) και *Automatic Recode* (αυτόματη επανακωδικοποίηση).

Προαπαιτούμενη γνώση

Βασικές έννοιες διεξαγωγής μιας ποσοτικής έρευνας και βασικές λειτουργίες του περιβάλλοντος: Κεφάλαιο 1 του συγγράμματος.

2.1 Επέκταση των δεδομένων

Πολλές φορές πριν ή και κατά τη διάρκεια της ανάλυσης των δεδομένων χρειάζεται να προβούμε στη δημιουργία καινούριων μεταβλητών οι οποίες βασίζονται στις υπάρχουσες. Για παράδειγμα, μπορούμε να υπολογίσουμε συνολικά σκορ από διάφορες μεταβλητές, να επανακωδικοποιήσουμε κατηγορικές μεταβλητές κ.λπ. Οι διαδικασίες που οδηγούν στην επέκταση των δεδομένων της έρευνας στο περιβάλλον του SPSS βρίσκονται στο μενού **Transform**. Το SPSS περιέχει πολλές εντολές για τον μετασχηματισμό των δεδομένων. Παρακάτω, περιγράφονται μόνο οι εντολές **Automatic Recode** (αυτόματη επανακωδικοποίηση), **Recode into Same Variables** (επανακωδικοποίηση στις ίδιες μεταβλητές), **Recode into Different Variables** (επανακωδικοποίηση σε διαφορετικές μεταβλητές), και **Compute Variable** (δημιουργία νέων μεταβλητών με την εφαρμογή αριθμητικών εκφράσεων).

2.1.1 Αυτόματη επανακωδικοποίηση – Automatic Recode

Όταν οι κατηγορικές μεταβλητές του αρχείου δεδομένων είναι αλφαριθμητικές (string), δηλαδή οι τιμές τους δεν είναι αριθμητικές, αλλά περιλαμβάνουν μια σειρά χαρακτήρων/συμβόλων ή όταν οι κωδικοί των κατηγοριών δεν είναι διαδοχικοί ακέραιοι αριθμοί, δημιουργούνται διάφορα προβλήματα. Αυτό συμβαίνει, για παράδειγμα, όταν μεταφέρουμε δεδομένα από ένα αρχείο Excel στο SPSS, όπως στην περίπτωση του ηλεκτρονικού ερωτηματολογίου, όπου όλες οι μεταβλητές λαμβάνονται από το SPSS ως αλφαριθμητικές. Υπάρχουν ορισμένες διαδικασίες κατά την ανάλυση δεδομένων, κατά τις οποίες δεν μπορούμε να επεξεργαστούμε αλφαριθμητικές μεταβλητές ή απαιτείται οι κατηγορίες των κατηγορικών μεταβλητών, να έχουν κωδικοποιηθεί σε διαδοχικούς ακέραιους αριθμούς (Norris et al., 2017).

Με την εντολή **Automatic Recode** μπορούμε να επανακωδικοποιήσουμε τις τιμές αριθμητικών ή αλφαριθμητικών μεταβλητών σε διαδοχικούς ακέραιους αριθμούς με αφετηρία το 1. Οι νέες μεταβλητές που δημιουργούνται από την εντολή Automatic Recode διατηρούν τις ετικέτες μεταβλητών και τιμών (variable labels και value labels) των παλιών μεταβλητών. Για οποιεσδήποτε τιμές των αρχικών μεταβλητών χωρίς καθορισμένες ετικέτες χρησιμοποιούνται οι αρχικές τιμές ως ετικέτες για τις επανακωδικοποιημένες τιμές. Όταν εκτελεστεί η εντολή Automatic Recode, στο αρχείο αποτελεσμάτων εμφανίζεται ένας πίνακας που απεικονίζει τις παλιές και τις νέες τιμές, όπως επίσης και τα value labels που αντιστοιχούν στις νέες τιμές.

Οι τιμές των αλφαριθμητικών μεταβλητών επανακωδικοποιούνται με αλφαβητική σειρά, με τα κεφαλαία γράμματα να προηγούνται από τα αντίστοιχα μικρά. Με την επανακωδικοποίηση οι system-missing values παραμένουν οι ίδιες. Οι user-missing values επανακωδικοποιούνται ξανά ως user-missing values, έχοντας ως τιμές τους αμέσως μεγαλύτερους ακέραιους αριθμούς από τη μέγιστη τιμή του καινούριου συστήματος κωδικοποίησης, με τη σειρά τους όμως να διατηρείται.

Έχουμε επίσης τη δυνατότητα να χρησιμοποιήσουμε το ίδιο σχήμα επανακωδικοποίησης για όλες τις μεταβλητές που έχουμε επιλέξει («[Use the same recoding scheme for all variables](#)»). Για παράδειγμα, αν έχουμε επιλέξει η επανακωδικοποίηση της πρώτης μεταβλητής να ξεκινά από τη χαμηλότερη τιμή, τότε η επιλογή αυτή επαναλαμβάνεται αυτόματα για όλες τις επιλεγμένες μεταβλητές.

Με την επιλογή [Treat blank string values as user-missing](#) οι κενές ή οι μηδενικές τιμές των μεταβλητών που θα επανακωδικοποιηθούν μετατρέπονται σε user-missing values. Η επιλογή αυτή θα κωδικοποιήσει αυτόματα τις κενές τιμές, σε τιμές υψηλότερες από τη μέγιστη τιμή που δεν λείπει (έγκυρη τιμή).

Επιπλέον, μπορούμε να αποθηκεύσουμε ένα σχήμα αυτόματης κωδικοποίησης σε ένα αρχείο προτύπου με την επιλογή [Save Template as](#) και να το εφαρμόσουμε σε άλλες μεταβλητές ή αρχεία δεδομένων, επιλέγοντας [Apply Template from](#).

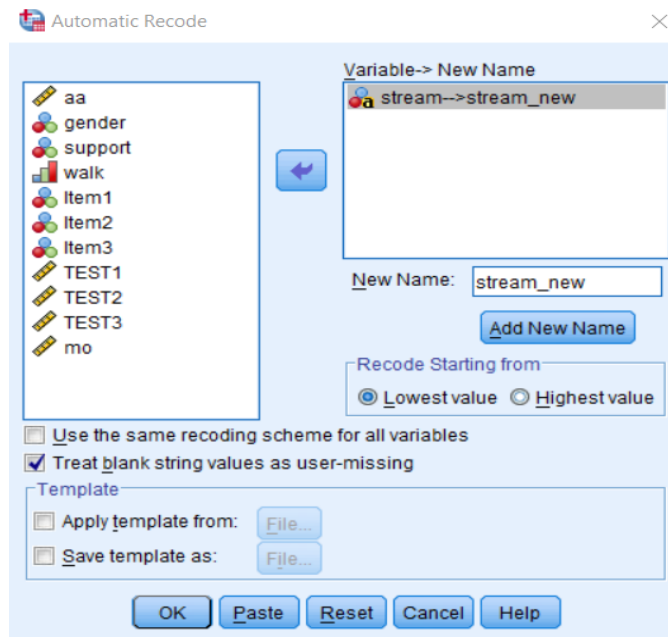
Στις ενότητες αυτού του κεφαλαίου θα χρησιμοποιηθεί το αρχείο δεδομένων του SPSS «marks_trans_2.sav». Στο αρχείο αυτό διερευνώνται οι σχέσεις του φύλου μαθητών Λυκείου (μεταβλητή «gender»), της ενισχυτικής διδασκαλίας (μεταβλητή «support»), της οικονομικής κατάστασης των μαθητών (μεταβλητή «walk») και της κατεύθυνσης σπουδών (μεταβλητή «stream»), με την αντιλαμβανόμενη ικανότητα των μαθητών στα μαθηματικά (μεταβλητές «Item1», «Item2» και «Item3»).

Η κωδικοποίηση των κατηγορικών μεταβλητών γίνεται ως εξής:

- Μεταβλητή «φύλο»: 0 για το κορίτσι και 1 για το αγόρι.
- Με τη μεταβλητή «support» καταγράφεται αν οι μαθητές παρακολουθούν την ενισχυτική διδασκαλία ή όχι. Οπότε παίρνουμε δύο τιμές: 0 για το Όχι και 1 για το Ναι.
- Με τη μεταβλητή «walk» καταγράφεται η οικονομική κατάσταση των μαθητών. Έχουμε τρεις τιμές: 0 - χαμηλή, 1 - μέτρια και 2 - καλή οικονομική κατάσταση.
- Με τη μεταβλητή «stream» καταγράφεται η κατεύθυνση σπουδών των μαθητών (θετική, θεωρητική και τεχνολογική). Είναι μια αλφαριθμητική μεταβλητή τύπου string.
- Η μεταβλητή «Item1» αντιστοιχεί στη δήλωση «είμαι πολύ καλός στα μαθηματικά». Η μεταβλητή «Item2» αντιστοιχεί στη δήλωση «δυσκολεύομαι στα μαθηματικά». Η μεταβλητή «Item3» αντιστοιχεί στη δήλωση «σπάνια κάνω λάθη στα μαθηματικά». Οι τρεις αυτές μεταβλητές παίρνουν τις τιμές 1 - διαφωνώ απόλυτα, 2 - διαφωνώ, 3 - Ούτε συμφωνώ ούτε διαφωνώ, 4 - συμφωνώ, 5 - συμφωνώ απόλυτα.

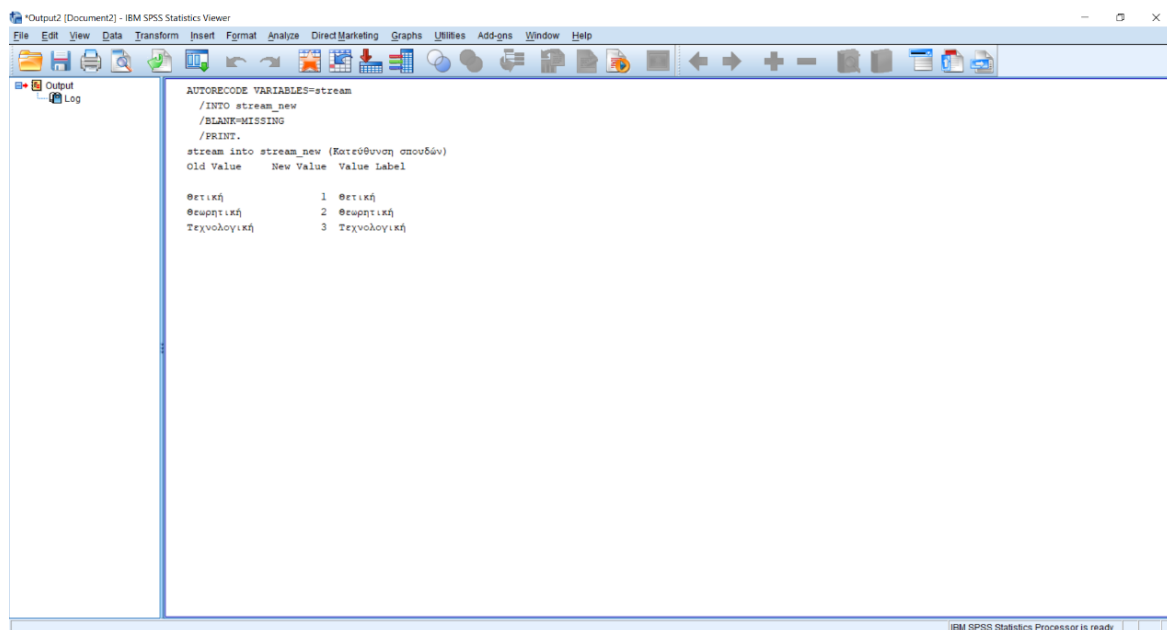
Στη συνέχεια, εφαρμόζουμε την αυτόματη επανακωδικοποίηση στην αλφαριθμητική μεταβλητή «stream» ως εξής:

- Επιλέγουμε από το μενού [Transform => Automatic Recode](#) και ανοίγει το πλαίσιο διαλόγου [Automatic Recode](#) (Εικόνα 2.1).
- Επιλέγουμε τη μεταβλητή «stream» και την εισάγουμε στο πλαίσιο [Variable => New Name](#) πατώντας το δεξί βέλος.
- Στο πεδίο [New Name](#) πληκτρολογούμε το όνομα της νέας μεταβλητής (stream_new) με την οποία θα επανακωδικοποιηθεί η αρχική μεταβλητή «stream» και κάνουμε κλικ στο κουμπί [Add New Name](#).
- Στο πλαίσιο [Recode Starting from](#), τσεκάρουμε την επιλογή [Lowest value](#), ώστε η κωδικοποίηση να αρχίζει από τη χαμηλότερη τιμή και στη συνέχεια τσεκάρουμε την επιλογή [Treat blank string values as user-missing](#) ώστε οι κενές τιμές της μεταβλητής να μετατραπούν σε user-missing values.
- Τέλος, κάνουμε κλικ στο κουμπί [OK](#).



Εικόνα 2.1 Πλαίσιο διαλόγου «Automatic Recode».

Το αρχείο αποτελεσμάτων απεικονίζει τις παλιές και τις νέες τιμές, όπως και τα value labels που αντιστοιχούν στις νέες τιμές (Εικόνα 2.2).



Εικόνα 2.2 Αρχείο αποτελεσμάτων για την αυτόματη επανακωδικοποίηση της μεταβλητής «stream».

2.1.2 Επανακωδικοποίηση στις ίδιες μεταβλητές – Recode into Same Variables

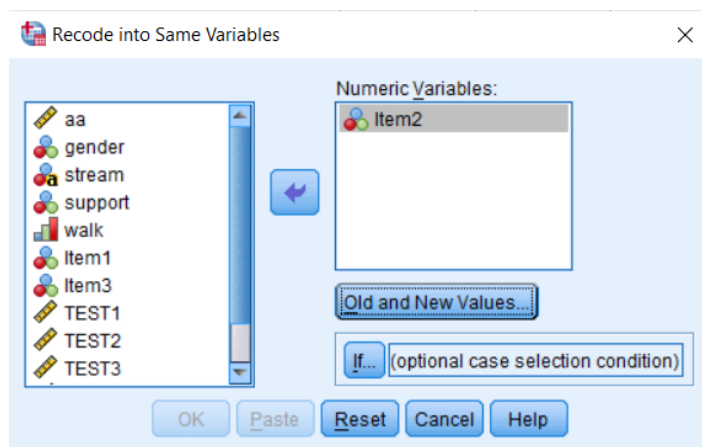
Μια ιδιαίτερα χρήσιμη δυνατότητα που μας παρέχει το SPSS, όπως είδαμε και προηγουμένως, είναι η επανακωδικοποίηση των τιμών μιας μεταβλητής. Η επανακωδικοποίηση μπορεί να γίνει είτε αντικαθιστώντας τις τιμές της αρχικής μεταβλητής είτε δημιουργώντας μια νέα μεταβλητή. Τα αποτελέσματα της επανακωδικοποίησης στην ίδια μεταβλητή (Recode into Same Variable) καταχωρίζονται ως νέες τιμές στην ίδια μεταβλητή και αντικαθιστούν ένα σύνολο τιμών της ίδιας μεταβλητής με κάποιες άλλες τιμές. Η καινούρια εκδοχή της μεταβλητής που δημιουργείται καταργεί την παλιότερη μεταβλητή, η οποία δεν είναι πλέον

διαθέσιμη. Υπάρχει επίσης η δυνατότητα να επανακωδικοποιηθούν πολλές μεταβλητές του ίδιου τύπου με την εντολή [Recode into Same Variables](#).

Οι μεταβλητές «Item1», «Item2» και «Item3», όπως αναφέρθηκε και προηγουμένως, μετρούν την αντιλαμβανόμενη ικανότητα των μαθητών στα Μαθηματικά, με μία πέντε βαθμών κλίμακα Likert. Οι τρεις αυτές μεταβλητές απαρτίζουν την κλίμακα «αντιλαμβανόμενη ικανότητα στα μαθηματικά». Όμως, η μεταβλητή «Item2», η οποία αναφέρεται στην ερώτηση «δυσκολεύομαι στα μαθηματικά», είναι αρνητικής διατύπωσης, δηλαδή οι χαμηλές τιμές αντιστοιχούν σε υψηλή αντιλαμβανόμενη ικανότητα στα μαθηματικά, ενώ οι υψηλές τιμές αντιστοιχούν σε χαμηλή αντιλαμβανόμενη ικανότητα στα μαθηματικά. Πρέπει να γίνει αντιστροφή των τιμών της μεταβλητής «Item2», έτσι ώστε να δημιουργηθεί μια νέα μεταβλητή, η οποία να έχει την ίδια κατεύθυνση με τις άλλες δύο. Αυτό μπορεί να γίνει με τη διαδικασία [Recode into Same Variables](#), με την επανακωδικοποίηση των τιμών ως εξής: Το 1 να γίνει 5, το 2 να γίνει 4, το 4 να γίνει 2 και το 5 να γίνει 1.

Εφαρμόζουμε τη διαδικασία «επανακωδικοποίηση στην ίδια μεταβλητή» για τη μεταβλητή «Item2» ως εξής:

- Επιλέγουμε από το μενού [Transform => Recode into Same Variables](#) και ανοίγει το αντίστοιχο πλαίσιο διαλόγου (Εικόνα 2.3).
- Επιλέγουμε τη μεταβλητή «Item2» και την εισάγουμε στο πεδίο [Numeric Variables](#) πατώντας το δεξί βέλος.



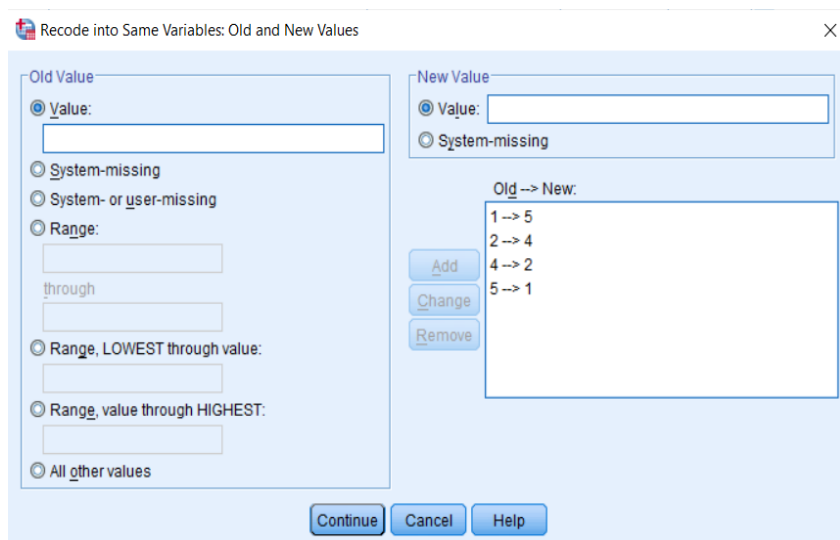
Εικόνα 2.3 Πλαίσιο διαλόγου «Recode into Same Variables».

- Κάνουμε κλικ στο κουμπί [Old and New Values](#), ώστε να εμφανιστεί το πλαίσιο διαλόγου [Old and New Values](#), στο οποίο ορίζουμε τον τρόπο που θα γίνουν οι επανακωδικοποιήσεις των τιμών μιας μεταβλητής (Εικόνα 2.4).

Το πλαίσιο αυτό περιλαμβάνει τα παρακάτω πεδία:

- **Old Value:** Στο πεδίο αυτό ορίζουμε τις παλιές τιμές που θα επανακωδικοποιηθούν. Οι επανακωδικοποιήσεις αυτές αναφέρονται σε απλές τιμές (Values), σε system-missing ή user-missing values και σε εύρος τιμών (Range).
- **New Value:** Στο πεδίο αυτό ορίζουμε τη νέα τιμή, η οποία θα αντικαταστήσει την παλιά τιμή ή ένα συγκεκριμένο εύρος παλιών τιμών. Επίσης, υπάρχει περίπτωση η system-missing value να αντικαταστήσει την παλιά τιμή.
- **Old -> New:** Στη συγκεκριμένη λίστα, κάθε φορά που έχουμε ορίσει την παλιά και τη νέα τιμή, πατάμε το πλήκτρο [Add](#) και εμφανίζεται η αντίστοιχη επανακωδικοποίηση. Έχουμε τη δυνατότητα να αλλάξουμε μια επανακωδικοποίηση πατώντας το πλήκτρο [Change](#) ή να την αφαιρέσουμε πατώντας το πλήκτρο [Remove](#).

Στο παρακάτω πλαίσιο διαλόγου [Old and New Values](#) (Εικόνα 2.4) φαίνονται οι επανακωδικοποιήσεις της μεταβλητής «Item2», προκειμένου να αντιστραφούν οι τιμές της.



Εικόνα 2.4 Πλαίσιο διαλόγου «Old and New Values».

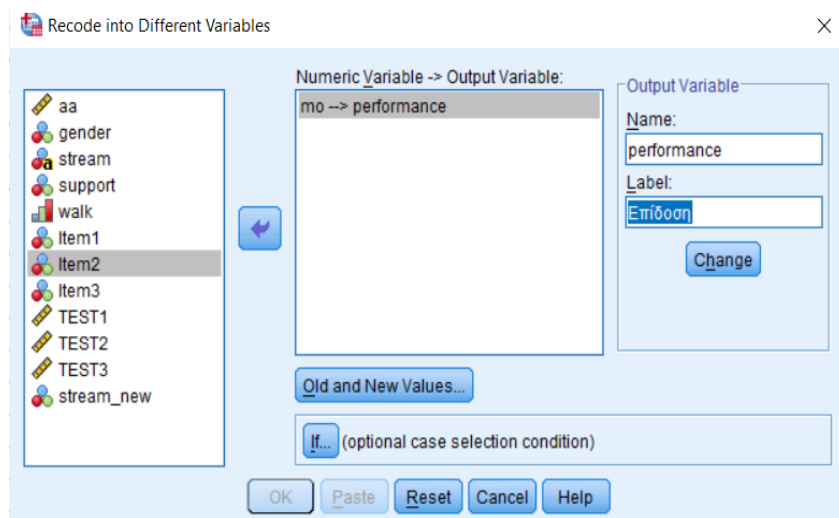
2.1.3 Επανακωδικοποίηση σε διαφορετική μεταβλητή – Recode into Different Variables

Μπορούμε να επανακωδικοποιήσουμε ποσοτικές και αλφαριθμητικές μεταβλητές, όπως επίσης μπορούμε να επανακωδικοποιήσουμε ποσοτικές μεταβλητές σε αλφαριθμητικές και το αντίστροφο. Επιπλέον, έχουμε τη δυνατότητα να επιλέξουμε πολλές μεταβλητές ώστε να τις επανακωδικοποιήσουμε όλες μαζί με την ίδια διαδικασία, αρκεί να είναι όλες του ίδιου τύπου. Δεν μπορούμε όμως να επανακωδικοποιήσουμε ποσοτικές και αλφαριθμητικές μεταβλητές μαζί (Field, 2013).

Τα αποτελέσματα της επανακωδικοποίησης σε διαφορετική μεταβλητή (Recode into Different Variables) καταχωρίζονται ως νέες τιμές σε μια άλλη καινούρια μεταβλητή που ορίζει ο ερευνητής. Η αρχική μεταβλητή παραμένει η ίδια. Για παράδειγμα, θέλουμε να δημιουργήσουμε μία νέα μεταβλητή με το όνομα «performance», η οποία θα προκύψει από την επανακωδικοποίηση των τιμών της μεταβλητής «τελικός βαθμός» (mo) σε «μεταβλητή με κλάσεις - ομάδες τιμών» ως εξής: (μικρότερος ή ίσος του 60) = Απέτυχε που θα αντιστοιχεί στο 1, (60 έως 70) = Μέτριος που θα αντιστοιχεί στο 2, (70 έως 80) = Καλός που θα αντιστοιχεί στο 3, (80 έως 90) = Πολύ καλός που θα αντιστοιχεί στο 4 και (90 έως 100) = Άριστος που θα αντιστοιχεί στο 5.

Η διαδικασία της επανακωδικοποίησης σε διαφορετική μεταβλητή γίνεται ως εξής:

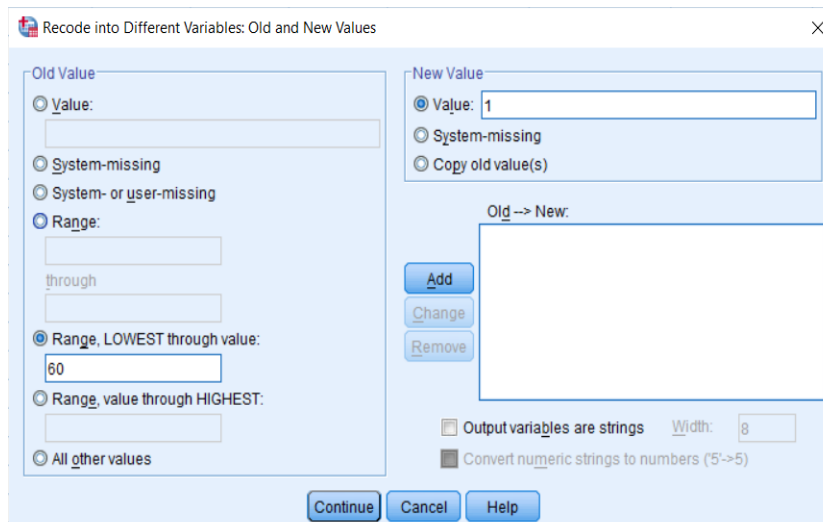
- Επιλέγουμε από το μενού **Transform => Recode into Different Variables** και ανοίγει το αντίστοιχο πλαίσιο διαλόγου.
- Επιλέγουμε από αριστερά τη μεταβλητή ή τις μεταβλητές που θέλουμε να επανακωδικοποιήσουμε και τις εισάγουμε στη λίστα προς τα δεξιά **Numeric Variable => Output Variable**, πατώντας το δεξί βέλος (Εικόνα 2.5). Εδώ, επιλέγουμε τη μεταβλητή «mo» και την εισάγουμε στη λίστα **Numeric Variable => Output Variable** πατώντας το δεξί βέλος.
- Για κάθε μεταβλητή που θα επανακωδικοποιήσουμε, πληκτρολογούμε το όνομα (Name) και την ετικέτα (Label) της καινούριας μεταβλητής. Έπειτα, πατάμε το κουμπί **Change**. Εδώ, στο πλαίσιο **Name** πληκτρολογούμε το όνομα «performance» και στο πλαίσιο **Label** πληκτρολογούμε την ετικέτα «Επίδοση». Στη συνέχεια, κάνουμε κλικ στο κουμπί **Change**.



Εικόνα 2.5 Πλαίσιο διαλόγου «Recode into Different Variables».

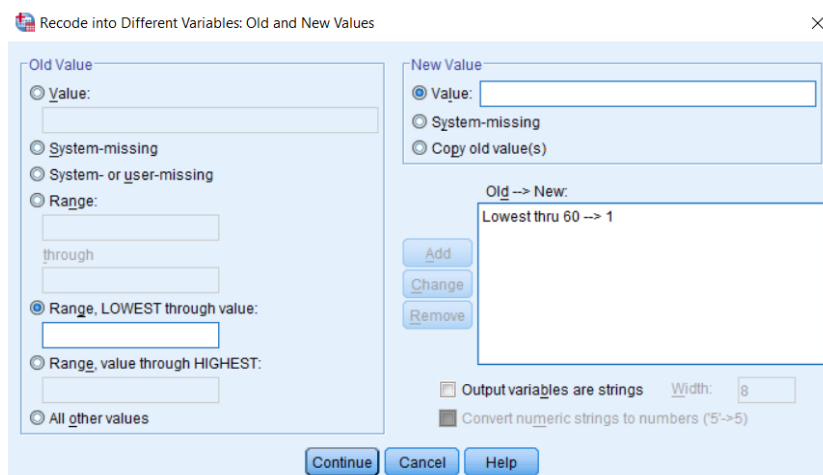
- Κάνουμε κλικ στο κουμπί **Old and New Values...**. Στο πλαίσιο διαλόγου που ανοίγει, ορίζουμε τις επανακωδικοποιήσεις των τιμών της αρχικής μεταβλητής που θα έχει η νέα μεταβλητή.
- Στο πεδίο **Old Value**, στο πλαίσιο **Value**, ορίζουμε τις επανακωδικοποιήσεις των τιμών της αρχικής μεταβλητής. Οι τιμές αυτές μπορεί να είναι είτε απλές τιμές (Values), είτε ένα εύρος τιμών (Range), είτε missing values.
- Στο πεδίο **New Value**, στο πλαίσιο **Value**, ορίζουμε τη νέα τιμή, η οποία πρόκειται να αντικαταστήσει είτε μια παλιά τιμή είτε ένα εύρος παλιών τιμών. Η νέα τιμή, αν είναι η system-missing value, πρέπει να δηλωθεί ως **System-missing**, κάνοντας κλικ στην αντίστοιχη επιλογή. Επιπλέον, η επιλογή **Copy old value(s)** διατηρεί την παλιά τιμή. Εάν ορισμένες τιμές δεν απαιτούν επανακωδικοποίηση, κάνουμε κλικ στην επιλογή **All other values**, στο πεδίο **Old Value** και έπειτα κάνουμε κλικ στην επιλογή **Copy old value(s)**.
- Στη λίστα **Old -->New**, κάθε φορά που ορίζουμε την παλιά και τη νέα τιμή και πατάμε το κουμπί **Add**, εμφανίζεται η αντίστοιχη επανακωδικοποίηση. Έχουμε τη δυνατότητα να αλλάξουμε κάποια επανακωδικοποίηση, πατώντας το κουμπί **Change**, ή να καταργήσουμε κάποια επανακωδικοποίηση πατώντας το κουμπί **Remove**.

Για να ορίσουμε την πρώτη επανακωδικοποίηση, δηλαδή αν η τιμή της μεταβλητής «mo» είναι μικρότερη ή ίση του 60 θα αντιστοιχεί στο 1, στο πεδίο **Old Value**, στο πλαίσιο **Range, LOWEST through value**, πληκτρολογούμε την τιμή 60. Στη συνέχεια, στο πεδίο **New Value**, στο πλαίσιο **Value**, πληκτρολογούμε την τιμή 1 (Εικόνα 2.6).



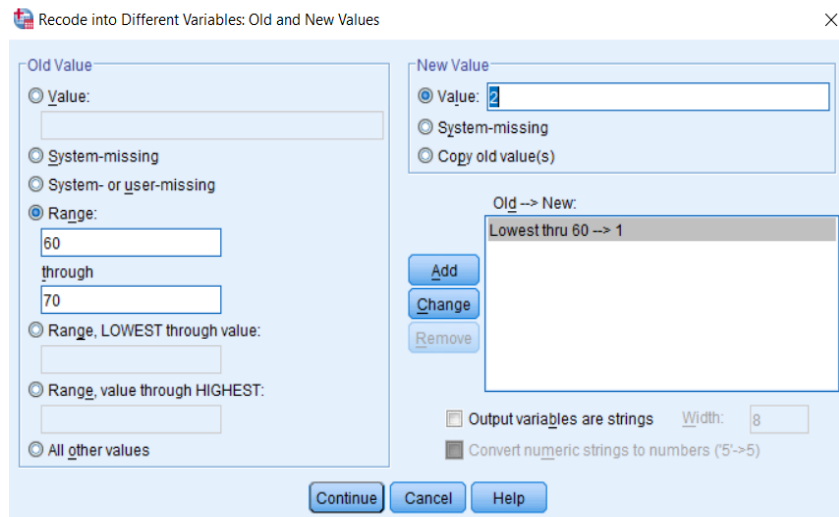
Εικόνα 2.6 Πλαίσιο διαλόγου «Old and New Values», για την επανακωδικοποίηση της τιμής ≤ 60 .

Στη συνέχεια, πατάμε το κουμπί **Add**. Στη λίστα **Old --> New** εμφανίζεται η επανακωδικοποίηση Lowest thru 60 --> 1 (Εικόνα 2.7).



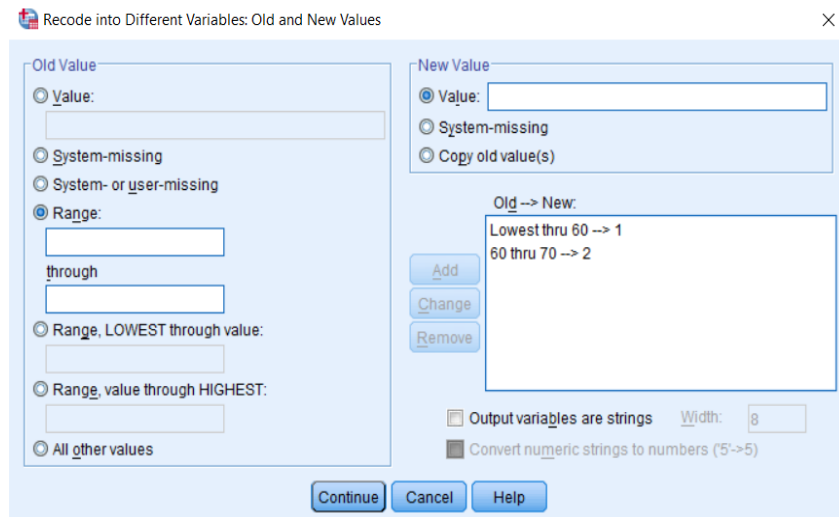
Εικόνα 2.7 Πλαίσιο διαλόγου «Old and New Values», για την πρόσθεση της επανακωδικοποίησης της τιμής ≤ 60 .

Για να ορίσουμε τη δεύτερη επανακωδικοποίηση, δηλαδή αν η τιμή της μεταβλητής «μο» κυμαίνεται πάνω από 60 έως και 70 ή αλλιώς (60, 70] θα αντιστοιχεί στο 2, στο πεδίο **Old Value**, στο πλαίσιο **Range**, στο πρώτο πεδίο, πληκτρολογούμε την τιμή 60 και στο δεύτερο πεδίο **through** πληκτρολογούμε την τιμή 70. Στη συνέχεια, στο πεδίο **New Value**, στο πλαίσιο **Value**, πληκτρολογούμε την τιμή 2 (Εικόνα 2.8).



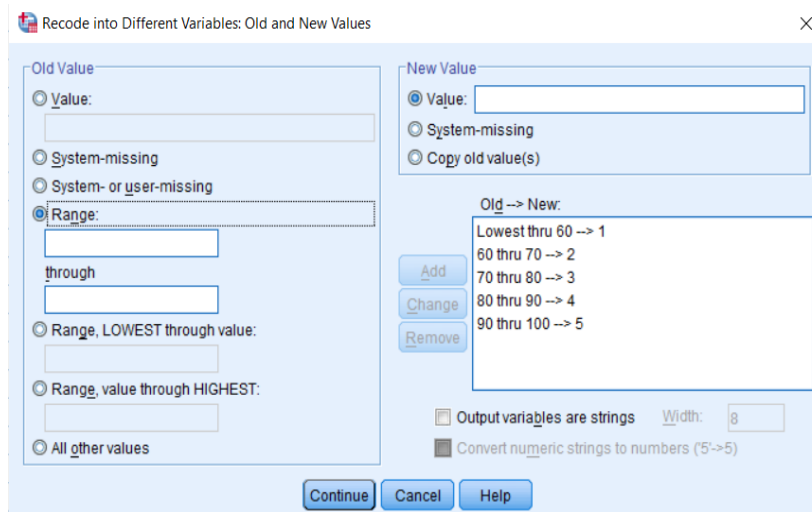
Εικόνα 2.8 Πλαίσιο διαλόγου «Old and New Values», για την επανακωδικοποίηση του εύρους τιμών 60-70.

Στη συνέχεια, πατάμε το κουμπί **Add**. Στη λίστα **Old --> New** εμφανίζεται η επανακωδικοποίηση 60 thru 70 -->2 (Εικόνα 2.9).



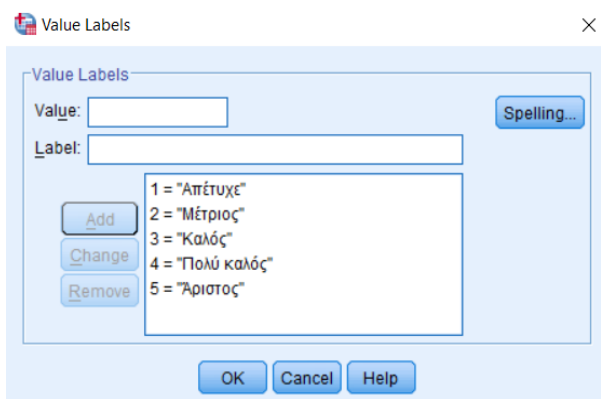
Εικόνα 2.9 Πλαίσιο διαλόγου «Old and New Values», για την πρόσθεση της επανακωδικοποίησης της τιμής 60-70.

Οι υπόλοιπες επανακωδικοποιήσεις των τιμών της μεταβλητής «μο», (70-80], (80-90] και (90-100] γίνονται ακριβώς με τον ίδιο τρόπο. Για το διάστημα τιμών (90-100], θα μπορούσαμε εναλλακτικά στο πεδίο **Range, value through HIGHEST**, να πληκτρολογήσουμε την τιμή 90, υποδεικνύοντας ότι αναφερόμαστε σε όλες τις τιμές μετά το 90. Όταν ολοκληρωθεί η διαδικασία, το πλαίσιο διαλόγου **Old and New Values** θα έχει την παρακάτω μορφή (Εικόνα 2.10):



Εικόνα 2.10 Πλαίσιο διαλόγου «Old and New Values», για την τελική επανακωδικοποίηση της μεταβλητής «πο».

Αυτό που υπολείπεται είναι να ορίσουμε στη νέα μεταβλητή «performance» τις ετικέτες των τιμών, δηλαδή να μεταβούμε στην καρτέλα **Variable View** και να δώσουμε τα εξής (Εικόνα 2.11):



Εικόνα 2.11 Ετικέτες της νέας μεταβλητής «performance» στο πλαίσιο διαλόγου «Value Labels», για την τελική επανακωδικοποίηση της μεταβλητής.

2.1.4 Διαδικασία υπολογισμού τιμών νέας μεταβλητής – Compute Variable

Αρκετά συχνά τα δεδομένα που συλλέγονται δεν έχουν την κατάλληλη μορφή ώστε να μπορούν άμεσα να αναλυθούν στατιστικά. Προκύπτει λοιπόν η ανάγκη να μετασχηματιστούν μαθηματικά ώστε να είναι συμβατά με την εκάστοτε μεθοδολογία.

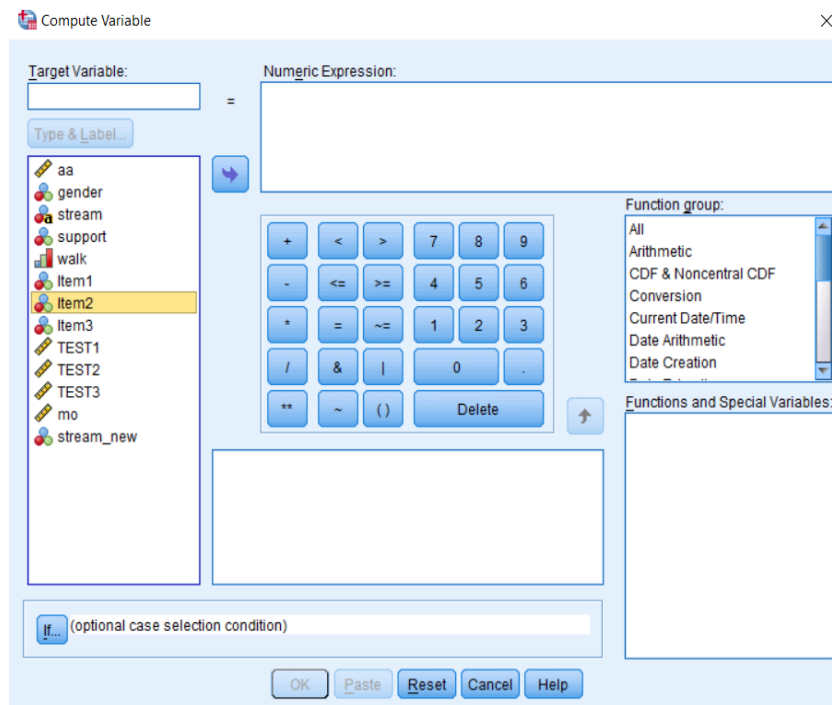
Με την εντολή **Compute Variable** του μενού **Transform** εκτελούνται οι κατάλληλοι μετασχηματισμοί των μεταβλητών ως εξής: υπολογίζονται οι τιμές μιας νέας μεταβλητής ή μετασχηματίζονται οι ήδη υπάρχουσες μεταβλητές με βάση τους υπολογισμούς των τιμών κάποιων άλλων μεταβλητών.

Η εντολή **Compute Variable** μας δίνει τις εξής δυνατότητες (Γναρδέλλης, 2013):

- Οι μεταβλητές που θα μετασχηματιστούν μπορεί να είναι ποσοτικές ή αλφαριθμητικές.
- Οι υπολογισμοί των τιμών μιας μεταβλητής μπορούν να γίνουν στο σύνολο των δεδομένων που περιέχει ή να γίνουν επιλεκτικά σε ένα μόνο υποσύνολο των δεδομένων.
- Για να υπολογιστούν οι τιμές μιας μεταβλητής μπορούμε να εκτελέσουμε δικούς μας μαθηματικούς υπολογισμούς ή να χρησιμοποιήσουμε διάφορες συναρτήσεις πολλών τύπων από ένα μεγάλο πλήθος διαθέσιμων μαθηματικών συναρτήσεων.

Η διαδικασία του υπολογισμού των τιμών μιας μεταβλητής γίνεται ως εξής:

- Επιλέγουμε από το μενού **Transform => Compute Variable** και ανοίγει το αντίστοιχο πλαίσιο διαλόγου (Εικόνα 2.12).



Εικόνα 2.12 Πλαίσιο διαλόγου «Compute Variable».

- Στο πεδίο **Target Variable** πληκτρολογούμε το όνομα της νέας μεταβλητής ή της ήδη υπάρχουσας μεταβλητής στην περίπτωση που θέλουμε να την υπολογίσουμε ξανά.
- Στο πεδίο **Numeric Expression** ορίζουμε τη μεταβλητή με τη χρήση μιας αριθμητικής έκφρασης. Η αριθμητική έκφραση μπορεί να καταχωριστεί με τη χρήση της αριθμομηχανής που βρίσκεται στο πλαίσιο διαλόγου, όπως φαίνεται στην παραπάνω εικόνα (Εικόνα 2.12). Μπορεί επίσης να καταχωριστεί χρησιμοποιώντας το πληκτρολόγιο του υπολογιστή.
- Οι αριθμητικές πράξεις με τους αντίστοιχους τελεστές τους, οι οποίες εισάγονται στην έκφραση για τον υπολογισμό των τιμών της μεταβλητής, είναι: η πρόσθεση (+), η αφαίρεση (-), ο πολλαπλασιασμός (*), η διαίρεση (/) και η ύψωση σε δύναμη (**).

Η προτεραιότητα εκτέλεσης των πράξεων είναι: εκφράσεις μέσα σε παρενθέσεις, ύψωση σε δύναμη και ρίζες, πολλαπλασιασμοί και διαιρέσεις και, τέλος, προσθέσεις και αφαιρέσεις.

Η αριθμητική έκφραση μπορεί επίσης να περιέχει συναρτήσεις, οι οποίες είναι ενσωματωμένες ανά κατηγορία στη λίστα **Function group**. Όταν επιλέγουμε μια συγκεκριμένη ομάδα, εμφανίζονται όλες οι συναρτήσεις που περιλαμβάνει η ομάδα αυτή, στη λίστα **Functions and Special Variables**. Κάνοντας κλικ σε μια συνάρτηση και πατώντας το πάνω βέλος, η συνάρτηση αυτή εισάγεται στο πεδίο **Numeric Expression**. Εκεί, πρέπει να συντάξουμε τα ορίσματα της συνάρτησης, τα οποία διαχωρίζονται με αγγλικά ερωτηματικά (?) και κόμματα (.). Ο υπολογισμός των τιμών μιας μεταβλητής μπορεί να εκτελεστεί είτε στο σύνολο των περιπτώσεων είτε σε ένα υποσύνολο περιπτώσεων. Μπορούμε να ορίσουμε το υποσύνολο των περιπτώσεων του αρχείου δεδομένων μας πατώντας το κουμπί **If** (optional case selection condition), που φαίνεται παρακάτω.

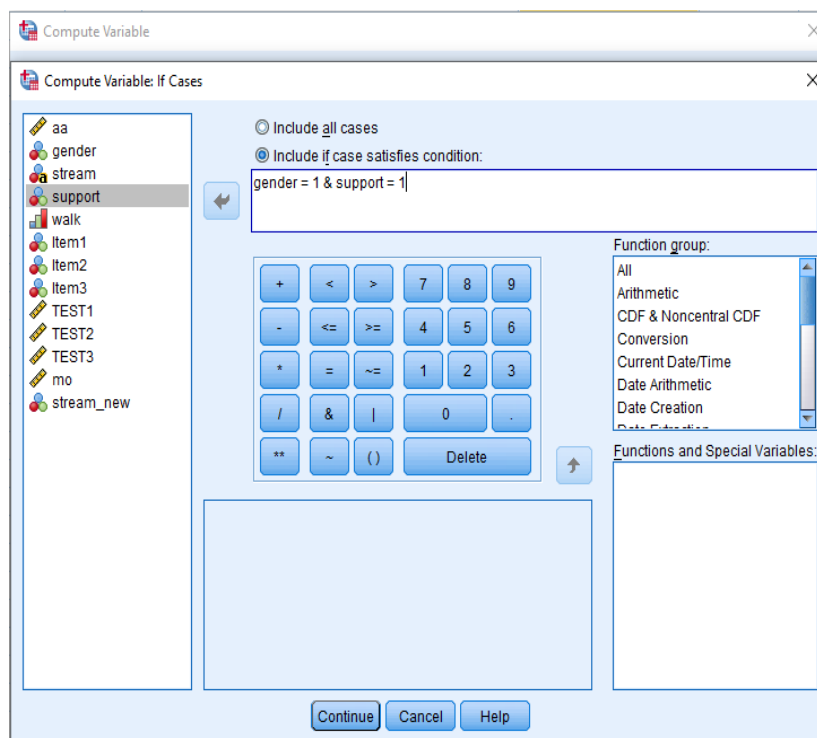
Στη συνέχεια, εμφανίζεται το πλαίσιο διαλόγου **If Cases** (Εικόνα 2.13), στο οποίο μπορούμε να επιλέξουμε κάποιες περιπτώσεις (cases) του αρχείου μας με βάση μια *συνθήκη*. Η συνθήκη αυτή έχει δύο πιθανές εκβάσεις για κάθε περίπτωση του αρχείου δεδομένων: να ικανοποιείται (αληθής) ή να μην ικανοποιείται (ψευδής). Αν η συνθήκη ικανοποιείται, επιλέγεται η περίπτωση, ενώ, αντίθετα, αν η συνθήκη δεν ικανοποιείται, τότε δεν επιλέγεται η περίπτωση. Η έκφραση της συνθήκης μπορεί να περιέχει ονόματα μεταβλητών, σταθερές, αριθμούς, αριθμητικές πράξεις, συναρτήσεις και λογικούς ή σχεσιακούς τελεστές.

Οι λογικοί τελεστές είναι η πράξη «AND» και η πράξη «OR». Ο τελεστής AND έχει ως αποτέλεσμα η λογική παράσταση να είναι αληθής μόνο όταν και τα δύο μέλη της παράστασης είναι αληθή. Σε όλες τις άλλες

περιπτώσεις η παράσταση είναι ψευδής. Ο τελεστής OR έχει ως αποτέλεσμα η λογική παράσταση να είναι ψευδής μόνο όταν και τα δύο μέλη της παράστασης είναι ψευδή. Σε όλες τις άλλες περιπτώσεις η παράσταση είναι αληθής. Μπορούμε να πληκτρολογήσουμε τις λέξεις αυτές ή να πατήσουμε στην αριθμομηχανή το σύμβολο (&) για την πράξη AND ή το σύμβολο (|) για την πράξη OR.

Οι σχεσιακοί τελεστές είναι οι: μικρότερο (<), μεγαλύτερο (>), μικρότερο ή ίσο (<=), μεγαλύτερο ή ίσο (>=), ίσο (=) και διάφορο (<>). Αν οι αριθμοί περιέχουν δεκαδικά ψηφία, τότε το δεκαδικό σύμβολο είναι η τελεία (.)

Για να γίνει καλύτερα κατανοητός ο υπολογισμός των τιμών μιας μεταβλητής σε ένα υποσύνολο περιπτώσεων, ας υποθέσουμε ότι στο αρχείο δεδομένων «Marks_trans_2.sav» θέλουμε να επιλέξουμε μόνο τα αγόρια (μεταβλητή «gender») που παρακολουθούν την ενισχυτική διδασκαλία (μεταβλητή «support»). Θα δημιουργήσουμε τη συνθήκη «gender = 1 & support = 1», όπως φαίνεται στην παρακάτω εικόνα (Εικόνα 2.13):



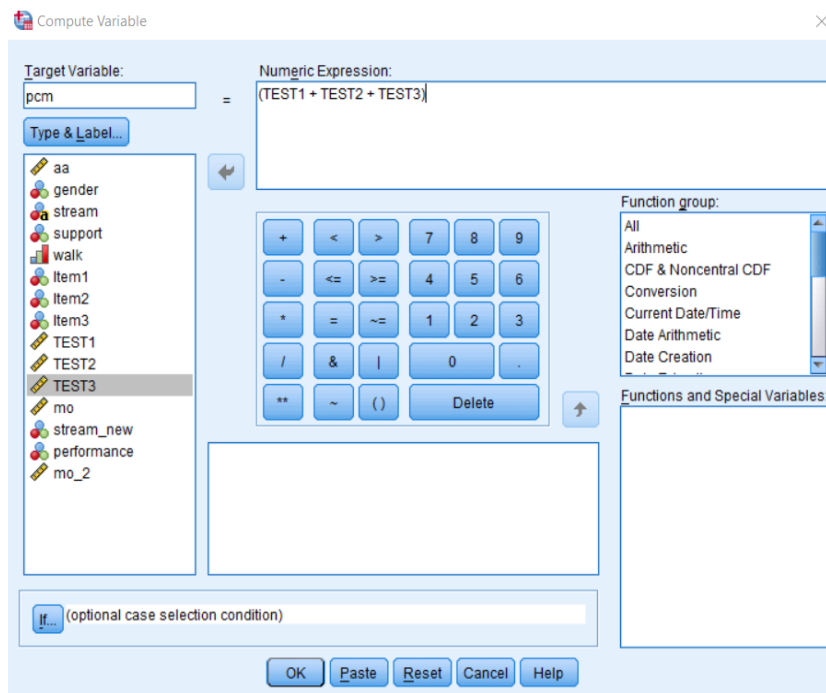
Εικόνα 2.13 Πλαίσιο διαλόγου «Compute Variable: If Cases».

Ας δούμε τώρα ένα παράδειγμα για τον υπολογισμό μιας νέας μεταβλητής με τη διαδικασία **Compute Variable**. Στο αρχείο δεδομένων «Marks_trans_2.sav» θέλουμε να δημιουργήσουμε μια καινούρια μεταβλητή, οι τιμές της οποίας θα προκύψουν από τον υπολογισμό της μέσης τιμής των τριών μεταβλητών TEST1, TEST2 και TEST3, ως εξής: $(TEST1+TEST2+TEST3) / 3$. Το όνομα της νέας μεταβλητής θα είναι «rcm».

Στην περίπτωση του υπολογισμού της μέσης τιμής με τη συγκεκριμένη αριθμητική πράξη, αν υπάρχουν ελλείπουσες τιμές (missing values) ακόμα και σε μία από τις μεταβλητές που εισέρχονται στην πράξη, στο αποτέλεσμα του υπολογισμού της τελικής μεταβλητής θα υπάρχουν επίσης missing values (θα δούμε κενό). Στην περίπτωση που χρησιμοποιηθεί η συνάρτηση της Μέσης Τιμής (Mean), αν υπάρχουν missing values, στο αποτέλεσμα του υπολογισμού της τελικής μεταβλητής θα χρησιμοποιηθούν μόνο οι μεταβλητές οι οποίες δεν έχουν missing values.

Δημιουργούμε τη νέα μεταβλητή με το όνομα «rcm» ως εξής (Εικόνα 2.14):

- Επιλέγουμε από το μενού **Transform => Compute Variable** και ανοίγει το αντίστοιχο πλαίσιο διαλόγου.
- Στο πεδίο **Target Variable** πληκτρολογούμε το όνομα «rcm».
- Στο πεδίο **Numeric Expression** ανοίγουμε παρένθεση και είτε πληκτρολογούμε την αριθμητική έκφραση $(TEST1+TEST2+TEST3) / 3$ είτε μετακινούμε μία προς μία τις μεταβλητές από αριστερά στα δεξιά μέσα στην παρένθεση, πατώντας το δεξί βέλος, είτε πατώντας διπλό κλικ πάνω σε κάθε μεταβλητή.



Εικόνα 2.14 Πλαίσιο διαλόγου «*Compute Variable*» για τον υπολογισμό της μεταβλητής «*pcm*».

Βιβλιογραφία

- Γναρδέλλης, Χ. (2013). *Ανάλυση δεδομένων με το IBM SPSS Statistics 21*. Αθήνα: Εκδ. Παπαζήση.
- Field, A. (2013). *Discovering statistics using IBM SPSS statistics* (5th ed.). London: SAGE.
- Norris, G., Qureshi, F., Howitt, D., & Cramer, D. (2017). *Εισαγωγή στη Στατιστική με το SPSS για τις Κοινωνικές Επιστήμες*. Επιμ. Εμμανουηλίδης, Χ. Αθήνα: Εκδ. Κλειδάριθμος.

Κεφάλαιο 3

Χειρισμός Δεδομένων

Σύνοψη

Στο κεφάλαιο αυτό περιγράφονται εντολές χειρισμού δεδομένων από το μενού *Data*. Συγκεκριμένα παρουσιάζονται οι εντολές *Sort Cases* (ταξινόμηση περιπτώσεων), *Sort Variables* (ταξινόμηση μεταβλητών), *Split File* (διαχωρισμός αρχείου σε υποομάδες), *Select Cases* (επιλογή περιπτώσεων), *Weight Cases* (στάθμιση περιπτώσεων), *Merge Files – Add Cases* (συγχώνευση αρχείων με τη μέθοδο της πρόσθεσης περιπτώσεων) και *Merge Files – Add Variables* (συγχώνευση αρχείων με τη μέθοδο της πρόσθεσης μεταβλητών).

Προαπαιτούμενη γνώση

Βασικές λειτουργίες του περιβάλλοντος του SPSS: Κεφάλαια 1 και 2 του συγγράμματος.

3.1 Επεξεργασία και χειρισμός δεδομένων

Το SPSS περιέχει πολλές εντολές για τον χειρισμό και την τροποποίηση των δεδομένων, που είναι απαραίτητες σε αρκετές αναλύσεις. Οι εντολές αυτές επιτρέπουν να επιλέξουμε πολλές επιλογές που αφορούν την παρουσίαση των δεδομένων έως και την επιλογή ορισμένων περιπτώσεων για συγκεκριμένες αναλύσεις. Έτσι, μεταξύ άλλων, έχουμε στη διάθεσή μας εντολές: ταξινόμησης περιπτώσεων ή μεταβλητών, διαχωρισμού αρχείου σε υποομάδες που αντιστοιχούν σε συγκεκριμένες κατηγορίες κατηγορικών μεταβλητών, επιλογής υποσυνόλου του δείγματος σύμφωνα με κάποια κριτήρια, στάθμισης των περιπτώσεων του δείγματος και συγχώνευσης αρχείων με τη μέθοδο της επέκτασης των περιπτώσεων ή των μεταβλητών.

3.1.1 Ταξινόμηση περιπτώσεων – Sort Cases

Ορισμένες διαδικασίες στο SPSS, πριν από την εκτέλεσή τους, απαιτούν τα δεδομένα μας να ταξινομηθούν με έναν συγκεκριμένο τρόπο (Γναρδέλλης, 2013; IBM SPSS Statistics Base 25, n.d.). Η διαδικασία της ταξινόμησης περιπτώσεων (**Sort Cases**) μας δίνει τη δυνατότητα να αναδιατάξουμε τις περιπτώσεις του αρχείου δεδομένων μας. Μπορούμε δηλαδή να ταξινομήσουμε τα δεδομένα των γραμμών σε σχέση με τις τιμές μίας ή και περισσότερων μεταβλητών. Οι τιμές των μεταβλητών που θα επιλέξουμε μπορούν να ταξινομηθούν σε αύξουσα (από τη μικρότερη προς τη μεγαλύτερη τιμή ή αλφαβητικά) ή φθίνουσα σειρά (από τη μεγαλύτερη προς τη μικρότερη τιμή ή αντίστροφα αλφαβητικά). Η ταξινόμηση των περιπτώσεων μπορεί να γίνει με βάση τόσο τις ποσοτικές (Scale) όσο και τις κατηγορικές μεταβλητές (Nominal).

Όταν η ταξινόμηση των περιπτώσεων γίνεται σε σχέση με τις τιμές πολλών μεταβλητών, πρώτα ταξινομούνται οι περιπτώσεις με βάση τις τιμές της πρώτης μεταβλητής. Κατόπιν, η ταξινόμηση γίνεται με βάση τις τιμές της δεύτερης μεταβλητής για καθεμία ταξινομημένη τιμή της πρώτης μεταβλητής κ.ο.κ. Όταν έχουμε ίδιες τιμές στην πρώτη μεταβλητή που έχει ταξινομηθεί, ελέγχονται οι τιμές της δεύτερης μεταβλητής κ.ο.κ.

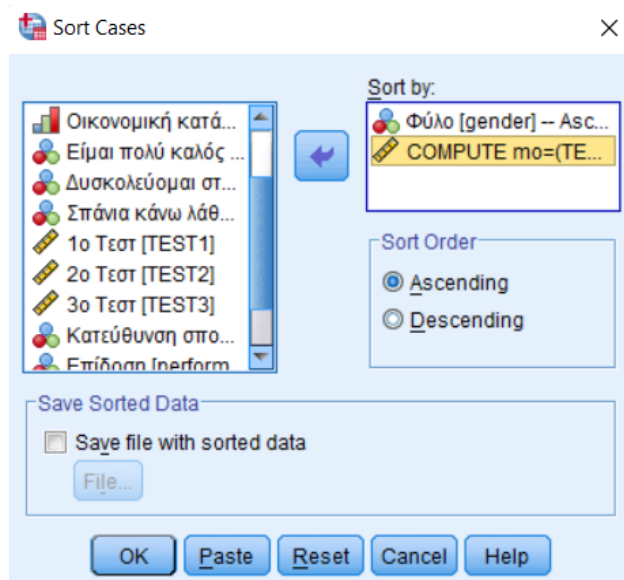
Η σειρά με την οποία εμφανίζονται τα δεδομένα δεν επηρεάζει ποτέ την ανάλυση. Ταξινομώντας τις περιπτώσεις, μπορούμε να έχουμε μία γρήγορη αίσθηση του τι συμβαίνει στα δεδομένα μας. Για παράδειγμα, αν εκτελέσουμε μία φθίνουσα ταξινόμηση πρώτα ως προς το φύλο και στη συνέχεια ως προς την επίδοση σε ένα μάθημα, θα εντοπίσουμε γρήγορα ποιος είναι ο μεγαλύτερος βαθμός στο συγκεκριμένο μάθημα και αν ανήκει σε αγόρι ή κορίτσι.

Στις ενότητες αυτού του κεφαλαίου θα χρησιμοποιηθεί το αρχείο δεδομένων του SPSS «marks_trans4_reverse_item4_recode_diff». Θέλουμε, για παράδειγμα, να ταξινομήσουμε τις περιπτώσεις του αρχείου με βάση πρώτα τη μεταβλητή φύλο (gender) κατά αύξουσα σειρά και έπειτα με βάση τη μεταβλητή της μέσης τιμής των 3 τεστ (mo), επίσης κατά αύξουσα σειρά. Επειδή στη μεταβλητή «gender» η κατηγορία «κορίτσι» είναι κωδικοποιημένη με το 0 και η κατηγορία «αγόρι» είναι κωδικοποιημένη με το 1, αφού εκτελεστεί η συγκεκριμένη διαδικασία ταξινόμησης, θα εμφανίζονται πρώτα τα κορίτσια από τη χαμηλότερη

προς την υψηλότερη μέση τιμή και στη συνέχεια θα εμφανίζονται τα αγόρια, επίσης από τη χαμηλότερη προς την υψηλότερη μέση τιμή.

Η διαδικασία της ταξινόμησης περιπτώσεων γίνεται ως εξής:

- Επιλέγουμε από το μενού Data => Sort Cases και ανοίγει το αντίστοιχο πλαίσιο διαλόγου (Εικόνα 3.1).



Εικόνα 3.1 Πλαίσιο διαλόγου «Sort Cases».

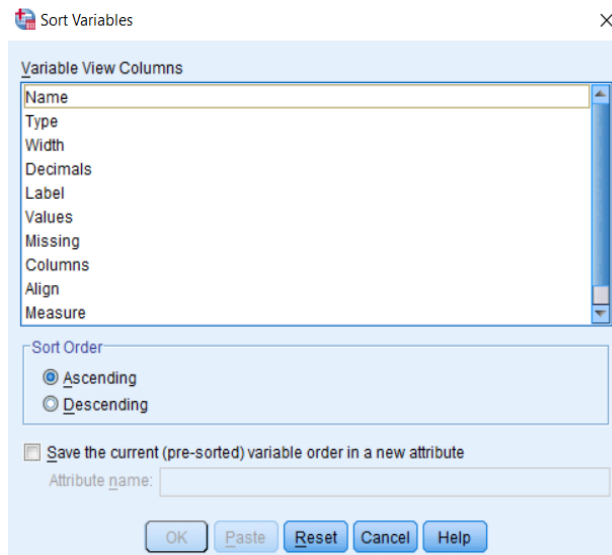
- Επιλέγουμε από αριστερά τις μεταβλητές σε σχέση με τις οποίες θα γίνει η ταξινόμηση των περιπτώσεων και τις εισάγουμε στο πεδίο **Sort by** πατώντας το δεξί βέλος. Δηλαδή, επιλέγουμε πρώτα τη μεταβλητή «gender» και μετά τη μεταβλητή «mo».
- Η σειρά ταξινόμησης επιλέγεται στο πλαίσιο Sort Order. Η αύξουσα ταξινόμηση αντιστοιχεί στην επιλογή Ascending, ενώ η φθίνουσα ταξινόμηση στην επιλογή Descending.
- Έχουμε τη δυνατότητα (προαιρετικά) να αποθηκεύσουμε τα πρόσφατα ταξινομημένα δεδομένα μας σε ένα νέο αρχείο επιλέγοντας το πλαίσιο ελέγχου **Save file with sorted data**. Στη συνέχεια, κάνουμε κλικ στο κουμπί **File...**, για να ορίσουμε ένα όνομα και μια τοποθεσία για το νέο αρχείο δεδομένων.
- Τέλος, κάνουμε κλικ στο κουμπί **OK**.

3.1.2 Ταξινόμηση μεταβλητών – Sort Variables

Ένα αρχείο μπορεί να ταξινομηθεί επίσης σε σχέση με τις μεταβλητές του. Η ταξινόμηση των μεταβλητών θα αναδιατάξει τη σειρά των στηλών του αρχείου δεδομένων μας. Οι μεταβλητές ταξινομούνται με βάση τις τιμές μόνο μίας ιδιότητάς τους τους τη φορά, δηλαδή: Όνομα, Τύπος, Πλάτος, Δεκαδικά ψηφία, Ετικέτες, Ελλείπουσες τιμές, Στήλες, Στοιχισι, Μέτρηση ή Προσαρμοσμένο Χαρακτηριστικό. Οι μεταβλητές μπορούν να ταξινομηθούν σε αύξουσα ή φθίνουσα σειρά σε σχέση με την επιλεγμένη ιδιότητα (Γναρδέλλης, 2013; IBM SPSS Statistics Base 25, n.d.). Θα χρησιμοποιηθεί το ίδιο αρχείο δεδομένων του SPSS «marks_trans4_reverse_item4_recode_diff».

Η διαδικασία της ταξινόμησης μεταβλητών γίνεται ως εξής:

- Επιλέγουμε από το μενού Data => Sort Variables και ανοίγει το αντίστοιχο πλαίσιο διαλόγου (Εικόνα 3.2).



Εικόνα 3.2 Πλαίσιο διαλόγου «Sort Variables».

- Επιλέγουμε την ιδιότητα των μεταβλητών βάσει της οποίας θέλουμε να γίνει η ταξινόμηση.
- Επιλέγουμε τη σειρά ταξινόμησης από το πλαίσιο **Sort Order**. Η αύξουσα ταξινόμηση αντιστοιχεί στην επιλογή **Ascending**, ενώ η φθίνουσα ταξινόμηση στην επιλογή **Descending**.
- Μπορούμε προαιρετικά να αποθηκεύσουμε την τρέχουσα (προ-ταξινομημένη) σειρά μεταβλητών σαν μια καινούρια ιδιότητα, κάνοντας κλικ στην επιλογή **Save the current (pre-sorted) variable order in a new attribute** και πληκτρολογώντας ένα όνομα στο πλαίσιο ελέγχου **Attribute name**.
- Τέλος, κάνουμε κλικ στο κουμπί **OK** (Εικόνα 3.2).

Εναλλακτικά, για να ταξινομήσουμε τις μεταβλητές του αρχείου δεδομένων μας, μπορούμε απλώς να κάνουμε δεξί κλικ στο όνομα της στήλης που αντιστοιχεί στην ιδιότητα των μεταβλητών βάσει της οποίας θα γίνει η ταξινόμηση και να επιλέξουμε τη σειρά ταξινόμησης.

3.1.3 Διαχωρισμός αρχείου σε υποομάδες – Split File

Όταν πραγματοποιούμε ανάλυση των δεδομένων, μερικές φορές χρειάζεται να διαχωρίσουμε προσωρινά τα δεδομένα μας προκειμένου να συγκρίνουμε τα αποτελέσματα σε διαφορετικές υποομάδες περιπτώσεων. Αυτό μπορεί να φανεί πολύ χρήσιμο, όταν θέλουμε να συγκρίνουμε κατανομές συχνοτήτων ή περιγραφικά στατιστικά σε σχέση με τις κατηγορίες κάποιας μεταβλητής (π.χ. το φύλο), ειδικά στην περίπτωση που θέλουμε ξεχωριστούς πίνακες αποτελεσμάτων για κάθε υποομάδα (Γναρδέλλης, 2013; IBM SPSS Statistics Base 25, n.d.).

Η εντολή **Split File** διαχωρίζει τα δεδομένα σε υποομάδες περιπτώσεων με βάση τις κατηγορίες μίας ή και περισσότερων κατηγορικών μεταβλητών ελέγχου. Ο διαχωρισμός αυτός έχει να κάνει μόνο με τα αποτελέσματα των αναλύσεων και όχι με τα δεδομένα. Στην περίπτωση που η διαίρεση του αρχείου θα γίνει με βάση τις κατηγορίες πολλών μεταβλητών, αυτή πραγματοποιείται ιεραρχικά (Bryman, 2016). Πιο συγκεκριμένα, ο διαχωρισμός θα ξεκινήσει από την πρώτη μεταβλητή της αντίστοιχης λίστας, στη συνέχεια εσωτερικά στην κάθε κατηγορία της αρχικής μεταβλητής, σε σχέση με τις τιμές της δεύτερης μεταβλητής κ.ο.κ. Για παράδειγμα, εάν επιλέξουμε το φύλο ως την πρώτη μεταβλητή ομαδοποίησης και την εθνικότητα ως τη δεύτερη μεταβλητή ομαδοποίησης, οι περιπτώσεις θα ομαδοποιηθούν ανά ομάδα εθνικότητας σε κάθε κατηγορία του φύλου.

Η εντολή **Split File** μας παρέχει δύο επιλογές ως προς τη μορφή που θα εμφανιστούν τα αποτελέσματα της ανάλυσης:

- **Compare groups**. Με αυτή την επιλογή τα αποτελέσματα μιας ανάλυσης εμφανίζονται ξεχωριστά σε έναν κοινό πίνακα για καθεμία κατηγορία, προκειμένου να συγκρίνουμε τις

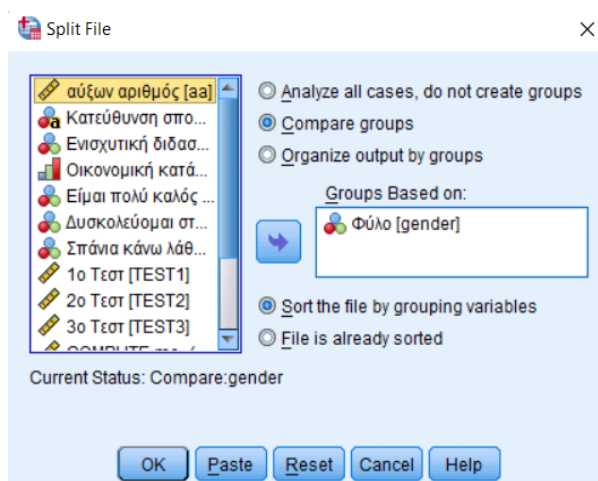
κατηγορίες μεταξύ τους. Σχετικά με τα γραφήματα, δημιουργείται ένα ξεχωριστό γράφημα για κάθε ομάδα και όλα τα γραφήματα εμφανίζονται μαζί στον Output Viewer.

- **Organize output by groups.** Με τη συγκεκριμένη επιλογή τα αποτελέσματα της ανάλυσης οργανώνονται και εμφανίζονται χωριστά, δηλαδή σε διαφορετικούς πίνακες για καθεμία κατηγορία.

Για τον διαχωρισμό αρχείου σε υποομάδες, θα χρησιμοποιηθεί το ίδιο αρχείο δεδομένων του SPSS «marks_trans4_reverse_item4_recode_diff.sav».

Η διαδικασία διαχωρισμού του αρχείου γίνεται ως εξής:

- Επιλέγουμε από το μενού **Data => Split File** και ανοίγει το αντίστοιχο πλαίσιο διαλόγου (Εικόνα 3.3).



Εικόνα 3.3 Πλαίσιο διαλόγου «Split File».

- Επιλέγουμε τη μεταβλητή ή τις μεταβλητές με βάση τις τιμές των οποίων πραγματοποιείται ο διαχωρισμός των παρατηρήσεων του αρχείου σε ομάδες.
- Επιλέγουμε τη μορφή εμφάνισης των αποτελεσμάτων (**Compare groups** ή **Organize output by groups**).
- Οι περιπτώσεις του αρχείου δεδομένων πρέπει να είναι ταξινομημένες κατά αύξουσα σειρά, κατά τις τιμές των μεταβλητών διαχωρισμού και με την ίδια σειρά που παρατίθενται οι μεταβλητές στη λίστα **Groups Based On**. Εάν το αρχείο δεδομένων δεν είναι ήδη ταξινομημένο, επιλέγουμε την επιλογή **Sort the file by grouping variables**.
- Τέλος, κάνουμε κλικ στο κουμπί **OK** (Εικόνα 3.3).

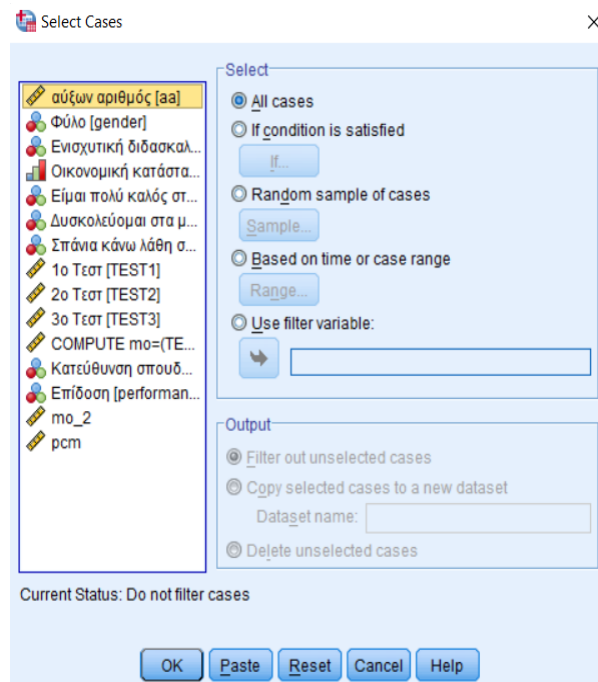
Αν, για παράδειγμα, θέλουμε να συγκρίνουμε ποσοστιαία τα δύο φύλα όσον αφορά το αν παρακολουθούν την ενισχυτική διδασκαλία ή όχι, θα κάνουμε τη διαδικασία **Split File**, όπως φαίνεται στην Εικόνα 3.3. Επιλέγοντας ως μορφή εμφάνισης των αποτελεσμάτων την **Compare groups**, το SPSS μας εμφανίζει τον παρακάτω πίνακα:

Ενισχυτική διδασκαλία						
Φύλο			Frequency	Percent	Valid Percent	Cumulative Percent
Κορίτσι	Valid	Όχι	13	39,4	39,4	39,4
		Ναι	20	60,6	60,6	100,0
		Total	33	100,0	100,0	
Αγόρι	Valid	Όχι	17	60,7	60,7	60,7
		Ναι	11	39,3	39,3	100,0
		Total	28	100,0	100,0	

Πίνακας 3.1 Πίνακας αποτελεσμάτων με τη μορφή εμφάνισης «Compare groups» από το «Split File».

3.1.4 Επιλογή περιπτώσεων – Select Cases

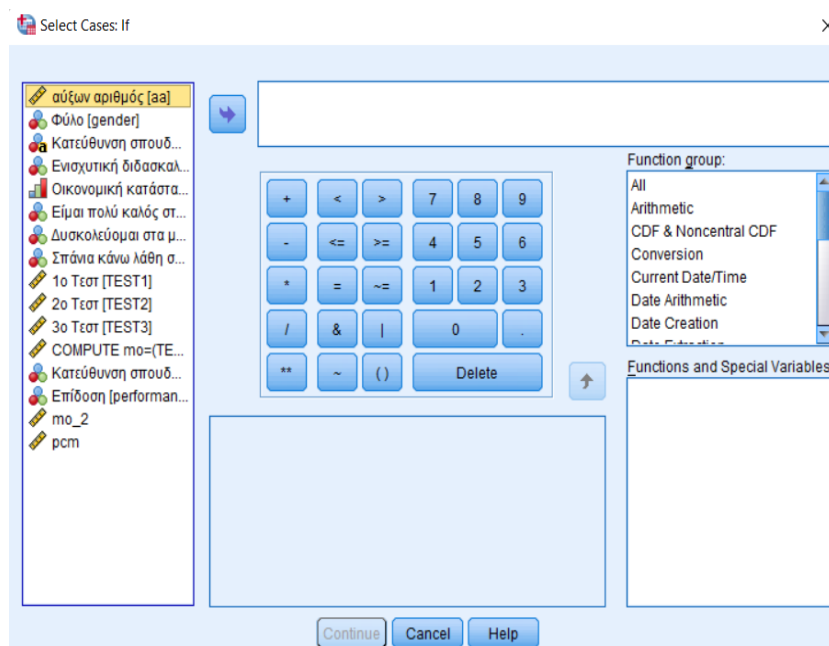
Με τη χρήση της εντολής **Select Cases** έχουμε τη δυνατότητα να επιλέξουμε ένα ή περισσότερα υποσύνολα περιπτώσεων ενός αρχείου δεδομένων που θέλουμε να αναλύσουμε (Γναρδέλλης, 2013; Bryman, 2016; IBM SPSS Statistics Base 25, n.d.). Οι περιπτώσεις αυτές πρέπει να πληρούν κάποια κριτήρια. Για παράδειγμα, θέλουμε να δημιουργήσουμε έναν πίνακα συχνοτήτων μίας μεταβλητής μόνο για μία συγκεκριμένη ηλικιακή ομάδα. Επιλέγοντας από το μενού **Data** την επιλογή **Select Cases**, εμφανίζεται το παρακάτω πλαίσιο διαλόγου (Εικόνα 3.4).



Εικόνα 3.4 Πλαίσιο διαλόγου «Select Cases».

Στο δεξί μέρος του μενού αυτού, έχουμε τις εξής λύσεις με τις οποίες θα γίνει η επιλογή των περιπτώσεων:

- **All cases.** Επιλογή όλων των περιπτώσεων.
- **If condition is satisfied.** Η επιλογή αυτή ενεργοποιείται πατώντας το πλήκτρο **if**.
- Εμφανίζεται ένα νέο πλαίσιο διαλόγου, το **Select Cases: If** (Εικόνα 3.5). Το υποπλαίσιο αυτό είναι παραπλήσιο με το μενού της **Transform => Compute Variable**, που είδαμε στο Κεφάλαιο 2. Στο πλαίσιο πάνω δεξιά πληκτρολογούμε τη συνθήκη, βάσει της οποίας θα γίνει η επιλογή των περιπτώσεων. Από αριστερά επιλέγουμε τις μεταβλητές που θα εισαχθούν στη συνθήκη. Η συνθήκη ελέγχεται, αν ικανοποιείται ή όχι. Μία περίπτωση επιλέγεται μόνο όταν ικανοποιείται η συγκεκριμένη συνθήκη. Αν η συνθήκη δεν ικανοποιείται, τότε η περίπτωση δεν επιλέγεται.



Εικόνα 3.5 Υποπλάσιο διαλόγου «Select Cases: If».

Η συνθήκη μπορεί να εκφράζεται με ονόματα μεταβλητών, με τιμές ή εύρος τιμών μιας μεταβλητής, αριθμούς, συναρτήσεις και αριθμητικές ή λογικές πράξεις. Οι αριθμητικοί τελεστές είναι: η πρόσθεση (+), η αφαίρεση (-), ο πολλαπλασιασμός (*), η διαίρεση (/) και η ύψωση σε δύναμη (**). Οι σχεσιακοί τελεστές είναι οι: μικρότερο (<), μεγαλύτερο (>), μικρότερο ή ίσο (<=), μεγαλύτερο ή ίσο (>=), ίσο (=) και διάφορο (~=). Οι λογικοί τελεστές είναι το «AND» και το «OR». Ο τελεστής AND έχει ως αποτέλεσμα η λογική συνθήκη να είναι αληθής μόνο όταν και τα δύο μέλη της συνθήκης είναι αληθή. Σε όλες τις άλλες περιπτώσεις η συνθήκη είναι ψευδής. Ο τελεστής OR έχει ως αποτέλεσμα η λογική συνθήκη να είναι ψευδής μόνο στην περίπτωση όπου και τα δύο μέλη της συνθήκης είναι ψευδή. Σε όλες τις άλλες περιπτώσεις η συνθήκη είναι αληθής. Μπορούμε να πληκτρολογήσουμε τις λέξεις αυτές ή να πατήσουμε στην αριθμομηχανή το σύμβολο (&) για τον τελεστή AND ή το σύμβολο (|) για τον τελεστή OR.

Για παράδειγμα, στο αρχείο δεδομένων «marks_trans4_reverse_item4_recode_diff», αν η συνθήκη είναι: «gender = 1 & TEST1 >= 60», τότε επιλέγονται μόνο οι περιπτώσεις στις οποίες η μεταβλητή «gender» παίρνει την τιμή 1 για τα αγόρια και η τιμή της μεταβλητής «TEST1» είναι μεγαλύτερη ή ίση από 60 (Εικόνα 3.6).

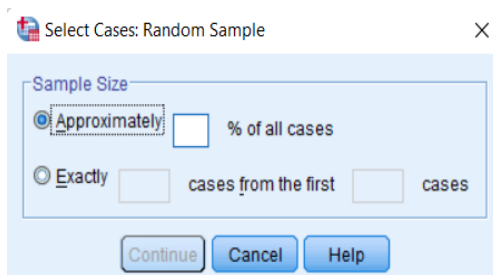
	aa	gender	stream	support	walk	Item1	Item2	Item3	TEST1	TEST2	TEST3	filter_\$	var
30	50	0	Τεχνολογ...	1	1	5	4	4	87	89	98	0	
31	41	0	Θεωρητική	0	0	1	3	1	98	90	91	0	
32	43	0	Τεχνολογ...	1	1	5	5	4	87	97	95	0	
33	51	0	Τεχνολογ...	0	1	2	2	2	87	97	95	0	
34	2	1	Θετική	1	0	4	4	4	60	45	40	1	
35	16	1	Θετική	0	1	3	2	3	47	45	67	0	
36	1	1	Θετική	0	0	1	3	1	55	56	58	0	
37	18	1	Θετική	0	0	2	4	1	51	75	45	0	
38	54	1	Τεχνολογ...	0	1	1	3	1	56	56	67	0	
39	3	1	Θετική	1	1	4	4	4	67	46	78	1	
40	28	1	Θεωρητική	0	2	3	1	3	75	54	65	1	
41	59	1	Τεχνολογ...	0	2	1	2	5	51	63	82	0	
42	56	1	Τεχνολογ...	0	0	1	3	1	51	71	81	0	
43	6	1	Θετική	1	0	4	4	5	78	56	77	1	
44	60	1	Τεχνολογ...	1	2	3	5	4	51	78	82	0	
45	55	1	Τεχνολογ...	0	0	2	2	1	75	70	67	1	
46	17	1	Θετική	0	0	1	3	2	75	83	67	1	
47	21	1	Θετική	0	2	2	3	2	89	56	82	1	
48	35	1	Θεωρητική	0	1	4	5	1	78	78	73	1	
49	33	1	Θεωρητική	0	2	3	5	5	78	82	70	1	

Εικόνα 3.6 Ενδεικτική μορφή της καρτέλας Data View με την εκτέλεση της εντολής «Select Cases».

Οι περιπτώσεις που δεν επιλέγονται (unselected cases) παραμένουν στο αρχείο δεδομένων και δεν διαγράφονται, εκτός και αν αποφασίσουμε ότι δεν μας χρειάζονται πλέον. Οι περιπτώσεις που δεν θα επιλεγθούν δεν χρησιμοποιούνται στις αναλύσεις. Κάθε φορά που επιλέγουμε περιπτώσεις με τη χρήση της εντολής **Select Cases** το SPSS δημιουργεί μία καινούρια μεταβλητή που ονομάζεται *φίλτρο* (filter_\$). Η μεταβλητή «filter_\$» παίρνει την τιμή 1 για τις μεταβλητές που επιλέγονται και την τιμή 0 για τις μεταβλητές που δεν επιλέγονται. Οι μη επιλεγμένες περιπτώσεις διακρίνονται με μια πλάγια γραμμή πάνω στον αύξοντα αριθμό τους στην καρτέλα Data View. Επίσης, στη γραμμή κατάστασης, κάτω δεξιά, εμφανίζεται η ένδειξη **Filter On**.

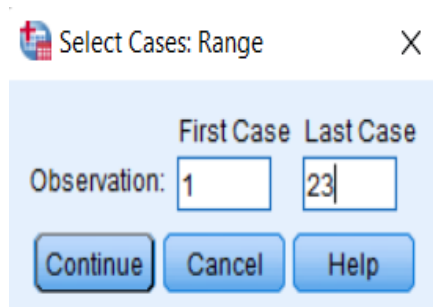
Στο πλαίσιο **Output**, επιλέγοντας **Filter out unselected cases**, οι μη επιλεγμένες περιπτώσεις μπορούν να ανακληθούν αργότερα, ενεργοποιώντας την επιλογή **All cases** χωρίς να χρειάζεται να ανοίξουμε ξανά το αρχείο. Επιλέγοντας **Delete unselected cases** οι μη επιλεγμένες περιπτώσεις διαγράφονται. Για να μην τις χάσουμε πρέπει να αποθηκεύσουμε το αρχείο με ένα διαφορετικό όνομα. Εναλλακτικά, επιλέγοντας **Copy selected cases to a new dataset**, μπορούμε να αντιγράψουμε τις επιλεγμένες περιπτώσεις σε ένα νέο αρχείο δεδομένων (dataset). Στο πλαίσιο **Dataset name** θα πληκτρολογήσουμε το όνομα του νέου αρχείου.

- **Random sample of cases.** Η επιλογή αυτή ενεργοποιείται πατώντας το πλήκτρο **Sample**. Εμφανίζεται ένα νέο πλαίσιο διαλόγου, το **Select Cases: Random Sample** (Εικόνα 3.7). Εδώ μπορούμε να επιλέξουμε ένα υποσύνολο περιπτώσεων με τυχαία δειγματοληψία. Το μέγεθος του τυχαίου δείγματος που θα επιλέξουμε μπορεί να προκύψει κατά προσέγγιση μέσω της επιλογής **Approximately**, ορίζοντας ένα ποσοστό των περιπτώσεων, π.χ. περίπου το 20%. Επιπλέον, μέσω της επιλογής **Exactly**, μπορούμε να επιλέξουμε ένα τυχαίο δείγμα ακριβώς ίσο με το πλήθος των περιπτώσεων που θα ορίσουμε. Το πλήθος αυτό πρέπει να είναι μικρότερο ή ίσο του αριθμού των περιπτώσεων του αρχικού αρχείου δεδομένων.



Εικόνα 3.7 Υποπλαίσιο διαλόγου «Select Cases: Random Sample»

- **Based on time or case range.** Η επιλογή αυτή ενεργοποιείται πατώντας το πλήκτρο **Range**. Εμφανίζεται ένα νέο πλαίσιο διαλόγου, το **Select Cases: Range** (Εικόνα 3.8).



Εικόνα 3.8 Υποπλαίσιο διαλόγου «Select Cases: Range»

Στο υποπλαίσιο διαλόγου **Select Cases: Range**, όπως φαίνεται στην Εικόνα 3.8, μπορούμε να επιλέξουμε ένα εύρος υποσυνόλου παρατηρήσεων πληκτρολογώντας στο πρώτο κουτί (*First Case*) τον αύξοντα αριθμό της πρώτης περίπτωσης και πληκτρολογώντας στο δεύτερο κουτί (*Last Case*) τον αύξοντα αριθμό της τελευταίας περίπτωσης που θα περιέχει το νέο υποσύνολο παρατηρήσεων. Αν, για παράδειγμα, θέλουμε να επιλέξουμε ένα δείγμα, το οποίο θα περιέχει ένα εύρος από την 1^η έως την 23^η περίπτωση, θα πληκτρολογήσουμε τις τιμές, όπως φαίνονται στην Εικόνα 3.8. Η διαδικασία της επιλογής περιπτώσεων ολοκληρώνεται πατώντας το πλήκτρο **OK**.

Για να ακυρώσουμε την επιλογή περιπτώσεων και να επιλέξουμε πάλι όλες τις περιπτώσεις του αρχείου δεδομένων μας, αρκεί να επιλέξουμε **Data => Select Cases => All cases**. Εναλλακτικά, μπορούμε να μεταβούμε στην καρτέλα Variable View και να διαγράψουμε τη μεταβλητή «filter_\$», κάνοντας δεξί κλικ πάνω στον αύξοντα αριθμό της και επιλέγοντας **Clear**.

3.1.5 Στάθμιση περιπτώσεων – Weight Cases

Η εντολή **Weight Cases** σταθμίζει τις περιπτώσεις σύμφωνα με τις τιμές μιας άλλης αριθμητικής μεταβλητής συχνοτήτων που θα δημιουργήσουμε (μεταβλητή στάθμισης). Οι τιμές της μεταβλητής στάθμισης δηλώνουν τον αριθμό των επαναλήψεων για κάθε περίπτωση της αρχικής μεταβλητής, την οποία θέλουμε να σταθμίσουμε (Γναρδέλλης, 2013; Bryman, 2016; IBM SPSS Statistics Base 25, n.d.). Η μεταβλητή στάθμισης μπορεί να είναι μια μεταβλητή του αρχικού αρχείου ή μια άλλη μεταβλητή (για παράδειγμα η μεταβλητή μιας στατιστικής ανάλυσης).

Η μεταβλητή στάθμισης μπορεί να έχει δεκαδικά ψηφία. Για παράδειγμα, αν η κατηγορία μιας μεταβλητής έχει περισσότερες παρατηρήσεις και θέλουμε να της δώσουμε χαμηλότερη βαρύτητα (π.χ. να την υποδιπλασιάσουμε), τότε η μεταβλητή στάθμισης μπορεί να πάρει την τιμή 0,5.

Οι μεταβλητές οι οποίες έχουν μηδενικές, αρνητικές ή ελλείπουσες τιμές δεν χρησιμοποιούνται σε μια ανάλυση. Τα γραφήματα που έχουν κατασκευαστεί με μεταβλητές στάθμισης αναγράφουν κάτω δεξιά με ποια μεταβλητή σταθμίστηκαν οι περιπτώσεις (cases weighted by..).

Επιπλέον, η εντολή **Weight Cases** χρησιμοποιείται, όταν πραγματοποιούμε μια ποιοτική έρευνα. Σε μια ποιοτική έρευνα ο ερευνητής συνήθως κατασκευάζει έναν πίνακα διασταύρωσης, ο οποίος περιέχει τις συχνότητες των κατηγοριών, δύο ποιοτικών μεταβλητών, δηλαδή το πλήθος των τιμών των μεταβλητών και όχι τις τιμές όλων των περιπτώσεων. Για να ελέγξουμε τη σχέση δύο ποιοτικών μεταβλητών, αφού έχει κατασκευαστεί ένας πίνακας διασταύρωσης με τα δεδομένα, χρησιμοποιούμε τη στάθμιση περιπτώσεων.

Για παράδειγμα, σε μια ποιοτική έρευνα, που αφορούσε 29 διδακτικές προσεγγίσεις που εφαρμόστηκαν από άνδρες και γυναίκες εκπαιδευτικούς, καταγράφηκε κατά πόσο αυτές οι προσεγγίσεις εμπίπτουν σε δασκαλοκεντρική προσέγγιση, μαθητοκεντρική προσέγγιση και μεικτή προσέγγιση. Δημιουργήθηκε ένας πίνακας διασταύρωσης λαμβάνοντας υπόψη τις μεταβλητές «διδασκτική προσέγγιση» και «φύλο» (Πίνακας 3.2).

Για να προχωρήσουμε στην ανάλυση στο περιβάλλον του SPSS, μία λύση είναι να τοποθετήσουμε τα δεδομένα σε δύο στήλες, μία για καθεμία από τις δύο μεταβλητές, ή να υποδείξουμε στο SPSS, λαμβάνοντας υπόψη τη διάταξη που φαίνεται στον Πίνακα διασταύρωσης 3.2, τις αντίστοιχες μετρήσεις για τις δύο μεταβλητές. Η πρώτη επιλογή φαντάζει πάρα πολύ δύσκολη, αφού σε κάθε στήλη θα πρέπει να πληκτρολογήσουμε 29 τιμές, δηλαδή τις τιμές 1, 2 και 3 για τη μεταβλητή «διδασκτική προσέγγιση» και τις τιμές 1 και 2 για τη μεταβλητή «φύλο». Η δεύτερη επιλογή, και πιο σύντομη, βασίζεται στη δυνατότητα που μας δίνει το SPSS να σταθμίσουμε περιπτώσεις (weight cases).

		Gender Φύλο		Total
		1 Άνδρες	2 Γυναίκες	
Teaching Διδακτική προσέγγιση	1 Δασκαλοκεντρική προσέγγιση	7	3	10
	2 Μαθητοκεντρική προσέγγιση	3	6	9
	3 Μεικτή προσέγγιση	5	5	10
Total		15	14	29

Πίνακας 3.2 Πίνακας διασταύρωσης για τις μεταβλητές «Φύλο» και «Διδακτική προσέγγιση».

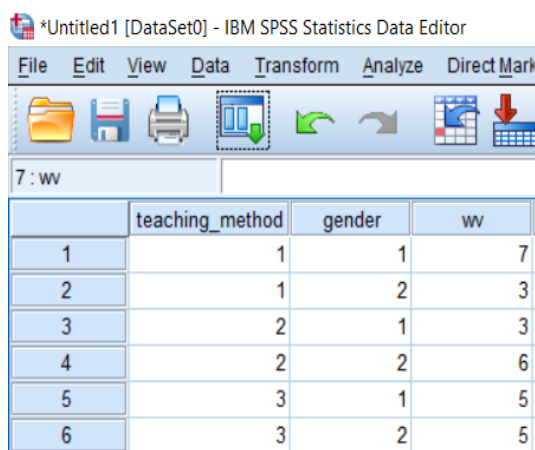
Για την εισαγωγή των δεδομένων στην περίπτωση αυτή, πάλι θα δημιουργήσουμε δύο στήλες, μία για καθεμία από τις δύο μεταβλητές. Η πρώτη στήλη θα αντιστοιχεί στη μεταβλητή «Διδακτική προσέγγιση», που θα παίρνει τιμές 1, 2 και 3, ενώ η δεύτερη στη μεταβλητή «Φύλο», που θα παίρνει τιμές 1 και 2. Επίσης θα πρέπει να δημιουργήσουμε μία τρίτη στήλη που θα αντιστοιχεί στις συχνότητες στάθμισης, όπως αυτές διακρίνονται στον πίνακα διασταύρωσης, που φαίνεται παραπάνω. Με αυτόν τον τρόπο θα κατασκευαστεί ο Πίνακας 3.3 που ακολουθεί.

Η διαδικασία γίνεται ως εξής:

- Κατασκευάζουμε στο SPSS το αρχείο «στάθμιση_περιπτώσεων.sav» με τις μεταβλητές «teaching method» (διδασκτική προσέγγιση), «gender» (φύλο) και «wn» (συχνότητες) (Εικόνα 3.9). Δίνουμε σε αυτές τις μεταβλητές τις τιμές που φαίνονται στον Πίνακα 3.3
- Επιλέγουμε από το μενού **Data => Weight Cases** και ανοίγει το αντίστοιχο πλαίσιο διαλόγου.
- Επιλέγουμε **Weight cases by**.
- Τέλος, επιλέγουμε από αριστερά τη μεταβλητή στάθμισης «wn» και την εισάγουμε δεξιά, στο πλαίσιο **Frequency Variable**, όπως φαίνεται στην Εικόνα 3.10.

Διδακτική προσέγγιση	Φύλο	Συχνότητες
1	1	7
1	2	3
2	1	3
2	2	6
3	1	5
3	2	5

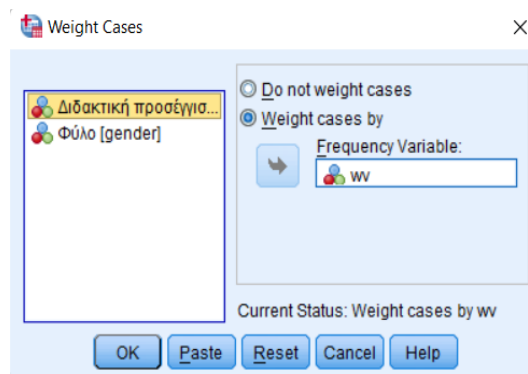
Πίνακας 3.3 Πίνακας διασταύρωσης με τις συχνότητες στάθμισης για τις μεταβλητές «Φύλο» και «Διδακτική προσέγγιση».



*Untitled1 [DataSet0] - IBM SPSS Statistics Data Editor

	teaching_method	gender	wn
1	1	1	7
2	1	2	3
3	2	1	3
4	2	2	6
5	3	1	5
6	3	2	5

Εικόνα 3.9 Τιμές (συχνότητες) της μεταβλητής στάθμισης «wn».



Εικόνα 3.10 Πλαίσιο διαλόγου «Weight Cases».

Αν θέλουμε να υπολογίσουμε τον πίνακα συχνοτήτων της μεταβλητής «διδασκτική προσέγγιση» και εκτελέσουμε την εντολή **Frequencies** (βλ. Κεφ. 6), θα εμφανιστεί ο παρακάτω Πίνακας.

Διδασκτική προσέγγιση		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	1 Δασκαλοκεντρική προσέγγιση	10	34,5	34,5	34,5
	2 Μαθητοκεντρική προσέγγιση	9	31,0	31,0	65,5
	3 Μεικτή προσέγγιση	10	34,5	34,5	100,0
	Total	29	100,0	100,0	

Πίνακας 3.4 Πίνακας συχνοτήτων της μεταβλητής «διδασκτική προσέγγιση» μετά τη στάθμιση των περιπτώσεων με τη μεταβλητή «wn».

Όπως παρατηρούμε στον Πίνακα 3.4 στη στήλη Frequency, η «δασκαλοκεντρική προσέγγιση» εμφανίζεται 10 φορές, η «μαθητοκεντρική προσέγγιση» εμφανίζεται 9 φορές και η «μεικτή προσέγγιση» εμφανίζεται 10 φορές. Αυτό συμβαίνει γιατί, σύμφωνα με τις συχνότητες του Πίνακα 3.3, η διδασκτική προσέγγιση με τον αριθμό 1, δηλαδή η «δασκαλοκεντρική προσέγγιση», θα επαναληφθεί $7 + 3 = 10$ φορές, η διδασκτική προσέγγιση με τον αριθμό 2, δηλαδή η «μαθητοκεντρική προσέγγιση», θα επαναληφθεί $3 + 6 = 9$ φορές, ενώ η διδασκτική προσέγγιση με τον αριθμό 3, δηλαδή η «μεικτή προσέγγιση», θα επαναληφθεί $5 + 5 = 10$ φορές.

3.1.6 Συγχώνευση αρχείων – Merge Files

Μερικές φορές σε μια ομαδική έρευνα, επειδή συνεργάζονται κάποιοι ερευνητές, ο καθένας από αυτούς πιθανόν να εισαγάγει τα δικά του δεδομένα στο SPSS. Με αυτόν τον τρόπο προκύπτουν περισσότερα του ενός αρχεία δεδομένων. Τα αρχεία αυτά πρέπει να ενωθούν μεταξύ τους (Γναρδέλλης, 2013; Bryman, 2016; IBM SPSS Statistics Base 25, n.d.).

Η συγχώνευση αυτή γίνεται με δύο τρόπους:

- Τα αρχεία περιέχουν τις ίδιες μεταβλητές, αλλά έχουν διαφορετικές περιπτώσεις. Στην περίπτωση αυτή θα ενωθούν οι περιπτώσεις (add cases).
- Τα αρχεία περιέχουν τις ίδιες περιπτώσεις, αλλά έχουν διαφορετικές μεταβλητές. Στην περίπτωση αυτή θα ενωθούν οι μεταβλητές (add variables).

3.1.6.1 Merge Files – Add Cases

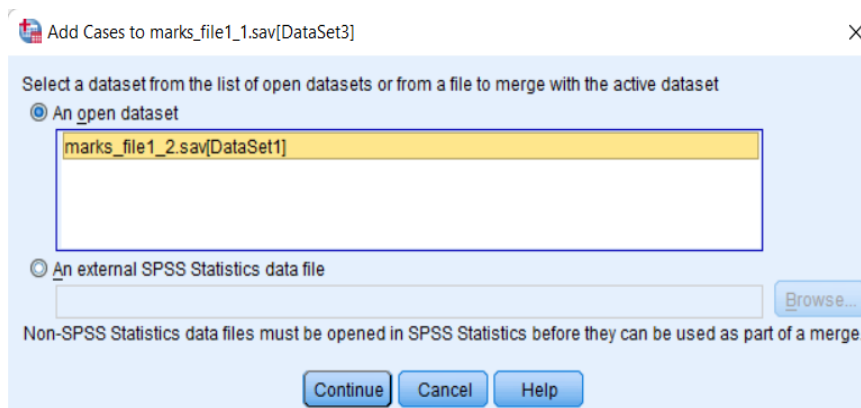
Με την επιλογή **Add Cases** συγχωνεύεται το ενεργό αρχείο δεδομένων με ένα δεύτερο αρχείο δεδομένων ή ένα εξωτερικό αρχείο. Το δεύτερο αρχείο περιέχει τις ίδιες μεταβλητές με το πρώτο αρχείο, δηλαδή τις στήλες, αλλά διαφορετικές περιπτώσεις, δηλαδή γραμμές.

Αν οι μεταβλητές οποιουδήποτε αρχείου δεδομένων έχουν διαφορετικά ονόματα (unpaired variables), τότε το νέο αρχείο που θα συγχωνευτεί θα περιέχει μόνο τις μεταβλητές που έχουν το ίδιο όνομα και στα δύο αρχεία. Υπάρχει η δυνατότητα να δημιουργήσουμε ζεύγη δύο μεταβλητών που έχουν διαφορετικά ονόματα και να τις συμπεριλάβουμε στο νέο, συγχωνευμένο αρχείο.

Οι αριθμητικές μεταβλητές (Numeric) δεν μπορούν να συγχωνευτούν με αλφαριθμητικές μεταβλητές (String), ακόμα και αν έχουν το ίδιο όνομα. Επίσης, δεν μπορούν να συγχωνευτούν αλφαριθμητικές μεταβλητές άνισου πλάτους (Width). Το καθορισμένο πλάτος μίας αλφαριθμητικής μεταβλητής πρέπει να είναι το ίδιο και στα δύο αρχεία δεδομένων.

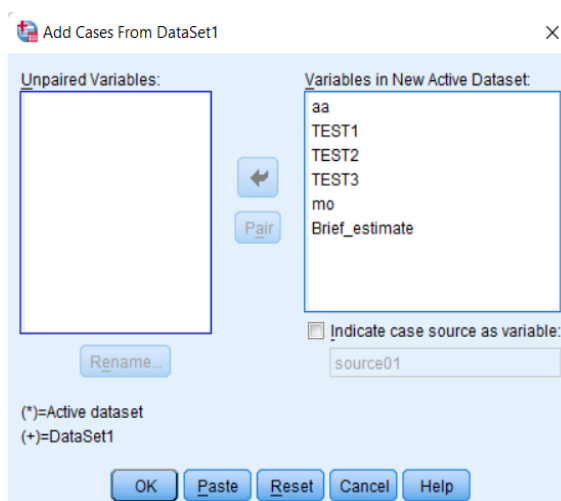
Ας υποθέσουμε ότι θέλουμε να συγχωνεύσουμε το αρχείο «marks_file1_1.sav» με το αρχείο «marks_file1_2.sav». Η διαδικασία της συγχώνευσης των δύο αρχείων με τη μέθοδο **Add Cases** θα γίνει ως εξής:

- Ανοίγουμε τα δύο αρχεία που πρόκειται να συγχωνεύσουμε.
- Επιλέγουμε από το μενού **Data => Merge Files => Add Cases** και ανοίγει το αντίστοιχο πλαίσιο διαλόγου (Εικόνα 3.11).



Εικόνα 3.11 Πλαίσιο διαλόγου «Add Cases to».

- Επιλέγουμε το δεύτερο αρχείο που θα συγχωνευτεί με το πρώτο αρχείο (αρχείο «marks_file1_2.sav») και κάνουμε κλικ στο κουμπί **Continue**. Στη συνέχεια, ανοίγει το παρακάτω πλαίσιο διαλόγου (Εικόνα 3.12).



Εικόνα 3.12 Πλαίσιο διαλόγου «Add Cases From».

- Στο δεξί μέρος, στη λίστα **Variables in New Active DataSet**, φαίνονται οι μεταβλητές που θα περιληφθούν στο νέο αρχείο δεδομένων. Μπορούμε να μην επιλέξουμε κάποιες από αυτές, επιλέγοντάς τις και πατώντας το κουμπί (βελάκι). Στο παράδειγμά μας, κρατάμε όλες τις μεταβλητές.
- Στο αριστερό μέρος, στη λίστα **Unpaired Variables**, θα εμφανιστούν, εάν υπάρχουν, οι μεταβλητές που περιλαμβάνονται μόνο σε ένα από τα δύο αρχεία που πρόκειται να συγχωνευτούν. Στις μεταβλητές που υπάρχουν στο ενεργό αρχείο, δεξιά από το όνομά τους, εμφανίζεται το σύμβολο (*), ενώ στις μεταβλητές που υπάρχουν στο δεύτερο αρχείο, δεξιά από το όνομά τους, εμφανίζεται το σύμβολο (+). Εδώ, όλες οι μεταβλητές περιλαμβάνονται και στα δύο αρχεία, για αυτό και η λίστα είναι κενή.
- Από τη λίστα **Unpaired Variables** μπορούμε να επιλέξουμε ταυτόχρονα τις δύο μεταβλητές (πατώντας το πλήκτρο **Ctrl**) και έπειτα να πατήσουμε το κουμπί **Pair**, ώστε το ζεύγος των μεταβλητών να εισαχθεί στο συγχωνευμένο αρχείο ως μία μεταβλητή. Το όνομα της νέας μεταβλητής θα είναι το ίδιο με το όνομα της μεταβλητής που υπάρχει στο ενεργό αρχείο. Μπορούμε να μετονομάσουμε τις μεταβλητές του ζεύγους πατώντας το κουμπί **Rename**.
- Η επιλογή **Indicate case source as variable** μάς δίνει τη δυνατότητα να κατασκευάσουμε στο συγχωνευμένο αρχείο μια καινούρια μεταβλητή, η οποία θα δηλώνει την πηγή του αρχείου από το οποίο προέρχονται οι περιπτώσεις της. Δηλαδή, η νέα μεταβλητή «πηγή» θα έχει την τιμή 0, όταν οι περιπτώσεις της υπάρχουν στο ενεργό αρχείο, και την τιμή 1, όταν οι περιπτώσεις της υπάρχουν στο δεύτερο αρχείο.
- Η διαδικασία της συγχώνευσης ολοκληρώνεται κάνοντας κλικ στο κουμπί **OK**.

Όσον αφορά την αποθήκευση του συγχωνευμένου αρχείου, αυτό γίνεται με την ονομασία «συγχωνευμένο_αρχείο_με_add_cases_marks_file_1&2».

Εναλλακτικά, προκειμένου να συγχωνεύσουμε δύο αρχεία δεδομένων που έχουν όμως τις ίδιες ακριβώς μεταβλητές, μπορούμε να αντιγράψουμε και να επικολλήσουμε τα δεδομένα από το ένα αρχείο στο άλλο.

3.1.6.2 Merge Files – Add Variables

Με την επιλογή **Add Variables** συγχωνεύεται το ενεργό αρχείο δεδομένων με ένα δεύτερο αρχείο δεδομένων ή ένα εξωτερικό αρχείο. Το δεύτερο αρχείο περιέχει τις ίδιες περιπτώσεις με το πρώτο αρχείο, δηλαδή τις γραμμές, αλλά διαφορετικές μεταβλητές, δηλαδή στήλες. Αυτή η συγχώνευση είναι χρήσιμη όταν, για παράδειγμα, θέλουμε να ενώσουμε ένα αρχείο δεδομένων, που περιέχει τις μετρήσεις ενός ελέγχου πριν μια διδακτική παρέμβαση, με ένα άλλο αρχείο, που περιέχει τις μετρήσεις ενός ελέγχου μετά τη διδακτική παρέμβαση.

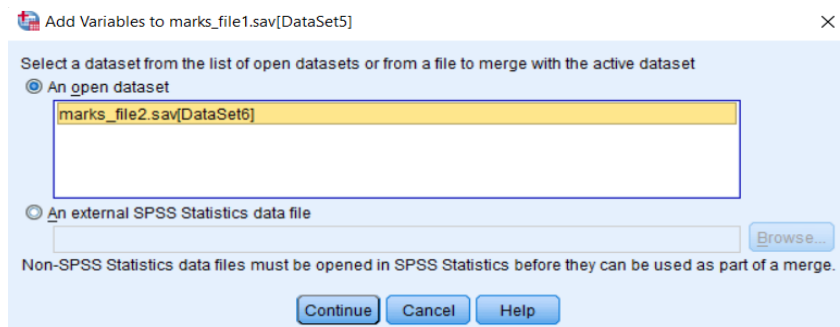
Για να γίνει σωστά η συγχώνευση των αρχείων πρέπει οι περιπτώσεις τους να είναι ταξινομημένες με την ίδια σειρά και στα δύο αρχεία. Επιπλέον, πρέπει να έχουν τις ίδιες περιπτώσεις, που αντιστοιχούν στα υποκείμενα της έρευνας. Αν, για παράδειγμα, τα υποκείμενα της έρευνας είναι φοιτητές, πρέπει οι μετρήσεις για κάθε περίπτωση (γραμμή) των δύο αρχείων να αντιστοιχούν στον ίδιο φοιτητή.

Η ταύτιση των περιπτώσεων των δύο αρχείων μπορεί να πραγματοποιηθεί με τη χρήση μίας ή και περισσότερων μεταβλητών-κλειδιά. Βασική προϋπόθεση είναι ότι η μεταβλητή-κλειδί πρέπει να υπάρχει και στα δύο αρχεία, να είναι δηλαδή κοινή. Επιπλέον, πρέπει να προσέξουμε, πριν από τη συγχώνευση, τα δεδομένα των δύο αρχείων να είναι ταξινομημένα κατά αύξουσα σειρά ως προς την κοινή μεταβλητή-κλειδί.

Όταν τα ονόματα των μεταβλητών του δεύτερου αρχείου είναι ίδια με τα ονόματα των μεταβλητών του ενεργού, δηλαδή του πρώτου αρχείου, εξαιρούνται οι μεταβλητές αυτές από τη συγχώνευση επειδή η μέθοδος της προσθήκης μεταβλητών (**Add Variables**) προϋποθέτει ότι αυτές οι μεταβλητές περιέχουν διπλότητες πληροφορίας.

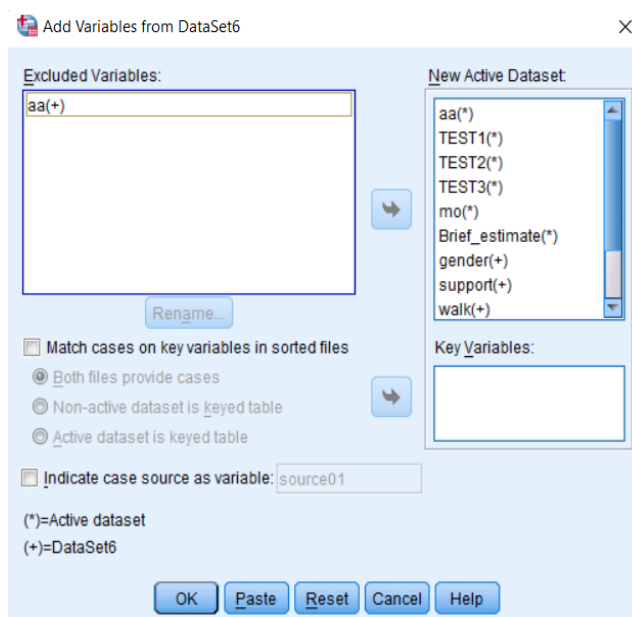
Ας υποθέσουμε ότι θέλουμε να συγχωνεύσουμε τις μεταβλητές του αρχείου «marks_file1_1.sav» με τις μεταβλητές του αρχείου «marks_file1_2.sav». Η διαδικασία της συγχώνευσης των δύο αρχείων με τη μέθοδο **Add Variables** θα γίνει ως εξής:

- Ανοίγουμε τα δύο αρχεία που πρόκειται να συγχωνεύσουμε.
- Επιλέγουμε από το μενού **Data => Merge Files => Add Variables** και ανοίγει το αντίστοιχο πλαίσιο διαλόγου (Εικόνα 3.13).



Εικόνα 3.13 Πλαίσιο διαλόγου «Add Variables to».

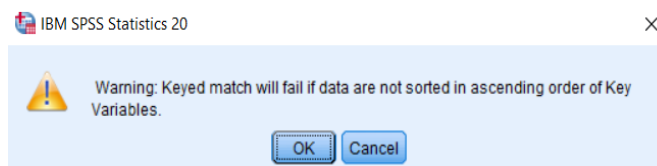
- Επιλέγουμε το δεύτερο αρχείο που θα συγχωνευτεί με το πρώτο αρχείο (αρχείο «marks_file1_2.sav») και κάνουμε κλικ στο κουμπί **Continue**. Στη συνέχεια, ανοίγει το παρακάτω πλαίσιο διαλόγου (Εικόνα 3.14).



Εικόνα 3.14 Πλαίσιο διαλόγου «Add Variables from».

- Στο δεξί μέρος, στη λίστα **New Active DataSet**, φαίνονται οι μεταβλητές που θα περιληφθούν στο νέο αρχείο δεδομένων. Στη λίστα αυτή περιέχονται μόνο οι μεταβλητές, τα ονόματα των οποίων υπάρχουν μόνο σε ένα από τα δύο αρχεία. Στο παράδειγμά μας υπάρχουν οι μεταβλητές από το πρώτο αρχείο, οι οποίες έχουν το σύμβολο (*), και οι μεταβλητές από το δεύτερο αρχείο, οι οποίες έχουν το σύμβολο (+). Παρατηρούμε ότι οι μεταβλητές αυτές δεν είναι κοινές.
- Στο αριστερό μέρος, στη λίστα **Excluded Variables**, θα εμφανιστούν οι μεταβλητές που εξαιρούνται από τη συγχώνευση επειδή υπάρχουν και στα δύο αρχεία. Εδώ υπάρχει η κοινή μεταβλητή «aa(+)». Μπορούμε να μετονομάσουμε κάποια μεταβλητή του δεύτερου αρχείου που υπάρχει και στο πρώτο αρχείο, ώστε να εισαχθεί στο αρχείο που θα συγχωνευτεί, κάνοντας κλικ στην επιλογή **Rename**.
- Κάτω από τη λίστα **New Active DataSet** υπάρχει η λίστα **Key Variables**. Εδώ θα εισάγουμε τη ή τις μεταβλητές-κλειδιά με τις οποίες θα γίνει η ακριβής ταύτιση των περιπτώσεων των δύο αρχείων. Όπως αναφέρθηκε, οι μεταβλητές-κλειδιά πρέπει να υπάρχουν και στα δύο αρχεία με το ίδιο όνομα και πριν από τη συγχώνευση, ενώ τα δεδομένα των δύο αρχείων να είναι ταξινομημένα κατά αύξουσα σειρά ως προς τις κοινές μεταβλητές-κλειδιά. Αν δεν έχει γίνει η αύξουσα ταξινόμηση, τότε το SPSS εμφανίζει το παρακάτω προειδοποιητικό μήνυμα (Εικόνα 3.15). Στο παράδειγμά μας, και στα δύο αρχεία περιέχεται η μεταβλητή «aa», η οποία και ορίζεται ως η μεταβλητή-κλειδί. Πριν όμως γίνει αυτό, αν προσέξουμε το αρχείο

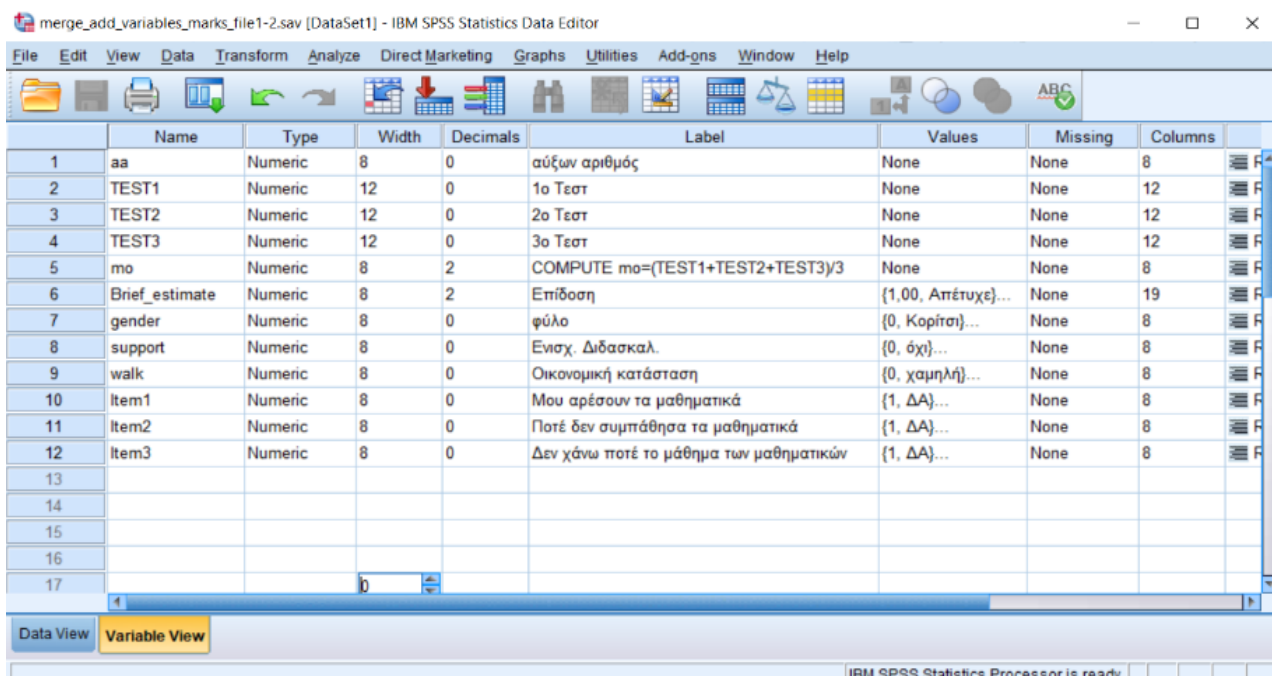
«marks_file1_2.sav», θα παρατηρήσουμε ότι η μεταβλητή «aa» είναι ταξινομημένη κατά φθίνουσα και όχι κατά αύξουσα σειρά. Η ταξινόμηση της μεταβλητής «aa» γίνεται από το μενού **Data => Sort Cases**, όπως αναφέρθηκε στην Ενότητα 3.1. Επιλέγουμε από αριστερά τη μεταβλητή «aa» (αύξων αριθμός), την εισάγουμε στο πεδίο **Sort by** και στο πλαίσιο **Sort Order** επιλέγουμε **Ascending**.



Εικόνα 3.15 Προειδοποιητικό μήνυμα ότι δεν έχει γίνει αύξουσα ταξινόμηση των δεδομένων στις μεταβλητές-κλειδιά.

- Πριν εισάγουμε τη μεταβλητή «aa» στη λίστα **Key Variables**, πρέπει να τσεκάρουμε την επιλογή **Match cases on key variables in sorted files** ώστε να ενεργοποιηθεί η λίστα **Key Variables** (Εικόνα 3.16). Οι επιλογές **Non-active dataset is keyed table** και **Active dataset is keyed table** αφορούν το αν είναι ενεργός ή όχι ένας συγκεντρωτικός πίνακας δεδομένων (keyed table).
- Αφού εισαγάγουμε τη μεταβλητή «aa» στη λίστα **Key Variables**, κάνουμε κλικ στο κουμπί **OK** για να ολοκληρωθεί η διαδικασία της συγχώνευσης των δύο αρχείων.
- Μόλις ολοκληρωθεί η διαδικασία της συγχώνευσης των δύο αρχείων με τη μέθοδο **Add Variables**, προκύπτει ένα τελικό συγχωνευμένο αρχείο, το οποίο περιέχει συνολικά τις μεταβλητές των δύο αρχείων (Εικόνα 3.16).

Όσον αφορά την αποθήκευση του συγχωνευμένου αρχείου, αυτό γίνεται με την ονομασία «συγχωνευμένο_αρχείο_με_add_variables_marks_file_1&2».



Εικόνα 3.16 Το αποτέλεσμα της συγχώνευσης των δύο αρχείων με τη μέθοδο «Merge-Add Variables».

Βιβλιογραφία

Bryman, A. (2016). *Social research methods*. London: Oxford University Press.

Γναρδέλλης, Χ. (2013). *Ανάλυση δεδομένων με το IBM SPSS Statistics 21*. Αθήνα: Εκδ. Παπαζήση.

IBM SPSS Statistics Base 25 (n.d.). Retrieved from

https://faculty.ksu.edu.sa/sites/default/files/ibm_spss_statistics_base.pdf

Κεφάλαιο 4

Περιγραφική στατιστική στο SPSS

Σύνοψη

Στο κεφάλαιο αυτό γίνεται αρχικά μια σύντομη παρουσίαση των βασικών στατιστικών που περιγράφουν την κατανομή μιας μεταβλητής. Στη συνέχεια, με την αξιοποίηση του SPSS υπολογίζονται και εξάγονται πίνακες συχνοτήτων, καθώς και περιγραφικά στατιστικά μιας ποσοτικής μεταβλητής. Τέλος, το παρόν κεφάλαιο περιλαμβάνει τη σύνθετη παρουσίαση των περιγραφικών στατιστικών δύο μεταβλητών.

Προαπαιτούμενη γνώση

Βασικές γνώσεις του περιβάλλοντος του SPSS για την εισαγωγή και την οργάνωση των δεδομένων: Κεφάλαια 1, 2, και 3 του συγγράμματος.

4.1 Εισαγωγή στην περιγραφική στατιστική ανάλυση

Η περιγραφική στατιστική ανάλυση είναι πολύ σημαντική στην αρχή μιας ανάλυσης. Αποτελεί μια συνοπτική παρουσίαση των δεδομένων της έρευνας και μας επιτρέπει να αναδείξουμε το βασικό μήνυμα, δηλαδή τα βασικά σημεία, αλλά και να εντοπίσουμε τα χαρακτηριστικά που αναδεικνύονται, ώστε να μπορούμε να τα εξετάσουμε σε σχέση με άλλες περιπτώσεις. Επίσης, με την περιγραφή αυτή δείχνουμε σε ποιο σημείο έχουμε φτάσει σχετικά με την κατανόησή μας, αφού δεν μπορούμε να δείξουμε την κατανόησή μας χωρίς να την περιγράψουμε.

Αν βρεθούμε αντιμέτωποι με μια μάζα δεδομένων, δεν μπορούμε εύκολα να την αφομοιώσουμε. Η άμεση απάντησή μας είναι να περιγράψουμε τα δεδομένα αυτά πιο απλοποιημένα. Ένας τρόπος για να γίνει αυτό είναι η «συμπύεση» των τιμών σε δείκτες (περιγραφικά στατιστικά), που θα δούμε στη συνέχεια του Κεφαλαίου. Αυτό θα επέτρεπε τη σύγκριση των δεικτών και επομένως τη σύντομη επεξεργασία των δεδομένων, αλλά και τον προσδιορισμό μιας στρατηγικής ανάλυσης των δεδομένων. Στην κατεύθυνση αυτή, θα αναζητηθούν σύνδεσμοι μεταξύ των μεταβλητών που θα μας επιτρέπουν την εξήγηση του ερευνητικού φαινομένου (Γναρδέλλης, 2013; Bryman, 2016; IBM SPSS Statistics Base 25, n.d.). Για παράδειγμα, η σύγκριση των περιγραφικών στατιστικών της ακαδημαϊκής εμπλοκής αγοριών και κοριτσιών μπορεί να μας οδηγήσει στη διερεύνηση της υπόθεσης «σχέση του φύλου με την ακαδημαϊκή εμπλοκή των φοιτητών».

Τα περιγραφικά στατιστικά, που μας βοηθούν να εξετάσουμε σύντομα την κατανομή των τιμών μιας μεταβλητής, κυρίως είναι ο πίνακας συχνοτήτων (frequency table) και τα περιγραφικά στατιστικά, όπως τα μέτρα θέσης ή κεντρικής τάσης (measure of central tendency), τα μέτρα διασποράς (measure of dispersion) και τα μέτρα μορφής (measure of shape).

4.1.1 Πίνακες συχνοτήτων

Πριν ξεκινήσουμε να παρουσιάζουμε την κατανομή των τιμών μιας μεταβλητής, κρίνεται προτιμότερο να κάνουμε πρώτα την αντιστοίχιση των δυνατών τιμών σε πραγματικούς αριθμούς (Κωδικοποίηση), (δείτε Κεφάλαιο 1). Η κωδικοποίηση αφορά μόνο τις ποιοτικές μεταβλητές. Για τις ονομαστικές μεταβλητές, οι αριθμοί που χρησιμοποιούνται είναι αυθαίρετοι, για παράδειγμα, για τη μεταβλητή «φύλο» (1 για Αγόρι και 2 για Κορίτσι). Όσον αφορά τις μεταβλητές διάταξης, οι αριθμοί θα πρέπει να διατάσσονται παρόμοια με τις τιμές της κλίμακας, που τις έχουμε αντιστοιχίσει. Για παράδειγμα, αν σε μια δήλωση πεποιθήσεων ο ερωτώμενος έπρεπε να απαντήσει σε μια κλίμακα τεσσάρων σημείων, η αντιστοίχιση θα μπορούσε να είναι: 1 Καθόλου, 2 Μέτρια, 3 Πολύ και 4 Πάρα πολύ.

Τα δεδομένα μιας μεταβλητής που προκύπτουν από κάποιο πείραμα ή έρευνα, είναι τις περισσότερες φορές συλλογές αριθμών, όπου σε αυτή τη μορφή πολλά ενδιαφέροντα ερωτήματα είναι δύσκολο να απαντηθούν. Για παράδειγμα, «ποια είναι η μεγαλύτερη και ποια η μικρότερη τιμή;» «ποια τιμή έχει τη

συχνότερη εμφάνιση;» «πόσες τιμές βρίσκονται κάτω από αυτόν τον αριθμό;» «πόσες ανήκουν σ' εκείνο το διάστημα τιμών;» κ.λπ. Υπάρχει συνεπώς η αναγκαιότητα μιας οργανωμένης παρουσίασης των αριθμών, ώστε να ληφθούν ευκολότερα οι πληροφορίες που μας ενδιαφέρουν. Η κατανομή συχνοτήτων απαντώντας σε αυτή την ανάγκη αποτελεί τον βασικό τρόπο μελέτης των τιμών μιας μεταβλητής. Ο πίνακας συχνοτήτων των τιμών μιας μεταβλητής μας πληροφορεί για το ποιες είναι οι διακριτές τιμές της μεταβλητής και πώς αυτές κατανέμονται.

Για τη δημιουργία του πίνακα συχνοτήτων θα πρέπει να υπολογίσουμε τα παρακάτω:

1. Απόλυτη συχνότητα (frequency) ή σκέτο συχνότητα μιας τιμής x_i (το i παίρνει τιμές από 1 έως και k , όπου k οι δυνατές τιμές της τυχαίας μεταβλητής). Είναι ο αριθμός επαναλήψεων της τιμής της μεταβλητής. Επομένως, το άθροισμα των συχνοτήτων των τιμών της μεταβλητής είναι ίσο με το μέγεθος του δείγματος (συμβολίζεται με n), δηλαδή με το σύνολο των τιμών της μεταβλητής.
2. Σχετική συχνότητα (relative frequency) μιας τιμής x_i είναι ο λόγος της συχνότητας της τιμής προς το μέγεθος του δείγματος. Επομένως, i) οι σχετικές συχνότητες είναι μικρότερες ή ίσες της μονάδας (μετακινώντας δύο θέσεις δεξιά την υποδιαστολή γίνονται ποσοστά), και ii) το άθροισμα των σχετικών συχνοτήτων των τιμών της μεταβλητής είναι ίσο με τη μονάδα.
3. Απόλυτη αθροιστική συχνότητα (cumulative frequency) ή σκέτο αθροιστική συχνότητα μιας τιμής x_i είναι το άθροισμα των συχνοτήτων των τιμών που είναι μικρότερες ή ίσες με την τιμή αυτή. Επομένως, i) η απόλυτη αθροιστική συχνότητα της μικρότερης τιμής ισούται με την απόλυτη συχνότητα της τιμής, και ii) η απόλυτη αθροιστική συχνότητα της μεγαλύτερης τιμής ισούται με το μέγεθος του δείγματος.
4. Αθροιστική σχετική συχνότητα (cumulative relative frequency) μιας τιμής x_i είναι το άθροισμα των σχετικών συχνοτήτων των τιμών που είναι μικρότερες ή ίσες με την τιμή αυτή. Επομένως, i) η σχετική αθροιστική συχνότητα της μικρότερης τιμής ισούται με τη σχετική συχνότητα της τιμής, και ii) η σχετική αθροιστική συχνότητα της μεγαλύτερης τιμής ισούται με τη μονάδα.

Αξίζει να τονίσουμε ότι η παρουσίαση της κατανομής των τιμών μιας μεταβλητής μέσω ενός πίνακα συχνοτήτων αφορά και ποιοτικές (ονομαστικής κλίμακας ή ιεραρχικής κλίμακας) και ποσοτικές μεταβλητές. Ωστόσο, στην περίπτωση της ποσοτικής μεταβλητής που οι τιμές της έχουν μικρή συχνότητα εμφάνισης, δεν προσφέρει μια σύντομη αποτύπωση της κατανομής της. Στην περίπτωση αυτή, καλύτερα να ομαδοποιούνται οι τιμές της ποσοτικής μεταβλητής σε *ισοπλατείς* ή *ανισοπλατείς κλάσεις* (μπορείτε να δείτε πώς μπορεί να γίνει αυτό με το SPSS στο Κεφάλαιο 2, ενότητα 2.1.2 «Επανακωδικοποίηση στις ίδιες μεταβλητές – Recode into Same Variables»). Τέλος, τόσο η απόλυτη αθροιστική συχνότητα όσο και η αθροιστική σχετική συχνότητα έχει νόημα να υπολογιστούν και να παρουσιαστούν μόνο όταν οι τιμές της μεταβλητής διατάσσονται.

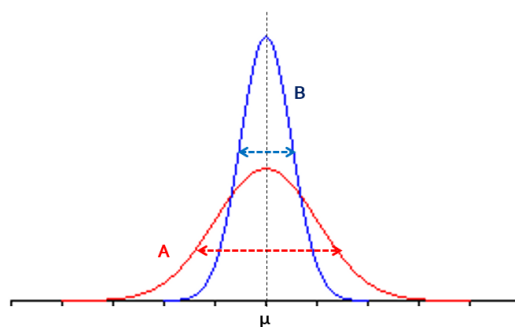
Για παράδειγμα, στον Πίνακα 4.1 παρουσιάζουμε την κατανομή συχνοτήτων των βαθμών των φοιτητών ενός παιδαγωγικού τμήματος στο μάθημα της γλώσσας. Όπως μπορείτε να δείτε, στην έρευνα αυτή πήραν μέρος 200 φοιτητές (η αθροιστική συχνότητα του μεγαλύτερου βαθμού). Επιπλέον, κάποιες σημαντικές πληροφορίες που εξάγονται άμεσα από τον πίνακα και μπορούν να συνοδεύσουν τη σύντομη παρουσίαση της κατανομής των τιμών της μεταβλητής είναι: α) ο βαθμός με τη μεγαλύτερη συχνότητα (60) είναι το 6, β) η πλειονότητα των φοιτητών ($74\% = 100\% - 26\%$) πέρασε το μάθημα (με βαθμό τουλάχιστον 5) και γ) το 14% ($100\% - 86\%$) των φοιτητών πήρε βαθμό τουλάχιστον 8.

Βαθμοί	Συχνότητα	Σχετική συχνότητα (%)	Αθροιστική Συχνότητα	Αθροιστική σχετική συχνότητα (%)
3	12	6,0%	12	6,0%
4	40	20,0%	52	26,0%
5	50	25,0%	102	51,0%
6	60	30,0%	162	81,0%
7	10	5,0%	172	86,0%
8	18	9,0%	190	95,0%
9	8	4,0%	198	99,0%
10	2	1,0%	200	100,0%
Σύνολο	200	100,0%		

Πίνακας 4.1 Κατανομή συχνοτήτων της επίδοσης των φοιτητών ενός παιδαγωγικού τμήματος στο μάθημα της γλώσσας.

Μια κατανομή των τιμών μιας ποσοτικής μεταβλητής που θα μας απασχολήσει πολύ στα επόμενα είναι η κανονική κατανομή (normal distribution). Η κατανομή αυτή είναι πολύ σημαντική αφού πολλά φαινόμενα που συμβαίνουν γύρω μας κατανέμονται κατά προσέγγιση κανονικά ή ακολουθούν την κανονική κατανομή. Χαρακτηριστικά όπως: το ύψος, το βάρος, η βαθμολογία σε διαγώνισμα, κ.λπ. αλλά και τυχαία σφάλματα που εμφανίζονται κατά τη διάρκεια διάφορων μετρήσεων, ακολουθούν κατά προσέγγιση την κανονική κατανομή. Επίσης, όπως θα δούμε στη συνέχεια, η κανονική κατανομή αποτελεί θεωρητικό υπόβαθρο της στατιστικής συμπερασματολογίας ή της επαγωγικής στατιστικής.

Στην περίπτωση της κανονικής κατανομής ο πίνακας συχνοτήτων των τιμών της μεταβλητής είναι συμμετρικά κατανεμημένος γύρω από μια κλάση τιμών στην οποία θα βρίσκεται η μέση τιμή. Έτσι το ιστόγραμμα (γραφική απεικόνιση συχνοτήτων τιμών μιας ομαδοποιημένης ποσοτικής μεταβλητής) αυτής της μεταβλητής θα έχει άξονα συμμετρίας τον ιστό στον οποίο αντιστοιχεί η μεσαία κλάση. Επομένως αν θεωρητικά ο αριθμός των κλάσεων για μια συνεχή μεταβλητή είναι αρκετά μεγάλος και το πλάτος των κλάσεων είναι αρκετά μικρό (πολύ κοντά στο μηδέν), τότε η πολυγωνική γραμμή συχνοτήτων «πολύγωνο συχνοτήτων» (που οριοθετεί το ιστόγραμμα) τείνει να πάρει τη μορφή μιας ομαλής καμπύλης που μοιάζει με καμπάνα. Στην εικόνα 4.1 φαίνεται ότι η κατανομή της βαθμολογίας των μαθητών σε δύο διαφορετικά τμήματα (Α και Β) είναι κανονική. Πιο συγκεκριμένα στα δύο τμήματα η μέση (μ) βαθμολογία είναι ίδια, και οι τιμές της βαθμολογίας των μαθητών γύρω από τη μέση τιμή κατανέμονται μεν συμμετρικά αλλά με διαφορετική τυπική απόκλιση (σ).



Εικόνα 4.1 Η μορφή της κανονικής κατανομής.

4.1.2 Περιγραφικά στατιστικά

Τα Περιγραφικά στατιστικά χρησιμοποιούνται για τη συνοπτική παρουσίαση των τιμών μιας ποσοτικής μεταβλητής, συμπληρωματικά ή και μαζί με τον πίνακα συχνοτήτων. Τα μέτρα αυτά είναι πολύ χρήσιμα στην περιγραφική στατιστική ανάλυση. Αν υπολογίζονται για το δείγμα λέγονται *στατιστικά* ή *δείκτες*, ενώ αν υπολογίζονται για τον πληθυσμό λέγονται *παράμετροι* (Γναρδέλλης 2013; Bryman, 2016; IBM SPSS Statistics Base 25, n.d.).

Τα μέτρα αυτά είναι:

1. Μέτρα θέσης ή κεντρικής τάσης

Μία από τις ιδιότητες μιας ομάδας αριθμών είναι η *κεντρική τάση*. Είναι η αριθμητική έκφραση της τάσης των τιμών σχετικά με τη συγκέντρωση γύρω από μια τιμή της κλίμακας μέτρησής τους. Χρησιμοποιείται για να εκφράσει την περιοχή μεγαλύτερης συγκέντρωσης των τιμών, η οποία περιοχή κατά κάποιο τρόπο, αντιπροσωπεύει την ομάδα. Τα μέτρα αυτά παίρνουν τιμές εντός του εύρους των τιμών της μεταβλητής. Τα κυριότερα μέτρα κεντρικής τάσης είναι: i) η *Μέση τιμή* (Mean), ii) η *Διάμεσος* (Median), *Εκατοστιαίες τιμές* (Percentiles) και τα *Τεταρτημόρια* (Quartiles), και iii) η *Επικρατούσα* (Mode) ή *Δεσπόζουσα* τιμή.

Η *Μέση τιμή* (μ) έχει μεγάλη εφαρμογή για περαιτέρω στατιστική ανάλυση και ορίζεται ως το άθροισμα των τιμών, διαιρεμένο από τον αριθμό των τιμών της ομάδας. Για τον υπολογισμό της, χρησιμοποιούνται όλες οι τιμές της μεταβλητής, επομένως, επηρεάζεται πολύ από τις ακραίες τιμές της μεταβλητής. Η μέση τιμή δεν έχει νόημα να υπολογιστεί για ποιοτικά δεδομένα. Ένα περισσότερο ανθεκτικό μέτρο θέσης είναι ο *περικομμένος μέσος* ($\kappa\%$ trimmed mean). Υπολογίζεται αν αφαιρέσουμε τις $\kappa\%$ αριστερά (μικρές τιμές) και τις $\kappa\%$ δεξιά (μεγάλες τιμές) και για τις υπόλοιπες τιμές αν υπολογίσουμε τη μέση τιμή.

Η *Διάμεσος* είναι μια τιμή τέτοια, ώστε ο αριθμός των παρατηρήσεων που είναι μεγαλύτερες απ' αυτήν να είναι ίσος με τον αριθμό των παρατηρήσεων που είναι μικρότερες απ' αυτήν. Αυτός ο ορισμός εκφράζει

ισοδύναμα ότι η διάμεσος διαιρεί την κατανομή μιας μεταβλητής σε δύο ίσα μέρη. Η διάμεσος μπορεί να μην αντιστοιχεί σε δυνατή τιμή της μεταβλητής και δεν υπολογίζεται σε ονομαστική ποιοτική μεταβλητή. Οι Εκατοστηαίες τιμές ή Εκατοστημόρια (Percentiles) αποτελούν γενίκευση της έννοιας της διαμέσου. Ως *κ-εκατοστηαία τιμή* ή *Ρκ εκατοστημόριο* ενός συνόλου παρατηρήσεων ονομάζουμε την τιμή εκείνη για την οποία, αν οι μετρήσεις διαταχθούν από τη μικρότερη στη μεγαλύτερη, το πολύ κ% των παρατηρήσεων είναι μικρότερες ή ίσες του Ρκ και το πολύ (100-κ)% των παρατηρήσεων είναι μεγαλύτερες ή ίσες από την τιμή αυτή. Τα εκατοστημόρια P25, P50, P75, καλούνται *τεταρτημόρια* (quartiles) και συμβολίζονται με Q1 (αριστερά της τιμής αυτής υπάρχει το πολύ 25% των παρατηρήσεων και δεξιά το πολύ 75%), Q2 = P50 είναι η διάμεσος, και Q3 (αριστερά της τιμής αυτής υπάρχει το πολύ 75% των παρατηρήσεων και δεξιά το πολύ 25% των παρατηρήσεων).

Η *Επικρατούσα τιμή*, είναι η συχνότερη τιμή στα δεδομένα. Η τιμή αυτή δεν αξιοποιείται εύκολα για περαιτέρω στατιστική ανάλυση και δεν ορίζεται πάντα μονοσήμαντα. Μπορούμε να έχουμε πολλές επικρατούσες τιμές ή και καμία. Όταν όλες οι τιμές έχουν την ίδια συχνότητα, δεν υπάρχει επικρατούσα τιμή. Αντίστοιχα όταν δύο τιμές έχουν την ίδια μεγαλύτερη συχνότητα τότε η κατανομή χαρακτηρίζεται ως *δικόρυφη*. Η επικρατούσα τιμή, υπολογίζεται και σε ποσοτικές αλλά και σε ποιοτικές μεταβλητές.

2. Μέτρα διασποράς ή μεταβλητότητας

Οι τιμές μιας μεταβλητής συνήθως βρίσκονται διασκορπισμένες μέσα σ' ένα διάστημα της κλίμακας, έχουν δηλαδή μια μεταβλητότητα, που άλλες φορές είναι μικρή και άλλες φορές είναι μεγάλη. Στόχος της έρευνας είναι να εξηγήσουμε τη μεταβλητότητα αυτή των δεδομένων και η στατιστική μάς βοηθά να την μελετήσουμε. Ως *μεταβλητότητα* (variability) της μέτρησης ορίζουμε τον βαθμό «απλώματος» των δεδομένων γύρω από τη μέση τιμή. Τα Μέτρα μεταβλητότητας (Variability) ή διασποράς (Dispersion) είναι απαραίτητα για τον προσδιορισμό της μεταβλητότητας που παρουσιάζει ένα σύνολο μετρήσεων. Τα μέτρα αυτά απαντούν σε ερωτήματα του πόσο όμοια ή ανόμοια είναι τα μέλη μιας ομάδας, λαμβάνοντας υπόψη τις μετρήσεις αυτές. Επίσης, αξιοποιούνται σε συνδυασμό με τα μέτρα θέσης και περιγράφουν συμπληρωματικά τις κατανομές των δεδομένων. Βασικά μέτρα μεταβλητότητας ή διασποράς είναι: i) το *Εύρος* (Range), ii) το *Ενδοτεταρτημοριακό εύρος* (Interquartile Range), και iii) η *Διακύμανση* (Variance) και η *Τυπική απόκλιση* (Standard Deviation).

Εύρος (Range) είναι η διαφορά μεταξύ της μέγιστης (max) και της ελάχιστης τιμής (min) της κατανομής. Μπορεί να χρησιμοποιηθεί για την εκτίμηση της τυπικής απόκλισης (περίπου το 1/6 του εύρους), στην περίπτωση που οι τιμές της μεταβλητής ακολουθούν την κανονική κατανομή. Ωστόσο, δεν μας λέει τίποτα για τη διασπορά των τιμών της κατανομής γύρω από τον μέσο όρο και δεν χρησιμοποιείται στην περιγραφική στατιστική.

Ενδοτεταρτημοριακό εύρος (Interquartile Range) αποτελεί το εύρος του μεσαίου τμήματος της κατανομής (50% των τιμών μιας κατανομής) και υποδεικνύει το «άπλωμα» των μεσαίων παρατηρήσεων. Προσδιορίζεται μέσω των *τεταρτημορίων* (quartiles): Interquartile Range: IQR= Q3-Q1. Μεγάλες τιμές IQR σημαίνει ότι το 1ο και 3ο τεταρτημόριο απέχουν, υποδεικνύοντας υψηλό επίπεδο μεταβλητότητας. Το ενδοτεταρτημοριακό εύρος είναι ένας καλός εκπρόσωπος των μεσαίων τιμών της κατανομής και δεν επηρεάζεται από τις ακραίες τιμές της. Ωστόσο, δεν περιγράφει κανέναν από τους δείκτες, οι οποίοι είναι βασικοί για την περιγραφική στατιστική.

Η *Διακύμανση* και η *Τυπική απόκλιση* (σ) είναι δύο σχεδόν ταυτόσημες έννοιες που περιγράφουν τη μεταβλητότητα. Ορίζουμε ως διακύμανση ενός πληθυσμού N τιμών με μέση τιμή μ , τη μέση τετραγωνική απόκλιση (απόσταση) των N μετρήσεων από τη μέση τιμή μ του πληθυσμού. Ως μέτρο μεταβλητότητας συνήθως χρησιμοποιείται η Τυπική Απόκλιση (Standard Deviation) δηλαδή, η τετραγωνική ρίζα της διακύμανσης. Η τυπική απόκλιση έχει την ίδια μονάδα μέτρησης με τις τιμές της μεταβλητής. Μικρές τιμές της τυπικής απόκλισης υποδεικνύουν μικρότερη μεταβλητότητα γύρω από τη μέση τιμή. Η τυπική απόκλιση χρησιμοποιείται στην περιγραφική στατιστική και μπορεί να χρησιμοποιηθεί για την εκτίμηση των παραμέτρων του πληθυσμού. Ωστόσο, είναι πολύ ευαίσθητη στις ακραίες τιμές της κατανομής.

Αναφορικά με την τυπική απόκλιση και τη μέση τιμή, αν η κατανομή των τιμών μιας ποσοτικής μεταβλητής είναι κατά προσέγγιση κανονική, μπορούμε κατά προσέγγιση να προσδιορίσουμε τι ποσοστό των τιμών βρίσκονται σε συγκεκριμένα διαστήματα τιμών. Συγκεκριμένα, το 68% των τιμών βρίσκονται στο διάστημα $[\mu-\sigma, \mu+\sigma]$, το 95% στο διάστημα $[\mu-2\sigma, \mu+2\sigma]$ και το 99,7% στο διάστημα $[\mu-3\sigma, \mu+3\sigma]$. Το γεγονός αυτό, λαμβάνοντας υπόψη ότι η κανονική κατανομή έχει άξονα συμμετρίας τη μέση τιμή (μ), μας επιτρέπει να έχουμε μια καλή εικόνα του πώς κατανέμονται οι τιμές της κατανομής. Για παράδειγμα, αν το ύψος των μαθητών κατανέμεται κατά προσέγγιση κανονικά και έχει $\mu=167$ cm και $\sigma=7$ cm, τότε γνωρίζουμε ότι περίπου το 2,35% των μαθητών έχουν ύψος στο διάστημα $[\mu+2\sigma, \mu+3\sigma]=[167+2*7, 167+3*7]=[181, 188]$.

3. Μέτρα μορφής

Όλες οι καμπύλες που περιγράφουν την κατανομή των τιμών μιας συνεχούς μεταβλητής δεν έχουν τη μορφή της κανονικής κατανομής (καμπάνα). Οι κυριότερες διαφορές εδράζονται στην έλλειψη συμμετρίας της καμπύλης αλλά και στον τρόπο συσσώρευσης των τιμών γύρω από τη μέση τιμή. Δύο συντελεστές που μας βοηθούν να κατανοούμε καλύτερα κατά πόσο αποκλίνει η καμπύλη των τιμών της μεταβλητής μας από την καμπύλη της κανονικής κατανομής, είναι οι συντελεστές: i) *Ασυμμετρίας* (Skewness), και ii) *Κύρτωσης* (Kurtosis).

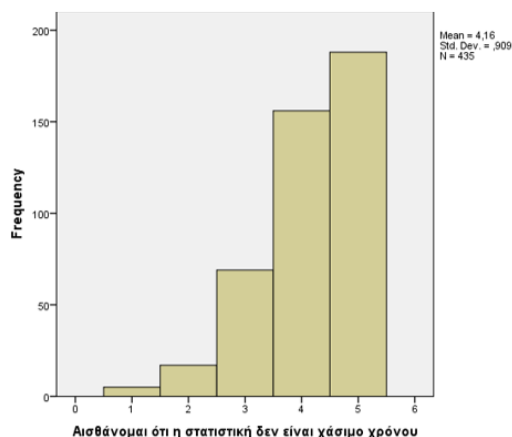
Όταν ο συντελεστής *ασυμμετρίας* είναι αρνητικός (Εικόνα 4.2), τότε οι τιμές της κατανομής είναι τραβηγμένες αριστερά, δηλαδή, η κατανομή είναι αρνητικά ασύμμετρη. Όταν ο συντελεστής είναι θετικός, τότε οι τιμές της κατανομής είναι τραβηγμένες δεξιά, δηλαδή, η κατανομή είναι θετικά ασύμμετρη. Για παράδειγμα, στην περίπτωση μέτρησης των στάσεων, όταν μια δήλωση είναι θετικά διατυπωμένη (π.χ. «Αισθάνομαι ότι η στατιστική δεν είναι χάσιμο χρόνου», με δυνατές απαντήσεις τύπου λίκερτ, Διαφωνώ απόλυτα έως και Συμφωνώ Απόλυτα), τότε οι δυνατές απαντήσεις κυρίως κινούνται στις υψηλές τιμές της κλίμακας και μόνο λίγες αποτυπώνονται χαμηλά, δηλαδή η κατανομή τραβιέται από αριστερά και επομένως έχει αρνητική ασυμμετρία (skewness=-1,007) (Εικόνα 4.2).

Ο συντελεστής *κύρτωσης* προσδιορίζει τον τρόπο συγκέντρωσης των τιμών της κατανομής γύρω από τη μέση τιμή. Όταν ο συντελεστής είναι θετικός, τότε οι τιμές είναι περισσότερο συγκεντρωμένες γύρω από τη μέση τιμή και η κατανομή αποκαλείται *λεπτόκυρτη*. Όταν ο συντελεστής είναι αρνητικός, τότε οι τιμές είναι περισσότερο διάσπαρτες και η κατανομή αποκαλείται *πλατύκυρτη*.

Στην ενότητα 6.3.1 «Έλεγχος της κανονικότητας των τιμών μιας ποσοτικής μεταβλητής» παρουσιάζονται αναλυτικά διάφοροι τρόποι ελέγχου της κατανομής των τιμών μιας ποσοτικής μεταβλητής.

Αισθάνομαι ότι η στατιστική δεν είναι χάσιμο χρόνου				
	Frequency	Percent	Valid Percent	Cumulative Percent
Valid 1	5	1,1	1,1	1,1
2	17	3,9	3,9	5,1
3	69	15,9	15,9	20,9
4	156	35,9	35,9	56,8
5	188	43,2	43,2	100,0
Total	435	100,0	100,0	

Descriptive Statistics					
	N	Skewness		Kurtosis	
	Statistic	Statistic	Std. Error	Statistic	Std. Error
Αισθάνομαι ότι η στατιστική δεν είναι χάσιμο χρόνου	435	-1,007	,117	,701	,234
Valid N (listwise)	435				



Εικόνα 4.2 Κατανομή συχνοτήτων με αρνητική ασυμμετρία.

4.1.3 Κατασκευή πίνακα κατανομής συχνοτήτων στο περιβάλλον του SPSS

Για τα παραδείγματα αυτού του κεφαλαίου θα χρησιμοποιηθεί το αρχείο δεδομένων του SPSS «marks_trans_2.sav». Στο αρχείο αυτό διερευνώνται οι σχέσεις του φύλου μαθητών Λυκείου (μεταβλητή «gender»), της ενισχυτικής διδασκαλίας (μεταβλητή «support»), της οικονομικής κατάστασης των μαθητών (μεταβλητή «walk») και της κατεύθυνσης σπουδών (μεταβλητή «stream») με την αντιλαμβανόμενη ικανότητα των μαθητών στα μαθηματικά (μεταβλητές «Item1», «Item2» και «Item3»).

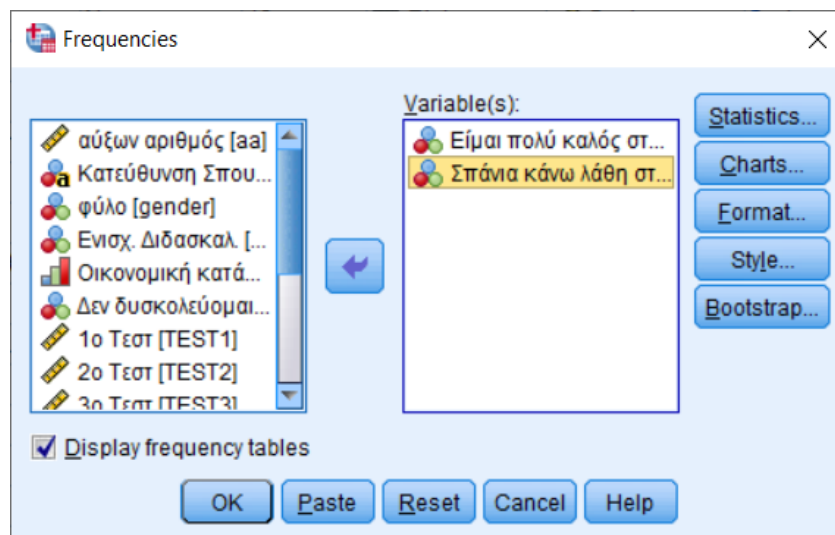
Η κωδικοποίηση των κατηγορικών μεταβλητών γίνεται ως εξής:

- Μεταβλητή «φύλο»: 0 για το κορίτσι και 1 για το αγόρι.
- Με τη μεταβλητή «support» καταγράφεται αν οι μαθητές παρακολουθούν την ενισχυτική διδασκαλία ή όχι. Οπότε παίρνουμε δύο τιμές: 0 για το Όχι και 1 για το Ναι.
- Με τη μεταβλητή «walk» καταγράφεται η οικονομική κατάσταση των μαθητών. Έχουμε τρεις τιμές: 0 - χαμηλή, 1 - μέτρια και 2 - καλή οικονομική κατάσταση.

- Με τη μεταβλητή «stream» καταγράφεται η κατεύθυνση σπουδών των μαθητών (θετική, θεωρητική και τεχνολογική). Είναι μια αλφαριθμητική μεταβλητή (τύπου string).
- Η μεταβλητή «Item1» αντιστοιχεί στη δήλωση «Είμαι πολύ καλός στα μαθηματικά». Η μεταβλητή «Item2» αντιστοιχεί στη δήλωση «Δυσκολεύομαι στα μαθηματικά». Η μεταβλητή «Item3» αντιστοιχεί στη δήλωση «Σπάνια κάνω λάθη στα μαθηματικά». Οι τρεις αυτές μεταβλητές παίρνουν τις τιμές 1 - διαφωνώ απόλυτα, 2 - διαφωνώ, 3 - Ούτε συμφωνώ ούτε διαφωνώ, 4 – συμφωνώ, 5 - συμφωνώ απόλυτα.

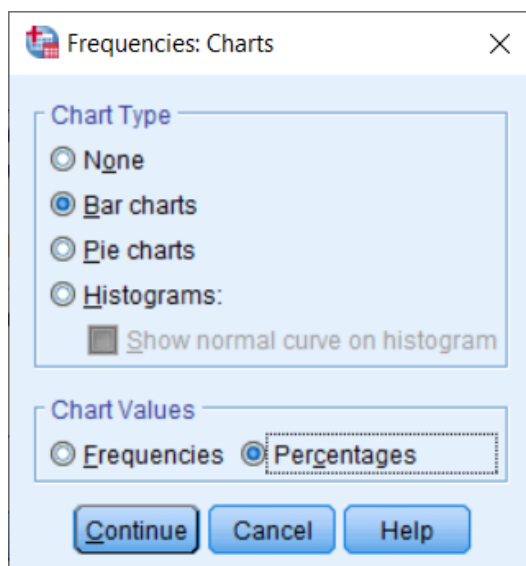
Αξιοποιώντας τα δεδομένα αυτά θέλουμε να κατασκευάσουμε πίνακες κατανομής συχνοτήτων για τις μεταβλητές: «Item1» (Είμαι πολύ καλός στα μαθηματικά) και «Item3» (Σπάνια κάνω λάθη στα μαθηματικά).

Από το μενού επιλέγουμε διαδοχικά: **Analyze =>Descriptive statistics=>Frequencies**. Στο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 4.3) επιλέγουμε τις δύο μεταβλητές διαδοχικά και τις τοποθετούμε, πατώντας το δεξί βέλος ή σύροντάς τες, στην περιοχή **Variable(s)**. Προσοχή, στο ότι εξ ορισμού η επιλογή **Display frequency tables** είναι επιλεγμένη. Αυτό σημαίνει ότι θα παραχθεί και θα εμφανιστεί ένας πίνακας συχνοτήτων.



Εικόνα 4.3 Πλαίσιο διαλόγου «Frequencies».

Αν θέλουμε ταυτόχρονα να υπολογιστούν και να εμφανιστούν και συγκεκριμένα περιγραφικά στατιστικά, τότε θα επιλέξουμε **Statistics**, όπου εκεί έχουμε τη δυνατότητα να ζητήσουμε την εμφάνιση διάφορων περιγραφικών στατιστικών. Αναλυτικότερη παρουσίαση αυτών μπορείτε να δείτε στις επόμενες ενότητες του Κεφαλαίου. Στην περίπτωση που επιθυμούμε ταυτόχρονα με τον πίνακα κατανομής συχνοτήτων να εμφανιστεί και αντίστοιχο γράφημα, επιλέγουμε **Charts** όπως φαίνεται και στο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 4.4). Έχουμε επίσης τη δυνατότητα να δημιουργήσουμε ραβδόγραμμα (**Bar charts**), κυκλικό διάγραμμα (**Pie charts**) και Ιστόγραμμα (**Histograms**) με την ταυτόχρονη εμφάνιση και της καμπύλης της κανονικής κατανομής (**Show normal curve on histogram**). Επιπλέον, στην επιλογή **Chart values**, μπορούμε να ζητήσουμε η κατανομή στο γράφημα να βασιστεί στις συχνότητες (**Frequencies**) ή στις σχετικές συχνότητες % (**Percentages**). Να τονίσουμε ότι τόσο τα ραβδογράμματα όσο και τα κυκλικά διαγράμματα χρησιμοποιούνται κυρίως για ποιοτικές μεταβλητές. Επίσης, η χρήση των ραβδογραμμάτων ενδείκνυται κυρίως σε ποιοτικές μεταβλητές ιεραρχικής κλίμακας και όταν οι συχνότητες ή οι σχετικές συχνότητες της κατανομής των τιμών έχουν μικρές διαφορές. Τέλος, το ιστόγραμμα χρησιμοποιείται στις ποσοτικές μεταβλητές. Στην περίπτωσή μας, επειδή οι μεταβλητές αυτές είναι διατάξιμες επιλέξαμε αρχικά το ραβδόγραμμα να είναι βασισμένο στις σχετικές συχνότητες % (**Percentages**) και έπειτα πατήσαμε **Continue**.



Εικόνα 4.4 Πλαίσιο διαλόγου «Frequencies:Charts».

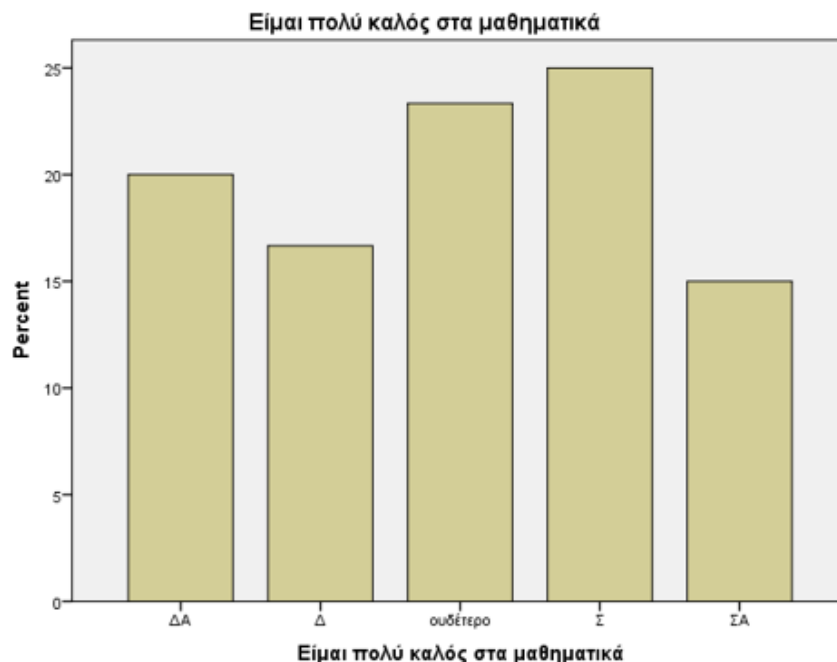
Παρακάτω, θα σχολιαστούν μόνο ο πίνακας συχνοτήτων και το ραβδόγραμμα για την πρώτη μεταβλητή. Οι εξαγόμενοι πίνακες συχνοτήτων (Εικόνα 4.5) αποτελούνται από τον πίνακα που συνοψίζει τον αριθμό των μετρήσεων με τις έγκυρες τιμές (valid) και χαμένες ή ελλείπουσες τιμές (missing) και για τις δύο μεταβλητές. Τονίζουμε ότι υπάρχουν δύο ειδών χαμένες τιμές: α) όταν δεν έχει δοθεί απάντηση από τον ερωτώμενο, δηλαδή υπάρχει κενό, τότε η τιμή χαρακτηρίζεται ως missing system, και β) όταν έχει οριστεί από εμάς κάποια τιμή ως ελλείπουσα, δηλαδή έχουμε δηλώσει την τιμή αυτή ως τιμή που δεν πρέπει να ληφθεί υπόψη (για παράδειγμα, μια πολύ ακραία τιμή της μεταβλητής που θέλουμε να αποκλειστεί από την ανάλυση), χαρακτηρίζεται ως missing τοποθετώντας δίπλα την/τις τιμή/-ές που έχουμε αποκλείσει.

Ο δεύτερος πίνακας αποτελείται από πέντε στήλες. Στην πρώτη παρουσιάζονται οι τιμές της μεταβλητής ταξινομημένες ως προς τον κωδικό (αριθμό που αντιστοιχήσαμε σε κάθε κατηγορία), καθώς και στην περίπτωση που υπάρχουν χαμένες τιμές υπάρχει η ένδειξη *missing*. Στη δεύτερη στήλη (Frequency) παρουσιάζονται οι συχνότητες των τιμών της μεταβλητής αλλά και των ελλειπουσών τιμών. Στην τρίτη στήλη (Percent) παρουσιάζονται οι σχετικές συχνότητες εκφρασμένες σε ποσοστά των τιμών της μεταβλητής και των ελλειπουσών τιμών. Στην τέταρτη στήλη (Valid Percent) παρουσιάζονται τα έγκυρα ποσοστά, δηλαδή οι σχετικές συχνότητες εκφρασμένες σε ποσοστά των τιμών της μεταβλητής χωρίς να λαμβάνονται υπόψη οι ελλείπουσες τιμές. Στην περίπτωση αυτή για τον υπολογισμό των σχετικών συχνοτήτων, το μέγεθος του δείγματος βασίζεται μόνο στις έγκυρες τιμές. Τέλος, στην πέμπτη στήλη (Cumulative Percent) παρουσιάζονται οι αθροιστικές σχετικές συχνότητες εκφρασμένες σε ποσοστά των έγκυρων τιμών (χωρίς της ελλείπουσες τιμές) της μεταβλητής. Υπενθυμίζουμε ότι η στήλη αυτή έχει νόημα μόνο όταν οι τιμές της μεταβλητής διατάσσονται.

Τέλος, στο ραβδόγραμμα που εμφανίζεται παρουσιάζονται οι σχετικές συχνότητες εκφρασμένες σε ποσοστά των τιμών της μεταβλητής.

Statistics			
		Είμαι πολύ καλός στα μαθηματικά	Σπάνια κάνω λάθη στα μαθηματικά
N	Valid	60	59
	Missing	1	2

Είμαι πολύ καλός στα μαθηματικά					
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	1 ΔΑ	12	19,7	20,0	20,0
	2 Δ	10	16,4	16,7	36,7
	3 ουδέτερο	14	23,0	23,3	60,0
	4 Σ	15	24,6	25,0	85,0
	5 ΣΑ	9	14,8	15,0	100,0
	Total	60	98,4	100,0	
Missing	System	1	1,6		
Total		61	100,0		



Εικόνα 4.5 Πίνακες αποτελεσμάτων «Frequencies».

4.1.3.1 Περιγραφικά στατιστικά μέσω της διαδικασίας «Frequencies»

Αξιοποιώντας το αρχείο δεδομένων του SPSS «marks_trans_2.sav», θέλουμε να υπολογίσουμε τα περιγραφικά στατιστικά για τη μεταβλητή «mo». Η μεταβλητή «mo» ($mo = (TEST1 + TEST2 + TEST3) / 3$) αποτελεί μια ποσοτική σύνοψη της επίδοσης των μαθητών στα τρία τεστ.

Από το μενού επιλέγουμε διαδοχικά: **Analyze=>Descriptive statistics=>Frequencies**. Στο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 4.3) επιλέγουμε τη μεταβλητή και την τοποθετούμε, πατώντας το δεξί βέλος ή εναλλακτικά σύροντάς τη, στην περιοχή **Variable(s)**. Προσοχή, στο ότι εξ ορισμού η επιλογή **Display frequency tables** είναι επιλεγμένη. Αφού η μεταβλητή είναι ποσοτική δεν έχει νόημα να εξαχθεί και πίνακας συχνοτήτων. Επομένως, αφαιρούμε την επιλογή αυτή κάνοντας κλικ πάνω στο κουτάκι. Στη συνέχεια, επιλέγουμε **Statistics** και στο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 4.6) μπορούμε να ζητήσουμε τον υπολογισμό και επομένως την εμφάνιση διάφορων περιγραφικών στατιστικών.

Τα βασικά στατιστικά που συνήθως επιλέγουμε είναι: **Mean** (Μέση τιμή), **Median** (Διάμεσος), **Minimum** (ελάχιστη τιμή), **Maximum** (Μέγιστη τιμή), και **Standard deviation** (Τυπική απόκλιση). Κάποια επιπλέον στατιστικά που μπορούμε να ζητήσουμε είναι: **Quartiles** (τεταρτημόρια), **Mode** (επικρατούσα τιμή), **Variance**

(διακύμανση), **Percentiles** (εκατοστημόρια), **Skewness** (συντελεστής ασυμμετρίας), και **Kurtosis** (συντελεστής κύρτωσης). Αφού επιλέξουμε τα βασικά στατιστικά, πατάμε **Continue** και τέλος **OK**.

Εικόνα 4.6 Πλαίσιο διαλόγου «Frequencies Statistics».

Το αποτέλεσμα της παραπάνω διαδικασίας είναι ένας πίνακας όπου τα στατιστικά είναι τοποθετημένα κατακόρυφα (Εικόνα 4.7). Αρχικά παρουσιάζεται ο αριθμός των έγκυρων τιμών στη συνέχεια ο αριθμός των χαμένων τιμών και στη συνέχεια η μέση τιμή, η διάμεσος, η τυπική απόκλιση, η ελάχιστη και η μέγιστη τιμή.

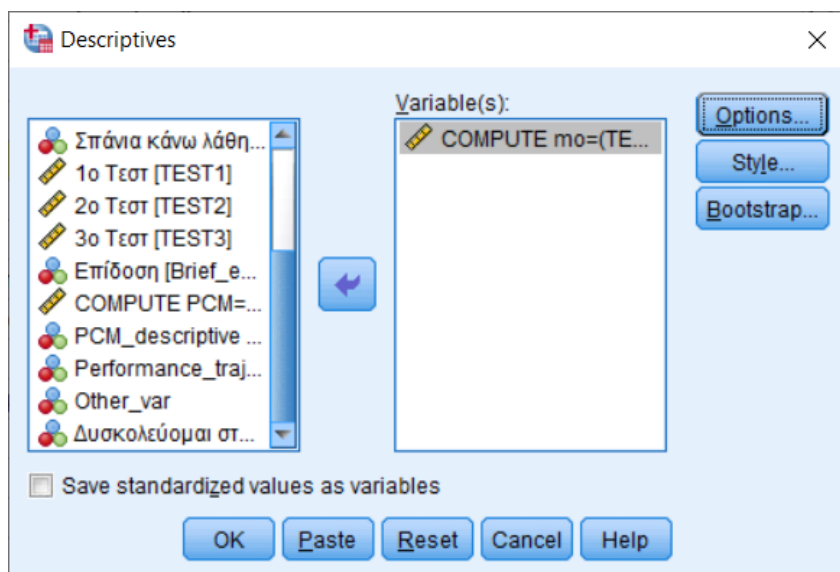
Statistics		
COMPUTE mo = (TEST1+TEST2+TEST3) /3		
N	Valid	61
	Missing	0
Mean		74,1585
Median		77,6667
Std. Deviation		13,24185
Minimum		35,67
Maximum		93,00

Εικόνα 4.7 Πίνακας της διαδικασίας «Frequencies Statistics».

4.1.3.2 Περιγραφικά στατιστικά μέσω της διαδικασίας «Descriptives»

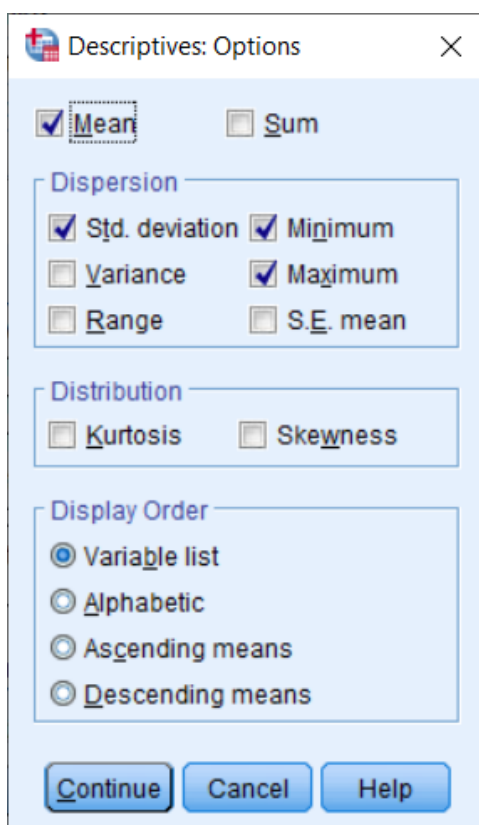
Αξιοποιώντας το αρχείο δεδομένων του SPSS «marks_trans_2.sav», θέλουμε να υπολογίσουμε τα περιγραφικά στατιστικά για τη μεταβλητή «mo». Από το μενού επιλέγουμε διαδοχικά: **Analyze=>Descriptive statistics=>Descriptives**. Στο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 4.8) επιλέγουμε τη μεταβλητή και την τοποθετούμε, πατώντας το δεξί βέλος ή σύροντάς τη στην περιοχή **Variable(s)**.

Προσοχή, στην περίπτωση που χρειάζεται να υπολογίσουμε τις τυπικές τιμές (standardized values) των τιμών της μεταβλητής, δημιουργώντας καινούριες μεταβλητές, επιλέγουμε στο κάτω μέρος και αριστερά την επιλογή **Save standardized values as variables**. Οι τυπικές τιμές, όπως θα δούμε και στη συνέχεια (ενότητα 6.3 Έλεγχος δεδομένων (Data screening) με το SPSS), είναι πολύ χρήσιμες στην υπόδειξη των ακραίων τιμών.



Εικόνα 4.8 Πλαίσιο διαλόγου «Descriptives».

Στη συνέχεια επιλέγουμε **Options** και στο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 4.9) μπορούμε να ζητήσουμε τον υπολογισμό και επομένως την εμφάνιση διάφορων περιγραφικών στατιστικών. Όπως φαίνεται, τα στατιστικά που μπορούμε να επιλέξουμε είναι πολύ περιορισμένα σε σχέση με τη δυνατότητα που είχαμε μέσω της επιλογής «frequencies». Τα βασικά στατιστικά που επιλέγουμε είναι: **Mean** (Μέση τιμή), **Minimum** (ελάχιστη τιμή), **Maximum** (Μέγιστη τιμή), και **Standard deviation** (Τυπική απόκλιση).



Εικόνα 4.9 Πλαίσιο διαλόγου «Descriptives:Options».

Το αποτέλεσμα της παραπάνω διαδικασίας είναι ένας πίνακας όπου τα στατιστικά είναι τοποθετημένα οριζόντια (Εικόνα 4.10). Αρχικά παρουσιάζεται ο αριθμός των έγκυρων τιμών και στη συνέχεια η ελάχιστη, η μέγιστη τιμή, η μέση τιμή και η τυπική απόκλιση.

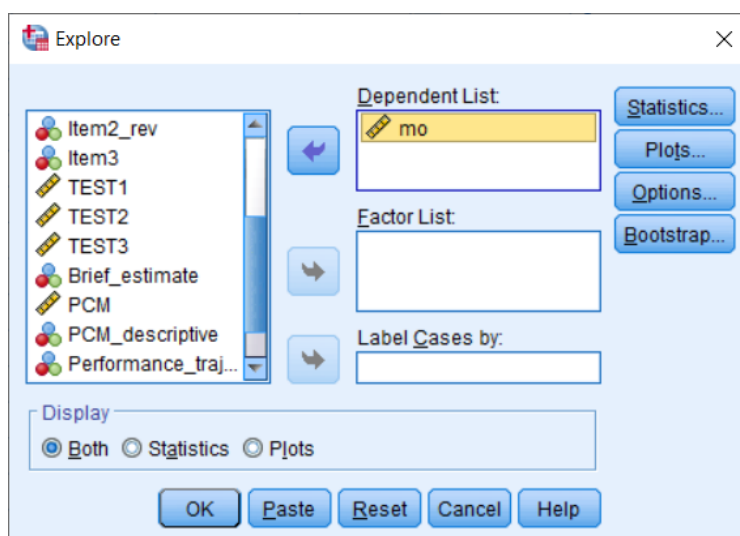
Descriptive Statistics					
	N	Minimum	Maximum	Mean	Std. Deviation
COMPUTE mo = (TEST1+TEST2+TEST3)/3	61	35,67	93,00	74,1585	13,24185
Valid N (listwise)	61				

Εικόνα 4.10 Πίνακας που εξάγει η διαδικασία «Descriptives».

4.1.3.3 Περιγραφικά στατιστικά μέσω της διαδικασίας «Explore»

Η διαδικασία **Explore** παράγει λεπτομερή στατιστικά στοιχεία και γραφήματα για ποσοτικές μεταβλητές για ένα ολόκληρο δείγμα ή για υποσύνολα ενός δείγματος, όπως αυτά διακρίνονται μέσω μιας ποιοτικής μεταβλητής. Υπάρχουν πολλοί λόγοι για τη χρήση της διαδικασίας Explore όπως: εξέταση των δεδομένων και έλεγχος προϋποθέσεων (Γναρδέλλης 2013; Bryman, 2016; IBM SPSS Statistics Base 25, n.d.). Η εξέταση των δεδομένων μπορεί να δείξει ότι έχουμε ασυνήθιστες τιμές, ακραίες τιμές, κενά στα δεδομένα ή άλλες ιδιαιτερότητες. Ο έλεγχος των προϋποθέσεων μπορεί να μας βοηθήσει να προσδιορίσουμε τις κατάλληλες στατιστικές τεχνικές που θα χρησιμοποιήσουμε για την ανάλυση των δεδομένων. Μπορούμε, δηλαδή, να τη χρησιμοποιήσουμε για τον έλεγχο της κανονικότητας της κατανομής των τιμών μιας ποσοτικής μεταβλητής με την αξιοποίηση συγκεκριμένων στατιστικών αλλά και διαγνωστικών διαγραμμάτων. Για παράδειγμα, όπως θα δούμε και στην αξιοποίηση παραμετρικών ελέγχων, θα πρέπει να ελέγξουμε αν η κατανομή των τιμών της ποσοτικής μεταβλητής ακολουθεί την κανονική κατανομή. Παρακάτω παρουσιάζεται ο υπολογισμός περιγραφικών στατιστικών με τη διαδικασία Explore και στη συνέχεια παρουσιάζονται διάφοροι τρόποι (και με τη διαδικασία Explore) εξέτασης των δεδομένων και ελέγχου της κανονικότητας.

Αξιοποιώντας το αρχείο δεδομένων του SPSS «marks_trans_2.sav», θέλουμε να υπολογίσουμε τα περιγραφικά στατιστικά για τη μεταβλητή «mo». Από το μενού επιλέγουμε διαδοχικά: **Analyze=>Descriptive statistics=>Explore**. Στο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 4.11) επιλέγουμε τη μεταβλητή και την τοποθετούμε, πατώντας το δεξί βέλος ή σύροντάς τη, στην περιοχή **Dependent list**. Κάτω και δεξιά μπορούμε να επιλέξουμε να εμφανιστούν μόνο αναλυτικά περιγραφικά στατιστικά (**Statistics**) για την/τις ποσοτική/-ές μεταβλητή/-ές, μόνο συγκεκριμένα γραφήματα (**Plots**) αλλά και τα δύο (**Both**). Στην περιοχή **Factor List** τοποθετούμε την/τις ποιοτική/-ές μεταβλητή/-ές που θα επιτρέψουν την παρουσίαση των περιγραφικών στατιστικών της κάθε ποσοτικής μεταβλητής ανά ομάδα, όπως αυτές διακρίνονται από τις τιμές των ποιοτικών μεταβλητών. Στην περιοχή **Label Cases by**, τοποθετούμε την ποιοτική μεταβλητή, που υποδεικνύει την ταυτότητα των υποκειμένων, για παράδειγμα τον αύξοντα αριθμό των περιπτώσεων του δείγματος. Είναι χρήσιμη η υπόδειξη του **Label Cases by** μόνο για την περίπτωση που θα ζητήσουμε την υπόδειξη των ακραίων περιπτώσεων (outliers). Στην περίπτωση που δεν τοποθετήσουμε κάποια μεταβλητή «ταυτότητας των υποκειμένων», το περιβάλλον χρησιμοποιεί για την υπόδειξη των ακραίων περιπτώσεων τον αριθμό γραμμής των περιπτώσεων.



Εικόνα 4.11 Πλαίσιο διαλόγου «Descriptives: Explore».

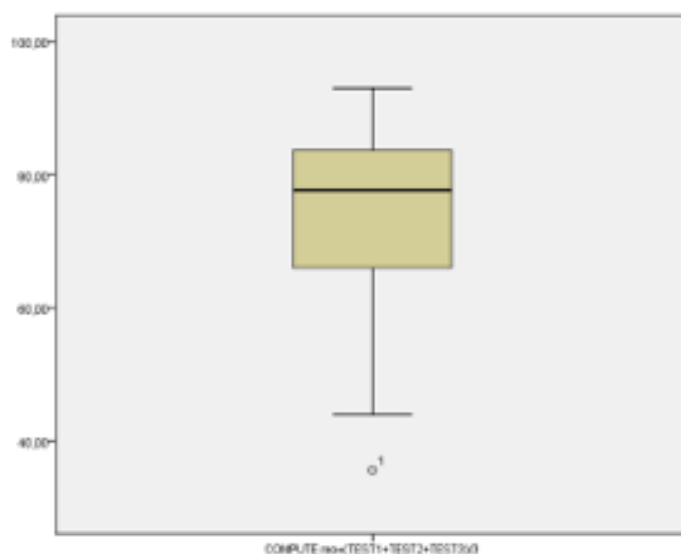
Το αποτέλεσμα της παραπάνω διαδικασίας είναι δύο πίνακες και δύο γραφήματα (Εικόνα 4.12). Στον πρώτο πίνακα, παρουσιάζεται συνοπτικά ο αριθμός των έγκυρων και των ελλειπουσών τιμών. Στον δεύτερο πίνακα, παρουσιάζονται αναλυτικά σχεδόν όλα τα βασικά περιγραφικά στατιστικά της κατανομής των τιμών της ποσοτικής μεταβλητής. Αναφορικά με τα δύο γραφήματα, το πρώτο λέγεται *φυλλογράφημα* ή αλλιώς *γράφημα μίσχου και φύλλου* (Stem and leaf). Στο γράφημα αυτό, παρουσιάζονται οι αναλυτικές τιμές της ποσοτικής μεταβλητής σε δύο μέρη: α) το *στέλεχος* ή *μίσχος* (stem) που αποτελείται από το/τα πρώτο/-α ψηφίο/-α μιας τιμής και β) το *φύλλο* (leaf) που αποτελείται από τα υπόλοιπα ψηφία μιας τιμής. Έτσι παρουσιάζεται σύντομα η κατανομή των τιμών μιας ποσοτικής μεταβλητής, όπως θα την βλέπαμε σε ένα ιστόγραμμα και ταυτόχρονα παρουσιάζονται αναλυτικά όλες οι τιμές της μεταβλητής. Παρατηρώντας το γράφημα αυτό, διαφαίνεται μια αρνητική συμμετρία η οποία υποστηρίζεται και από τον συντελεστή ασυμμετρίας (Skewness= -0,960).

Τέλος, το δεύτερο γράφημα λέγεται *θηκόγραμμα* ή *διάγραμμα πλαισίου και απολήξεων* (box-and whisker plot) και αποτελεί μια σύνοψη των πέντε δεικτών (min, Q1, Q2, Q3, max) και γενικά μια αποτύπωση της διασποράς και των ακραίων τιμών μιας ποσοτικής μεταβλητής. Αποτελείται από ένα ορθογώνιο με βάσεις το πρώτο (Q1) και το τρίτο (Q3) τεταρτημόριο, ενώ ενδιάμεσα τοποθετείται η διάμεσος (δεύτερο τεταρτημόριο Q2). Από τα μέσα των βάσεων αναπτύσσονται γραμμές οι οποίες συνδέουν τις οριακές τιμές (Min & Max) της μεταβλητής. Στην περίπτωση που υπάρχουν ακραίες τιμές, εμφανίζονται πέραν του βασικού κορμού των πέντε δεικτών. Στο συγκεκριμένο γράφημα διαφαίνεται μια ακραία τιμή που αντιστοιχεί στην περίπτωση με αριθμό γραμμής το 1.

Case Processing Summary						
	Cases					
	Valid		Missing		Total	
	N	Percent	N	Percent	N	Percent
COMPUTE mo=(TEST1+TEST2+TEST3)/3	61	100,0%	0	0,0%	61	100,0%

Descriptives					Statistic	Std. Error
COMPUTE mo = (TEST1+TEST2+TEST3)/3	Mean				74,1585	1,69545
	95% Confidence Interval for Mean		Lower Bound		70,7671	
			Upper Bound		77,5499	
	5% Trimmed Mean				74,9183	
	Median				77,6667	
	Variance				175,347	
	Std. Deviation				13,24185	
	Minimum				35,67	
	Maximum				93,00	
	Range				57,33	
	Interquartile Range				17,83	
	Skewness				-,960	,306
	Kurtosis				,401	,604

mo COMPUTE mo = (TEST1+TEST2+TEST3) /3
 COMPUTE mo=(TEST1+TEST2+TEST3)/3 Stem-and-Leaf Plot
 Frequency Stem & Leaf
 1,00 Extremes (= <36)
 3,00 4 . 458
 7,00 5 . 2236789
 6,00 6 . 345667
 20,00 7 . 0002355666777777889
 19,00 8 . 0011223334444455669
 5,00 9 . 01333
 Stem width: 10,00
 Each leaf: 1 case(s)



Εικόνα 4.12 Πίνακες και γραφήματα που εξάγει η διαδικασία «Explore».

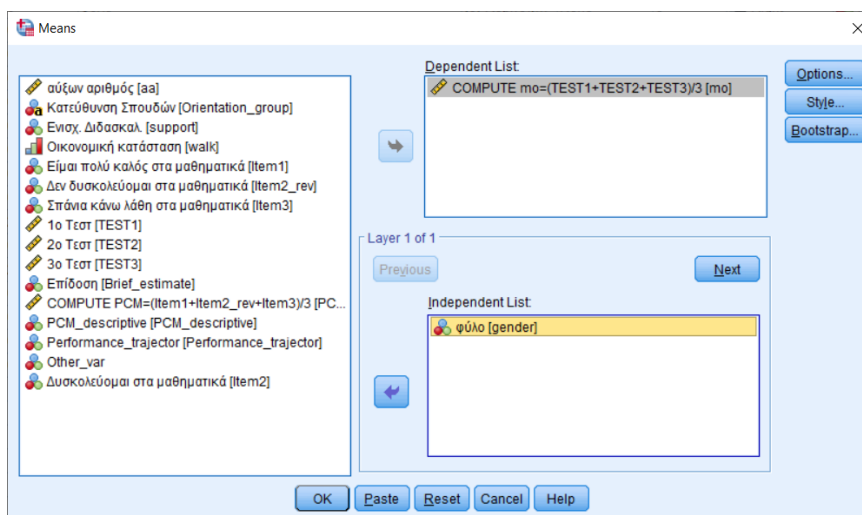
4.1.4 Σύνθετη παρουσίαση δύο μεταβλητών

Η σύνθετη παρουσίαση δύο μεταβλητών χρησιμοποιείται στην περίπτωση που μας ενδιαφέρει η παρουσίαση, με τη βοήθεια περιγραφικών στατιστικών ή σχετικών συχνοτήτων, της σχέσης μιας ποιοτικής μεταβλητής (ανεξάρτητη) και μιας ποσοτικής ή ποιοτικής μεταβλητής (εξαρτημένη) (Γναρδέλλης 2013; Bryman, 2016; IBM SPSS Statistics Base 25, n.d.). Για παράδειγμα, θέλουμε να συγκρίνουμε τα περιγραφικά στατιστικά της επίδοσης των αντρών και των γυναικών σε ένα μάθημα. Η σύγκριση αυτή μας παρέχει μια αρχική εκτίμηση των διαφορών των δύο ομάδων. Παρακάτω, παρουσιάζονται με το SPSS τα εξής: α) η σύνθετη παρουσίαση μιας ποιοτικής και μιας ποσοτικής μεταβλητής και β) η σύνθετη παρουσίαση δύο ποιοτικών μεταβλητών.

4.1.4.1 Σύνθετη παρουσίαση μιας ποιοτικής και μιας ποσοτικής μεταβλητής: «Compare means»

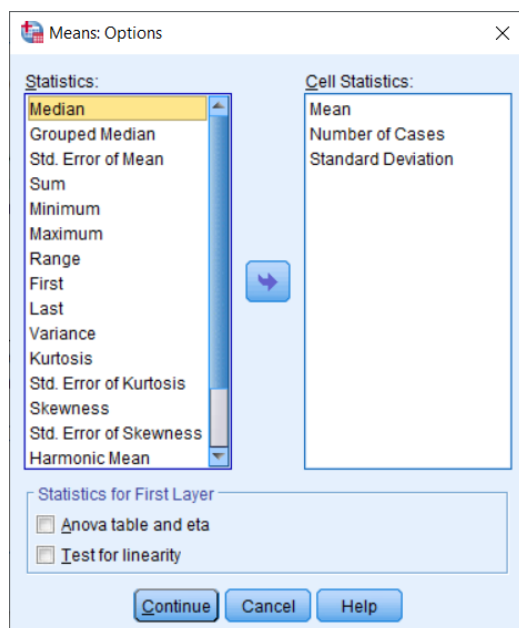
Χρησιμοποιείται στην περίπτωση που μας ενδιαφέρει η παρουσίαση, με τη βοήθεια περιγραφικών στατιστικών, της σχέσης μιας ποιοτικής μεταβλητής (ανεξάρτητη) και μιας ποσοτικής μεταβλητής (εξαρτημένη). Για να το πετύχουμε αυτό, διαλέγουμε διαδοχικά από το μενού [Analyze=>Descriptive Statistics=>Explore](#) και όπως δείξαμε και πιο πάνω (στην ενότητα «Περιγραφικά στατιστικά μέσω της επιλογής Explore») τοποθετούμε την ποσοτική μεταβλητή στο [Dependent List](#) και την ποιοτική μεταβλητή στο [Factor List](#).

Εναλλακτικά, μπορούμε να χρησιμοποιήσουμε τη διαδικασία [Compare means](#). Από το μενού επιλέγουμε [Analyze=>Compare means=>Means](#). Στο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 4.13), τοποθετούμε την ποσοτική μεταβλητή στην περιοχή [Dependent List](#) και την ποιοτική μεταβλητή στην περιοχή [Independent List](#).



Εικόνα 4.13 Πλαίσιο διαλόγου «Means».

Στη συνέχεια, στην επιλογή [Options](#) μπορούμε στο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 4.14) να επιλέξουμε συγκεκριμένα περιγραφικά στατιστικά από την περιοχή των στατιστικών ([Statistics](#)) και να τα τοποθετήσουμε στην περιοχή [Cell Statistics](#). Αν δεν επιλεγούν κάποια στατιστικά, ο εξαγόμενος πίνακας θα περιλαμβάνει για κάθε ομάδα, τη μέση τιμή ([Mean](#)), των αριθμό των τιμών της μεταβλητής σε κάθε ομάδα ([Number of Cases](#)) και την τυπική απόκλιση ([Standard Deviation](#)).



Εικόνα 4.14 Πλαίσιο διαλόγου «Means:Options».

Το αποτέλεσμα της παραπάνω διαδικασίας είναι δύο πίνακες (Εικόνα 4.15). Στον πρώτο πίνακα, παρουσιάζεται συνοπτικά ο αριθμός των έγκυρων και των ελλειπουσών τιμών. Στον δεύτερο πίνακα, παρουσιάζονται σε τρεις στήλες μετά τις υποδείξεις των ομάδων, η μέση τιμή, το μέγεθος του δείγματος και η τυπική απόκλιση, για κάθε ομάδα τιμών αλλά και συνολικά. Στο συγκεκριμένο παράδειγμα, φαίνεται μια υπεροχή της μέσης τιμής της επίδοσης των κοριτσιών έναντι των αγοριών.

Case Processing Summary						
	Cases					
	Included		Excluded		Total	
	N	Percent	N	Percent	N	Percent
COMPUTE mo = (TEST1+TEST2+TEST3)/3 * φύλο	61	100,0%	0	0,0%	61	100,0%

Report			
COMPUTE mo = (TEST1+TEST2+TEST3) /3			
Φύλο	Mean	N	Std. Deviation
0 Κορίτσι	75,8990	33	15,26854
1 Αγόρι	72,1071	28	10,26122
Total	74,1585	61	13,24185

Εικόνα 4.15 Εξαγόμενοι πίνακες της διαδικασίας «Means».

4.1.4.2 Σύνθετη παρουσίαση μιας ποιοτικής και μιας ποσοτικής μεταβλητής: «Custom Tables»

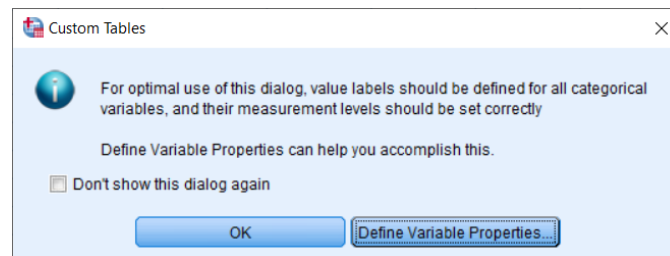
Αξιοποιώντας το αρχείο δεδομένων του SPSS «marks_trans_2.sav», θέλουμε να παρουσιάσουμε σε έναν πίνακα τη γνωστική πορεία των μαθητών στα τρία τεστ ανάλογα με την οικονομική τους κατάσταση. Στον πίνακα αυτόν θα πρέπει να εμφανίζονται περιγραφικά στατιστικά (Ελάχιστη και Μέγιστη τιμή, Μέση τιμή, και Τυπική Απόκλιση) ξεχωριστά για κάθε ομάδα μαθητών, όπως αυτή διαμορφώνεται από την οικονομική τους κατάσταση.

Τα τρία τεστ αποτελούν ποσοτικές μεταβλητές (scale) και η οικονομική κατάσταση (κακή, μέτρια, και καλή) αποτελεί ποιοτική μεταβλητή διάταξης (ordinal). Για την παρουσίαση σε πίνακα, των επιδόσεων των τριών ομάδων (κακή, μέτρια, και καλή οικονομική κατάσταση) στα τρία τεστ, μπορούμε να δημιουργήσουμε

έναν συνθετικό πίνακα (διασταύρωσης), όπου στις στήλες θα παρουσιάζονται τα περιγραφικά στατιστικά των τριών ποσοτικών μεταβλητών για κάθε τιμή (ομάδα) της ποιοτικής μεταβλητής (οικονομική κατάσταση).

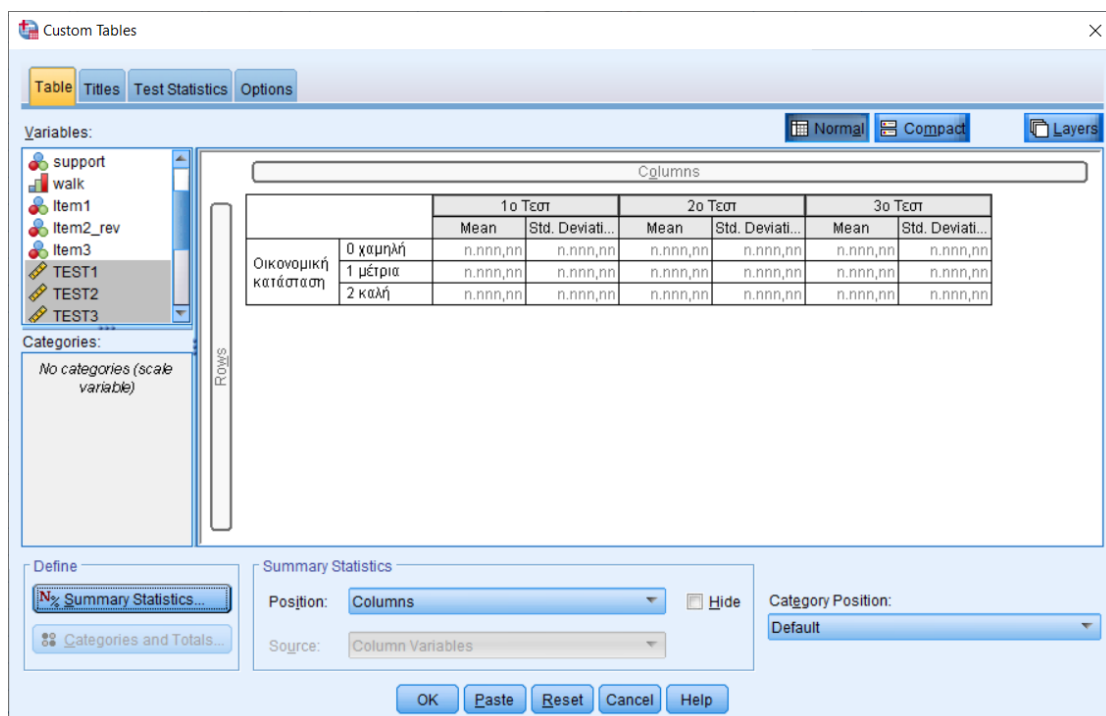
Ο πίνακας αυτός μπορεί να δημιουργηθεί στο SPSS από το μενού **Analyze=>Tables=>Custom Tables**. Το πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 4.16) μας ενημερώνει ότι οι προσαρμοσμένοι πίνακες λειτουργούν καλύτερα εάν έχουν καταχωριστεί οι ετικέτες για όλες τις τιμές των κατηγορικών μεταβλητών. Εάν το έχουμε ήδη κάνει, μπορούμε απλώς να κάνουμε κλικ στο **OK**.

Θα λάβουμε ένα μήνυμα σχετικά με το επίπεδο μέτρησης εάν κάποια από τις μεταβλητές εξακολουθεί να είναι ρυθμισμένη σε άγνωστο επίπεδο μέτρησης. Το SPSS μας δίνει μερικές επιλογές, αλλά μπορεί να θελήσουμε να ακυρώσουμε αυτήν την περίπτωση και να μεταβούμε στην καρτέλα **Variable View**, ώστε να ελέγξουμε ποιες μεταβλητές έχουν ορίσει το **Measure** ως άγνωστο και να τις αλλάξουμε στο ορθό επίπεδο μέτρησης.



Εικόνα 4.16 Πλαίσιο διαλόγου ενημέρωσης για την καλύτερη χρήση του «Custom Tables».

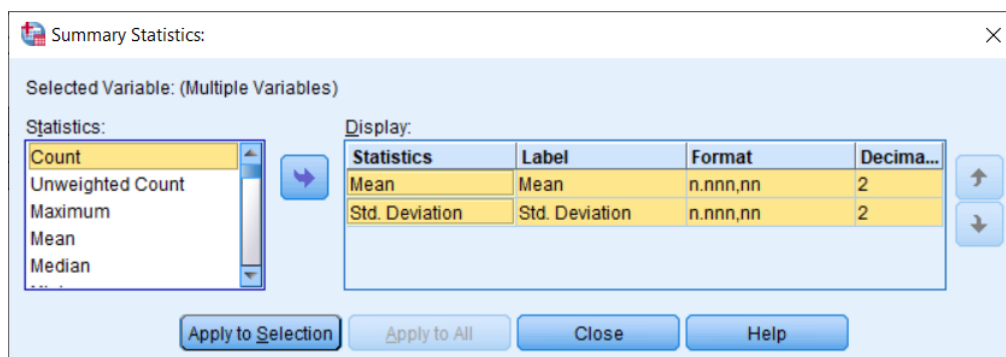
Αμέσως μετά στο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 4.17), τις μεταβλητές που θέλουμε να συμπεριλάβουμε στον πίνακα, τις σύρουμε και τις αφήνουμε (drag and drop) δίπλα στην επιλογή **Columns** ή στην επιλογή **Rows**. Στη συγκεκριμένη περίπτωση τοποθετήσαμε τις ποσοτικές μεταβλητές στις **Rows** και την ποιοτική μεταβλητή στις **Columns**.



Εικόνα 4.17 Πλαίσιο διαλόγου «Custom Tables».

Από την επιλογή **N% Summary Statistics** στο πλαίσιο διαλόγου **Summary Statistics** που εμφανίζεται (Εικόνα 4.17) μπορούμε να επιλέξουμε τα στατιστικά (**Count**, **Mean**, **Minimum** κ.λπ.) που μας ενδιαφέρουν για την ομάδα των μεταβλητών, καθώς και τη μορφή της παρουσίασης (επιθυμητή ακρίβεια κ.λπ.).

Αμέσως μετά επιλέγουμε **Apply to All**. Όπως φαίνεται στην Εικόνα 4.18, μπορούμε στο **Summary Statistics** να επιλέξουμε το πλαίσιο **Hide** για να μην παρουσιάζονται στον πίνακα οι λέξεις Mean κ.λπ. Στη θέση **Position** η επιλογή **Column**, υποδεικνύει ότι τα στατιστικά που επιλέξαμε να εμφανιστούν θα φαίνονται σε στήλες. Τέλος, στο **Category Position** επιλέγουμε **Row Labels in Columns**, δηλαδή οι κατηγορίες των μεταβλητών να τοποθετηθούν σε στήλες.



Εικόνα 4.18 Πλαίσιο διαλόγου «Summary Statistics».

Στον εξαγόμενο πίνακα που ακολουθεί (Εικόνα 4.19), παρουσιάζονται οι μέσες τιμές (Mean) και οι τυπικές αποκλίσεις (Standard Deviation) των ποσοτικών μεταβλητών για κάθε τιμή (ομάδα) της ποιοτικής μεταβλητής. Η πορεία της επίδοσης σε κάθε ομάδα (σύμφωνα με την οικονομική κατάσταση) υποδεικνύεται κάθετα συγκρίνοντας τις μέσες τιμές σε κάθε ομάδα. Επομένως, λαμβάνοντας υπόψη τις μέσες τιμές στα τρία τεστ σε κάθε ομάδα (κατακόρυφη ανάγνωση), παρατηρούμε σχεδόν παρόμοια γνωστική εξέλιξη στα τρία τεστ.

		1ο Τεστ		2ο Τεστ		3ο Τεστ	
		Mean	Standard Deviation	Mean	Standard Deviation	Mean	Standard Deviation
Οικονομική κατάσταση	0 χαμηλή	70,80	16,38	74,36	13,72	78,76	16,82
	1 μέτρια	69,94	17,49	74,06	21,80	81,78	13,23
	2 καλή	66,39	15,44	72,78	17,36	78,00	14,23

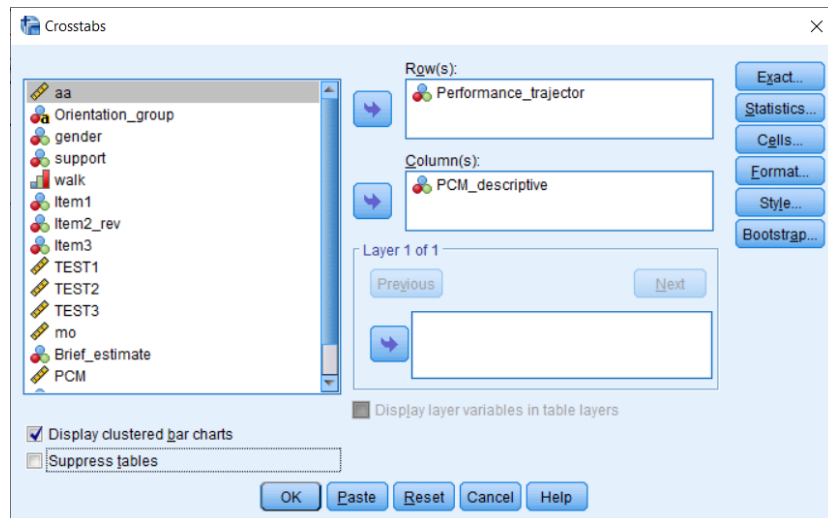
Εικόνα 4.19 Εξαγόμενος πίνακας της διαδικασίας «Custom Tables».

4.1.4.3 Σύνθετη παρουσίαση δύο ποιοτικών μεταβλητών: «Crosstabs»

Αξιοποιώντας το αρχείο δεδομένων του SPSS «marks_trans_2.sav» και τις μεταβλητές «PCM_descriptive» και «Performance_trajectory», δηλαδή την εκτίμηση της αντιλαμβανόμενης ικανότητας των μαθητών στα μαθηματικά (δύο τιμές) και τη σύνοψη της γνωστικής τους πορείας (τρεις τιμές), θέλουμε να παρουσιάσουμε σε πίνακα διασταύρωσης τη διαφοροποίηση της γνωστικής τους πορείας ανάλογα με τα δύο διαφορετικά επίπεδα της αντιλαμβανόμενης ικανότητάς τους στα μαθηματικά.

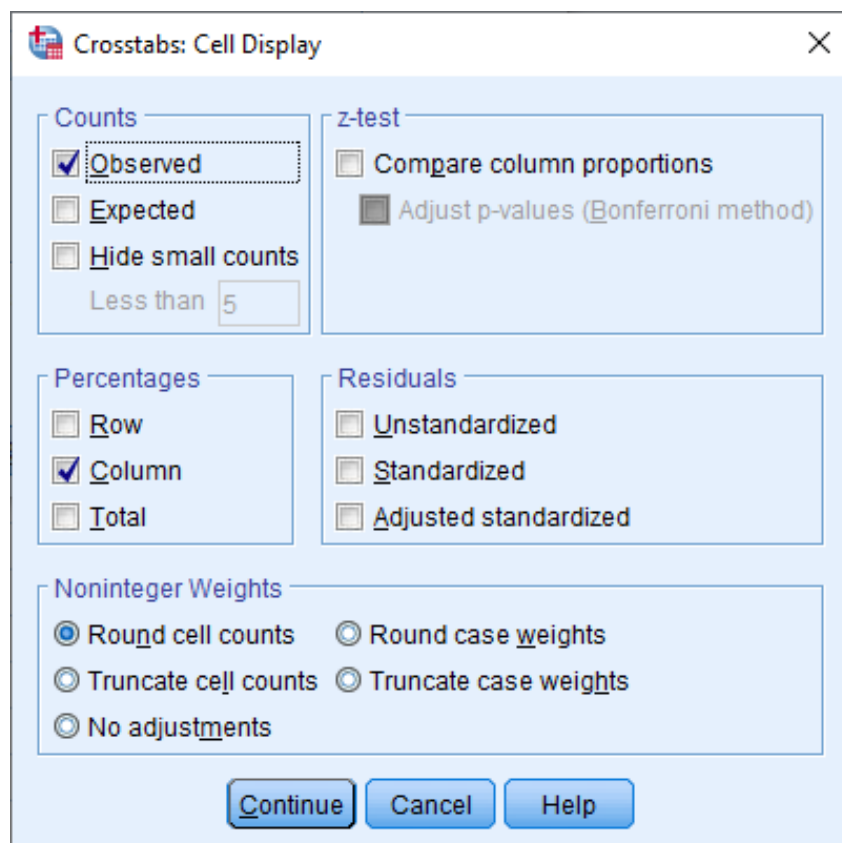
Οι δύο αυτές μεταβλητές είναι ποιοτικές διάταξης. Για να παρουσιάσουμε συνθετικά τις δύο αυτές μεταβλητές θα πρέπει να φτιάξουμε έναν πίνακα διασταύρωσης. Στις στήλες τοποθετούμε συνήθως τις τιμές της ανεξάρτητης μεταβλητής και στις γραμμές τις τιμές της εξαρτημένης μεταβλητής. Είναι σημαντικό ότι για να μπορέσουμε να καταλάβουμε την όποια σχέση, θα πρέπει να ορίσουμε, έστω και καταχρηστικά, τη μία μεταβλητή ως ανεξάρτητη και την άλλη ως εξαρτημένη. Στην περίπτωση αυτή η «Performance_trajectory» και η «PCM_descriptive» είναι η εξαρτημένη και η ανεξάρτητη μεταβλητή αντίστοιχα. Υπενθυμίζουμε ότι η έρευνα πραγματοποιήθηκε σε δύο στάδια: στο πρώτο στάδιο (χρονικά) καταγράψαμε την αντιλαμβανόμενη ικανότητα στα μαθηματικά «PCM_descriptive» και στο δεύτερο καταγράψαμε την πορεία της επίδοσης «Performance_trajectory».

Ο πίνακας αυτός μπορεί να δημιουργηθεί στο SPSS, από το μενού **Analyze =>Descriptive Statistics =>Crosstabs**. Στο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 4.20), τοποθετούμε την ανεξάρτητη μεταβλητή στην περιοχή **Column(s)** και την εξαρτημένη στην περιοχή **Row(s)**. Αν θέλουμε να εμφανιστεί αντίστοιχο ραβδόγραμμα σύνθετης παρουσίασης δύο ποιοτικών μεταβλητών, επιλέγουμε την επιλογή **Display clustered bar charts**. Με την επιλογή **Suppress table**, έχουμε τη δυνατότητα να μην εμφανιστεί ο αντίστοιχος πίνακας.



Εικόνα 4.20 Πλαίσιο διαλόγου «Crosstabs».

Στην επιλογή **Cells** και στο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 4.21), επιλέγουμε στην περιοχή **Percentages** την επιλογή **Column** και στη συνέχεια επιλέγουμε **Continue** και τέλος **OK**. Με το Percentages-Column δηλώνουμε ότι θέλουμε να υπολογιστούν οι σχετικές συχνότητες (%) των τιμών της μεταβλητής που έχουμε τοποθετήσει στις γραμμές, για κάθε τιμή (ομάδα τιμών) που έχουμε τοποθετήσει στις στήλες (Column).



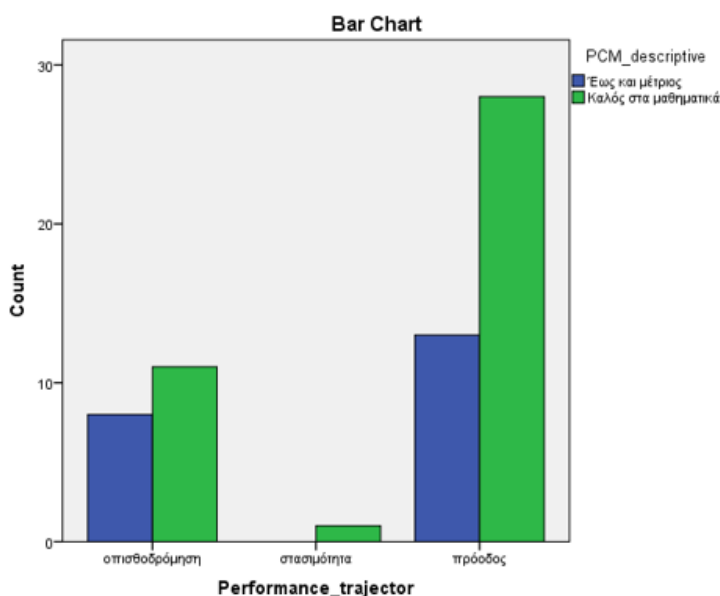
Εικόνα 4.21 Πλαίσιο διαλόγου «Crosstabs: Cell Display».

Στον εξαγόμενο πίνακα (Εικόνα 4.22), παρουσιάζονται οι σχετικές συχνότητες (%) των τιμών της εξαρτημένης μεταβλητής που έχουμε τοποθετήσει στις γραμμές, για κάθε τιμή (ομάδα τιμών) της ανεξάρτητης μεταβλητής που έχουμε τοποθετήσει στις στήλες. Συγκρίνοντας τα ποσοστά των τιμών (οριζόντια) κάθε τιμής της εξαρτημένης μεταβλητής για κάθε τιμή της ανεξάρτητης μεταβλητής, μπορούμε να περιγράψουμε τη σχέση

των δύο αυτών μεταβλητών. Όπως φαίνεται στον εξαγόμενο πίνακα, το ποσοστό των μαθητών που δήλωσαν μέτρια αντιλαμβανόμενη ικανότητα στα μαθηματικά και παρουσίασαν πρόοδο είναι μικρότερο (61,9%) από το αντίστοιχο ποσοστό της άλλης ομάδας (70,0%) που δήλωσαν καλή αντιλαμβανόμενη ικανότητα στα μαθηματικά.

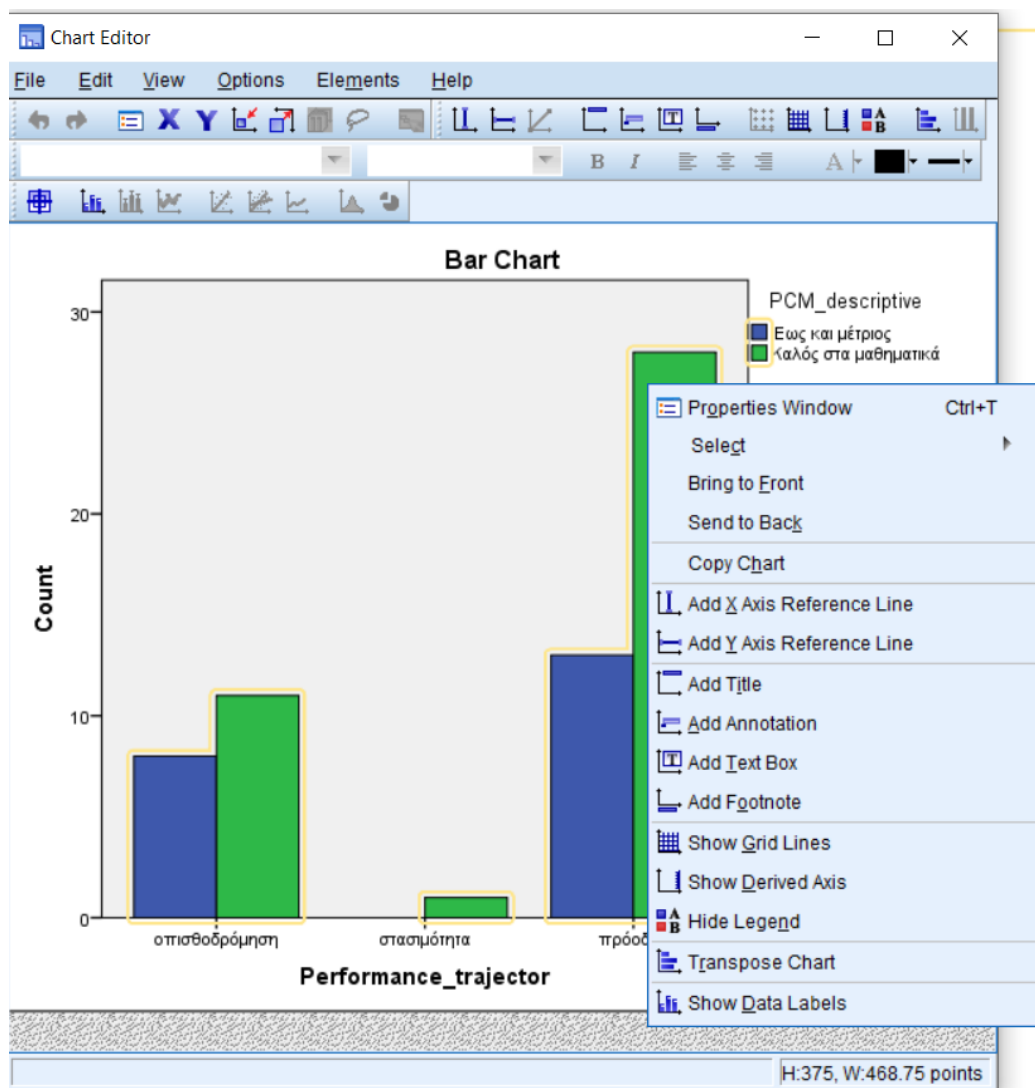
Στο εξαγόμενο γράφημα, αναπαρίστανται μόνο οι συχνότητες κάθε διασταύρωσης των τιμών του πίνακα.

Performance_trajector * PCM_descriptive Crosstabulation				
		PCM_descriptive		Total
		1 Έως και μέτριος	2 Καλός στα μαθηματικά	
Performance_trajector	0 οπισθοδρόμηση	Count	8	11
		% within PCM_descriptive	38,1%	27,5%
	1 στασιμότητα	Count	0	1
		% within PCM_descriptive	0,0%	2,5%
	2 πρόοδος	Count	13	28
		% within PCM_descriptive	61,9%	70,0%
Total		Count	21	40
		% within PCM_descriptive	100,0%	100,0%



Εικόνα 4.22 Εξαγόμενος πίνακας της διαδικασίας «Crosstabs».

Βέβαια, προκειμένου να εμφανιστούν στις στήλες οι συχνότητες αυτές θα πρέπει να επεξεργαστούμε το γράφημα. Αυτό γίνεται αν κάνουμε δύο αριστερά κλικ με το ποντίκι πάνω στο γράφημα, όπου αμέσως εμφανίζεται το παράθυρο επεξεργασίας του γραφήματος ([Chart Editor](#)) (Εικόνα 4.23) και εκεί με δεξί κλικ του ποντικιού πάνω σε οποιαδήποτε στήλη επιλέγουμε [Show Data Labels](#).



Εικόνα 4.23 Παράθυρο επεξεργασίας γραφήματος.

4.2 Προβλήματα για εξάσκηση και ανατροφοδότηση

Άσκηση 1

Δίνονται τα έτη φοίτησης των 25 φοιτητών που παρακολουθούν το μάθημα Στατιστική Ι: 3°, 3°, 3°, 3°, 3°, 4°, 3°, 4°, 3°, 3°, ΕΠ, 3°, ΕΠ, 3°, 3°, 3°, 3°, 3°, 4°, 3°, ΕΠ, 3°, 3°, 3°, 4°.

α) Να περάσετε τα δεδομένα στο περιβάλλον του SPSS, οργανωμένα σε μια μεταβλητή.

β) Να φτιάξετε τον πίνακα συχνοτήτων για τα έτη φοίτησης των φοιτητών.

γ) Να γράψετε ένα συνοδευτικό κείμενο στο οποίο θα φαίνονται: i) πόσοι από τους φοιτητές που παρακολουθούν το μάθημα φοιτούν μέχρι και το 4° έτος σπουδών και ii) ποιο ποσοστό των φοιτητών που παρακολουθούν το μάθημα φοιτούν μέχρι και το 4° έτος σπουδών.

Άσκηση 2

Για τα υποκείμενα της έρευνας που διεξάγουμε έχουμε καταγράψει την κατεύθυνση σπουδών τους στο Λύκειο: 1, 2, 2, 3, 1, 2, 1, 2, 1, 2, 3, 3, 3, 3, 2, 1, 1, 1, 1, 2, 2, 2, 1, 1, 1, 3, 2, 3, 2, 3, 2, 3, 2, 3, 3, 2, 2, 1, 1, 2. όπου 1=θεωρητική, 2=τεχνολογική, 3=θετική.

α) Να περάσετε τα δεδομένα στο περιβάλλον του SPSS, οργανωμένα σε μια μεταβλητή.

β) Να φτιάξετε τον πίνακα κατανομής συχνοτήτων για τις τιμές της μεταβλητής «Κατεύθυνση σπουδών των φοιτητών».

γ) Να γράψετε ένα συνοδευτικό κείμενο (μία πρόταση) που να συνοψίζει την κατανομή των τιμών της συγκεκριμένης μεταβλητής.

Άσκηση 3

Λαμβάνοντας υπόψη τους επόμενους πίνακες συχνοτήτων, να γράψετε ένα μικρό κείμενο που να περιγράφει τα δημογραφικά χαρακτηριστικά των συμμετεχόντων εκπαιδευτικών σύμφωνα με την κατανομή των αντίστοιχων μεταβλητών στους πίνακες.

q0001 Φύλο

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Άντρας	14	2,1	2,1	2,1
	Γυναίκα	655	97,9	97,9	100,0
	Total	669	100,0	100,0	

age_group

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	<=40	169	25,3	25,3	25,3
	>40 & <=50	424	63,4	63,4	88,6
	>50	76	11,4	11,4	100,0
	Total	669	100,0	100,0	

q0004 Επίπεδο ολοκληρωμένης εκπαίδευσης

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Παιδαγωγική Ακαδημία και εξομοίωση σε ΑΕΙ	136	20,3	20,3	20,3
	Πτυχίο ΑΕΙ	383	57,2	57,2	77,6
	Μεταπτυχιακό Δίπλωμα	136	20,3	20,3	97,9
	Διδακτορικό Δίπλωμα	14	2,1	2,1	100,0
	Total	669	100,0	100,0	

q0005 Έτη διδακτικής εμπειρίας

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	1-10	153	22,9	22,9	22,9
	11-20	393	58,7	58,8	81,7
	21-30	106	15,8	15,9	97,6
	30+	16	2,4	2,4	100,0
	Total	668	99,9	100,0	
Missing	System	1	,1		
Total		669	100,0		

Άσκηση 4

Στο αρχείο «Marks_trans_2.sav» η μεταβλητή «PCM» αναφέρεται στη αντιλαμβανόμενη ικανότητα των μαθητών για τα μαθηματικά.

α) Να υπολογίσετε τα περιγραφικά στατιστικά (μέση τιμή, διάμεσος και τυπική απόκλιση) της μεταβλητής «PCM»:

i) σε όλο το δείγμα,

ii) για τις δύο ομάδες τιμών, όπως αυτές διακρίνονται από τη μεταβλητή «support»,

iii) για τις τρεις ομάδες τιμών, όπως αυτές διακρίνονται από τη μεταβλητή «Walk».

β) Να γράψετε ένα συνοδευτικό κείμενο για όλα τα παραπάνω.

Βιβλιογραφία

Bryman, A. (2016). *Social research methods*. London: Oxford University Press.

Γιαλαμάς, Β. (2005). *Στατιστικές Τεχνικές και Εφαρμογές στις Επιστήμες της Αγωγής*. Αθήνα: Εκδ. Πατάκη.

Γναρδέλλης, Χ. (2013). *Ανάλυση δεδομένων με το IBM SPSS Statistics 21*. Αθήνα: Εκδ. Παπαζήση.

Field, A. (2018). *Discovering statistics using IBM SPSS statistics* (5th ed.). London: SAGE.

IBM SPSS Statistics Base 25 (n.d.). Retrieved from
https://faculty.ksu.edu.sa/sites/default/files/ibm_spss_statistics_base.pdf

Κεφάλαιο 5

Ανάλυση ερωτήσεων πολλαπλών απαντήσεων

Σύνοψη

Στο κεφάλαιο αυτό γίνεται αρχικά μια σύντομη παρουσίαση των ερωτήσεων πολλαπλών απαντήσεων. Στη συνέχεια, υποδεικνύονται τρόποι οργάνωσης των δεδομένων των ερωτήσεων πολλαπλών απαντήσεων στο περιβάλλον του SPSS. Τέλος, αξιοποιώντας τα Custom Tables παρουσιάζονται τρόποι εμφάνισης πινάκων συχνοτήτων των ερωτήσεων πολλαπλών απαντήσεων.

Προαπαιτούμενη γνώση

Βασικές γνώσεις του περιβάλλοντος του SPSS για την εισαγωγή και την οργάνωση των δεδομένων και πίνακες κατανομής συχνοτήτων μιας μεταβλητής: Κεφάλαια 1 έως και 4 του συγγράμματος.

5.1 Εισαγωγή στην ανάλυση ερωτήσεων πολλαπλών απαντήσεων

Πολλαπλές απαντήσεις προκύπτουν για ένα θέμα όταν έχουμε ένα σύνολο σχετικών επιλογών ή χαρακτηριστικών στα οποία κάποιο υποκείμενο μπορεί να διαθέτει ένα ή περισσότερα από αυτά τα χαρακτηριστικά (Field, 2018; IBM SPSS Statistics Base 25, n.d.). Σε αυτό το κεφάλαιο θα επικεντρωθούμε στις ερωτήσεις πολλαπλών απαντήσεων. Μια ερώτηση πολλαπλών απαντήσεων παρουσιάζει μια λίστα πιθανών επιλογών και ο ερωτώμενος επιλέγει εκείνες τις επιλογές που είναι αληθινές για αυτόν.

Για παράδειγμα, ας υποθέσουμε ότι μας ενδιαφέρει να ερευνήσουμε την προηγούμενη εμπειρία μάθησης των μεταπτυχιακών φοιτητών με τη Στατιστική. Τότε μια ενδεικτική ερώτηση που θα πρέπει να συμπεριληφθεί στην ερευνά μας είναι αυτή που φαίνεται στην Εικόνα 5.1.

Ποια από τα επόμενα έχετε διδαχθεί; (μπορείτε να επιλέξετε περισσότερα από ένα)

- Περιγραφική Στατιστική ☐
- Επαγωγική Στατιστική ☐
- Πολυμεταβλητή Ανάλυση ☐

Εικόνα 5.1 Ενδεικτική ερώτηση πολλαπλών απαντήσεων.

Στην περίπτωση αυτή, ο κάθε ερωτώμενος (έστω οι απαντήσεις τριών συμμετεχόντων) μπορεί να επιλέξει κανένα ή κάποια από τις πολλές απαντήσεις (Εικόνα 5.2).

Η απάντηση του 1^{ου} ερωτώμενου

Ποια από τα επόμενα έχετε διδαχθεί; (μπορείτε να επιλέξετε περισσότερα από ένα)

- Περιγραφική Στατιστική ☒
- Επαγωγική Στατιστική ☐
- Πολυμεταβλητή Ανάλυση ☐

Η απάντηση του 2^{ου} ερωτώμενου

Ποια από τα επόμενα έχετε διδαχθεί; (μπορείτε να επιλέξετε περισσότερα από ένα)

- Περιγραφική Στατιστική ☐
- Επαγωγική Στατιστική ☐
- Πολυμεταβλητή Ανάλυση ☐

Η απάντηση του 3^{ου} ερωτώμενου

Ποια από τα επόμενα έχετε διδαχθεί; (μπορείτε να επιλέξετε περισσότερα από ένα)

- Περιγραφική Στατιστική ☒
- Επαγωγική Στατιστική ☒
- Πολυμεταβλητή Ανάλυση ☐

Εικόνα 5.2 Ενδεικτικές απαντήσεις τριών συμμετεχόντων στην ερώτηση πολλαπλών απαντήσεων.

Για να αναλύσουμε σωστά τις ερωτήσεις πολλαπλών απαντήσεων στο SPSS, το σύνολο δεδομένων μας θα πρέπει να έχει την ακόλουθη δομή (Field, 2018; IBM SPSS Statistics Base 25, n.d.):

- Κάθε σειρά (περίπτωση) θα πρέπει να αντιπροσωπεύει ένα υποκείμενο της έρευνας.
- Για μία δεδομένη ερώτηση πολλαπλών απαντήσεων κάθε δυνατή επιλογή απάντησης πρέπει να αναπαρίσταται σε ξεχωριστή στήλη (μεταβλητή).
- Οι τιμές δεδομένων πρέπει να είναι αριθμητικοί κωδικοί (συνήθως 1), εάν υπάρχει απάντηση, και κενό ή μηδέν, εάν δεν υπάρχει απάντηση στη συγκεκριμένη επιλογή.

Λαμβάνοντας υπόψη ότι κάθε δυνατή απάντηση είτε θα την επιλέξει κάποιος είτε όχι, η οργάνωση των δεδομένων για τις απαντήσεις των ερωτώμενων βασίζεται στη δημιουργία τόσων δίτιμων μεταβλητών όσες είναι και οι δυνατές επιλογές που έχει ο ερωτώμενος. Στην περίπτωση της παραπάνω ερώτησης, θα δημιουργήσουμε τρεις δίτιμες μεταβλητές που θα παίρνουν τις τιμές 0 και 1. Το 1 στην περίπτωση που η επιλογή αυτή επιλέγεται από τους ερωτώμενους, και 0 ή τίποτα στην περίπτωση που δεν επιλέγεται από τους ερωτώμενους.

Στην Εικόνα 5.3 μπορούμε να δούμε τη μορφή των δεδομένων οργανωμένη σε τρεις μεταβλητές, όσες είναι και οι δυνατές επιλογές που είχαν οι ερωτώμενοι. Έτσι, φαίνεται ότι ο πρώτος ερωτώμενος έχει επιλέξει ότι έχει διδαχθεί μόνο περιγραφική στατιστική (μεταβλητή «descriptive_stat»), ο δεύτερος τίποτα από τα προτεινόμενα και ο τρίτος έχει επιλέξει ότι έχει διδαχθεί τα δύο πρώτα, δηλαδή επαγωγική στατιστική (μεταβλητή «inferential_stat») και πολυμεταβλητή ανάλυση (μεταβλητή «multivariate_anal»).

aa	descriptive_stat	inferential_stat	multivariate_anal
1	1	0	0
2	0	0	0
3	1	1	0
4	1	1	1
5	1	0	0
6	0	1	1
7	0	1	1
8	0	1	0
9	1	1	0
10	1	0	0
11	1	1	1

Εικόνα 5.3 Μορφή και οργάνωση δεδομένων που συλλέγονται με τις ερωτήσεις πολλαπλών επιλογών.

Το παραπάνω μοτίβο των ερωτήσεων πολλαπλών απαντήσεων μπορεί να χρησιμοποιηθεί για να μετασχηματιστούν και οι ερωτήσεις ανοικτού τύπου (Bryman, 2016). Ας υποθέσουμε για παράδειγμα ότι μας ενδιαφέρει να καταγράψουμε τους λόγους για τους οποίους οι εκπαιδευτικοί συμμετέχουν σε ένα πρόγραμμα μεταπτυχιακών σπουδών. Μία ενδεικτική ερώτηση ανοικτού τύπου θα μπορούσε να είναι: «Παρουσιάστε τους λόγους για τη συμμετοχή σας στο πρόγραμμα μεταπτυχιακών σπουδών». Στην περίπτωση αυτή ο ερωτώμενος θα γράψει από κανέναν έως και κάποιον αριθμό λόγων. Η οργάνωση αυτών των δεδομένων βασίζεται στη λογική των ερωτήσεων πολλαπλών απαντήσεων. Εντοπίζουμε αρχικά σε κάθε ερωτώμενο τους διάφορους λόγους και στη συνέχεια φτιάχνουμε τόσες δίτιμες μεταβλητές (1 επιλέγουν τον λόγο και 0 δεν τον επιλέγουν) όσοι και οι διαφορετικοί λόγοι που δήλωσαν οι μεταπτυχιακοί φοιτητές.

Αν, για παράδειγμα, η προσεκτική ανάγνωση και η μελέτη όλων των απαντήσεων των ερωτώμενων στην ανοικτού τύπου ερώτηση οδήγησαν σε 5 διαφορετικούς λόγους, τότε θα φτιάξουμε 5 δίτιμες μεταβλητές (Εικόνα 5.4)

Έτσι, για παράδειγμα, ο πρώτος ερωτώμενος δήλωσε μέσα από το κείμενο που έγραψε τους δύο πρώτους λόγους. Ο δεύτερος αντίστοιχα τον πρώτο, τον τρίτο και τον πέμπτο λόγο κ.λπ. Τέλος, ο ενδέκατος φοιτητής δεν απάντησε στην ερώτηση αυτή δηλώνοντας κάποιον συγκεκριμένο λόγο.

aa	reason1	reason2	reason3	reason4	reason5
1	1	1	0	0	0
2	1	0	1	0	1
3	0	0	0	1	1
4	0	0	0	0	1
5	0	1	1	1	0
6	0	1	0	0	0
7	0	1	0	1	1
8	1	1	1	1	1
9	0	0	1	1	0
10	0	1	0	0	1
11	0	0	0	0	0

Εικόνα 5.4 Μορφή και οργάνωση των απαντήσεων των ερωτώμενων στην ερώτηση ανοικτού τύπου που οδηγεί σε πέντε διαφορετικές περιπτώσεις.

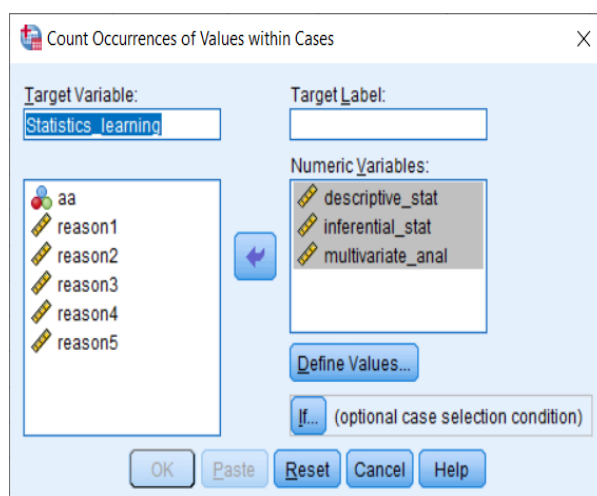
5.1.1 Καταμέτρηση επιλεγμένων επιλογών και εντοπισμός μη ερωτηθέντων

Ένα σημαντικό ζήτημα όταν εργαζόμαστε με ερωτήσεις πολλαπλών απαντήσεων είναι να μπορούμε να δηλώσουμε πόσα υποκείμενα δεν απάντησαν στην ερώτηση. Στην περίπτωση των ερωτήσεων πολλαπλών απαντήσεων, η απάντηση αυτή δεν είναι τόσο απλή όπως στις ερωτήσεις μιας επιλογής, όπου μπορούμε απλώς να μετρήσουμε τον αριθμό των τιμών που λείπουν σε μία στήλη. Η δυσκολία αυτή έγκειται στο γεγονός ότι, στις ερωτήσεις πολλαπλών απαντήσεων, οι απαντήσεις κατανέμονται σε πολλές δίτιμες μεταβλητές που μπορούν να επιλεγούν ανεξάρτητα (Field, 2018; IBM SPSS Statistics Base 25, n.d.). Επομένως, δεν απάντησε στην ερώτηση κάποιος που δεν έχει επιλέξει καμία από τις δυνατές απαντήσεις ή, αλλιώς, σε όλες τις δίτιμες μεταβλητές έχουμε πληκτρολογήσει το μηδέν. Επιπλέον, μπορεί να θέλουμε να μάθουμε πόσες επιλογές είχαν την τάση να επιλέγουν οι ερωτηθέντες. Οι περισσότεροι άνθρωποι επέλεξαν μόνο 1 επιλογή ή οι περισσότεροι είχαν την τάση να επιλέγουν 2 ή 3 επιλογές;

Μπορούμε να απαντήσουμε και στις δύο αυτές ερωτήσεις χρησιμοποιώντας τη διαδικασία **Count Values Within Cases** στο SPSS. Αυτή η διαδικασία, χρησιμοποιώντας ένα σύνολο μεταβλητών, μετράει πόσες φορές εμφανίζεται μια συγκεκριμένη τιμή για μία δεδομένη περίπτωση (υποκείμενο της έρευνας). Αυτή η μέτρηση συμπεριλαμβάνεται ως νέα μεταβλητή στο σύνολο δεδομένων, την οποία μπορούμε στη συνέχεια να χρησιμοποιήσουμε απαντώντας στα παραπάνω ερωτήματα.

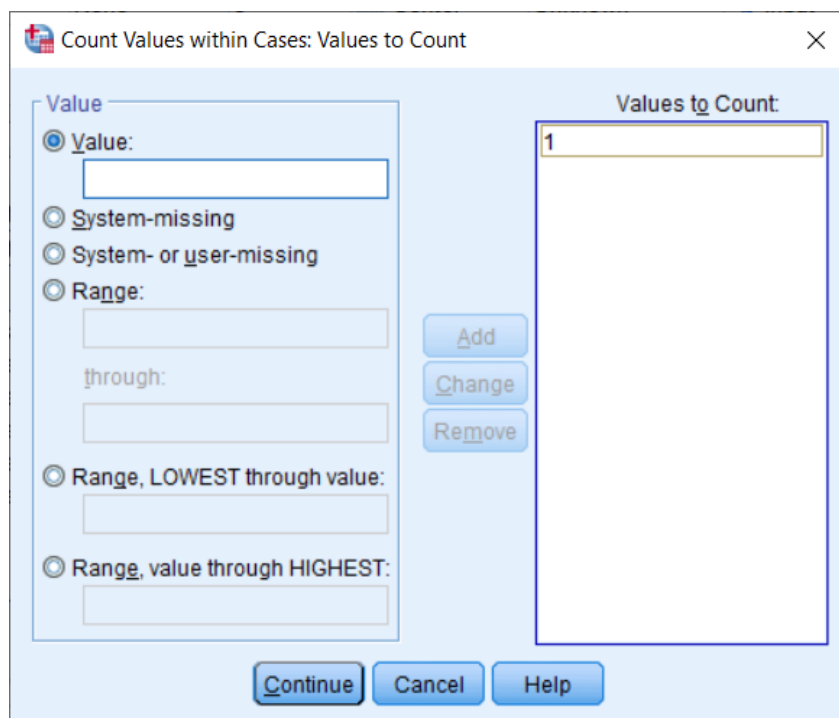
Στα δεδομένα του παραδείγματός μας (αρχείο «multiple responses.sav» και μεταβλητές «descriptive_stat», «inferential_stat» και «multivariate_anal») χρησιμοποιήσαμε τον αριθμό 1 για να υποδείξουμε την επιλογή, επομένως θέλουμε να μετρήσουμε τον αριθμό των 1 (όπου 1=επιλογή) που έχει ένα υποκείμενο στις τρεις μεταβλητές απόκρισης. Αν κάποιος δεν έχει κανένα 1 σε καμία από τις τρεις επιλογές, θα έχει μέτρηση 0.

Από το μενού επιλέγουμε διαδοχικά **Transform=> Count Values Within Cases**. Στο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 5.5) στις περιοχές **Target Variable** και **Target Label** πληκτρολογούμε το όνομα και τη σύντομη περιγραφή της μεταβλητής, αντίστοιχα. Επίσης, στην περιοχή **Numeric Variables** τοποθετούμε τις δίτιμες μεταβλητές «descriptive_stat», «inferential_stat» και «multivariate_anal».



Εικόνα 5.5 Πλαίσιο διαλόγου «Count Occurrence of Values within Cases».

Στη συνέχεια, επιλέγουμε το **Define Values** και στο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 5.6), και πιο συγκεκριμένα στην περιοχή **Value**, πληκτρολογούμε την τιμή 1 και πατάμε **Add**, τοποθετώντας πρακτικά την τιμή αυτή στην περιοχή **Values to Count**. Με αυτόν τον τρόπο δηλώνουμε ότι σε κάθε τριάδα απαντήσεων θα μετρηθεί πόσες φορές εμφανίζεται ο αριθμός 1 και επομένως πόσες απαντήσεις επέλεξε ο ερωτώμενος.



Εικόνα 5.6 Πλαίσιο διαλόγου «Count Values within Cases: Values to Count».

Το αποτέλεσμα της παραπάνω διαδικασίας είναι η δημιουργία μιας καινούριας μεταβλητής (Εικόνα 5.7), στην οποία εμφανίζεται ο αριθμός των επιλογών που επέλεξε ο κάθε ερωτώμενος.

aa	descriptive_stat	inferential_stat	multivariate_anal	Statistics_learning
1	1	0	0	1
2	0	0	0	0
3	1	1	0	2
4	1	1	1	3
5	1	0	0	1
6	0	1	1	2
7	0	1	1	2
8	0	1	0	1
9	1	1	0	2
10	1	0	0	1
11	1	1	1	3

Εικόνα 5.7 Η νέα μεταβλητή «Statistics_learning» που δημιούργησε η διαδικασία «Count Values Within Cases».

Η δημιουργία ενός πίνακα κατανομής συχνοτήτων για αυτή τη μεταβλητή, θα μας δώσει τις απαντήσεις που θέσαμε. Όπως φαίνεται στην Εικόνα 5.8, ένας μόνο δεν απάντησε στην ερώτηση αυτή των πολλαπλών απαντήσεων και οι περισσότεροι ερωτώμενοι ($36,4\%+18,2\%=54,6\%$) είχαν την τάση να επιλέγουν 2 τουλάχιστον επιλογές.

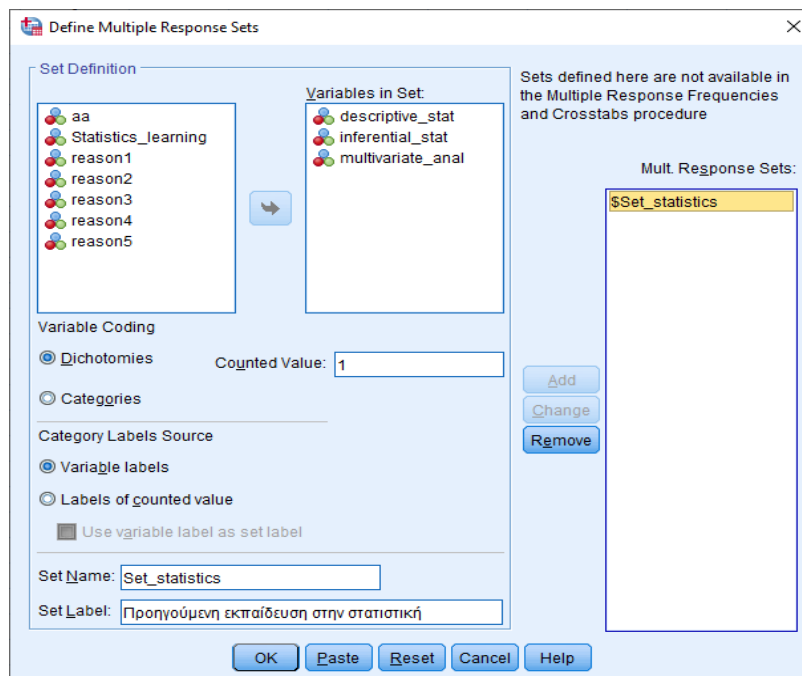
Statistics_learning					
		Frequenc	Percent	Valid Percent	Cumulative Percent
Valid	0	1	9,1	9,1	9,1
	1	4	36,4	36,4	45,5
	2	4	36,4	36,4	81,8
	3	2	18,2	18,2	100,0
	Total	11	100,0	100,0	

Εικόνα 5.8 Πίνακας κατανομής συχνοτήτων της μεταβλητής «Statistics_learning».

5.1.2 Δημιουργία σετ πολλαπλών απαντήσεων στο SPSS

Το σετ πολλαπλών απαντήσεων αποτελεί μια οντότητα σαν μια μεταβλητή που δημιουργείται στο SPSS και συγκεντρώνει το σύνολο των απαντήσεων που δόθηκαν από τους ερωτώμενους σε όλες τις δυνατές επιλογές της ερώτησης (IBM SPSS Statistics Base 25, n.d.).

Για τη δημιουργία του σετ πολλαπλών απαντήσεων, από το μενού επιλέγουμε διαδοχικά **Analyze** => **Custom Tables** => **Multiple Response Sets**. Στο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 5.9) μετακινούμε αρχικά τις δίτιμες μεταβλητές που αντιστοιχούν στις δυνατές επιλογές της ερώτησης πολλαπλών απαντήσεων (ακολουθώντας το προηγούμενο παράδειγμα, στην περίπτωσή μας από το αρχείο «multiple responses.sav», διαλέγουμε τις μεταβλητές «descriptive_stat», «inferential_stat» και «multivariate_anal») στο πεδίο **Variables in set**. Στη συνέχεια, στο πεδίο **Variable Coding** επιλέγουμε **Dichotomies** και υποδεικνύουμε ποιος αριθμός θα πρέπει να μετράει την επιλογή της απάντησης στο πεδίο **Counted Value**. Θυμίζουμε ότι ο αριθμός αυτός είναι το 1. Στη συνέχεια, στα πεδία **Set Name** και **Set Label** πληκτρολογούμε το όνομα και την πιο αναλυτική περιγραφή του σετ των πολλαπλών απαντήσεων και πατώντας **Add** βλέπουμε στην περιοχή **Mult. Response Sets** να εμφανίζεται το σετ που δημιουργήσαμε.



Εικόνα 5.9 Πλαίσιο διαλόγου «Define Multiple Response Sets».

Πατώντας **OK** θα δούμε στο αρχείο αποτελεσμάτων (output) τον επόμενο Πίνακα (Εικόνα 5.10) με πληροφορίες για το σετ και τις δίτιμες μεταβλητές (Elementary Variables) που το αποτελούν, καθώς και την τιμή (Counted Value) που αντιστοιχεί στην επιλογή των ερωτώμενων.

Multiple Response Sets				
Name	Coded As	Counted Value	Data Type	Elementary Variables
\$Set_statistics	Dichotomies	1	Numeric	descriptive_stat inferential_stat multivariate_anal

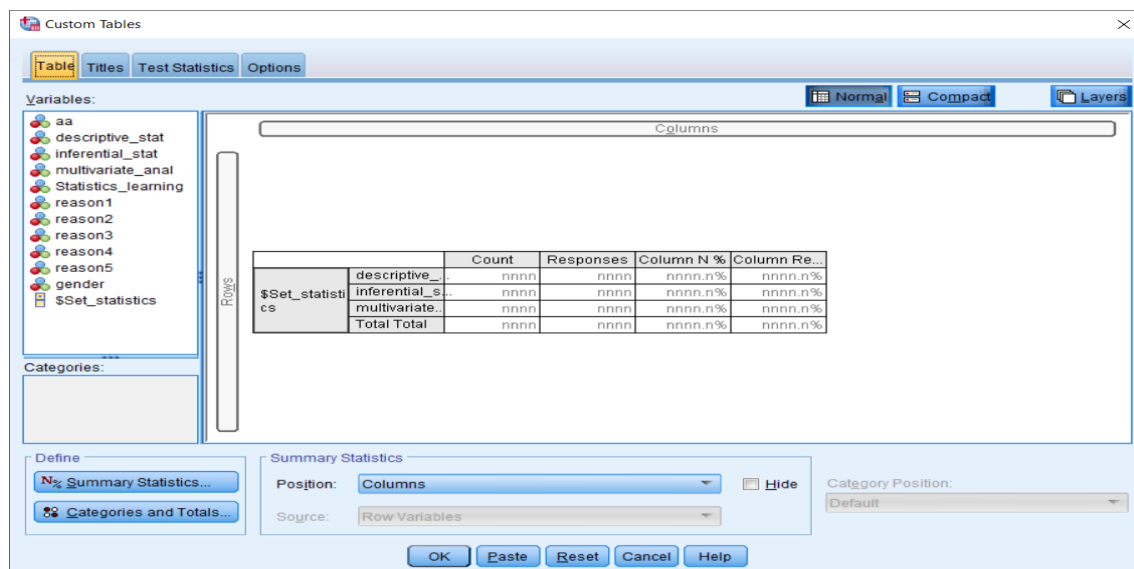
Εικόνα 5.10 Εξαγόμενος πίνακας με πληροφορίες για το σετ και τις δίτιμες μεταβλητές.

Το σετ πολλαπλών απαντήσεων που δημιουργήσαμε στο πλαίσιο διαλόγου στο πάνω μέρος και αριστερά στην υπόδειξη «[Sets defined here are not available in the Multiple Response Frequencies and Crosstabs procedure](#)» δεν είναι διαθέσιμο στις διαδικασίες «Frequencies» και «Crosstabs» του μενού [Analyze=>Descriptive statistics](#).

Τέλος, η δυνατότητα που έχουμε να δημιουργήσουμε σετ πολλαπλών απαντήσεων μέσω της επιλογής [Analyze => Multiple Response => Define Variable Sets](#) θα δημιουργήσει ένα προσωρινό σετ που θα μπορεί να χρησιμοποιηθεί μόνο όσο είναι ενεργό το συγκεκριμένο αρχείο δεδομένων στο περιβάλλον του SPSS (IBM SPSS Statistics Base 25, n.d.). Επιπλέον, μέσω αυτής της επιλογής έχουμε τη δυνατότητα να χρησιμοποιήσουμε μόνο τις διαδικασίες «Frequency» και «Crosstabs» που ανοίγουν μετά την επιλογή [Analyze => Multiple Response](#).

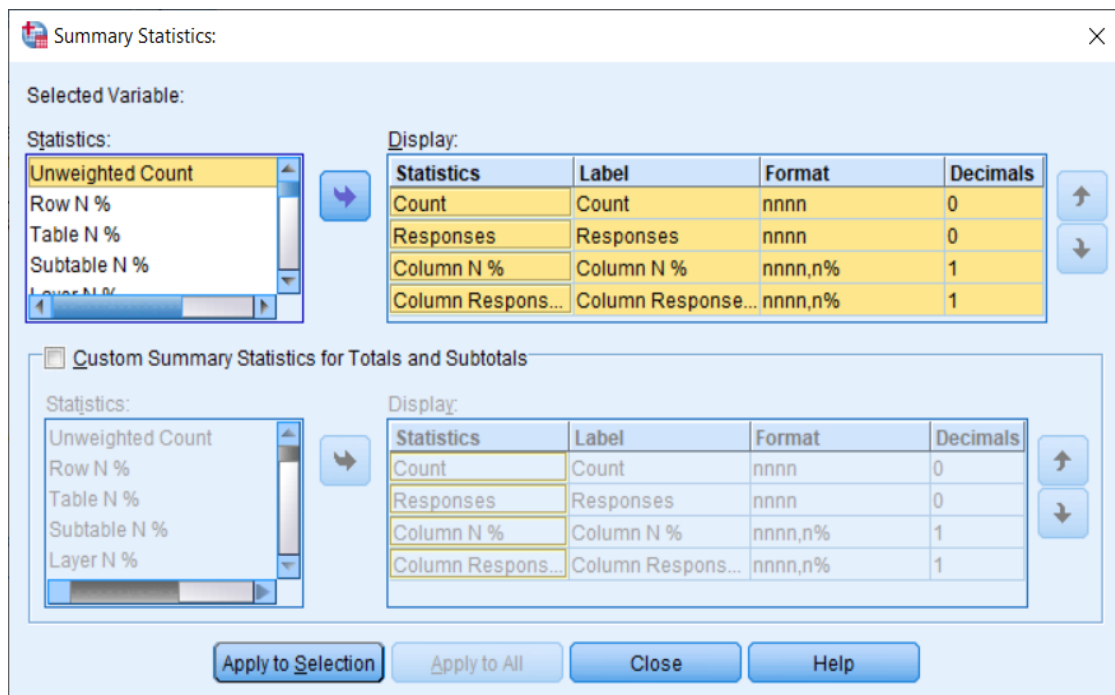
5.1.3 Δημιουργία πίνακα συχνοτήτων για το σετ πολλαπλών απαντήσεων

Για να μπορέσουμε να πραγματοποιήσουμε οποιαδήποτε ανάλυση με αυτό το σετ πολλαπλών απαντήσεων, θα πρέπει να χρησιμοποιήσουμε από το μενού διαδοχικά τις επιλογές [Analyze=>Custom Tables=>Custom Tables](#). Για το σετ πολλαπλών απαντήσεων που δημιουργήσαμε πριν (\$Set_statistics) θα δημιουργήσουμε έναν πίνακα κατανομής συχνοτήτων. Στο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 5.11), τις μεταβλητές που θέλουμε να συμπεριλάβουμε στον πίνακα, μπορούμε να τις σύρουμε και να τις αφήσουμε (drag and drop) δίπλα στην επιλογή [Columns](#) ή στην επιλογή [Rows](#). Στη συγκεκριμένη περίπτωση, τοποθετούμε το σετ πολλαπλών απαντήσεων (εικονίδιο)  στις [Rows](#).



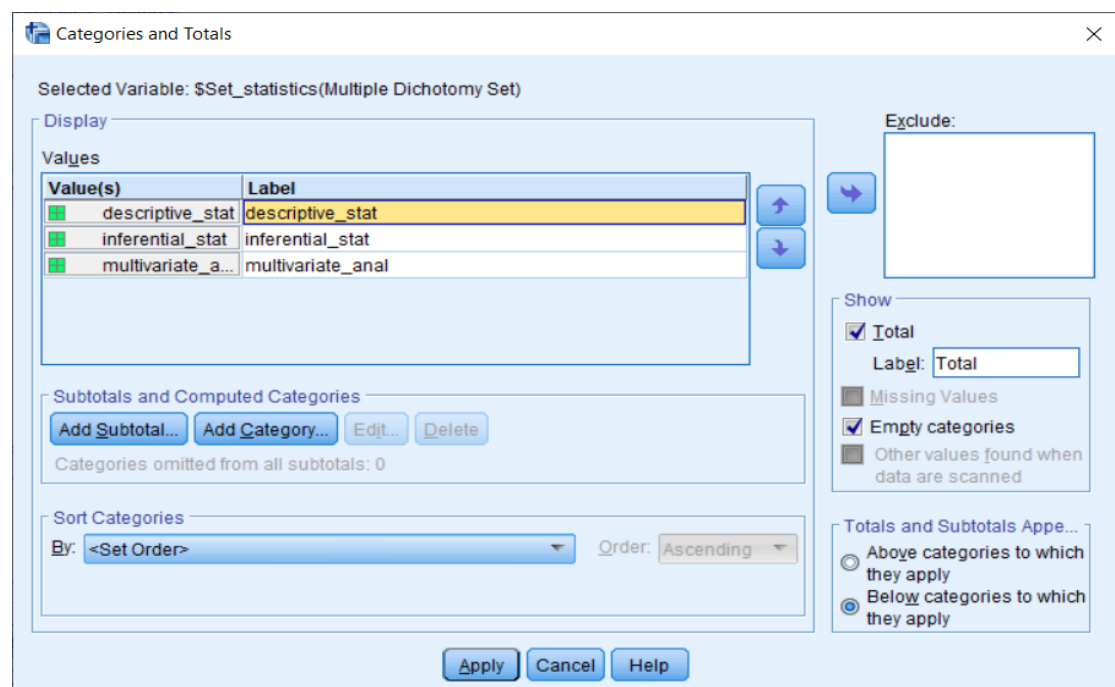
Εικόνα 5.11 Πλαίσιο διαλόγου της διαδικασίας «Custom Tables».

Από την επιλογή [N% Summary Statistics](#) στο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 5.12) μπορούμε να επιλέξουμε τα στατιστικά που μας ενδιαφέρουν για τη μεταβλητή, καθώς και τη μορφή της παρουσίασης (επιθυμητή ακρίβεια κ.λπ.). Στη συγκεκριμένη περίπτωση, θα επιλέξουμε: [Count](#), [Responses](#), [Column N%](#) και [Column Responses%](#). Αμέσως μετά επιλέγουμε [Apply to All](#) ή [Selection](#).



Εικόνα 5.12 Πλαίσιο διαλόγου της επιλογής «N% Summary Statistics».

Τέλος, στην επιλογή **Category and Totals** και στο αντίστοιχο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 5.13) επιλέγουμε την επιλογή **Total**, προκειμένου να εμφανιστεί κάτω από κάθε στήλη το αντίστοιχο άθροισμα.



Εικόνα 5.13 Πλαίσιο διαλόγου «Category and Totals».

Στην Εικόνα 5.14 μπορούμε να δούμε τον εξαγόμενο πίνακα της παραπάνω διαδικασίας. Στις στήλες Count ή Responses βλέπουμε τη συχνότητα των απαντήσεων σε καθεμία από τις επιλογές. Το πρώτο άθροισμα αντιστοιχεί στο σύνολο των συμμετεχόντων που απάντησαν (επέλεξαν) έστω και μία από τις επιλογές της ερώτησης πολλαπλών απαντήσεων. Το δεύτερο άθροισμα (στήλη Responses) αντιστοιχεί στο σύνολο των απαντήσεων (επιλογών) σε όλες τις επιλογές της ερώτησης πολλαπλών απαντήσεων. Η στήλη Column N% παρουσιάζει το ποσοστό απαντήσεων σε κάθε επιλογή, των ερωτώμενων που απάντησαν σε τουλάχιστον μία επιλογή από τις τρεις (συχνότητα επιλογής της απάντησης προς το σύνολο των ερωτώμενων που απάντησαν σε

μία τουλάχιστον επιλογή). Απαντά στο ερώτημα τι ποσοστό των ερωτώμενων που επέλεξαν τουλάχιστον έναν τύπο στατιστικής έχει διδαχθεί κάθε τύπο από τους προτεινόμενους. Αν και το άθροισμα αυτών των ποσοστών μπορεί να ξεπερνά και το 100%, το άθροισμα που φαίνεται αντιστοιχεί στο ποσοστό όλων των ερωτώμενων (Total του Count: 10) που απάντησαν σε τουλάχιστον μία επιλογή από τις τρεις, δηλαδή το 100%. Τέλος, η στήλη Column Responses% παρουσιάζει το ποσοστό των απαντήσεων της κάθε επιλογής (συχνότητα επιλογής της απάντησης προς το σύνολο των απαντήσεων). Απαντά στο ερώτημα τι ποσοστό των επιλογών (δυνατών απαντήσεων) αφορά τον κάθε τύπο της στατιστικής που έχουν διδαχθεί οι φοιτητές. Εδώ το άθροισμα αυτών των ποσοστών είναι 100%.

		Count	Responses	Column N %	Column Responses %
\$Set_statistics	descriptive_stat	7	7	70,0%	38,9%
	inferential_stat	7	7	70,0%	38,9%
	multivariate_anal	4	4	40,0%	22,2%
	Total	10	18	100,0%	100,0%

Εικόνα 5.14 Εξαγόμενος πίνακας της διαδικασίας «Custom tables», για το σετ πολλαπλών απαντήσεων.

5.2 Προβλήματα για εξάσκηση και ανατροφοδότηση

Άσκηση 1

Λαμβάνοντας υπόψη την ερώτηση πολλαπλών απαντήσεων αναφορικά με τους λόγους που καθόρισαν τη συμμετοχή των φοιτητών στο πρόγραμμα μεταπτυχιακών σπουδών (βλ. αρχείο «multiple responses.sav» και τις δίτιμες μεταβλητές «reason1», «reason2», «reason3», «reason4» και «reason5» που αντιστοιχούν στους πέντε διαφορετικούς λόγους) να:

α) δημιουργήσετε ένα σετ πολλαπλών απαντήσεων για την ερώτηση αυτή,

β) παρουσιάσετε σε έναν πίνακα συχνοτήτων:

i) τη συχνότητα των απαντήσεων σε καθεμία από τις επιλογές,

ii) το ποσοστό απαντήσεων σε καθεμία από τις πέντε επιλογές των ερωτώμενων που απάντησαν σε τουλάχιστον μία επιλογή από τις πέντε,

γ) παρουσιάσετε σε έναν πίνακα συχνοτήτων τα i) και ii) πώς διακρίνονται στους άντρες και πώς στις γυναίκες (μεταβλητή «gender»).

Τέλος, να γράψετε ένα συνοδευτικό κείμενο για κάθε εξαγόμενο πίνακα.

Βιβλιογραφία

Bryman, A. (2016). *Social research methods*. London: Oxford University Press.

Field, A. (2018). *Discovering statistics using IBM SPSS statistics* (5th ed.). London: SAGE.

IBM SPSS Statistics Base 25 (n.d.). Retrieved from:
https://faculty.ksu.edu.sa/sites/default/files/ibm_spss_statistics_base.pdf

Κεφάλαιο 6

Εισαγωγή στην Εκτιμητική και στον Έλεγχο Υποθέσεων

Σύνοψη

Στο κεφάλαιο αυτό γίνεται αρχικά μια σύντομη παρουσίαση των βασικών προσεγγίσεων της επαγωγικής στατιστικής. Παρουσιάζεται η εκτίμηση σημείου και διαστήματος, με έμφαση στον υπολογισμό διαστήματος εμπιστοσύνης με το SPSS. Ακολουθεί μια εισαγωγή στον έλεγχο υποθέσεων, με έμφαση στη διάκριση μεταξύ αμφίπλευρου και μονόπλευρου ελέγχου. Επιπλέον, παρουσιάζονται οι ερευνητικοί σχεδιασμοί ανεξάρτητων και εξαρτημένων δειγμάτων. Τέλος, παρουσιάζονται τρόποι ελέγχου των δεδομένων αλλά και ο έλεγχος της κανονικότητας των τιμών των ποσοτικών μεταβλητών πριν τη διεξαγωγή της ανάλυσης.

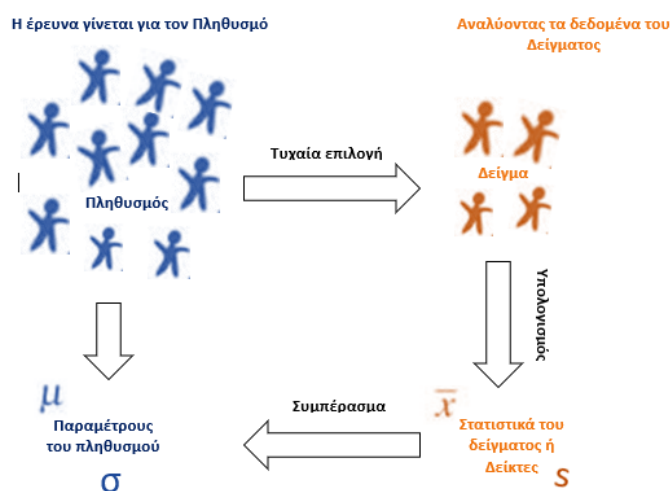
Προαπαιτούμενη γνώση

Βασικές γνώσεις του περιβάλλοντος του SPSS, καθώς και έννοιες που συζητήθηκαν αναφορικά με την περιγραφική στατιστική: Κεφάλαια 1 έως και 5 του συγγράμματος.

6.1 Εισαγωγή στην επαγωγική στατιστική

Για να δοθούν απαντήσεις σε ένα ερευνητικό ερώτημα πρέπει να γίνουν παρατηρήσεις και να καταγραφούν μετρήσεις που αφορούν χαρακτηριστικά των ατόμων που μελετώνται. Τα αποτελέσματα των παρατηρήσεων ή μετρήσεων αποτελούν τις *ακατέργαστες τιμές*. Μια συλλογή από τέτοιες τιμές που προέρχεται από μια ομάδα ατόμων αποτελεί τα *δεδομένα*. Οι στατιστικές μέθοδοι και τεχνικές που χρησιμοποιούνται για την οργάνωση και άντληση πληροφοριών από τα δεδομένα ταξινομούνται σε δύο μεγάλες κατηγορίες: στην περιγραφική στατιστική και στην επαγωγική στατιστική (Field, 2018; IBM SPSS Statistics Base 25, n.d.).

Η περιγραφική στατιστική περιλαμβάνει μεθόδους συλλογής, παρουσίασης και περιγραφής δεδομένων. Πιο συγκεκριμένα, εκφράζει αυτό που οι περισσότεροι άνθρωποι σκέφτονται όταν ακούν τη λέξη «Στατιστική». Υπάρχουν πολλές στατιστικές διαδικασίες, όπως η κατανομή συχνοτήτων, οι υπολογισμοί κεντρικής θέσης (όπως η μέση τιμή), γραφικές παραστάσεις κ.λπ., τις οποίες είδαμε στα προηγούμενα Κεφάλαια του συγγράμματος. Η επαγωγική στατιστική μπορεί να οριστεί απλώς ως το σύνολο των μεθόδων που χρησιμοποιούν δειγματικά δεδομένα, προκειμένου να ληφθούν αποφάσεις και να εξαχθούν συμπεράσματα για τον πληθυσμό (Γιαλαμάς, 2005; Γναρδέλλης, 2013; Bryman, 2016). Επιλέγεται δηλαδή ένα τυχαίο υποσύνολο του πληθυσμού, το δείγμα, στο οποίο συλλέγονται δεδομένα και υπολογίζονται στατιστικά (δείκτες) με στόχο να προσεγγίσουμε τις παραμέτρους του πληθυσμού και επομένως να αποκτήσουμε γνώση για το τι είναι πιθανό να συμβεί στον πληθυσμό (Εικόνα 6.1).



Εικόνα 6.1 Επαγωγική στατιστική, από το δείγμα στην εκτίμηση παραμέτρων του πληθυσμού.

Λαμβάνοντας υπόψη ότι από έναν πληθυσμό θα μπορούσαμε να επιλέξουμε άπειρα τυχαία δείγματα του ίδιου μεγέθους, για να μπορέσουμε να εκτιμήσουμε το τι συμβαίνει στον πληθυσμό είναι απαραίτητο να σκεφτούμε πώς τα στατιστικά του δείγματος σχετίζονται με τις παραμέτρους του πληθυσμού που θέλουμε να εκτιμήσουμε. Επομένως, θα πρέπει να εντάξουμε στην προσέγγισή μας την κατανομή των στατιστικών (για παράδειγμα δειγματική μέση τιμή) όλων των πιθανών δειγμάτων ίδιου μεγέθους, που μπορούν να επιλεγούν από τον πληθυσμό.

Η δειγματοληπτική κατανομή (sampling distribution) μιας δειγματικής στατιστικής συνάρτησης (π.χ. μέση τιμή, διακύμανση, μια αναλογία κ.λπ.) δημιουργείται αν θεωρητικά επιλέξουμε όλα τα δυνατά δείγματα ίδιου μεγέθους από έναν πληθυσμό και για καθένα από αυτά υπολογίσουμε τη στατιστική συνάρτηση. Για παράδειγμα, η δειγματοληπτική κατανομή της δειγματικής μέσης τιμής ($\bar{x}$) προσδιορίζεται, αν επιλέξουμε όλα τα δυνατά δείγματα ίδιου μεγέθους (n) από έναν πληθυσμό μεγέθους $N > n$ και υπολογίσουμε τη μέση τιμή όλων των δειγμάτων που επιλέξαμε.

Ειδικότερα, αν ο πληθυσμός ακολουθεί την κανονική κατανομή με μέση τιμή (μ) και διακύμανση (σ^2), τότε για τη δειγματοληπτική κατανομή της δειγματικής μέσης τιμής ($\bar{x}$) όλων των δυνατών δειγμάτων ισχύουν τα εξής:

- είναι κανονικά κατανομημένα,
- η μέση τιμή της δειγματοληπτικής κατανομής των μέσων τιμών ισούται με τη μέση τιμή (μ) του πληθυσμού,
- η διακύμανση της δειγματοληπτικής κατανομής ισούται με (σ^2/n).

Στην περίπτωση που ο πληθυσμός δεν ακολουθεί ή δεν γνωρίζουμε αν ακολουθεί την κανονική κατανομή, ισχύει το Κεντρικό Οριακό Θεώρημα (ΚΟΘ/ Central Limit Theorem - CLT) (Field, 2018).

Σύμφωνα με το Κεντρικό Οριακό Θεώρημα σε έναν οποιασδήποτε κατανομής πληθυσμό μεγέθους N , με μέση τιμή (μ) και τυπική απόκλιση (σ), η δειγματοληπτική κατανομή της δειγματικής μέσης τιμής όλων των δυνατών δειγμάτων μεγέθους (n) είναι:

- κατά προσέγγιση κανονική,
- με μέση τιμή (μ) και διακύμανση (σ^2/n), εφόσον το μέγεθος των δειγμάτων είναι επαρκώς μεγάλο: στις περισσότερες των περιπτώσεων $n \geq 30$.

Λαμβάνοντας υπόψη τη δειγματοληπτική κατανομή μιας στατιστικής συνάρτησης, είμαστε σε θέση να μελετούμε την απόσταση της στατιστικής τιμής από την αντίστοιχη παράμετρο και επομένως να προσεγγίζουμε τι μπορεί να συμβαίνει στον πληθυσμό. Παρακάτω θα αναφερθούμε στις δύο κατευθύνσεις αξιοποίησης της επαγωγικής στατιστικής, την εκτιμητική και τον έλεγχο των υποθέσεων.

6.1.1 Η Εκτιμητική

Η Εκτιμητική αφορά την εκτίμηση παραμέτρων του πληθυσμού μέσω των δεικτών ή στατιστικών του δείγματος (Εικόνα 6.1). Η εκτίμηση αυτή μπορεί να βασιστεί σε δύο προσεγγίσεις: τη *Σημειακή εκτίμηση* (point estimation) και την *Εκτίμηση διαστήματος* (interval estimation) (Γιαλαμάς, 2005; Bryman, 2016).

Με τη *σημειακή εκτίμηση* ο προσδιορισμός κάποιας παραμέτρου του πληθυσμού γίνεται μέσω ενός στατιστικού δείκτη. Για παράδειγμα, η δειγματοληπτική μέση τιμή (δείκτης) χρησιμοποιείται για να εκτιμήσουμε τη μέση τιμή του πληθυσμού (παραμέτρος). Προσοχή στο ότι η σημειακή εκτίμηση δεν προσδιορίζει ακριβώς την όποια παράμετρο, αλλά μόνο κατά προσέγγιση. Μάλιστα, καλό θα είναι να συνοδεύεται η όποια εκτίμηση της πληθυσμιακής μέσης τιμής μαζί με την τυπική απόκλιση αλλά και το μέγεθος του δείγματος. Έτσι, αν η μέση επίδοση των φοιτητών του δείγματος στη στατιστική είναι $M=6,2$ με $T.A.=1,2$ σε δείγμα μεγέθους $n=100$, τότε η τιμή αυτή, σύμφωνα με τη σημειακή εκτίμηση, αποτελεί μια προσέγγιση της μέσης τιμής του πληθυσμού από τον οποίο προήλθε το προηγούμενο αντιπροσωπευτικό δείγμα.

Με την *εκτίμηση διαστήματος* (interval estimation) ο προσδιορισμός κάποιας παραμέτρου του πληθυσμού γίνεται με τον υπολογισμό ενός διαστήματος δυνατών τιμών μέσα στο οποίο περιλαμβάνεται πιθανόν η τιμή της παραμέτρου του πληθυσμού. Το διάστημα αυτό καλείται *διάστημα εμπιστοσύνης* (ΔE) (confidence interval) σε επίπεδο εμπιστοσύνης ή ακρίβειας: $1-\alpha$, όπου α το επίπεδο σημαντικότητας ή η πιθανότητα να κάνουμε λάθος και το ΔE να μην περιέχει την εκτιμώμενη τιμή της παραμέτρου του πληθυσμού. Συνήθως, η πιθανότητα

αυτή είναι $\alpha=5\%$, άρα το επίπεδο εμπιστοσύνης είναι 95% . Επομένως, το διάστημα εμπιστοσύνης της μέσης τιμής του πληθυσμού μας λέει ότι, αν πάρουμε 100 δείγματα ίδιου μεγέθους από τον ίδιο πληθυσμό, τότε περιμένουμε ότι τα 95 από αυτά τα αντίστοιχα διαστήματα εμπιστοσύνης θα περιέχουν τη μέση τιμή του πληθυσμού και 5 από αυτά ότι δεν θα την περιέχουν.

Για παράδειγμα, ας υποθέσουμε ότι θέλουμε να εκτιμήσουμε τη μέση διδακτική εμπειρία (σε μήνες) των αδιόριστων εκπαιδευτικών, ειδικότητας φυσικών επιστημών. Έστω ένα δείγμα εκπαιδευτικών ($n=30$, επαρκώς μεγάλο), οι οποίοι έχουν μέση διδακτική εμπειρία $\bar{x}=11$ μήνες. Γνωρίζουμε ότι η τυπική απόκλιση της διδακτικής εμπειρίας των εκπαιδευτικών ΦΕ στον πληθυσμό είναι $\sigma=3$ μήνες. Μέσω της τυπικής κανονικής κατανομής (με μέση τιμή 0 και τυπική απόκλιση 1) το 95% διάστημα, στο οποίο θα βρίσκεται με βεβαιότητα (η πιθανότητα να βρίσκεται στο διάστημα αυτό είναι 95%) η μέση τιμή (μ) του πληθυσμού είναι:

$$\bar{x} - 1,96 * \frac{\sigma}{\sqrt{n}} \leq \mu \leq \bar{x} + 1,96 * \frac{\sigma}{\sqrt{n}}$$

δηλαδή

$$11 - 1,96 * \frac{3}{\sqrt{30}} \leq \mu \leq 11 + 1,96 * \frac{3}{\sqrt{30}}$$

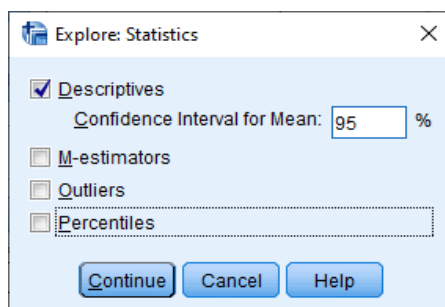
δηλαδή

$$9,93 \leq \mu \leq 12,07.$$

6.1.1.1 Υπολογισμός του διαστήματος εμπιστοσύνης με το SPSS

Για τον υπολογισμό του διαστήματος εμπιστοσύνης μιας ποσοτικής μεταβλητής στο περιβάλλον του SPSS θα χρησιμοποιήσουμε τη διαδικασία **Explore** (Αναλυτικές οδηγίες χρήσης της συγκεκριμένης διαδικασίας μπορείτε να δείτε στην ενότητα του κεφαλαίου της Περιγραφικής Στατιστικής, Κεφάλαιο 4). Στην κατεύθυνση αυτή, αξιοποιώντας το αρχείο δεδομένων του SPSS «marks_trans_2.sav», θέλουμε να υπολογίσουμε το διάστημα εμπιστοσύνης για τη μεταβλητή «mo».

Επιπρόσθετα της προηγούμενης παρουσίασης του πλαισίου διαλόγου της διαδικασίας **Explore**, στην επιλογή **Statistics** (Εικόνα 6.2) μπορούμε να επιλέξουμε, μεταξύ άλλων, το ποσοστιαίο διάστημα (στο πεδίο **Descriptives - Confidence Interval for Mean**), στο οποίο θα βασιστεί ο υπολογισμός του διαστήματος εμπιστοσύνης. Το συνηθισμένο ποσοστιαίο διάστημα είναι 95% .



Εικόνα 6.2 Πλαίσιο διαλόγου της επιλογής «Statistics» της διαδικασίας «Explore».

Στα εξαγόμενα αποτελέσματα που εμφανίζονται στον πίνακα (Εικόνα 6.3) μπορείτε να δείτε τα δύο άκρα Lower Bound (70.77) και Upper Bound (77,55) του 95% Confidence Interval for Mean.

Descriptives			Statistic	Std. Error
COMPUTE mo=(TEST1+TEST2+TE ST3)/3	Mean		74,1585	1,69545
	95% Confidence Interval	Lower Bound	70,7671	
		Upper Bound	77,5499	
	5% Trimmed Mean		74,9183	
	Median		77,6667	
	Variance		175,347	
	Std. Deviation		13,24185	
	Minimum		35,67	
	Maximum		93,00	
	Range		57,33	
	Interquartile Range		17,83	
	Skewness		-,960	,306
	Kurtosis		,401	,604

Εικόνα 6.3 Εξαγόμενος πίνακας της διαδικασίας «Explore».

6.1.2 Έλεγχος Υποθέσεων

Τα συμπεράσματα αναφορικά με τις παραμέτρους του πληθυσμού προκύπτουν και από τη διατύπωση αλλά και τον έλεγχο συγκεκριμένων υποθέσεων που αφορούν πληθυσμιακές παραμέτρους. Με τον έλεγχο υποθέσεων υποστηρίζεται η λήψη απόφασης σχετικά με έναν πληθυσμό, εξετάζοντας ένα δείγμα από τον πληθυσμό αυτό. *Ερευνητική υπόθεση* είναι μια σύντομη και ακριβής πρόταση στην οποία περιγράφουμε τι θεωρούμε ότι θα συμβεί στην έρευνα ή στο πείραμά μας. Με τις διαδικασίες ελέγχου υποθέσεων υποστηρίζουμε κατά πόσο οι υποθέσεις ικανοποιούνται με τα υπάρχοντα δεδομένα. Προσοχή στο ότι για τη διατύπωση των υποθέσεων βασιζόμαστε στην προηγούμενη έρευνα. Οι υποθέσεις αυτές καλούνται ερευνητικές και για να ελεγχθούν θα πρέπει να μετασχηματιστούν, χρησιμοποιώντας τη γλώσσα και την ορολογία της στατιστικής, σε στατιστικές υποθέσεις (Γιαλαμάς, 2005).

Για παράδειγμα, ερευνητικές υποθέσεις είναι οι ακόλουθες:

- Οι γυναίκες ηλικίας 20 ετών μπορούν να συγκρατήσουν περισσότερες λέξεις σε σχέση με τους άντρες της ίδιας ηλικίας.
- Οι γυναίκες ηλικίας 20 ετών μπορούν να συγκρατήσουν διαφορετικό αριθμό λέξεων σε σχέση με τους άντρες της ίδιας ηλικίας.
- Όσο μεγαλύτερη είναι η ακαδημαϊκή εμπλοκή του φοιτητή, τόσο μεγαλύτερη είναι και η ακαδημαϊκή επίδοσή του.
- Η ακαδημαϊκή εμπλοκή του φοιτητή σχετίζεται με την ακαδημαϊκή επίδοσή του.

Οι αντίστοιχες στατιστικές υποθέσεις θα μπορούσαν να είναι οι ακόλουθες:

- Η μέση τιμή του αριθμού των λέξεων που μπορούν οι γυναίκες ηλικίας 20 ετών να συγκρατήσουν είναι μεγαλύτερη από αυτή των αντρών.
- Η μέση τιμή του αριθμού των λέξεων που μπορούν οι γυναίκες ηλικίας 20 ετών να συγκρατήσουν είναι διαφορετική από αυτή των αντρών.
- Το σκορ της ακαδημαϊκής εμπλοκής του φοιτητή σχετίζεται θετικά με το σκορ της ακαδημαϊκής επίδοσής του.
- Το σκορ της ακαδημαϊκής εμπλοκής του φοιτητή σχετίζεται αρνητικά με το σκορ της ακαδημαϊκής επίδοσής του.

Όπως θα δούμε στη συνέχεια, αν και οι υποθέσεις αυτές ανά δύο φαίνεται να μοιάζουν, στην πραγματικότητα δεν είναι ίδιες, και στις δύο περιπτώσεις, είτε πρόκειται για ερευνητική είτε για στατιστική υπόθεση. Η (α) και η (γ) είναι υποθέσεις που έχουν κατεύθυνση ή, αλλιώς, μονόπλευρου ελέγχου, ενώ η (β) και η (δ) είναι υποθέσεις που δεν έχουν κατεύθυνση ή, αλλιώς, δίπλευρου ελέγχου.

Γενικότερα, για τον έλεγχο των ερευνητικών υποθέσεων χρειάζεται να διατυπωθούν δύο ειδών υποθέσεις:

Η πρώτη είναι η *Μηδενική υπόθεση* (null hypothesis) (συμβολίζεται με H_0). Η μηδενική είναι η υπόθεση που ελέγχεται (Γιαλαμάς, 2005; Field, 2018). Αναπαριστά την επίδραση της τύχης στη διαμόρφωση του αποτελέσματος. Είναι η υπόθεση που υποστηρίζει ότι δεν υπάρχει σχέση (ή διαφορά) μεταξύ των μεταβλητών που μελετώνται.

Για παράδειγμα, η μηδενική στατιστική υπόθεση θα μπορούσε να είναι:

- Δεν υπάρχει σχέση μεταξύ δύο μεταβλητών ($r_{12} = 0$, δηλαδή ο αντίστοιχος συντελεστής συσχέτισης που αναπαριστά τη συσχέτιση των μεταβλητών 1 και 2 είναι μηδέν).
- Δεν υπάρχει διαφορά, δηλαδή: α) η παράμετρος του πληθυσμού είναι ίση με κάποια τιμή ($\mu = \mu_0$, για παράδειγμα η μέση τιμή μ του πληθυσμού παίρνει συγκεκριμένη τιμή μ_0) και β) η παράμετρος του 1^{ου} πληθυσμού είναι ίση με την παράμετρο του 2^{ου} πληθυσμού ($\mu_1 = \mu_2$, για παράδειγμα η μέση τιμή του 1^{ου} πληθυσμού μ_1 είναι ίση με τη μέση τιμή του 2^{ου} πληθυσμού) κ.λπ.

Η δεύτερη υπόθεση είναι η *Εναλλακτική υπόθεση* (Alternative hypothesis) ή *Δηλωτική* (Declarative) (συμβολίζεται με H_A ή H_1). Δεν αναπαριστά την επίδραση της τύχης στη διαμόρφωση του αποτελέσματος. Η υπόθεση αυτή αναφέρεται στη θεώρηση (βασίζεται στην προηγούμενη έρευνα) που κάνει ο ερευνητής αναφορικά με τη σχέση που υπάρχει μεταξύ των μεταβλητών που μελετά. Με τον έλεγχο που θα πραγματοποιήσουμε, η μηδενική υπόθεση είτε απορρίπτεται είτε δεν απορρίπτεται. Εάν δεν την απορρίψουμε, λέμε ότι είτε δεν απορρίπτεται γενικότερα στον πληθυσμό είτε τα συγκεκριμένα δεδομένα δεν είναι ικανά να την απορρίψουν. Ωστόσο, αν ο έλεγχος οδηγήσει στην απόρριψή της, τότε συμπεραίνουμε ότι τα δεδομένα υποστηρίζουν την εναλλακτική υπόθεση (Γιαλαμάς, 2005; Field, 2018). Η εναλλακτική υπόθεση μπορεί να διατυπωθεί ως υπόθεση χωρίς κατεύθυνση ή ως υπόθεση με κατεύθυνση. Για παράδειγμα, η εναλλακτική στατιστική υπόθεση χωρίς ή με κατεύθυνση θα μπορούσε να είναι:

Υπόθεση χωρίς κατεύθυνση ή Αμφίπλευρου ελέγχου (Two tailed)

- Υπάρχει σχέση μεταξύ δύο μεταβλητών, π.χ. $r_{12} \neq 0$.
- Υπάρχει διαφορά, δηλαδή: α) η παράμετρος του πληθυσμού είναι διαφορετική από (κάποια τιμή) ($\mu \neq \mu_0$), β) η παράμετρος του 1^{ου} πληθυσμού είναι διαφορετική από την παράμετρο του 2^{ου} πληθυσμού ($\mu_1 \neq \mu_2$) κ.λπ.

Υπόθεση με κατεύθυνση ή Μονόπλευρου ελέγχου (One tailed)

- Υπάρχει θετική $r_{12} > 0$ ή αρνητική $r_{12} < 0$ συσχέτιση μεταξύ δύο μεταβλητών.
- Η παράμετρος του πληθυσμού είναι μεγαλύτερη από κάποια τιμή ($\mu > \mu_0$) ή η παράμετρος του πληθυσμού είναι μικρότερη από κάποια τιμή ($\mu < \mu_0$).
- Η παράμετρος του 1^{ου} πληθυσμού είναι μικρότερη ($\mu_1 < \mu_2$) από την παράμετρο του 2^{ου} πληθυσμού ή η παράμετρος του 1^{ου} πληθυσμού είναι μεγαλύτερη ($\mu_1 > \mu_2$) από την παράμετρο του 2^{ου} πληθυσμού κ.λπ.

6.1.2.1 Διαδικασία ελέγχου υποθέσεων

Μια συνηθισμένη πορεία ελέγχου της μηδενικής υπόθεσης είναι η εξής:

1. Διατύπωση των δύο υποθέσεων, μηδενικής και εναλλακτικής.
2. Επιλέγεται ένα τυχαίο δείγμα.
3. Εφαρμόζουμε κάποια στατιστική συνάρτηση ελέγχου (κριτήριο: η κατανομή που θεωρούμε ότι ακολουθεί το φαινόμενο που μελετάμε, π.χ. κανονική κατανομή). Βρίσκουμε τη στατιστική τιμή η οποία είναι κοινή και για τους δύο τύπους ελέγχου υποθέσεων: Αμφίπλευρου και Μονόπλευρου ελέγχου.

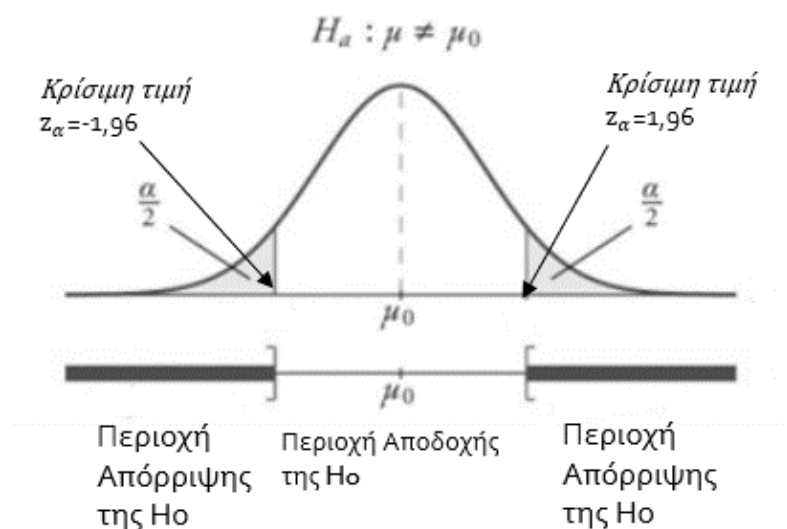
4. Στη συνέχεια, με βάση μια κρίσιμη τιμή καθορίζουμε την περιοχή απόρριψης της H_0 πέρα από την κρίσιμη τιμή. Η κρίσιμη τιμή είναι η τιμή της στατιστικής συνάρτησης ελέγχου για συγκεκριμένο επίπεδο σημαντικότητας (συνήθως $\alpha=5\%$).

Οι δυνατές τιμές που μπορεί να πάρει η στατιστική συνάρτηση χωρίζονται σε δύο ομάδες. Η πρώτη ομάδα αποτελεί την περιοχή της απόρριψης και η άλλη την περιοχή της αποδοχής (Εικόνες 6.3 & 6.4). Οι τιμές της περιοχής απόρριψης είναι αυτές που έχουν μικρή πιθανότητα να ληφθούν, όταν η μηδενική υπόθεση αληθεύει. Αντίθετα, οι τιμές της περιοχής αποδοχής έχουν μεγαλύτερη πιθανότητα να ληφθούν υπό τη μηδενική υπόθεση. Σύμφωνα με τον κανόνα απόφασης, αν η τιμή του στατιστικού ελέγχου (στατιστική τιμή) που υπολογίζεται από το δείγμα που διαθέτουμε ανήκει στην περιοχή απόρριψης, τότε απορρίπτουμε τη μηδενική υπόθεση. Αν, αντίθετα, η τιμή βρίσκεται στην περιοχή αποδοχής, τότε η μηδενική υπόθεση δεν απορρίπτεται.

Ο καθορισμός των δύο περιοχών γίνεται με βάση το επιθυμητό επίπεδο σημαντικότητας, που συμβολίζεται με α . Ο όρος *επίπεδο σημαντικότητας* προέρχεται από το γεγονός ότι μια τιμή του κριτηρίου ελέγχου, που ανήκει στην περιοχή απόρριψης, λέγεται σημαντική. Το επίπεδο σημαντικότητας εκφράζει το εμβαδόν που βρίσκεται ανάμεσα στην καμπύλη της κατανομής και το τμήμα του οριζώντιου άξονα, που αποτελεί την περιοχή απόρριψης της μηδενικής υπόθεσης. Εκφράζει, δηλαδή, την πιθανότητα να έχει εμφανιστεί το αποτέλεσμα από τυχαίους παράγοντες. Είναι φανερό ότι το α εκφράζει την πιθανότητα απόρριψης μιας αληθινής μηδενικής υπόθεσης. Επειδή το να απορριφθεί μια αληθινή μηδενική υπόθεση αποτελεί σφάλμα, είναι λογικό να ζητήσουμε μικρή πιθανότητα απόρριψης της μηδενικής υπόθεσης, όταν αυτή αληθεύει. Για τον λόγο αυτόν θα επιλέξουμε μια μικρή τιμή για το α . Οι πιο συχνά επιλεγόμενες τιμές του α είναι $0,01=1\%$, $0,05=5\%$ και σπάνια μεγαλύτερες, όπως το $0,10=10\%$ (Γιαλαμάς, 2005; Bryman, 2018; Field, 2018).

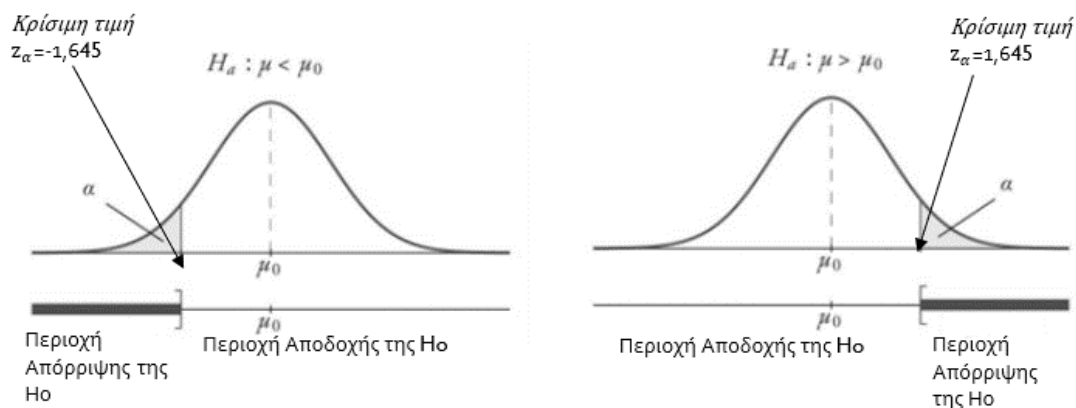
Στην περίπτωση που απορρίπτουμε μια μηδενική υπόθεση, αναφέρουμε ότι το αποτέλεσμα είναι στατιστικά σημαντικό. Αυτό σημαίνει πρακτικά ότι το φαινόμενο που εμείς μελετάμε δεν ακολουθεί μια τυχαία συμπεριφορά, αλλά εξηγείται βάσει κάποιου παράγοντα. Για παράδειγμα, αν υποθέσουμε ότι υπάρχει στατιστικά σημαντική σχέση μεταξύ του φύλου των μαθητών και της επίδοσής τους στα μαθηματικά, αυτό σημαίνει ότι η επίδοση των μαθητών στα δύο φύλα δεν κατανέμεται τυχαία και επομένως ανεξάρτητα του φύλου. Μάλιστα αξίζει να πούμε ότι, αν στην προηγούμενη σχέση το δείγμα ήταν και αντιπροσωπευτικό του πληθυσμού, το φύλο εξηγεί την προηγούμενη επίδοση στον αντίστοιχο πληθυσμό.

Επειδή η περιοχή απόρριψης αποτελείται από δύο μέρη, πρέπει ένα μέρος του α να συσχετιστεί με μεγάλες τιμές, ενώ το υπόλοιπο με μικρές. Έτσι είναι προφανές ότι θα πρέπει το α να διαιρεθεί σε ίσα μέρη, δηλαδή να συσχετιστεί το $\alpha/2=0,025$ με τις μικρές ακραίες τιμές και το $\alpha/2=0,025$ με τις μεγάλες ακραίες τιμές (Εικόνα 6.4).



Εικόνα 6.4 Κρίσιμες τιμές, περιοχές «αποδοχής» και «απόρριψης» της μηδενικής υπόθεσης για αμφίπλευρο έλεγχο με κανονική κατανομή σε επίπεδο σημαντικότητας α .

Ένας μονόπλευρος έλεγχος είναι λιγότερο αυστηρός απ' ότi ένας αμφίπλευρος. Η περιοχή απόρριψης ορίζεται εξ ολοκλήρου στη μία μόνο πλευρά της κατανομής της συνάρτησης του ελέγχου. Έτσι, το επίπεδο σημαντικότητας είναι διπλάσιο από το αντίστοιχο του αμφίπλευρου ελέγχου (Εικόνα 6.5). Επομένως, αν απορρίψουμε τη μηδενική υπόθεση αμφίπλευρου ελέγχου, έχουμε απορρίψει ταυτόχρονα και τη μηδενική υπόθεση μονόπλευρου ελέγχου. Για παράδειγμα, αν υποθέσουμε ότi η στατιστική συνάρτηση ελέγχου είναι η κανονική κατανομή, όπως φαίνεται στις Εικόνες 6.3 και 6.4, η κρίσιμη τιμή του αμφίπλευρου ελέγχου (1,96) είναι μεγαλύτερη της κρίσιμης τιμής του μονόπλευρου ελέγχου (1,645). Άρα, αν η στατιστική τιμή που υπολογίζουμε για τον έλεγχο της μηδενικής υπόθεσης είναι μεγαλύτερη του 1,96, τότε θα είναι μεγαλύτερη και του 1,645 και επομένως, θα απορρίπτεται η μηδενική υπόθεση αμφίπλευρου ελέγχου και ταυτόχρονα και η μηδενική υπόθεση μονόπλευρου ελέγχου. Ωστόσο, χρειάζεται προσοχή, αφού σε πολλές περιπτώσεις μονόπλευρου ελέγχου μπορεί να απορρίπτεται η μηδενική υπόθεση, ενώ στον αντίστοιχο αμφίπλευρο έλεγχο να μην απορρίπτεται. Υπό το πρίσμα αυτό ο μονόπλευρος έλεγχος πρέπει να αξιοποιείται σε περιπτώσεις που θέλουμε να υποστηρίξουμε συγκεκριμένη κατεύθυνση και δεν υπάρχει λόγος να ελέγχεται η αντίθετη. Για παράδειγμα, αν μας ενδιαφέρει να υποστηρίξουμε μια καινούρια διδακτική παρέμβαση έναντι μιας παραδοσιακής ή μια νέα φαρμακευτική αγωγή έναντι μιας προηγούμενης, είναι προφανές ότi χρειάζεται να διατυπώσουμε υποθέσεις μονόπλευρου ελέγχου, αναδεικνύοντας την καινούρια διδακτική παρέμβαση και τη νέα φαρμακευτική αγωγή έναντι της παραδοσιακής διδακτικής παρέμβασης και της προηγούμενης φαρμακευτικής αγωγής αντίστοιχα.



Εικόνα 6.5 Κρίσιμες τιμές, περιοχές «αποδοχής» και «απόρριψης» της μηδενικής υπόθεσης για μονόπλευρο έλεγχο με κανονική κατανομή σε επίπεδο σημαντικότητας α .

Το σφάλμα της λαθεμένης απόρριψης της μηδενικής υπόθεσης λέγεται *σφάλμα τύπου I* και συμβολίζεται με α (Εικόνα 6.6). Είναι, δηλαδή, η πιθανότητα απόρριψης της μηδενικής υπόθεσης ενώ αυτή ισχύει «P(απόρριψης της H_0 / H_0 είναι αληθής)». Το σφάλμα αυτό θεωρείται ως σφάλμα στο οποίο δεν πρέπει να υποπέσουμε και καθορίζεται από την αρχή για τον έλεγχο της μηδενικής υπόθεσης. Προσοχή στο ότi το μεγάλο δείγμα μπορεί να οδηγήσει σε σφάλμα τύπου I.

Το σφάλμα που γίνεται όταν δεχθούμε μια εσφαλμένη μηδενική υπόθεση λέγεται *σφάλμα τύπου II* και συμβολίζεται με β (Εικόνα 6.6). Είναι, δηλαδή, η πιθανότητα, αποδοχής της μηδενικής υπόθεσης, ενώ αυτή δεν ισχύει «P(μη απόρριψη της H_0 / H_0 είναι ψευδής)». Προσοχή στο ότi το μικρό δείγμα μπορεί να οδηγήσει σε σφάλμα τύπου II.

Κατά τον έλεγχο της μηδενικής υπόθεσης μας ενδιαφέρει να απορρίψουμε μια εσφαλμένη μηδενική υπόθεση. Η πιθανότητα αυτή λέγεται *ισχύς του ελέγχου* και συμβολίζεται $1-\beta$ (Εικόνα 6.6). Η ισχύς του ελέγχου αυξάνεται όταν αυξάνουμε το μέγεθος του δείγματος και μειώνεται όταν μειώνεται το επίπεδο σημαντικότητας του ελέγχου. Τέλος, όσο πιο μεγάλη είναι η διαφορά μεταξύ της «πραγματικής» τιμής της παραμέτρου και της τιμής που καθορίζεται στη μηδενική υπόθεση, τόσο πιο μεγάλη είναι η ισχύς του ελέγχου (Γιαλαμάς, 2005).

		Δυνατή κατάσταση για H_0	
		H_0 αληθεύει	H_0 εσφαλμένη
Στατιστική απόφαση	Απόρριψη H_0	Σφάλμα τύπου I (α)	Σωστή απόφαση ($1-\beta$)
	Αποδοχή H_0	Σωστή απόφαση ($1-\alpha$)	Σφάλμα τύπου II (β)

Εικόνα 6.6 Είδη σφαλμάτων κατά την απόρριψη ή μη της H_0 .

Στο περιβάλλον του SPSS, όπως θα δούμε στα επόμενα Κεφάλαια, δεν χρειάζεται να πραγματοποιήσουμε υπολογισμούς για τον έλεγχο των υποθέσεων. Η διαδικασία αυτή είναι αρκετά απλουστευμένη και βασίζεται κυρίως στην επιλογή του κατάλληλου στατιστικού ελέγχου που ταιριάζει κάθε φορά στον τύπο των μεταβλητών που εμπλέκονται στον έλεγχο. Στην επόμενη ενότητα του κεφαλαίου αυτού θα αποσαφηνίσουμε τους σχεδιασμούς ανεξάρτητων και εξαρτημένων δειγμάτων στα οποία βασίζεται ένα μεγάλο μέρος των ελέγχων της επαγωγικής στατιστικής.

6.2 Σχεδιασμοί ανεξάρτητων και εξαρτημένων δειγμάτων

Στον σχεδιασμό ανεξάρτητων (independent) δειγμάτων η επιλογή των μονάδων ανάλυσης σε ένα δείγμα δεν προσδιορίζεται από την επιλογή των μονάδων ανάλυσης στα άλλα δείγματα. Πρακτικά, τα εξεταζόμενα υποκείμενα από τα δείγματα είναι διαφορετικά και δεν υπάρχει κάποια σχέση μεταξύ τους. Κλασικό παράδειγμα αποτελεί η πειραματική διαδικασία για τη σύγκριση μιας μέτρησης που πραγματοποιείται στην πειραματική ομάδα (experiment) και στην ομάδα ελέγχου (control). Σύμφωνα με τη διαδικασία αυτή, κάθε υποκείμενο έχει τοποθετηθεί με τυχαίο τρόπο σε μία μόνο από τις παραπάνω ομάδες (Γιαλαμάς, 2005; Bryman, 2018; Field, 2018).

Σε συνέχεια του προηγούμενου παραδείγματος, αν μας ενδιαφέρει να υποδείξουμε ποια από τις δύο προηγούμενες ομάδες έχει υψηλότερη επίδοση σε ένα τεστ, τότε η βασική μας μέτρηση είναι «η επίδοση στο τεστ» (εξαρτημένη μεταβλητή), όπου οι τιμές αυτού του τεστ διακρίνονται σε δύο ανεξάρτητα δείγματα (Εικόνα 6.7) με τη βοήθεια της μεταβλητής «ομάδα» (ανεξάρτητη μεταβλητή).

Ομάδα	Epidosh_test
Ομάδα Ελέγχου	89
Ομάδα Ελέγχου	76
Ομάδα Ελέγχου	56
Ομάδα Ελέγχου	34
Ομάδα Ελέγχου	67
Ομάδα Ελέγχου	79
Ομάδα Ελέγχου	80
Ομάδα Ελέγχου	69
Ομάδα Ελέγχου	90
Ομάδα Ελέγχου	51
Πειραματική Ομάδα	78
Πειραματική Ομάδα	45
Πειραματική Ομάδα	89
Πειραματική Ομάδα	90
Πειραματική Ομάδα	98
Πειραματική Ομάδα	97
Πειραματική Ομάδα	97
Πειραματική Ομάδα	69
Πειραματική Ομάδα	80
Πειραματική Ομάδα	67

Εικόνα 6.7 Τα δύο ανεξάρτητα δείγματα της επίδοσης σύμφωνα με τη μεταβλητή «ομάδα».

Στον σχεδιασμό εξαρτημένων (dependent) δειγμάτων οι μονάδες ανάλυσης από το ένα δείγμα αντιστοιχούν στις μονάδες ανάλυσης στα άλλα δείγματα. Η μέτρηση είναι πολλαπλή (τουλάχιστον δύο μετρήσεις) για τα ίδια υποκείμενα. Κλασικό παράδειγμα αποτελεί η σύγκριση των μετρήσεων πριν και μετά κάποιας μεταχείρισης (Treatment). Αν η μεταχείριση αυτή είναι μια διδακτική παρέμβαση, θα επιθυμούσαμε για τα ίδια άτομα, μετά την παρέμβαση, η μέτρηση της επίδοσής τους να είναι υψηλότερη από την επίδοσή τους σε παρόμοιο ή ισοδύναμο τεστ πριν την παρέμβαση. Επομένως, η σύγκριση πραγματοποιείται στις δύο διαφορετικές μετρήσεις (μεταβλητές) της επίδοσης για τα ίδια άτομα, πριν και μετά. Τα δεδομένα των μετρήσεων του παραδείγματος θα μπορούσαν να είναι αυτά που φαίνονται στην Εικόνα 6.8. Άρα έχουμε δύο μεταβλητές, τις μετρήσεις της επίδοσης πριν και μετά για τα ίδια υποκείμενα.

Metrshsh_prin	Metrshsh_meta
56	78
57	85
45	75
78	90
87	85
67	65
56	70
84	90
74	95
90	89

Εικόνα 6.8 Τα εξαρτημένα δείγματα μέτρησης της επίδοσης πριν και μετά.

Σε πολλές περιπτώσεις μπορούμε να εφαρμόσουμε και σχεδιασμό ανεξάρτητων αλλά και σχεδιασμό εξαρτημένων δειγμάτων. Αξίζει να τονίσουμε, όμως, ότι ο σχεδιασμός εξαρτημένων δειγμάτων είναι καλύτερος, αφού με τον σχεδιασμό αυτόν ελέγχονται με μεγαλύτερη ασφάλεια οι όποιες εξωγενείς μεταβλητές. Στον σχεδιασμό εξαρτημένων δειγμάτων το ποσό της μεταβλητότητας που εξηγείται λόγω των διαφορετικών χαρακτηριστικών των υποκειμένων (εξωγενείς μεταβλητές) εξαλείφεται και έτσι μειώνεται η πιθανότητα να απορρίψουμε εσφαλμένα τη μηδενική υπόθεση (Σφάλμα τύπου I).

Για παράδειγμα, αν θέλω να ελέγξω τη διαφορά στην επίδοση μεταξύ αγοριών και κοριτσιών στη γλώσσα, τότε έχω τη δυνατότητα να σχεδιάσω μια μελέτη ανεξάρτητων ή μια μελέτη εξαρτημένων δειγμάτων. Στον σχεδιασμό ανεξάρτητων δειγμάτων θα διαλέξουμε τυχαία k αγόρια και k κορίτσια από έναν μεγάλο πληθυσμό μαθητών. Ωστόσο, η όποια διαφορά της επίδοσης μεταξύ αγοριών και κοριτσιών δεν θα εξηγείται μόνο μέσω του φύλου. Το σίγουρο είναι ότι τα αγόρια και τα κορίτσια αυτά έχουν πολλά άλλα διαφορετικά χαρακτηριστικά (π.χ. μορφωτικό και κοινωνικό επίπεδο γονέων κ.λπ.), που πιθανόν να εξηγούν την όποια διαφορά στην επίδοση στη γλώσσα. Η άλλη περίπτωση είναι να διαλέξουμε τυχαία k οικογένειες με δίδυμα διαφορετικού φύλου (αρκετά δύσκολο να βρούμε τέτοιο δείγμα, αλλά όχι αδύνατον), δηλαδή μιλάμε για σχεδιασμό εξαρτημένων δειγμάτων. Τότε η όποια διαφορά της επίδοσης μεταξύ αγοριών και κοριτσιών θα εξηγείται κυρίως μόνο μέσω του φύλου, αφού κάθε ζευγάρι παιδιών (αγόρι και κορίτσι), από τη στιγμή που μεγαλώνουν στην ίδια οικογένεια, θα έχουν πολλά κοινά χαρακτηριστικά.

6.3 Έλεγχος δεδομένων (Data Screening) με το SPSS

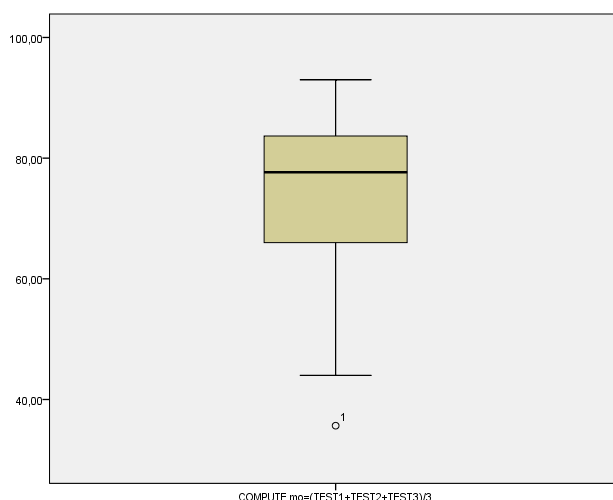
Μόλις εισαχθούν τα δεδομένα στο SPSS, πρέπει να τα ελέγξουμε για τυχόν σφάλματα κατά την εισαγωγή των δεδομένων αλλά και για τυχόν ακραίες τιμές. Ο έλεγχος της εισαγωγής των δεδομένων γίνεται συνήθως με τη δημιουργία πινάκων κατανομής συχνοτήτων όλων των μεταβλητών προκειμένου να εντοπίσουμε τιμές που δεν ταιριάζουν με το πλαίσιο της μέτρησης που χρησιμοποιήθηκε ή είναι ακραίες. Για παράδειγμα, στη μέτρηση

της ηλικίας των φοιτητών σε έρευνα μπορεί να δηλωθεί ως ηλικία μια τιμή (π.χ. 12 έτη) που είναι εκτός του εύρους των ηλικιών των φοιτητών.

Ο προσδιορισμός των ακραίων τιμών (outliers) είναι πολύ σημαντικός, αφού η μέση τιμή που είναι ένας βασικός εκτιμητής των τιμών των ποσοτικών μεταβλητών επηρεάζεται έντονα από τις ακραίες τιμές (Γναρδέλλης, 2013; Bryman, 2016; IBM SPSS Statistics Base 25, n.d.). Ενδεικτικό παράδειγμα είναι η λανθασμένη χρήση της μέσης τιμής για τον προσδιορισμό του μέσου εισοδήματος των πολιτών μιας χώρας. Η περίπτωση, δηλαδή, που τα λίγα υψηλά εισοδήματα (θετική ασυμμετρία) «τραβούν τον μέσο όρο ψηλά» και παρουσιάζουν έτσι μια λανθασμένη εικόνα του εισοδήματος των πολιτών μιας χώρας. Υπό αυτή την έννοια, οι ακραίες τιμές μπορεί να διαστρεβλώσουν τα αποτελέσματα της έρευνας και πολλές φορές πρέπει να εξαιρούνται από την ανάλυση. Στην περίπτωση που κάποιες ακραίες τιμές δεν μπορούν να εξαιρεθούν από την ανάλυση συμπληρωματικά θα μπορούσε να χρησιμοποιηθεί η διάμεσος ως εκτιμητής των τιμών της μεταβλητής. Ο προσδιορισμός των ακραίων τιμών θα πρέπει να γίνεται κυρίως σε όλες τις ποσοτικές μεταβλητές (τύπου scale). Τιμές που βρίσκονται εντός του διαστήματος: $[Q1-1,5*IQR, Q3+1,5*IQR]$ είναι συνήθως αποδεκτές. Όπου $Q1$, $Q3$, το πρώτο και το τρίτο τεταρτημόριο και IQR το ενδοτεταρτημοριακό εύρος της κατανομής των τιμών.

Για να εντοπίσουμε τις ακραίες τιμές:

- **1ος τρόπος:** Με τη βοήθεια του θηκόγραμματος (boxplot). Από το μενού επιλέγουμε διαδοχικά **Graphs=> Legacy Dialogs=> boxplot** και μετά επιλέγουμε **Simple** και **Summaries of separate variables**. Οι ακραίες τιμές που υποδεικνύονται με έναν κύκλο ή αστεράκι είναι αυτές που βρίσκονται εκτός των ορίων $[Q1-1,5*IQR, Q3+1,5*IQR]$. Στο θηκόγραμμα που φαίνεται στην Εικόνα 6.9 μπορούμε να δούμε μια ακραία τιμή της μεταβλητής «mo» (για τη δημιουργία του γραφήματος χρησιμοποιήθηκε η μεταβλητή «mo» από το αρχείο «marks_trans_2.sav») που βρίσκεται στον αριθμό γραμμής 1.



Εικόνα 6.9 Εξαγόμενο θηκόγραμμα από τη διαδικασία «boxplot».

- **2ος τρόπος:** Μέσω της εντολής διερεύνησης της κατανομής ποσοτικών μεταβλητών, δηλαδή τη διαδικασία **Explore**. Από το μενού επιλέγουμε διαδοχικά **Analyze=> Descriptive Statistics=>Explore**, έπειτα επιλέγουμε **Statistics** και στη συνέχεια **Outliers**. Εμφανίζονται σε πίνακα οι πιθανές ακραίες τιμές, καθώς επίσης και ο αριθμός γραμμής για να μπορέσουμε να τις εντοπίσουμε.

Για την ίδια μεταβλητή «mo» μέσω της διαδικασίας **Explore** στα εξαγόμενα αποτελέσματα μπορούμε να δούμε έναν πίνακα και ένα θηκόγραμμα, όπως το προηγούμενο. Στον εξαγόμενο Πίνακα (Εικόνα 6.10) θα δούμε τις πέντε μεγάλες και τις πέντε μικρές τιμές (Value) της μεταβλητής (δηλαδή τις πιθανές ακραίες τιμές), καθώς επίσης και τον αριθμό γραμμής (Case number) κάθε τιμής στο φύλλο δεδομένων του SPSS. Η μικρότερη τιμή, 35,67, σημειώθηκε με σύμβολο κύκλου στο προηγούμενο θηκόγραμμα.

Extreme Values				
			Case Number	Value
COMPUTE mo=(TEST1+TEST2+TEST3)/3	Highest	1	59	93,00
		2	60	93,00
		3	61	93,00
		4	58	91,33
		5	57	90,00
	Lowest	1	1	35,67
		2	2	44,00
		3	3	45,00
		4	4	48,33
		5	5	52,00

Εικόνα 6.10 Εξαγόμενος πίνακας της διαδικασίας «Explore => Outliers».

- **3ος τρόπος:** Υπολογίζουμε τις τυπικές τιμές (z-values) των μεταβλητών. Δημιουργούνται αυτόματα καινούριες μεταβλητές με τις τυπικές τιμές των μεταβλητών. Από το μενού επιλέγουμε διαδοχικά [Analyze => Descriptive Statistics => Descriptives](#) και έπειτα επιλέγουμε κάτω και αριστερά [Save standardized values as variables](#). Οι τυπικές τιμές αντιστοιχούν στην απόσταση της τιμής της κατανομής από τη μέση τιμή της κατανομής, με μονάδα μέτρησης αυτής της απόστασης την τυπική απόκλιση της κατανομής. Οι τυπικές τιμές συνήθως εκτός του ορίου (-3, 3) υποδεικνύουν τις πιθανές ακραίες τιμές.

Ακολουθώντας αυτή τη διαδικασία για την ίδια μεταβλητή «mo», δημιουργείται μια καινούρια μεταβλητή με όνομα «zmo». Στην Εικόνα 6.11 μπορούμε να δούμε την καινούρια μεταβλητή με τις αντίστοιχες τυπικές τιμές και επιπλέον μπορούμε να δούμε ότι η τιμή με αριθμό γραμμής 1 (η πρώτη στη στήλη), έχει τυπική τιμή πολύ κοντά στο -3.

Zmo
-2,90683
-2,27751
-2,20199
-1,95027
-1,67337
-1,64819
-1,59785
-1,34612
-1,29578
-1,16991
-1,09439
-,79232
-,71680
-,66646

Εικόνα 6.11 Καινούρια μεταβλητή «zmo» με τις τυπικές τιμές της μεταβλητής «mo».

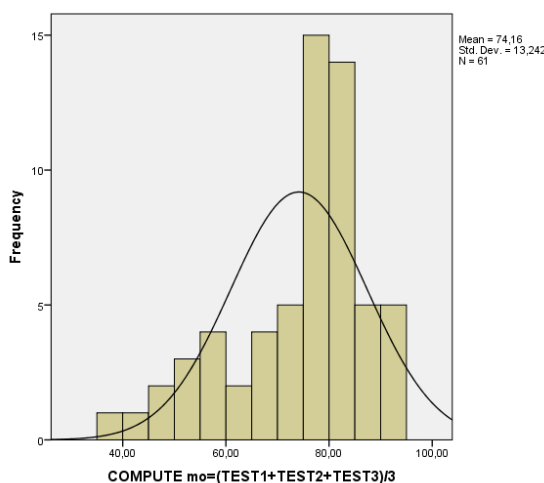
6.3.1 Έλεγχος της κανονικότητας των τιμών μιας ποσοτικής μεταβλητής

Μία από τις κύριες προϋποθέσεις για την αξιοποίηση μεταβλητών σε πολλούς στατιστικούς ελέγχους είναι ότι οι τιμές αυτών των μεταβλητών πρέπει να είναι κανονικά κατανοημένες. Το εάν τα δεδομένα προέρχονται από κανονικά κατανοημένο πληθυσμό ή όχι, μπορεί να ελεγχθεί με διαφορετικές μεθόδους. Κυρίως οι μέθοδοι αυτοί μας βοηθούν να στηρίξουμε τη συμμετρία, η οποία είναι προϋπόθεση της κανονικότητας (Γιαλαμάς, 2005; Γναρδέλλης, 2013; Field, 2018).

Οι συνηθέστερες μέθοδοι ελέγχου της κανονικότητας ενός συνόλου δεδομένων είναι:

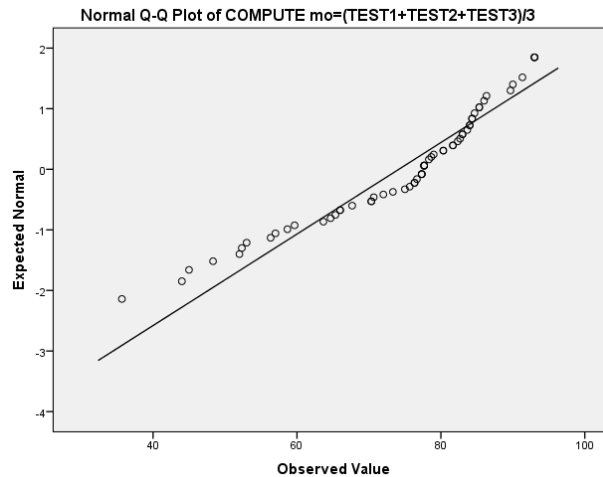
- Ελέγχουμε κατά πόσο το ιστόγραμμα των τιμών της ποσοτικής μεταβλητής είναι συμμετρικό. Για τη δημιουργία του ιστογράμματος επιλέγουμε από το μενού διαδοχικά: **Graphs => Legacy Dialogs => Histogram**, τοποθετούμε τη μεταβλητή (στο παράδειγμά μας τη μεταβλητή «mo»), καθώς επίσης επιλέγουμε να εμφανιστεί η καμπύλη της κανονικής κατανομής (**Display Normal Curve**). Η αυτόματη εμφάνιση της καμπύλης της κανονικής κατανομής (normal curve) των τιμών της μεταβλητής (γραφική παράσταση της εκθετικής συνάρτησης, λαμβάνοντας υπόψη τη μέση τιμή και την τυπική απόκλιση των τιμών της μεταβλητής), συμβάλλει στον κατά προσέγγιση έλεγχο του αν και κατά πόσο η κατανομή των τιμών αποκλίνει από την καμπύλη της κανονικής κατανομής.

Στην Εικόνα 6.12 παρουσιάζεται το γράφημα που ζητήσαμε. Όπως φαίνεται, η κατανομή των τιμών της μεταβλητής παρουσιάζει μια μικρή αρνητική ασυμετρία (κάποιες μικρές τιμές «τραβούν» την κατανομή αριστερά).



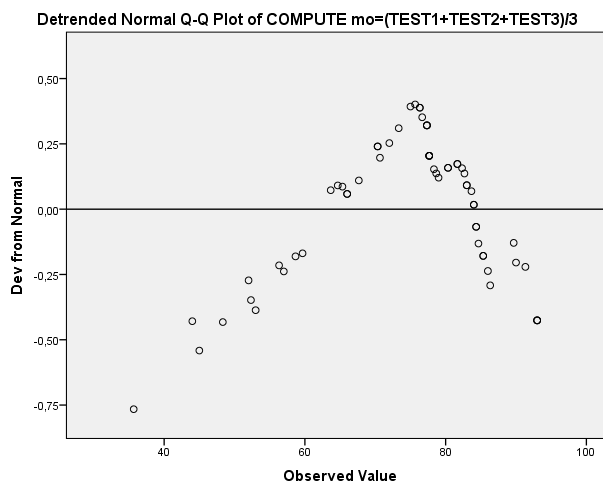
Εικόνα 6.12 Το εξαγόμενο ιστόγραμμα και η καμπύλη της κανονικής κατανομής.

- Ελέγχουμε το θηκόγραμμα (Box plot). Κυρίως, θα πρέπει η διάμεσος να χωρίζει το κουτί στη μέση και να μην υπάρχουν ακραίες παρατηρήσεις.
- Θα πρέπει οι λόγοι λοξότητας (Skewness) προς το αντίστοιχο τυπικό της σφάλμα (Std. Error) και κυρτότητας (Kurtosis) προς το αντίστοιχο τυπικό της σφάλμα (standard error) να είναι μεταξύ [-2, 2]. Τις τιμές αυτές μπορείτε να τις υπολογίσετε μέσω της διαδικασίας **Explore**.
- Στο γράφημα Normal Q-Q Plot που εμφανίζεται αυτόματα με τη διαδικασία **Analyze => Descriptive Statistic => Explore => plots** ελέγχουμε αν τα σημεία βρίσκονται στη διχοτόμο της γωνίας. Στην Εικόνα 6.13 παρουσιάζεται το γράφημα που ζητήσαμε. Όπως φαίνεται, η κατανομή των τιμών της μεταβλητής «mo» δεν ακολουθεί την κανονική κατανομή, αφού πολλά σημεία βρίσκονται σε απόσταση από τη διχοτόμο (ευθεία γραμμή στο γράφημα).



Εικόνα 6.13 Normal Q-Q Plot, για τον έλεγχο της κανονικότητας των τιμών της μεταβλητής «mo».

- Στο γράφημα Detrended Normal Q-Q Plot, που εμφανίζεται αυτόματα με τη διαδικασία [Analyze => Descriptive Statistic => Explore => plots](#), τα σημεία εκατέρωθεν της ευθείας πρέπει να μην σχηματίζουν κάποιο συγκεκριμένο πρότυπο. Στην Εικόνα 6.14 παρουσιάζεται το γράφημα που ζητήσαμε. Όπως φαίνεται, η κατανομή των τιμών της μεταβλητής «mo» δεν ακολουθεί την κανονική κατανομή, αφού τα σημεία φαίνεται να αναπαριστούν μια παραβολή «ανάποδο ποτήρι».



Εικόνα 6.14 Detrended Normal Q-Q Plot για τον έλεγχο της κανονικότητας των τιμών της μεταβλητής «mo».

- Έλεγχοι Kolmogorov – Smirnov, Lilliefors και Shapiro – Wilk. Από το μενού επιλέγουμε [Analyze => Descriptive Statistic => Explore => plots](#) και, στη συνέχεια, επιλέγουμε [Normality test with plots](#). Για να μην έχουμε απόκλιση από την κανονικότητα, θα πρέπει να μην απορρίπτεται η μηδενική υπόθεση του ελέγχου: Αν το δείγμα μας έχει μέγεθος μικρότερο ή ίσο του 50, τότε το καταλληλότερο στατιστικό κριτήριο, για να ελέγξει την ύπαρξη (αποδοχή της μηδενικής υπόθεσης) της κανονικότητας, είναι το Shapiro – Wilk.
- Στην περίπτωση που το δείγμα μας ξεπερνά το 50, τότε το καταλληλότερο στατιστικό κριτήριο για να ελεγχθεί η ύπαρξη της κανονικότητας είναι το Kolmogorov – Smirnov. Το στατιστικό κριτήριο Lilliefors που θα βλέπουμε στους αντίστοιχους πίνακες είναι μια διόρθωση του Kolmogorov – Smirnov.

Στην Εικόνα 6.15 παρουσιάζεται ο πίνακας ελέγχου που ζητήσαμε. Όπως φαίνεται η κατανομή των τιμών της μεταβλητής «mo» δεν ακολουθεί την κανονική κατανομή, αφού απορρίπτεται η μηδενική υπόθεση

(sig.<0,05). Η μηδενική υπόθεση που ελέγχεται είναι ότι οι τιμές της κατανομής της συγκεκριμένης μεταβλητής (mo) ταιριάζουν με τις τιμές που τυχαία θα προκύψουν από μια κανονική κατανομή με μέση τιμή και τυπική απόκλιση ίδια με αυτή της ελεγχόμενης μεταβλητής.

Tests of Normality						
	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
COMPUTE mo = (TEST1+TEST2+TEST3) /3	,172	61	,000	,919	61	,001
a. Lilliefors Significance Correction						

Εικόνα 6.15 Εξαγόμενος πίνακας ελέγχου της κανονικότητας των τιμών της μεταβλητής «mo» με τους ελέγχους Kolmogorov – Smirnov, Lilliefors και Shapiro – Wilk.

6.4 Προβλήματα για εξάσκηση και ανατροφοδότηση

Άσκηση 1

Επιλέξτε τον κατάλληλο σχεδιασμό, ανεξάρτητων ή εξαρτημένων δειγμάτων, για τις επόμενες ερευνητικές υποθέσεις. Σε κάθε περίπτωση να δικαιολογήσετε την απάντησή σας.

- Η μέση επίδοση των συμμετεχόντων στο posttest είναι διαφορετική από τη μέση επίδοση των ίδιων συμμετεχόντων στο pretest.
- Οι άνθρωποι μιλούν περισσότερη ώρα στο τηλέφωνο σε άντρες πωλητές απ' ότι σε γυναίκες πωλήτριες.
- Απόφοιτοι μηχανικοί είχαν υψηλότερο μέσο μισθό 10 χρόνια μετά τις σπουδές απ' ότι είχαν 5 χρόνια μετά τις σπουδές.
- Οι πατεράδες επικοινωνούν λιγότερο με τον εκπαιδευτικό του παιδιού τους απ' ότι οι μητέρες.
- Το βάρος των μαθητών μειώθηκε από την τρίτη λυκείου στο πρώτο έτος του πανεπιστημίου.
- Τα κορίτσια μελετούν περισσότερες ώρες από τα αγόρια.
- Το προσδόκιμο ζωής των ανθρώπων που δεν έχουν ποτέ παντρευτεί είναι διαφορετικό από αυτό των ανθρώπων που έχουν παντρευτεί.
- Οι άντρες σύζυγοι κοιμούνται περισσότερες ώρες τη νύχτα απ' ότι οι γυναίκες σύζυγοι.
- Οι άνθρωποι που ζουν στον Βόρειο Πόλο έχουν μεγαλύτερη διάρκεια ζωής από τους ανθρώπους που ζουν στον Νότιο Πόλο.
- Πρώτο παιδί διδύμων έχει διαφορετική επίδοση στο σχολείο από το δεύτερο παιδί διδύμων.
- Οι καθολικές γυναίκες γεννούν περισσότερα παιδιά από τις ορθόδοξες γυναίκες.
- Οι γιατροί στην Ελλάδα έχουν διαφορετικό ετήσιο εισόδημα από τους γιατρούς στην Ιταλία.
- Σε όλο τον κόσμο το μέσο πλήθος μελών ΔΕΠ των Τμημάτων Προσχολικής Εκπαίδευσης είναι μεγαλύτερο από το μέσο πλήθος μελών ΔΕΠ των Τμημάτων Φιλολογίας.

Άσκηση 2

Δίνεται το 95% διάστημα εμπιστοσύνης [25, 29] της ηλικίας των μεταπτυχιακών φοιτητών ενός τμήματος φοιτητών. Να απαντήστε στα επόμενα ερωτήματα:

- Τι αντιπροσωπεύει το διάστημα αυτό;
- Αν μας ενδιαφέρει να ελέγχουμε κατά πόσο η μέση ηλικία των μεταπτυχιακών φοιτητών του πληθυσμού από τον οποίο επιλέχθηκε το δείγμα είναι μεγαλύτερη των 30 ετών, ποιο νομίζετε ότι θα είναι το αποτέλεσμα του ελέγχου;

Αν το διάστημα εμπιστοσύνης της διαφοράς της ηλικίας αντρών και γυναικών (Αντρών-Γυναικών) των παραπάνω μεταπτυχιακών φοιτητών είναι [2, 2,9], να ελέγξετε:

- Αν οι άντρες στον αντίστοιχο πληθυσμό των μεταπτυχιακών φοιτητών έχουν μεγαλύτερη ηλικία από τις γυναίκες.

Άσκηση 3

Στο αρχείο «marks_trans_2.sav» και μόνο για τις γυναίκες του δείγματος, να πραγματοποιήσετε τα επόμενα:

A. Να ελέγξετε την κανονικότητα των τιμών της μεταβλητής PCM με τους παρακάτω τρόπους:

- Υπολογίστε τους λόγους λοξότητας προς το αντίστοιχο τυπικό της σφάλμα και κυρτότητας προς το αντίστοιχο τυπικό της σφάλμα και ελέγξτε αν είναι μεταξύ $[-2, 2]$.
- Δημιουργήστε το αντίστοιχο ιστόγραμμα και ελέγξτε κατά πόσο είναι συμμετρικό, υποδεικνύοντας μια περίπου κωδωνοειδή καμπύλη.

B. Να υπολογίστε τις τυπικές τιμές αυτής της μεταβλητής για τις γυναίκες του δείγματος και υποδείξτε μέσω αυτών αν υπάρχουν ακραίες τιμές της κατανομής της.

Βιβλιογραφία

Bryman, A. (2016). *Social research methods*. London: Oxford University Press.

Γιαλαμάς, Β. (2005). *Στατιστικές Τεχνικές και Εφαρμογές στις Επιστήμες της Αγωγής*. Αθήνα: Εκδ. Πατάκη.

Γναρδέλλης, Χ. (2013). *Ανάλυση δεδομένων με το IBM SPSS Statistics 21*. Αθήνα: Εκδ. Παπαζήση.

Field, A. (2018). *Discovering statistics using IBM SPSS statistics* (5th ed.). London: SAGE.

IBM SPSS Statistics Base 25 (n.d.). Retrieved from
https://faculty.ksu.edu.sa/sites/default/files/ibm_spss_statistics_base.pdf

Κεφάλαιο 7

Έλεγχος μέσων τιμών (T-Tests)

Σύνοψη

Στο κεφάλαιο αυτό εξετάζονται τα τρία είδη παραμετρικών ελέγχων *t-test*, τουτέστιν ο έλεγχος *t-test* ενός δείγματος, ο έλεγχος *t-test* δύο ανεξάρτητων δειγμάτων και ο έλεγχος *t-test* δύο εξαρτημένων δειγμάτων ή δειγμάτων κατά ζεύγη. Παρουσιάζονται και ελέγχονται οι προϋποθέσεις εκτέλεσης των ελέγχων αυτών με τη βοήθεια των τεστ «Kolmogorov-Smirnov» ή/και «Shapiro-Wilk», καθώς και με τη χρήση θηκογραμμάτων. Επίσης, παρουσιάζονται και ερμηνεύονται τα αποτελέσματα των αναλύσεων τόσο των περιγραφικών στατιστικών όσο και των αντίστοιχων ελέγχων. Τέλος, περιγράφεται ο τρόπος συγγραφής των αποτελεσμάτων για κάθε έλεγχο *t-test* στο πλαίσιο μιας επιστημονικής εργασίας.

Προαπαιτούμενη γνώση

Μετασχηματισμός δεδομένων, περιγραφική στατιστική και εισαγωγή στην εκτιμητική και στον έλεγχο των υποθέσεων: Κεφάλαια 2, 4, και 6 του συγγράμματος.

7.1 Εισαγωγή στο *t-test*

Το *t-test*, γνωστό και ως *t-test* του Student, είναι ένα εργαλείο για την αξιολόγηση των μέσων τιμών ενός ή δύο πληθυσμών ακολουθώντας τον έλεγχο υποθέσεων. Ένα *t-test* μπορεί να χρησιμοποιηθεί για να αξιολογηθεί εάν μια μέση τιμή μιας μεμονωμένης ομάδας διαφέρει από μια γνωστή τιμή (*t-test* ενός δείγματος), εάν δύο ομάδες διαφέρουν μεταξύ τους (*t-test* δύο ανεξάρτητων δειγμάτων) ή εάν υπάρχει σημαντική διαφορά σε ζευγαρωτές μετρήσεις (*t-test* ζευγαρωτών ή εξαρτημένων δειγμάτων).

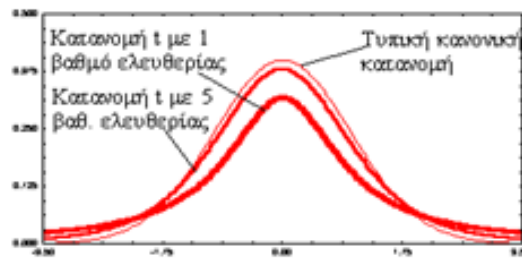
7.1.1 Έλεγχος *t* ενός δείγματος – One-Sample T-Test

Στον έλεγχο *t* ενός δείγματος (One-Sample T-Test) υπολογίζουμε τη μέση τιμή ($\bar{x}$) μιας ποσοτικής μεταβλητής του δείγματος της έρευνας που διεξάγουμε και ελέγχουμε τη μηδενική υπόθεση ότι η πληθυσμιακή μέση τιμή μ είναι ίση με την τιμή μ_0 . Λόγω του ότι στην πράξη σπάνια γνωρίζουμε την τυπική απόκλιση ενός άγνωστου κατά τα άλλα πληθυσμού (σ), πρέπει να εκτιμήσουμε την τιμή της (σ) από τα δειγματικά δεδομένα (Γιαλαμάς, 2005; Γναρδέλλης, 2013). Ένας πολύ καλός εκτιμητής της (σ) είναι η δειγματική τυπική απόκλιση (s). Σε αυτή την περίπτωση, θα γίνει χρήση της θεωρητικής κατανομής με το όνομα Student ή *t*-κατανομή. Λέγεται έτσι, διότι το 1908 ο Gosset δημοσίευσε τη μέθοδο αυτή στο περιοδικό Biometrika με το ψευδώνυμο «student» (https://en.wikipedia.org/wiki/William_Sealy_Gosset).

Το στατιστικό του ελέγχου ονομάζεται *t* και δίνεται από τη σχέση:

$$t = \frac{\bar{x} - \mu_0}{\frac{s}{\sqrt{N}}} \quad (7.1)$$

Η κατανομή *t* είναι μια κωδωνοειδής κατανομή με μικρή διαφορά από την τυπική κανονική κατανομή, με μέση τιμή 0 και τυπική απόκλιση λίγο μεγαλύτερη από τη μονάδα (Εικόνα 7.1). Για κάθε μέγεθος δείγματος *N* η ποσότητα *N*-1 εκφράζει τους βαθμούς ελευθερίας της *t*. Για κάθε τιμή *N*-1 έχουμε διαφορετική κατανομή *t* και διαφορετικές κρίσιμες τιμές. Όσο το μέγεθος του δείγματος ή οι βαθμοί ελευθερίας αυξάνονται, τόσο η κατανομή *t* προσεγγίζει την τυπική κανονική κατανομή (Γιαλαμάς, 2005).



Εικόνα 7.1 Τυπική κανονική κατανομή και κατανομές t με 1 και 5 βαθμούς ελευθερίας.

Ο έλεγχος, αν η μέση τιμή του πληθυσμού διαφέρει από μια συγκεκριμένη προκαθορισμένη τιμή μ_0 , γίνεται μέσω της μηδενικής υπόθεσης:

$$H_0: \mu = \mu_0,$$

έναντι της εναλλακτικής υπόθεσης:

$$H_A: \mu \neq \mu_0$$

Η ποσότητα t που είδαμε παραπάνω ακολουθεί την κατανομή t με $N-1$ βαθμούς ελευθερίας (df), όταν ισχύει η μηδενική υπόθεση.

Αφού είναι γνωστή η κατανομή του στατιστικού του ελέγχου, μπορεί να υπολογιστεί η πιθανότητα p , αν πάρουμε μια τιμή του στατιστικού το ίδιο ή και περισσότερο ακραία από την τιμή που υπολογίζεται από τα δεδομένα του δείγματος. Η πιθανότητα p ονομάζεται *παρατηρούμενο επίπεδο σημαντικότητας* και μπορεί να υπολογιστεί με ακρίβεια από τα σύγχρονα στατιστικά πακέτα, όπως για παράδειγμα το SPSS (Γιαλαμάς, 2005). Λόγω του ότι κατά την απόρριψη της μηδενικής υπόθεσης η τιμή του στατιστικού είναι περισσότερο ακραία από την κρίσιμη τιμή, το εμβαδόν της επιφάνειας που ορίζεται από τη θέση της υπολογισμένης τιμής έως το άκρο της κατανομής είναι μικρότερο από το επίπεδο σημαντικότητας α . Επομένως, στις τιμές του στατιστικού του ελέγχου που είναι μεγαλύτερες από την κρίσιμη τιμή αντιστοιχούν παρατηρούμενα επίπεδα σημαντικότητας p μικρότερα από το θεωρητικό επίπεδο σημαντικότητας α . Αντίστοιχα, όταν η τιμή p (p -value) είναι μικρότερη ή ίση με το α , τότε απορρίπτεται η μηδενική υπόθεση (Γιαλαμάς, 2005). Στην περίπτωση αυτή, ο έλεγχος ονομάζεται *στατιστικά σημαντικός*. Η πιθανότητα p του ελέγχου εμφανίζεται στα λογισμικά ως *Sig. (2-tailed)*. Η έκβαση του ελέγχου στο SPSS, στην περίπτωση του αμφίπλευρου ελέγχου, θα κρίνεται πάντα από τη σύγκριση της τιμής *Sig. (2-tailed)* με την τιμή του α . Για παράδειγμα, εάν το α ισούται με 5% ($\alpha = 0,05$) και η τιμή p (ή *Sig. (2-tailed)*) που υπολογίστηκε ισούται με 0,04, τότε απορρίπτουμε τη μηδενική υπόθεση και κατά συνέπεια ο έλεγχος θεωρείται στατιστικά σημαντικός. Όσο μικρότερη είναι η τιμή p , τόσο πιο σημαντικό είναι το αποτέλεσμα του ελέγχου. Ουσιαστικά, η τιμή p εκφράζει την πιθανότητα σφάλματος για τη γενίκευση του στατιστικού αποτελέσματος στον πληθυσμό. Λόγω του ότι στο SPSS χρησιμοποιείται μόνο ο αμφίπλευρος έλεγχος, στην περίπτωση που θέλουμε να εκτελέσουμε μονόπλευρο έλεγχο, η τιμή p του ελέγχου είναι το 1/2 της τιμής *Sig. (2-tailed)* που δίνεται από το SPSS, δηλαδή του αμφίπλευρου ελέγχου.

Αλλάζοντας την τιμή του επιπέδου σημαντικότητας α , δηλαδή την τιμή της πιθανότητας του σφάλματος τύπου I, ένας έλεγχος μπορεί να καταστεί από στατιστικά μη σημαντικός σε στατιστικά σημαντικός και το αντίστροφο. Επομένως η τιμή του α πρέπει να καθορίζεται εκ των προτέρων (*a priori*), δηλαδή πριν την εκτέλεση του ελέγχου, σύμφωνα με την ηθική και δεοντολογία της έρευνας.

Ο έλεγχος t -test ανήκει στην κατηγορία των παραμετρικών ελέγχων. Τα παραμετρικά τεστ αποτελούν πολύ ισχυρά στατιστικά κριτήρια, που για να χρησιμοποιηθούν απαιτούν την ικανοποίηση συγκεκριμένων προϋποθέσεων αναφορικά με συγκεκριμένες παραμέτρους του πληθυσμού από τον οποίο προέρχεται το δείγμα (Γιαλαμάς, 2005; Γναρδέλλης, 2013). Αυτές οι προϋποθέσεις είναι:

- Οι τιμές της μεταβλητής που μας ενδιαφέρει, θα πρέπει να προέρχονται από ανεξάρτητες παρατηρήσεις. Δύο παρατηρήσεις είναι ανεξάρτητες, αν δεν υπάρχει ένας σταθερός προβλεπτικός παράγοντας που να επηρεάζει τη σχέση ανάμεσά τους. Για παράδειγμα, αν οι 200 τιμές της επίδοσης στα μαθηματικά προέρχονται από μαθητές που φοιτούν σε ένα τυχαίο δείγμα 7 τάξεων Α' Γυμνασίου ενός νομού, οι τιμές αυτές δεν αποτελούν ανεξάρτητες

παρατηρήσεις λόγω του παράγοντα «τάξη» (διαφορετικές τάξεις έχουν πιθανώς διαφορετικό τρόπο διδασκαλίας, διαφορετικές επιλογές αξιολόγησης του διδάσκοντα κ.λπ.), που ενδεχομένως να επιδρά στην επίδοση των μαθητών. Η προϋπόθεση των ανεξάρτητων παρατηρήσεων ικανοποιείται τις περισσότερες φορές με την επιλογή τυχαίου δείγματος. Επίσης, το τυχαίο δείγμα σημαίνει ότι είναι αντιπροσωπευτικό του πληθυσμού και κατά συνέπεια τα αποτελέσματα μπορούν να γενικευτούν στον πληθυσμό.

- Η μέτρηση της μεταβλητής που μας ενδιαφέρει θα πρέπει να γίνεται σε κλίμακα τουλάχιστον ίσων διαστημάτων.
- Το δείγμα των τιμών της μεταβλητής που μας ενδιαφέρει θα πρέπει να προέρχεται από πληθυσμό που ακολουθεί κατά προσέγγιση την κανονική κατανομή. Η παραβίαση της προϋπόθεσης αυτής δεν δημιουργεί προβλήματα στην αξία των αποτελεσμάτων, ειδικά όταν έχουμε μεγάλο μέγεθος δείγματος. Δείγματα επαρκώς μεγάλα θεωρούνται τα δείγματα που είναι τουλάχιστον μεγέθους 30 (Γναρδέλλης, 2013). Στα πολύ μικρά δείγματα η κανονικότητα της κατανομής είναι σημαντική και σε αυτά η κανονικότητα πρέπει να εξασφαλίζεται κατά προσέγγιση.
- Οι ομάδες τιμών της μεταβλητής που μας ενδιαφέρει (όπως διακρίνονται από κάποια ποιοτική μεταβλητή με τουλάχιστον δύο τιμές) θα πρέπει να έχουν ίσες διακυμάνσεις.

Η εκτέλεση του ελέγχου t ενός δείγματος στο SPSS θα γίνει με το ακόλουθο παράδειγμα. Παρακάτω παρουσιάζεται ο Πίνακας με τα δεδομένα της μεταβλητής «stress». Οι τιμές της μεταβλητής αυτής εκφράζουν το επίπεδο άγχους εννέα (9) φοιτητών ανθρωπιστικών σπουδών στη Στατιστική. Οι μετρήσεις προήλθαν χρησιμοποιώντας υποσύνολο των δηλώσεων της κλίμακας άγχους για τη στατιστική Stars (Lavidas, Manesis & Gialamas, 2021) και η επιλογή των εννέα φοιτητών έγινε χρησιμοποιώντας πίνακα τυχαίων αριθμών. Με βάση τα παραπάνω θα πρέπει να γίνει έλεγχος της μηδενικής υπόθεσης, αν η μέση τιμή του επιπέδου άγχους των φοιτητών στη Στατιστική, από τον πληθυσμό από τον οποίο έχουμε επιλέξει τα δεδομένα, είναι 35.

Stress
39
31
40
26
18
32
27
25
24

Πίνακας 7.1 Πίνακας δεδομένων.

Κατ' αρχάς εισάγουμε τα δεδομένα στο SPSS με τις διαδικασίες που περιγράφηκαν στο Κεφάλαιο 1. Η διατύπωση των ερευνητικών υποθέσεων (αμφίπλευρου ελέγχου) γίνεται ως εξής:

- Μηδενική υπόθεση: Η μέση τιμή του πληθυσμού από τον οποίο προέρχεται το δείγμα είναι 35, δηλαδή $H_0: \mu = 35$.
- Εναλλακτική υπόθεση: Η μέση τιμή του πληθυσμού από τον οποίο προέρχεται το δείγμα διαφέρει από 35, δηλαδή $H_A: \mu \neq 35$.

Επισήμανση: Σύμφωνα με την παραπάνω διατύπωση των ερευνητικών υποθέσεων δεν μας ενδιαφέρει η κατεύθυνση της διαφοράς. Η μ μπορεί να είναι είτε μεγαλύτερη είτε μικρότερη από 35. Εάν όμως μας ενδιαφέρει η κατεύθυνση της διαφοράς, τότε πρέπει να εκτελέσουμε μονόπλευρο έλεγχο. Για παράδειγμα, αν θέλαμε να γίνει έλεγχος της μηδενικής υπόθεσης ότι η μέση τιμή του επιπέδου άγχους των φοιτητών στη στατιστική του πληθυσμού από τον οποίο έχουμε επιλέξει τα δεδομένα είναι μέχρι και 35, τότε η διατύπωση των ερευνητικών υποθέσεων θα γινόταν:

- $H_0: \mu \geq 35$
- $H_A: \mu < 35$

Στη συνέχεια, πρέπει να ελέγξουμε αν πληρούνται οι προϋποθέσεις για να χρησιμοποιηθεί το t-test ενός δείγματος. Οι προϋποθέσεις αυτές είναι:

- Η μεταβλητή της οποίας ελέγχεται η μέση τιμή θα πρέπει να είναι ποσοτική. Η μεταβλητή «stress» είναι ποσοτική και η επιλογή των φοιτητών ήταν τυχαία.
- Αν το δείγμα είναι μικρό ($N < 30$), θα πρέπει να ελεγχθεί η κανονικότητα της μεταβλητής. Στην περίπτωσή μας το δείγμα είναι μικρότερο του 30, επομένως είναι απαραίτητος ο έλεγχος της κανονικότητας της μεταβλητής «stress». Ο έλεγχος αυτός γίνεται με τη βοήθεια των ελέγχων «Kolmogorov-Smirnov» ή/και «Shapiro-Wilk». Αυτά τα δύο τεστ ελέγχουν τη μηδενική υπόθεση ότι η πληθυσμιακή κατανομή της μεταβλητής «stress» είναι κανονική (Γναρδέλλης, 2013). Με λίγα λόγια:

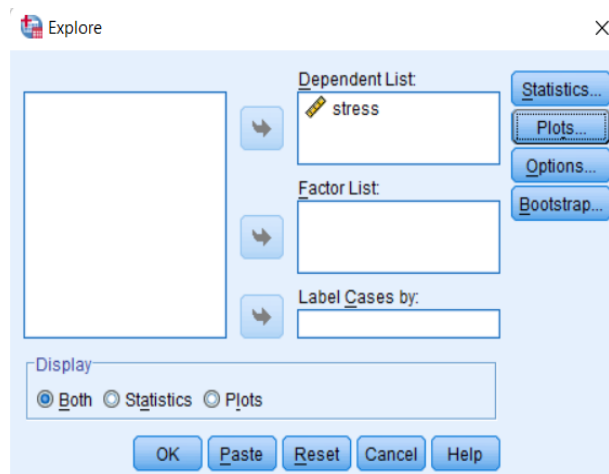
H_0 : Τα δεδομένα του δείγματος προέρχονται από έναν πληθυσμό που ακολουθεί την κανονική κατανομή.

H_A : Τα δεδομένα του δείγματος δεν προέρχονται από έναν πληθυσμό που ακολουθεί την κανονική κατανομή.

- Για μεγάλα δείγματα χρησιμοποιείται το «Kolmogorov-Smirnov» τεστ, ενώ για δείγματα με λιγότερες από 50 περιπτώσεις χρησιμοποιείται συμπληρωματικά και το «Shapiro-Wilk» τεστ.

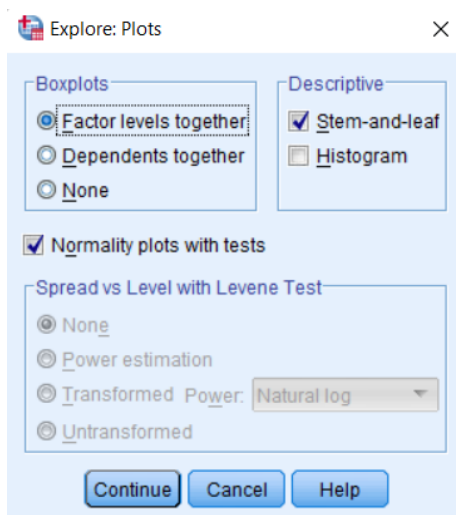
Η διαδικασία εκτέλεσης των παραπάνω ελέγχων γίνεται ως εξής:

- Επιλέγουμε Analyze => Descriptive statistics => Explore και ανοίγει το αντίστοιχο πλαίσιο διαλόγου. Στο πεδίο Dependent List μετακινούμε την ποσοτική μεταβλητή που επιλέγουμε από αριστερά, που στην προκειμένη περίπτωση είναι η μεταβλητή «stress» (Εικόνα 7.2).



Εικόνα 7.2 Πλαίσιο διαλόγου «Explore».

- Στη συνέχεια, κάνουμε κλικ στο κουμπί Plots και στο πλαίσιο διαλόγου που θα εμφανιστεί, τσεκάρουμε την επιλογή Normality plots with tests (Εικόνα 7.3). Τέλος, πατάμε Continue.



Εικόνα 7.3 Πλαίσιο διαλόγου «Explore:Plots».

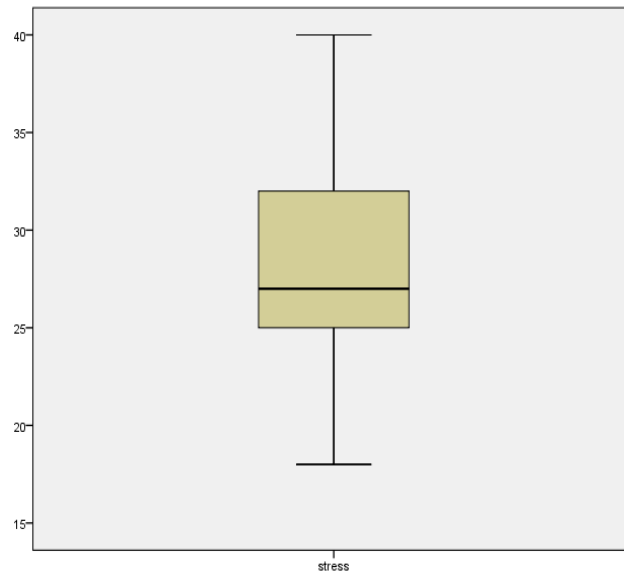
Στο αρχείο αποτελεσμάτων ο 3^{ος} κατά σειρά πίνακας (Tests of Normality) που απεικονίζει τα αποτελέσματα των ελέγχων Kolmogorov-Smirnov και Shapiro-Wilk φαίνεται παρακάτω (Πίνακας 7.1).

Όπως παρατηρούμε, και στους δύο ελέγχους η τιμή Sig. (p value) είναι μεγαλύτερη από το 0.05 ($p = .200$ για τον έλεγχο Kolmogorov-Smirnov και $p = .657$ για τον έλεγχο Shapiro-Wilk). Επομένως, δεν απορρίπτεται η μηδενική υπόθεση. Αυτό σημαίνει ότι η κατανομή της μεταβλητής «stress» είναι κανονική.

Επισήμανση: Τα συγκεκριμένα τεστ είναι αυστηρά, δηλαδή μπορούν να οδηγήσουν ευκολότερα στην απόρριψη της μηδενικής υπόθεσης ότι η πληθυσμιακή κατανομή της μεταβλητής που εξετάζεται είναι κανονική. Αυτό μπορεί να ισχύει και σε περιπτώσεις που η κατανομή είναι κατά προσέγγιση κανονική. Για τον λόγο αυτό καλό είναι να ελέγχουμε και με τη βοήθεια γραφημάτων, όπως για παράδειγμα του θηκογράμματος, την κανονικότητα της κατανομής. Εκτελώντας τις εντολές που αναφέρθηκαν παραπάνω, στο τέλος του αρχείου αποτελεσμάτων απεικονίζεται το θηκόγραμμα της μεταβλητής που εξετάζουμε. Εδώ, όπως φαίνεται και στο θηκόγραμμα της Εικόνας 7.4, η διάμεσος των παρατηρήσεων (μεσαία οριζόντια γραμμή στο ορθογώνιο) δεν φαίνεται να αποκλίνει πολύ από τη διάμεσο του ορθογώνιου, επομένως δεν αναδεικνύεται σοβαρή απόκλιση από την κανονική κατανομή.

Tests of Normality						
	Kolmogorov-Smirnova			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
stress	,172	9	,200*	,947	9	,657
*. This is a lower bound of the true significance.						
a. Lilliefors Significance Correction						

Πίνακας 7.1 Αποτελέσματα ελέγχων «Kolmogorov-Smirnov» και «Shapiro-Wilk».

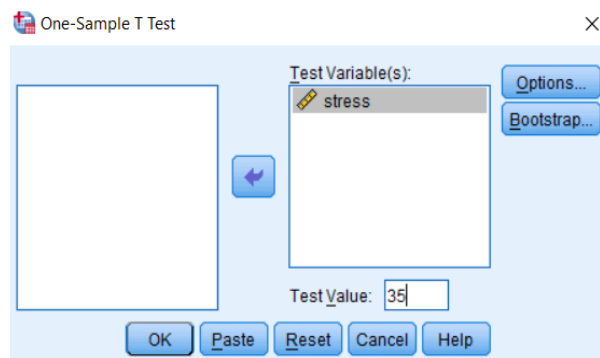


Εικόνα 7.4 Θηκόγραμμα της μεταβλητής «stress».

Επομένως, με βάση τα παραπάνω συμπεραίνουμε ότι ισχύουν οι προϋποθέσεις για την εκτέλεση του t-test ενός δείγματος. Στην προκειμένη περίπτωση, το t-test αφορά τη μηδενική υπόθεση ότι η μέση τιμή του επιπέδου άγχους των φοιτητών στη Στατιστική από τον πληθυσμό, από τον οποίο έχουμε επιλέξει τα δεδομένα, είναι 35.

Τα βήματα για την εκτέλεση του t-test ενός δείγματος είναι τα εξής:

- Επιλέγουμε **Analyze => Compare Means => One-Sample T-Test** και ανοίγει το αντίστοιχο πλαίσιο διαλόγου. Στη λίστα **Test Variable(s)** μετακινούμε τη μεταβλητή που επιλέγουμε από αριστερά, που στην προκειμένη περίπτωση είναι η μεταβλητή «stress».
- Στη συνέχεια, στο πεδίο **Test Value**, εισάγουμε κάθε φορά την τιμή που αναφέρεται στη μηδενική υπόθεση. Δηλαδή, εδώ θα ελέγξουμε τη μεταβλητή «stress» ως προς την τιμή 35 (Εικόνα 7.5).



Εικόνα 7.5 Πλαίσιο διαλόγου «One-Sample T-Test».

Τα αποτελέσματα της ανάλυσης απεικονίζονται στην Εικόνα 7.6.

One-Sample Statistics				
	N	Mean	Std. Deviation	Std. Error Mean
stress	9	29,11	7,149	2,383

Περιγραφικά στατιστικά του δείγματος

One-Sample Test						
	Test Value = 35					
	t	df	Sig. (2-tailed)	Mean Difference	95% Confidence Interval of the Difference	
					Lower	Upper
stress	-2,471	8	,039	-5,889	-11,38	-,39

Η τιμή Sig. (2-tailed), δηλαδή το p value ισούται με 0,039. Επομένως, απορρίπτεται η μηδενική υπόθεση.

Εικόνα 7.6 Πίνακες αποτελεσμάτων του «One-Sample T-Test».

Ο πρώτος Πίνακας, που ονομάζεται *One Sample Statistics*, απεικονίζει το μέγεθος του δείγματος ($N = 9$), καθώς και τα περιγραφικά στατιστικά του επιπέδου άγχους των φοιτητών στη Στατιστική, δηλαδή τη μέση τιμή (Mean = 29,11), την τυπική απόκλιση (Std. Deviation = 7,15) και το τυπικό σφάλμα του μέσου (Std. Error Mean = 2,38).

Ο δεύτερος Πίνακας με τίτλο *One-Sample Test* αναφέρει τα αποτελέσματα του ελέγχου σχετικά με το ότι η μέση τιμή του επιπέδου άγχους των φοιτητών στη Στατιστική από τον πληθυσμό, από τον οποίο έχουμε επιλέξει τα δεδομένα, είναι 35. Η τιμή του στατιστικού του ελέγχου t είναι -2,471 και οι βαθμοί ελευθερίας (df) είναι $8 = (9 - 1)$. Το Sig. (2-tailed), δηλαδή το p value ισούται με 0,039. Όπως αναφέρθηκε, προκειμένου να αποφασίσουμε αν θα απορριφθεί η μηδενική υπόθεση ή όχι σε αμφίπλευρο έλεγχο, συγκρίνουμε το p value με το επίπεδο σημαντικότητας α (συνήθως το 0,05). Εάν $p \leq \alpha$ απορρίπτεται η μηδενική υπόθεση. Άρα εδώ, επειδή $0,039 \leq 0,05$ απορρίπτεται, η μηδενική υπόθεση ότι η πληθυσμιακή τιμή του επιπέδου άγχους των φοιτητών στη Στατιστική είναι ίση με 35. Υπενθυμίζουμε ότι, όταν απορριφθεί η μηδενική υπόθεση σε αμφίπλευρο έλεγχο, απορρίπτεται ταυτόχρονα και η μηδενική υπόθεση του μονόπλευρου ελέγχου (Προσοχή στο ότι το αντίθετο δεν ισχύει πάντα). Επομένως, για να αποφασίσουμε την κατεύθυνση της απόρριψης, δηλαδή εάν $\mu > 35$ ή $\mu < 35$, συγκρίνουμε τη μέση τιμή του δείγματός μας με την τιμή της μηδενικής υπόθεσης, δηλαδή εδώ το 35. Πράγματι, επειδή η μέση τιμή του δείγματός μας είναι 29,11, δηλαδή μικρότερη του 35, αποφασίζουμε ότι η μέση τιμή του επιπέδου άγχους των φοιτητών στη Στατιστική είναι μικρότερη του 35. Οι δύο τελευταίες στήλες του Πίνακα (που φαίνεται στην εικόνα 7.6) δείχνουν το 95% του διαστήματος εμπιστοσύνης όσον αφορά τη διαφορά της μέσης τιμής του πληθυσμού από την τιμή 35.

Το παραπάνω αποτέλεσμα μπορεί να συνταχθεί ως αναφορά στο πλαίσιο μιας εργασίας ως εξής: «Ελέγξαμε την υπόθεση ότι η μέση τιμή του επιπέδου άγχους των φοιτητών στη Στατιστική είναι 35, με τη βοήθεια του ελέγχου t -test ενός δείγματος. Από το αποτέλεσμα του ελέγχου προκύπτει ότι η μέση τιμή του επιπέδου άγχους των φοιτητών στη Στατιστική είναι σημαντικά μικρότερη από 35 [$t(8) = 2,47, p = 0,039$].».

7.1.2 Έλεγχος t ανεξάρτητων δειγμάτων – Independent-Samples T-Test

Σε πειραματικές έρευνες και σε έρευνες παρατήρησης γίνεται χρήση του ελέγχου υποθέσεων όταν έχουμε τις μέσες τιμές δύο ανεξάρτητων δειγμάτων και με βάση τις τιμές αυτές θέλουμε να ελέγξουμε αν οι μέσες τιμές μ_1 και μ_2 των αντίστοιχων πληθυσμών, από τους οποίους προέρχονται τα δείγματα, είναι ίσες ή διαφορετικές

(Γιαλαμάς, 2005). Επομένως, όταν εκτελείται ένας έλεγχος μέσω των τιμών δύο ανεξάρτητων δειγμάτων, χρησιμοποιούνται δύο μεταβλητές: μία κατηγορική μεταβλητή που περιέχει τις δύο ομάδες ανεξάρτητων δειγμάτων και μία ποσοτική μεταβλητή (Γναρδέλλης, 2013).

Οι υποθέσεις οι οποίες μπορούν να διατυπωθούν είναι:

- $H_0: \mu_1 - \mu_2 = 0, \quad H_A: \mu_1 - \mu_2 \neq 0$
- $H_0: \mu_1 - \mu_2 \geq 0, \quad H_A: \mu_1 - \mu_2 < 0$
- $H_0: \mu_1 - \mu_2 \leq 0, \quad H_A: \mu_1 - \mu_2 > 0$

Επιπλέον, μπορούν να γίνουν οι ίδιοι έλεγχοι με άλλες τιμές αντί του 0. Για παράδειγμα, η ερευνητική υπόθεση που θέτει ένας ερευνητής είναι η μ_1 να είναι μεγαλύτερη από τη μ_2 κατά 3 ποσοστιαίες μονάδες. Οι στατιστικές υποθέσεις διατυπώνονται ως εξής:

- $H_0: \mu_1 - \mu_2 \leq 3\%, \quad H_A: \mu_1 - \mu_2 > 3\%$

Υπάρχουν τρεις περιπτώσεις συνθηκών (Γιαλαμάς, 2005; Γναρδέλλης, 2013):

- Τα δείγματα προέρχονται από κανονικά κατανεμημένους πληθυσμούς με γνωστές τις διακυμάνσεις των πληθυσμών σ_1^2 και σ_2^2 . Το στατιστικό για τον έλεγχο της μηδενικής υπόθεσης το οποίο ακολουθεί την τυπική κανονική κατανομή είναι:

$$Z = \frac{(\bar{x}_1 - \bar{x}_2) - (\mu_1 - \mu_2)}{\sqrt{\frac{\sigma_1^2}{N_1} + \frac{\sigma_2^2}{N_2}}} \quad (7.2)$$

Όπου,

$\bar{x}_1, \bar{x}_2$ και N_1, N_2 οι μέσες τιμές και τα μεγέθη των δύο δειγμάτων

- Τα δείγματα προέρχονται από κανονικά κατανεμημένους πληθυσμούς με άγνωστες τις διακυμάνσεις των πληθυσμών. Επειδή πολύ σπάνια είναι γνωστές οι διακυμάνσεις των πληθυσμών από τους οποίους προέρχονται τα δείγματα, η εκτίμησή τους γίνεται από τα διαθέσιμα δεδομένα του δείγματος της έρευνας. Μπορούν να προκύψουν δύο περιπτώσεις: οι διακυμάνσεις των δύο πληθυσμών να είναι ίσες ή άνισες (Γιαλαμάς, 2005). Όταν οι διακυμάνσεις των δύο πληθυσμών είναι ίσες μεταξύ τους, τότε χρησιμοποιούμε τον εξής τύπο για την εκτίμηση της κοινής διακύμανσης, s_p^2 :

$$s_p^2 = \frac{s_1^2(N_1 - 1) + s_2^2(N_2 - 1)}{N_1 + N_2 - 2} \quad (7.3)$$

Το στατιστικό για τον έλεγχο της μηδενικής υπόθεσης είναι:

$$t = \frac{(\bar{x}_1 - \bar{x}_2) - (\mu_1 - \mu_2)}{\sqrt{\frac{s_p^2}{N_1} + \frac{s_p^2}{N_2}}} \quad (7.4)$$

Όπου $\bar{x}_1, \bar{x}_2$ και N_1, N_2 οι μέσες τιμές και τα μεγέθη των δύο δειγμάτων.

Το στατιστικό αυτό ακολουθεί την κατανομή t με $N_1 + N_2 - 2$ βαθμούς ελευθερίας.

- Όταν έχουμε επιλέξει δύο ανεξάρτητα και τυχαία δείγματα από δύο κανονικούς πληθυσμούς των οποίων οι διακυμάνσεις διαφέρουν πολύ, τότε το στατιστικό για τον έλεγχο της μηδενικής υπόθεσης είναι:

$$t' = \frac{(\bar{x}_1 - \bar{x}_2) - (\mu_1 - \mu_2)}{\sqrt{\frac{s_1^2}{N_1} + \frac{s_2^2}{N_2}}} \quad (7.5)$$

Δηλαδή οι διακυμάνσεις των δύο πληθυσμών έχουν αντικατασταθεί από τις γνωστές εκτιμώμενες διακυμάνσεις S_1^2, S_2^2 των δύο δειγμάτων.

Επιπλέον, οι βαθμοί ελευθερίας df προσαρμόζονται (Γναρδέλλης, 2013) στην τιμή:

$$df' = \frac{\left(\frac{s_1^2}{N_1} + \frac{s_2^2}{N_2}\right)^2}{\frac{1}{N_1 - 1} \left(\frac{s_1^2}{N_1}\right)^2 + \frac{1}{N_2 - 1} \left(\frac{s_2^2}{N_2}\right)^2} \quad (7.6)$$

Για τις δύο παραπάνω περιπτώσεις ελέγχου με άγνωστες τις διακυμάνσεις των δύο πληθυσμών, στον αμφίπλευρο έλεγχο, θα απορρίψουμε τη μηδενική υπόθεση, αν η υπολογισμένη τιμή t ή t' είναι είτε μεγαλύτερη είτε ίση από την κρίσιμη τιμή $t_{\kappa\rho, (1-\alpha/2)}$ ή $t'_{\kappa\rho, (1-\alpha/2)}$ αντίστοιχα ή είναι είτε μικρότερη είτε ίση από την αρνητική αυτής της τιμής. Σε μονόπλευρο έλεγχο θα υπολογιστούν με ανάλογο τρόπο η κρίσιμες τιμές $t_{\kappa\rho, (1-\alpha)}$ ή $t'_{\kappa\rho, (1-\alpha)}$, για df ή df' βαθμούς ελευθερίας αντίστοιχα.

- Τα δείγματα προέρχονται από πληθυσμούς που δεν κατανέμονται κανονικά. Αν το μέγεθος των δειγμάτων είναι μεγάλο (N τουλάχιστον 30), δηλαδή επαρκώς μεγάλα, τότε σύμφωνα με το «Κεντρικό Οριακό Θεώρημα» η δειγματοληπτική κατανομή της διαφοράς των μέσων τιμών θα κατανέμεται κανονικά (Γιαλαμάς, 2005).

Το στατιστικό που χρησιμοποιείται είναι το εξής:

$$Z = \frac{(\bar{x}_1 - \bar{x}_2) - (\mu_1 - \mu_2)}{\sqrt{\frac{\sigma_1^2}{N_1} + \frac{\sigma_2^2}{N_2}}} \quad (7.7)$$

Το στατιστικό ακολουθεί την τυπική κανονική κατανομή, αν αληθεύει η μηδενική υπόθεση. Αν οι διακυμάνσεις των δύο πληθυσμών είναι γνωστές, τότε μπορούν να χρησιμοποιηθούν. Αν είναι άγνωστες, τότε μπορούν να αντικατασταθούν από τις εκτιμήσεις τους, δηλαδή από κάθε μία δειγματική διακύμανση.

Το στατιστικό του ελέγχου είναι:

$$Z = \frac{(\bar{x}_1 - \bar{x}_2) - (\mu_1 - \mu_2)}{\sqrt{\frac{s_1^2}{N_1} + \frac{s_2^2}{N_2}}} \quad (7.8)$$

Αν βρεθεί παραβίαση της κανονικότητας των πληθυσμών από τους οποίους προέρχονται τα δείγματα και το μέγεθος των δειγμάτων είναι μικρό, τότε δεν μπορεί να χρησιμοποιηθεί ούτε η κατανομή z ούτε η κατανομή t . Στην περίπτωση αυτή θα εκτελεστούν *μη παραμετρικοί έλεγχοι* (Nonparametric Tests) (Γιαλαμάς, 2005; Γναρδέλλης, 2013).

Για την εκτέλεση του *t-test* ανεξάρτητων δειγμάτων πρέπει να πληρούνται ορισμένες προϋποθέσεις (Γιαλαμάς, 2005; Γναρδέλλης, 2013):

- Οι πληθυσμοί από τους οποίους προέρχονται τα δείγματα των ποσοτικών τιμών θα πρέπει να κατανέμονται κανονικά. Το *t test* είναι ανθεκτικό σε μέτριες αποκλίσεις από την κανονικότητα,

ιδιαίτερα στα μεγάλα δείγματα. Ουσιαστικά, απαιτείται μια κατά προσέγγιση συμμετρία των κατανομών των δύο δειγμάτων και η απουσία πολύ ακραίων τιμών. Ο έλεγχος της ύπαρξης ακραίων τιμών για κάθε δείγμα γίνεται κυρίως γραφικά με τη βοήθεια του θηκογράμματος (boxplot). Όταν τα δείγματα είναι μικρά ($N < 30$) και υπάρχει έντονη ασυμμετρία, συνιστάται η χρήση του μη παραμετρικού ελέγχου *Mann-Whitney U*, ο οποίος αποτελεί τη μη παραμετρική εκδοχή του *t-test* ανεξάρτητων δειγμάτων.

- Οι διακυμάνσεις των δύο πληθυσμών θα πρέπει να είναι ίσες. Η παραβίαση της προϋπόθεσης της ομοιογένειας της διακύμανσης είναι αμελητέα, όταν τα μεγέθη των δύο δειγμάτων είναι ίσα. Στην περίπτωση που δεν υπάρχει ισότητα των διακυμάνσεων, δημιουργείται πρόβλημα, όταν $S^2_{\text{μέγιστη}} > 10S^2_{\text{ελάχιστη}}$ ή, σε μικρότερες διαφορές, όταν τα μεγέθη των δειγμάτων διαφέρουν σημαντικά. Ο έλεγχος της υπόθεσης για την ισότητα των διακυμάνσεων των δύο πληθυσμών είναι ο εξής:

$$H_0: \sigma_1^2 = \sigma_2^2$$

$$H_A: \sigma_1^2 \neq \sigma_2^2$$

- Η ισότητα των διακυμάνσεων ελέγχεται με το τεστ του *Levene*. Το τεστ αυτό είναι περισσότερο απαλλαγμένο από την επίδραση της έλλειψης κανονικότητας των δεδομένων. Το αποτέλεσμα του συγκεκριμένου τεστ θα κρίνει τον τρόπο με τον οποίο θα γίνει ο έλεγχος *t* των δύο μέσων τιμών. Αν το «παρατηρούμενο επίπεδο σημαντικότητας p » βρεθεί μεγαλύτερο από το 0,05 (αντίστοιχη ένδειξη στο SPSS, Sig. > 0,05), δηλαδή δεν απορριφθεί η μηδενική υπόθεση, τότε ο έλεγχος των μέσων τιμών γίνεται θεωρώντας ότι οι διακυμάνσεις των δύο πληθυσμών είναι ίσες. Το στατιστικό που θα χρησιμοποιηθεί είναι το 7.4. Αντίθετα, αν το p βρεθεί μικρότερο ή ίσο από το 0,05 (Sig. $\leq$ 0,05), δηλαδή απορριφθεί η μηδενική υπόθεση, τότε ο έλεγχος των μέσων τιμών γίνεται θεωρώντας ότι οι διακυμάνσεις των δύο πληθυσμών είναι άνισες. Το στατιστικό που θα χρησιμοποιηθεί είναι το 7.5.
- Η επιλογή μιας παρατήρησης μέσα σε μια ομάδα θα πρέπει να είναι ανεξάρτητη από την επιλογή των υπόλοιπων παρατηρήσεων.
- Τα δείγματα θα πρέπει να έχουν επιλεχτεί με τυχαία δειγματοληψία και να είναι ανεξάρτητα μεταξύ τους. Ανεξαρτησία δειγμάτων έχουμε όταν η επιλογή ενός δείγματος δεν επηρεάζεται από την επιλογή του άλλου. Ουσιαστικά, τα δείγματα είναι ανεξάρτητα μεταξύ τους όταν μετράμε το ίδιο χαρακτηριστικό σε διαφορετικές πειραματικές ομάδες. Δεν υπάρχει κάποιος κοινός παράγοντας ανάμεσα στις δύο ομάδες που να μπορεί να επηρεάσει τη μέτρηση. Για παράδειγμα, όταν θέλουμε να ελέγξουμε αν υπάρχει διαφορά στις μέσες τιμές ανάμεσα στα δύο φύλα όσον αφορά τη συχνότητα χρήσης των μέσων κοινωνικής δικτύωσης, θα χρησιμοποιηθεί ο έλεγχος ανεξάρτητων δειγμάτων. Δεν υπάρχει κάποιος κοινός παράγοντας που επηρεάζει τα δύο φύλα σχετικά με το πόσο συχνά χρησιμοποιούν τα μέσα κοινωνικής δικτύωσης.

Η εκτέλεση του ελέγχου *t* ανεξάρτητων δειγμάτων στο SPSS θα γίνει με το ακόλουθο παράδειγμα. Θα χρησιμοποιηθεί το αρχείο «interest_statistics.sav». Το συγκεκριμένο αρχείο περιλαμβάνει δύο μεταβλητές: το φύλο των φοιτητών (gender) και το ενδιαφέρον τους (interest) για το μάθημα της Στατιστικής. Πρόκειται για φοιτητές ανθρωπιστικών σπουδών. Η μέτρηση του ενδιαφέροντος για το μάθημα της Στατιστικής πραγματοποιήθηκε με την υποκλίμακα *interest* του SATS 36. Τα δείγματα είναι ανεξάρτητα μεταξύ τους, γιατί μετράμε το ίδιο χαρακτηριστικό σε δύο διαφορετικές πειραματικές ομάδες: φοιτητές ή άνδρες και φοιτήτριες ή γυναίκες.

Η διατύπωση των ερευνητικών υποθέσεων γίνεται ως εξής:

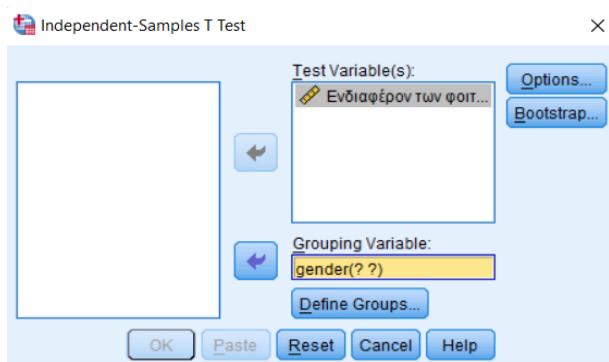
- Μηδενική υπόθεση: Η μέση τιμή του πληθυσμού των φοιτητών όσον αφορά το ενδιαφέρον τους για το μάθημα της Στατιστικής ισούται με τη μέση τιμή του πληθυσμού των φοιτητριών, δηλαδή $H_0: \mu_A = \mu_B$.
- Εναλλακτική υπόθεση: Η μέση τιμή του πληθυσμού των φοιτητών όσον αφορά το ενδιαφέρον τους για το μάθημα της Στατιστικής διαφέρει από τη μέση τιμή του πληθυσμού των φοιτητριών, δηλαδή $H_A: \mu_A \neq \mu_B$.

Επισημάνση: Σύμφωνα με την παραπάνω διατύπωση των ερευνητικών υποθέσεων, πρόκειται για έναν αμφίπλευρο έλεγχο διότι δεν μας ενδιαφέρει η κατεύθυνση της διαφοράς.

Επειδή το δείγμα είναι μικρό ($N = 23$), θα πρέπει να ελεγχθεί η κανονικότητα της μεταβλητής «interest». Εκτελούμε λοιπόν τους ελέγχους «Kolmogorov-Smirnov» και «Shapiro-Wilk», όπως ακριβώς περιγράφεται στην προηγούμενη ενότητα. Τα αποτελέσματα των ελέγχων αυτών, καθώς και τα αντίστοιχα θηκογράμματα, δεν έδειξαν σοβαρή απόκλιση από την κανονικότητα.

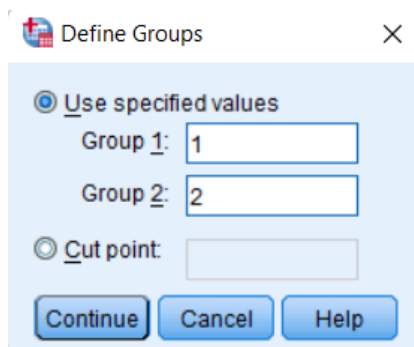
Τα βήματα για την εκτέλεση του t-test ανεξάρτητων δειγμάτων είναι τα εξής:

- Επιλέγουμε **Analyze => Compare Means => Independent-Samples T-Test** και ανοίγει το αντίστοιχο πλαίσιο διαλόγου (Εικόνα 7.7). Στη λίστα **Test Variable(s)** μετακινούμε την ποσοτική μεταβλητή που επιλέγουμε από αριστερά, που στην προκειμένη περίπτωση είναι η μεταβλητή «interest».
- Στη συνέχεια, στο πεδίο **Grouping Variable** εισάγουμε την κατηγορική μεταβλητή «gender».



Εικόνα 7.7 Πλαίσιο διαλόγου «Independent-Samples T-Test».

- Έπειτα, πατώντας το πλήκτρο **Define Groups**, ορίζουμε τους αριθμούς που αντιστοιχούν στις δύο κατηγορίες της μεταβλητής (1 = Άνδρας, 2 = Γυναίκα) (Εικόνα 7.8). Προσοχή στο ότι αν δεν θυμόμαστε την κωδικοποίηση της κατηγορικής μεταβλητής, μπορούμε να κάνουμε δεξί κλικ πάνω στη μεταβλητή «gender» και να πατήσουμε την επιλογή **Variable Information**.



Εικόνα 7.8 Πλαίσιο διαλόγου «Define Groups».

Τα αποτελέσματα της ανάλυσης που εμφανίζονται στο αρχείο αποτελεσμάτων του SPSS απεικονίζουν δύο πίνακες. Ο πρώτος πίνακας με το όνομα *Group Statistics* (Εικόνα 7.9) απεικονίζει τα περιγραφικά στατιστικά για κάθε ομάδα (μέγεθος δείγματος, μέση τιμή, τυπική απόκλιση και τυπικό σφάλμα).

Group Statistics					
	Φύλο	N	Mean	Std. Deviation	Std. Error Mean
Ενδιαφέρον των φοιτητών για το μάθημα της στατιστικής	Άνδρας	11	20,36	2,656	,801
	Γυναίκα	12	18,17	2,855	,824

Εικόνα 7.9 Πίνακας αποτελεσμάτων Group Statistics.

Ο δεύτερος πίνακας με τίτλο *Independent Samples Test* (Εικόνα 7.10) απεικονίζει τα στατιστικά στοιχεία σχετικά με τον έλεγχο. Ο Πίνακας αυτός αποτελείται από δύο μέρη. Το πρώτο μέρος αριστερά με τίτλο *Levene's Test for Equality of Variances* δίνει τα αποτελέσματα του ελέγχου της υπόθεσης της ισότητας των διακυμάνσεων των δύο πληθυσμών. Το στατιστικό του ελέγχου ακολουθεί την κατανομή *F*. Εδώ, το $F = 0,220$ και το p value (Sig.) = 0,644. Αυτό σημαίνει ότι δεν απορρίπτεται η μηδενική υπόθεση και επομένως οι διακυμάνσεις των ανδρών και των γυναικών δεν φαίνεται να είναι διαφορετικές.

Το δεύτερο μέρος του Πίνακα (βασικό τμήμα του ελέγχου) με τίτλο *t-test for Equality of Means* (Εικόνα 7.10), αποτελείται από δύο γραμμές. Η πρώτη γραμμή αντιστοιχεί στο *Equal variances assumed*, ενώ η δεύτερη στο *Equal variances not assumed*. Ο τρόπος ανάγνωσης του συγκεκριμένου μέρους του πίνακα είναι ο εξής: αν η τιμή του Sig. της πρώτης γραμμής *Equal variances assumed* είναι $> 0,05$, τότε συνεχίζουμε στην ίδια γραμμή και διαβάζουμε το επόμενο κομμάτι *t-test for Equality of Means* του Πίνακα. Αν η συγκεκριμένη τιμή είναι $\leq 0,05$, τότε διαβάζουμε τη δεύτερη γραμμή *Equal variances not assumed* (αν δεν ισχύει η ισότητα των διακυμάνσεων, τότε ακολουθούμε μια διόρθωση των βαθμών ελευθερίας) του επόμενου τμήματος *t-test for Equality of Means* του Πίνακα.

Εδώ, επειδή η τιμή του Sig. της πρώτης γραμμής *Equal variances assumed* είναι 0,644, χρησιμοποιούμε το *t-test* της πρώτης γραμμής. Το στατιστικό του ελέγχου *t* ισούται με 1,906 και η πιθανότητα να απορριφθεί η υπόθεση της ισότητας των μέσων τιμών μεταξύ ανδρών και γυναικών είναι Sig. (2-tailed) = 0,070 $> \alpha = 0,05$. Επομένως, δεν απορρίπτεται η μηδενική υπόθεση.

Το παραπάνω αποτέλεσμα μπορεί να συνταχθεί ως αναφορά στο πλαίσιο μιας εργασίας ως εξής: Ελέγξαμε την υπόθεση της διαφοράς των μέσων τιμών μεταξύ ανδρών και γυναικών όσον αφορά το ενδιαφέρον τους για το μάθημα της Στατιστικής με τη χρήση ενός *t-test* ανεξάρτητων δειγμάτων. Το μέσο ενδιαφέρον των ανδρών για το μάθημα της Στατιστικής (M.T. = 20,36, T.A. = 2,656) είναι μεγαλύτερο από το αντίστοιχο μέσο ενδιαφέρον των γυναικών (M.T. = 18,17, T.A. = 2,855), ωστόσο η διαφορά αυτή δεν είναι στατιστικά σημαντική ($t(21) = 1,906$, $p = 0,070$).

Στην περίπτωση που μας ενδιαφέρει η κατεύθυνση της διαφοράς, τότε πρέπει να εκτελέσουμε μονόπλευρο έλεγχο.

- Μηδενική υπόθεση: Η μέση τιμή του πληθυσμού των φοιτητών, όσον αφορά το ενδιαφέρον τους για το μάθημα της Στατιστικής, είναι μικρότερη ή ίση από τη μέση τιμή του πληθυσμού των φοιτητριών, δηλαδή $H_0: \mu_A \leq \mu_B$.
- Εναλλακτική υπόθεση: Η μέση τιμή του πληθυσμού των φοιτητών, όσον αφορά το ενδιαφέρον τους για το μάθημα της Στατιστικής, είναι υψηλότερη από τη μέση τιμή του πληθυσμού των φοιτητριών, δηλαδή $H_A: \mu_A > \mu_B$.

Αφού εκτελέσουμε κανονικά τον αμφίπλευρο έλεγχο βρίσκουμε την τιμή του p , διαιρώντας την αρχική με το 2. Στο παράδειγμά μας, η τιμή p του αμφίπλευρου ελέγχου ισούται με 0,070. Η τιμή του στατιστικού t είναι θετική (1,906) και μας ενδιαφέρει ο πρώτος πληθυσμός να έχει μέση τιμή μεγαλύτερη από τον δεύτερο πληθυσμό.

Σε αυτή την περίπτωση, η τιμή p του μονόπλευρου ελέγχου είναι το μισό της τιμής p του αμφίπλευρου ελέγχου: $p = 0,070 / 2 = 0,035 < 0,05$. Επομένως, αντίθετα με τον αμφίπλευρο έλεγχο, στον μονόπλευρο έλεγχο απορρίπτεται η μηδενική υπόθεση. Προσοχή: πολλές φορές μπορεί να απορρίπτεται ο μονόπλευρος έλεγχος και να μην απορρίπτεται ο αντίστοιχος αμφίπλευρος. Επομένως, η διατύπωση μονόπλευρου ελέγχου πρέπει να γίνεται με αυστηρότητα μόνο στην αρχή της έρευνας και πάντα με βάση την προηγούμενη έρευνα.

Το παραπάνω αποτέλεσμα μπορεί να συνταχθεί ως αναφορά, στο πλαίσιο μιας επιστημονικής εργασίας, ως εξής: Ελέγξαμε την υπόθεση της διαφοράς των μέσων τιμών μεταξύ των φοιτητών και των φοιτητριών σχετικά με το ενδιαφέρον τους για το μάθημα της Στατιστικής με τη χρήση ενός t-test ανεξάρτητων δειγμάτων.

Ο έλεγχος του Levene δεν ανέδειξε παραβίαση της προϋπόθεσης των διακυμάνσεων των δύο πληθυσμών. Ο μονόπλευρος έλεγχος t ανεξάρτητων δειγμάτων αποκάλυψε ότι το μέσο ενδιαφέρον των φοιτητών (M.T.= 20,36, T.A.=2,656) για το μάθημα της Στατιστικής είναι σημαντικά μεγαλύτερο από το αντίστοιχο μέσο ενδιαφέρον των φοιτητριών (M.T.= 18,17, T.A.=2,855) ($t(21) = 1,906$, $p = 0,035$).

Independent Samples Test										
		Levene's Test for Equality of Variances		t-test for Equality of Means						
		F	Sig.	t	df	Sig. (2-tailed)	Mean Difference	Std. Error Difference	95% Confidence Interval of the Difference	
									Lower	Upper
Ενδιαφέρον των φοιτητών για το μάθημα της στατιστικής	Equal variances assumed	,220	,644	1,906	21	,070	2,197	1,153	-,201	4,595
	Equal variances not assumed			1,912	21,0	,070	2,197	1,149	-,193	4,587

Η τιμή Sig. (2-tailed), δηλαδή το p value όταν οι διακυμάνσεις των δύο πληθυσμών είναι ίσες (0,070). Επομένως, δεν απορρίπτεται η μηδενική υπόθεση.

Εικόνα 7.10 Πίνακας αποτελεσμάτων ελέγχου t ανεξάρτητων δειγμάτων.

Τέλος, αξίζει να τονίσουμε ότι το αποτέλεσμα οποιουδήποτε ελέγχου είναι εξαρτώμενο της κατανομής των τιμών του δείγματος. Αυτό σημαίνει ότι είναι αδύνατον με το συγκεκριμένο δείγμα να δείξουμε ότι το ενδιαφέρον των φοιτητριών για τη Στατιστική είναι στατιστικά σημαντικά υψηλότερο του ενδιαφέροντος των φοιτητών, αφού στα δειγματικά δεδομένα παρατηρείται το αντίθετο. Για να αποτυπωθεί αυτό και στον έλεγχο του t-test χρησιμοποιώντας το p -value, χρειάζεται να κάνουμε το εξής: αφού θέλουμε ο πρώτος πληθυσμός (φοιτητές) να έχει μικρότερη μέση τιμή από τον δεύτερο πληθυσμό (φοιτήτριες), δηλαδή αρνητική τιμή της διαφοράς των μέσων τιμών (MTφοιτητές - MTφοιτήτριες), και το αντίστοιχο t είναι θετικό, δηλαδή ετερόσημο της διαφοράς των μέσων τιμών, τότε το καινούριο p -value υπολογίζεται ως εξής: $1 - \llbracket p\text{-value αμφίπλευρου ελέγχου} \rrbracket / 2 = 1 - 0,070 / 2 = 1 - 0,035 = 0,965 > 0,05$.

7.1.3 Έλεγχος t δειγμάτων κατά ζεύγη – Paired-Samples T-Test

Ένα άλλο σχέδιο ελέγχου που χρησιμοποιείται συχνά για την αξιολόγηση μιας πειραματικής διαδικασίας είναι αυτό των σχετιζόμενων κατά ζεύγη παρατηρήσεων (Γιαλαμάς, 2005). Στον σχεδιασμό δύο εξαρτημένων δειγμάτων, όπως αλλιώς λέγεται, οι δύο μετρήσεις (εξαρτημένα δείγματα) γίνονται για τα ίδια άτομα. Χαρακτηριστικό παράδειγμα αποτελεί η μέτρηση της επίδοσης των ίδιων υποκειμένων σε δύο διαφορετικές χρονικές στιγμές (pre & post test). Αν, για παράδειγμα, θέλουμε να συγκρίνουμε την αποτελεσματικότητα δύο διδακτικών μεθόδων εκμάθησης Στατιστικής και επιλέξουμε δύο δείγματα φοιτητών, όπου στο καθένα από αυτά θα εφαρμοστεί μία μέθοδος, τότε θα εκτελέσουμε έναν έλεγχο t ανεξάρτητων δειγμάτων. Στην περίπτωση όμως, που επιλέξουμε ένα δείγμα φοιτητών, όπου ο καθένας από αυτούς θα διδαχτεί Στατιστική με τον ίδιο τρόπο και μετράμε την επίδοσή τους με ένα τεστ πριν και μετά, αν θέλουμε να ελέγξουμε τη διαφορά των μέσων τιμών των δύο επιδόσεων, θα εκτελέσουμε έλεγχο t δειγμάτων κατά ζεύγη.

Επίσης, ο έλεγχος t δειγμάτων κατά ζεύγη εκτελείται όταν υπάρχει μια κοινή ιδιότητα ανάμεσα στα υποκείμενα της έρευνας, η οποία μπορεί να επηρεάσει τη μέτρηση των μέσων τιμών. Για παράδειγμα, εάν θέλουμε να διερευνήσουμε τις διαφορές στις δαπάνες των αγοριών και των κοριτσιών, οικογενειών με δίδυμα (αγόρι και κορίτσι) που βρίσκονται στην εφηβική ηλικία, θα εκτελέσουμε έναν έλεγχο t δειγμάτων κατά ζεύγη, γιατί τα δύο αδέλφια ανήκουν στην ίδια οικογένεια και επομένως έχουν την ίδια οικονομική δυνατότητα που επηρεάζει τις δαπάνες τους. Για τη δημιουργία του πίνακα δεδομένων που θα χρησιμοποιηθεί στην περίπτωση

αυτή καταγράφουμε για κάθε οικογένεια από ένα τυχαίο δείγμα οικογενειών τη δαπάνη του κοριτσιού και τη δαπάνη του αγοριού (Γιαλαμάς, 2005).

Ο σχεδιασμός εξαρτημένων δειγμάτων και επομένως ο έλεγχος t δειγμάτων κατά ζεύγη έχει δύο πλεονεκτήματα (Γιαλαμάς, 2005):

- Είναι πιο ισχυρός έλεγχος σε σχέση με τον έλεγχο t ανεξάρτητων δειγμάτων. Αυτό σημαίνει ότι η πιθανότητα να απορριφθεί η μηδενική υπόθεση είναι μεγαλύτερη στις συγκρίσεις των μέσων τιμών κατά ζεύγη. Αυτό το χαρακτηριστικό είναι προς όφελος του ερευνητή στην πλειονότητα των πειραματικών διαδικασιών. Το πλεονέκτημα αυτό προκύπτει επειδή η τιμή του τυπικού σφάλματος είναι μικρότερη στον έλεγχο κατά ζεύγη και κατά συνέπεια καταλήγουμε σε τιμή t μεγαλύτερη από την αντίστοιχη του ερευνητικού σχεδίου ανεξάρτητων δειγμάτων.
- Η εκτέλεση του συγκεκριμένου ελέγχου χρησιμοποιείται όταν χρειάζεται να εξαλειφθούν κάποια εξωγενή χαρακτηριστικά των συμμετεχόντων, τα οποία μπορούν να επηρεάσουν τα αποτελέσματα μιας έρευνας. Επομένως, ελαχιστοποιείται η πιθανότητα να απορριφθεί λανθασμένα η μηδενική υπόθεση. Για παράδειγμα, ένας δάσκαλος διεξάγει ένα πείραμα, με το οποίο συγκρίνει δύο μεθόδους απομνημόνευσης αριθμών. Δύο διαφορετικές ομάδες μαθητών, μία για κάθε διαφορετική μέθοδο, εξασκούνται για ένα χρονικό διάστημα. Ας υποθέσουμε ότι στο τέλος συγκρίνεται η ικανότητα απομνημόνευσης αριθμών των δύο διαφορετικών ομάδων και απορρίπτεται η μηδενική υπόθεση, δηλαδή η διαφορά βρέθηκε σημαντική. Όμως, υπάρχει πιθανότητα η διαφορά αυτή να μην οφείλεται στη μεγαλύτερη αποτελεσματικότητα της μιας μεθόδου σε σχέση με την άλλη, αλλά στο ότι οι δύο ομάδες μαθητών είχαν διαφορετικές ικανότητες απομνημόνευσης πριν αρχίσουν την πειραματική διαδικασία. Στην αντίθετη περίπτωση, όταν μετρήσουμε δύο φορές το ίδιο υποκείμενο της έρευνας, η πιθανή διαφορά των δύο μετρήσεων θα οφείλεται αποκλειστικά και μόνο στην αποτελεσματικότητα των δύο μεθόδων.

Στον έλεγχο t κατά ζεύγη χρησιμοποιούνται οι διαφορές d των μέσων τιμών κάθε ζεύγους. Οι διαφορές αυτές ορίζουν μία καινούρια μεταβλητή, μέσω της οποίας θα πραγματοποιηθεί ο έλεγχος. Δεχόμαστε ότι οι δειγματικές διαφορές αποτελούν ένα τυχαίο δείγμα από έναν κανονικά κατανομημένο πληθυσμό διαφορών, με μέση τιμή μ_d και τυπική απόκλιση σ_d (Γιαλαμάς, 2005).

Ο έλεγχος της μηδενικής υπόθεσης της ισότητας των μέσων τιμών των δύο συνθηκών (π.χ. πριν και μετά τη θεραπεία) είναι ο εξής:

$$H_0: \mu_1 = \mu_2$$

$$H_A: \mu_1 \neq \mu_2$$

Επειδή όμως για τη μέση τιμή μ_d των διαφορών που μπορεί να υπολογιστούν για κάθε ζεύγος τιμών ισχύει $\mu_d = \mu_1 - \mu_2$, οι υποθέσεις μπορεί ισοδύναμα να διατυπωθούν ως εξής:

$$H_0: \mu_d = 0$$

$$H_A: \mu_d \neq 0$$

Επομένως, ο έλεγχος ανάγεται σε έναν έλεγχο t ενός δείγματος για τη μέση τιμή των διαφορών του πληθυσμού. Το στατιστικό του ελέγχου της μηδενικής υπόθεσης είναι:

$$t = \frac{\bar{d} - \mu_d}{\frac{s_d}{\sqrt{N}}} \quad (7.9)$$

όπου: $\bar{d}$ είναι η μέση τιμή του δείγματος των διαφορών,

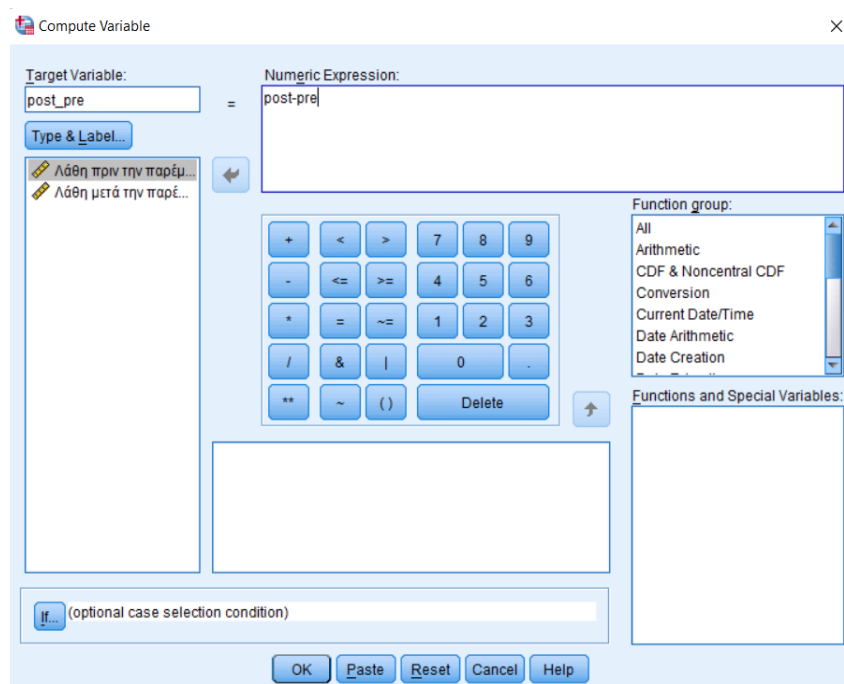
s_d είναι η τυπική απόκλιση του δείγματος των διαφορών,

N το μέγεθος του δείγματος των διαφορών.

Η εκτέλεση του ελέγχου t δειγμάτων κατά ζεύγη στο SPSS θα γίνει με το ακόλουθο παράδειγμα. Θα χρησιμοποιηθεί το αρχείο «paradeigma_2d_recognition_(pre-post)_2.sav». Στο συγκεκριμένο αρχείο διερευνάται η επίδραση μιας διδακτικής παρέμβασης στην τάξη του νηπιαγωγείου, όσον αφορά την αναγνώριση γεωμετρικών σχημάτων. Επομένως, ο έλεγχος που θα πραγματοποιηθεί είναι μονόπλευρος. Αξίζει να τονίσουμε ότι στις περιπτώσεις που μας ενδιαφέρει να ελέγχουμε κατά πόσο μια διδακτική παρέμβαση επέφερε γνωστικό όφελος στους μαθητές, η υπόθεση που πρέπει να ελεγχθεί θα πρέπει να είναι μονόπλευρου ελέγχου. Ας υποθέσουμε, λοιπόν, ότι μετρούνται για τα ίδια νήπια πριν την παρέμβαση (pre-test) και μετά από μία εβδομάδα (post-test) τα λάθη στην αναγνώριση γεωμετρικών σχημάτων δύο διαστάσεων. Έχουμε δύο μετρήσεις (μεταβλητές) δηλαδή την «pre» και την «post».

Κατ' αρχάς, δημιουργούμε μία νέα μεταβλητή με το όνομα «post_pre», η οποία προκύπτει από τη διαφορά των δύο μεταβλητών, (post-pre). Αυτό γίνεται μέσω των εντολών **Transform => Compute Variable** (βλ. Κεφ. 2). Στο πεδίο **Target Variable** πληκτρολογούμε το όνομα της καινούριας μεταβλητής «post_pre» και στη λίστα **Numeric Expression** πληκτρολογούμε **post-pre** (Εικόνα 7.11).

Στη συνέχεια, λόγω του ότι το δείγμα είναι πολύ μικρό ($N = 9$), πρέπει να ελεγχθεί η κανονικότητα της μεταβλητής «post_pre». Εκτελούμε λοιπόν τους ελέγχους «Kolmogorov-Smirnov» και «Shapiro-Wilk», όπως ακριβώς περιγράφεται στην Ενότητα 7.1. Τα αποτελέσματα των ελέγχων αυτών, καθώς και το αντίστοιχο θηκόγραμμα, δεν έδειξαν σοβαρή απόκλιση από την κανονικότητα.



Εικόνα 7.11 Δημιουργία μεταβλητής «post_pre».

Η διατύπωση των ερευνητικών υποθέσεων γίνεται ως εξής:

- Μηδενική υπόθεση: Δεν υπάρχει διαφορά στη μέση τιμή των λαθών στα οποία υποπίπτουν τα νήπια στην αναγνώριση διδιάστατων γεωμετρικών σχημάτων πριν και μετά τη διδακτική παρέμβαση: $H_0: \mu_1 \leq \mu_2$.
- Εναλλακτική υπόθεση: Η μέση τιμή των λαθών στα οποία υποπίπτουν τα νήπια στην αναγνώριση διδιάστατων γεωμετρικών σχημάτων πριν τη διδακτική παρέμβαση είναι μεγαλύτερη από τη μέση τιμή των λαθών μετά τη διδακτική παρέμβαση: $H_A: \mu_1 > \mu_2$ (μονόπλευρου ελέγχου).

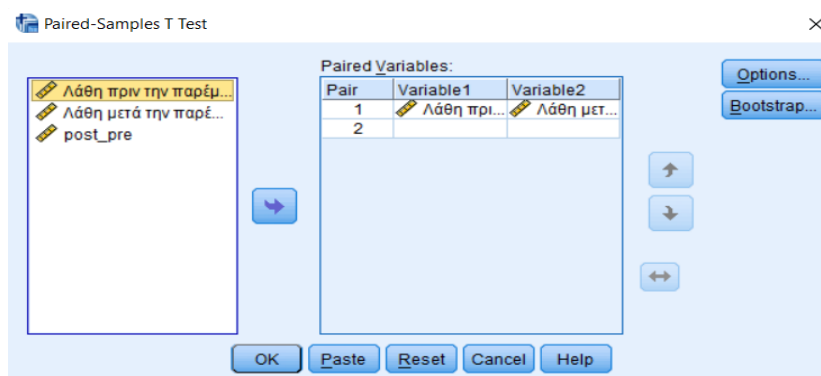
Στο περιβάλλον του SPSS δεν χρειάζεται να ακολουθήσουμε τη διαδικασία δημιουργίας μιας καινούριας μεταβλητής (διαφορά των δύο μετρήσεων). Η διαδικασία βασίζεται στην επιλογή των δύο μετρήσεων

ξεχωριστά. Το SPSS, ωστόσο, για τον έλεγχο αυτόν δημιουργεί μία καινούρια μεταβλητή, την οποία δεν τη βλέπουμε και η όλη διαδικασία ελέγχου βασίζεται σε αυτή.

Αναλυτικότερα, τα βήματα για την εκτέλεση του t-test δειγμάτων κατά ζεύγη είναι τα εξής:

- Επιλέγουμε **Analyze => Compare Means => Paired-Samples T-Test** και ανοίγει το αντίστοιχο πλαίσιο διαλόγου. Επιλέγουμε από αριστερά το ζεύγος των μεταβλητών κάνοντας κλικ σε καθεμία από αυτές, για να τις τοποθετήσουμε στη λίστα **Paired Variables**. Εδώ, θα επιλέξουμε τις μεταβλητές **pre** (Λάθη πριν την παρέμβαση) και **post** (Λάθη μετά την παρέμβαση) (Εικόνα 7.12).

Προσοχή όμως γιατί η σειρά τοποθέτησης καθορίζει το πρόσημο της τιμής του στατιστικού t . Αν, για παράδειγμα, η μέση τιμή της δεύτερης μεταβλητής που θα μετακινήσουμε είναι μεγαλύτερη από τη μέση τιμή της πρώτης μεταβλητής, η τιμή του t θα είναι αρνητική. Έχουμε τη δυνατότητα να επιλέξουμε και άλλα ζεύγη μεταβλητών και να τα τοποθετήσουμε στη λίστα **Paired Variables**, ώστε να εκτελέσουμε ξεχωριστούς ελέγχους t για καθένα ζεύγος.



Εικόνα 7.12 Πλαίσιο διαλόγου «Paired-Samples T-Test».

Τα αποτελέσματα της ανάλυσης που εμφανίζονται στο αρχείο αποτελεσμάτων του SPSS απεικονίζονται σε τρεις πίνακες: Ο πρώτος πίνακας με το όνομα *Paired Sample Statistics* (Εικόνα 7.13) απεικονίζει τα περιγραφικά στατιστικά για τις δύο μετρήσεις πριν και μετά την παρέμβαση (μέσες τιμές, μέγεθος δείγματος, τυπικές αποκλίσεις και τυπικά σφάλματα). Η μέση τιμή των λαθών πριν την παρέμβαση είναι 4,22 με τυπική απόκλιση 0,67, ενώ η μέση τιμή των λαθών μετά την παρέμβαση είναι 3,11 με τυπική απόκλιση 0,78.

Paired Samples Statistics				
		Mean	N	Std. Deviation
Pair 1	Λάθη πριν την παρέμβαση	4,2222	9	,66667
	Λάθη μετά την παρέμβαση	3,1111	9	,78174
				Std. Error Mean
				,22222
				,26058

Εικόνα 7.13 Πίνακας αποτελεσμάτων *Paired Samples Statistics*.

Ο δεύτερος πίνακας με τίτλο *Paired Samples Correlations* (Εικόνα 7.14) εμφανίζει τον συντελεστή συσχέτισης *Pearson* των δύο μεταβλητών (Correlation = 0,426).

Paired Samples Correlations			
		N	Correlation
Pair 1	Λάθη πριν την παρέμβαση & Λάθη μετά την παρέμβαση	9	,426
			Sig.
			,252

Εικόνα 7.14 Πίνακας αποτελεσμάτων *Paired Samples Correlations*.

Ο τρίτος πίνακας με τίτλο *Paired Samples Test* (Εικόνα 7.15) απεικονίζει τα στατιστικά στοιχεία σχετικά με τον έλεγχο. Η μέση τιμή των διαφορών των δύο μετρήσεων (Mean) είναι 1,11 με τυπική απόκλιση 0,78. Το τυπικό σφάλμα είναι 0,26 και το 95% διάστημα εμπιστοσύνης των διαφορών είναι [0,51, -1,71].

Ο έλεγχος της διαφοράς των δύο μετρήσεων είναι μονόπλευρος, αφού μας ενδιαφέρει να ελέγξουμε εάν τα λάθη των νηπίων μειώθηκαν μετά τη διδακτική παρέμβαση. Όμως, όπως έχουμε ήδη αναφέρει, το SPSS υπολογίζει την πιθανότητα για αμφίπλευρο έλεγχο με εναλλακτική υπόθεση $H_A: \mu_d \neq 0$. Στον παρακάτω πίνακα η τιμή του στατιστικού του ελέγχου t είναι θετική και ισούται με 4,264 και η πιθανότητα να απορριφθεί η μηδενική υπόθεση είναι $\text{Sig. (2-tailed)} = 0,003$. Η τιμή p για τον μονόπλευρο έλεγχο είναι το μισό της τιμής p του αμφίπλευρου ελέγχου, δηλαδή $p = 0,003/2 = 0,0015 < 0,05$. Επομένως, απορρίπτεται η μηδενική υπόθεση.

Το παραπάνω αποτέλεσμα μπορεί να συνταχθεί ως αναφορά, στο πλαίσιο μιας επιστημονικής εργασίας, ως εξής: Ελέγξαμε την υπόθεση της διαφοράς των μέσων τιμών των λαθών των νηπίων στην αναγνώριση διδιάστατων γεωμετρικών σχημάτων, μετά την επίδραση μίας διδακτικής παρέμβασης στην τάξη του νηπιαγωγείου, με τη χρήση ενός t -test δειγμάτων κατά ζεύγη. Όπως ανέδειξε ο μονόπλευρος έλεγχος, η μέση τιμή των λαθών στα οποία υποπίπτουν τα νήπια στην αναγνώριση 2D γεωμετρικών σχημάτων πριν τη διδακτική παρέμβαση ($M.T. = 4,22$, $T.A. = 0,67$) είναι σημαντικά μεγαλύτερη από τη μέση τιμή των λαθών μετά τη διδακτική παρέμβαση ($M.T. = 3,11$, $T.A. = 0,78$). [$t_{(8)} = 4,264$, $p = 0,0015$]. Επομένως, μπορούμε να υποστηρίξουμε ότι πέτυχε η διδακτική παρέμβαση, αφού τα λάθη των νηπίων μειώθηκαν μετά από την εφαρμογή της παρέμβασης αυτής στην τάξη.

Paired Samples Test									
		Paired Differences				t	df	Sig. (2-tailed)	
		Mean	Std. Deviation	Std. Error Mean	95% Confidence Interval of the Difference				
					Lower				Upper
Pair 1	Λάθη πριν την παρέμβαση - Λάθη μετά την παρέμβαση	1,111	,78174	,26058	,51022	1,71201	4,264	8	,003

Εικόνα 7.15 Πίνακας αποτελεσμάτων *Paired Samples Test*.

7.2 Προβλήματα για εξάσκηση και ανατροφοδότηση

Άσκηση 1

Χρησιμοποιώντας το αρχείο δεδομένων «educ_environment.sav», να διερευνηθεί εάν διαφοροποιούνται οι απόψεις των κατοίκων εντός μιας πόλης σχετικά με τον βαθμό ενημέρωσης που παρέχουν τα σχολεία στη δημιουργία της πράσινης συνείδησης σε σχέση με αυτούς που ζουν εκτός πόλης.

Στο αρχείο αυτό υπάρχουν δύο μεταβλητές:

«q1»: Περιοχή κατοικίας των συμμετεχόντων με τιμές (1. Στο κέντρο της πόλης, 2. Κοντά στο κέντρο της πόλης, 3. Σχετικά μακριά από το κέντρο της πόλης και 4. Πολύ μακριά από το κέντρο της πόλης).

«q2»: Βαθμός ενημέρωσης που παρέχουν τα σχολεία με σκοπό να καλλιεργηθεί η πράσινη συνείδηση. Η μέτρηση πραγματοποιήθηκε με αυτοσχέδια κλίμακα επτά δηλώσεων στην οποία οι συμμετέχοντες δήλωναν αν συμφωνούν (1) ή όχι (0). Για τη μεταβλητή αυτή υπολογίστηκε το άθροισμα των δηλώσεων όπου οι συμμετέχοντες συμφωνούν.

Να διατυπώσετε τις υποθέσεις αμφίπλευρου ελέγχου.

Να ελέγξετε τις προϋποθέσεις πραγματοποίησης του t -test ανεξάρτητων δειγμάτων και να πραγματοποιήσετε τον έλεγχο στο περιβάλλον του SPSS ,καθώς και να γράψετε τα αποτελέσματα.

Άσκηση 2

Χρησιμοποιώντας το αρχείο δεδομένων «spreadsheets.sav», να διερευνηθεί εάν βελτιώνεται η επίδοση των μαθητών μετά την εμπλοκή τους σε εκπαιδευτικές δραστηριότητες επίλυσης λεκτικού προβλήματος με τη χρήση των Ηλεκτρονικών Λογιστικών Φύλλων (ΗΛΦ). Στο αρχείο αυτό υπάρχουν δύο μεταβλητές «pretest» και «posttest» μέτρησης πριν και μετά, σχετικά με τη βελτίωση της επίδοσης των μαθητών να λύσουν μαθηματικά λεκτικά προβλήματα με τη χρήση ΗΛΦ.

α) Να διατυπώσετε τις υποθέσεις μονόπλευρου ελέγχου. Να ελέγξετε τις προϋποθέσεις πραγματοποίησης του t -test κατά ζεύγη.

β) Να πραγματοποιήσετε τον έλεγχο στο περιβάλλον του SPSS και να γράψετε τα αποτελέσματα.

Άσκηση 3

Να χρησιμοποιηθεί το αρχείο δεδομένων «attitudes_toward_maths.sav» για να απαντηθούν τα παρακάτω:

1. Βελτιώνονται ή χειροτερεύουν οι στάσεις των αγοριών απέναντι στα Μαθηματικά από την Πρώτη στην Τρίτη τάξη Γυμνασίου; Οι μεταβλητές που χρησιμοποιούνται είναι:

Φύλο (sex), Τάξη (class), Άγχος για Μαθηματικά (a), Χρησιμότητα των Μαθηματικών (u), Στάση για επιτυχία στα Μαθηματικά (as), Κίνητρο αποτελεσματικότητας (e) (Σημ. : Οι μεγάλες τιμές, στις 4 μεταβλητές στάσεων που μελετώνται, εκφράζουν κατάλληλη στάση).

2. Ισχύουν οι ίδιες διαπιστώσεις για τον πληθυσμό των κοριτσιών;

Άσκηση 4

Να χρησιμοποιηθεί το αρχείο δεδομένων «schools.sav» για να απαντηθούν τα παρακάτω:

- Να μελετηθούν οι αλλαγές ανάμεσα στις σχολικές χρονιές 1993 και 1994 σε ποσοστά επιτυχίας (μεταβλητές grad93 και grad94) και σε βαθμούς στο τεστ ACT (act93 και act94). Μπορεί να υποστηριχτεί ότι το εκπαιδευτικό σύστημα στην περιοχή αυτή βελτιώθηκε; Ποια σχολεία αποτελούν ακραίες περιπτώσεις;

Βιβλιογραφία

- Γιαλαμάς, Β. (2004). *Στατιστικές Τεχνικές και Εφαρμογές στις Επιστήμες της Αγωγής*. Αθήνα: Εκδ. Πατάκη.
- Γναρδέλλης, Χ. (2013). *Ανάλυση δεδομένων με το IBM SPSS Statistics 21*. Αθήνα: Εκδ. Παπαζήση.
- Lavidas, K., Manesis, D., & Gialamas, V. (2021). Investigation of the Statistical Anxiety Rating Scale psychometric properties with a sample of Greek students. *International Journal of Educational Psychology*, 10(2), 116-142. <http://dx.doi.org/10.17583/ijep.2021.6032>

Κεφάλαιο 8

Έλεγχος χ^2

Σύνοψη

Στις περιπτώσεις εκείνες στις οποίες απαιτείται σύγκριση ανάμεσα στις παρατηρούμενες συχνότητες και στις αναμενόμενες συχνότητες χρησιμοποιείται η θεωρητική κατανομή χ^2 . Στο κεφάλαιο αυτό, αρχικά παρουσιάζεται μια σύνοψη της θεωρητικής κατανομής χ^2 . Στη συνέχεια, στο περιβάλλον του SPSS γίνεται η παρουσίαση της αξιοποίησης της θεωρητικής κατανομής χ^2 για τον έλεγχο α) της καλής προσαρμογής, δηλαδή του πόσο καλά ταιριάζει η κατανομή των παρατηρήσεων με μια γνωστή θεωρητική κατανομή και β) της ανεξαρτησίας δύο ποιοτικών μεταβλητών, δηλαδή του ελέγχου της σχέσης μεταξύ δύο ποιοτικών μεταβλητών.

Προαπαιτούμενη γνώση

Μετασχηματισμός δεδομένων, περιγραφική στατιστική και εισαγωγή στην εκτιμητική και στον έλεγχο των υποθέσεων: Κεφάλαια 2, 4, και 6 του συγγράμματος.

8.1 Έλεγχος χ^2

Μια θεωρητική κατανομή, με μεγάλη πρακτική αξία, είναι η κατανομή χ^2 (chi-square). Η κατανομή αυτή χρησιμοποιείται σε περιπτώσεις όπου απαιτείται σύγκριση ανάμεσα στις παρατηρούμενες συχνότητες και στις θεωρητικές ή αναμενόμενες συχνότητες. Οι αναμενόμενες συχνότητες δημιουργούνται στη βάση κάποιας υπόθεσης (θεωρίας), η οποία είναι ανεξάρτητη από τα διαθέσιμα δεδομένα (Γιαλαμάς, 2004). Το ερώτημα είναι κατά πόσο η διαφορά ανάμεσα στις παρατηρούμενες και αναμενόμενες συχνότητες είναι σημαντική. Εδώ, η μηδενική υπόθεση είναι ότι δεν υπάρχει διαφορά ανάμεσα σε παρατηρούμενες και αναμενόμενες συχνότητες και η εναλλακτική ότι υπάρχει διαφορά. Αν οι παρατηρούμενες συχνότητες διαφέρουν σημαντικά από τις αναμενόμενες, αυτό αποτελεί βάση για την απόρριψη της μηδενικής υπόθεσης.

Ας δούμε ένα απλό παράδειγμα. Ρίχνοντας ένα νόμισμα 100 φορές και, καταμετρώντας κάθε φορά το αποτέλεσμα, παίρνουμε 42 φορές κεφαλή και 58 γράμματα. Αυτές είναι οι παρατηρούμενες συχνότητες (observed frequencies) και τις συμβολίζουμε f_o . Αν το νόμισμα είναι αμερόληπτο (ομοιόμορφη κατανομή), θα έπρεπε να έχουμε 50 φορές κεφαλή και 50 φορές γράμματα. Αυτές τις συχνότητες τις αποκαλούμε αναμενόμενες (expected frequencies) και τις συμβολίζουμε f_e . Τα ερωτήματα που τίθενται είναι με ποιο τρόπο μπορούν να συγκριθούν αυτές οι δύο κατανομές συχνοτήτων και αν είναι λογικό να απορριφθεί η υπόθεση ότι το νόμισμα είναι αμερόληπτο. Όπως θα δούμε στη συνέχεια του Κεφαλαίου αυτού, η εφαρμογή της κατανομής δίνει την απάντηση σε τέτοια ερωτήματα.

Όσον αφορά την κατανομή χ^2 η ποσότητα που εκφράζει τον βαθμό απόκλισης των παρατηρούμενων συχνοτήτων από τις αναμενόμενες συχνότητες είναι η:

$$\chi^2 = \sum \frac{(f_o - f_e)^2}{f_e} \quad (8.1)$$

όπου

f_o = η παρατηρούμενη συχνότητα,

f_e = η αναμενόμενη συχνότητα.

Παρατηρούμε ότι η ποσότητα αυτή είναι πάντα μεγαλύτερη ή ίση με το 0. Το μέγεθός της εξαρτάται από το πόσο διαφέρουν οι αναμενόμενες από τις παρατηρούμενες συχνότητες, αφού η διαφορά $f_o - f_e$ για κάθε αποτέλεσμα εμπλέκεται στον Τύπο 8.1.

Η δειγματοληπτική κατανομή αυτής της ποσότητας θα σκιαγραφηθεί με τη βοήθεια του παραδείγματος της συνάρτησης στο 8.1. Ας δούμε τώρα πώς υπολογίζεται η τιμή της χ^2 στον Πίνακα 8.1.

	f_o	f_e	$f_o - f_e$	$(f_o - f_e)^2$	$\frac{(f_o - f_e)^2}{f_e}$
Κεφάλι	42	50	-8	64	1,28
Γράμματα	58	50	+8	64	1,28
Σύνολο					$\chi^2 = 2,56$

Πίνακας 8.1 Υπολογισμός της στατιστικής τιμής χ^2 .

Με 100 ρίψεις προκύπτουν δύο συχνότητες –μία για τα γράμματα και μία για τις κεφαλές. Αν η συχνότητα στα γράμματα είναι 42, στα κεφάλια είναι 58. Αν αντίστοιχα η συχνότητα είναι 47 στα γράμματα, τότε θα είναι 53 στις κεφαλές. Παρατηρούμε ότι, όταν μεταβάλλεται η μία συχνότητα, η άλλη είναι καθορισμένη. Μία μόνο συχνότητα έχουμε την ελευθερία να μεταβάλουμε. Σε αυτή την περίπτωση ένας βαθμός ελευθερίας συνδέεται με την τιμή της χ^2 .

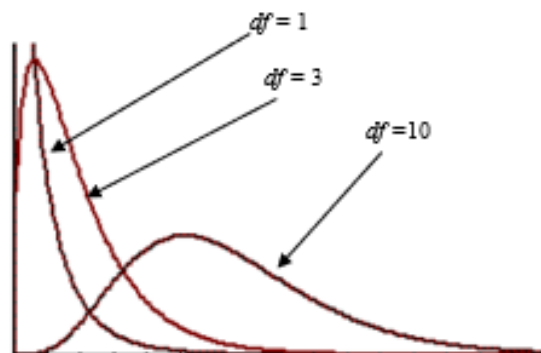
Γενικεύοντας, οι βαθμοί ελευθερίας ορίζονται από τον τύπο:

$$df = k - 1 \quad (8.2)$$

όπου

k = το πλήθος των συχνοτήτων.

Αν ρίχναμε άλλες 100 φορές το νόμισμα, θα είχαμε μια νέα τιμή για την χ^2 . Αν επαναλαμβάναμε τη διαδικασία των 100 ρίψεων πολλές φορές θα είχαμε ένα μεγάλο αριθμό τιμών για την χ^2 . Η δειγματοληπτική κατανομή αυτών των τιμών είναι κατά προσέγγιση η θεωρητική κατανομή χ^2 , με έναν βαθμό ελευθερίας, ο οποίος είναι γνωστής συναρτησιακής μορφής. Πρόκειται, ουσιαστικά, για μια οικογένεια κατανομών, αφού για κάθε τιμή των βαθμών ελευθερίας έχουμε μια διαφορετική κατανομή χ^2 . Το μόνο που χρειάζεται είναι οι βαθμοί ελευθερίας προκειμένου να χρησιμοποιηθεί ο τύπος της κατανομής και να βρεθεί η πιθανότητα μιας συγκεκριμένης τιμής ή, προκειμένου να εκτιμηθεί (με τη βοήθεια της ολοκλήρωσης της συνάρτησης $f(\chi^2)$), η πιθανότητα να επιλεγεί τυχαία μια τιμή μεγαλύτερη ή ίση από τη συγκεκριμένη τιμή. Στην Εικόνα 8.1 δίνονται ορισμένα μέλη της οικογένειας για διάφορους βαθμούς ελευθερίας.



Εικόνα 8.1 Η μορφή της κατανομής για διάφορους βαθμούς ελευθερίας.

Ακολουθεί η παρουσίαση της χρήσης της κατανομής χ^2 σε δύο τύπους ελέγχου: στον έλεγχο καλής προσαρμογής και στον έλεγχο ανεξαρτησίας.

8.1.1 Έλεγχος καλής προσαρμογής

Ο έλεγχος καλής προσαρμογής (goodness of fit) χρησιμοποιείται προκειμένου να συγκριθεί η κατανομή ενός δείγματος με μια θεωρητική κατανομή. Οι αναμενόμενες ή θεωρητικές συχνότητες αναφέρονται σε δείγμα ίσου μεγέθους με εκείνο της έρευνάς μας το οποίο κατανέμεται σε απόλυτη συμφωνία με τη θεωρητική κατανομή που ορίζεται από τη μηδενική υπόθεση. Η μηδενική υπόθεση συνήθως προκύπτει από προηγούμενη έρευνα

(Γιαλαμάς, 2004; Γναρδέλλης, 2003). Η διαδικασία του ελέγχου υπόθεσης καλής προσαρμογής θα δοθεί με τη βοήθεια ενός παραδείγματος, το οποίο συνδέεται με τη χρήση του διαδικτύου.

Πιο συγκεκριμένα, θα διερευνηθεί κατά πόσο τα δεδομένα του τυχαίου δείγματος των 2715 συμμετεχόντων στην κοινωνική έρευνα του 2011 επαρκούν για το συμπέρασμα ότι στην Ελλάδα η πλειονότητα των κατοίκων ηλικίας 15 και άνω χρησιμοποιούσε το διαδίκτυο. Η κατανομή συχνοτήτων των απαντήσεων των συμμετεχόντων παρουσιάζεται στον Πίνακα 8.2.

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	μη χρήστης	1523	56,1	56,2	56,2
	χρήστης	1187	43,7	43,8	100,0
	Total	2710	99,8	100,0	
Missing	System	5	,2		
Total		2715	100,0		

Πίνακας 8.2 Χρήση διαδικτύου και ηλεκτρονικού ταχυδρομείου.

Η διαδικασία του ελέγχου της υπόθεσης γίνεται ως εξής:

1. Στατιστικές Υποθέσεις: Επειδή η έννοια της πλειονότητας αναφέρεται ευθέως σε ένα ποσοστό άνω του 50%, το ερευνητικό ερώτημα σε αυτή την περίπτωση μπορεί να απαντηθεί από τον έλεγχο των παρακάτω στατιστικών υποθέσεων:
 - H_0 : Στον πληθυσμό δεν υπάρχει διαφορά ανάμεσα στις δύο συχνότητες: Μη χρήστης: 50% και Χρήστης: 50%.
 - H_A : Υπάρχει διαφορά μεταξύ των δύο συχνοτήτων.
 Ορίζουμε ως επίπεδο σημαντικότητας $\alpha=0,05$.
2. Το κριτήριο ελέγχου είναι αυτό που δίνεται από την εξίσωση στο (8.1). Η κατανομή που ακολουθεί το κριτήριο ελέγχου, είναι η κατανομή χ^2 , όταν η μηδενική υπόθεση αληθεύει. Οι βαθμοί ελευθερίας εδώ είναι $2-1=1$, αφού υπάρχουν 2 συχνότητες, οι οποίες δεσμεύονται από 1 σχέση ανάμεσά τους –το άθροισμά τους να είναι 2710. Για να ορίσουμε την περιοχή απόρριψης, αναζητούμε την κρίσιμη τιμή στον πίνακα της κατανομής χ^2 με 1 βαθμό ελευθερίας. Αριστερά της ζητούμενης τιμής βρίσκεται το $1-\alpha = 0,95$ του εμβადού, που περικλείεται ανάμεσα στον άξονα χ και την καμπύλη της κατανομής. Η τιμή είναι 3,84 και, επομένως, αν η τιμή που θα προκύψει από τον υπολογισμό του κριτηρίου με τα δεδομένα του δείγματος είναι μεγαλύτερη ή ίση από αυτή, απορρίπτουμε τη μηδενική υπόθεση.
3. Υπολογισμός του κριτηρίου χ^2 για το δείγμα γίνεται στον Πίνακα 8.3

	f_0	f_e	$f_0 - f_e$	$(f_0 - f_e)^2$	$\frac{(f_0 - f_e)^2}{f_e}$
μη χρήστης	1523	1355	168,0	28224	20,83
χρήστης	1187	1355	-168,0	28224	20,83
					$\chi^2 = 41,66$

Πίνακας 8.3 Υπολογισμός της στατιστικής τιμής χ^2 .

4. Επειδή $\chi^2 = 41,66 > 3,84$, απορρίπτεται η μηδενική υπόθεση.

Συμπεραίνουμε, λοιπόν, ότι οι απαντήσεις της χρήσης του διαδικτύου δεν κατανέμονται ομοιόμορφα στον πληθυσμό των πολιτών της Ελλάδας του 2011. Επειδή στο δείγμα η συχνότητα της απάντησης «μη χρήστης» είναι μεγαλύτερη από τη συχνότητα των περιπτώσεων που δήλωσαν «χρήστης», το συμπέρασμα είναι ότι η πλειονότητα των πολιτών δεν χρησιμοποιούσε το διαδίκτυο.

8.1.1.1 Έλεγχος χ^2 καλής προσαρμογής με χρήση λογισμικού (SPSS)

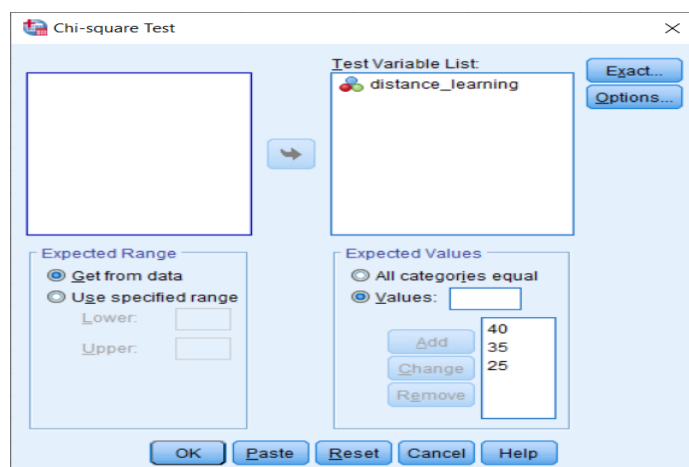
Στον Πίνακα 8.4 παρουσιάζονται οι απαντήσεις των 150 μαθητών σε έρευνα επισκόπησης που πραγματοποιήθηκε με ερώτημα «ποιο τύπο εκπαίδευσης προτιμούν».

Δυνατές απαντήσεις	Συχνότητα
Εξ αποστάσεως	45
Διά Ζώσης	65
Μεικτή	40

Πίνακας 8.4 Κατανομή των απαντήσεων των 150 μαθητών.

Ο ερευνητής καλείται να απαντήσει αν είναι σημαντικά διαφορετικές οι απόψεις των μαθητών για το θέμα σε σχέση με την προηγούμενη έρευνα (θεωρητική κατανομή), στην οποία οι μαθητές είχαν δηλώσει τα ακόλουθα στην ίδια ερώτηση: «Εξ αποστάσεως 40%», «Διά Ζώσης 35%» και «Μεικτή 25%». Επομένως, ο ερευνητής θα πρέπει να συγκρίνει τις συχνότητες στο δείγμα με τις αναμενόμενες συχνότητες, δηλαδή αυτές που προκύπτουν αν εφαρμοστεί στο δείγμα η κατανομή σχετικών συχνοτήτων (%) της προηγούμενης έρευνας.

Προκειμένου να προχωρήσουμε στην ανάλυση στο περιβάλλον του SPSS, θα χρησιμοποιήσουμε το αρχείο «goodness of fit example.sav». Οι εντολές που πρέπει να δώσουμε είναι: **Analyze=> Nonparametric tests => Legacy Dialogs => Chi-square**. Στο παράθυρο διαλόγου που εμφανίζεται (Εικόνα 8.2) εισάγουμε τη μεταβλητή «distance_learning» στο πλαίσιο **Test Variable list**. Στο πλαίσιο **Expected Values**, αφού θέλουμε να ελέγξουμε την κατανομή του δείγματος με τη θεωρητική και όχι με την ομοιόμορφη κατανομή (all categories equal), επιλέγουμε **Values** και εισάγουμε διαδοχικά τα αντίστοιχα ποσοστά της θεωρητικής κατανομής για κάθε κατηγορία.



Εικόνα 8.2. Παράθυρο διαλόγου του ελέγχου χ^2 καλής προσαρμογής.

Μετά το «κλικ» στο κουμπί **OK** του παραθύρου, εμφανίζονται στο περιβάλλον του **Output** του SPSS οι παρακάτω δύο Πίνακες 8.5 και 8.6.

Distance learning			
	Observed N	Expected N	Residual
1 Εξ αποστάσεως	45	60,0	-15,0
2 Διά ζώσης	65	52,5	12,5
3 Μεικτή	40	37,5	2,5
Total	150		

Πίνακας 8.5 Κατανομή παρατηρούμενων και αναμενόμενων συχνοτήτων.

Test Statistics	
	distance learning
Chi-Square	6,893 ^a
df	2
Asymp. Sig.	,032

Chi-Square: η τιμή του στατιστικού χ^2
df: βαθμοί ελευθερίας
Asymp. Sig.: σημαντικότητα της τιμής χ^2

a. 0 cells (0,0%) have expected frequencies less than 5. The minimum expected cell frequency is 37,5.

Πίνακας 8.6 Έλεγχος χ^2 .

Παρατηρούμε στον Πίνακα 8.6 ότι αποκαλύπτεται σημαντική διαφορά παρατηρούμενων και θεωρητικών συχνοτήτων, αφού το p (Asymp. Sig.) = 0,032 < 0,05.

Όσον αφορά την παρουσίαση του αποτελέσματος, γράφουμε τα εξής:

Ελέγξαμε την υπόθεση της θεωρητικής κατανομής των απαντήσεων των μαθητών σχετικά με τον τύπο της εκπαίδευσης που προτιμούν, με τη βοήθεια του ελέγχου χ^2 καλής προσαρμογής. Από τα αποτελέσματα του ελέγχου προκύπτει ότι υπάρχει σημαντική διαφοροποίηση της κατανομής του δείγματος από τη θεωρητική κατανομή, εφόσον $\chi^2(2, N=150) = 6,893$ $p < 0,05$. Οι διαφορές αυτές εντοπίζονται κυρίως στις κατηγορίες «δια ζώσης» και «εξ αποστάσεως» εκπαίδευση. Διαφαίνεται έτσι μια σημαντική μείωση των προτιμήσεων για «εξ αποστάσεως» εκπαίδευση και μια σημαντική αύξηση των προτιμήσεων για «διά ζώσης» εκπαίδευση.

8.1.2 Έλεγχος ανεξαρτησίας δύο ποιοτικών μεταβλητών

Ο έλεγχος ανεξαρτησίας δύο ποιοτικών μεταβλητών μας βοηθά να ελέγξουμε τη συμμεταβολή των τιμών των δύο μεταβλητών, δηλαδή ποιες αλλαγές παρατηρούνται στις τιμές της μιας μεταβλητής όταν μεταβάλλονται οι τιμές της άλλης. Στη διερεύνηση των παραγόντων που συνδέονται με τη συμμετοχή ή όχι στις βουλευτικές εκλογές θα μπορούσε να τεθεί το ερώτημα: πώς διαφοροποιούνται, ως προς το ποσοστό συμμετοχής στις εκλογές, άνδρες και γυναίκες στην Ελλάδα του 2015; Με διαφορετική διατύπωση, ποια είναι η συνάφεια φύλου και συμμετοχής στις εκλογές; Για την περιγραφή της σχέσης χρησιμοποιούνται οι πίνακες σύμπτωσης ή συνάφειας μεταξύ δύο μεταβλητών, καθώς και η κατάλληλη γραφική απεικόνιση της συνάφειας. Η κατανομή χ^2 θα χρησιμοποιηθεί για τη γενίκευση της σχέσης από το δείγμα στον πληθυσμό. Επίσης, θα εκφράσουμε ποσοτικά τον βαθμό της συνάφειας δύο μεταβλητών με τη μορφή των συντελεστών συνάφειας.

8.1.2.1 Πίνακες συνάφειας

Με τον πίνακα συνάφειας αναζητούμε τη σχέση ανάμεσα σε δύο μεταβλητές με ποιοτικά ή κατηγορικά δεδομένα. Συνεχίζοντας με τα δεδομένα από την ευρωπαϊκή κοινωνική έρευνα του 2011, θα μελετηθεί η σχέση ανάμεσα στις μεταβλητές «ηλικιακή κατηγορία» και «Χρήση διαδικτύου και ηλεκτρονικού ταχυδρομείου». Η μεταβλητή «Ηλικιακή κατηγορία» ομαδοποιεί τους συμμετέχοντες σε 5 κατηγορίες, ενώ η μεταβλητή «Χρήση διαδικτύου και ηλεκτρονικού ταχυδρομείου» ομαδοποιεί τους συμμετέχοντες σε δύο κατηγορίες –ένα μέρος των συμμετεχόντων ήταν χρήστες (User) και ένα άλλο δεν ήταν χρήστες (No user).

Αν θελήσουμε όμως να συνδυάσουμε τις δύο μεταβλητές, δηλαδή να αναζητήσουμε πόσοι συμμετέχοντες, οι οποίοι είναι χρήστες διαδικτύου, είναι στην ηλικία των ≤ 35 ή πόσοι είναι 36-45 ετών κ.λπ., πρέπει να λάβουμε υπόψη την ταξινόμηση των συμμετεχόντων κάθε κατηγορίας της μιας μεταβλητής στις κατηγορίες της άλλης. Στο συγκεκριμένο παράδειγμα, οι 2710 συμμετέχοντες θα ταξινομηθούν σε $2 \times 5 = 10$ ομάδες, οι οποίες μπορούν να εκφραστούν από τις διασταυρώσεις γραμμών και στηλών του πίνακα, που έχει ως γραμμές τις 2 κατηγορίες της χρήσης διαδικτύου και ως στήλες τις 5 κατηγορίες ηλικίας.

Ο Πίνακας 8.7, που εξάγεται από το περιβάλλον του SPSS, δίνει μια εικόνα της συμμεταβολής των δύο μεταβλητών. Κάθε κελί του Πίνακα 8.7 περιέχει τον αριθμό των ατόμων που έχουν συγχρόνως τα χαρακτηριστικά της αντίστοιχης γραμμής και στήλης. Για παράδειγμα, στο κελί της δεύτερης γραμμής και της δεύτερης στήλης βρίσκεται ο αριθμός 484, που αντιπροσωπεύει το πλήθος των χρηστών που έχουν ηλικία από 31 έως 25 έτη. Οι αριθμοί στα κελιά ονομάζονται *συνδυαστικές συχνότητες*.

netis e1 Χρήση διαδικτύου και ηλεκτρονικού ταχυδρομείου * agea1 Ηλικιακή κατηγορία								
Crosstabulation								
			agea1 Ηλικιακή κατηγορία					Total
			<= 30	31 - 45	46 - 60	61 - 75	76+	
netuse1 Χρήση διαδικτύου και ηλεκτρονικού ταχυδρομείου	No user	Count	96	317	390	498	221	1522
	User	Count	479	484	188	30	6	1187
Total		Count	575	801	578	528	227	2709

Πίνακας 8.7 Συνάφεια ηλικιακής κατηγορίας και χρήσης διαδικτύου.

Στον ίδιο Πίνακα έχουμε προσθέσει μία στήλη στο δεξί μέρος και μία γραμμή στο κάτω μέρος που ονομάζονται *σύνολα* (Total). Κάθε στοιχείο της στήλης «σύνολο» είναι το άθροισμα των συνδυαστικών συχνοτήτων της γραμμής στην οποία ανήκει. Κάθε στοιχείο της γραμμής «σύνολο» είναι το άθροισμα των συχνοτήτων της στήλης κάτω από την οποία βρίσκεται. Η γραμμή «σύνολο» δείχνει την κατανομή συχνοτήτων της μεταβλητής «ηλικιακή κατηγορία», ενώ η στήλη «σύνολο» την κατανομή της μεταβλητής «χρήση διαδικτύου» στο δείγμα των 2709 ατόμων. Οι απλές κατανομές των μεταβλητών, που δίνονται από τα «σύνολα», λέγονται *περιθώριες κατανομές*.

8.1.2.2 Σχετικές συχνότητες κατηγορικών συμμεταβλητών

Ο στόχος της κατασκευής ενός πίνακα σύμπτωσης είναι η διερεύνηση της σχέσης μεταξύ των δύο μεταβλητών. Η σχέση αυτή εκφράζεται από το αποτέλεσμα της σύγκρισης διάφορων συχνοτήτων που βρίσκονται σε μία γραμμή ή σε μία στήλη. Στον Πίνακα 8.7, για παράδειγμα, συγκρίνοντας τις συνδυαστικές συχνότητες της γραμμής «Χρήστης» συμπεραίνουμε ότι η ηλικιακή κατηγορία 31 - 45 υπερτερεί ελαφρώς σε αριθμό χρηστών της κατηγορίας <=30. Συνεχίζοντας τη σύγκριση των δύο ηλικιακών κατηγοριών, στη γραμμή «μη χρήστης» παρατηρούμε ότι η ηλικιακή κατηγορία 31 - 45 υπερτερεί σημαντικά σε αριθμό μη χρηστών της κατηγορίας <=30. Συνεπώς, στις περιπτώσεις που οι συχνότητες της γραμμής «Σύνολο» έχουν σημαντικές διαφορές μεταξύ τους (575 και 801) η διαδικασία της σύγκρισης δημιουργεί σύγχυση. Η ρύθμιση αυτού του προβλήματος γίνεται με τη δημιουργία σχετικών συχνοτήτων όπου οι συχνότητες της γραμμής «Σύνολο», στον νέο πίνακα, είναι 100%. Οι συνδυαστικές συχνότητες μπορούν να μετατραπούν σε σχετικές συχνότητες με τους εξής τρόπους:

1. Το περιεχόμενο μιας στήλης αποτελεί τον πίνακα απλής κατανομής της μεταβλητής «χρήση διαδικτύου» για τη συγκεκριμένη ηλικιακή κατηγορία στην οποία αντιστοιχεί η στήλη. Από τον Πίνακα 8.7 υπολογίζονται οι σχετικές συχνότητες για την πρώτη στήλη (<=30):
 $1 \text{ τετραγώνου} = 100 \cdot (96/575) = 16,7\%$
 $2 \text{ τετραγώνου} = 100 \cdot (479/575) = 83,3\%$
 Το 83,3% των συμμετεχόντων ηλικίας <=30 ετών ήταν χρήστες του διαδικτύου, ενώ το υπόλοιπο 16,7% της ίδιας ηλικιακής κατηγορίας δεν ήταν χρήστες. Με τον υπολογισμό των σχετικών συχνοτήτων για όλες τις ηλικιακές κατηγορίες η σύγκριση όλων των ποσοστών μιας γραμμής του πίνακα οδηγεί σε συμπέρασμα αναφορικά με τη σχέση των δύο μεταβλητών ή την επίδραση της ηλικίας στη χρήση του διαδικτύου.
2. Με τον ίδιο τρόπο εργαζόμαστε για τις γραμμές, υπολογίζοντας τις σχετικές συχνότητες των απλών κατανομών συχνοτήτων, τις οποίες εκφράζουν οι γραμμές του Πίνακα 8.7. Για κάθε κατηγορία της μεταβλητής «ηλικιακή κατηγορία» έχουμε μια κατανομή συχνοτήτων της μεταβλητής «Χρήση διαδικτύου και ηλεκτρονικού ταχυδρομείου».

Οι κατανομές σχετικών συχνοτήτων μπορούν να δοθούν σε ξεχωριστούς πίνακες, έναν για κάθε περίπτωση, ή να γραφούν στον ίδιο πίνακα. Συνήθως, στις στήλες τοποθετούνται οι κατηγορίες της μεταβλητής η οποία θεωρείται ανεξάρτητη, ενώ στις γραμμές οι κατηγορίες της μεταβλητής η οποία θεωρείται εξαρτημένη. Τότε χρησιμοποιούνται οι σχετικές συχνότητες των στηλών, δηλαδή της θέσης που έχουν τοποθετηθεί οι κατηγορίες της ανεξάρτητης μεταβλητής. Στον Πίνακα 8.8 δίνονται οι σχετικές συχνότητες των στηλών, καθώς

και οι σχετικές συχνότητες για τη στήλη «περιθώριο» που εκφράζει την κατανομή της μεταβλητής «Χρήση διαδικτύου και ηλεκτρονικού ταχυδρομείου» για το σύνολο των 2709 συμμετεχόντων.

8.1.2.3 Περιγραφή της συνάφειας κατηγορικών μεταβλητών από πίνακα σύμπτωσης

Για την περιγραφή της συνάφειας ανάμεσα στις δύο μεταβλητές ελέγχεται ο εξαγόμενος από το SPSS Πίνακας σύμπτωσης 8.8. Αν οι σχετικές συχνότητες κάθε κατηγορίας της εξαρτημένης μεταβλητής που έχουν τοποθετηθεί στις γραμμές είναι ίσες μεταξύ τους, τότε δεν υπάρχει συσχέτιση μεταξύ των μεταβλητών. Αν, αντίθετα, παρατηρηθεί αξιόλογη διαφορά ανάμεσά τους, συμπεραίνουμε ότι υπάρχει συνάφεια μεταξύ τους. Η μελέτη του πίνακα οδηγεί στο συμπέρασμα, αναφορικά με το δείγμα των 2709 περιπτώσεων, ότι η ηλικία συνδέεται ισχυρά με τη χρήση του διαδικτύου, με το ποσοστό χρήσης να μειώνεται διαρκώς από τη μικρότερη στη μεγαλύτερη ηλικία (από 83% έως 2,6%). Αργότερα θα δούμε πώς υπολογίζουμε τον βαθμό αυτής της συνάφειας με τη βοήθεια κάποιου δείκτη.

8.1.2.4 Γραφική απεικόνιση συνάφειας

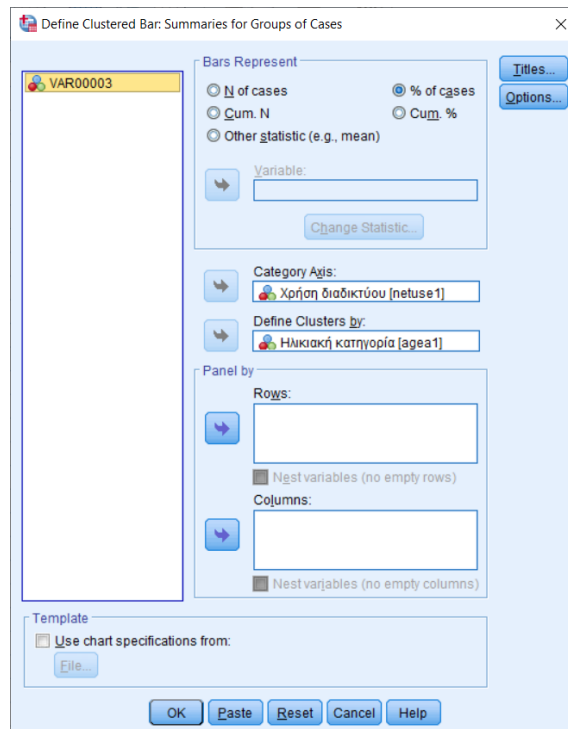
Με πολλούς τρόπους μπορούμε να απεικονίσουμε γραφικά το αποτέλεσμα της προηγούμενης παραγράφου. Σε αυτήν την υποενότητα θα παρουσιαστούν τα σύνθετα ραβδογράμματα. Στο μενού του περιβάλλοντος του SPSS επιλέγουμε διαδοχικά **Graphs=>Legacy Dialogs=>Bar=>clustered "Summaries for groups of cases"=>Define**. Στο παράθυρο διαλόγου που εμφανίζεται (Εικόνα 8.3) τοποθετούμε τη μεταβλητή «netuse1» (Χρήση διαδικτύου) στο **Category axis** και τη μεταβλητή «agea1» (Ηλικιακή κατηγορία) στο **Define clusters by**. Επίσης, επιλέγουμε **Bars Represent: % of cases**, δηλαδή να εμφανίζουμε τα ποσοστά των τιμών της μεταβλητής

netuse1 Χρήση διαδικτύου και ηλεκτρονικού ταχυδρομείου * agea1 Ηλικιακή κατηγορία								
			agea1 Ηλικιακή κατηγορία					Total
			<= 30	31 - 45	46 - 60	61 - 75	76+	
netuse1 Χρήση διαδικτύου και ηλεκτρονικού ταχυδρομείου	No user							
		% within agea1 Ηλικιακή κατηγορία	16,7%	39,6%	67,5%	94,3%	97,4%	56,2%
	User							
		% within agea1 Ηλικιακή κατηγορία	83,3%	60,4%	32,5%	5,7%	2,6%	43,8%
Total								
		% within agea1 Ηλικιακή κατηγορία	100,0%	100,0 %	100,0 %	100,0 %	100,0 %	100,0 %

Each subscript letter denotes a subset of agea1 Ηλικιακή κατηγορία categories whose column proportions do not differ significantly from each other at the ,05 level.

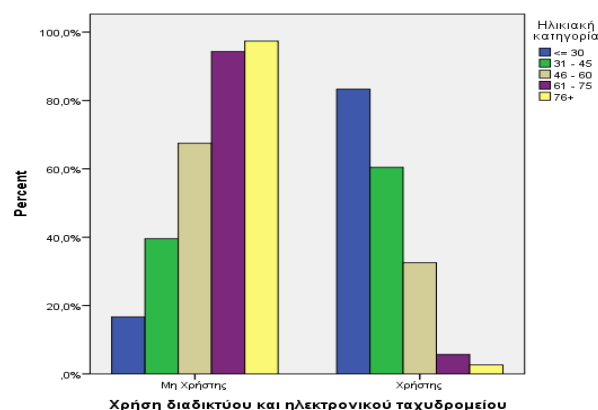
«netuse1» για κάθε τιμή της μεταβλητής «agea1».

Πίνακας 8.8 Συνάφεια ηλικιακής κατηγορίας και χρήσης διαδικτύου: Σχετικές συχνότητες στηλών.



Εικόνα 8.3 Πλαίσιο διαλόγου για τη γραφική απεικόνιση της συνάφειας.

Το τελικό σχήμα ονομάζεται *σύνθετο ραβδόγραμμα*. Στην Εικόνα 8.4 απεικονίζεται γραφικά το περιεχόμενο του Πίνακα 8.8.



Εικόνα 8.4 Συνάφεια ηλικιακής κατηγορίας και χρήσης διαδικτύου.

8.1.2.5 Υπολογισμός της τιμής χ^2 σε πίνακες σύμπτωσης και βαθμοί ελευθερίας

Η τιμή χ^2 που δίνει την απόκλιση των παρατηρούμενων συχνοτήτων του πίνακα σύμπτωσης, από τις αναμενόμενες συχνότητες, που εκφράζουν την ανεξαρτησία των δύο μεταβλητών, είναι η ίδια που χρησιμοποιήθηκε στον έλεγχο καλής προσαρμογής (Εξίσωση 8.1, $\chi^2 = \sum \frac{(f_o - f_e)^2}{f_e}$).

Ας επανέλθουμε στον Πίνακα 8.7, για να ορίσουμε τη αναμενόμενη και παρατηρούμενη συχνότητα.

Η παρατηρούμενη f_o για κάθε τετράγωνο είναι αυτή που φαίνεται σε κάθε τετράγωνο του Πίνακα 8.7. Για παράδειγμα, στη διασταύρωση πρώτης γραμμής και δεύτερης στήλης η παρατηρούμενη συχνότητα είναι 317, ενώ στη διασταύρωση της δεύτερης γραμμής και πρώτης στήλης είναι 479. Η αναμενόμενη συχνότητα f_e ενός κελιού ορίζεται ως η συχνότητα που βρίσκεται σε απόλυτη συμφωνία με τη μηδενική υπόθεση. Η μηδενική υπόθεση στην περίπτωση του πίνακα σύμπτωσης μπορεί να εκφραστεί ως εξής: *H₀*: Η κατανομή σχετικών συχνοτήτων της χρήσης διαδικτύου έχει την ίδια μορφή σε όλες της ηλικιακές κατηγορίες

Τώρα, θα πρέπει να οριστεί πώς θα έπρεπε να κατανέμεται το δείγμα σύμφωνα με τη μηδενική υπόθεση. Αναφέρθηκε προηγουμένως ότι σύμφωνα με τη μηδενική υπόθεση η κατανομή των σχετικών συχνοτήτων της χρήσης πρέπει να είναι η ίδια για όλες τις ηλικιακές κατηγορίες. Αλλά ποια είναι αυτή η κατανομή; Η καλύτερη εκτίμηση που διαθέτουμε για την κατανομή αυτή προέρχεται από τα δεδομένα του δείγματος των 2709 περιπτώσεων και δίνεται στη στήλη του περιθωρίου του Πίνακα 8.8. Επομένως, αν τα άτομα της ηλικίας ≤ 30 ακολουθούσαν αυτήν την κατανομή, οι χρήστες διαδικτύου θα έπρεπε να αποτελούν το 43,8% περίπου του συνόλου των 575 περιπτώσεων αυτής της ηλικίας.

Άρα ο αριθμός των χρηστών αυτών θα ήταν:

$$f_e = \frac{43,817 \cdot 575}{100} = 251,95.$$

Με τις ίδιες ακριβώς σκέψεις μπορούν να υπολογιστούν οι αναμενόμενες συχνότητες των υπόλοιπων κελιών. Η διαδικασία αυτή μπορεί να γραφτεί με τη μορφή ενός απλού τύπου. Η αναμενόμενη συχνότητα του κελιού K_{ij} στη διασταύρωση της i γραμμής και j στήλης είναι:

$$f_e = \frac{r_i \cdot c_j}{N} \quad (8.3)$$

όπου

r_i και c_j τα σύνολα της γραμμής i και στήλης j ,

N το μέγεθος του δείγματος.

Ως παράδειγμα δίνεται ο υπολογισμός του κελιού (χρήστης), (≤ 35) που συμφωνεί με τον υπολογισμό που είδαμε παραπάνω:

$$f_e = \frac{r_2 \cdot c_1}{N} = \frac{1187 \cdot 575}{2709} = 251,95.$$

Ολόκληρη η διαδικασία υπολογισμού της τιμής χ^2 αποτυπώνεται στον Πίνακα 8.9.

Κελί	f_o	f_e	$(f_o - f_e)^2 / f_e$
Γραμμή 1, στήλη 1	96	323,053	159,581
Γραμμή 1, στήλη 2	317	450,027	39,322
Γραμμή 1, στήλη 3	390	324,738	13,116
Γραμμή 1, στήλη 4	498	296,647	136,671
Γραμμή 1, στήλη 5	221	127,536	68,495
Γραμμή 2, στήλη 1	479	251,947	204,619
Γραμμή 2, στήλη 2	484	350,973	50,42
Γραμμή 2, στήλη 3	188	253,262	16,817
Γραμμή 2, στήλη 4	30	231,353	175,243
Γραμμή 2, στήλη 5	6	99,464	87,826
			$\chi^2 = 952,11$

Πίνακας 8.9 Υπολογισμός της τιμής χ^2 για τα δεδομένα του Πίνακα 8.7.

Προκειμένου να καθοριστεί κατά πόσο η τιμή του στατιστικού χ^2 είναι σημαντικά μεγάλη, πρέπει να βρεθούν οι βαθμοί ελευθερίας df και στη συνέχεια να ανατρέξουμε στον πίνακα της κατανομής χ^2 . Για τον έλεγχο της χ^2 ανεξαρτησίας θα πρέπει να υπολογιστεί ο αριθμός των κελιών για τα οποία μπορούν να οριστούν ελεύθερα αναμενόμενες συχνότητες. Όπως είδαμε, η αναμενόμενη συχνότητα υπολογίζεται με τη βοήθεια του μεγέθους του δείγματος και των συνόλων των γραμμών και στηλών του αρχικού πίνακα σύμπτωσης. Αυτά τα σύνολα

περιορίζουν την ελευθερία στην επιλογή των αναμενόμενων συχνοτήτων. Στον Πίνακα 8.7 διαπιστώνεται ότι μόνο 4 κελιά μπορούν να ορισθούν ελεύθερα. Γενικεύοντας, σε κάθε πίνακα μπορούν να καθοριστούν όλα τα κελιά μίας γραμμής εκτός από ένα και όλα τα κελιά μίας στήλης εκτός από ένα. Άρα ο συνολικός αριθμός των f_e που μπορούν να ορισθούν ελεύθερα και οι βαθμοί ελευθερίας df είναι:

$$df = (R - 1) \cdot (C - 1) \quad (8.4)$$

όπου

R = το πλήθος γραμμών,

C = το πλήθος στηλών.

8.1.2.6 Διαδικασία ελέγχου ανεξαρτησίας

Η πιο συχνή εφαρμογή της κατανομής χ^2 γίνεται με σκοπό τον έλεγχο ανεξαρτησίας δύο κατηγορικών μεταβλητών που παρατηρούνται στον ίδιο πληθυσμό. Στην προκειμένη περίπτωση θα παρακολουθήσουμε την πορεία ελέγχου μέσα από το παράδειγμα που ήδη μελετήσαμε και για το οποίο υπολογίστηκε η τιμή χ^2 στην προηγούμενη ενότητα. Η διαδικασία ελέγχου είναι η παρακάτω:

Έστω ότι δεχόμαστε ότι το δείγμα είναι τυχαίο:

1. Υποθέσεις αναφορικά με τον πληθυσμό των πολιτών στην Ελλάδα του 2011:
 H_0 : Η χρήση διαδικτύου και η ηλικιακή κατηγορία είναι ανεξάρτητες.
 H_A : Υπάρχει σχέση ανάμεσα στις δύο μεταβλητές.
 Το επίπεδο σημαντικότητας του ελέγχου είναι $\alpha=0,05$.
2. Το στατιστικό είναι αυτό που δίνεται από την εξίσωση. Η κατανομή του κριτηρίου είναι η χ^2 . Οι βαθμοί ελευθερίας υπολογίζονται στην περίπτωση ελέγχου ανεξαρτησίας από τον τύπο:
 $df = (R-1)(C-1) = (5-1)(2-1) = 4$.
 Για $df=4$ και $\alpha=0,05$, με τη βοήθεια του πίνακα της κατανομής, βρίσκουμε ότι η κρίσιμη τιμή είναι $\chi^2_{\kappa\rho} = 9,49$. Θα απορριφθεί η μηδενική υπόθεση, αν $\chi^2 \geq 9,49$.
3. Ο υπολογισμός του στατιστικού για το δείγμα μάς δόθηκε στον Πίνακα 8.3. Έχουμε: $\chi^2=952,11$.
4. Επειδή $\chi^2=952,11 > \chi^2_{\kappa\rho} = 9,49$, απορρίπτουμε τη μηδενική υπόθεση. Συμπεραίνουμε την ύπαρξη συνάφειας ανάμεσα σε χρήση διαδικτύου και ηλικιακή κατηγορία. Αυτό μπορεί να αναφερθεί ως σημαντική σχέση με $\chi^2(4, N=2709) = 952,11, p<0,05$. Μετά από όσα αναφέρθηκαν, παρατηρώντας τον πίνακα σχετικών συχνοτήτων και το σύνθετο ραβδόγραμμα διαπιστώνεται ότι όσο μεγαλύτερη είναι η ηλικία τόσο μικρότερο είναι το ποσοστό χρήσης του διαδικτύου ή το ποσοστό χρήσης του διαδικτύου φθίνει σημαντικά από τη μια ηλικιακή κατηγορία στην αμέσως επόμενη.

8.1.2.7 Έλεγχος ποσοστών στηλών

Μετά τη διαπίστωση της σημαντικότητας της συνάφειας στον έλεγχο ανεξαρτησίας με την χ^2 κατανομή, στην περίπτωση πινάκων σύμπτωσης μεγαλύτερων διαστάσεων από τον πίνακα 2x2 απαιτείται παραπέρα διερεύνηση σχετικά με τις διαφορές των ποσοστών μιας γραμμής. Για παράδειγμα, ως προς το ποσοστό χρηστών διαδικτύου ποιες από τις ηλικιακές κατηγορίες διαφέρουν σημαντικά μεταξύ τους;

Για κάθε γραμμή του πίνακα, αν π_i και π_j τα ποσοστά δύο κατηγοριών i και j της ανεξάρτητης μεταβλητής στον πληθυσμό, θα πρέπει να ελεγχθούν ως $C(C-1)/2$ οι μηδενικές υποθέσεις της μορφής: $H_0: \pi_i = \pi_j$.

Το στατιστικό z του ελέγχου αυτών των υποθέσεων, το οποίο ακολουθεί την τυπική κανονική κατανομή (Γναρδέλλης, 2003) είναι:

$$z = \frac{p_i - p_j}{\sqrt{p_{ij}(1-p_{ij})\left(\frac{1}{c_i} + \frac{1}{c_j}\right)}} \quad (8.5)$$

$$\text{όπου } p_{ij} = \frac{c_i \cdot p_i + c_j \cdot p_j}{c_i + c_j},$$

p_{ij} είναι το σταθμισμένο μέσο ποσοστό των δύο στηλών,

p_i και p_j τα ποσοστά των στηλών i και j ,
 c_i και c_j τα σύνολα των στηλών.

Επειδή διεξάγονται $C(C-1)/2$ έλεγχοι προτείνεται η διόρθωση *Bonferroni* σχετικά με το επίπεδο σημαντικότητας α . Το διορθωμένο επίπεδο δίνεται από τη σχέση $\alpha_B = \alpha/[C(C-1)/2]$.

Για τη σύγκριση των ποσοστών των χρηστών διαδικτύου ανάμεσα στις κατηγορίες ≤ 30 και 31-45 από τα στοιχεία των Πινάκων 8.8 και 8.9 υπολογίζουμε:

$$P_{ij} = (0,833 \cdot 575 + 0,604 \cdot 801) / (575 + 801) = 0,7$$

$$z = \frac{0,833 - 0,604}{\sqrt{0,7(1 - 0,7) \left(\frac{1}{575} + \frac{1}{801} \right)}} = 9,14$$

όπου

$$\alpha_B = 0,5/[5(5-1)/2] = 0,005.$$

Από τον πίνακα της κανονικής κατανομής βρίσκουμε $Z_{\alpha_B} = \pm 2,81$ και απορρίπτουμε τη μηδενική υπόθεση, διότι το στατιστικό $z = 9,14 > 2,81$. Συνεπώς, το ποσοστό χρηστών στην ηλικία ≤ 30 (83,3%) είναι σημαντικά μεγαλύτερο του αντίστοιχου ποσοστού της ηλικίας 31-45 (60,4%).

8.1.2.8 Προϋποθέσεις στην εκτέλεση του ελέγχου χ^2

Μία από τις πιο σημαντικές προϋποθέσεις για τη διεξαγωγή ελέγχου με την κατανομή χ^2 αφορά το μέγεθος των αναμενόμενων συχνοτήτων. Ενδιαφέρουσα συζήτηση για τους λόγους που οδηγούν στο πρόβλημα των μικρών αναμενόμενων συχνοτήτων μπορούν να βρουν οι ενδιαφερόμενοι στον Howell (1997). Αν και δεν υπάρχει απόλυτη συμφωνία για το πώς θα αντιμετωπίζεται το πρόβλημα των μικρών αναμενόμενων συχνοτήτων, πολλοί συγγραφείς ακολουθούν την προσέγγιση του Cochran (1954). Ο συγγραφέας αυτός προτείνει να μη χρησιμοποιείται η χ^2 σε πίνακες σύμπτωσης με περισσότερους από έναν βαθμούς ελευθερίας, αν η μικρότερη αναμενόμενη συχνότητα είναι μικρότερη του 1 ή το ποσοστό των κελιών με αναμενόμενη συχνότητα, η οποία είναι μικρότερη από 5, ξεπερνά το 20%. Για να ικανοποιηθεί αυτή η προϋπόθεση, είναι δυνατόν σε μεγάλους πίνακες να συγχωνευτούν γειτονικές κατηγορίες γραμμών ή στηλών, αρκεί αυτή η ομαδοποίηση να έχει λογική στήριξη. Για παράδειγμα, θα ήταν δυνατόν να αποτελέσουν μία κατηγορία εισοδήματος οι κατηγορίες «πολύ χαμηλό» και «χαμηλό» εισόδημα, μιας κλίμακας με πέντε κατηγορίες («πολύ χαμηλό», «χαμηλό», «μεσαίο», «υψηλό», «πολύ υψηλό»). Η πρακτική αυτή, παρά το ότι χρησιμοποιήθηκε συχνά στο παρελθόν αλλά χρησιμοποιείται και σήμερα, είναι προτιμότερο να αποφεύγεται, αν διαθέτουμε λογισμικό που εκτελεί διαδικασίες ακριβούς ελέγχου της ίδιας αρχικής μηδενικής υπόθεσης. Μια τέτοια τεχνική ακριβούς ελέγχου, για πίνακες 2x2, είναι αυτή του Fisher. Σε εξειδικευμένα στατιστικά λογισμικά υπάρχουν ακριβείς τεχνικές για πίνακες μεγαλύτερων διαστάσεων. Προς μία αναλυτική παρουσίαση και συζήτηση προτείνονται πιο εξειδικευμένα κείμενα (Γναρδέλλης, 2003; Metha, & Patel, 1983, 1986). Δηλαδή, ο ακριβής έλεγχος του Fisher προϋποθέτει σταθερές συχνότητες περιθωρίων, πράγμα που σπάνια είναι ρεαλιστικό και αυτός είναι και ο λόγος που δεν προτείνεται από πολλούς συγγραφείς.

8.1.2.9 Συντελεστές συνάφειας

Ο έλεγχος με την κατανομή χ^2 οδηγεί στην απόρριψη ή μη της μηδενικής υπόθεσης που αναφέρεται στην ανεξαρτησία μεταξύ δύο μεταβλητών, όταν τα δεδομένα δίνονται με τη μορφή ενός πίνακα σύμπτωσης. Αλλά και μετά την απόρριψη της μηδενικής υπόθεσης, δεν γνωρίζουμε πολλά πράγματα για τον βαθμό αυτής της συνάφειας. Η αξιολόγηση της τιμής χ^2 και μόνο θα οδηγούσε σε εσφαλμένα συμπεράσματα, αφού η τιμή αυτή εξαρτάται άμεσα από το μέγεθος του δείγματος το οποίο κατανέμεται στον πίνακα σύμπτωσης. Σε αυτή την ενότητα θα αναφερθούν ορισμένοι συντελεστές συνάφειας που εκφράζουν με τυποποιημένο τρόπο το μέγεθος της συνάφειας μεταξύ δύο μεταβλητών, χρησιμοποιώντας τις συχνότητες του πίνακα της συνάφειας (Γιαλαμάς, 2004). Αυτοί οι συντελεστές ονομάζονται *μέτρα συνάφειας*.

Ένας από τους πιο γνωστούς συντελεστές, που χρησιμοποιεί την τιμή του στατιστικού χ^2 και υπολογίζεται από τον πίνακα συμπτώσεως δύο μεταβλητών, ονομάζεται *συντελεστής συμπτώσεως* (C) (Γναρδέλλης, 2003).

Μπορεί να υπολογιστεί σε πίνακα οποιασδήποτε διάστασης και δίνεται από τον τύπο:

$$C = \sqrt{\frac{\chi^2}{\chi^2 + N}} \quad (8.6)$$

όπου,

χ^2 = η ποσότητα του τύπου και

N = το μέγεθος του δείγματος.

Ο συντελεστής C παίρνει τιμές από 0 έως 1 στο πλαίσιο ενός διαστήματος. Γίνεται κατανοητό ότι η τιμή του εξαρτάται άμεσα από τη διαφορά ανάμεσα στις αναμενόμενες και τις παρατηρούμενες συχνότητες. Η σύμπτωση των δύο συχνοτήτων για κάθε τετράγωνο του πίνακα σύμπτωσης ισοδυναμεί με $C = 0$ και εκφράζει την απόλυτη έλλειψη συνάφειας. Όσο απομακρύνεται η τιμή του από το 0, τόσο μεγαλύτερη είναι η συνάφεια μεταξύ των μεταβλητών.

Στην περίπτωση, τώρα, των πινάκων 2×2 χρησιμοποιείται ο συντελεστής ϕ , προκειμένου να εκφραστεί η σχέση ανάμεσα σε δύο μεταβλητές. Ο συγκεκριμένος συντελεστής περιγράφεται από τον τύπο:

$$\phi = \sqrt{\frac{\chi^2}{N}}$$

Ο συντελεστής ϕ οδηγεί στην ίδια ακριβώς τιμή με τον συντελεστή συσχέτισης r , ($|r|$), ο οποίος μελετάται στο Κεφάλαιο 11 της γραμμικής συσχέτισης. Όταν, για παράδειγμα, κωδικοποιηθεί το φύλο με 1 ή 2 για τα αγόρια και για τα κορίτσια, αντίστοιχα, και η μεταβλητή «Λήψη απόφασης» για κάποιο επάγγελμα, κωδικοποιηθεί με 1 για το όχι και 2 για το ναι, ο συντελεστής συσχέτισης r (Pearson) δίνει την τιμή του ϕ και το νόημα του συντελεστή προκύπτει από το γεγονός ότι η τιμή r^2 , άρα και ϕ^2 , εκφράζει το ποσοστό της διακύμανσης της μιας μεταβλητής που ερμηνεύεται από τη σχέση της με την άλλη.

Ένα μεγάλο πρόβλημα του συντελεστή ϕ που συζητήθηκε παραπάνω είναι ο περιορισμός του σε πίνακες διαστάσεων 2×2 . Μια λύση του προβλήματος προτάθηκε από τον Cramér (1946) μέσω του συντελεστή *Cramér's V*:

$$V = \sqrt{\frac{\chi^2}{N(k-1)}} \quad (8.7)$$

όπου

N = το μέγεθος του δείγματος και

k = η τιμή της μικρότερης διάστασης του πίνακα.

Ο συντελεστής *Cramér's V* μπορεί να θεωρηθεί ως μια απλή επέκταση του ϕ για $k > 1$ και ισχύει, φυσικά, το γεγονός ότι σε πίνακα 2×2 ο *Cramér's V* είναι ϕ (Γιαλαμάς, 2004). Το σημαντικό εδώ, σε αντίθεση με τον συντελεστή σύμπτωσης C , είναι ότι η μέγιστη τιμή του *Cramér's V* είναι 1 ανεξάρτητα από τις διαστάσεις του πίνακα. Στην περίπτωση που θέλαμε να κρατήσουμε μόνο ένα μέτρο συνάφειας, θα επιλέγαμε τον συντελεστή *Cramér's V*, αφού δεν περιορίζεται από τις διαστάσεις του πίνακα και έχει μία ερμηνεία για την περίπτωση τουλάχιστον του πίνακα 2×2 .

Ο συντελεστής V για το παράδειγμα της συνάφειας «χρήση διαδικτύου» και «ηλικιακής κατηγορίας» με δεδομένα του Πίνακα 8.8 είναι:

$$V = \sqrt{\frac{\chi^2}{N(k-1)}} = \sqrt{\frac{952,11}{2709(2-1)}} = 0,593.$$

Η τιμή $V = 0,59$ είναι μια επίσης ιδιαίτερα υψηλή τιμή, όπως εκείνη του συντελεστή C .

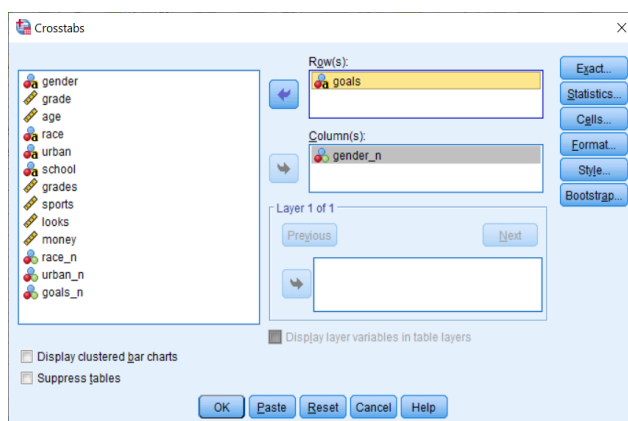
8.1.2.10 Πίνακες συνάφειας και έλεγχος χ^2 ανεξαρτησίας στο περιβάλλον του SPSS

Σχετικά με τη διερεύνηση του παράγοντα φύλου στην εξήγηση της προτεραιότητας στόχων των μαθητών θα μπορούσε να τεθεί το ερώτημα: Διαφέρουν και, αν ναι, με ποιον τρόπο οι προτεραιότητες στόχων των αγοριών και των κοριτσιών; Με μια διαφορετική διατύπωση, ποια είναι η συνάφεια φύλου και προτεραιότητας στόχων των μαθητών; Προκειμένου να απαντηθούν τα προηγούμενα ερωτήματα χρησιμοποιούνται οι πίνακες συνάφειας ή σύμπτωσης μεταξύ δύο μεταβλητών, καθώς και η κατάλληλη γραφική απεικόνιση της συνάφειας. Επιπλέον, η κατανομή χ^2 θα χρησιμοποιηθεί για τη γενίκευση της σχέσης από το δείγμα στον πληθυσμό. Τέλος, θα εκφράσουμε ποσοτικά τον βαθμό της συνάφειας των δύο μεταβλητών με την παρουσίαση συγκεκριμένων συντελεστών συνάφειας.

Για τη διερεύνηση του προηγούμενου ελέγχου χρησιμοποιούμε το αρχείο «data_kids.sav». Τα δεδομένα του προέρχονται από την έρευνα των Chase και Dummer (1992), οι οποίοι μελέτησαν τον ρόλο των σπορ ως καθοριστικό κοινωνικό παράγοντα για τους μαθητές 10-12 ετών. Οι μεταβλητές που θα χρησιμοποιηθούν είναι οι: φύλο (μεταβλητή «gender_n») και η προτεραιότητα στόχων του μαθητή (μεταβλητή «goals»). Οι μαθητές του δείγματος δήλωσαν ποιο ήταν πιο σημαντικό μεταξύ των: «καλός βαθμός», «δημοτικότητα», «σπορ».

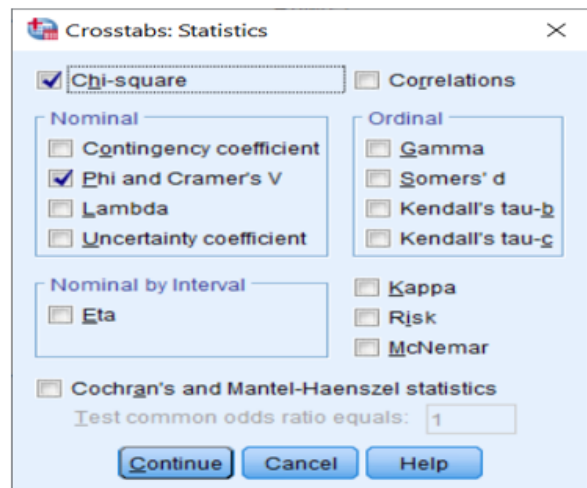
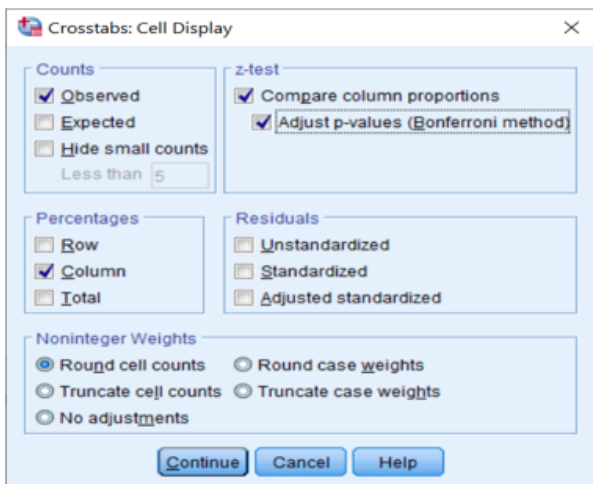
Προκειμένου να διερευνηθεί η σχέση δύο κατηγορικών μεταβλητών κατασκευάζεται ο πίνακας συνάφειας με n , X , k αριθμό κελιών, όπου n ο αριθμός των κατηγοριών της μίας μεταβλητής και k ο αριθμός κατηγοριών της άλλης μεταβλητής. Στις στήλες του πίνακα βάζουμε την ανεξάρτητη μεταβλητή, ενώ στις γραμμές την εξαρτημένη.

Ο πίνακας συνάφειας στο SPSS κατασκευάζεται ως εξής: Επιλέγουμε από το μενού **Analyze** => **Descriptive Statistics** => **Crosstabs**, εισάγουμε στο πλαίσιο **Row** τη μεταβλητή «goals» και στο πλαίσιο **Column** τη μεταβλητή «gender_n» (Εικόνα 8.5).



Εικόνα 8.5 Πλαίσιο διαλόγου για την κατασκευή του πίνακα συνάφειας.

Χρησιμοποιώντας τα κουμπιά επιλογής έχουμε: Στην επιλογή **Cells** στο πλαίσιο **Percentages** τσεκάρουμε **Column**. Επίσης, στην ίδια επιλογή στο πλαίσιο **z-test** τσεκάρουμε την επιλογή **Compare column proportions** και την επιλογή **Adjust p-values (Bonferroni method)** και έπειτα πατάμε **Continue**. Από την άλλη πλευρά, στην επιλογή **Statistics** τσεκάρουμε **Chi-square** και στο πλαίσιο **Nominal** επιλέγουμε **Phi and Cramer's V** και έπειτα πατάμε **Continue** (Εικόνα 8.6).



Εικόνα 8.6 Επιλογές ελέγχου χ^2 .

Μετά το «κλικ» στο κουμπί **OK** του παραθύρου εμφανίζονται στο περιβάλλον του **Output** του SPSS τέσσερις πίνακες:

Ο Πίνακας 8.10, παρουσιάζει το μέγεθος του δείγματος, όπως επίσης και τις περιπτώσεις που έχουν επισημανθεί ως χαμένες τιμές. Στην περίπτωση του παραδείγματός μας, δεν έχουμε χαμένες τιμές και συνολικά αναλύονται οι απαντήσεις 478 μαθητών.

Case Processing Summary						
	Cases					
	Valid		Missing		Total	
	N	Percent	N	Percent	N	Percent
goals * gender_n	478	100,0%	0	0,0%	478	100,0%

Πίνακας 8.10 Συνοπτικός πίνακας περιπτώσεων.

Στην Εικόνα 8.7 παρουσιάζεται ο πίνακας συνάφειας των δύο μεταβλητών (φύλου και προτεραιότητα στόχων). Η επιλογή μας **percentages – column** υπέδειξε στο SPSS να υπολογιστούν τα ποσοστά των τιμών της μεταβλητής «προτεραιότητα στόχων» για κάθε τιμή της μεταβλητής «φύλο». Έτσι, στις στήλες περιλαμβάνονται τα ποσοστά (σχετικές συχνότητες) των τιμών της προτεραιότητα στόχων σε κάθε φύλο.

goals * gender_n Crosstabulation					
		gender_n		Total	
		1 Αγόρι	2 Κορίτσι		
goals 1 καλό βαθμό	Count	117 ^a	130 ^a	247	
	% within gender_n	51,5%	51,8%	51,7%	
2 δημοτικότητα	Count	50 ^a	91 ^b	141	
	% within gender_n	22,0%	36,3%	29,5%	
3 σπορ	Count	60 ^a	30 ^b	90	
	% within gender_n	26,4%	12,0%	18,8%	
Total	Count	227	251	478	
	% within gender_n	100,0%	100,0%	100,0%	

Each subscript letter denotes a subset of gender_n categories whose column proportions do not differ significantly from each other at the ,05 level.

Προκειμένου να κατανοήσουμε τη σχέση των δύο μεταβλητών, συγκρίνονται σε κάθε γραμμή όλα τα ποσοστά μεταξύ τους (π.χ. το 22% των αγοριών επιλέγουν Δημοτικότητα, ενώ ένα πολύ μεγαλύτερο ποσοστό, 36,3% των κοριτσιών επιλέγουν την ίδια κατηγορία).

Εικόνα 8.7 Πίνακας συνάφειας μεταξύ των δύο μεταβλητών.

Η εικόνα 8.8 παρουσιάζει το αποτέλεσμα του ελέγχου χ^2 της σημαντικότητας της σχέσης μεταξύ των δύο μεταβλητών ή τη γενίκευση της σχέσης από το δείγμα στον πληθυσμό, στην περίπτωση που έχουμε πραγματοποιήσει δειγματοληψία που οδηγεί σε αντιπροσωπευτικό δείγμα.

Οι δύο υποθέσεις που διατυπώνονται είναι:

- H_0 , μηδενική υπόθεση: Το φύλο και η προτεραιότητα στόχων των μαθητών είναι ανεξάρτητες.
- H_A , εναλλακτική υπόθεση: Υπάρχει σχέση ανάμεσα στις δύο μεταβλητές.

Η μηδενική υπόθεση ελέγχεται σε επίπεδο σημαντικότητας $\alpha=0,05$.

Απαιτείται ιδιαίτερη προσοχή στο ότι, για να αξιολογηθεί το αποτέλεσμα του ελέγχου χ^2 , θα πρέπει να διαπιστωθεί κατά πόσο ο έλεγχος είναι έγκυρος. Η εγκυρότητά του εξαρτάται από την ικανοποίηση δύο προϋποθέσεων (βλ. στο κάτω μέρος του Πίνακα, Εικ. 8.8): α) το ποσοστό των κελιών με αναμενόμενη συχνότητα μικρότερη του 5, δεν πρέπει να ξεπερνά το 20%. β) η ελάχιστη αναμενόμενη συχνότητα πρέπει να είναι μεγαλύτερη από 1.

Chi-Square Tests			
	Value	df	Asymptotic Significance (2-sided)
Pearson Chi-Square	21,455 ^a	2	,000
Likelihood Ratio	21,769	2	,000
Linear-by-Linear Association	4,322	1	,038
N of Valid Cases	478		

a. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 42,74.

Ο έλεγχος χ^2 εμφανίζεται στην πρώτη γραμμή του πίνακα. Απορρίπτεται η μηδενική υπόθεση της ανεξαρτησίας των δύο μεταβλητών, αφού $p < 0.001$.

Ικανοποιούνται οι προϋποθέσεις για την εκτέλεση του ελέγχου χ^2 , αφού το ποσοστό των κελιών με αναμενόμενη συχνότητα μικρότερη του 5 είναι 0% και η ελάχιστη αναμενόμενη συχνότητα είναι 42,74.

Εικόνα 8.8 Πίνακας ελέγχου χ^2 .

Εφόσον απορριφθεί η μηδενική υπόθεση (εδώ απορρίφθηκε), τότε και μόνο τότε διενεργούνται (Πίνακας 8.5) όλες οι δυνατές ανά ζεύγη συγκρίσεις και το αποτέλεσμά τους δίνεται με τη μορφή συμβόλων (a, b, c... ως δείκτες των απόλυτων συχνοτήτων). Εάν στους δείκτες δύο συχνοτήτων υπάρχει ένα κοινό γράμμα-σύμβολο, τότε δεν απορρίπτεται η H_0 της ισότητας των δύο ποσοστών σε επίπεδο $\alpha = 0,05$. Για παράδειγμα, στην πρώτη γραμμή δεν έχουμε σημαντική διαφορά, αφού τόσο για τα αγόρια όσο και για τα κορίτσια ο δείκτης είναι ο ίδιος «α». Αντιθέτως, στη δεύτερη και τρίτη γραμμή έχουμε σημαντικές διαφορές, αφού οι δείκτες είναι διαφορετικοί.

Τέλος, στον Πίνακα 8.11 εμφανίζονται οι συντελεστές συσχέτισης που έχουμε ζητήσει. Υπενθυμίζουμε ότι για περιπτώσεις μεταβλητών με περισσότερες από δύο τιμές προτείνεται ο συντελεστής *Cramer's V*. Τιμές των συντελεστών αυτών κοντά στο 0,2 υποδεικνύουν μικρής έντασης συσχετίσεις.

Symmetric Measures			
		Value	Approximate Significance
Nominal by Nominal	Phi	,212	,000
	Cramer's V	,212	,000
N of Valid Cases		478	

Πίνακας 8.11 Πίνακας συντελεστών συνάφειας.

Όσον αφορά την παρουσίαση του αποτελέσματος, παραπάνω ελέγξαμε την υπόθεση της ανεξαρτησίας ανάμεσα στο φύλο και τους στόχους που θέτει ο μαθητής με τη βοήθεια του ελέγχου χ^2 ανεξαρτησίας. Από τα αποτελέσματα του ελέγχου προκύπτει ότι υπάρχει σημαντική και μικρής έντασης σχέση ανάμεσα στο φύλο και τους στόχους του μαθητή, $\chi^2(2, N=478) = 21.455, p < 0.01, V=0,212$. Από τις επιμέρους ανά ζεύγη συγκρίσεις σε κάθε γραμμή με τον έλεγχο z προέκυψε ότι τα κορίτσια προτιμούν σημαντικά περισσότερο από τα αγόρια τη δημοτικότητα (36,3% έναντι 22%), ενώ αντιθέτως τα αγόρια προτιμούν σε σημαντικά υψηλότερο ποσοστό τα σπορ σε σχέση με τα κορίτσια (26,4% έναντι 12%).

8.2 Προβλήματα για εξάσκηση και ανατροφοδότηση

Άσκηση 1

Στο αρχείο «children game.sav» υπάρχουν μετρήσεις για τη σχολική ηλικία των μαθητών και τη σημασία που δίνουν στο παιχνίδι. Ο ερευνητής θέλει να διερευνήσει τη σχέση μεταξύ της σχολικής ηλικίας και της σημασίας που δίνουν οι μαθητές στο παιχνίδι.

- Να κατασκευαστεί κατάλληλος πίνακας σχετικών συχνοτήτων που οδηγεί στη διερεύνηση της σχέσης. Τι συμπεραίνετε;
- Να γίνει ο κατάλληλος έλεγχος της συνάφειας αυτών των μεταβλητών. Αν δεν ισχύουν οι προϋποθέσεις, να συγχωνεύσετε τις τιμές της μεταβλητής «σημασία που δίνουν οι μαθητές στο παιχνίδι».
- Τέλος, να υπολογιστεί ο καταλληλότερος συντελεστής συνάφειας.

Άσκηση 2

Η προηγούμενη έρευνα υποδεικνύει ότι 2 στους 5 μαθητές θεωρούν απαραίτητη τη χρήση της βίας για την επιβολή της τάξης. Να ελέγξετε αυτή την υπόθεση, αν από τους 75 ερωτώμενους μαθητές οι 40 απάντησαν ότι η χρήση της βίας δεν είναι απαραίτητη για την επιβολή της τάξης. Υπόδειξη: Για την εισαγωγή των δεδομένων στο περιβάλλον του SPSS να χρησιμοποιήσετε τη στάθμιση περιπτώσεων (weight cases), για την οποία μπορείτε να ανατρέξετε στο Κεφάλαιο 3.

Άσκηση 3

Να ερευνηθεί η επίδραση του φύλου, του μορφωτικού επιπέδου και της περιοχής που ζουν οι μαθητές στη χρήση του διαδικτύου (από το αρχείο «SS_EKKE_2011_for x2.sav» της ευρωπαϊκής κοινωνικής έρευνας του 2011, οι μεταβλητές «gender», «education», «residential area» και «netuse», αντίστοιχα):

- Να περιγράψτε τις επιδράσεις με κατάλληλους πίνακες.
- Να δοθεί η στατιστική σημαντικότητα των επιδράσεων.
- Να εξηγηθεί το σχήμα της επίδρασης στην περίπτωση σημαντικών επιδράσεων με τη βοήθεια κατάλληλου ελέγχου.
- Να συγκριθούν οι σημαντικές επιδράσεις ως προς το μέγεθός τους με τη βοήθεια κατάλληλου συντελεστή συνάφειας.

Άσκηση 4

Στον πίνακα σύμπτωσης φαίνονται οι απαντήσεις αγοριών και κοριτσιών για τους στόχους που θέτουν.

		Φύλο	
		Αγόρι	Κορίτσι
Στόχοι	Καλοί βαθμοί	117	130
	Δημοτικότητα	50	91
	Σπορ	60	30

- Να περάσετε τα δεδομένα στο SPSS χρησιμοποιώντας τη στάθμιση περιπτώσεων (weight cases).
- Να γίνει ο κατάλληλος έλεγχος της συνάφειας αυτών των μεταβλητών.
- Να κατασκευαστεί κατάλληλος πίνακας σχετικών συχνοτήτων που οδηγεί στη διερεύνηση της σχέσης. Τι συμπεραίνετε;
- Να υπολογιστεί ο καταλληλότερος συντελεστής συνάφειας.

Άσκηση 5

Οι Λεονταρή και Γιαλαμάς (1994) σε μελέτη σχετική με τους παράγοντες που συνδέονται με το άγχος των μαθητών ηλικίας 10-12 ετών διερεύνησαν, μεταξύ άλλων, τη σχέση άγχους και επίδοσης. Ο παρακάτω πίνακας εκφράζει αυτή τη σχέση στο δείγμα των κοριτσιών, τα οποία εμφανίζουν σημαντικά υψηλότερο επίπεδο άγχους από τα αγόρια:

		Επίδοση					Σύνολο
		Χαμηλή	Μάλλον χαμηλή	Μέτρια	Μάλλον Υψηλή	Υψηλή	
Άγχος	Χαμηλό	0	4	3	6	6	19
	Μέτριο	0	3	10	8	9	30
	Υψηλό	4	7	6	7	6	30
Σύνολο		4	14	19	21	21	79

- Ισχύουν οι προϋποθέσεις αναφορικά με τις αναμενόμενες συχνότητες προκειμένου να εκτελεστεί ο έλεγχος χ^2 ανεξαρτησίας; Εξηγήστε την απάντησή σας.
- Αν δεν ισχύουν τελικά οι προϋποθέσεις, πώς θα επιλυθεί το πρόβλημα;

Βιβλιογραφία

- Γιαλαμάς, Β. (2004). *Στατιστικές Τεχνικές και Εφαρμογές στις Επιστήμες της Αγωγής*. Αθήνα: Εκδ. Πατάκη.
- Γναρδέλλης, Χ. (2003). *Εφαρμοσμένη Στατιστική*. Αθήνα: Εκδ. Παπαζήση.
- Howell, C. D. (1997). *Statistical Methods in Psychology*. Fourth Edition. Duxbury Press, by Wadsworth Publishing Co. International Thomson Publishing Inc.

Κεφάλαιο 9

Ανάλυση Διακύμανσης (ANOVA)

Σύνοψη

Στο κεφάλαιο αυτό εξετάζονται τα τρία είδη ελέγχου ανάλυσης διακύμανσης, η ανάλυση διακύμανσης προς έναν παράγοντα – *One-Way ANOVA*, η ανάλυση διακύμανσης επαναλαμβανόμενων μετρήσεων και η ανάλυση διακύμανσης με δύο ή περισσότερους παράγοντες. Παρουσιάζονται και ερμηνεύονται τα αποτελέσματα των αναλύσεων, τόσο των περιγραφικών στατιστικών όσο και των αντίστοιχων ελέγχων. Επίσης, στο πλαίσιο μιας επιστημονικής εργασίας, παρουσιάζεται ο τρόπος συγγραφής των αποτελεσμάτων για κάθε έλεγχο ανάλυσης διακύμανσης.

Προαπαιτούμενη γνώση

Μετασχηματισμός δεδομένων, περιγραφική στατιστική και εισαγωγή στην εκτιμητική και στον έλεγχο των υποθέσεων: Κεφάλαια 2, 4, και 6 του συγγράμματος.

9.1 Εισαγωγή στην ανάλυση διακύμανσης

Η ανάλυση διακύμανσης αποτελεί έναν ισχυρό παραμετρικό στατιστικό έλεγχο, διότι μας δίνει τη δυνατότητα να διερευνήσουμε ταυτόχρονα την επίδραση μίας ή και περισσότερων ανεξάρτητων μεταβλητών στην εξαρτημένη μεταβλητή. Μας επιτρέπει επίσης να υπολογίζουμε όχι μόνο τις κύριες επιδράσεις της καθεμίας ανεξάρτητης μεταβλητής στην εξαρτημένη, αλλά και τις αλληλεπιδράσεις των τιμών των ανεξάρτητων μεταβλητών και της εξαρτημένης μεταβλητής. Επιπλέον, μπορούμε να μετρήσουμε και να αξιολογήσουμε τη συμμετοχή κάθε ανεξάρτητης μεταβλητής στη συνολική μεταβλητότητα/διακύμανση των δεδομένων μας (Γναρδέλλης, 2009).

9.1.1 Ανάλυση διακύμανσης προς έναν παράγοντα – *One-Way ANOVA*

Οι έλεγχοι των μέσων τιμών που είδαμε μέχρι στιγμής περιορίζονται στη σύγκριση δειγματικών μέσων, προερχόμενων από δύο το πολύ τυχαία δείγματα (π.χ. επίδοση στη στατιστική ανδρών και γυναικών ή επίδοση των μαθητών στα Μαθηματικά, πριν και μετά από μια διδακτική παρέμβαση). Η πραγματοποίηση αυτών των ελέγχων έγινε με τη χρήση στατιστικών συναρτήσεων που βασίζονται στην κατανομή *t* του *Student*. Σε περιπτώσεις, όμως, όπου ο έλεγχος γίνεται για περισσότερες από δύο μέσες τιμές, η χρήση της δοκιμασίας του *Student* αποκλείεται και η κύρια εναλλακτική λύση που εφαρμόζεται είναι αυτή της *Ανάλυσης Διακύμανσης* (Analysis of Variance - ANOVA) (Γιαλαμάς, 2004). Για παράδειγμα, όταν θέλουμε να διερευνήσουμε τις διαφορές μεταξύ των μέσων τιμών της επίδοσης των μαθητών στο μάθημα της γλώσσας σε τρεις διαφορετικές ομάδες προσανατολισμού: Ανθρωπιστικών σπουδών, Θετικών σπουδών και Σπουδών Τ.Π.Ε. θα χρησιμοποιήσουμε τον έλεγχο ANOVA.

Σε γενικές γραμμές, το ζήτημα του ελέγχου των μέσων τιμών μπορεί να ιδωθεί μέσα από τη διερεύνηση μιας μηδενικής υπόθεσης, η οποία έχει την εξής μορφή:

$$H_0: \mu_1 = \mu_2 = \mu_3 \dots = \mu_k$$

Μία πρώτη, αλλά εσφαλμένη αντιμετώπιση ενός τέτοιου ελέγχου, μπορεί να γίνει με απλή προέκταση του ελέγχου δύο μέσων τιμών. Θα μπορούσε δηλαδή κάποιος να θεωρήσει ότι ελέγχοντας όλους τους δυνατούς συνδυασμούς των μέσων τιμών ανά δύο, μπορεί να απορρίψει ή να αποδεχθεί την παραπάνω μηδενική υπόθεση. Μία τέτοια όμως προσέγγιση υποκρύπτει σοβαρούς κινδύνους για πιθανή στρέβλωση του επιπέδου σημαντικότητας του ζητούμενου ελέγχου. Με λίγα λόγια, αυξάνεται η πιθανότητα να απορρίψουμε τη μηδενική υπόθεση, ενώ αυτή ισχύει. Πράγματι, όταν πολλές μέσες τιμές ελέγχονται ανά δύο, η πιθανότητα να προκύψει σημαντική διαφορά μεταξύ δύο μέσων τιμών «από τύχη» αυξάνεται αναλόγως του αριθμού των συγκρινόμενων μέσων τιμών.

Εάν, παραδείγματος χάρι, για έλεγχο της ισότητας των μέσων τιμών μ_α και μ_β χρησιμοποιηθεί ως περιοχή απόρριψης η τιμή 0,05 και για τον έλεγχο των μέσων μ_γ και μ_δ χρησιμοποιηθεί επίσης η τιμή 0,05, τότε η συνολική περιοχή αποδοχής και για τους δύο ελέγχους (με βάση τον πολλαπλασιαστικό κανόνα των πιθανοτήτων) είναι $0,95 \times 0,95 = 0,9025$. Δηλαδή, η περιοχή αποδοχής είναι περίπου 0,90 και επομένως το συνολικό επίπεδο σημαντικότητας του ελέγχου είναι το $p = 0,10$ και όχι το επιθυμητό $p \leq 0,05$, που απαιτεί η στατιστική συμπερασματολογία (η πιθανότητα σφάλματος Τύπου Ι). Για έλεγχο έξι ζευγών μέσων τιμών με επίπεδο σημαντικότητας κάθε ελέγχου $p = 0,05$, η συνολική πιθανότητα σφάλματος Τύπου Ι και για τους έξι ελέγχους δεν είναι πλέον η 0,05, αλλά η $1 - (0,95)^6 = 0,265$. Είναι επομένως εμφανές ότι τέτοιου είδους έλεγχοι μέσων τιμών δεν μπορούν να γίνονται με πολλαπλά t-tests, αλλά με μία άλλη διαδικασία ελέγχου, η οποία θα διασφαλίζει τη διατήρηση του επιπέδου σημαντικότητας στα αποδεκτά επίπεδα που ορίζει η στατιστική συμπερασματολογία.

Η ανάλυση διακύμανσης προς έναν παράγοντα είναι η απλούστερη μορφή ελέγχου των διαφορών των μέσων τιμών, περισσότερων από δύο ανεξάρτητων δειγμάτων. Χρησιμοποιεί τα ανεξάρτητα δείγματα που ορίζονται μόνο με μία ανεξάρτητη μεταβλητή (ποιοτική), η οποία στην ορολογία της ανάλυσης διακύμανσης ονομάζεται *Παράγοντας* (Factor). Οι μέσες τιμές υπολογίζονται σε μία ποσοτική μεταβλητή, που ονομάζεται *εξαρτημένη*.

Προκειμένου να εισαγάγουμε τις στατιστικές υποθέσεις στην ανάλυση διακύμανσης, θα χρησιμοποιήσουμε το εξής παράδειγμα: Θέλουμε να εξετάσουμε την επίδραση της κατεύθυνσης σπουδών των μαθητών στο Λύκειο, στην επίδοση στο μάθημα της Βιολογίας Γενικής Παιδείας (ΓΠ). Το ερευνητικό ερώτημα είναι κατά πόσο διαφοροποιείται η επίδοση των μαθητών στη Βιολογία Γενικής Παιδείας ανάλογα με την κατεύθυνση σπουδών των μαθητών στη Γ' Λυκείου.

Οι μεταβλητές είναι:

- Ανεξάρτητη = Η κατεύθυνση σπουδών (κατηγορική μεταβλητή με τρεις τιμές – ομάδες).
- Εξαρτημένη = Η επίδοση (ποσοτική μεταβλητή).

Έχουμε τρεις ομάδες (ανεξάρτητα δείγματα τιμών) επίδοσης στο μάθημα της Βιολογίας:

- 1η ομάδα - 1^ο ανεξάρτητο δείγμα: επίδοση στη Βιολογία ΓΠ των μαθητών της Θεωρητικής κατεύθυνσης.
- 2η ομάδα - 2^ο ανεξάρτητο δείγμα: επίδοση στη Βιολογία ΓΠ των μαθητών της Θετικής κατεύθυνσης.
- 3η ομάδα – 3^ο ανεξάρτητο δείγμα: επίδοση στη Βιολογία ΓΠ των μαθητών της Τεχνολογικής κατεύθυνσης.

Μπορούμε να γράψουμε με σύμβολα:

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

Αυτό σημαίνει πως δεν υπάρχει διαφορά ανάμεσα στους μέσους των τριών πληθυσμών· οι πληθυσμιακοί μέσοι είναι όλοι ίδιοι. Για την εναλλακτική υπόθεση μπορούμε να γράψουμε:

H_A : Τουλάχιστον 2 μέσες τιμές από τις παραπάνω τρεις διαφέρουν.

Παρατηρούμε λοιπόν, ότι δεν δίνεται μία συγκεκριμένη εναλλακτική υπόθεση. Αυτό συμβαίνει επειδή ακριβώς πολλές εναλλακτικές υποθέσεις είναι πιθανές. Μια εναλλακτική, για παράδειγμα, θα ήταν ότι οι δύο πρώτοι μέσοι είναι ίσοι, αλλά ο τρίτος διαφέρει. Μια άλλη θα μπορούσε να είναι ότι και οι τρεις μέσοι διαφέρουν μεταξύ τους:

$$H_A: \mu_1 = \mu_2 \neq \mu_3 \text{ και } H_A: \mu_1 \neq \mu_2 \neq \mu_3$$

Είναι αναμενόμενο ότι οι τιμές θα διαφέρουν μεταξύ τους και οι διαφορές αυτές θα είναι προς δύο κατευθύνσεις.

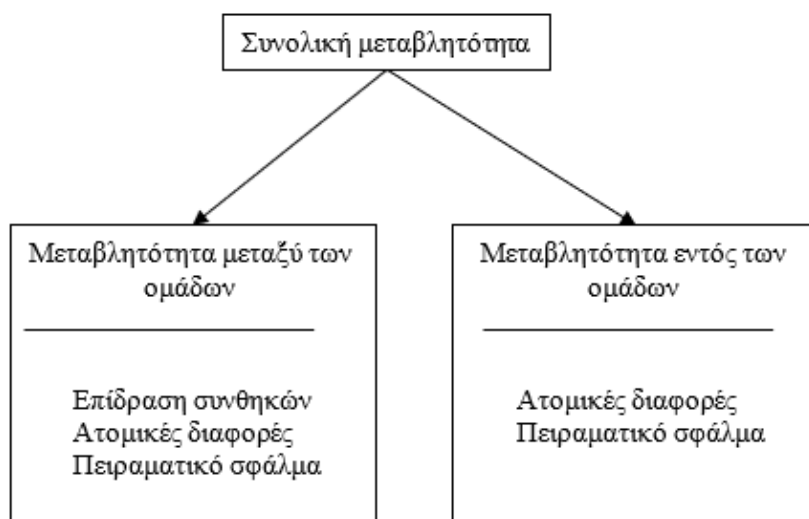
Η πρώτη κατεύθυνση είναι ότι οι μέσες επιδόσεις κάθε ομάδας θα διαφέρουν μεταξύ τους (μεταβλητότητα μεταξύ «between» των ομάδων). Υπάρχουν τρεις δυνατές ερμηνείες για τη μεταβλητότητα μεταξύ των δειγματικών μέσων:

1. Επίδραση ανεξάρτητης μεταβλητής: Είναι πιθανόν οι διαφορετικές κατευθύνσεις σπουδών να έχουν προκαλέσει διαφορές μεταξύ των τριών δειγμάτων.
2. Ατομικές Διαφορές: Τα άτομα που επιλέγονται σε μια ερευνητική διαδικασία ή σε ένα πείραμα έχουν διαφορές στην κοινωνική προέλευση, στην ικανότητα και τη στάση. Κάθε άτομο λοιπόν είναι μοναδικό. Είναι οπότε πιθανό, όταν συγκρίνουμε διαφορετικές ομάδες ατόμων, οι παρατηρούμενες διαφορές τους να οφείλονται ακριβώς στις ατομικές διαφορές.
3. Πειραματικό (τυχαίο) σφάλμα: Υπάρχουν πολλές πηγές πειραματικού σφάλματος. Τα εργαλεία μέτρησης μπορεί να μην είναι αξιόπιστα. Κάποια άτομα μπορεί να μην δώσουν την απαραίτητη προσοχή στο άκουσμα των οδηγιών της έρευνας ή του πειράματος ή μπορεί να μην είναι σταθερός ο τρόπος με τον οποίο δίνονται οι οδηγίες από τον ερευνητή. Αυτό το είδος των ανεξέλεγκτων και απρόβλεπτων διαφορών ονομάζεται πειραματικό ή τυχαίο σφάλμα και μπορεί από μόνο του να προκαλέσει διαφορές μεταξύ των δειγμάτων.

Η δεύτερη κατεύθυνση είναι ότι οι μέσες επιδόσεις μπορεί να διαφέρουν εντός της κάθε ομάδας (μεταβλητότητα εντός (within) των ομάδων). Μέσα σε κάθε δείγμα υπάρχει μεταβλητότητα μεταξύ της επίδοσης. Για παράδειγμα, όλοι οι μαθητές της θεωρητικής κατεύθυνσης δεν θα έχουν σημειώσει την ίδια επίδοση, με λίγα λόγια οι επιδόσεις μεταβάλλονται. Υπάρχουν δύο δυνατές ερμηνείες για τη μεταβλητότητα εντός των δειγματικών μέσων:

1. Ατομικές Διαφορές: Οι τιμές αντιστοιχούν σε διαφορετικά άτομα και για τους λόγους που εξηγήθηκαν παραπάνω αυτές οι τιμές διαφέρουν ανεξάρτητα από το γεγονός ότι τα άτομα αυτά έχουν επιλέξει την ίδια κατεύθυνση σπουδών.
2. Πειραματικό σφάλμα: Η μεταβλητότητα μεταξύ των τιμών μιας ομάδας μπορεί να οφείλεται στο πειραματικό σφάλμα που είναι το ίδιο με αυτό που παρατηρείται στη μεταβλητότητα μεταξύ των ομάδων.

Σημειώνεται δε ότι η μεταβλητότητα στο εσωτερικό μιας ομάδας δεν μπορεί να αποδοθεί στην επίδραση των συνθηκών, αφού όλα τα άτομα της ομάδας μελετώνται κάτω από την ίδια συνθήκη. Η ανάλυση της μεταβλητότητας δίνεται στην Εικόνα 9.1.



Εικόνα 9.1 Διαίρεση της συνολικής μεταβλητότητας σε δύο συνιστώσες.

Αφού ορίσαμε τη συνολική μεταβλητότητα σε δύο βασικές συνιστώσες, αυτό που θα κάνουμε στη συνέχεια είναι να τις συγκρίνουμε. Δηλαδή, θα συγκρίνουμε τη διακύμανση μεταξύ των ομάδων (between groups) και

τη διακύμανση εντός των ομάδων (within groups). Για τη σύγκριση αυτή θα υπολογιστεί το κριτήριο-πηλίκιο F .

Για την περίπτωση ανάλυσης διακύμανσης ανεξάρτητων δειγμάτων το πηλίκιο F έχει την παρακάτω δομή (Γιαλαμάς, 2005):

$$F = \frac{\text{διακύμανση μεταξύ των ομάδων}}{\text{διακύμανση εντός των ομάδων}} \quad (9.1)$$

Όταν εκφράσουμε την κάθε συνιστώσα της μεταβλητότητας με τις πηγές προέλευσής της, το πηλίκιο F παίρνει τη μορφή:

$$F = \frac{\text{επίδραση συνθηκών} + \text{ατομικές διαφορές} + \text{πειραματικό σφάλμα}}{\text{ατομικές διαφορές} + \text{πειραματικό σφάλμα}} \quad (9.2)$$

Παρατηρούμε ότι μεταξύ της μεταβλητότητας εντός των ομάδων και της μεταβλητότητας μεταξύ των ομάδων η διαφορά είναι ο όρος που οφείλεται στην επίδραση των συνθηκών.

Ας εξετάσουμε τι συμβαίνει με το πηλίκιο F για καθεμία από τις δύο υποθέσεις:

1. Αν η H_0 αληθεύει, τότε δεν υπάρχει επίδραση των συνθηκών. Σε αυτή την περίπτωση, αριθμητής και παρονομαστής του πηλίκου εκφράζουν την ίδια διακύμανση:

$$F = \frac{0 + \text{ατομικές διαφορές} + \text{πειραματικό σφάλμα}}{\text{ατομικές διαφορές} + \text{πειραματικό σφάλμα}} \quad (9.3)$$

Όταν η H_0 αληθεύει το πηλίκιο F αναμένεται να είναι ίσο με 1.

1. Αν η H_0 δεν είναι αληθής, τότε υπάρχει επίδραση των συνθηκών, ο αριθμητής στην εξίσωση 9.2 είναι μεγαλύτερος του παρονομαστή και το πηλίκιο F αναμένεται να είναι μεγαλύτερο, από 1.

Η διασπορά (μεταβλητότητα) στο εσωτερικό των k ομάδων μπορεί να εκτιμηθεί από την ποσότητα:

$$SSW = \sum_{j=1}^k \sum_{i=1}^{n_j} (x_{ij} - \bar{x}_j)^2 \quad (9.4)$$

Όπου:

$j = 1, 2, \dots, k$ και

το i διατρέχει το σύνολο των ατόμων n_j της κάθε ομάδας.

$\bar{x}_j$ η μέση τιμή των ατόμων της κάθε ομάδας.

Η ποσότητα αυτή ονομάζεται *άθροισμα τετραγώνων στο εσωτερικό των ομάδων* (Within groups sum of squares).

Η διασπορά μεταξύ των ομάδων (θεωρώντας δηλαδή τις μέσες τιμές των ομάδων ως παρατηρήσεις) εκτιμάται από την ποσότητα

$$SSB = \sum_{j=1}^k n_j (\bar{x}_j - \bar{x})^2 \quad (9.5)$$

όπου:

$j = 1, 2, \dots, k$ και

n_j σύνολο των ατόμων της κάθε ομάδας,

$\bar{x}_j$ η μέση τιμή των ατόμων της κάθε ομάδας,

$\bar{x}$ η συνολική δειγματική μέση τιμή.

Η ανωτέρω ποσότητα ονομάζεται *άθροισμα τετραγώνων μεταξύ των ομάδων* (Between groups sum of squares).

Τέλος η συνολική διασπορά όλων των τιμών γύρω από τη συνολική πληθυσμιακή μέση τιμή εκτιμάται από την ποσότητα

$$TSS = \sum_{j=1}^k \sum_{i=1}^{n_j} (x_{ij} - \bar{x})^2 \quad (9.6)$$

όπου:

$j = 1, 2, \dots, k$ και

n_j σύνολο των ατόμων της κάθε ομάδας,

$\bar{x}$ η συνολική δειγματική μέση τιμή.

Το άθροισμα αυτό ονομάζεται *συνολικό άθροισμα τετραγώνων* (Total sum of squares).

Με απλές αλγεβρικές πράξεις αποδεικνύεται ότι οι τρεις παραπάνω ποσότητες *SSB*, *SSW* και *TSS* συνδέονται μεταξύ τους με την εξής σχέση

$$TSS = SSW + SSB \quad (9.7)$$

Ο έλεγχος της μηδενικής υπόθεσης βασίζεται στη σύγκριση της διασποράς μεταξύ των ομάδων (between groups) σε σχέση με τη διασπορά στο εσωτερικό (εντός) των ομάδων (within groups). Για να απορριφθεί η μηδενική υπόθεση δηλαδή να θεωρήσουμε ότι οι μέσες τιμές των ομάδων δεν είναι όλες ίσες μεταξύ τους, θα πρέπει οι κατανομές στο εσωτερικό κάθε ομάδας να διαφέρουν μεταξύ τους. Θα πρέπει επομένως, η καθεμία ομάδα να διαφοροποιείται σαφώς από τις υπόλοιπες, με μία συνθήκη, η οποία πρακτικά διασφαλίζεται από την ύπαρξη μικρής διασποράς στο εσωτερικό της κάθε ομάδας και μεγάλης διασποράς μεταξύ των ομάδων.

Με βάση επομένως τους ορισμούς των αντίστοιχων διασπορών που δώσαμε προηγουμένως θα πρέπει το κλάσμα

$$\frac{\text{Διασπορά μεταξύ των ομάδων}}{\text{Διασπορά στο εσωτερικό των ομάδων}}$$

να γίνεται όσο το δυνατόν μεγαλύτερο.

Εάν στο παραπάνω κλάσμα ο αριθμητής και ο παρονομαστής αντικατασταθούν αντίστοιχα με το μέσο άθροισμα τετραγώνων μεταξύ των ομάδων, δηλαδή

$$S_B^2 = \frac{SSB}{k-1} \quad (9.8)$$

και το μέσο άθροισμα τετραγώνων στο εσωτερικό των ομάδων, δηλαδή

$$S_W^2 = \frac{SSW}{n-k}, \quad (9.9)$$

τότε η ποσότητα η οποία προκύπτει

$$F = \frac{S_B^2}{S_W^2} \quad (9.10)$$

(υπό την προϋπόθεση ότι ισχύει η μηδενική υπόθεση) ακολουθεί την κατανομή *F*, με $k-1$ και $n-k$ βαθμούς ελευθερίας.

Η κατανομή *F* όμως, ως γνωστόν, είναι μία ασύμμετρη κωδωνοειδής κατανομή (Εικόνα 9.2), η οποία ξεκινάει από το 0 και εκτείνεται ασυμπτωτικά προς τις θετικές τιμές του άξονα τιμών. Η κατανομή αυτή καθορίζεται από το ζεύγος των βαθμών ελευθερίας $df_1=k-1$ και $df_2=n-k$, μεταξύ και εντός των ομάδων αντίστοιχα. Επομένως, η περιοχή της κατανομής που περιλαμβάνει μόλις το 5% των τιμών της και βρίσκεται στο δεξί της άκρο ορίζει την περιοχή απόρριψης $\alpha = 0,05$ για τη μηδενική υπόθεση:

$$H_0: \mu_1 = \mu_2 = \mu_3 \dots = \mu_k$$

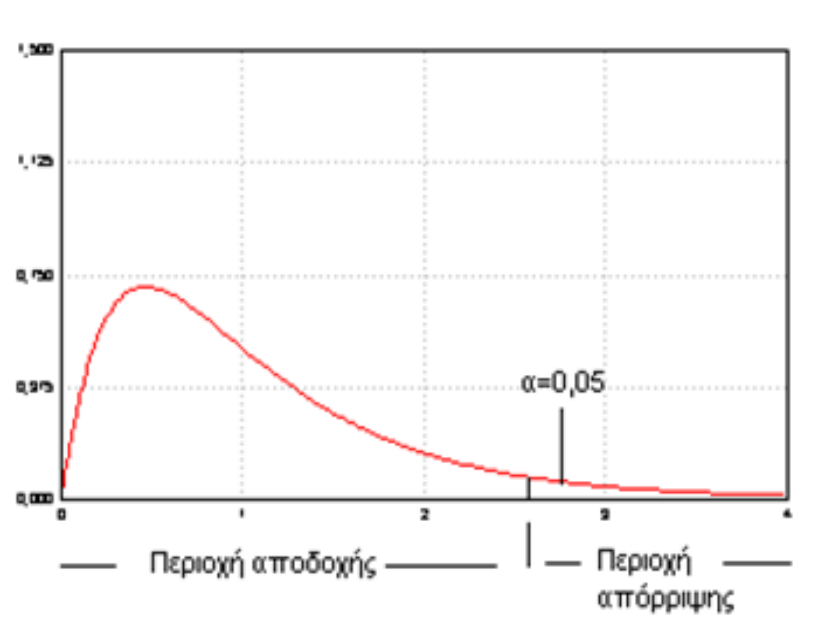
Επομένως, εάν η πιθανότητα να πάρει η κατανομή *F* μία τιμή μεγαλύτερη ή ίση από την αντίστοιχη τιμή F_S (εξίσωση 9.10), που προκύπτει από τα δειγματικά δεδομένα είναι πολύ μικρή, δηλαδή

$$P(F \geq F_S) \leq 0,05,$$

τότε απορρίπτεται η μηδενική υπόθεση και γίνεται αποδεκτή η εναλλακτική H_A .

Ακολουθώντας επομένως τα γνωστά βήματα των επαγωγικών ελέγχων, η παραπάνω διαδικασία μπορεί να συνοψιστεί ως εξής:

- Μηδενική Υπόθεση H_0 : $\mu_1 = \mu_2 = \dots = \mu_k$.
- Εναλλακτική Υπόθεση H_A : Δύο τουλάχιστον από τις $\mu_1, \mu_2, \dots, \mu_k$ διαφέρουν μεταξύ τους.
- Κριτήριο ελέγχου $F = \frac{s_B^2}{s_W^2}$, όπου $s_B^2 = \frac{SSB}{k-1}$ και $s_W^2 = \frac{SSW}{n-k}$, με:
- Για πιθανότητα σφάλματος Τύπου I, και βαθμούς ελευθερίας $k-1$ και $n-k$,
- Απορρίπτεται η H_0 εάν $P(F \geq F_s) \leq 0,05$, όπου F_s η δειγματική τιμή του κριτηρίου F .



Εικόνα 9.2 Κατανομή F με 4 και 25 βαθμούς ελευθερίας.

Προκειμένου η παραπάνω διαδικασία ελέγχου της ισότητας των μέσων τιμών να είναι έγκυρη θα πρέπει να διασφαλίζονται οι εξής τρεις προϋποθέσεις:

- Τα δείγματα να έχουν επιλεγεί τυχαία και ανεξάρτητα το καθένα από τα υπόλοιπα.
- Οι κατανομές στο εσωτερικό των αντίστοιχων πληθυσμών να είναι κατά προσέγγιση κανονικές. Όταν τα δείγματα είναι μικρά ($N < 30$), πρέπει να γίνει ο έλεγχος ότι τα δεδομένα από τις ομάδες προέρχονται από πληθυσμούς με κανονικές κατανομές. Όπως και στην περίπτωση του ελέγχου t , μέτριες αποκλίσεις από την κανονικότητα δεν επιδρούν στο αποτέλεσμα του ελέγχου. Αυτό είναι ιδιαίτερα αληθές όταν τα μεγέθη των δειγμάτων αυξάνουν. Όταν τα δείγματα είναι πολύ μικρά και υπάρχουν σοβαρές αμφιβολίες για την κανονικότητα των πληθυσμών, μια πιθανή εναλλακτική επιλογή είναι η *Kruskal-Wallis* (μη-παραμετρική ανάλυση διακύμανσης) (βλ. Κεφ. 10 του παρόντος).
- Οι διασπορές στο εσωτερικό των πληθυσμών να είναι ίσες μεταξύ τους.

Η εκτέλεση του ελέγχου της ανάλυσης διακύμανσης προς έναν παράγοντα (One-Way ANOVA) στο SPSS θα γίνει με το ακόλουθο παράδειγμα. Θα χρησιμοποιηθεί το αρχείο δεδομένων «biology_by_orient.sav». Το συγκεκριμένο αρχείο έχει δύο μεταβλητές: τη μεταβλητή «κατεύθυνση σπουδών (kat)» (κατηγορική με τρεις τιμές – ομάδες: θεωρητική κατεύθυνση, τεχνολογική κατεύθυνση και θετική κατεύθυνση) και την «επίδοση στο μάθημα της Βιολογίας (epid_biol)» (ποσοτική). Η έρευνα αυτή εξετάζει την επίδραση της κατεύθυνσης σπουδών των μαθητών στο Λύκειο, στην επίδοση στο μάθημα της Βιολογίας Γενικής Παιδείας (ΓΠ), δηλαδή κατά πόσο διαφοροποιείται η επίδοση των μαθητών στη Βιολογία Γενικής Παιδείας ανάλογα με την κατεύθυνση σπουδών των μαθητών στη Γ' Λυκείου. Θα πραγματοποιηθεί έλεγχος ανάλυσης διακύμανσης

προς έναν παράγοντα γιατί έχουμε μία ανεξάρτητη κατηγορική μεταβλητή τριών κατηγοριών (kat) και μία ποσοτική μεταβλητή (epid_biol).

Η διατύπωση των ερευνητικών υποθέσεων γίνεται ως εξής:

- H_0 : θεωρητικής κατεύθυνσης = τεχνολογικής κατεύθυνσης = θετικής κατεύθυνσης
- H_A : τουλάχιστον δύο από τις μέσες τιμές διαφέρουν μεταξύ τους.

Κατ' αρχάς πρέπει να ελέγξουμε αν πληρούνται οι τρεις προϋποθέσεις που αναφέρθηκαν παραπάνω για να χρησιμοποιηθεί η ανάλυση διακύμανσης:

1. Οι μετρήσεις είναι ανεξάρτητες μεταξύ τους και η εξαρτημένη μεταβλητή (βαθμολογία στη βιολογία) είναι ισοδιαστημική.
2. Επειδή και τα τρία δείγματα είναι μικρά ($N < 30$), πρέπει να γίνει ο έλεγχος ότι τα δεδομένα από τις τρεις ομάδες προέρχονται από πληθυσμούς με κανονικές κατανομές. Ο έλεγχος αυτός γίνεται με τη βοήθεια των ελέγχων «Kolmogorov-Smirnov» ή/και «Shapiro-Wilk» από το πλαίσιο διαλόγου «Explore», όπως έχουμε ήδη αναφέρει στο Κεφάλαιο 6. Αυτά τα δύο τεστ ελέγχουν τη μηδενική υπόθεση ότι η πληθυσμιακή κατανομή της μεταβλητής «epid_biol» είναι κανονική (Γναρδέλλης, 2009). Τα αποτελέσματα των ελέγχων αυτών δεν έδειξαν σοβαρή απόκλιση από την κανονικότητα (Εικόνα 9.3).

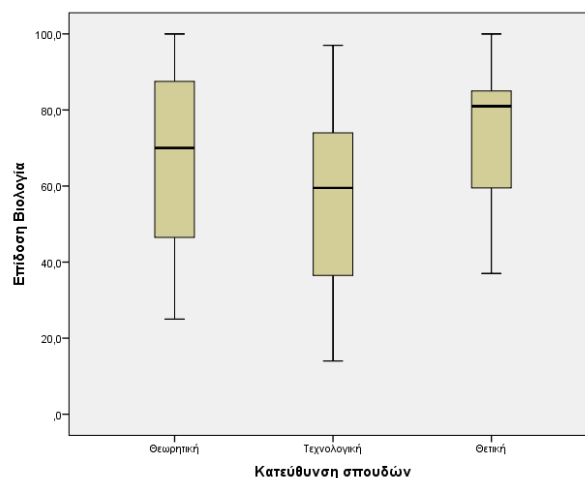
Tests of Normality						
kat		Kolmogorov-Smirnov ^a			Shapiro-Wilk	
		Statistic	df	Sig.	Statistic	df
epid_biol	1 Θεωρητική	,178	19	,116	,920	19
	2 Τεχνολογική	,123	21	,200*	,957	21
	3 Θετική	,167	22	,110	,941	22

*. This is a lower bound of the true significance.

a. Lilliefors Significance Correction

Εικόνα 9.3 Αποτελέσματα των ελέγχων «Kolmogorov-Smirnov» και «Shapiro-Wilk».

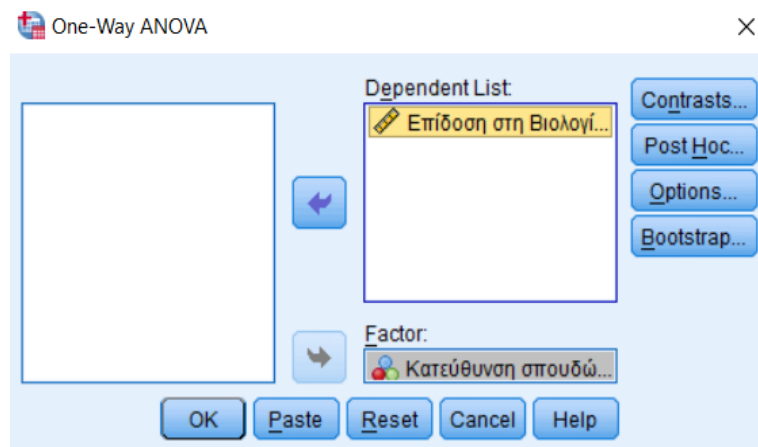
Το ίδιο έδειξαν και τα αντίστοιχα θηκογράμματα (Εικόνα 9.4), λαμβάνοντας υπόψη ότι διαφαίνεται μια συμμετρία στα αντίστοιχα ορθογώνια με άξονα συμμετρίας τη διάμεσο των παρατηρήσεων. Επιπλέον, το θηκογράμμα της ομάδας «θετική» ανέδειξε ότι στη θετική κατεύθυνση σπουδών σημειώθηκε η μεγαλύτερη επίδοση.



Εικόνα 9.4 Θηκογράμματα επίδοσης των τριών κατευθύνσεων σπουδών.

3. Η τρίτη προϋπόθεση, δηλαδή ότι οι διακυμάνσεις των πληθυσμών από τους οποίους προέρχονται τα δείγματα πρέπει να είναι ίσες, πραγματοποιείται με το τεστ του *Levene*. Το τεστ του *Levene* είναι περισσότερο απαλλαγμένο από την επίδραση της έλλειψης κανονικότητας στα δεδομένα σε σχέση με άλλες διαδικασίες. Παραβιάσεις αυτής της προϋπόθεσης είναι αμελητέες, όταν τα μεγέθη των δειγμάτων είναι περίπου ίσα. Η έλλειψη ομοιογένειας αποτελεί πρόβλημα, όταν $S^2_{\text{μέγιστη}} > 10S^2_{\text{ελάχιστη}}$ ή σε μικρότερες διαφορές, όταν τα μεγέθη των δειγμάτων έχουν μεγάλες διαφορές μεταξύ τους. Στην περίπτωση που δεν ισχύει η προϋπόθεση της ομοιογένειας των διακυμάνσεων, η διαδικασία της ανάλυσης διακύμανσης στο SPSS μας παρέχει τα τεστ των *Welch* και *Brown-Forsythe* (Γιαλαμάς, 2005).

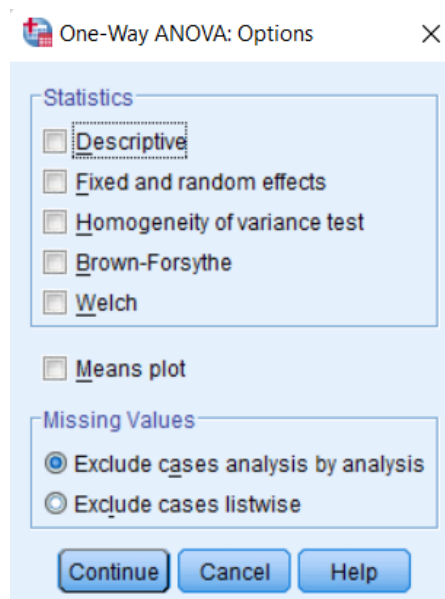
Η διαδικασία εκτέλεσης της ανάλυσης διακύμανσης προς έναν παράγοντα γίνεται στο SPSS ως εξής: Αρχικά επιλέγουμε **Analyze => Compare Means => One-Way ANOVA** και ανοίγει το αντίστοιχο πλαίσιο διαλόγου. Στο πεδίο **Dependent List** μετακινούμε την ποσοτική μεταβλητή που επιλέγουμε από αριστερά, στην προκειμένη περίπτωση τη μεταβλητή «*epid_biol*» (επίδοση στη Βιολογία) (Εικόνα 9.5). Μπορούμε βέβαια να επιλέξουμε και άλλες ποσοτικές μεταβλητές. Στη συνέχεια, επιλέγουμε από αριστερά την κατηγορική μεταβλητή, δηλαδή τη μεταβλητή «*kat*» (κατεύθυνση σπουδών) και τη μετακινούμε στο πεδίο **Factor**.



Εικόνα 9.5 Πλαίσιο διαλόγου «One-Way ANOVA».

Στη συνέχεια, πατώντας το πλήκτρο **Options**, ανοίγει ένα μενού επιλογών (Εικόνα 9.6), όπου μπορούμε να υπολογίσουμε τα περιγραφικά στατιστικά για καθεμία ομάδα (επιλογή **Descriptive**), να εκτιμήσουμε την επίδραση του παράγοντα (κατηγορική μεταβλητή) στις τιμές της εξαρτημένης (ποσοτικής) μεταβλητής (επιλογή **Fixed and random effects**), να εκτελέσουμε τον έλεγχο της προϋπόθεσης της ισότητας των διακυμάνσεων των πληθυσμών με το τεστ του *Levene* (επιλογή **Homogeneity of variance test**) και να υλοποιήσουμε τα εναλλακτικά τεστ των *Brown-Forsythe* και *Welch* (επιλογές **Brown-Forsythe** και **Welch**, αντίστοιχα). Επίσης μπορούμε μέσω αυτού του μενού επιλογών να κατασκευάσουμε τα γραφήματα που απεικονίζουν τις μέσες τιμές της ποσοτικής μεταβλητής για κάθε ομάδα (επιλογή **Means plot**) αλλά και να ορίσουμε τον τρόπο χειρισμού των ελλειπουσών τιμών (επιλογές **Exclude cases analysis by analysis** και **Exclude cases listwise**).

Οι παραπάνω επιλογές εντολών είναι προαιρετικές και αυτές που χρησιμοποιούμε κυρίως είναι η επιλογή **Descriptive** και η επιλογή **Homogeneity of variance test**. Και σε αυτήν την περίπτωση, θα επιλέξουμε τις παραπάνω εντολές.



Εικόνα 9.6 Πλαίσιο διαλόγου «One-Way ANOVA:Options».

Τα αποτελέσματα της ανάλυσης απεικονίζονται στο αρχείο αποτελεσμάτων του SPSS με τρεις πίνακες, όπως φαίνεται παρακάτω (Εικόνες 9.7, 9.8 και 9.9):

Descriptives								
epid_biol								
	N	Mean	Std. Deviation	Std. Error	95% Confidence Interval for Mean		Minimum	Maximum
					Lower Bound	Upper Bound		
1 Θεωρητική	19	69,447	22,3501	5,1275	58,675	80,220	25,0	100,0
2 Τεχνολογική	21	56,095	24,3555	5,3148	45,009	67,182	14,0	97,0
3 Θετική	22	74,295	18,3573	3,9138	66,156	82,435	37,0	100,0
Total	62	66,645	22,7985	2,8954	60,855	72,435	14,0	100,0

Εικόνα 9.7 Περιγραφικά στατιστικά για την «επίδοση στη Βιολογία».

Test of Homogeneity of Variances			
epid_biol			
Levene Statistic	df1	df2	Sig.
1,233	2	59	,299

Εικόνα 9.8 Έλεγχος της ισότητας των διακυμάνσεων των πληθυσμών.

Ο πρώτος Πίνακας (Εικόνα 9.7) εμφανίζει τα περιγραφικά στατιστικά για κάθε ομάδα σχετικά με την επίδοση στο μάθημα της Βιολογίας. Τα στατιστικά αυτά αφορούν το μέγεθος του δείγματος, τις μέσες τιμές, τις τυπικές αποκλίσεις, τα τυπικά σφάλματα, τα 95% διαστήματα εμπιστοσύνης και τις ελάχιστες και μέγιστες τιμές για καθεμία ομάδα κατεύθυνσης σπουδών. Όπως παρατηρούμε, η υψηλότερη μέση τιμή εμφανίζεται στην ομάδα «θετική» (Μ.Τ. = 74,295, Τ.Α. = 18,357). Τη χαμηλότερη μέση τιμή σημειώνει η ομάδα «τεχνολογική» (Μ.Τ. = 56,095, Τ.Α. = 24,356). Με λίγα λόγια, τη μεγαλύτερη επίδοση στο μάθημα της Βιολογίας την είχαν οι μαθητές με θετική κατεύθυνση σπουδών, ενώ τη μικρότερη επίδοση είχαν οι μαθητές με τεχνολογική κατεύθυνση σπουδών.

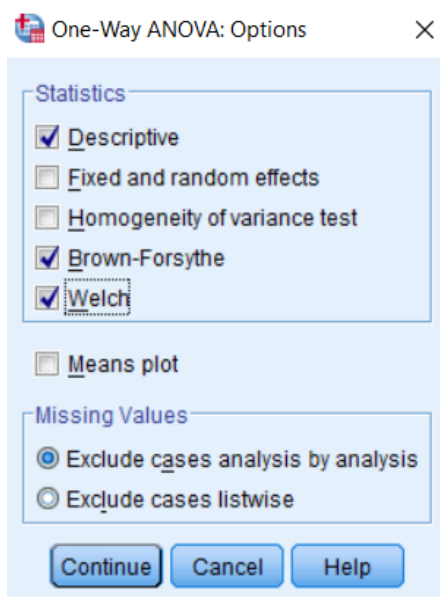
ANOVA					
epid_biol					
	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	3774,107	2	1887,054	3,986	,024
Within Groups	27932,086	59	473,425		
Total	31706,194	61			

Σημαντική διαφορά ανάμεσα στις τρεις κατευθύνσεις σπουδών ($p = 0,037 < 0,05$).

Εικόνα 9.9 «Επίδοση στη Βιολογία».

Ο δεύτερος Πίνακας (Εικόνα 9.8) εμφανίζει τα αποτελέσματα του ελέγχου της προϋπόθεσης της ισότητας των διακυμάνσεων των πληθυσμών, από τους οποίους προέρχονται τα δείγματα, εφαρμόζοντας το τεστ του *Levene*. Η τιμή του στατιστικού είναι 1,233 και η τιμή p (Sig.) είναι 0,299. Άρα, δεν απορρίπτεται η μηδενική υπόθεση και επομένως μπορούμε να ισχυριστούμε ότι οι διακυμάνσεις των πληθυσμών από τους οποίους προέρχονται τα δείγματα είναι ίσες. Αυτό σημαίνει ότι μπορούμε να εφαρμόσουμε την ανάλυση διακύμανσης.

Στην περίπτωση που παραβιάζεται η προϋπόθεση της ισότητας των διακυμάνσεων στους πληθυσμούς, δηλαδή η τιμή p είναι $\leq 0,05$, τότε δεν είμαστε σε θέση να εκτελέσουμε την ανάλυση διακύμανσης. Αντί της ανάλυσης διακύμανσης πρέπει να χρησιμοποιήσουμε τα εναλλακτικά τεστ των *Brown-Forsythe* και *Welch*, για να εκτιμήσουμε την ισότητα των μέσων τιμών των τριών πληθυσμών. Στην περίπτωση αυτή, θα τσεκάρουμε τις επιλογές *Brown-Forsythe* και *Welch* αντίστοιχα, στο πλαίσιο διαλόγου *One-Way ANOVA: Options* (Εικόνα 9.10).



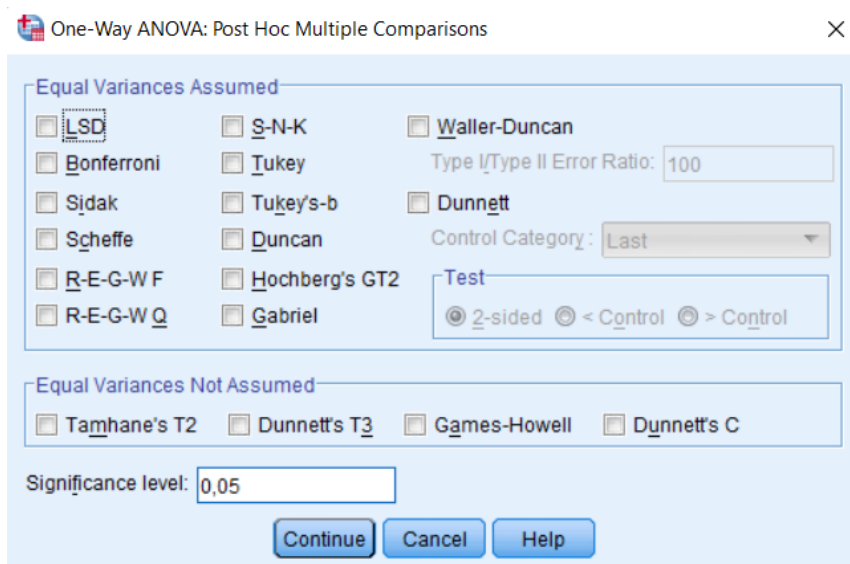
Εικόνα 9.10 Εναλλακτικά τεστ των *Brown-Forsythe* και *Welch*.

Ο τρίτος Πίνακας (Εικόνα 9.9) απεικονίζει τα αποτελέσματα της ανάλυσης διακύμανσης. Στις δύο πρώτες στήλες του Πίνακα αποτυπώνονται το άθροισμα των τετραγώνων των αποκλίσεων που υπάρχει ανάμεσα στις ομάδες, το άθροισμα των τετραγώνων των αποκλίσεων που υπάρχει εντός των ομάδων, καθώς και το συνολικό άθροισμα των τετραγώνων των αποκλίσεων (η συνολική διασπορά όλων των τιμών γύρω από τη συνολική πληθυσμιακή μέση). Η στήλη df μας δείχνει τους βαθμούς ελευθερίας, η στήλη *Mean Square* το μέσο τετράγωνο απόκλισης, η στήλη F την τιμή του στατιστικού και η στήλη *Sig.* το p value (σημαντικότητα του ελέγχου).

Όπως παρατηρούμε, η τιμή p στη στήλη *Sig.* = $0,024 < 0,05$ και επομένως απορρίπτεται η μηδενική υπόθεση. Αυτό σημαίνει ότι υπάρχουν τουλάχιστον δύο από τις τρεις πληθυσμιακές μέσες τιμές που διαφέρουν σημαντικά μεταξύ τους.

Το συγκεκριμένο αποτέλεσμα δεν μας προσδιορίζει όμως ποιες μέσες τιμές διαφέρουν σημαντικά μεταξύ τους ανά ομάδα. Προκειμένου να εξάγουμε αυτό το συμπέρασμα, θα πρέπει να εκτελέσουμε έναν επιπλέον

έλεγχος για να βρούμε τις ομάδες, στις οποίες διαφοροποιούνται οι μέσες τιμές της ποσοτικής μεταβλητής. Ο έλεγχος αυτός ονομάζεται *έλεγχος πολλαπλών συγκρίσεων (multiple comparisons test)* ή *εκ των υστέρων έλεγχος (post hoc test)*. Σε έναν έλεγχο πολλαπλών συγκρίσεων, με τη χρήση των δεδομένων συγκρίνουμε ανά δύο, όλες τις συνθήκες μεταξύ τους. Στο παρόν βιβλίο θα δούμε τους ελέγχους αυτούς μέσω των διαδικασιών του SPSS. Για περισσότερες λεπτομέρειες αναφορικά με τους στατιστικούς ελέγχους των πολλαπλών συγκρίσεων μπορείτε να μελετήσετε το βιβλίο του Γιαλαμά (2005). Για να εκτελέσουμε έναν έλεγχο πολλαπλών συγκρίσεων, από το πλαίσιο διαλόγου **One-Way ANOVA** κάνουμε κλικ στο πλήκτρο **Post Hoc** (Εικόνα 9.11).



Εικόνα 9.11 Πλαίσιο διαλόγου «One-Way ANOVA: Post Hoc Multiple Comparisons».

Στο πλαίσιο διαλόγου αυτό εμφανίζονται τα διαφορετικά είδη ελέγχων πολλαπλών συγκρίσεων που διαθέτει το SPSS. Προσφέρονται έλεγχοι όταν ισχύει η προϋπόθεση της ισότητας των διακυμάνσεων των πληθυσμών, αλλά και όταν έχουμε άνισες διακυμάνσεις.

Υπάρχουν δύο κατηγορίες τέτοιων δοκιμασιών. Η πρώτη κατηγορία περιλαμβάνει διαδικασίες ελέγχου μέσων τιμών ανά δύο (*pairwise multiple comparison tests*) και η δεύτερη περιλαμβάνει διαδικασίες προσδιορισμού ομογενών συνόλων μέσων τιμών, οι οποίες δεν διαφέρουν σημαντικά μεταξύ τους (*multiple range tests*).

Όταν ισχύει η προϋπόθεση της ισότητας των διακυμάνσεων των πληθυσμών, δύο από αυτές τις δοκιμασίες είναι οι πλέον συχνά χρησιμοποιούμενες: η μέθοδος του Bonferroni και η μέθοδος του Tukey. Η πρώτη αποτελεί ένα «*pairwise multiple comparison test*», ενώ η δεύτερη παρέχει και «*pairwise multiple comparison tests*» και «*multiple range tests*» (Γιαλαμάς, 2005; Γναρδέλλης, 2009).

Η μέθοδος (ή διόρθωση) *Bonferroni* είναι και η πιο απλή προς κατανόηση, διορθώνει την περιοχή απόρριψης για όλα τα απλά *t*-test που μπορούν να γίνουν μεταξύ των μέσων τιμών, διαιρώντας τη με τον αριθμό των δυνατών συγκρίσεων. Δηλαδή, θέτει ως περιοχή απόρριψης αυτών των ελέγχων την α/m , όπου α μία επιθυμητή περιοχή απόρριψης (συνήθως $\alpha=0,05$) και m ο αριθμός όλων των δυνατών συγκρίσεων μεταξύ των μέσων τιμών. Αντίστοιχα, υπολογίζει διορθωμένα επίπεδα σημαντικότητας για τους ελέγχους αυτούς. Ο έλεγχος *Bonferroni* χρησιμοποιείται συνήθως όταν ο αριθμός των συγκρινόμενων μέσων τιμών είναι σχετικά μικρός (δεν υπερβαίνει τις πέντε).

Ο έλεγχος *Tukey* είναι ο περισσότερο ασφαλής και αξιόπιστος, γιατί αντιμετωπίζει αποτελεσματικά τον κίνδυνο για σφάλμα Τύπου I. Ο συγκεκριμένος έλεγχος μας επιτρέπει να υπολογίσουμε μια τιμή που καθορίζει την ελάχιστη διαφορά που πρέπει να έχουν δύο δειγματικές μέσες τιμές, ώστε η διαφορά να είναι σημαντική. Η διαφορά αυτή ονομάζεται *HSD* (ευκρινά σημαντική διαφορά). Αν η διαφορά των μέσων είναι μεγαλύτερη του *HSD*, τότε συμπεραίνουμε ότι υπάρχει μια σημαντική διαφορά μεταξύ των δύο συνθηκών. Ο έλεγχος *Tukey* χρησιμοποιείται όταν ο αριθμός των συγκρινόμενων μέσων τιμών είναι σχετικά μεγάλος (υπερβαίνει τις πέντε).

Όταν διαπιστώνουμε άνισες διακυμάνσεις των πληθυσμών, το SPSS παρέχει τέσσερις ελέγχους πολλαπλών συγκρίσεων: *Tamhane's T2*, *Dunnett's T3*, *Games-Howell*, και τον έλεγχο *Dunnett's C*.

Προτείνεται συνήθως η διαδικασία Games–Howell, αφού θεωρείται πιο ισχυρή και ακριβής και όταν τα μεγέθη δειγμάτων είναι άνισα (Field, 2013).

Στην περίπτωση της ανάλυσης των δεδομένων του αρχείου (biology_by_orient.sav), επειδή προέκυψε ότι οι τρεις πληθυσμιακές μέσες τιμές διαφέρουν σημαντικά μεταξύ τους, από το πλαίσιο διαλόγου **One-Way ANOVA: Post Hoc Multiple Comparisons** επιλέγουμε τους ελέγχους **Bonferroni** και **Tukey**. Τα αποτελέσματα αυτών των ελέγχων παρουσιάζονται στην Εικόνα 9.12.

Multiple Comparisons						
Dependent Variable: epid_biol						
Tukey HSD						
(I) kat	(J) kat	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
1 Θεωρητική	2 Τεχνολογική	13,3521	6,8892	,137	-3,211	29,915
	3 Θετική	-4,8481	6,8144	,758	-21,232	11,535
2 Τεχνολογική	1 Θεωρητική	-13,3521	6,8892	,137	-29,915	3,211
	3 Θετική	-18,2002*	6,6380	,022	-34,160	-2,241
3 Θετική	1 Θεωρητική	4,8481	6,8144	,758	-11,535	21,232
	2 Τεχνολογική	18,2002*	6,6380	,022	2,241	34,160

*. The mean difference is significant at the 0.05 level.

epid_biol			
Tukey HSD ^{a,b}			
kat	N	Subset for alpha = 0.05	
		1	2
2 Τεχνολογική	21	56,095	
1 Θεωρητική	19	69,447	69,447
3 Θετική	22		74,295
Sig.		,129	,756

Means for groups in homogeneous subsets are displayed.

a. Uses Harmonic Mean Sample Size = 20,590.

b. The group sizes are unequal. The harmonic mean of the group sizes is used. Type I error levels are not guaranteed.

Εικόνα 9.12 Αποτελέσματα των ελέγχων Bonferroni και Tukey.

Για να μπορέσουμε να ερμηνεύσουμε τα αποτελέσματα των ελέγχων αυτών, ακολουθούμε την εξής διαδικασία: ο πρώτος Πίνακας (pairwise multiple comparison test) των αποτελεσμάτων είναι ουσιαστικά χωρισμένος σε δύο Υποπίνακες, ένας για τον έλεγχο *Tukey* και ένας για τον έλεγχο *Bonferroni*. Κάθε ομάδα συγκρίνεται με τις άλλες δύο, υπολογίζοντας τις διαφορές των μέσων τιμών τους (στήλη Mean Difference (I-J)). Οι υπόλοιπες στήλες αναφέρονται στα τυπικά σφάλματα (Std. Error), στις πιθανότητες απόρριψης της μηδενικής υπόθεσης (Sig.) και στα 95% διαστήματα εμπιστοσύνης (95% Confidence Interval). Οι αστερίσκοι στις διαφορές των μέσων τιμών μας δηλώνουν ποιες από αυτές είναι σημαντικές σε επίπεδο σημαντικότητας 0,05. Στο παράδειγμά μας, παρατηρούμε ότι οι διαφορές οι οποίες είναι σημαντικές τόσο στον έλεγχο *Tukey* όσο και στον έλεγχο *Bonferroni* είναι η θετική με την τεχνολογική κατεύθυνση σπουδών ($p = 0,030$ στον έλεγχο *Tukey* και $p = 0,033$ στον έλεγχο *Bonferroni*). Αυτό σημαίνει ότι η επίδοση των μαθητών της τεχνολογικής κατεύθυνσης στο μάθημα της Βιολογίας Γ.Π. φαίνεται να είναι σημαντικά μικρότερη της επίδοσης των μαθητών της θετικής κατεύθυνσης.

Ο δεύτερος Πίνακας (multiple range tests) απεικονίζει μια άλλη εκδοχή του ελέγχου *Tukey*. Οι ομαδοποιήσεις (δύο ομάδες) μέσω των τιμών γίνονται με βάση τη σημαντικότητα των διαφορών τους. Μέσες τιμές που βρίσκονται στην ίδια ομάδα δεν διαφέρουν μεταξύ τους. Αντίθετα, μέσες τιμές που δεν βρίσκονται μαζί σε καμία ομάδα (δεν αλληλοκαλύπτονται), διαφέρουν σημαντικά μεταξύ τους. Στην Εικόνα 9.12 παρατηρούμε ότι δεν αλληλοκαλύπτονται οι μέσες τιμές της τεχνολογικής με τις μέσες τιμές της θετικής κατεύθυνσης. Άρα αυτές μόνο οι δύο μέσες τιμές διαφέρουν σημαντικά. Το αποτέλεσμα αυτό είναι ίδιο με το αποτέλεσμα των ελέγχων πολλαπλών συγκρίσεων *Tukey* και *Bonferroni*.

Το παραπάνω αποτέλεσμα μπορεί να συνταχθεί ως αναφορά στο πλαίσιο μιας επιστημονικής εργασίας ως εξής: Από την ανάλυση διακύμανσης για τον έλεγχο της επίδοσης των μαθητών Λυκείου στο μάθημα της Βιολογίας Γενικής Παιδείας, ανά κατεύθυνση σπουδών, φαίνεται να υπάρχουν στατιστικά σημαντικές διαφορές [$F_{(2, 59)} = 3,956, p=0,024$]. Επιπλέον, από τον έλεγχο των πολλαπλών συγκρίσεων, χρησιμοποιώντας τα κριτήρια *Tukey* και *Bonferroni*, φαίνεται ότι η επίδοση των μαθητών της τεχνολογικής κατεύθυνσης (Μ.Τ.= 56,10 Τ.Α.= 24,36) είναι σημαντικά μικρότερη της επίδοσης των μαθητών της θετικής (Μ.Τ.=74,29, Τ.Α.=18,36).

9.1.2 Ανάλυση διακύμανσης επαναλαμβανόμενων μετρήσεων

Στο πειραματικό σχέδιο των επαναλαμβανόμενων μετρήσεων υλοποιούνται πολλαπλές μετρήσεις της ίδιας μεταβλητής στα ίδια υποκείμενα της έρευνας (Γναρδέλλης, 2009). Πιο συγκεκριμένα, κάθε ερευνητικό υποκείμενο εκτίθεται σε όλες τις πειραματικές συνθήκες. Η ανάλυση διακύμανσης επαναλαμβανόμενων μετρήσεων χρησιμοποιείται για να καθορίσει τον βαθμό στον οποίο τα δειγματικά δεδομένα επαληθεύουν διαφορές μέσω των τιμών ανάμεσα σε δύο ή περισσότερες πειραματικές συνθήκες (Γιαλαμάς, 2005).

Το βασικό πλεονέκτημα του σχεδίου των επαναλαμβανόμενων μετρήσεων (οδηγούν σε εξαρτημένα δείγματα. Βλ. επίσης Κεφ. 6) είναι ότι είναι ισχυρότερο από το σχέδιο των ανεξάρτητων μετρήσεων, επειδή στο πρώτο (επαναλαμβανόμενων μετρήσεων) δεν υπάρχουν ατομικές διαφορές. Δηλαδή, με αυτόν τον τρόπο ασκείται περισσότερος έλεγχος στις ατομικές μεταβλητές. *Ατομικές μεταβλητές* ονομάζονται τα σταθερά χαρακτηριστικά ενός ατόμου, δηλαδή τα σταθερά χαρακτηριστικά κατά την είσοδο του ατόμου στο πείραμα. Το φύλο, η επίδοση, το μορφωτικό επίπεδο πατέρα/μητέρας, τα χαρακτηριστικά της προσωπικότητας αποτελούν μερικά παραδείγματα ατομικών μεταβλητών. Όταν χρησιμοποιείται το σχέδιο των ανεξάρτητων μετρήσεων, οι ατομικές μεταβλητές μπορεί να δημιουργήσουν πρόβλημα. Ας υποθέσουμε ότι ένας ερευνητής ενδιαφέρεται για μια μέθοδο διδασκαλίας που θα βελτιώνει τη μάθηση. Αν από τύχη η πειραματική ομάδα έχει μαθητές με μεγαλύτερη ικανότητα, η ανώτερη απόδοση αυτής της ομάδας μπορεί να οφείλεται στην ικανότητά της και όχι στη νέα μέθοδο διδασκαλίας. Στην περίπτωση των επαναλαμβανόμενων μετρήσεων, επειδή σε κάθε συνθήκη του πειράματος χρησιμοποιούνται τα ίδια άτομα, ο μέσος βαθμός ικανότητας των υποκειμένων σε όλες τις πειραματικές συνθήκες είναι ο ίδιος (Γιαλαμάς, 2005).

Η ανάλυση διακύμανσης επαναλαμβανόμενων μετρήσεων έχει τα ίδια σύμβολα και υπολογισμούς με το σχέδιο ελέγχου ανεξάρτητων δειγμάτων. Όμως, επειδή απομακρύνονται οι ατομικές διαφορές, αλλάζει η δομή του πηλίκου F .

Η μεταβλητότητα μεταξύ των ομάδων στην περίπτωση ενός ερευνητικού σχεδίου επαναλαμβανόμενων μετρήσεων μπορεί να προκαλείται από:

1. Επίδραση συνθηκών. Οι διαφορετικές συνθήκες διαχείρισης μπορεί να έχουν διαφορετικές επιδράσεις και έτσι να προκαλούν υψηλότερες ή χαμηλότερες τιμές από συνθήκη σε συνθήκη.
2. Πειραματικό σφάλμα. Κάθε φορά που γίνονται μετρήσεις, μπορεί να υπάρξει σφάλμα. Το σφάλμα οφείλεται στην ακρίβεια των εργαλείων μέτρησης, σε ανεξέλεγκτες μεταβολές των εργαστηριακών συνθηκών κ.λπ.

Το πηλίκο στην ανάλυση διακύμανσης επαναλαμβανόμενων μετρήσεων θα έχει την εξής δομή (Γιαλαμάς, 2005):

$$F = \frac{\text{επίδραση συνθηκών} + \text{πειραματικό σφάλμα}}{\text{πειραματικό σφάλμα}} \quad (9.11)$$

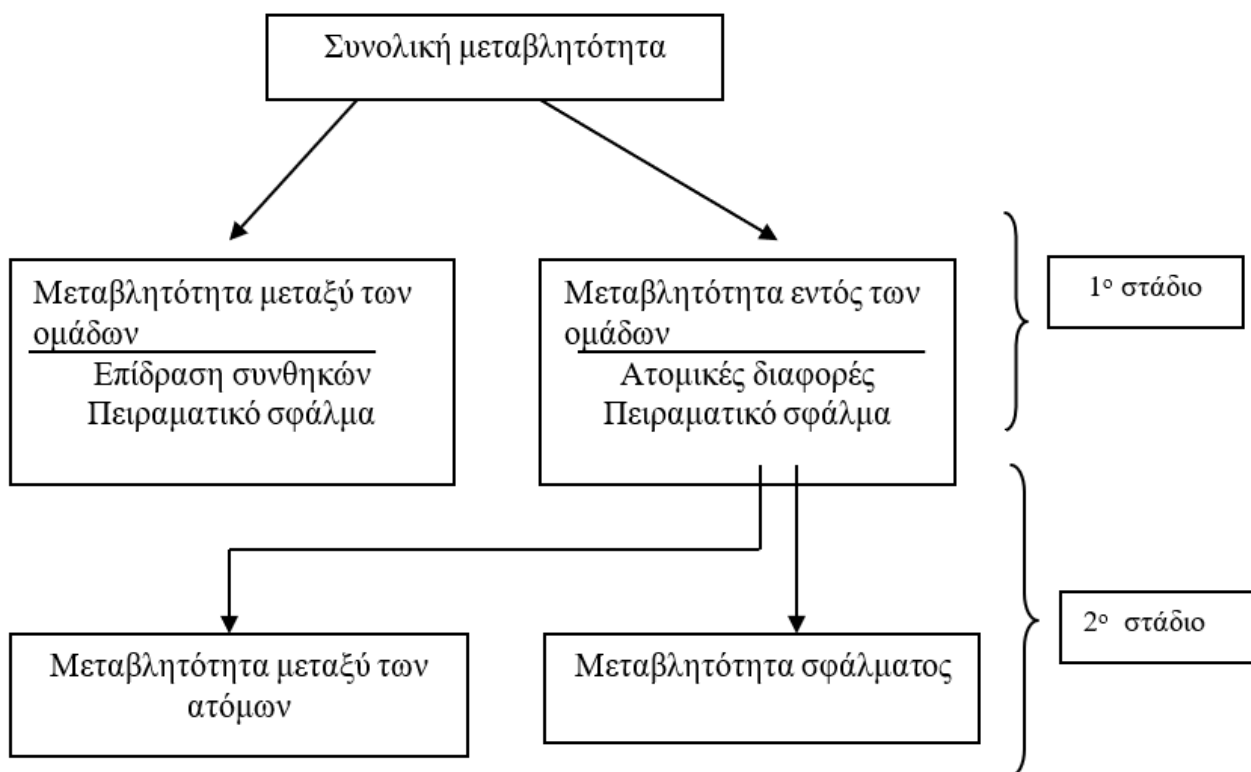
Προκειμένου να πραγματοποιήσουμε μια μέτρηση της διακύμανσης, η οποία οφείλεται στο πειραματικό σφάλμα, η μεταβλητότητα εντός των ομάδων μπορεί να αποδοθεί στα εξής:

- Ατομικές διαφορές. Σε κάθε πειραματική συνθήκη οι τιμές αντιστοιχούν σε διαφορετικά άτομα.
- Πειραματικό σφάλμα. Το μη συστηματικό σφάλμα υπάρχει πάντα ως συνιστώσα της μεταβλητότητας μεταξύ των τιμών.

Προκειμένου να μετρήσουμε το πειραματικό σφάλμα, θα αναλυθεί η εντός των ομάδων μεταβλητότητα που οφείλεται σε διακύμανση από ατομικές διαφορές και σε διακύμανση από πειραματικό σφάλμα.

Η ανάλυση διακύμανσης επαναλαμβανόμενων μετρήσεων γίνεται σε δύο στάδια (Εικόνα 9.13):

1. Στο πρώτο στάδιο η συνολική διακύμανση χωρίζεται σε δύο μέρη: (α) διακύμανση μεταξύ των ομάδων και (β) διακύμανση εντός των ομάδων. Αποτελεί δηλαδή το ίδιο στάδιο με αυτό της ανάλυσης διακύμανσης ανεξάρτητων δειγμάτων.
2. Στο δεύτερο στάδιο πρέπει να διαχωριστούν οι ατομικές διαφορές με το πειραματικό σφάλμα. Αυτό γίνεται με τον υπολογισμό της διακύμανσης μεταξύ των ατόμων και την αφαίρεσή της από τη διακύμανση εντός των ομάδων.



Εικόνα 9.13 Στάδια συνολικής ανάλυσης διακύμανσης επαναλαμβανόμενων μετρήσεων.

Για τη διακύμανση μεταξύ των ομάδων υπολογίζεται το μέσο τετράγωνο της απόκλισης MS_{BG} και για τη διακύμανση του σφάλματος το μέσο τετράγωνο απόκλισης του σφάλματος MS_{error} (Γιαλαμάς, 2005):

$$MS_{BG} = \frac{SS_{BG}}{df_{BG}}, \text{ όπου οι βαθμοί ελευθερίας είναι: } df_{BG} = N - 1$$

Το μέσο τετράγωνο απόκλισης του σφάλματος:

$$MS_{error} = \frac{SS_{error}}{df_{error}}, \text{ όπου οι βαθμοί ελευθερίας είναι: } df_{error} = (N - 1)(k - 1).$$

Όπου,

N = το μέγεθος δείγματος σε κάθε ομάδα και k = το πλήθος των συνθηκών.

Το πηλίκο $F = \frac{MS_{BG}}{MS_{error}}$ ακολουθεί την κατανομή $F_{(N-1, (N-1)(k-1))}$, όταν ισχύει η μηδενική υπόθεση.

Οι βασικές προϋποθέσεις για την εκτέλεση της ανάλυσης διακύμανσης επαναλαμβανόμενων μετρήσεων είναι οι παρακάτω (Γιαλαμάς, 2005):

- Η επιλογή μιας παρατήρησης πρέπει να είναι ανεξάρτητη από την επιλογή των υπόλοιπων παρατηρήσεων.
- Η κατανομή του πληθυσμού από τον οποίο προέρχονται τα δείγματα, σε κάθε πειραματική συνθήκη, πρέπει να είναι κατά προσέγγιση κανονική.
- Οι πληθυσμοί από τους οποίους προέρχονται τα δείγματα, δηλαδή έναν για κάθε πειραματική συνθήκη, πρέπει να έχουν ίσες διακυμάνσεις.
- Οι διακυμάνσεις σε όλες τις διαφορετικές μετρήσεις, καθώς και οι συνδιακυμάνσεις μεταξύ των ζευγών συνθηκών, πρέπει να είναι ίδιες. Πρακτικά, πρέπει να έχουν παρόμοιες σχέσεις ανά ζεύγη (προϋπόθεση της σφαιρικότητας του Mauchly - *Mauchly's test of sphericity*).

Ο έλεγχος του πηλίκου F στην ανάλυση διακύμανσης επαναλαμβανόμενων μετρήσεων είναι ανθεκτικός, όταν η δεύτερη και η τρίτη προϋπόθεση παραβιάζονται. Όμως, η ανθεκτικότητα του ελέγχου δεν ισχύει στην περίπτωση της παραβίασης της πρώτης προϋπόθεσης. Η ομοιογένεια της συνδιακύμανσης αναφέρεται στην προσδοκία ότι η σχετική θέση ενός ατόμου παραμένει σταθερή μέσα σε κάθε πειραματική συνθήκη. Αυτή η προϋπόθεση παραβιάζεται, όταν το αποτέλεσμα μιας παρέμβασης δεν είναι ίδιο για όλα τα άτομα ή όταν υπάρχουν επιδράσεις για μερικά, αλλά όχι για όλα τα υποκείμενα. Ο έλεγχος της τέταρτης προϋπόθεσης γίνεται με τη χρήση του πίνακα διακυμάνσεων-συνδιακυμάνσεων, ο οποίος θα πρέπει να εμφανίζει μία συμμετρία κυκλικής μορφής. Ο έλεγχος αυτός πραγματοποιείται με το τεστ του *Mauchly* (*Mauchly's test of sphericity*). Όταν κατά την εκτέλεση του ελέγχου, η τιμή p είναι μεγαλύτερη από το 0,05, ισχύει η συγκεκριμένη προϋπόθεση. Αν δεν ισχύει αυτή η προϋπόθεση, το SPSS κάνει μια διόρθωση των βαθμών ελευθερίας του στατιστικού F , χρησιμοποιώντας τρεις εκτιμήσεις (epsilon adjustment). Μία εναλλακτική προσέγγιση είναι να χρησιμοποιηθεί η πολυμεταβλητή ανάλυση διακύμανσης, η οποία περιλαμβάνεται στη διαδικασία των επαναλαμβανόμενων μετρήσεων (Repeated Measures) του SPSS (Γβαρδέλλης, 2009).

Η εκτέλεση του ελέγχου της ανάλυσης διακύμανσης επαναλαμβανόμενων μετρήσεων στο SPSS, θα γίνει με το ακόλουθο παράδειγμα. Θα χρησιμοποιηθεί το αρχείο δεδομένων «repeated_measures_3tests.sav». Στο αρχείο αυτό έχουμε τα δεδομένα μιας έρευνας 16 νηπίων, στα οποία πραγματοποιήθηκε μια σειρά διδακτικών παρεμβάσεων για τη μέτρηση με μη συμβατικά μέσα. Στην αρχή, στη μέση και στο τέλος αυτών των διδακτικών παρεμβάσεων μετρήθηκαν τα λάθη στα οποία υποπίπτουν οι μαθητές. Τα αποτελέσματα των λαθών που υποπίπτουν οι μαθητές στις τρεις αυτές μετρήσεις είναι καταχωρισμένα στις μεταβλητές «test1», «test2» και «test3».

Οι ερευνητικές υποθέσεις είναι οι εξής:

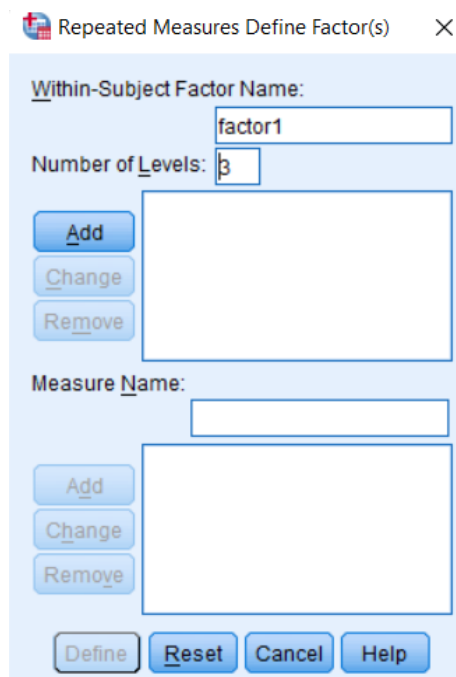
- Μηδενική υπόθεση: Δεν υπάρχει διαφορά στις μέσες τιμές των λαθών που υποπίπτουν οι μαθητές στις τρεις αυτές μετρήσεις.
- Εναλλακτική υπόθεση: Σε δύο τουλάχιστον από αυτές τις μετρήσεις υπάρχει διαφορά στις μέσες τιμές των λαθών που υποπίπτουν οι μαθητές.

Η διαδικασία εκτέλεσης της ανάλυσης διακύμανσης επαναλαμβανόμενων μετρήσεων γίνεται στο SPSS ως εξής: Επιλέγουμε **Analyze => General Linear Model => Repeated Measures** και ανοίγει το πλαίσιο διαλόγου **Repeated Measures Define Factor(s)** (Εικόνα 9.14). Στο πλαίσιο αυτό ορίζουμε έναν ή και περισσότερους παράγοντες, οι οποίοι προκύπτουν από τις επαναλήψεις των μετρήσεων των εξαρτημένων μεταβλητών. Στο παράδειγμά μας, στο πεδίο **Within-Subject Factor Name** πληκτρολογούμε το όνομα του παράγοντα, στην προκειμένη περίπτωση υπάρχει ήδη το όνομα **factor1**, και στο πεδίο **Number of Levels** πληκτρολογούμε τον αριθμό των επιπέδων, δηλαδή των επαναλήψεων του παράγοντα – στην προκειμένη περίπτωση πληκτρολογούμε **3**, όσες δηλαδή είναι οι διαφορετικές μετρήσεις. Έπειτα κάνουμε κλικ στο πλήκτρο **Add**. Το όνομα του παράγοντα και ο αριθμός των επιπέδων, δηλαδή **factor1(3)**, φαίνονται δίπλα δεξιά.

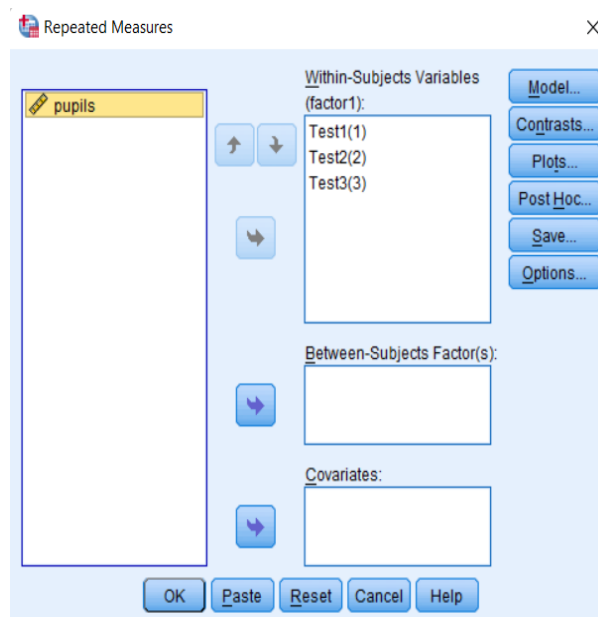
Στη συνέχεια, κάνουμε κλικ στο πλήκτρο **Define** και ανοίγει το πλαίσιο διαλόγου **Repeated Measures** (Εικόνα 9.15). Στο πλαίσιο αυτό επιλέγουμε από αριστερά μία ποσοτική μεταβλητή που αντιστοιχεί σε μία επανάληψη της ίδιας μέτρησης και πατάμε το δεξί βέλος ώστε να την εισάγουμε στη λίστα **Within-Subjects Variables**. Συνεχίζουμε τα ίδια βήματα για καθεμία ποσοτική μεταβλητή. Έχουμε τη δυνατότητα να αλλάξουμε τις θέσεις των μεταβλητών στη λίστα **Within-Subjects Variables**, δηλαδή τις επαναλήψεις μιας μέτρησης, πατώντας το πάνω και κάτω βέλος στη λίστα.

Στο ίδιο πλαίσιο διαλόγου (**Repeated Measures**), πατάμε το πλήκτρο **Options**, δεξιά κάτω, ώστε να ορίσουμε κάποιες επιλογές. Συγκεκριμένα, στη λίστα **Factor(s) and Factor Interactions**, επιλέγουμε τον

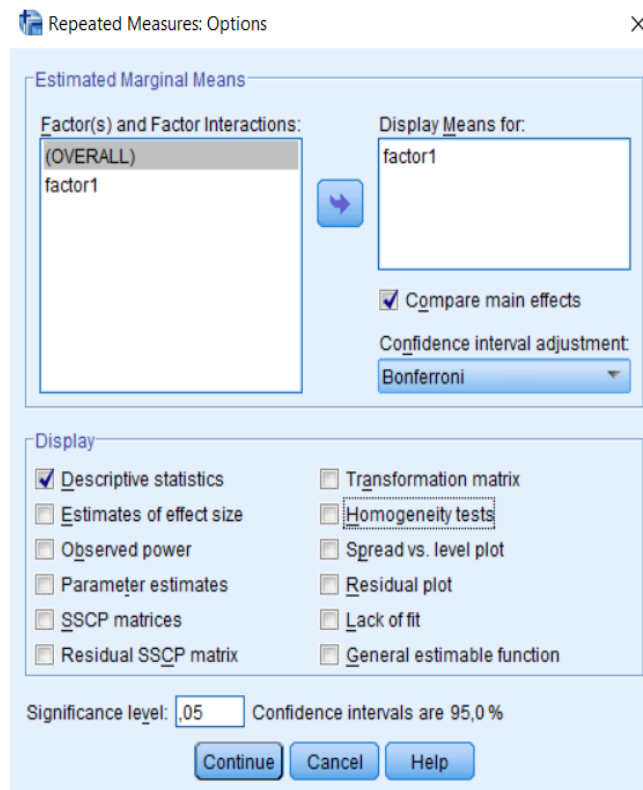
παράγοντα **factor1** και πατάμε το δεξί βέλος ώστε να εισαχθεί στη λίστα **Display Means for**. Στη συνέχεια, τσεκάρουμε την επιλογή **Compare main effects** και στη λίστα επιλογών **Confidence interval adjustment** τσεκάρουμε την επιλογή **Bonferroni**. Τέλος, στη λίστα επιλογών **Display** τσεκάρουμε την επιλογή **Descriptive statistics**. Όλες οι επιλογές του πλαισίου διαλόγου **Repeated Measures: Options** φαίνονται στην Εικόνα 9.16.



Εικόνα 9.14 Πλαίσιο διαλόγου «*Repeated Measures Define Factor(s)*».



Εικόνα 9.15 Πλαίσιο διαλόγου «*Repeated Measures*» του πλήκτρου «*Define*».



Εικόνα 9.16 Πλαίσιο διαλόγου «Repeated Measures» του πλήκτρου «Define».

Τα αποτελέσματα της ανάλυσης απεικονίζονται στο αρχείο αποτελεσμάτων του SPSS με τους εξής Πίνακες (Εικόνες 9.17, 9.18, 9.19 και 9.20):

Ο Πίνακας στην Εικόνα 9.17 με τίτλο «Descriptive Statistics» απεικονίζει τα περιγραφικά στατιστικά των λαθών των νηπίων. Παρατηρούμε ότι στο «test1» έχουμε τη μεγαλύτερη μέση τιμή (M.T. = 6,19, T.A. = 1,64) και στο «test3» τη μικρότερη (M.T. = 1,63, T.A. = 1,36).

Descriptive Statistics			
	Mean	Std. Deviation	N
test1	6,19	1,642	16
test2	2,94	1,237	16
test3	1,63	1,360	16

Εικόνα 9.17 Περιγραφικά στατιστικά.

Ο Πίνακας στην Εικόνα 9.18 με τίτλο «Mauchly's Test of Sphericity» μας δείχνει αν ισχύει η προϋπόθεση της σφαιρικότητας του *Mauchly*. Όταν η τιμή p είναι μεγαλύτερη από το 0,05, ισχύει η συγκεκριμένη προϋπόθεση. Εδώ, η τιμή *Sig.* ισούται με 0,042 και επομένως δεν ισχύει η προϋπόθεση αυτή. Για αυτό τον λόγο, στον πίνακα εμφανίζεται η διόρθωση «epsilon» των βαθμών ελευθερίας του στατιστικού F . Η διόρθωση αυτή γίνεται με τρεις διαφορετικές εκτιμήσεις (Greenhouse-Geisser, Huynh-Feldt και Lower-Bound). Λαμβάνοντας υπόψη ότι η μικρότερη διόρθωση παρέχεται από την Greenhouse-Geisser, αυτή τη διόρθωση θα ακολουθήσουμε για τον έλεγχο της μηδενικής υπόθεσης.

Ο Πίνακας στην Εικόνα 9.19 με τίτλο «Tests of Within-Subjects Effects» μας δείχνει τα αποτελέσματα της μονομεταβλητής ανάλυσης διακύμανσης επαναλαμβανόμενων μετρήσεων. Χρησιμοποιείται η γραμμή του Πίνακα «Greenhouse-Geisser», αφού δεν ισχύει η προϋπόθεση της σφαιρικότητας και επομένως αποδεχόμαστε

τον έλεγχο λαμβάνοντας υπόψη την αντίστοιχη διόρθωση των βαθμών ελευθερίας τόσο στην περιοχή «factor1» όσο και την περιοχή του «Error» (factor1) (σφάλμα).

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
factor1	,636	6,344	2	,042	,733	,793	,500

Tests the null hypothesis that the error covariance matrix of the orthogonalized transformed dependent variables is proportional to an identity matrix.

a. Design: Intercept
Within Subjects Design: factor1

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.

Δεν ισχύει η
προϋπόθεση της
σφαιρικότητας του
Mauchly ($p = 0,042$)

Διορθώσεις
epsilon

Εικόνα 9.18 Έλεγχος σφαιρικότητας του Mauchly.

Tests of Within-Subjects Effects						
Measure: MEASURE_1						
Source		Type III Sum of Squares	df	Mean Square	F	Sig.
factor1	Sphericity Assumed	176,542	2	88,271	68,265	,000
	Greenhouse-Geisser	176,542	1,466	120,434	68,265	,000
	Huynh-Feldt	176,542	1,585	111,369	68,265	,000
	Lower-bound	176,542	1,000	176,542	68,265	,000
Error(factor1)	Sphericity Assumed	38,792	30	1,293		
	Greenhouse-Geisser	38,792	21,988	1,764		
	Huynh-Feldt	38,792	23,778	1,631		
	Lower-bound	38,792	15,000	2,586		

Σημαντική διαφορά ανάμεσα στα τρία test ($p < 0,001$)

Εικόνα 9.19 Πίνακας μονομεταβλητής ανάλυσης διακύμανσης επαναλαμβανόμενων μετρήσεων.

Όπως και στο παράδειγμα της προηγούμενης ενότητας, το συγκεκριμένο αποτέλεσμα δεν μας προσδιορίζει ποιες μέσες τιμές διαφέρουν σημαντικά μεταξύ τους ανά ομάδα. Για να εξάγουμε αυτό το συμπέρασμα, πρέπει να εκτελέσουμε τον έλεγχο πολλαπλών συγκρίσεων (ανά δύο ομάδες), για να βρούμε τις ομάδες, στις οποίες διαφοροποιούνται οι μέσες τιμές της ποσοτικής μεταβλητής.

Το SPSS μας παρέχει τις εξής επιλογές:

- Τον έλεγχο *Bonferroni*, ο οποίος χρησιμοποιείται όταν παραβιάζεται η προϋπόθεση της σφαιρικότητας. Θεωρείται περισσότερο ισχυρός στον έλεγχο του σφάλματος Τύπου I.
- Όταν δεν παραβιάζεται η προϋπόθεση της σφαιρικότητας χρησιμοποιείται ο έλεγχος *Tukey - Least Significant Difference (LSD)*.
- Στην περίπτωση που μειωθεί η ισχύς του ελέγχου, χρησιμοποιείται ο έλεγχος *Sidak*.

Στο παράδειγμά μας, επειδή παραβιάζεται η προϋπόθεση της σφαιρικότητας, χρησιμοποιήθηκε ο έλεγχος *Bonferroni*. Στον παρακάτω Πίνακα (Pairwise Comparisons, Εικόνα 9.20) αναγράφονται τα αποτελέσματα των πολλαπλών συγκρίσεων ανά δύο, των επαναλαμβανόμενων μετρήσεων.

Pairwise Comparisons						
Measure: MEASURE_1						
(I) factor1	(J) factor1	Mean Difference (I-J)	Std. Error	Sig. ^b	95% Confidence Interval for Difference ^b	
					Lower Bound	Upper Bound
1	2	3,250 [*]	,452	,000	2,033	4,467
	3	4,563 [*]	,465	,000	3,309	5,816
2	1	-3,250 [*]	,452	,000	-4,467	-2,033
	3	1,313 [*]	,254	,000	,629	1,996
3	1	-4,563 [*]	,465	,000	-5,816	-3,309
	2	-1,313 [*]	,254	,000	-1,996	-,629

Based on estimated marginal means

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

b. Adjustment for multiple comparisons: Bonferroni.

Σημαντική διαφορά των σφαλμάτων των μαθητών ανάμεσα στη μέτρηση 1 (test) και στις μετρήσεις 2 και 3 ($p < 0,001$)

Εικόνα 9.20 Αποτελέσματα του ελέγχου *Bonferroni*.

Όπως παρατηρούμε, υπάρχει σημαντική διαφορά ανάμεσα στις μετρήσεις των σφαλμάτων των μαθητών. Πιο συγκεκριμένα, η μέση τιμή των σφαλμάτων των μαθητών στην πρώτη μέτρηση (στην αρχή των διδακτικών παρεμβάσεων) είναι σημαντικά μεγαλύτερη από τη μέση τιμή της δεύτερης μέτρησης (στη μέση των διδακτικών παρεμβάσεων) και από τη μέση τιμή της τρίτης μέτρησης (στο τέλος των διδακτικών παρεμβάσεων).

Το παραπάνω αποτέλεσμα μπορεί να συνταχθεί ως αναφορά στο πλαίσιο μιας επιστημονικής εργασίας ως εξής:

Ο έλεγχος σφαιρικότητας του Mauchly αποκάλυψε παραβίαση της σφαιρικότητας ($\chi^2(2)=6,344$, $p=0,042$) και έτσι οι βαθμοί ελευθερίας διορθώθηκαν ακολουθώντας τη διόρθωση ($\varepsilon=0,733$) του Greenhouse-Geisser. Από τον έλεγχο της ανάλυσης της διακύμανσης των επαναλαμβανόμενων μετρήσεων, παρατηρούνται σημαντικές διαφορές στα λάθη, στα οποία υποπίπτουν οι μαθητές σε τουλάχιστον δύο περιπτώσεις ($F_{(1,46, 21,98)} = 68,265$, $p < 0,001$). Επιπλέον, από τον έλεγχο των πολλαπλών συγκρίσεων *Bonferroni* διαπιστώθηκε ότι οι μαθητές υποπίπτουν σε σημαντικά ($p < 0,001$) λιγότερα λάθη στη μέτρηση 3, σε σχέση τόσο με τη μέτρηση 2 όσο και με τη μέτρηση 1. Με λίγα λόγια, οι μαθητές υποπίπτουν σε σημαντικά λιγότερα σφάλματα στη μέτρηση στο τέλος των διδακτικών παρεμβάσεων, σε σχέση με τη μέτρηση στη μέση αλλά και στην αρχή των παρεμβάσεων αυτών. Επομένως, φαίνεται ότι η συγκεκριμένη σειρά των διδακτικών παρεμβάσεων για τη μάθηση με μη συμβατικά μέσα είναι επιτυχής.

9.1.3 Ανάλυση διακύμανσης με δύο ή περισσότερους παράγοντες

Στις προηγούμενες ενότητες εξετάστηκαν πειραματικά σχέδια ή μελέτες συσχέτισης, με μία ανεξάρτητη μεταβλητή. Όμως, στην πραγματική έρευνα πολλές φορές τα πράγματα είναι διαφορετικά. Η ανθρώπινη συμπεριφορά, για παράδειγμα, επηρεάζεται από τον συνδυασμό επιδράσεων δύο ή περισσότερων ανεξάρτητων μεταβλητών ή ομάδων συνθηκών. Ένα παράδειγμα ανάλυσης διακύμανσης με δύο ανεξάρτητες μεταβλητές ή σχέδιο δύο παραγόντων στην ορολογία της ανάλυσης διακύμανσης δίνεται στον Πίνακα 9.1. Στο πείραμα αυτό συγκρίνονται δύο μέθοδοι διδασκαλίας των Μαθηματικών στη Δ' τάξη του Δημοτικού και τρεις διαφορετικοί τρόποι αξιολόγησης στη διάρκεια της χρονιάς. Η εξαρτημένη μεταβλητή είναι η επίδοση στα Μαθηματικά που επιτυγχάνει ο μαθητής στο τέλος της χρονιάς.

Παράγοντας Α (μέθοδος διδασκαλίας)			
		Μέθοδος Α	Μέθοδος Β
Παράγοντας Β (μέθοδος αξιολόγησης)	1 ^η : Εβδομαδιαία εργασία στο σπίτι	Βαθμοί για n = 10 μαθητές που διδάχτηκαν με τη μέθοδο Α και αξιολογούνταν με εβδομαδιαία εργασία στο σπίτι.	Βαθμοί για n = 10 μαθητές που διδάχτηκαν με τη μέθοδο Β και αξιολογούνταν με εβδομαδιαία εργασία στο σπίτι.
	2 ^η : Εβδομαδιαίο διαγώνισμα	Βαθμοί για n = 10 μαθητές που διδάχτηκαν με τη μέθοδο Α και αξιολογούνταν με εβδομαδιαίο διαγώνισμα.	Βαθμοί για n = 10 μαθητές που διδάχτηκαν με τη μέθοδο Β και αξιολογούνταν με εβδομαδιαίο διαγώνισμα.
	3 ^η : Διαγώνισμα σε κάθε ενότητα	Βαθμοί για n = 10 μαθητές που διδάχτηκαν με τη μέθοδο Α και αξιολογούνταν με διαγώνισμα σε κάθε ενότητα.	Βαθμοί για n = 10 μαθητές που διδάχτηκαν με τη μέθοδο Β και αξιολογούνταν με διαγώνισμα σε κάθε ενότητα.

Πίνακας 9.1 Δομή ενός πειράματος με δύο παράγοντες.

Οι διαφορετικές τιμές του κάθε παράγοντα ονομάζονται *επίπεδα του παράγοντα*. Στο παράδειγμα αυτό, ο παράγοντας «μέθοδος διδασκαλίας» έχει δύο επίπεδα (μέθοδος Α και μέθοδος Β) και ο παράγοντας «μέθοδος αξιολόγησης» διαθέτει τρία επίπεδα (εβδομαδιαία διαγωνίσματα, διαγωνίσματα σε κάθε ενότητα και εργασίες στο σπίτι).

Ένα πείραμα με αυτή τη δομή αποκαλείται *δύο – επί - τρία (2 X 3) πειραματικό σχέδιο*. Κάθε κελί του Πίνακα 9.1 αναπαριστά μία συγκεκριμένη συνθήκη του πειράματος. Για παράδειγμα, το κελί της τρίτης γραμμής που βρίσκεται δεξιότερα στον πίνακα περιέχει τους βαθμούς των δέκα μαθητών που διδάχτηκαν με τη μέθοδο Β και αξιολογήθηκαν με εβδομαδιαία διαγωνίσματα κατά τη διάρκεια του έτους. Για το ερευνητικό σχέδιο του Πίνακα υπάρχουν $2 \times 3 = 6$ διαφορετικές συνθήκες. Παρατηρούμε ότι σε κάθε κελί βρίσκεται ένα δείγμα από διαφορετικούς μαθητές. Με αυτόν τον τρόπο, η μελέτη είναι ένα σχέδιο ανεξάρτητων μετρήσεων.

Η ανάλυση διακύμανσης με δύο παράγοντες θα διεξάγει τον έλεγχο για διαφορές μέσων τιμών στο πείραμα. Οι έλεγχοι είναι οι εξής:

- Διαφορές μέσων τιμών ανάμεσα στις δύο διδακτικές μεθόδους.
- Διαφορές μέσων τιμών ανάμεσα στις τρεις μεθόδους αξιολόγησης.
- Οποιαδήποτε άλλη διαφορά μέσης τιμής που μπορεί να προκύψει από συνδυασμούς μιας συγκεκριμένης μεθόδου διδασκαλίας και μιας συγκεκριμένης μεθόδου αξιολόγησης. Για παράδειγμα, μπορεί η μέθοδος διδασκαλίας Β να είναι αποτελεσματικότερη της Α, μόνο στην περίπτωση που οι μαθητές αξιολογούνται με τη μέθοδο των εβδομαδιαίων εργασιών στο σπίτι;

Στη μέθοδο της ανάλυσης διακύμανσης με δύο παράγοντες, εκτός από τον έλεγχο υποθέσεων για τη σημαντικότητα του παράγοντα Α και του παράγοντα Β (κύριες επιδράσεις), θα γίνει και έλεγχος των αλληλεπιδράσεων (ΑxΒ).

Οι κύριες επιδράσεις (main effects) αφορούν τον βαθμό επίδρασης καθεμιάς κατηγορικής μεταβλητής ξεχωριστά στη μεταβλητότητα της μέσης τιμής της εξαρτημένης (ποσοτικής) μεταβλητής. Αναφορικά με το παραπάνω παράδειγμα, ένας βασικός σκοπός του πειράματος είναι να διαπιστωθεί κατά πόσο οι διαφορές στη μέθοδο διδασκαλίας (παράγοντας Α) οδηγούν σε διαφορές στην επίδοση των μαθητών. Για να επιτευχθεί ο παραπάνω στόχος, θα πρέπει να συγκρίνουμε τον μέσο βαθμό των μαθητών της Α μεθόδου με τον μέσο βαθμό των μαθητών της Β μεθόδου.

Για να γίνει καλύτερα κατανοητή η διαδικασία της σύγκρισης των μέσων τιμών, στον παρακάτω Πίνακα 9.2 δίνονται οι υποθετικές μέσες τιμές κάθε πειραματικής συνθήκης (κελιού), όπως και οι υποθετικές μέσες τιμές για κάθε στήλη (μέθοδος διδασκαλίας) και κάθε γραμμή (μέθοδος αξιολόγησης). Τα δεδομένα δείχνουν ότι οι μαθητές της μεθόδου διδασκαλίας Α πήραν έναν μέσο βαθμό $\bar{x}=60$, ενώ οι μαθητές της μεθόδου διδασκαλίας Β πήραν έναν μέσο βαθμό $\bar{x}=70$. Η διαφορά μεταξύ των δύο μέσων τιμών αποτελεί την κύρια επίδραση της μεθόδου διδασκαλίας ή την κύρια επίδραση του παράγοντα Α. Με τον ίδιο τρόπο η κύρια επίδραση της μεθόδου αξιολόγησης εκφράζεται από τις διαφορές των γραμμών του Πίνακα 9.2, δηλαδή των μέσων $\bar{x}=70$, $\bar{x}=65$ και $\bar{x}=60$, που αντιστοιχούν στα τρία επίπεδα του παράγοντα Β.

Από τα παραπάνω γίνεται σαφές ότι οι διαφορές μέσων μεταξύ των γραμμών ή των στηλών απλά περιγράφουν τις κύριες επιδράσεις στο πειραματικό σχέδιο με δύο παράγοντες. Η ύπαρξη διαφορών στις

δειγματικές μέσες τιμές δεν καθιστά τις διαφορές στατιστικά σημαντικές. Κάθε παρατηρούμενη κύρια επίδραση πρέπει να αξιολογηθεί με έναν έλεγχο υπόθεσης, για να διερευνηθεί αν υπάρχουν σημαντικές διαφορές. Αν ο έλεγχος δεν οδηγήσει σε σημαντικές διαφορές, τότε συμπεραίνουμε ότι οι παρατηρούμενες διαφορές των μέσων τιμών αποδίδονται στο δειγματοληπτικό σφάλμα. Στο παράδειγμά μας, ο παράγοντας Α αναφέρεται στη σύγκριση δύο μεθόδων διδασκαλίας. Η μηδενική υπόθεση εκφράζει την άποψη ότι η μέθοδος διδασκαλίας δεν επιδρά στην επίδοση:

$$H_0: \mu_{A1} = \mu_{A2}$$

Προκειμένου να αξιολογηθεί η παραπάνω υπόθεση πρέπει να υπολογιστεί το πηλίκο F με αριθμητή τις διαφορές των παρατηρούμενων μέσων τιμών των δύο μεθόδων διδασκαλίας και παρονομαστή τις διαφορές που αναμένονται λόγω δειγματοληπτικού σφάλματος.

$$F = \frac{\text{διακύμανση (διαφορές) μεταξύ μέσων τιμών του παράγοντα A}}{\text{αναμενόμενη διακύμανση (διαφορές) λόγω δειγματοληπτικού σφάλματος}},$$

δηλαδή

$$F = \frac{\text{διακύμανση (διαφορές) μεταξύ μέσων τιμών των στηλών}}{\text{αναμενόμενη διακύμανση (διαφορές) λόγω δειγματοληπτικού σφάλματος}}.$$

(9.12)

Με τον ίδιο τρόπο, στον παράγοντα Β η μηδενική υπόθεση είναι:

$$H_0: \mu_{B1} = \mu_{B2} = \mu_{B3}$$

Η εναλλακτική υπόθεση αυτού του ελέγχου είναι:

H_A : τουλάχιστον μία μέση τιμή διαφέρει από τις άλλες.

Και σε αυτήν την περίπτωση το F συγκρίνει τις διαφορές ανάμεσα στους μέσους βαθμούς των τριών μεθόδων αξιολόγησης, με τις διαφορές που αναμένονται λόγω της διακύμανσης των τυχαίων δειγμάτων.

$$F = \frac{\text{διακύμανση (διαφορές) μεταξύ μέσων τιμών του παράγοντα B}}{\text{αναμενόμενη διακύμανση (διαφορές) λόγω δειγματοληπτικού σφάλματος}},$$

Δηλαδή

$$F = \frac{\text{διακύμανση (διαφορές) μεταξύ μέσων τιμών των γραμμών}}{\text{αναμενόμενη διακύμανση (διαφορές) λόγω δειγματοληπτικού σφάλματος}}$$

(9.13)

Παράγοντας Α (μέθοδος διδασκαλίας)				
		Μέθοδος Α	Μέθοδος Β	
Παράγοντας Β (μέθοδος αξιολόγησης)	1 ^η : Εβδομαδιαία εργασία στο σπίτι	$\bar{x} = 65$	$\bar{x} = 75$	$\bar{x} = 70$
	2 ^η : Εβδομαδιαίο διαγώνισμα	$\bar{x} = 60$	$\bar{x} = 70$	$\bar{x} = 65$
	3 ^η : Διαγώνισμα σε κάθε ενότητα	$\bar{x} = 55$	$\bar{x} = 65$	$\bar{x} = 60$
		$\bar{x} = 60$	$\bar{x} = 70$	

Πίνακας 9.2 Υποθετικά δεδομένα για το πείραμα του Πίνακα 9.1. Υπάρχουν κύριες επιδράσεις των δύο παραγόντων, αλλά δεν υπάρχει αλληλεπίδραση μεταξύ τους.

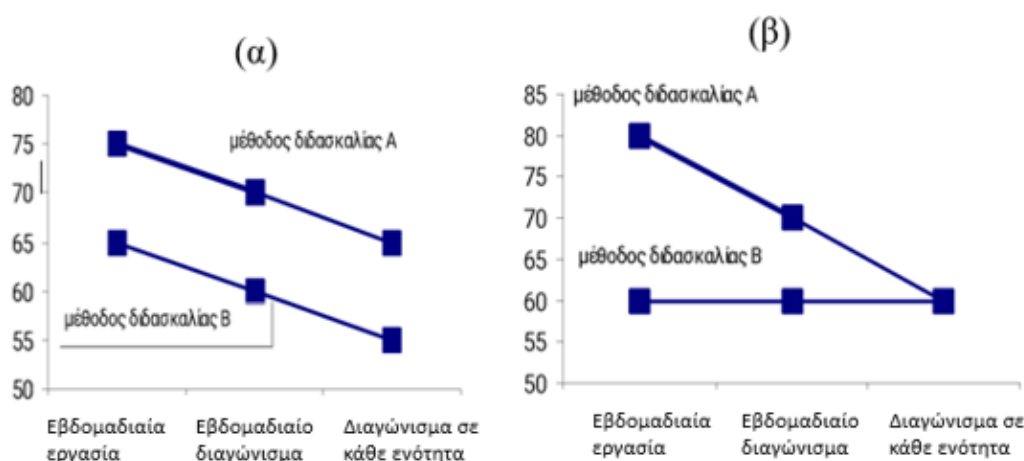
Από τις αλληλεπιδράσεις που αναφέρθηκαν παραπάνω, είναι δυνατόν να διαπιστωθούν άλλες διαφορές μέσων τιμών που οφείλονται σε ιδιαίτερους συνδυασμούς των δύο παραγόντων. Είναι δυνατόν μια πειραματική παρέμβαση να επιδρά ανάλογα με τις περιστάσεις στις οποίες διεξάγεται. Όταν η επίδραση ενός παράγοντα εξαρτάται από την επίδραση ενός άλλου παράγοντα, τότε υπάρχει αλληλεπίδραση (Γιαλαμάς, 2005).

Για να εμπεδωθεί η έννοια της αλληλεπίδρασης, θα εξεταστούν τα δεδομένα του Πίνακα 9.2. Παρατηρούμε την ύπαρξη μιας διαφοράς 10 βαθμών ανάμεσα στις μέσες βαθμολογίες των δύο μεθόδων διδασκαλίας. Αυτή η διαφορά αποτελεί την κύρια επίδραση της μεθόδου διδασκαλίας. Μελετώντας τις διαφορές των μέσων τιμών σε κάθε γραμμή του πίνακα παρατηρούμε παντού την ίδια διαφορά των 10 βαθμών. Συμπεραίνουμε, λοιπόν, πως η επίδραση των δύο μεθόδων διδασκαλίας δεν εξαρτάται από τον τρόπο αξιολόγησης των μαθητών κατά τη διάρκεια του έτους. Αυτό σημαίνει ότι δεν υπάρχει αλληλεπίδραση μεταξύ των δύο παραγόντων. Στη συνέχεια, θα μελετήσουμε τα δεδομένα του Πίνακα 9.3.

Παράγοντας Α (μέθοδος διδασκαλίας)				
		Μέθοδος Α	Μέθοδος Β	
Παράγοντας Β (μέθοδος αξιολόγησης)	1 ^η : Εβδομαδιαία εργασία στο σπίτι	$\bar{x} = 60$	$\bar{x} = 80$	$\bar{x} = 70$
	2 ^η : Εβδομαδιαίο διαγώνισμα	$\bar{x} = 60$	$\bar{x} = 70$	$\bar{x} = 65$
	3 ^η : Διαγώνισμα σε κάθε ενότητα	$\bar{x} = 60$	$\bar{x} = 60$	$\bar{x} = 60$
		$\bar{x} = 60$	$\bar{x} = 70$	

Πίνακας 9.3 Υποθετικά δεδομένα για το πείραμα του Πίνακα 9.1. Υπάρχουν κύριες επιδράσεις των δύο παραγόντων και αλληλεπίδραση μεταξύ τους.

Όπως παρατηρούμε, υπάρχουν ακριβώς οι ίδιες κύριες επιδράσεις, που συναντήσαμε στον προηγούμενο Πίνακα. Αυτό μπορεί να διαπιστωθεί, για παράδειγμα, από τη διαφορά των μέσων τιμών των περιθωρίων των δύο στηλών (εκφράζει την κύρια επίδραση του παράγοντα της «μεθόδου διδασκαλίας») που παραμένει στους 10 βαθμούς. Όμως, στον Πίνακα αυτόν η διαφορά των δέκα βαθμών δεν είναι σταθερή για όλα τα επίπεδα του παράγοντα «μέθοδος αξιολόγησης». Όταν οι μαθητές αξιολογούνται με εβδομαδιαία εργασία στο σπίτι, η διαφορά είναι 20 βαθμοί (80-60), όταν αξιολογούνται με εβδομαδιαίο διαγώνισμα, η διαφορά είναι στους 10 βαθμούς (70-60) και, όταν αξιολογούνται με διαγώνισμα σε κάθε ενότητα, η διαφορά είναι 0 (60-60). Με λίγα λόγια, υπάρχει αλληλεπίδραση μεταξύ των δύο παραγόντων. Η επίδραση της μεθόδου διδασκαλίας εξαρτάται από τη μέθοδο αξιολόγησης των μαθητών. Οι γραφικές παραστάσεις των δεδομένων των Πινάκων 9.3 και 9.4 δίνονται στην Εικόνα 9.21:



Εικόνα 9.21 Το γράφημα (α) αντιστοιχεί στα δεδομένα του Πίνακα 9.2 χωρίς αλληλεπίδραση. Το γράφημα (β) αντιστοιχεί στον Πίνακα 9.3 με αλληλεπίδραση.

Στην περίπτωση που δεν υπάρχει αλληλεπίδραση, οι γραμμές που ενώνουν τις μέσες τιμές των επιπέδων του παράγοντα A, για κάθε επίπεδο του παράγοντα B, είναι παράλληλες. Στο γράφημα (α) της Εικόνας 9.21 οι γραμμές έχουν σταθερή απόσταση μεταξύ τους, δείχνοντας ότι η επίδραση της μεθόδου διδασκαλίας παραμένει σταθερή (10 βαθμοί) για κάθε μέθοδο αξιολόγησης που δοκιμάστηκε. Στο γράφημα (β) η απόσταση μεταξύ των δύο γραμμών ποικίλλει. Με λίγα λόγια, οι γραμμές δεν είναι παράλληλες δείχνοντας με αυτόν τον τρόπο ότι η επίδραση της μεθόδου διδασκαλίας εξαρτάται από τη μέθοδο αξιολόγησης.

Προκειμένου να αξιολογηθεί η ύπαρξη αλληλεπίδρασης, η ανάλυση διακύμανσης αναζητεί διαφορές μέσων τιμών που δεν μπορούν να εξηγηθούν από την ύπαρξη των κύριων επιδράσεων. Για παράδειγμα, στον Πίνακα 9.3 ή στο γράφημα (β) της Εικόνας 9.21 υπάρχουν διαφορές μέσων βαθμών που δεν εξηγούνται από τους 10 βαθμούς διαφοράς της κύριας επίδρασης της μεθόδου διδασκαλίας. Ενώ στο δείγμα που αξιολογείται με διαγώνισμα ανά εβδομάδα, η διαφορά των μεθόδων διδασκαλίας είναι 10 βαθμοί, στο δείγμα που αξιολογήθηκε με διαγώνισμα κατά ενότητα, η διαφορά είναι 0, και στο δείγμα που αξιολογήθηκε με εβδομαδιαία εργασία, η διαφορά ανέρχεται στους 20 βαθμούς. Οι «ιδιαιτερες» αυτές διαφορές αξιολογούνται με ένα πηλίκιο F που έχει την παρακάτω δομή:

$$F = \frac{\text{διακύμανση (διαφορές) που δεν εξηγείται από τις κύριες επιδράσεις}}{\text{αναμενόμενη διακύμανση (διαφορές) λόγω δειγματοληπτικού σφάλματος}} \quad (9.14)$$

Η μηδενική υπόθεση που κρίνεται από το παραπάνω πηλίκιο εκφράζει την απουσία αλληλεπίδρασης:

H_0 : Δεν υπάρχει αλληλεπίδραση μεταξύ του παράγοντα A και του παράγοντα B.

Η εναλλακτική υπόθεση αυτού του ελέγχου είναι:

H_A : Υπάρχει αλληλεπίδραση μεταξύ του παράγοντα A και του παράγοντα B. Η επίδραση του ενός παράγοντα εξαρτάται από τα επίπεδα του άλλου.

Όπως διαπιστώσαμε, στην ανάλυση διακύμανσης με δύο παράγοντες, γίνονται τρεις διαφορετικοί έλεγχοι. Δύο έλεγχοι για τις κύριες επιδράσεις (ένας για κάθε παράγοντα) και ένας έλεγχος για την αλληλεπίδραση. Οι επιδράσεις αυτές ονομάζονται *επίδραση A*, *επίδραση B* και *αλληλεπίδραση A X B*.

Και στους τρεις ελέγχους το πηλίκιο F θα έχει την παρακάτω βασική δομή:

$$F = \frac{\text{διακύμανση μεταξύ των ομάδων}}{\text{διακύμανση εντός των ομάδων}} \quad (9.15)$$

Η διακύμανση μεταξύ των ομάδων (πειραματικών συνθηκών) υποθέτουμε ότι προέρχεται από τρεις πηγές:

1. Επίδραση συνθηκών (παράγοντας A ή παράγοντας B ή αλληλεπίδραση AxB).
2. Ατομικές διαφορές (υπάρχουν διαφορετικά άτομα σε κάθε συνθήκη).
3. Πειραματικό σφάλμα (πάντα υπάρχει η πιθανότητα σφάλματος στις μετρήσεις).

Η διακύμανση εντός των ομάδων δίνει ένα μέτρο της διακύμανσης που αναμένεται τυχαία, αφού με τη βοήθειά της μπορεί να εκτιμηθεί το δειγματοληπτικό σφάλμα. Οι διαφορές εντός των ομάδων δεχόμαστε ότι προέρχονται από:

1. Ατομικές διαφορές.
2. Πειραματικό σφάλμα.

Με τη βοήθεια αυτής της ανάλυσης για κάθε συνιστώσα και τα τρία πηλίκια F θα έχουν την παρακάτω μορφή:

$$F = \frac{\text{επίδραση συνθηκών} + \text{ατομικές διαφορές} + \text{πειραματικό σφάλμα}}{\text{ατομικές διαφορές} + \text{πειραματικό σφάλμα}} \quad (9.17)$$

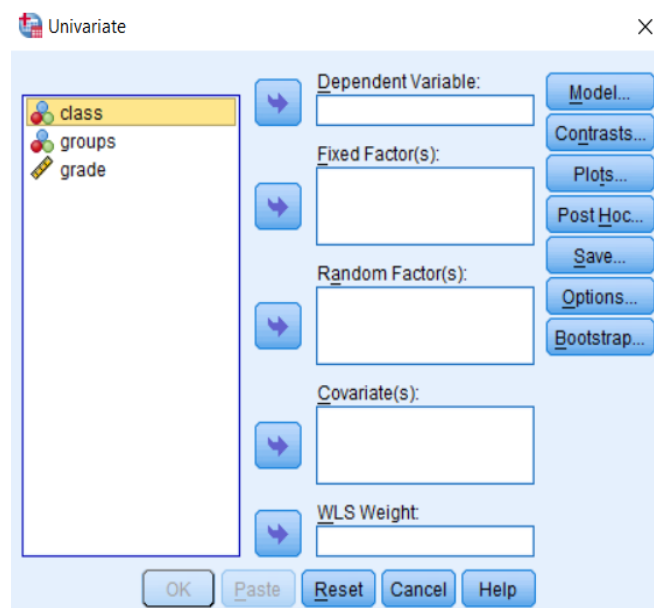
Όπως γνωρίζουμε από τις προηγούμενες μορφές της ανάλυσης διακύμανσης, μια τιμή του F κοντά στο 1 υποδηλώνει την απουσία επίδρασης. Μια τιμή του F πολύ μεγαλύτερη του 1, υποδηλώνει την ύπαρξη επίδρασης συνθηκών.

Η εκτέλεση του ελέγχου της ανάλυσης διακύμανσης με δύο παράγοντες στο SPSS, θα γίνει με το ακόλουθο παράδειγμα (Ρούσσος και Τσαούσης, 2002). Θα χρησιμοποιηθεί το αρχείο δεδομένων «pradeigma_two-way_anova.sav». Στο αρχείο αυτό έχουμε τα δεδομένα από μια έρευνα, η οποία διερευνά την επίλυση λεκτικών μαθηματικών προβλημάτων και θέλει να μελετήσει την επίδραση των οδηγιών στην επίδοση δύο ηλικιακών ομάδων μαθητών Δημοτικού. Εδώ, ο ερευνητής επιλέγει ίδιο πλήθος μαθητών από τη Γ' και τη ΣΤ' Δημοτικού, 42 μαθητές για κάθε τάξη αντίστοιχα και τους χωρίζει σε τρεις ομάδες: την ομάδα ελέγχου-ΟΕ (καμία οδηγία), την πειραματική ομάδα-ΠΟ 1 (πριν το πείραμα τους παρουσίασε παρόμοιο πρόβλημα με αυτό που θα έλυναν στο τεστ) και την πειραματική ομάδα-ΠΟ 2 (το ίδιο με την προηγούμενη ομάδα και επιπλέον τους διευκρίνισε ότι το πρόβλημα που θα έλυναν ήταν παρόμοιο με αυτό που τους παρουσίασε). Ο βασικός σκοπός του ερευνητή είναι να διερευνήσει πως διαμορφώνεται η επίδοση λαμβάνοντας υπόψη τον τύπο ενημέρωσης αναφορικά με το πρόβλημα που έπρεπε να λύσουν οι μαθητές, καθώς επίσης και την ηλικιακή τους ομάδα.

Οι ερευνητικές υποθέσεις είναι οι εξής:

- Μηδενική υπόθεση: Δεν υπάρχουν κύριες επιδράσεις της ηλικιακής ομάδας και του τύπου της πληροφόρησης ούτε αλληλεπιδράσεις αυτών στην επίδοση των μαθητών στην επίλυση του προβλήματος.
- Εναλλακτική υπόθεση: Υπάρχουν κύριες επιδράσεις της ηλικιακής ομάδας και του τύπου της πληροφόρησης αλλά και αλληλεπιδράσεις αυτών στην επίδοση των μαθητών στην επίλυση του προβλήματος.

Η διαδικασία εκτέλεσης της ανάλυσης διακύμανσης με δύο παράγοντες γίνεται στο SPSS ως εξής: Επιλέγουμε **Analyze => General Linear Model => Univariate** και ανοίγει το πλαίσιο διαλόγου **Univariate** (Εικόνα 9.22).

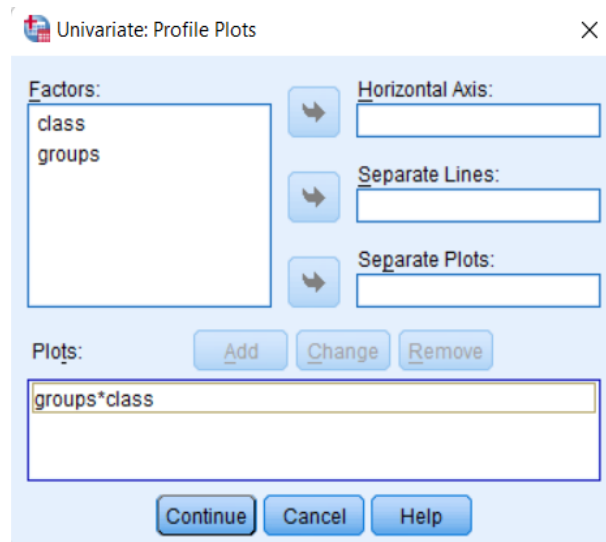


Εικόνα 9.22 Πλαίσιο διαλόγου «Univariate».

Επιλέγουμε από αριστερά την ποσοτική μεταβλητή που θέλουμε και πατάμε το δεξί βέλος ώστε να την εισάγουμε στο πεδίο **Dependent Variable**. Στην προκειμένη περίπτωση επιλέγουμε τη μεταβλητή «grade». Στη

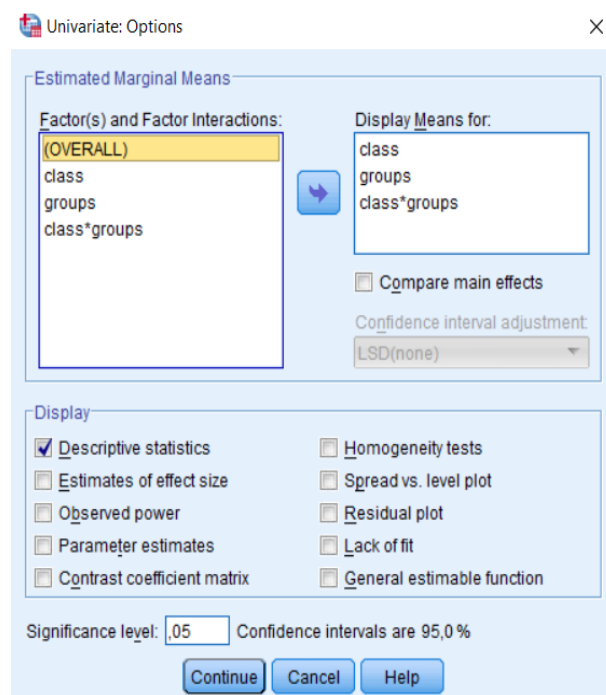
συνέχεια, επιλέγουμε τους παράγοντες της ανάλυσης, δηλαδή τις κατηγορικές μεταβλητές και τις εισάγουμε, πατώντας το δεξί βέλος, στο πεδίο **Fixed Factor(s)**. Στο παράδειγμά μας επιλέγουμε τις μεταβλητές «class» και «groups».

Για να κατασκευάσουμε τα διαγράμματα των προφίλ (profile plots), δηλαδή τα γραφήματα που περιγράφουν τις αλληλεπιδράσεις των μεταβλητών, κάνουμε κλικ στο κουμπί **Plots** δεξιά (Εικόνα 9.23). Από το πεδίο **Factors** επιλέγουμε τον παράγοντα με τις περισσότερες κατηγορίες και τον εισάγουμε στο πεδίο **Horizontal Axis** (οριζόντιος άξονας του διαγράμματος). Εδώ, επιλέγουμε τον παράγοντα «groups». Τον παράγοντα με τις λιγότερες κατηγορίες τον εισάγουμε στο πεδίο **Separate Lines**, ώστε να δημιουργηθούν ξεχωριστές γραμμές. Όταν ορίσουμε τους παράγοντες του διαγράμματος, κάνουμε κλικ στο κουμπί **Add** για να εισαχθούν στη λίστα **Plots** στο κάτω μέρος (**groups*class**).



Εικόνα 9.23 Πλαίσιο διαλόγου «διαγράμματα των προφίλ».

Στη συνέχεια, κάνουμε κλικ στο κουμπί **Options** στο πλαίσιο διαλόγου **Univariate** (Εικόνα 9.22) και εμφανίζεται το πλαίσιο διαλόγου **Options** (Εικόνα 9.24).



Εικόνα 9.24 Πλαίσιο διαλόγου «Univariate».

Στο κάτω μέρος, στην περιοχή **Display**, επιλέγουμε **Descriptive statistics**. Στο πάνω μέρος αριστερά, στην περιοχή **Estimated Marginal Means**, στη λίστα **Factor(s) and Factor Interactions**, επιλέγουμε από αριστερά τους παράγοντες «class», «groups» και «class*groups» και πατάμε το δεξί βέλος για να τους εισάγουμε στη λίστα **Display Means for**. Με αυτόν τον τρόπο θα εμφανιστούν οι μέσες τιμές των ομάδων τόσο από τους παράγοντες όσο και από τους συνδυασμούς των ομάδων των παραγόντων. Ολοκληρώνοντας τις παραπάνω διαδικασίες πατάμε το πλήκτρο **OK**. Τα αποτελέσματα της ανάλυσης είναι τα παρακάτω:

Ο πρώτος πίνακας (Between-Subjects Factors) απλώς μας δείχνει το μέγεθος του δείγματος για όλους τους δυνατούς συνδυασμούς των ομάδων των κατηγορικών μεταβλητών.

Ο δεύτερος πίνακας (Descriptive Statistics) της Εικόνας 9.25 μας δίνει τις μέσες τιμές, τις τυπικές αποκλίσεις και το μέγεθος του δείγματος για όλους τους συνδυασμούς των κατηγοριών τάξης και ομάδας.

Descriptive Statistics				
Dependent Variable: grade				
class	groups	Mean	Std. Deviation	N
ΣΤ' Δημοτ	ΟΕ	62,71	5,455	14
	ΠΟ 1	75,29	6,390	14
	ΠΟ 2	94,14	4,452	14
	Total	77,38	14,127	42
Γ' Δημοτ	ΟΕ	64,43	6,903	14
	ΠΟ 1	62,43	7,643	14
	ΠΟ 2	60,57	8,206	14
	Total	62,48	7,584	42
Total	ΟΕ	63,57	6,167	28
	ΠΟ 1	68,86	9,521	28
	ΠΟ 2	77,36	18,280	28
	Total	69,93	13,535	84

Εικόνα 9.25 Περιγραφικά στατιστικά.

Παρατηρούμε ότι οι μέσες τιμές της επίδοσης της ΣΤ' τάξης του Δημοτικού σε σχέση με την ομάδα ελέγχου (ΟΕ), την πειραματική ομάδα 1 (ΠΟ 1) και την πειραματική ομάδα 2 (ΠΟ 2), διαφέρουν αρκετά μεταξύ τους. Για παράδειγμα, η μέση τιμή της ΟΕ είναι 62,71, ενώ η μέση τιμή της ΠΟ 2 είναι 94,14.

Επιπλέον, η μέση επίδοση της ΣΤ' τάξης του Δημοτικού, ανεξαρτήτως των ομάδων, είναι 77,38, ενώ η μέση επίδοση της Γ' τάξης του Δημοτικού, ανεξαρτήτως των ομάδων, είναι 62,48. Επίσης, συνολικά και για τις δύο τάξεις η μέση τιμή της ΟΕ είναι 63,57, η μέση τιμή της ΠΟ 1 είναι 68,86 και η μέση τιμή της ΠΟ 2 είναι 77,36.

Ο τρίτος Πίνακας (Tests of Between-Subjects Effects) της Εικόνας 9.26 μας δίνει τα αποτελέσματα της ανάλυσης διακύμανσης των παραγόντων τάξης και ομάδων. Στην πρώτη γραμμή του πίνακα με τίτλο «Corrected Model» απεικονίζεται το άθροισμα τετραγώνων που αποδίδεται τόσο στις κύριες επιδράσεις όσο και στις αλληλεπιδράσεις, δηλαδή τη συνολική επίδραση των δύο παραγόντων. Το άθροισμα αυτό χρησιμοποιείται για να γίνει ο συνολικός έλεγχος του υποδείγματος. Η τιμή του στατιστικού του ελέγχου F (λόγος mean square των corrected model και error) ισούται με 53,562. Η τιμή p του ελέγχου είναι $< 0,001$. Αυτό σημαίνει ότι υπάρχουν σημαντικές κύριες επιδράσεις και αλληλεπιδράσεις των παραγόντων στις τιμές της επίδοσης.

Tests of Between-Subjects Effects					
Dependent Variable: grade					
Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	11775,857 ^a	5	2355,171	53,562	,000
Intercept	410760,429	1	410760,429	9341,686	,000
class	4665,190	1	4665,190	106,098	,000
groups	2708,857	2	1354,429	30,803	,000
class * groups	4401,810	2	2200,905	50,054	,000
Error	3429,714	78	43,971		
Total	425966,000	84			
Corrected Total	15205,571	83			

a. R Squared = ,774 (Adjusted R Squared = ,760)

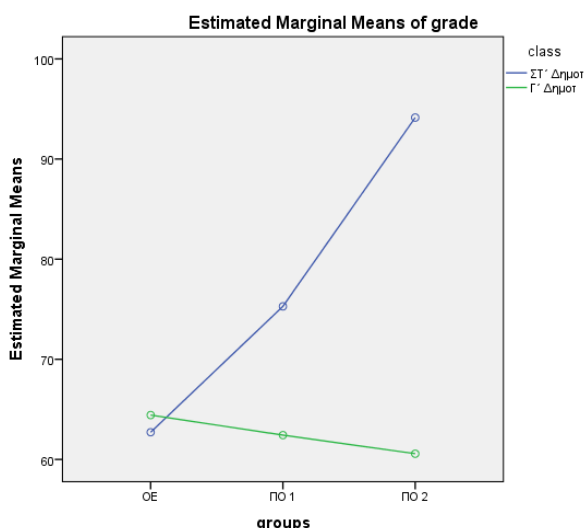
Σημαντικές κύριες επιδράσεις και αλληλεπιδράσεις “τάξης” και “ομάδων ελέγχου” ($p < 0,001$).

Εικόνα 9.26 Κύριες επιδράσεις και αλληλεπιδράσεις του υποδείγματος.

Από τις τιμές p των ελέγχων των κύριων επιδράσεων της τάξης (class) και των ομάδων (groups) αλλά και των αλληλεπιδράσεων (clas*groups) προκύπτει σημαντική επίδραση ($p < 0,001$) και των δύο μεταβλητών ξεχωριστά αλλά και των αλληλεπιδράσεων αυτών στις τιμές της επίδοσης. Στο κάτω μέρος του πίνακα ο συντελεστής «R Squared» εκφράζει το ποσοστό της συνολικής μεταβλητότητας της επίδοσης που εξηγούν οι κύριες επιδράσεις και οι αλληλεπιδράσεις των δύο παραγόντων. Εδώ οι δύο αυτές μεταβλητές εξηγούν το 77,4% της μεταβλητότητας της επίδοσης.

Το διάγραμμα των προφίλ φαίνεται στην Εικόνα 9.27. Οι δύο γραμμές που συνδέουν τα σημεία των μέσων τιμών, οι οποίες τέμνονται, επιβεβαιώνουν την ύπαρξη σημαντικής αλληλεπίδρασης.

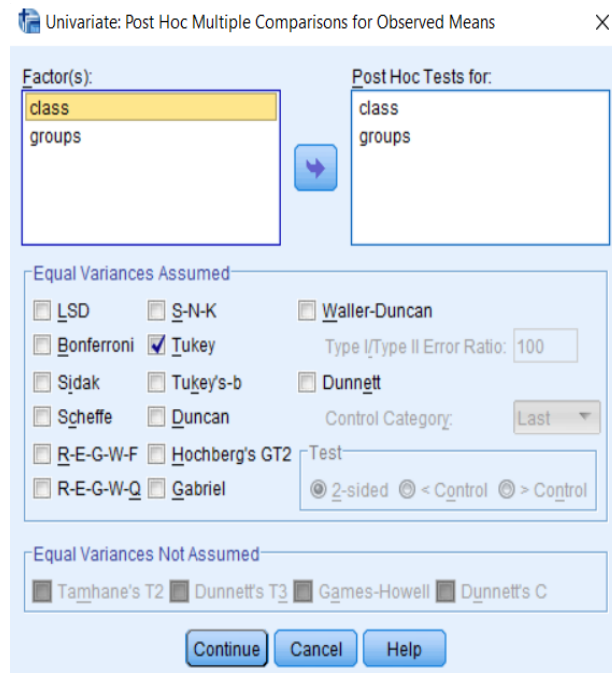
Όπως προέκυψε από τα παραπάνω αποτελέσματα της ανάλυσης διακύμανσης, ως σημαντικότητα των κύριων επιδράσεων του υποδείγματος του παραδείγματός μας μπορούν να συγκριθούν τα ζεύγη των μέσων τιμών των παραγόντων που διαφέρουν σημαντικά μεταξύ τους. Αυτή η διαδικασία γίνεται με τη χρήση των ελέγχων πολλαπλών συγκρίσεων (Post Hoc Tests). Επειδή ο αριθμός των συγκρίσεων είναι μεγάλος, θα χρησιμοποιηθεί ο έλεγχος *Tukey*.



Εικόνα 9.27 Διάγραμμα των προφίλ.

Για να εκτελέσουμε τους ελέγχους πολλαπλών συγκρίσεων, στο πλαίσιο διαλόγου *Univariate* κάνουμε κλικ στο κουμπί *Post Hoc*. Στο πλαίσιο διαλόγου *Univariate: Post Hoc Multiple Comparisons for Observed Means* που εμφανίζεται, επιλέγουμε από αριστερά τις μεταβλητές «class» και «groups» και τις εισάγουμε αριστερά, στη λίστα *Post Hoc Tests for*. Στη συνέχεια, από κάτω, κάνουμε κλικ στην επιλογή *Tukey* (Εικόνα 9.28) και έπειτα κάνουμε κλικ στο κουμπί *Continue*. Τα αποτελέσματα του ελέγχου πολλαπλών συγκρίσεων του *Tukey* απεικονίζονται στον Πίνακα της Εικόνας 9.29. Ο έλεγχος αυτός έχει νόημα μόνο στην περίπτωση που η ποιοτική μεταβλητή έχει τουλάχιστον τρεις τιμές, στην περίπτωση μας εδώ για τον παράγοντα ομάδες. Η σύγκριση των μέσων τιμών της επίδοσης στις δύο τάξεις αρκεί για να διατυπώσουμε τα

συμπεράσματά μας από τη στιγμή που αποκαλύπτονται στατιστικά σημαντικές επιδράσεις του παράγοντα αυτού στην επίδοση.



Εικόνα 9.28 Πλαίσιο διαλόγου «Univariate: Post Hoc Multiple Comparisons for Observed Means».

Multiple Comparisons						
Dependent Variable: grade						
Tukey HSD						
(I) groups	(J) groups	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
OE	ΠΟ 1	-5,29 [*]	1,772	,011	-9,52	-1,05
	ΠΟ 2	-13,79 [*]	1,772	,000	-18,02	-9,55
ΠΟ 1	OE	5,29 [*]	1,772	,011	1,05	9,52
	ΠΟ 2	-8,50 [*]	1,772	,000	-12,73	-4,27
ΠΟ 2	OE	13,79 [*]	1,772	,000	9,55	18,02
	ΠΟ 1	8,50 [*]	1,772	,000	4,27	12,73

Based on observed means.
The error term is Mean Square(Error) = 43,971.
*. The mean difference is significant at the ,05 level.

Εικόνα 9.29 Αποτελέσματα του ελέγχου πολλαπλών συγκρίσεων Tukey.

Από τους ελέγχους των μέσων τιμών των κατηγοριών της μεταβλητής «groups» παρατηρούμε ότι υπάρχει σημαντική διαφοροποίηση της μέσης τιμής της επίδοσης (μεταβλητή «grade») στην OE σε σχέση με την ΠΟ 1 ($p = 0,011$) και την ΠΟ 2 ($p < 0,001$). Η μέση τιμή της επίδοσης της OE είναι κατά 5,29 μονάδες μικρότερη από τη μέση τιμή της ΠΟ 1 και κατά 13,79 μονάδες μικρότερη από τη μέση τιμή της ΠΟ 2. Επίσης, υπάρχει σημαντική διαφοροποίηση της μέσης τιμής της επίδοσης στην ΠΟ 1 σε σχέση με την ΠΟ 2 ($p < 0,001$). Η μέση τιμή της επίδοσης της ΠΟ 2 είναι κατά 8,50 μονάδες μεγαλύτερη από τη μέση τιμή της ΠΟ 1.

Το παραπάνω αποτέλεσμα μπορεί να συνταχθεί ως αναφορά στο πλαίσιο μιας επιστημονικής εργασίας ως εξής: από τον έλεγχο της ανάλυσης της διακύμανσης με δύο παράγοντες παρατηρούνται τα παρακάτω:

- Υπάρχει στατιστικά σημαντική κύρια επίδραση της ηλικιακής ομάδας στην επίδοση των μαθητών [$F_{(1,78)} = 106,098, p = 0,001$].
- Υπάρχει στατιστικά σημαντική κύρια επίδραση της διαφορετικής ενημέρωσης σχετικά με το προς επίλυση πρόβλημα (πειραματικές συνθήκες) στην επίδοση των μαθητών [$F_{(2,78)} = 30,803, p = 0,001$].
- Υπάρχει στατιστικά σημαντική αλληλεπίδραση μεταξύ της διαφορετικής ενημέρωσης σχετικά με τη λύση του προβλήματος και της ηλικιακής ομάδας στην επίδοση των μαθητών [$F_{(2,78)} = 50,054, p = 0,001$].

Η αλληλεπίδραση αυτή δείχνει ότι η ενημέρωση για την ομοιότητα του προς επίλυση προβλήματος με κάποιο προηγούμενο πρόβλημα, επηρεάζει διαφορετικά την επίδοση στις δύο ηλικιακές ομάδες. Αναλυτικότερα, η επίδοση των μαθητών ήταν σχεδόν ίδια σε όλες τις ομάδες της Γ' Δημοτικού, δηλαδή ήταν ίδια και στην ομάδα χωρίς ενημέρωση (Μ.Τ. = 64,63, Τ.Α. = 6,90) και στις άλλες δύο με μικρότερη (Μ.Τ. = 62,43, Τ.Α. = 7,64) και μεγαλύτερη ενημέρωση (Μ.Τ. = 60,57, Τ.Α. = 7,58), αντίστοιχα. Ωστόσο, η επίδοση των μαθητών της ΣΤ' Δημοτικού χωρίς ενημέρωση ήταν σημαντικά μικρότερη (Μ.Τ. = 62,71, Τ.Α. = 5,46) από την επίδοση των μαθητών της ίδιας τάξης με μικρότερη (Μ.Τ. = 75,29, Τ.Α. = 6,39) και μεγαλύτερη ενημέρωση (Μ.Τ. = 94,14, Τ.Α. = 4,45), αντίστοιχα.

9.2 Προβλήματα για εξάσκηση και ανατροφοδότηση

Άσκηση 1

Να χρησιμοποιηθεί το αρχείο «αυτοαντίληψη ANOVA.sav». Αφού δημιουργήσετε τις τρεις καινούριες μεταβλητές («cognitive», «physical» και «social»), υπολογίζοντας τον μέσο όρο των απαντήσεων των ερωτώμενων, σε κάθε ομάδα που χρησιμοποιήθηκε για καθεμία από αυτές τις δηλώσεις (βλ. παρακάτω την προτεινομένη σύνθεση τους), να απαντηθούν τα παρακάτω ερωτήματα:

- Ποιος είναι ο μέσος όρος της ακαδημαϊκής αυτοαντίληψης και για τις τρεις διαστάσεις (γνωστική, σωματική και κοινωνική) που προκύπτει ανάλογα με την περιοχή κατοικίας;
- Δημιουργήστε αντίστοιχο γράφημα για το παραπάνω.
- Μπορούμε να ισχυριστούμε ότι η ακαδημαϊκή αυτοαντίληψη των μαθητών και στις τρεις διαστάσεις επηρεάζεται από την περιοχή κατοικίας;
- Κατά πόσο το αποτέλεσμα αυτό είναι στατιστικά σημαντικό;
- Αν είναι στατιστικά σημαντικό, μεταξύ ποιων κατηγοριών προσδιορίζεται η στατιστικά σημαντική διαφορά;
- Τεκμηριώστε την απάντησή σας για τη χρήση ή όχι παραμετρικού τεστ.

Μεταβλητές που θα χρησιμοποιήσετε:

- «region» = περιοχή κατοικίας.
- «cognitive» = γνωστική αυτοαντίληψη (Μεταβλητές: pcsc_1, pcsc_5, pcsc_9, pcsc_13, pcsc_17, pcsc_21, pcsc_25).
- «physical» = σωματική αυτοαντίληψη (Μεταβλητές: pcsc_3, pcsc_7, pcsc_15, pcsc_11, pcsc_19, pcsc_23, pcsc_27).
- «social» = κοινωνική αυτοαντίληψη (Μεταβλητές: pcsc_2, pcsc_6, pcsc_10, pcsc_14, pcsc_18, pcsc_22, pcsc_26).

Άσκηση 2

Έστω ότι μας ενδιαφέρει να ελέγξουμε κατά πόσο διαφοροποιείται η άποψη των κατοίκων για τον βαθμό ενημέρωσης που παρέχουν τα σχολεία σχετικά με τα περιβαλλοντικά προβλήματα (μεταβλητή: «q2») ανάλογα με την περιοχή κατοικίας τους (μεταβλητή: «q1»). Αρχείο «Educ_environment.sav».

Ερευνητικό ερώτημα: Διαφοροποιείται η άποψη των κατοίκων για τον βαθμό ενημέρωσης που παρέχουν τα σχολεία σχετικά με τα περιβαλλοντικά προβλήματα ανάλογα με την περιοχή κατοικίας τους;

- Ποιες είναι οι μεταβλητές; Ποια είναι ποιοτική και ποια ποσοτική; Ποια είναι η «ανεξάρτητη-independent ή παράγοντας-factor» και ποια είναι η «εξαρτημένη – dependent»;
- Ποιες είναι οι αντίστοιχες ερευνητικές υποθέσεις: H0 και H1;
- Προσδιορίστε μέτρα θέσης και διασποράς των τριών δειγμάτων με τη διαδικασία «Explore». Τι παρατηρείτε; Αναδεικνύεται κάποια σχέση, μεταξύ των δύο μεταβλητών;
- Στην ίδια διαδικασία ελέγξτε την κανονικότητα των αντίστοιχων πληθυσμών των δειγμάτων. Τι παρατηρείτε;
- Ελέγξτε την ισότητα των διακυμάνσεων.
- Ελέγξτε τη μηδενική υπόθεση.
- Εφαρμόστε ελέγχους πολλαπλών συγκρίσεων και παρουσιάστε το συμπέρασμα.

Άσκηση 3

Στο αρχείο δεδομένων «gss.sav» να μελετηθεί η επίδραση του φύλου (μεταβλητή «Sex»), καθώς και του επίπεδου σπουδών (μεταβλητή «Degree») στη μεταβλητή *ώρες εργασίας την περασμένη εβδομάδα* (μεταβλητή «hrs»).

- Ποιες είναι οι ανεξάρτητες μεταβλητές (παράγοντες) και ποια η εξαρτημένη;
- Να διατυπώσετε τις υποθέσεις για τις κύριες επιδράσεις και για τις αλληλεπιδράσεις.
- Ελέγξτε τη μηδενική υπόθεση για τις κύριες επιδράσεις και για τις αλληλεπιδράσεις και να παρουσιάσετε τα συμπεράσματα.

Βιβλιογραφία

- Γιαλαμάς, Β. (2004). *Στατιστικές Τεχνικές και Εφαρμογές στις Επιστήμες της Αγωγής*. Αθήνα: Εκδ. Πατάκη.
- Γναρδέλλης, Χ. (2009). *Ανάλυση δεδομένων με το PASW Statistics 17.0*. Αθήνα: Εκδ. Παπαζήση.
- Field, A. (2013). *Discovering statistics using IBM SPSS statistics* (5th ed.). London: SAGE.
- Ρούσσος, Π., & Τσαούσης, Γ. (2002). *Στατιστική εφαρμοσμένη στις Κοινωνικές Επιστήμες*. Αθήνα: Ελληνικά Γράμματα.

Κεφάλαιο 10

Μη Παραμετρικοί Έλεγχοι

Σύνοψη

Στο κεφάλαιο αυτό θα εξετάσουμε τους τέσσερις μη παραμετρικούς ελέγχους, *Mann-Whitney* για δύο ανεξάρτητα δείγματα, *Wilcoxon* για δύο εξαρτημένα δείγματα, *Kruskal-Wallis* για περισσότερα από δύο ανεξάρτητα δείγματα και *Friedman* για περισσότερα από δύο εξαρτημένα δείγματα. Παρουσιάζονται και ερμηνεύονται τα αποτελέσματα των αναλύσεων τόσο των περιγραφικών στατιστικών όσο και των αντίστοιχων ελέγχων. Επίσης, παρουσιάζεται ο τρόπος συγγραφής των αποτελεσμάτων για καθέναν από τους παραπάνω μη παραμετρικούς ελέγχους, στο πλαίσιο μιας επιστημονικής εργασίας.

Προαπαιτούμενη γνώση

Εκτιμητική και έλεγχος υποθέσεων, έλεγχος μέσων τιμών και ανάλυση διακύμανσης: Κεφάλαια 6, 7, και 9 του συγγράμματος.

10.1 Εισαγωγή στους μη παραμετρικούς ελέγχους

Πολλοί έλεγχοι της Επαγωγικής Στατιστικής, όπως για παράδειγμα ο έλεγχος t και η ανάλυση διακύμανσης, ονομάζονται *παραμετρικοί* γιατί πρέπει να πληρούνται κάποιες προϋποθέσεις που αφορούν τις παραμέτρους του πληθυσμού (βλ. Κεφ. 7 και 9). Για παράδειγμα, πρέπει να ισχύει η κανονικότητα των πληθυσμών (κυρίως σε μικρά δείγματα) κατά την εκτέλεση του ελέγχου t .

Οι μη παραμετρικοί έλεγχοι είναι κατάλληλοι όταν (Γιαλαμάς, 2005; Sprent & Smeeton, 2016):

- Παραβιάζονται οι προϋποθέσεις εφαρμογής των παραμετρικών ελέγχων.
- Δίνονται οι θέσεις ή η τάξη (rank) των ατόμων και όχι τιμές που εκφράζουν αριθμητική μεταβολή από άτομο σε άτομο. Π.χ. ένας εκπαιδευτικός διατάσσει τους 20 μαθητές του από τον λιγότερο δημοφιλή ως τον δημοφιλέστερο δίνοντας τους τιμές-θέσεις: 1, 2, 3, ... 20, προκειμένου να ελέγξει την υπόθεση ότι οι άριστοι μαθητές είναι πιο δημοφιλείς από τους υπόλοιπους συμμαθητές τους.

Οι μη παραμετρικοί έλεγχοι είναι λιγότερο ισχυροί από τους αντίστοιχους παραμετρικούς, δηλαδή οδηγούν δυσκολότερα σε στατιστικά σημαντική διαφορά, όταν πραγματικά αυτή υπάρχει, απ' ότι οι αντίστοιχοι παραμετρικοί έλεγχοι, όταν οι προϋποθέσεις για την εκτέλεση των τελευταίων ικανοποιούνται. Αξίζει να τονίσουμε ότι, επειδή υπάρχει ανθεκτικότητα του ελέγχου t και της ανάλυσης διακύμανσης απέναντι στην παραβίαση των προϋποθέσεων, υπάρχει μια τάση αποφυγής των μη παραμετρικών ελέγχων, οι οποίοι θα πρέπει να χρησιμοποιούνται κυρίως όταν υπάρχει σοβαρή παραβίαση των προϋποθέσεων στους αντίστοιχους παραμετρικούς.

Στο Κεφάλαιο αυτό θα εξετάσουμε τέσσερις μη παραμετρικούς ελέγχους με βάση την αντιστοίχισή τους (βλ. Πίνακα 10.1) με τους παραμετρικούς ελέγχους που είδαμε στα προηγούμενα Κεφάλαια, 7 και 9:

Μη Παραμετρικοί Έλεγχοι	Παραμετρικοί Έλεγχοι
Έλεγχος Mann-Whitney για δύο ανεξάρτητα δείγματα	Έλεγχος t ανεξάρτητων δειγμάτων
Έλεγχος Wilcoxon για δύο εξαρτημένα δείγματα	Έλεγχος t κατά ζεύγη
Έλεγχος Kruskal-Wallis για περισσότερα από δύο ανεξάρτητα δείγματα	Ανάλυση διακύμανσης προς έναν παράγοντα
Έλεγχος Friedman για περισσότερα από δύο εξαρτημένα δείγματα	Ανάλυση διακύμανσης επαναλαμβανόμενων μετρήσεων (εξαρτημένων δειγμάτων)

Πίνακας 10.1 Μη παραμετρικοί έλεγχοι με τους αντίστοιχους παραμετρικούς.

10.1.1 Έλεγχος Mann-Whitney για δύο ανεξάρτητα δείγματα

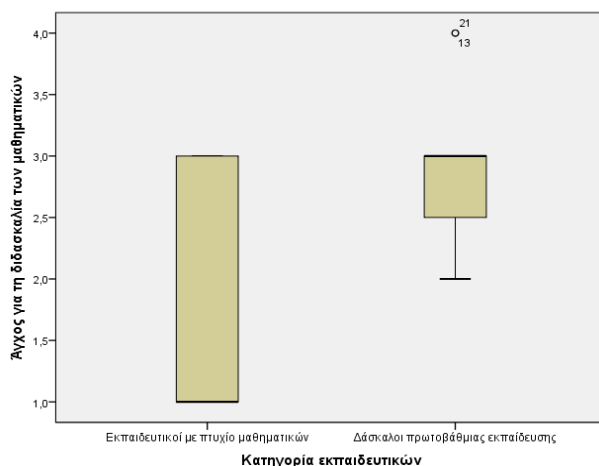
Ο μη παραμετρικός έλεγχος *Mann-Whitney* για δύο ανεξάρτητα δείγματα χρησιμοποιείται όταν δεν ικανοποιούνται οι προϋποθέσεις του αντίστοιχου παραμετρικού ελέγχου (έλεγχος *t* ανεξάρτητων δειγμάτων) (Κεφάλαιο 7).

Η εκτέλεση του ελέγχου *Mann-Whitney* για δύο ανεξάρτητα δείγματα στο SPSS θα γίνει με το ακόλουθο παράδειγμα. Θα χρησιμοποιηθεί το αρχείο δεδομένων «example_indep_non_parametric.sav». Στο συγκεκριμένο αρχείο διερευνάται το άγχος για τη διδασκαλία των Μαθηματικών (μεταβλητή «anxiety») ανάμεσα στους εκπαιδευτικούς με πτυχίο Μαθηματικών και στους δασκάλους της πρωτοβάθμιας εκπαίδευσης (μεταβλητή «group»).

Παρατηρούμε ότι οι μετρήσεις της μεταβλητής που μας ενδιαφέρει (εξαρτημένη μεταβλητή) πραγματοποιήθηκαν με τακτική κλίμακα και υπάρχουν μόνο 12 περιπτώσεις για κάθε κατηγορία εκπαιδευτικών. Λόγω του μικρού δείγματος πρέπει να ελεγχθεί η προϋπόθεση της κατά προσέγγιση κανονικότητας των δύο πληθυσμών. Με τη χρήση των θηκογραμμάτων (Εικόνα 10.1) διαπιστώθηκε ότι υπάρχει σοβαρή απόκλιση από την κανονικότητα, καθώς η γραμμή εντός του παραλληλογράμμου που αντιστοιχεί στη διάμεσο των παρατηρήσεων, δεν βρίσκεται στη μέση. Το γεγονός αυτό, καθώς επίσης και το γεγονός ότι οι παρατηρήσεις είναι ποιοτικές τιμές τακτικής κλίμακας, μας οδηγεί στο να εκτελέσουμε τον έλεγχο *Mann-Whitney* για δύο ανεξάρτητα δείγματα. Τα δύο δείγματα των παρατηρήσεων είναι ανεξάρτητα: (α. Εκπαιδευτικοί με πτυχίο Μαθηματικών που διδάσκουν Μαθηματικά στη δευτεροβάθμια εκπαίδευση και β. εκπαιδευτικοί που διδάσκουν Μαθηματικά στην πρωτοβάθμια εκπαίδευση).

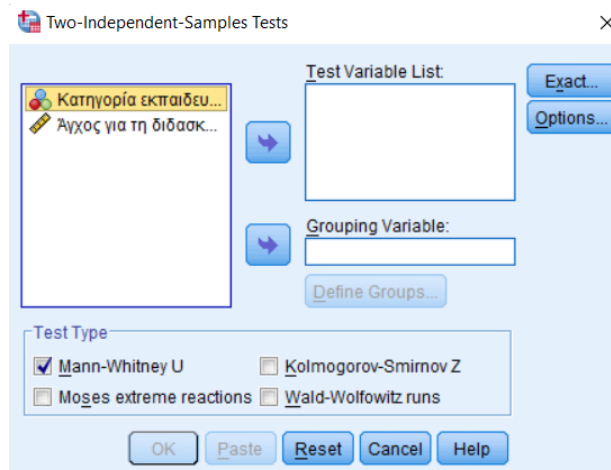
Οι υποθέσεις του ελέγχου είναι οι εξής:

- H_0 : οι κατανομές των τιμών του άγχους διδασκαλίας των Μαθηματικών στις δύο ομάδες εκπαιδευτικών είναι ίδιες, έχουν δηλαδή όλες οι τιμές την ίδια θέση (αριθμός θέσης, αν τις τοποθετήσουμε σε αύξουσα σειρά). Παρόμοια θα μπορούσαμε να πούμε ότι οι κατανομές έχουν τις ίδιες διαμέσους.
- H_A : οι κατανομές των τιμών του άγχους διδασκαλίας των Μαθηματικών στις δύο ομάδες εκπαιδευτικών διαφέρουν μεταξύ τους ως προς τις διαμέσους.



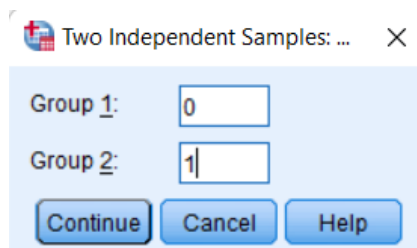
Εικόνα 10.1 Θηκογράμματα άγχους για τη διδασκαλία των Μαθηματικών.

Η διαδικασία εκτέλεσης του ελέγχου *Mann-Whitney* για δύο ανεξάρτητα δείγματα γίνεται στο SPSS ως εξής: Επιλέγουμε **Analyze => Nonparametric Tests => Legacy Dialogs => 2 Independent Samples** και ανοίγει το αντίστοιχο πλαίσιο διαλόγου (Εικόνα 10.2).



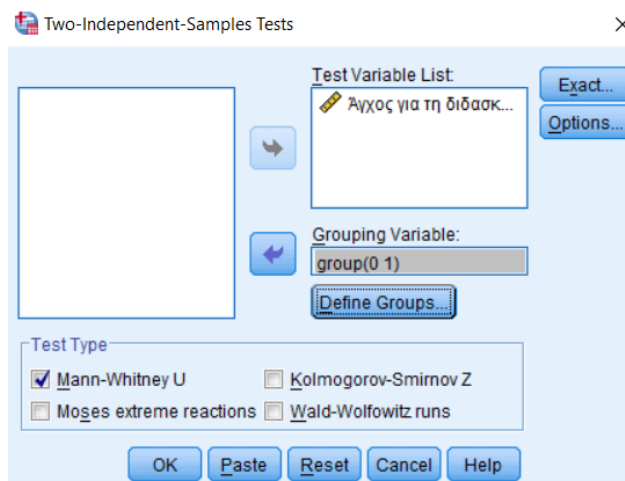
Εικόνα 10.2 Πλαίσιο διαλόγου «Two-Independent-Samples Tests».

Στο πεδίο **Test Variable List** μετακινούμε την ποσοτική μεταβλητή που επιλέγουμε από αριστερά, στην προκειμένη περίπτωση τη μεταβλητή «Άγχος για τη διδασκαλία των μαθηματικών» (anxiety). Την κατηγορική μεταβλητή «Κατηγορία Εκπαιδευτικών» (group) την τοποθετούμε στο πεδίο **Grouping Variable**. Έπειτα, πατάμε το πλήκτρο **Define Groups** και ορίζουμε τους κωδικούς (0 και 1) των δύο κατηγοριών (Εικόνα 10.3).



Εικόνα 10.3 Πλαίσιο διαλόγου «Two-Independent-Samples: Define Groups».

Στη συνέχεια, πατάμε το πλήκτρο **Continue** και μετά το πλήκτρο **OK**. Η τελική μορφή του πλαισίου διαλόγου **Two-Independent-Samples Tests** θα έχει την παρακάτω μορφή (Εικόνα 10.4).



Εικόνα 10.4 Τελική μορφή του πλαισίου διαλόγου «Two-Independent-Samples Tests».

Τα αποτελέσματα της ανάλυσης απεικονίζονται στο αρχείο αποτελεσμάτων του SPSS με τους πίνακες, όπως φαίνεται στην Εικόνα 10.5.

Ranks				
	Κατηγορία εκπαιδευτικών	N	Mean Rank	Sum of Ranks
Άγχος για τη διδασκαλία των μαθηματικών	Εκπαιδευτικοί με πτυχίο μαθηματικών	12	8,79	105,50
	Δάσκαλοι πρωτοβάθμιας εκπαίδευσης	12	16,21	194,50
	Total	24		

Test Statistics ^a	
	Άγχος για τη διδασκαλία των μαθηματικών
Mann-Whitney U	27,500
Wilcoxon W	105,500
Z	-2,746
Asymp. Sig. (2-tailed)	,006
Exact Sig. [2*(1-tailed Sig.)]	,008 ^b

a. Grouping Variable: Κατηγορία εκπαιδευτικών

b. Not corrected for ties.

Εικόνα 10.5 Αποτελέσματα του ελέγχου Mann-Whitney.

Ο πρώτος Πίνακας εμφανίζει το μέγεθος του δείγματος, τη μέση θέση (Mean Rank) και το άθροισμα των θέσεων (Sum of Ranks) της κάθε κατηγορίας. Η μέση θέση των δασκάλων πρωτοβάθμιας εκπαίδευσης (16,21) είναι σχεδόν διπλάσια από τη μέση θέση των εκπαιδευτικών με πτυχίο Μαθηματικών (8,79). Θα μπορούσαμε να παρουσιάσουμε συμπληρωματικά και τις διαμέσους, τις οποίες μπορούμε να υπολογίσουμε από την επιλογή του μενού [Analyze=>descriptive statistics=>explore](#) (βλ. Κεφ. 4).

Ο δεύτερος πίνακας απεικονίζει τα στατιστικά του ελέγχου «Mann-Whitney U» (27,500) και του ελέγχου «Wilcoxon W» (105,500). Ο έλεγχος αθροίσματος θέσεων (rank sum) του *Wilcoxon* είναι ισοδύναμος του *Mann-Whitney* (Field, 2018). Η τυπική τιμή του ελέγχου που υπολογίζεται είναι $Z = -2,746$. Η τιμή «Asymp. Sig. (2-tailed)» βασίζεται στην ασυμπτωτική μέθοδο του ελέγχου και μας δείχνει τη σημαντικότητα της τιμής Z και ισούται με 0,006. Αφού η τιμή αυτή είναι μικρότερη από την τιμή 0,05 ($0,006 < 0,005$), τότε απορρίπτεται η μηδενική υπόθεση της ισότητας των διαμέσων των δύο κατηγοριών. Η τιμή «Exact Sig. [2*(1-tailed Sig.)]» μας δίνει μια ακριβή τιμή για τη σημαντικότητα του ελέγχου, η οποία θα πρέπει να αξιοποιείται όταν τα δείγματα είναι μικρά (μικρότερα από 50 παρατηρήσεις) (Field, 2018). Αφού η τιμή αυτή είναι μικρότερη από την τιμή 0,05 ($0,008 < 0,005$), τότε απορρίπτεται η μηδενική υπόθεση της ισότητας των διαμέσων των δύο κατανομών. Επομένως, απορρίπτεται η μηδενική υπόθεση, δηλαδή υπάρχουν σημαντικές διαφορές ως προς τις μέσες θέσεις.

Το παραπάνω αποτέλεσμα μπορεί να συνταχθεί ως αναφορά στο πλαίσιο μιας επιστημονικής εργασίας ως εξής: Όπως αποκάλυψε ο αμφίπλευρος μη παραμετρικός έλεγχος ανεξάρτητων δειγμάτων *Mann-Whitney*, το άγχος των εκπαιδευτικών πρωτοβάθμιας εκπαίδευσης ($M.θ_{\text{επε}} = 16,21$) για τη διδασκαλία των Μαθηματικών είναι σημαντικά μεγαλύτερο από το άγχος των εκπαιδευτικών με πτυχίο Μαθηματικών ($M.θ_{\text{επι}} = 8,79$) (*Mann-Whitney U* = 27,5, $Z = -2,746$, $p = 0,008$).

10.1.2 Έλεγχος Wilcoxon για δύο εξαρτημένα δείγματα

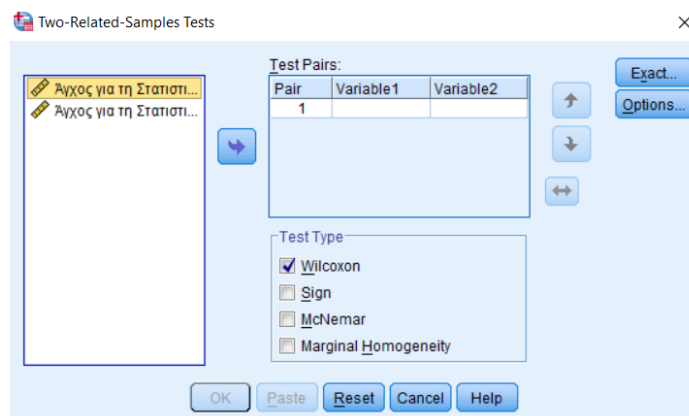
Ο μη παραμετρικός έλεγχος *Wilcoxon* αποτελεί μια εναλλακτική μέθοδο του ελέγχου *t* κατά ζεύγη (δύο εξαρτημένων δειγμάτων που είδαμε στο Κεφάλαιο 7), όταν δεν ικανοποιούνται οι προϋποθέσεις του ελέγχου *t*. Το ερευνητικό σχέδιο στην περίπτωση αυτή είναι οι επαναλαμβανόμενες μετρήσεις, κατά τις οποίες μία ομάδα ατόμων υπόκειται σε δύο μετρήσεις σε δύο διαφορετικά χρονικά διαστήματα. Τα διαστήματα αυτά μπορεί να αντιστοιχούν σε δύο μετρήσεις ενός πειράματος, όπως, για παράδειγμα, σε έναν έλεγχο πριν-μετά μιας διδακτικής παρέμβασης (Γιαλαμάς, 2005).

Η εκτέλεση του ελέγχου *Wilcoxon* για δύο εξαρτημένα δείγματα στο SPSS θα γίνει με το ακόλουθο παράδειγμα. Θα χρησιμοποιηθεί το αρχείο δεδομένων «example_dep_non_parametric.sav». Στο συγκεκριμένο αρχείο διερευνάται το άγχος των φοιτητών για τη Στατιστική πριν (μεταβλητή «anxiety_pre») και μετά τη διδασκαλία των μαθημάτων Στατιστικής (μεταβλητή «anxiety_post»). Οι μετρήσεις της μεταβλητής που μας ενδιαφέρει έχουν γίνει με τακτική κλίμακα. Λόγω του ότι τα δεδομένα προέρχονται από τακτικές κλίμακες θα εκτελεστεί ο μη παραμετρικός έλεγχος *Wilcoxon* για δύο εξαρτημένα δείγματα.

Οι υποθέσεις του ελέγχου είναι οι εξής:

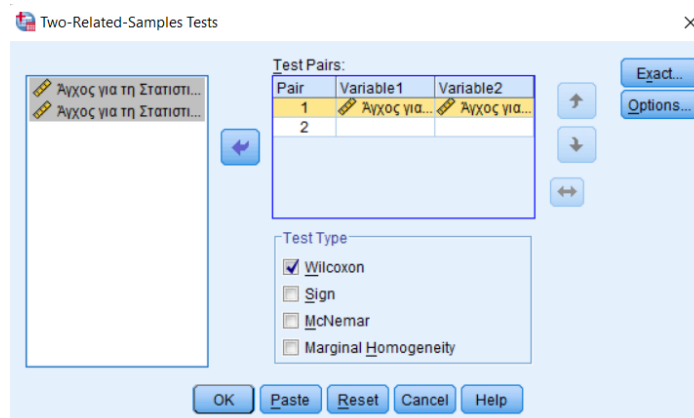
- H_0 : οι κατανομές των τιμών του άγχους για τη Στατιστική πριν και μετά τη διδασκαλία είναι ίδιες. Εναλλακτικά έχουν τις ίδιες διαμέσους.
- H_A : οι κατανομές των τιμών του άγχους για τη Στατιστική πριν και μετά τη διδασκαλία διαφέρουν ως προς τις διαμέσους.

Η διαδικασία εκτέλεσης του ελέγχου *Wilcoxon* για δύο εξαρτημένα δείγματα γίνεται στο SPSS ως εξής: Επιλέγουμε **Analyze => Nonparametric Tests => Legacy Dialogs => 2 Related Samples** και ανοίγει το αντίστοιχο πλαίσιο διαλόγου (Εικόνα 10.6).



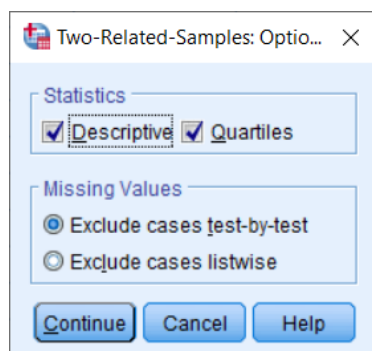
Εικόνα 10.6 Πλαίσιο διαλόγου «Two-Related-Samples Tests».

Επιλέγουμε από αριστερά τις δύο μεταβλητές μαζί, οι οποίες αντιστοιχούν στις επαναλαμβανόμενες μετρήσεις, και τις εισάγουμε διαδοχικά στη λίστα **Test Pairs** (Εικόνα 10.7).



Εικόνα 10.7 Τελική μορφή του πλαισίου διαλόγου «Two- Related -Samples Tests».

Έπειτα, κάνουμε κλικ στο πλήκτρο **Options** και μαρκάρουμε την επιλογή **Descriptive** (Εικόνα 10.8), ώστε να υπολογιστούν τα περιγραφικά στατιστικά (υπολογίζονται η μέση τιμή, η τυπική απόκλιση, η ελάχιστη και η μέγιστη τιμή των τιμών των δύο μετρήσεων) για καθεμία μεταβλητή του ζεύγους. Στην περίπτωση που θέλουμε να υπολογίσουμε και τη διάμεσο (δεύτερο τεταρτημόριο) επιλέγουμε την επιλογή **Quartiles** (τεταρτημόρια). Στη συνέχεια, πατάμε το πλήκτρο **Continue** και μετά το πλήκτρο **OK**.



Εικόνα 10.8 Πλαίσιο διαλόγου «Two-Related-Samples:Options».

Τα αποτελέσματα της ανάλυσης απεικονίζονται στο αρχείο αποτελεσμάτων του SPSS με τους Πίνακες, όπως φαίνεται στην παρακάτω Εικόνα 10.9.

Descriptive Statistics								
	N	Mean	Std. Deviation	Minimum	Maximum	Percentiles		
						25th	50th (Median)	75th
Anxiety_pre Anxiety statistics pre	36	2,58	,770	1	4	2,00	3,00	3,00
Anxiety_post Anxiety statistics post	36	1,75	,937	1	3	1,00	1,00	3,00

Ranks				
		N	Mean Rank	Sum of Ranks
Άγχος για τη Στατιστική μετά τη διδασκαλία - Άγχος για τη Στατιστική πριν τη διδασκαλία	Negative Ranks	21 ^a	15,71	330,00
	Positive Ranks	6 ^b	8,00	48,00
	Ties	9 ^c		
	Total	36		

a. Άγχος για τη Στατιστική μετά τη διδασκαλία < Άγχος για τη Στατιστική πριν τη διδασκαλία

b. Άγχος για τη Στατιστική μετά τη διδασκαλία > Άγχος για τη Στατιστική πριν τη διδασκαλία

c. Άγχος για τη Στατιστική μετά τη διδασκαλία = Άγχος για τη Στατιστική πριν τη διδασκαλία

Test Statistics ^a	
	Άγχος για τη Στατιστική μετά τη διδασκαλία - Άγχος για τη Στατιστική πριν τη διδασκαλία
Z	-3,474 ^b
Asymp. Sig. (2-tailed)	,001

a. Wilcoxon Signed Ranks Test

b. Based on positive ranks.

Εικόνα 10.9 Αποτελέσματα του ελέγχου Wilcoxon.

Ο πρώτος Πίνακας εμφανίζει τα περιγραφικά στατιστικά για καθεμία από τις δύο μετρήσεις: το μέγεθος του δείγματος, τη μέση τιμή, την τυπική απόκλιση, την ελάχιστη και τη μέγιστη τιμή, καθώς επίσης και τα τρία τεταρτημόρια. Το μέσο άγχος των φοιτητών για τη Στατιστική πριν τη διδασκαλία (Μ. Τ. = 2,58, Τ. Α. = 0,77) είναι μεγαλύτερο από το μέσο άγχος των φοιτητών για τη Στατιστική μετά τη διδασκαλία (Μ. Τ. = 1,75, Τ. Α. = 0,94).

Ο δεύτερος πίνακας απεικονίζει τις προσημασμένες θέσεις των διαφορών των τιμών του άγχους, ακολουθώντας την κατάταξή τους. Ο έλεγχος αυτός στηρίζεται στις διαφορές του ζεύγους των παρατηρήσεων στις οποίες βασίζεται η σύγκριση. Οι διαφορές αυτές (όχι τα ζεύγη που δεν διαφέρουν) τοποθετούνται σε σειρά, αγνοώντας το πρόσημο και υπολογίζονται οι θέσεις των διαφορών. Στη στήλη N εμφανίζεται ο αριθμός των παρατηρήσεων στην αντίστοιχη ομάδα. Οι διαφορές υπολογίζονται αφαιρώντας από τις τιμές της 2^{ης} μεταβλητής (άγχος για τη Στατιστική μετά τη διδασκαλία) τις αντίστοιχες τιμές της 1^{ης} μεταβλητής (άγχος για τη Στατιστική πριν τη διδασκαλία). Παρατηρούμε στη δεύτερη στήλη ότι από τις συνολικά 36 περιπτώσεις του δείγματος οι 21 εμφανίζουν αρνητικές διαφορές (Negative Ranks), 6 περιπτώσεις εμφανίζουν θετικές διαφορές (Positive Ranks) και 9 περιπτώσεις έχουν τις ίδιες τιμές (Ties). Στην τρίτη στήλη (Mean Rank) παρατηρούμε ότι η μέση (απόλυτη) θέση των αρνητικών διαφορών στην κατάταξη είναι 15,71 και η μέση θέση των θετικών διαφορών είναι 8,00. Με λίγα λόγια, οι αρνητικές διαφορές είναι μεγαλύτερες σε σχέση με τις θετικές. Στην τέταρτη στήλη (Sum of Ranks) εμφανίζονται το άθροισμα των αρνητικών θέσεων, το οποίο ισούται με 330, και το άθροισμα των θετικών θέσεων, που ισούται με 48.

Ο τρίτος Πίνακας απεικονίζει την τυποποιημένη τιμή Z (3,474) του ελέγχου, καθώς και τη σημαντικότητα της τιμής Z ($p = 0,001$) του αμφίπλευρου ελέγχου. Αφού η τιμή αυτή είναι μικρότερη από την τιμή 0,05 ($0,001 < 0,005$), απορρίπτεται η μηδενική υπόθεση της ισότητας των διαμέσων των δύο μετρήσεων. Επομένως, απορρίπτεται η μηδενική υπόθεση, δηλαδή υπάρχουν σημαντικές διαφορές μεταξύ των δύο μετρήσεων. Επιπλέον, όπως θα δείτε στην υποσημείωση του τελευταίου Πίνακα της Εικόνας 10.9, ο έλεγχος αυτός βασίστηκε στο άθροισμα θέσεων με θετικό πρόσημο, επομένως η συνάρτηση ελέγχου εδώ είναι το άθροισμα των θέσεων με θετικό πρόσημο, δηλαδή $T=48$.

Το παραπάνω αποτέλεσμα μπορεί να συνταχθεί ως αναφορά στο πλαίσιο μιας επιστημονικής εργασίας ως εξής: Όπως αποκάλυψε ο αμφίπλευρος μη παραμετρικός έλεγχος εξαρτημένων δειγμάτων *Wilcoxon*, το άγχος των φοιτητών για τη Στατιστική, πριν τη διδασκαλία της (διάμεσος=3), είναι σημαντικά μεγαλύτερο από το άγχος μετά τη διδασκαλία (διάμεσος=1) ($T=48$, $Z = 3,474$, $p = 0,001$).

10.1.3 Έλεγχος Kruskal-Wallis για περισσότερα από δύο ανεξάρτητα δείγματα

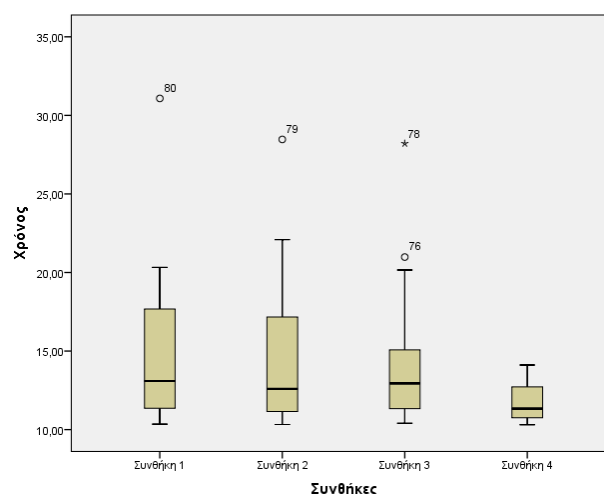
Ο μη παραμετρικός έλεγχος *Kruskal-Wallis* αποτελεί το αντίστοιχο του παραμετρικού ελέγχου της ανάλυσης διακύμανσης με έναν παράγοντα (Κεφάλαιο 9) για περισσότερα από δύο ανεξάρτητα δείγματα. Χρησιμοποιείται όταν οι κατανομές των πληθυσμών από τους οποίους προέρχονται τα δείγματα της έρευνάς μας δεν ακολουθούν την κανονική κατανομή και δεν έχουν ίσες διακυμάνσεις ή όταν η εξαρτημένη μεταβλητή είναι τακτικής (διατάξιμης) κλίμακας.

Η εκτέλεση του ελέγχου *Kruskal-Wallis* για περισσότερα από δύο ανεξάρτητα δείγματα στο SPSS, θα γίνει με το ακόλουθο παράδειγμα. Θα χρησιμοποιηθεί το αρχείο δεδομένων «k_ind_non_param.sav». Στο συγκεκριμένο αρχείο περιλαμβάνονται οι συνθήκες που ακολούθησαν τα νήπια και οι μετρήσεις των χρόνων (σε λεπτά) (μεταβλητή «time») που έκαναν τα νήπια, προκειμένου να σχεδιάσουν μια συγκεκριμένη εικόνα. Σε κάθε ομάδα - συνθήκη: 1, 2, 3, 4 (μεταβλητή «group») παρείχαμε όλο και πιο πολλούς βαθμούς ελευθερίας στα παιδιά για την αξιοποίηση διαφορετικών χρωμάτων. Στόχος μας είναι να ελεγχθεί η υπόθεση ότι οι χρόνοι που έκαναν τα νήπια είναι διαφορετικοί σε κάθε ομάδα.

Λόγω του ότι υπάρχουν 20 περιπτώσεις για κάθε συνθήκη, θα πρέπει να ελεγχθεί η προϋπόθεση της κατά προσέγγιση κανονικότητας των δύο πληθυσμών. Με τη χρήση τόσο των θηκογραμμάτων (Εικόνα 10.10) όσο και των ελέγχων για την κανονικότητα *Kolmogorov-Smirnov* και *Shapiro-Wilk* (Εικόνα 10.11) διαπιστώθηκε ότι υπάρχει σοβαρή απόκλιση από την κανονικότητα στις περισσότερες συνθήκες. Αυτό μας οδηγεί στο να εκτελέσουμε τον έλεγχο *Kruskal-Wallis* για περισσότερα από δύο ανεξάρτητα δείγματα.

Οι υποθέσεις του ελέγχου είναι οι εξής:

- H_0 : οι κατανομές των τιμών στις τέσσερις ομάδες είναι ίδιες. Εναλλακτικά έχουν τις ίδιες διαμέσους.
- H_A : σε δύο τουλάχιστον ομάδες οι κατανομές των τιμών διαφέρουν ως προς τις διαμέσους τους.



Εικόνα 10.10 Θηκογράμματα συνθηκών της έρευνας.

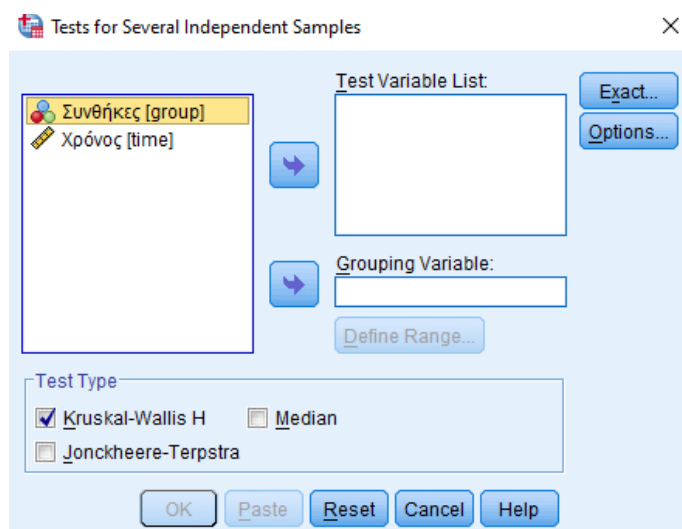
Tests of Normality						
Συνθήκες		Kolmogorov-Smirnov ^a			Shapiro-Wilk	
		Statistic	df	Sig.	Statistic	Sig.
Χρόνος	Συνθήκη 1	,181	20	,085	,805	,001
	Συνθήκη 2	,207	20	,024	,826	,002
	Συνθήκη 3	,267	20	,001	,743	,000
	Συνθήκη 4	,204	20	,028	,912	,071

a. Lilliefors Significance Correction

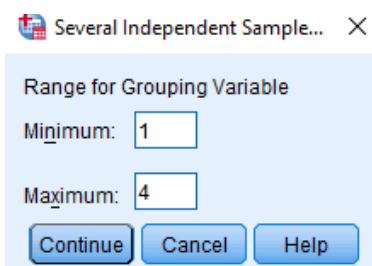
Εικόνα 10.11 Έλεγχοι κανονικότητας *Kolmogorov-Smirnov* και *Shapiro-Wilk*.

Η διαδικασία εκτέλεσης του ελέγχου *Kruskal-Wallis* για περισσότερα από δύο ανεξάρτητα δείγματα γίνεται στο SPSS ως εξής: Επιλέγουμε **Analyze => Nonparametric Tests => Legacy Dialogs => K Independent Samples** και ανοίγει το αντίστοιχο πλαίσιο διαλόγου (Εικόνα 10.12).

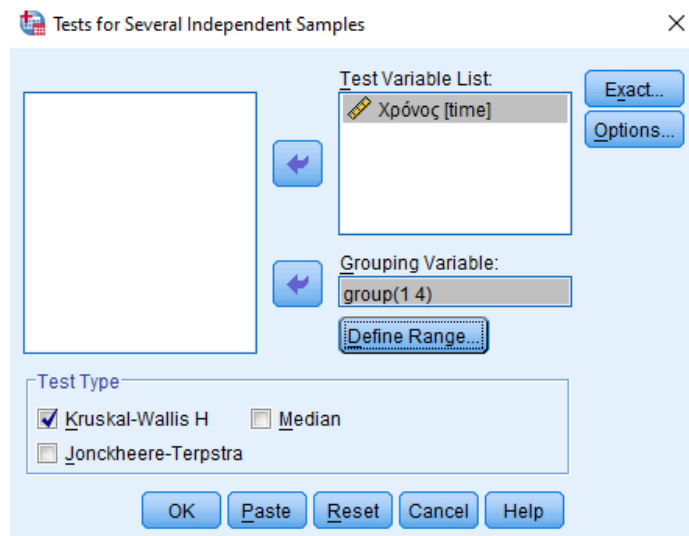
Στο πεδίο **Test Variable List** μετακινούμε την ποσοτική μεταβλητή που επιλέγουμε από αριστερά, στην προκειμένη περίπτωση τη μεταβλητή «Χρόνος» (time). Την κατηγορική μεταβλητή «Συνθήκες» (group) την τοποθετούμε στο πεδίο **Grouping Variable**. Έπειτα πατάμε το πλήκτρο **Define Range** και ορίζουμε την ελάχιστη (1) και τη μέγιστη τιμή (4) του εύρους των κωδικών των κατηγοριών (Εικόνα 10.13) και στη συνέχεια πατάμε το πλήκτρο **Continue**. Προσοχή στην περίπτωση που υπάρχουν κατηγορίες με κωδικούς εκτός του συγκεκριμένου εύρους που ορίζουμε, καθώς αυτές δεν θα συμπεριληφθούν στην ανάλυση. Για την ολοκλήρωση της διαδικασίας πατάμε το πλήκτρο **OK**. Η τελική μορφή του πλαισίου διαλόγου **Tests for Several Independent Samples** θα έχει την παρακάτω μορφή (Εικόνα 10.14).



Εικόνα 10.12 Πλαίσιο διαλόγου «*Tests for Several Independent Samples*».



Εικόνα 10.13 Πλαίσιο διαλόγου «*Several Independent Samples: Define Range*».



Εικόνα 10.14 Τελική μορφή του πλαισίου διαλόγου «Tests for Several Independent Samples».

Τα αποτελέσματα της ανάλυσης απεικονίζονται στο αρχείο αποτελεσμάτων του SPSS με τους πίνακες, όπως φαίνεται στην παρακάτω Εικόνα 10.15.

Ranks			
group Συνθήκες		N	Mean Rank
time Χρόνος	1 Συνθήκη 1	20	46,35
	2 Συνθήκη 2	20	44,15
	3 Συνθήκη 3	20	44,15
	4 Συνθήκη 4	20	27,35
Total		80	

Test Statistics ^{a,b}	
	time Χρόνος
Chi-Square	8,659
df	3
Asymp. Sig.	,034

a. Kruskal Wallis Test

b. Grouping Variable: group Συνθήκες

Εικόνα 10.15 Αποτελέσματα του ελέγχου Kruskal-Wallis.

Ο πρώτος Πίνακας παρουσιάζει το μέγεθος του δείγματος και τη μέση θέση (Mean Rank) κάθε κατηγορίας. Οι μέσες θέσεις των τεσσάρων συνθηκών διαφέρουν μεταξύ τους. Η μέση θέση της συνθήκης 1 είναι 46,35, της συνθήκης 2 και 3 είναι 44,15 και της συνθήκης 4 είναι 27,35.

Ο δεύτερος Πίνακας απεικονίζει το στατιστικό του ελέγχου Kruskal-Wallis, το οποίο ισούται με 8,66 και την τιμή «Asymp. Sig. (2-tailed)», η οποία μας δείχνει τη σημαντικότητα της κατανομής χ^2 και ισούται με 0,034. Αφού η τιμή αυτή είναι μικρότερη από την τιμή 0,05 ($0,034 < 0,005$), απορρίπτεται η μηδενική υπόθεση

της ισότητας των διαμέσων σε όλες τις συνθήκες. Επομένως, υπάρχουν σημαντικές διαφορές ως προς τις μέσες θέσεις σε τουλάχιστον δύο συνθήκες.

Το SPSS δεν παρέχει ελέγχους πολλαπλών συγκρίσεων για τον μη παραμετρικό έλεγχο *Kruskal-Wallis*. Εναλλακτικά, μπορούμε να πραγματοποιήσουμε συγκρίσεις ανά δύο ανεξάρτητα δείγματα μέσω του μη παραμετρικού ελέγχου *Mann-Whitney*. Ωστόσο, για να μην υποπέσουμε σε σφάλμα Τύπου I, θα πρέπει να προβούμε στη διόρθωση *Bonferroni*. Για τον σκοπό αυτό, αλλάζουμε το επίπεδο σημαντικότητας του ελέγχου ως εξής:

- Διαιρούμε το $\alpha = 0,05$ με το πλήθος των συγκρίσεων που θα πραγματοποιήσουμε.
- Αν έχουμε k ομάδες, τότε οι δυνατές συγκρίσεις είναι $k(k-1)/2$ και επομένως το επίπεδο σημαντικότητας των ελέγχων είναι $\alpha' = \alpha / [k(k-1)/2] = 0,05/6 = 0,00833$.

Βέβαια, αν ζητούμε τον έλεγχο συγκεκριμένων ομάδων (αφορά συγκεκριμένες ερευνητικές υποθέσεις) έχουμε τη δυνατότητα να πραγματοποιήσουμε μόνο κάποιες από τις δυνατές συγκρίσεις. Στην περίπτωση αυτή, ο παρονομαστής για τον υπολογισμό του α' θα είναι το πλήθος των συγκρίσεων που μας ενδιαφέρουν να διεξάγουμε.

Στο παράδειγμά μας, μας ενδιαφέρει να πραγματοποιήσουμε τρεις δυνατές συγκρίσεις μεταξύ της πρώτης ομάδας (1^η συνθήκη) και των άλλων τριών. Επομένως, το $\alpha' = 0,05/3 = 0,0167$. Θα εκτελέσουμε τρεις ελέγχους *Mann-Whitney* (βλ. 10.1 Έλεγχος Mann-Whitney για δύο ανεξάρτητα δείγματα), οι οποίοι μας εμφανίζουν τους Πίνακες της Εικόνας 10.16.

Ranks				
	group Συνθήκες	N	Mean Rank	Sum of Ranks
time Χρόνος	1 Συνθήκη 1	20	20,95	419,00
	2 Συνθήκη 2	20	20,05	401,00
	Total	40		

Test Statistics ^a	
	time Χρόνος
Mann-Whitney U	191,000
Wilcoxon W	401,000
Z	-,243
Asymp. Sig. (2-tailed)	,808
Exact Sig. [2*(1-tailed Sig.)]	,820 ^b

a. Grouping Variable: group Συνθήκες
b. Not corrected for ties.

Ranks				
	group Συνθήκες	N	Mean Rank	Sum of Ranks
time Χρόνος	1 Συνθήκη 1	20	21,10	422,00
	3 Συνθήκη 3	20	19,90	398,00
	Total	40		

Test Statistics ^a	
	time Χρόνος
Mann-Whitney U	188,000
Wilcoxon W	398,000
Z	-,325
Asymp. Sig. (2-tailed)	,745
Exact Sig. [2*(1-tailed Sig.)]	,758 ^b

a. Grouping Variable: group Συνθήκες
b. Not corrected for ties.

Ranks				
	group Συνθήκες	N	Mean Rank	Sum of Ranks
time Χρόνος	1 Συνθήκη 1	20	25,30	506,00
	4 Συνθήκη 4	20	15,70	314,00
	Total	40		

Test Statistics ^a	
	time Χρόνος
Mann-Whitney U	104,000
Wilcoxon W	314,000
Z	-2,597
Asymp. Sig. (2-tailed)	,009
Exact Sig. [2*(1-tailed Sig.)]	,009 ^b

a. Grouping Variable: group Συνθήκες
b. Not corrected for ties.

Εικόνα 10.16 Τρεις συγκρίσεις μεταξύ της πρώτης ομάδας και των άλλων τριών.

Από τα παραπάνω παρατηρούμε ότι απορρίπτεται η μηδενική υπόθεση, με επίπεδο σημαντικότητας $\alpha' = 0,0167$ στη σύγκριση της συνθήκης 1 με τη συνθήκη 4, αφού το p (exact sig)=0,009 είναι μικρότερο από 0,0167.

Το παραπάνω αποτέλεσμα μπορεί να συνταχθεί ως αναφορά στο πλαίσιο μιας επιστημονικής εργασίας ως εξής: Σύμφωνα με τον μη παραμετρικό έλεγχο ανεξάρτητων δειγμάτων *Kruskal-Wallis* ο χρόνος των νηπίων για την πραγματοποίηση του σχεδίου διαφέρει σε τουλάχιστον δύο συνθήκες ($H(3) = 8,66$, $p = 0,034$). (Προσοχή στο ότι η ελεγχοσυνάρτηση στον έλεγχο αυτόν, αν και δεν φαίνεται, είναι η H και όχι η χ^2).

Για τον έλεγχο των συγκρίσεων της πρώτης ομάδας με τις άλλες τρεις, αξιοποιώντας τη διόρθωση *Bonferroni*, χρησιμοποιήσαμε σε τρεις περιπτώσεις τον έλεγχο *Mann-Whitney*. Από τον έλεγχο αυτό φαίνεται ότι η μέση θέση της επίδοσης των μαθητών της 4^{ης} συνθήκης (15,70) είναι σημαντικά μικρότερη της μέσης θέσης της επίδοσης των μαθητών της 1^{ης} συνθήκης (25,30) (*Mann-Whitney U*=104, $Z=-2,567$, $p=0,009$). Φαίνεται, δηλαδή, ότι η επίδοση των μαθητών στην τέταρτη συνθήκη είναι στατιστικά σημαντικά υψηλότερη (μικρότερος χρόνος αποπεράτωσης του ίδιου έργου) από την επίδοση των μαθητών της πρώτης συνθήκης.

10.1.4 Έλεγχος Friedman για περισσότερα από δύο εξαρτημένα δείγματα

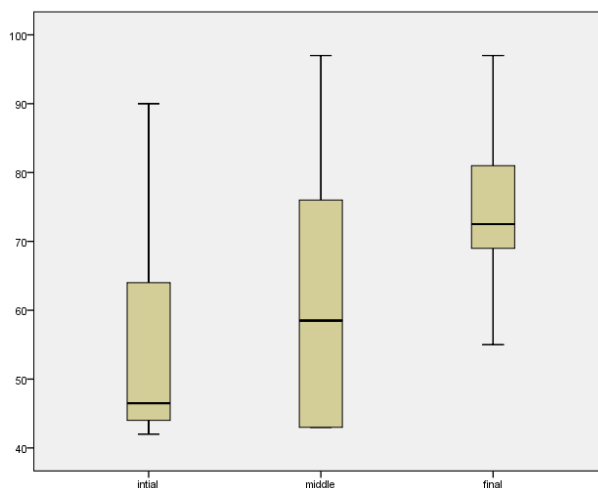
Ο έλεγχος *Friedman* αποτελεί τον μη παραμετρικό έλεγχο της ανάλυσης διακύμανσης με έναν παράγοντα για περισσότερα από δύο εξαρτημένα δείγματα. Αποτελεί την εναλλακτική προσέγγιση του παραμετρικού ελέγχου *Ανάλυση διακύμανσης επαναλαμβανόμενων μετρήσεων*. Χρησιμοποιείται, όταν οι κατανομές των πληθυσμών από τους οποίους προέρχονται τα δείγματα της έρευνάς μας δεν είναι κανονικές και δεν έχουν ίσες διακυμάνσεις ή όταν η εξαρτημένη μεταβλητή είναι τακτικής (διατάξιμης) κλίμακας.

Η εκτέλεση του ελέγχου *Friedman* για περισσότερα από δύο εξαρτημένα δείγματα στο SPSS θα γίνει με το ακόλουθο παράδειγμα. Θα χρησιμοποιηθεί το αρχείο δεδομένων «repeated meas non par.sav». Αφορά στη μέτρηση της επίδοσης των μαθητών Γυμνασίου με ένα ισοδύναμο τεστ στην αρχή (initial), σε κάποια ενδιάμεση χρονική στιγμή (middle) και στο τέλος (final) μιας διδακτικής παρέμβασης για τη διδασκαλία της ισότητας τριγώνων με την υποστήριξη των ΤΠΕ.

Λόγω του ότι υπάρχουν 10 περιπτώσεις για κάθε συνθήκη (μέτρηση των ίδιων υποκειμένων), θα πρέπει να ελεγχθεί η προϋπόθεση της κατά προσέγγιση κανονικότητας των δύο πληθυσμών. Με τη χρήση τόσο των θηκογραμμάτων (Εικόνα 10.17) όσο και των ελέγχων για την κανονικότητα *Kolmogorov-Smirnov* και *Shapiro-Wilk* (Εικόνα 10.18) διαπιστώθηκε ότι υπάρχει σοβαρή απόκλιση από την κανονικότητα στις δύο από τις τρεις συνθήκες. Επομένως, δεν μπορούμε να χρησιμοποιήσουμε παραμετρικό τεστ ανάλυσης διακύμανσης επαναλαμβανόμενων συγκρίσεων. Αυτό μας οδηγεί στο να εκτελέσουμε τον έλεγχο *Friedman* για περισσότερα από δύο εξαρτημένα δείγματα.

Οι υποθέσεις του ελέγχου είναι οι εξής:

- H_0 : οι κατανομές των τιμών στις τρεις συνθήκες είναι ίδιες. Εναλλακτικά, έχουν τις ίδιες διαμέσους.
- H_A : σε δύο τουλάχιστον συνθήκες οι κατανομές των τιμών διαφέρουν ως προς τις διαμέσους τους.



Εικόνα 10.17 Θηκογράμματα των τριών συνθηκών της έρευνας.

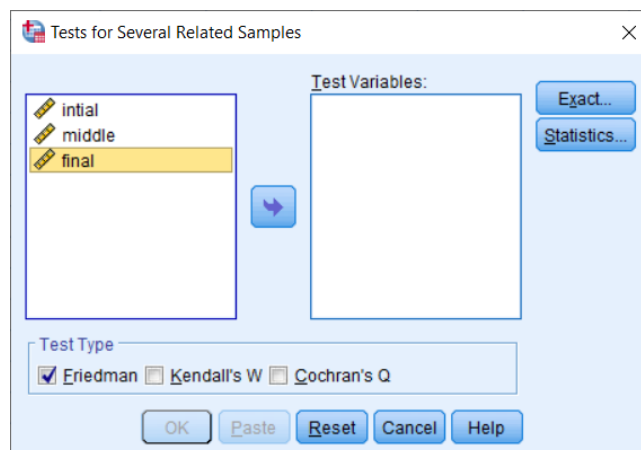
Tests of Normality						
	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
initial	,291	10	,016	,778	10	,008
middle	,271	10	,036	,841	10	,046
final	,188	10	,200 [*]	,907	10	,259

*. This is a lower bound of the true significance.

a. Lilliefors Significance Correction

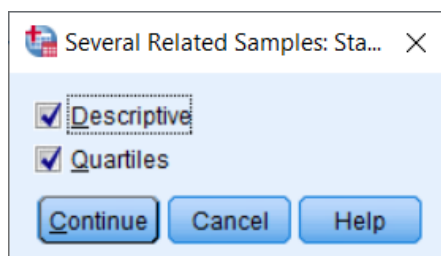
Εικόνα 10.18 Έλεγχοι κανονικότητας *Kolmogorov-Smirnov* και *Shapiro-Wilk*.

Η διαδικασία εκτέλεσης του ελέγχου *Friedman* για περισσότερα από δύο ανεξάρτητα δείγματα γίνεται στο SPSS ως εξής: Επιλέγουμε **Analyze** => **Nonparametric Tests** => **Legacy Dialogs** => **K Related Samples** και ανοίγει το αντίστοιχο πλαίσιο διαλόγου (Εικόνα 10.19). Στο πεδίο **Test Variables** μετακινούμε τις ποσοτικές μεταβλητές που επιλέγουμε από αριστερά, στην προκειμένη περίπτωση τις μεταβλητές «initial», «middle» και «final».



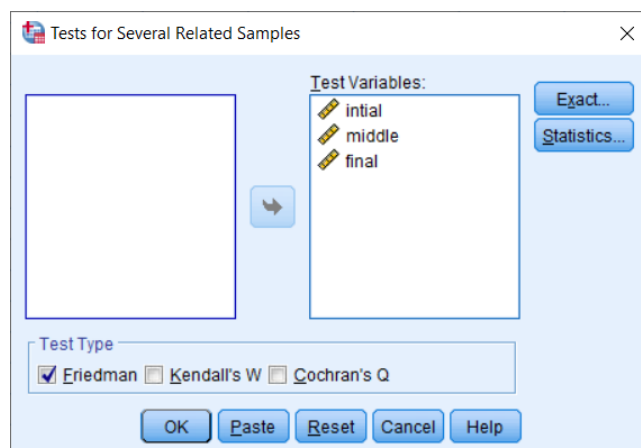
Εικόνα 10.19 Πλαίσιο διαλόγου «*Tests for Several Related Samples*».

Κάνουμε «κλικ» στο πλήκτρο **Statistics** και μαρκάρουμε την επιλογή **Descriptive** (Εικόνα 10.20), ώστε να υπολογιστούν τα περιγραφικά στατιστικά (υπολογίζονται η μέση τιμή, η τυπική απόκλιση, η ελάχιστη και η μέγιστη τιμή των τιμών των δύο μετρήσεων) για καθεμία μεταβλητή του ζεύγους. Στην περίπτωση που θέλουμε να υπολογίσουμε και τη διάμεσο (δεύτερο τεταρτημόριο) επιλέγουμε την επιλογή **Quartiles** (τεταρτημόρια). Έπειτα πατάμε το πλήκτρο **Continue** και μετά το πλήκτρο **OK**.



Εικόνα 10.20 Πλαίσιο διαλόγου «*Several Related Samples: Statistics*».

Η τελική μορφή του πλαισίου διαλόγου **Tests for Several Independent Samples** θα έχει την παρακάτω μορφή (Εικόνα 10.21).



Εικόνα 10.21 Τελική μορφή του πλαισίου διαλόγου «Tests for Several Independent Samples».

Τα αποτελέσματα της ανάλυσης απεικονίζονται στο αρχείο αποτελεσμάτων του SPSS με τους πίνακες, όπως φαίνεται στην παρακάτω Εικόνα 10.22.

Descriptive Statistics								
	N	Mean	Std. Deviation	Minimum	Maximum	Percentiles		
						25th	50th (Median)	75th
initial	10	55,80	17,599	42	90	43,50	46,50	68,75
middle	10	62,30	20,645	43	97	43,00	58,50	78,50
final	10	75,90	13,076	55	97	68,75	72,50	85,00

Ranks	
	Mean Rank
initial	1,40
middle	2,00
final	2,60

Test Statistics ^a	
N	10
Chi-Square	7,784
df	2
Asymp. Sig.	,020
Exact Sig.	,018
Point Probability	,001
a. Friedman Test	

Εικόνα 10.22 Αποτελέσματα του ελέγχου Friedman.

Ο πρώτος Πίνακας παρουσιάζει τα περιγραφικά στατιστικά της κατανομής της κάθε μέτρησης. Ο δεύτερος Πίνακας παρουσιάζει τη μέση θέση (Mean Rank) κάθε μέτρησης. Οι μέσες θέσεις των τριών μετρήσεων διαφέρουν μεταξύ τους. Η μέση θέση της πρώτης μέτρησης (initial) είναι 1,40, της δεύτερης (middle) είναι 2,00 και της τρίτης (final) είναι 2,60.

Ο τρίτος Πίνακας απεικονίζει το στατιστικό του ελέγχου *Friedman*, το οποίο ισούται με 7,784, και την τιμή «Asymp. Sig.», η οποία μας δείχνει τη σημαντικότητα της κατανομής χ^2 και ισούται με 0,020. Η τιμή «Exact Sig.» μας δίνει μια ακριβή τιμή για τη σημαντικότητα του ελέγχου. Αφού η τιμή αυτή είναι μικρότερη από την τιμή 0,05 ($0,020 < 0,005$), απορρίπτεται η μηδενική υπόθεση της ισότητας των διαμέσων σε όλες τις μετρήσεις. Επομένως, υπάρχουν σημαντικές διαφορές ως προς τις μέσες θέσεις σε τουλάχιστον δύο μετρήσεις. Αυτό πρακτικά σημαίνει ότι υπάρχουν σημαντικές αλλαγές σε δύο τουλάχιστον από τις μετρήσεις που πραγματοποιήθηκαν στη διάρκεια της διδακτικής παρέμβασης για τη διδασκαλία της ισότητας τριγώνων σε μαθητές Γυμνασίου.

Το SPSS δεν παρέχει ελέγχους πολλαπλών συγκρίσεων για τον μη παραμετρικό έλεγχο *Friedman*. Εναλλακτικά, μπορούμε να πραγματοποιήσουμε συγκρίσεις ανά δύο εξαρτημένων δειγμάτων μέσω του μη παραμετρικού ελέγχου *Wilcoxon*. Ωστόσο, για να μην υποπέσουμε σε σφάλμα Τύπου Ι, θα πρέπει να προβούμε στη διόρθωση *Bonferroni*. Για τον σκοπό αυτό, αλλάζουμε το επίπεδο σημαντικότητας του ελέγχου ως εξής:

- Διαιρούμε το $\alpha = 0,05$ με το πλήθος των συγκρίσεων που θα πραγματοποιήσουμε.
- Αν έχουμε k ομάδες, τότε οι δυνατές συγκρίσεις είναι $k(k-1)/2$ και επομένως το επίπεδο σημαντικότητας των ελέγχων είναι $\alpha' = \alpha/[k(k-1)/2] = 0,05/6 = 0,00833$.

Βέβαια, αν ζητούμε τον έλεγχο συγκεκριμένων ομάδων (αφορά συγκεκριμένες ερευνητικές υποθέσεις) έχουμε τη δυνατότητα να πραγματοποιήσουμε μόνο κάποιες από τις δυνατές συγκρίσεις. Στην περίπτωση αυτή, ο παρονομαστής για τον υπολογισμό του α' θα είναι το πλήθος των συγκρίσεων που μας ενδιαφέρουν να διεξάγουμε.

Στο παράδειγμά μας, μας ενδιαφέρει να πραγματοποιήσουμε όλες τις δυνατές συγκρίσεις (συνολικά τρεις). Επομένως, το $\alpha' = 0,05/3 = 0,0167$. Θα εκτελέσουμε τρεις ελέγχους *Wilcoxon* (βλ. 10.2 Έλεγχος *Wilcoxon* για δύο εξαρτημένα δείγματα), οι οποίοι μας εμφανίζουν τους Πίνακες της Εικόνας 10.23.

Ranks				
		N	Mean Rank	Sum of Ranks
middle - intial	Negative Ranks	2 ^a	5,25	10,50
	Positive Ranks	7 ^b	4,93	34,50
	Ties	1 ^c		
	Total	10		
final - intial	Negative Ranks	1 ^d	2,00	2,00
	Positive Ranks	8 ^e	5,38	43,00
	Ties	1 ^f		
	Total	10		
final - middle	Negative Ranks	2 ^g	2,50	5,00
	Positive Ranks	7 ^h	5,71	40,00
	Ties	1 ⁱ		
	Total	10		

a. middle < intial
b. middle > intial
c. middle = intial
d. final < intial
e. final > intial
f. final = intial
g. final < middle
h. final > middle
i. final = middle

Test Statistics ^a			
	middle - intial	final - intial	final - middle
Z	-1,428 ^b	-2,431 ^b	-2,079 ^b
Asymp. Sig. (2-tailed)	,153	,015	,038

a. Wilcoxon Signed Ranks Test
b. Based on negative ranks.

Εικόνα 10.23 Τρεις συνολικά συγκρίσεις μεταξύ τριών μετρήσεων.

Παρατηρούμε από τα παραπάνω ότι απορρίπτεται η μηδενική υπόθεση με επίπεδο σημαντικότητας $\alpha' = 0,0167$ στη σύγκριση της 1^{ης} μέτρησης με την τελευταία, αφού το p (Asymp. Sig) = 0,015 είναι μικρότερο από 0,0167. Επιπλέον, όπως θα δείτε στην υποσημείωση του τελευταίου Πίνακα της Εικόνας 10.23, ο έλεγχος αυτός βασίστηκε στο άθροισμα θέσεων με αρνητικό πρόσημο, επομένως η συνάρτηση ελέγχου εδώ είναι το άθροισμα των θέσεων με αρνητικό πρόσημο, δηλαδή στην περίπτωση μας $T=2$. Το παραπάνω αποτέλεσμα μπορεί να συνταχθεί ως αναφορά στο πλαίσιο μιας επιστημονικής εργασίας ως εξής: Σύμφωνα με τον μη παραμετρικό έλεγχο εξαρτημένων δειγμάτων *Friedman* η επίδοση των μαθητών διαφέρει σε τουλάχιστον δύο μετρήσεις ($\chi^2(2) = 7,784$, $p = 0,020$). Για τον έλεγχο των συγκρίσεων των τριών μετρήσεων, αξιοποιώντας τη διόρθωση

Bonferroni, χρησιμοποιήσαμε σε τρεις περιπτώσεις τον έλεγχο *Wilcoxon*. Από τον έλεγχο αυτό φαίνεται ότι η μέση θέση της επίδοσης των μαθητών της πρώτης μέτρησης (1,40) είναι σημαντικά μικρότερη της μέσης θέσης της επίδοσης των μαθητών της τρίτης μέτρησης (2,60) (*Wilcoxon* $T=2$, $Z=-2,431$, $p=0,015$). Φαίνεται, δηλαδή, ότι η επίδοση των μαθητών στην τελευταία μέτρηση είναι στατιστικά σημαντικά υψηλότερη από την επίδοση των μαθητών της αρχικής μέτρησης.

10.2 Προβλήματα για εξάσκηση και ανατροφοδότηση

Άσκηση 1

Στο αρχείο «*Klasmata.sav*» έχουμε καταχωρίσει τις μετρήσεις των μαθητών της Γ' Δημοτικού σε δύο τεστ, πριν και μετά τη διδακτική παρέμβαση, για τη διδασκαλία των κλασμάτων. Με το πρώτο τεστ (μετρήσεις: πριν1 και μετά1) ελέγχθηκε η εξοικείωση των μαθητών με τη σύνδεση της εικονικής με τη συμβολική αναπαράσταση. Με το δεύτερο τεστ (μετρήσεις: πριν2 και μετά2) ελέγχθηκε η εξοικείωση των μαθητών με τη σύγκριση ομώνυμων και ετερόνυμων κλασμάτων.

Χρησιμοποιώντας μη παραμετρικό τεστ *Wilcoxon* για δύο εξαρτημένα δείγματα, να ελέγξετε κατά πόσο βελτιώθηκε η επίδοση των μαθητών μετά τη διδακτική παρέμβαση σε σχέση με πριν, ξεχωριστά στα δύο τεστ.

- Να τεκμηριώσετε την αξιοποίηση των μη παραμετρικών τεστ.
- Να γράψετε τις αντίστοιχες υποθέσεις.
- Να γράψετε τα συμπεράσματά σας.

Άσκηση 2

Στο αρχείο «*anxiety_teaching_science.sav*» έχουμε καταχωρίσει μετρήσεις του άγχους φοιτητών/-τριών προσχολικής εκπαίδευσης για τη διδασκαλία του μαθήματος των Φυσικών Επιστημών ως μελλοντικοί/ές Νηπιαγωγοί. Επίσης, έχουμε στη διάθεσή μας για τον κάθε φοιτητή/-τρια την κατεύθυνση που είχαν παρακολουθήσει στη Γ' Λυκείου (1 Ανθρωπιστικών, 2 Θετικών σπουδών και 3 Πληροφορικής και Οικονομίας). Χρησιμοποιώντας μη παραμετρικό τεστ *Kruskal-Wallis* για ανεξάρτητα δείγματα, να ελέγξετε κατά πόσο διαφοροποιείται το άγχος των φοιτητών/-τριών προσχολικής εκπαίδευσης για τη διδασκαλία του μαθήματος των Φυσικών Επιστημών στις τρεις αυτές ομάδες (1 Ανθρωπιστικών, 2 Θετικών σπουδών και 3 Πληροφορικής και Οικονομίας).

- Να τεκμηριώσετε την αξιοποίηση των μη παραμετρικών τεστ.
- Να γράψετε τις υποθέσεις.
- Να καταγράψετε τα συμπεράσματά σας.
- Να πραγματοποιήσετε κατά ζεύγη συγκρίσεις του άγχους διδασκαλίας των Φυσικών Επιστημών στις τρεις ομάδες, χρησιμοποιώντας τη διόρθωση *Bonferroni*. Τι παρατηρείτε;

Βιβλιογραφία

- Γιαλαμάς, Β. (2005). *Στατιστικές Τεχνικές και Εφαρμογές στις Επιστήμες της Αγωγής*. Αθήνα: Εκδ. Πατάκη.
- Γναρδέλλης, Χ. (2009). *Ανάλυση δεδομένων με το PASW Statistics 17.0*. Αθήνα: Εκδ. Παπαζήση.
- Field, A. (2018). *Discovering statistics using IBM SPSS statistics* (5th ed.). London: SAGE.
- Lavidas, K., Barkatsas, T., Manesis, D., & Gialamas, V. (2020). A Structural Equation Model Investigating The Impact Of Tertiary Students' Attitudes Toward Statistics, Perceived Competence At Mathematics, And Engagement On Statistics Performance. *Statistics Education Research Journal*, 19(2).
- Sprent, P., & Smeeton, N. C. (2016). *Applied nonparametric statistical methods*. CRC press.

Κεφάλαιο 11

Συσχέτιση και Απλή Γραμμική Παλινδρόμηση

Σύνοψη

Παρουσιάζεται η γραμμική συσχέτιση μεταξύ ποσοτικών μεταβλητών με τη βοήθεια του διαγράμματος σκεδασμού. Δίνεται η ποσοτικοποίηση της κατεύθυνσης και της έντασης μιας γραμμικής συσχέτισης με τον συντελεστή συσχέτισης Pearson r . Η αξιολόγηση του μεγέθους και της στατιστικής σημαντικότητας του r παρουσιάζεται τόσο με τη βοήθεια πινάκων όσο και με τη χρήση του SPSS. Ακόμα, αναλύονται ορισμένα προβλήματα στην ερμηνεία του r που προκύπτουν από την επίδραση των ακραίων τιμών. Εξετάζεται επίσης το υπόδειγμα της απλής γραμμικής παλινδρόμησης που εκφράζεται από την εξίσωση $\hat{Y} = bX + a$ και δίνεται ο τρόπος υπολογισμού της σταθεράς a και της κλίσης b της εξίσωσης. Τέλος, γίνεται αναφορά στη σημαντική συμβολή του συντελεστή στον έλεγχο των ψυχομετρικών ιδιοτήτων των κλιμάκων.

Προαπαιτούμενη γνώση

Ιδιότητες κανονικής κατανομής, μέτρα κεντρικής θέσης και διακύμανσης, έλεγχος υποθέσεων με τις κατανομές F και t : Κεφάλαια 4, 6, 7 και 9 του συγγράμματος.

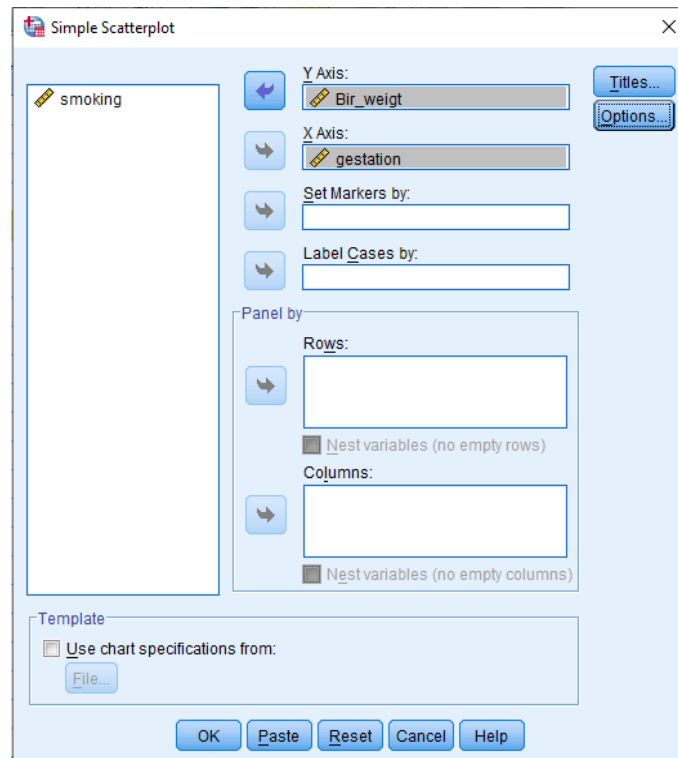
11.1 Εισαγωγή στη συσχέτιση

Η συσχέτιση και ο συντελεστής συσχέτισης (Pearson r) χρησιμοποιείται για την περιγραφή της σχέσης ανάμεσα σε δύο ποσοτικές μεταβλητές. Πολλές φορές οι μεταβλητές προέρχονται από παρατήρηση και δεν είναι αποτέλεσμα κάποιου χειρισμού εκ μέρους του ερευνητή. Για παράδειγμα, αν κάποιος ερευνητής ενδιαφέρεται για τη σχέση ανάμεσα στο άγχος του μαθητή και την επίδοσή του, θα μπορούσε να καταγράψει τους βαθμούς που έδωσε ο δάσκαλος μιας ομάδας μαθητών και να μετρήσει το επίπεδο του άγχους τους με τη βοήθεια ενός ερωτηματολογίου. Σε αυτή την περίπτωση, ο ερευνητής δεν χειρίζεται ούτε την επίδοση ούτε το άγχος των μαθητών, αλλά παρατηρεί αυτό που συμβαίνει εκ του φυσικού. Όπως ήδη αντιλαμβανόμαστε, η συσχέτιση απαιτεί δύο τιμές για κάθε άτομο (μία τιμή για την επίδοση και μία για το άγχος). Αυτές οι τιμές συνήθως συμβολίζονται με X και Y . Τα ζεύγη τιμών μπορεί να δοθούν με τη μορφή ενός πίνακα ή να αναπαρασταθούν γραφικά με ένα *διάγραμμα σκεδασμού* που τοποθετείται πάνω σε ένα σύστημα ορθογωνίων αξόνων. Στον οριζόντιο άξονα παίρνει τιμές η μεταβλητή X και στον κάθετο άξονα παίρνει τιμές η μεταβλητή Y . Κάθε υποκείμενο σημειώνεται, πάνω στο διάγραμμα, με τελεία (ή άλλο σύμβολο), που έχει συντεταγμένες τις τιμές του υποκειμένου για τις δύο μεταβλητές. Συνήθως, πάνω στα διαγράμματα, το «νέφος» των παρατηρήσεων προσομοιώνεται από μία ευθεία γραμμή που έχει την ιδιότητα να περνά πλησιέστερα στα σημεία (παρατηρήσεις) του νέφους από οποιαδήποτε άλλη ευθεία. Αυτή η ευθεία λέγεται *γραμμή παλινδρόμησης*. Με τον τρόπο κατασκευής της θα ασχοληθούμε αργότερα στο παρόν Κεφάλαιο.

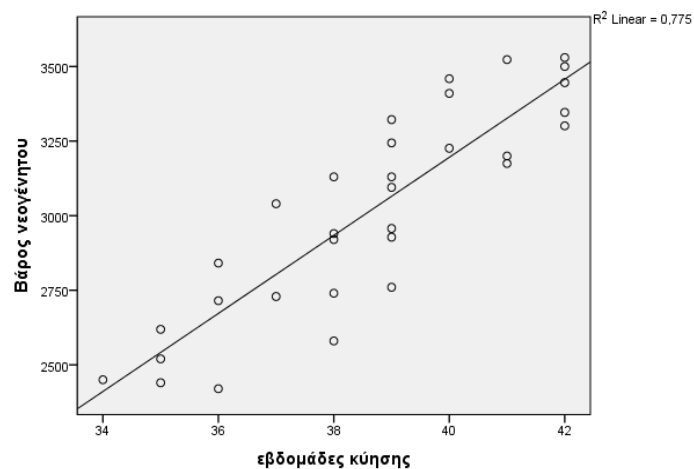
Ενδεικτικά, στην Εικόνα 11.2, δίνεται η σχέση των τιμών του βάρους του νεογνού σε γραμμάρια και της εβδομάδας κύησης κατά την οποία γεννήθηκε, σε ένα τυχαίο δείγμα $N=32$ νεογνών. Για να δημιουργήσουμε το διάγραμμα σκεδασμού, επιλέγουμε διαδοχικά από το μενού **Graphs: Graphs=>Legacy Dialogs=>Scatter/Dot...=>Simple Scatter=>Define**, εισάγουμε στο πλαίσιο **Y Axis** τη μεταβλητή «Birth_Weight» (Βάρος νεογνού) και στο πλαίσιο **X Axis** τη μεταβλητή «gestation» (εβδομάδα κύησης). Έπειτα, πατάμε το εικονίδιο **OK** (Εικόνα 11.1).

Αρχικά, δημιουργείται το διάγραμμα σκεδασμού, το οποίο περιλαμβάνει μόνο τα σημεία - περιπτώσεις και, στη συνέχεια, με κατάλληλη επεξεργασία σχηματίζεται η ευθεία γραμμή (γραμμή παλινδρόμησης), η οποία παρουσιάζεται στην ενότητα της παλινδρόμησης παρακάτω.

Για την προσθήκη της γραμμής παλινδρόμησης μεταφερόμαστε στον επεξεργαστή γραφήματος με διπλό κλικ πάνω στο γράφημα, το οποίο βρίσκεται στο αρχείο αποτελεσμάτων του SPSS και, στη συνέχεια, κάνουμε «κλικ» στην επιλογή **Add Fit Line at Total**, που βρίσκεται στην τρίτη ράβδο εργαλείων. Οι ράβδοι εργαλείων, βρίσκονται τοποθετημένες στο πάνω μέρος του παραθύρου του επεξεργαστή. Επίσης, από το εικονίδιο  (**Insert a Title**) δημιουργούμε τον τίτλο του γραφήματος.



Εικόνα 11.1 Επιλογές μεταβλητών για το διάγραμμα σκεδασμού.



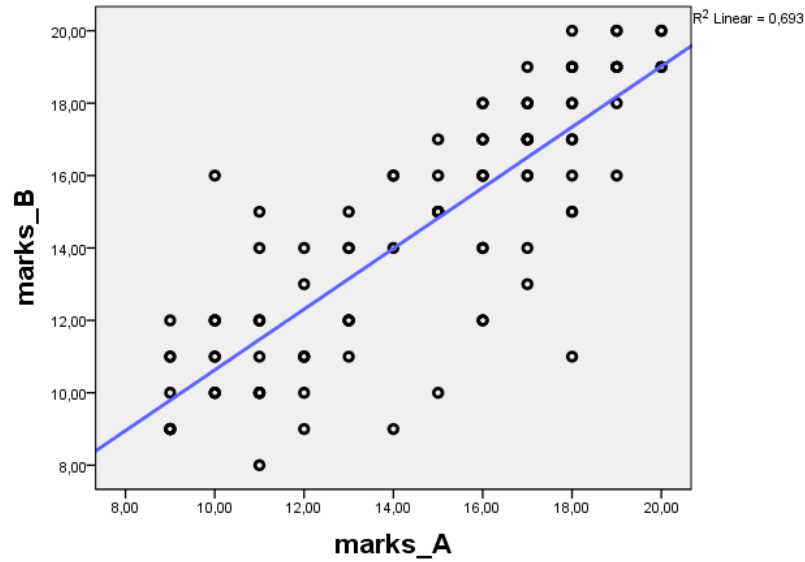
Εικόνα 11.2 Το διάγραμμα σκεδασμού και η γραμμή παλινδρόμησης βάρους του νεογνού και της εβδομάδας κύησης.

11.1.1 Χαρακτηριστικά της συνάφειας

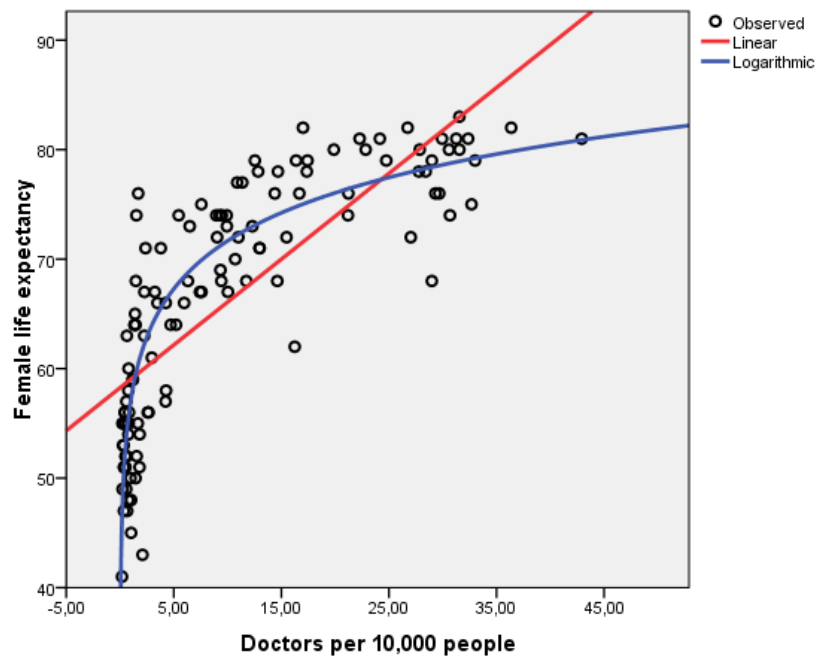
Τα βασικά χαρακτηριστικά της συνάφειας είναι η *μορφή*, ο *βαθμός* και η *διεύθυνση*.

- Η μορφή της συσχέτισης καθορίζεται από τη μορφή της γραμμής που προσεγγίζει καλύτερα τα σημεία του νέφους στο διάγραμμα σκεδασμού. Αν η γραμμή είναι ευθεία, έχουμε *ευθύγραμμη συσχέτιση* ή *απλώς γραμμική συσχέτιση*, ενώ, αν είναι κάποια καμπύλη, η συσχέτιση ονομάζεται *καμπυλόγραμμη*. Υπάρχει η περίπτωση όμως να μην υπάρχει σαφής τύπος γραμμής που να ακολουθείται από τα σημεία του διαγράμματος. Στις Εικόνες 11.3 και 11.5 έχουμε γραμμική συσχέτιση. Αντίθετα, στην Εικόνα 11.4 η καμπύλη που δημιουργείται από τη συνάρτηση του

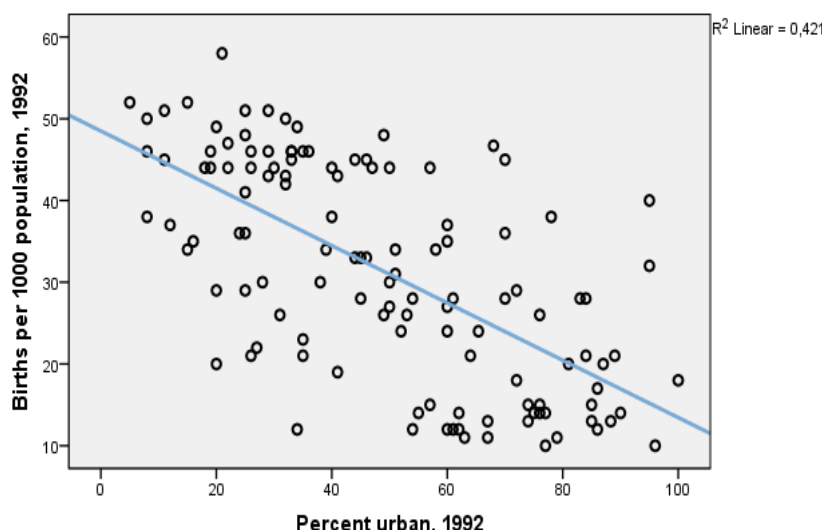
λογάριθμου (μπλε γραμμή) προσεγγίζει σαφώς καλύτερα τα σημεία από εκείνη της ευθείας γραμμής (κόκκινη γραμμή) και η συσχέτιση δεν είναι γραμμική.



Εικόνα 11.3 Συσχέτιση Βαθμών Α' (άξονας Χ) και Β' τριμήνου (άξονας Υ).



Εικόνα 11.4 Προσδόκιο επιβίωσης γυναικών (άξονας Υ) και αριθμός γιατρών σε 10.000 κατοίκους (άξονας Χ) σε ομάδα 122 χωρών.



Εικόνα 11.5 Συσχέτιση ποσοστού αστικοποίησης (άξονας X) και αριθμού γεννήσεων σε 1000 κατοίκους (άξονας Y) σε ομάδα 122 χωρών.

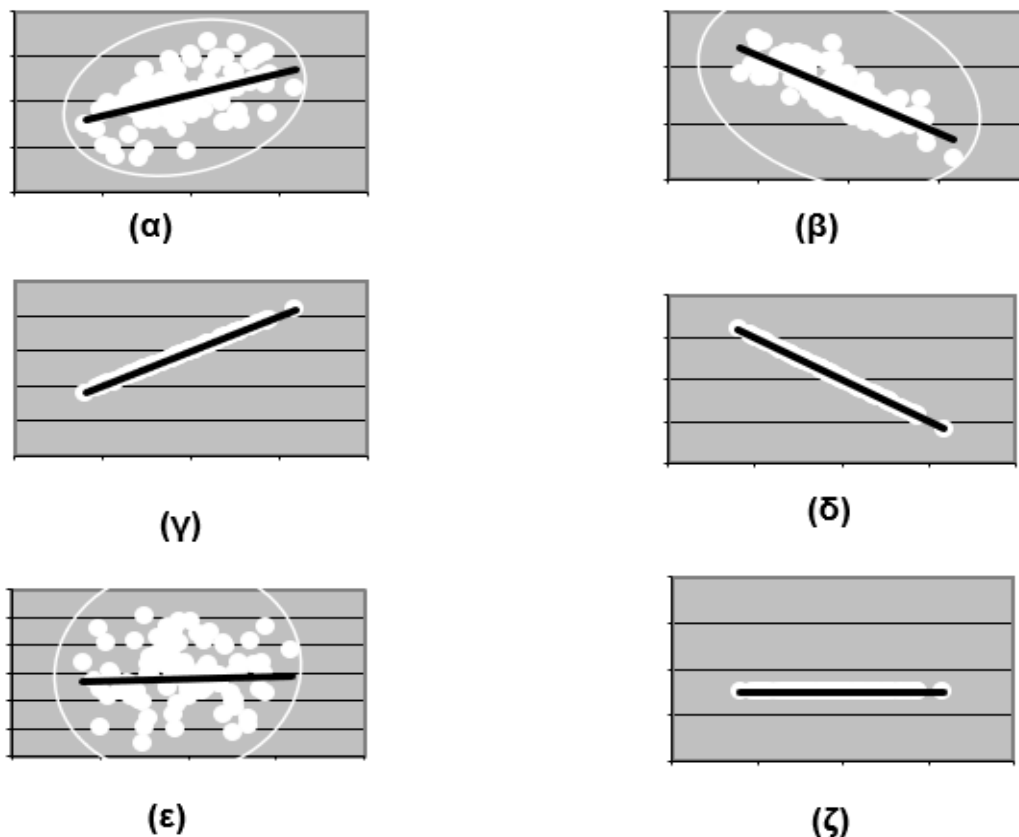
- Σχετικά με τη διεύθυνση της συσχέτισης διακρίνουμε:
 1. *Θετική συσχέτιση* όταν οι μεγαλύτερες τιμές της μιας μεταβλητής αντιστοιχούν σε μεγαλύτερες τιμές της άλλης και όταν οι μικρότερες τιμές της πρώτης αντιστοιχούν σε μικρότερες τιμές της δεύτερης. Η συσχέτιση των βαθμών του Α' τριμήνου με τους βαθμούς του Β' τριμήνου στην Εικόνα 11.3 είναι θετική. Είναι προφανές ότι οι βαθμοί ενός μαθητή είναι γενικώς υψηλοί ή χαμηλοί παραμένοντας στο ίδιο περίπου επίπεδο.
 2. *Αρνητική συσχέτιση* έχουμε όταν οι μεγαλύτερες τιμές της πρώτης αντιστοιχούν στις μικρότερες της δεύτερης μεταβλητής και αντίστροφα, οι μικρότερες τιμές της πρώτης αντιστοιχούν στις μεγαλύτερες τιμές της δεύτερης. Στην Εικόνα 11.5 δίνεται η σχέση ανάμεσα στο ποσοστό αστικοποίησης μιας χώρας και τον αριθμό γεννήσεων στους 1000 κατοίκους της. Η συσχέτιση είναι αρνητική. Αυξημένα ποσοστά αστικοποίησης, που αντιστοιχούν προφανώς σε αναπτυγμένες χώρες, συνδέονται με μικρό αριθμό γεννήσεων. Παρατηρούμε ότι στο 11.6(α), που απεικονίζει τη θετική συσχέτιση, η ευθεία παλινδρόμησης κατευθύνεται από την αρχή των αξόνων προς το άνοιγμα της γωνίας που σχηματίζουν, ενώ στο 11.6(β) γενικά διευθύνεται από το άκρο του άξονα y προς το άκρο του άξονα x .
- Σε ό,τι αφορά τον βαθμό συσχέτισης ανάμεσα στις μεταβλητές X και Y έχουμε δύο οριακές δυνατές περιπτώσεις, αν θεωρήσουμε τις θέσεις των τιμών σύμφωνα με το μέγεθός τους:
 1. Η πρώτη τιμή της X αντιστοιχεί στην πρώτη τιμή της Y , η δεύτερη τιμή της X στη δεύτερη τιμή της Y . . . και, τέλος, η τελευταία τιμή της X στην τελευταία τιμή της Y .
 2. Η πρώτη τιμή της X αντιστοιχεί στην τελευταία τιμή της Y , η δεύτερη τιμή της X στην προτελευταία τιμή της Y . . . και η τελευταία τιμή της X στην πρώτη της Y .

Και στις δύο πιο πάνω περιπτώσεις μιλάμε για *απόλυτη συνάφεια*. Στην πρώτη έχουμε απόλυτη θετική και στη δεύτερη απόλυτη αρνητική. Σε κάθε άλλη ενδιάμεση κατάσταση έχουμε *σχετική συνάφεια*.

Στην περίπτωση της γραμμικής συσχέτισης, αν οι δύο μεταβλητές συνδέονται με συναρτησιακή σχέση της μορφής $Y = a + bX$ η απόλυτη συσχέτιση ερμηνεύεται γραφικά (Εικόνα 11.6 (γ) και 11.6 (δ)) με τέλεια σύμπτωση νέφους παρατηρήσεων και ευθείας παλινδρόμησης. Γραφικά περιπτώσεις σχετικής συνάφειας δίνονται στις Εικόνες 11.6 (α) και 11.6 (β).

Αν δεν υπάρχει σαφής τάση συσώρευσης των παρατηρήσεων γύρω από κάποια ευθεία, αλλά, αντίθετα, τα σημεία κατανέμονται μάλλον στην τύχη σε όλη την περιοχή που περικλείεται από τους άξονες, δεν υπάρχει συνάφεια ανάμεσα στις δύο μεταβλητές (Εικόνα 11.6 ε). Το ίδιο συμβαίνει και στην περίπτωση που η μία

μεταβλητή παραμένει σταθερή, δηλαδή η τυπική της απόκλιση $S=0$ (Εικόνα 11.6 ζ). Στις Εικόνες 11.6 (α), (β) και (ε) έχει σχεδιαστεί μια καμπύλη γύρω από τα σημεία. Τέτοιες γραμμές βοηθούν στο να γίνει κατανοητή η διαγραφόμενη σχέση.



Εικόνα 11.6 Χαρακτηριστικές περιπτώσεις διαγραμμάτων σκεδασμού.

11.1.2 Ο Συντελεστής συσχέτισης Pearson και η αξιολόγηση της τιμής του

Στην ενότητα αυτή θα δοθεί μία αριθμητική έκφραση, όπως ήδη αναφέραμε, αυτή της συσχέτισης ανάμεσα σε δύο ποσοτικές μεταβλητές. Οι ποσότητες που εκφράζουν τη συσχέτιση λέγονται *συντελεστές συσχέτισης*. Στη συνέχεια, θα δοθεί ο έλεγχος υπόθεσης για την τιμή του συντελεστή στον πληθυσμό. Ο πιο γνωστός συντελεστής είναι ο συντελεστής *Pearson*, ο οποίος εκφράζει τον βαθμό, την κατεύθυνση και την ένταση της γραμμικής σχέσης μεταξύ δύο μεταβλητών.

Ο συντελεστής συσχέτισης *Pearson*, που συμβολίζεται με r , εννοιολογικά μπορεί να εκφραστεί ως εξής:

$$r = \frac{\text{βαθμός συµµεταβολής } X \text{ και } Y}{\text{βαθμός } X \text{ και } Y \text{ χωριστά}}$$

Οι τιμές του συντελεστή ανήκουν στο διάστημα $[-1, 1]$. Η τιμή 1 αντιπροσωπεύει την απόλυτη θετική συνάφεια, ενώ η -1, την απόλυτη αρνητική και, τέλος, η τιμή 0 δηλώνει την έλλειψη συνάφειας. Η μεταβλητότητα των « X και Y χωριστά» έχει ήδη μελετηθεί και εκφράζεται από τις τυπικές αποκλίσεις S_X και S_Y .

Για τον υπολογισμό του αριθμητή είναι απαραίτητο να εισαχθεί μια νέα έννοια: το *άθροισμα των γινομένων των αποκλίσεων*. Προηγουμένως έχει χρησιμοποιηθεί η έννοια του αθροίσματος των τετραγώνων των αποκλίσεων για τον υπολογισμό της μεταβλητότητας μιας μεταβλητής. Το άθροισμα των γινομένων των αποκλίσεων SP είναι ένας ανάλογος τρόπος υπολογισμού του βαθμού της συµµεταβολής μεταξύ δύο μεταβλητών.

Ο τύπος ορισμού του είναι:

$$SP = \sum(x - \bar{x}) \cdot (y - \bar{y}) \quad (11.1)$$

Ο τύπος υπολογισμού SP είναι κι αυτός ανάλογος με αυτόν του τύπου υπολογισμού του SS αθροίσματος τετραγώνων των αποκλίσεων και συνήθως οδηγεί σε ευκολότερους υπολογισμούς:

$$SP = \sum xy - \frac{\sum x \sum y}{N} \quad (11.2)$$

Στο πηλίκο που αναφέραμε στον εννοιολογικό ορισμό του *Pearson r* θα χρησιμοποιηθεί το SP στον αριθμητή ως μέτρο της συμμεταβολής και τα SS_x , SS_y ως μέτρα της μεταβολής των μεταβλητών X και Y .

Ο συντελεστής *Pearson r* δίνεται από τη σχέση:

$$r = \frac{SP}{\sqrt{SS_x \cdot SS_y}} \quad (11.3)$$

Ο υπολογιστικός τύπος του r που συνήθως χρησιμοποιείται είναι ο παρακάτω:

$$r = \frac{N(\sum xy) - (\sum x)(\sum y)}{\sqrt{[N(\sum x^2) - (\sum x)^2][N(\sum y^2) - (\sum y)^2]}} \quad (11.4)$$

όπου

$\sum x^2$ = το άθροισμα των τετραγώνων όλων των τιμών x ,

$\sum y^2$ = το άθροισμα των τετραγώνων όλων των τιμών y ,

$\sum xy$ = το άθροισμα των γινομένων των ζευγών,

$(\sum x)^2$ = το τετράγωνο του αθροίσματος των τιμών x ,

$(\sum y)^2$ = το τετράγωνο του αθροίσματος των τιμών y ,

N = ο αριθμός των παρατηρήσεων.

Στη συνέχεια, θα υπολογίσουμε τον συντελεστή *Pearson* για τα δεδομένα που χρησιμοποιήθηκαν στην Εικόνα 11.2. Προκειμένου να μειωθεί το μέγεθος των τιμών στους υπολογισμούς, οι τιμές της μεταβλητής «βάρος» διαιρέθηκαν με το 1000 και έτσι η νέα μονάδα της είναι το κιλό.

Η διαδικασία υπολογισμού των απαιτούμενων αθροισμάτων δίνεται παρακάτω στον Πίνακα 11.1 και η τιμή του συντελεστή συσχέτισης είναι:

$$\begin{aligned} r &= \frac{N(\sum xy) - (\sum x)(\sum y)}{\sqrt{[N(\sum x^2) - (\sum x)^2][N(\sum y^2) - (\sum y)^2]}} \\ r &= \frac{32(3740,9) - (1237)(96,2)}{\sqrt{[32(47987) - (1237)^2][32(292,9) - (96,2)^2]}} \\ &= \frac{119708,8 - 118999,4}{\sqrt{[1535584 - 1530169][9372,8 - 9254,44]}} \\ &= \frac{709,4}{\sqrt{(5415)(118,36)}} \\ r &= 0,886 \end{aligned}$$

Το μέγεθος της τιμής $r=0,88$ δηλώνει πολύ ισχυρή γραμμική θετική συσχέτιση μεταξύ του βάρους του νεογνού και της εβδομάδας κύησης κατά την οποία γεννήθηκε. Σύμφωνα με το *Choen* (1977) η ένταση της σχέσης χαρακτηρίζεται από την απόλυτη τιμή του r ως εξής:

$|r| \leq 0,2$ χαμηλή,

$0,2 < |r| \leq 0,5$ μέτρια,

$0,5 < |r| \leq 0,8$ ισχυρή,

$0,8 < |r| \leq 1$ πολύ ισχυρή.

Ένας πολύ δημοφιλής τρόπος αναφοράς στην ένταση της γραμμικής συσχέτισης μεταξύ δύο μεταβλητών είναι το τετράγωνο του συντελεστή συσχέτισης r^2 , του οποίου οι τιμές ανήκουν στο διάστημα $[0,1]$ ή ισοδύναμα στο διάστημα $[0\%, 100\%]$. Η ερμηνεία του ποσοστού αναφέρεται στο ότι αποτελεί το ποσοστό της διακύμανσης της μίας μεταβλητής που οφείλεται στη διακύμανση των τιμών της άλλης. Στην περίπτωση της συσχέτισης μεταξύ του βάρους του νεογνού και της εβδομάδας κύησης κατά την οποία γεννήθηκε, που υπολογίστηκε παραπάνω, η ποσότητα $r^2=(0,886)^2=0,78$ ή 78%. Συμπεραίνεται ότι ο βασικότερος παράγοντας μεταβλητότητας του βάρους των νεογνών, ο οποίος ερμηνεύει ένα πολύ υψηλό και στατιστικά σημαντικό 78% της διακύμανσης του βάρους τους, είναι η εβδομάδα κύησης κατά την οποία γεννήθηκαν. Η τιμή του r^2 δίνεται από το SPSS πάνω και δεξιά στο περιθώριο του γραφήματος της Εικόνας 11.2. Η μικρή διαφορά που παρατηρείται οφείλεται στις στρογγυλοποιήσεις των τιμών του Πίνακα 11.1.

	x	y	xy	x^2	y^2
	38	2,9	110,2	1444	8,41
	38	3,1	117,8	1444	9,61
	36	2,4	86,4	1296	5,76
	34	2,5	85,0	1156	6,25
	39	2,8	109,2	1521	7,84
	35	2,4	84,0	1225	5,76
	40	3,2	128,0	1600	10,24
	42	3,3	138,6	1764	10,89
	37	2,7	99,9	1369	7,29
	40	3,4	136,0	1600	11,56
	36	2,7	97,2	1296	7,29
	39	3,1	120,9	1521	9,61
	39	3,1	120,9	1521	9,61
	39	3,2	124,8	1521	10,24
	35	2,5	87,5	1225	6,25
	39	2,9	113,1	1521	8,41
	41	3,5	143,5	1681	12,25
	42	3,4	142,8	1764	11,56
	38	2,9	110,2	1444	8,41
	39	3	117,0	1521	9,00
	42	3,5	147,0	1764	12,25
	38	2,6	98,8	1444	6,76
	37	3	111,0	1369	9,00
	42	3,5	147,0	1764	12,25
	41	3,2	131,2	1681	10,24
	39	3,3	128,7	1521	10,89
	40	3,5	140,0	1600	12,25
	42	3,3	138,6	1764	10,89
	35	2,6	91,0	1225	6,76
	41	3,2	131,2	1681	10,24
	38	2,7	102,6	1444	7,29
	36	2,8	100,8	1296	7,84
Σύνολα	1237	96,2	3740,9	47987	292,9

Πίνακας 11.1 Τιμές και αθροίσματα για τον υπολογισμό του συντελεστή συσχέτισης r .

11.1.3 Έλεγχος υποθέσεων για την τιμή του συντελεστή συσχέτισης Pearson r

Για τον υπολογισμό του συντελεστή συσχέτισης r μπορούν να χρησιμοποιηθούν ζεύγη τιμών από δύο ποσοτικά χαρακτηριστικά τόσο από ένα δείγμα όσο και από ολόκληρο πληθυσμό. Στην περίπτωση που τα δεδομένα προέρχονται από ολόκληρο τον πληθυσμό, ο συντελεστής συμβολίζεται με ρ . Όπως συμβαίνει συνήθως, η συσχέτιση που εμφανίζεται σε ένα τυχαίο δείγμα αποτελεί τη βάση προκειμένου να εξαχθούν συμπεράσματα σχετικά με την αντίστοιχη συσχέτιση του πληθυσμού. Για παράδειγμα, η διερεύνηση της σχέσης ανάμεσα στην επίδοση του μαθητή και στη στάση του ως προς την επιτυχία στα σχολικά Μαθηματικά αφορά τον πληθυσμό των μαθητών. Το βασικό ερώτημα είναι αν υπάρχει ή όχι μια συσχέτιση στον πληθυσμό. Με λίγα λόγια, η μηδενική υπόθεση θα αντιπροσωπεύει την πρόταση «Δεν υπάρχει συσχέτιση στον πληθυσμό» και η εναλλακτική την πρόταση «Υπάρχει συσχέτιση στον πληθυσμό». Αυτά μπορούν να γραφτούν με τη βοήθεια του συντελεστή συσχέτισης ρ για τον πληθυσμό ως εξής:

$$H_0: \rho = 0$$

$$H_A: \rho \neq 0$$

Το κατάλληλο κριτήριο ελέγχου της μηδενικής υπόθεσης το οποίο ακολουθεί την κατανομή *student* με $N - 2$ βαθμούς ελευθερίας είναι:

$$t = \frac{r\sqrt{N-2}}{\sqrt{1-r^2}} \quad (11.5)$$

Επειδή η τιμή t εξαρτάται από τους βαθμούς ελευθερίας $N-2$, για τον συντελεστή r δεν χρειάζεται να υπολογιστεί η τιμή t του τύπου (11.5) αφού υπάρχουν Πίνακες (Παράρτημα: Πίνακας Ζ) που δίνουν τις κρίσιμες τιμές r , μετασχηματίζοντας τον Πίνακα των t σε πίνακα r με τη βοήθεια του Τύπου στο (11.5). Τα μοναδικά στοιχεία που απαιτούνται για τη χρήση του Πίνακα είναι οι βαθμοί ελευθερίας και το επίπεδο σημαντικότητας α . Τα επίπεδα α βρίσκονται στην πρώτη γραμμή του Πίνακα. Οι βαθμοί ελευθερίας βρίσκονται στην πρώτη στήλη του Πίνακα. Το σώμα του Πίνακα αποτελείται από τις κρίσιμες τιμές. Για να απορριφθεί η μηδενική υπόθεση και να υιοθετηθεί η εναλλακτική, πρέπει η απόλυτη τιμή της δειγματικής συσχέτισης να είναι ίση ή μεγαλύτερη από την κατάλληλη κρίσιμη τιμή του Πίνακα.

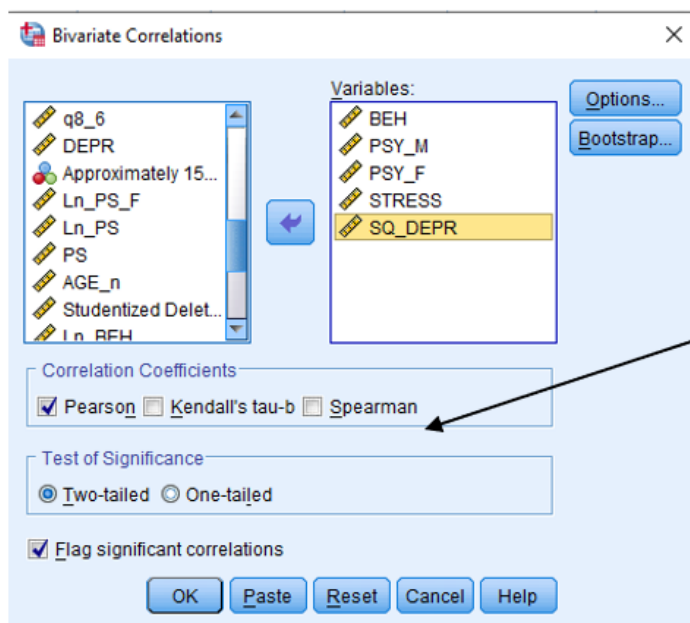
Για τον έλεγχο της συσχέτισης του βάρους του νεογνού και της εβδομάδας κύησης κατά την οποία γεννήθηκε σε επίπεδο $\alpha=0,05$, για $32-2 = 30$ βαθμούς ελευθερίας, από τον Πίνακα βρίσκουμε κρίσιμη τιμή $r_{\text{kp}}=0,349$. Επειδή ο συντελεστής που υπολογίστηκε για το τυχαίο δείγμα είναι ο $r=0,88 > 0,349$, συμπεραίνεται η ύπαρξη στατιστικά σημαντικής ισχυρής θετικής συσχέτισης ανάμεσα στο βάρος του νεογνού και της εβδομάδας κύησης κατά την οποία γεννήθηκε.

11.1.4 Πίνακας συσχετίσεων και έλεγχος υποθέσεων για την τιμή του συντελεστή συσχέτισης Pearson στο SPSS

Σε πολλές μελέτες απαιτείται ο υπολογισμός των ανά δύο συσχετίσεων ενός μεγάλου αριθμού μεταβλητών. Στο πλαίσιο μιας έρευνας για τις σχέσεις μεταξύ των εφήβων και των γονέων τους (Λεονταρή Α., & Γιαλαμάς Β., 2004), αναζητήθηκε η σχέση ανάμεσα στον ψυχολογικό έλεγχο (PSY_F) που ασκεί ο πατέρας (Υψηλή τιμή αντιπροσωπεύει μεγάλο έλεγχο), στον ψυχολογικό έλεγχο (PSY_M) που ασκεί η μητέρα, στον έλεγχο συμπεριφοράς (BEH) που ασκείται από τους δύο γονείς (Υψηλή τιμή αντιπροσωπεύει καλή ενημέρωση και εποπτεία του τρόπου ζωής), στο επίπεδο κατάθλιψης (SQ_DEPR,) του εφήβου (Υψηλές τιμές εκδηλώνουν υψηλό βαθμό κατάθλιψης της/του εφήβου), καθώς και στο επίπεδο άγχους (STRESS) του εφήβου. Όπως διευκρινίζεται στη συνέχεια, οι τιμές της μεταβλητής «SQ_DEPR» έχουν προκύψει από τις τετραγωνικές ρίζες των τιμών της αρχικής μεταβλητής «DEPR», προκειμένου να αποκατασταθεί η απόκλιση από την κανονική κατανομή. Θα κατασκευαστεί ο πίνακας των συντελεστών συσχέτισης *Pearson* (r) μεταξύ των πέντε μεταβλητών και θα συζητηθεί κατά πόσο επαληθεύεται η γενικότερη υπόθεση της θετικής επίδρασης του ελέγχου συμπεριφοράς και της αρνητικής επίδρασης του ψυχολογικού ελέγχου επί των δύο συνιστωσών της ψυχικής υγείας των εφήβων.

Από το βασικό μενού του SPSS επιλέγονται διαδοχικά **Analyze=>Correlate=>Bivariate**. Στη συνέχεια, επιλέγουμε τις πέντε μεταβλητές για τις οποίες θέλουμε τις συσχετίσεις, όπως φαίνεται στην Εικόνα 11.7, και

πατάμε το εικονίδιο **OK**. Ο πίνακας συσχετίσεων που εμφανίζεται στην Εικόνα 11.8 ως ένας Πίνακας ΚxΚ, όπου Κ είναι ο αριθμός των μεταβλητών. Στην περίπτωση μας, σε κάθε κελί του 5x5 Πίνακα μάς δίνονται 3 ποσότητες για το ζεύγος των μεταβλητών στο οποίο αντιστοιχεί ο συντελεστής γραμμικής συσχέτισης *Pearson r* και η παρατηρούμενη πιθανότητα *p* (να απορριφθεί εσφαλμένα η μηδενική υπόθεση $\rho=0$ και το μέγεθος Ν του δείγματος). Από τα στοιχεία του Πίνακα, για παράδειγμα, προκύπτουν σημαντικές συσχετίσεις ($p<0,05$) του επιπέδου κατάθλιψης των εφήβων (SQ_DEPR) τόσο με τις μεταβλητές του ψυχολογικού ελέγχου (PSY_F, PSY_M) όσο και με τον έλεγχο συμπεριφοράς (BEH). Η διεύθυνση της σχέσης, η οποία προσδιορίζεται από το πρόσημο της τιμής του *r*, είναι η αναμενόμενη από τη βιβλιογραφία. Ενδεικτικά, ο συντελεστής συσχέτισης μεταξύ του «BEH» και της «SQ_DEPR» είναι [$r=-0,209$, $p=0,003$] δηλώνοντας μέτριου βαθμού στατιστικά σημαντική αρνητική συσχέτιση. Όσο υψηλότερος ο έλεγχος συμπεριφοράς από τους γονείς απέναντι στον/στην έφηβο/η τόσο χαμηλότερο το επίπεδο της κατάθλιψης του/της.



Οι προκαθορισμένες επιλογές εδώ είναι:

Pearson: ο συντελεστής συσχέτισης *Pearson r*

Two-tailed: αμφίπλευρος έλεγχος

Flag significant correlations: σημειώνεται με κατάλληλα σύμβολα «**» πάνω στις τιμές του *r* το παρατηρούμενο επίπεδο σημαντικότητας.

Εικόνα 11.7 Επιλογές μεταβλητών στο πλαίσιο «Bivariate Correlations».

		Correlations				
		BEH	PSY_M	PSY_F	STRESS	SQ_DEPR
BEH	Pearson Correlation	1	-,286**	-,269**	-,080	-,209**
	Sig. (2-tailed)		,000	,000	,250	,003
	N	207	206	202	207	206
PSY_M	Pearson Correlation	-,286**	1	,504**	,120	,255**
	Sig. (2-tailed)	,000		,000	,086	,000
	N	206	206	201	206	205
PSY_F	Pearson Correlation	-,269**	,504**	1	,272**	,341**
	Sig. (2-tailed)	,000	,000		,000	,000
	N	202	201	202	202	201
STRESS	Pearson Correlation	-,080	,120	,272**	1	,608**
	Sig. (2-tailed)	,250	,086	,000		,000
	N	207	206	202	207	206
SQ_DEPR	Pearson Correlation	-,209**	,255**	,341**	,608**	1
	Sig. (2-tailed)	,003	,000	,000	,000	
	N	206	205	201	206	206

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

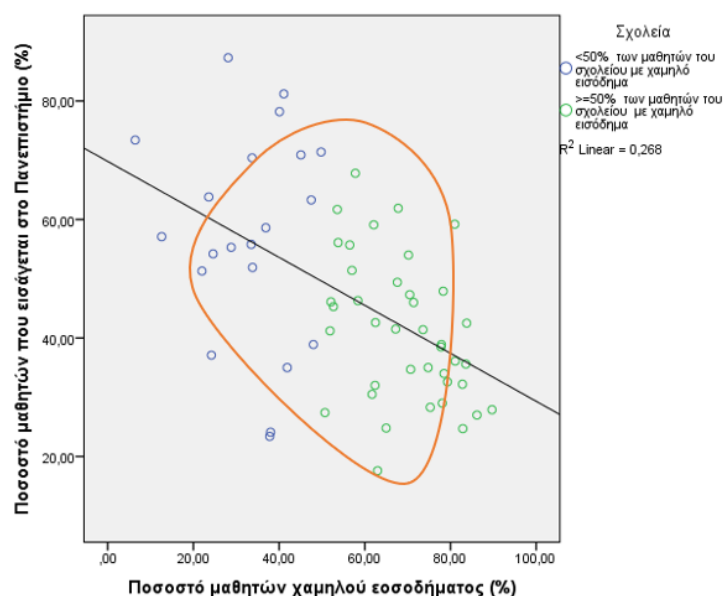
Pearson Correlation: Ο συντελεστής συσχέτισης μεταξύ του «BEH» και της «SQ_DEPR» είναι [$r=-0,209$].
Sig. (2-tailed): Η τιμή *p* για τον έλεγχο των υποθέσεων ($H_0: \rho = 0$ και $H_A: \rho \neq 0$) είναι $0,003 < 0,05$ και συνεπώς απορρίπτεται η H_0 .
N: Το μέγεθος του δείγματος που χρησιμοποιείται για τον έλεγχο αυτής της συσχέτισης είναι 206.

Εικόνα 11.8 Οι Συσχετίσεις (*Pearson r*) και η σημαντικότητά τους.

11.1.5 Ερμηνεία του συντελεστή Pearson

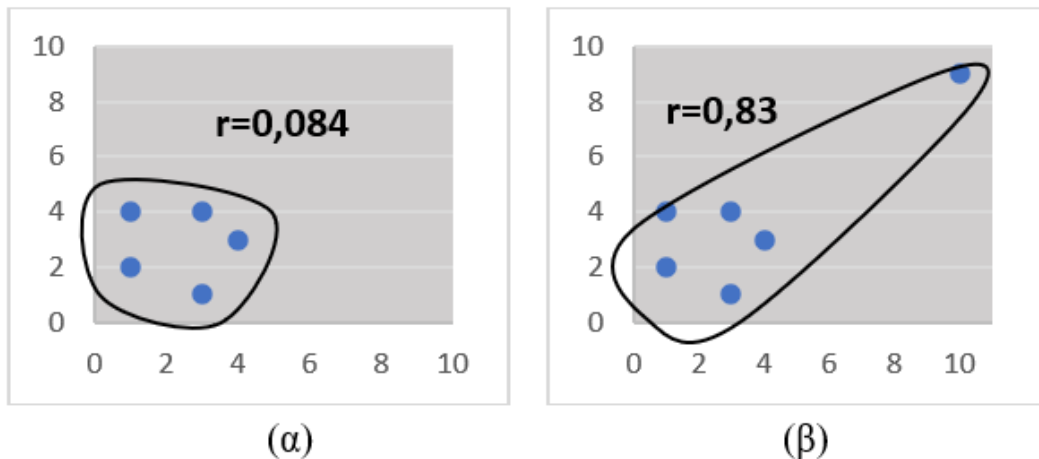
Όταν σε μια μελέτη χρησιμοποιείται ο συντελεστής *Pearson*, πρέπει να δίνεται προσοχή στα παρακάτω:

- Γενικά, ένα από τα πολύ κοινά σφάλματα στην ερμηνεία των συσχετίσεων είναι το συμπέρασμα ότι υπάρχει μια σχέση αιτίου-αποτελέσματος ανάμεσα στις δύο μεταβλητές. Η συσχέτιση περιγράφει τη σχέση μεταξύ δύο μεταβλητών και όχι τον λόγο για τον οποίο αυτές συνδέονται μεταξύ τους. Πολλές φορές ανακοινώνονται διάφορες συσχετίσεις, όπως για παράδειγμα στον τομέα της υγείας: *το κάπνισμα συνδέεται με τον καρκίνο του πνεύμονα*, ή στην εκπαίδευση: *τα έτη εκπαίδευσης του πατέρα σχετίζονται με τον βαθμό του υποψηφίου στις Πανελλαδικές Εξετάσεις*. Μπορεί λοιπόν να εξαχθεί το συμπέρασμα ότι το κάπνισμα προκαλεί καρκίνο του πνεύμονα ή ότι η επιτυχία του μαθητή οφείλεται στο επίπεδο μόρφωσης του πατέρα. Στις παραπάνω περιπτώσεις, αν και η σχέση μπορεί να είναι αιτιώδης, η ύπαρξη της συσχέτισης από μόνη της δεν το αποδεικνύει.
- Η τιμή του συντελεστή συσχέτισης μπορεί να αλλοιώνεται από το εύρος των τιμών στο οποίο ανήκουν τα δεδομένα. Ως παράδειγμα ας δούμε τη συσχέτιση ανάμεσα στο ποσοστό των μαθητών με χαμηλό οικογενειακό εισόδημα και το ποσοστό μαθητών που εισάγονται στα Πανεπιστήμια για ένα δείγμα Λυκείων μιας χώρας. Αν το δείγμα μας είχε περιοριστεί σε συνοικίες που προτιμώνται από τα ανώτερα κοινωνικά στρώματα, είναι λογικό το ποσοστό των μαθητών με χαμηλό εισόδημα να μην ξεπερνάει το 50%. Η συσχέτιση σε ένα τέτοιο εύρος τιμών 0%-50% μπορεί να είναι πολύ διαφορετική από τη συσχέτιση που υπάρχει αν χρησιμοποιούσαμε όλο το εύρος των ποσοστών μαθητών 0%-100% με χαμηλό εισόδημα. Στην Εικόνα 11.9 εμφανίζεται μια αρκετά ισχυρή αρνητική σχέση ($r = -0,51$) μεταξύ των δύο μεταβλητών που μελετώνται, όταν η δειγματοληψία συμπεριλαμβάνει μαθητές από όλες τις συνοικίες. Η συσχέτιση είναι περίπου ανύπαρκτη ($r = -0,03$), αν περιοριστούμε σε εύρος ποσοστών χαμηλού εισοδήματος 0%-50%, όπως φαίνεται και από τη μορφή της καμπύλης που περιλαμβάνει αυτά τα σημεία-Λύκεια στην Εικόνα 6.9.



Εικόνα 11.9 Συσχέτιση με περιορισμένο και με ολόκληρο το εύρος των τιμών της.

- Η ύπαρξη ακραίων ατόμων-σημείων μπορεί να επιδράσει έντονα την τιμή της συσχέτισης. Ακραίο θεωρείται το άτομο που έχει τιμή (στη μεταβλητή X και/ή τη μεταβλητή Y) σημαντικά μεγαλύτερη ή μικρότερη από τις τιμές των άλλων ατόμων. Αυτό το φαινόμενο παρουσιάζεται στην Εικόνα 11.9. Στην Εικόνα 11.10(α) έχουμε μια συσχέτιση μεταξύ των X και Y κοντά στο 0 ($r = 0,084$). Στην Εικόνα 11.10(β), όπου έχει προστεθεί ένα ακραίο σημείο, οι δύο μεταβλητές παρουσιάζουν μια ισχυρή θετική συσχέτιση ($r = 0,83$).



Εικόνα 11.10 Επίδραση μια ακραίας τιμής στον συντελεστή r .

11.1.6 Πού χρησιμοποιείται ο συντελεστής γραμμικής συσχέτισης;

Αν και οι συντελεστές συσχέτισης εφαρμόζονται σε έναν μεγάλο αριθμό περιπτώσεων, θα παρουσιαστούν οι συχνότερα χρησιμοποιούμενες εφαρμογές τους:

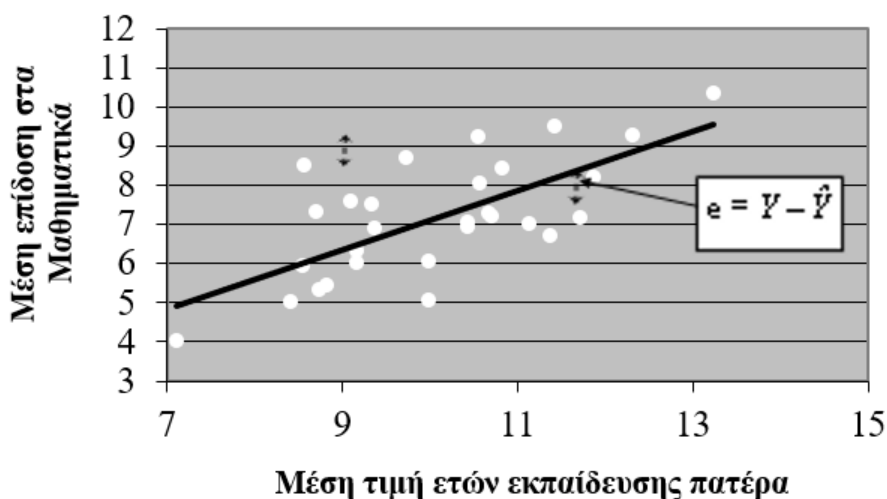
- **Εγκυρότητα:** Ένας εκπαιδευτικός δημιούργησε μια δοκιμασία στο μάθημα της Ιστορίας, με στόχο τη χρησιμοποίησή της σε έναν πειραματισμό κατά τον οποίο ελέγχεται η αποτελεσματικότητα μιας νέας διδακτικής προσέγγισης του μαθήματος. Κατά πόσο όμως μπορεί να υποστηριχτεί ότι αυτή η δοκιμασία αξιολογεί αυτό που υποστηρίζεται ότι αξιολογεί; Η συσχέτιση των βαθμών στη δοκιμασία με τους βαθμούς στο μάθημα της Ιστορίας που δίνει ο δάσκαλος των μαθητών, είναι μια τεχνική που χρησιμοποιείται για να αναδείξει την εγκυρότητα. Πράγματι, αν η νέα δοκιμασία αξιολογεί την ικανότητα στην Ιστορία, αναμένεται ισχυρή θετική συσχέτιση μεταξύ των βαθμών της δοκιμασίας και της επίδοσης στην Ιστορία που μετρείται από τον δάσκαλο των μαθητών (Εικόνα 11.6(α)).
- **Αξιοπιστία:** Μια διαδικασία μέτρησης θεωρείται αξιόπιστη, όταν παράγει τις ίδιες ή σχεδόν τις ίδιες τιμές, όταν τα ίδια άτομα μετριοούνται κάτω από τις ίδιες συνθήκες. Για παράδειγμα, αν μετρηθεί με την τιμή 16 η ακαδημαϊκή αυτοαντίληψη του μαθητή σήμερα, με τη βοήθεια ενός ερωτηματολογίου, αναμένεται μια τιμή 16 ή πολύ κοντά στο 16 στη μέτρηση της επόμενης εβδομάδας, με τη συμπλήρωση του ίδιου ερωτηματολογίου. Ένας τρόπος αξιολόγησης της αξιοπιστίας είναι να χρησιμοποιηθεί η συσχέτιση μεταξύ δύο ομάδων μετρήσεων. Όταν η αξιοπιστία είναι υψηλή, η συσχέτιση ανάμεσα σε δύο διαφορετικές μετρήσεις είναι θετική και ισχυρή (Εικόνα 11.6(α)).
- **Επαλήθευση της θεωρίας:** Πολλές θεωρίες προβλέπουν σχέσεις μεταξύ δύο μεταβλητών. Για παράδειγμα, ένας Κοινωνικός Ψυχολόγος ανέπτυξε μια θεωρία που υποστηρίζει τη σχέση ανάμεσα στον τύπο της προσωπικότητας και την κοινωνική συμπεριφορά σε μια κοινωνική κατάσταση. Ένας Παιδαγωγός προβλέπει την ύπαρξη σχέσης ανάμεσα στο πλήθος των μαθητών της τάξης και τη μέση επίδοσή τους (Εικόνα 11.6(α)). Μια θεωρία της Σχολικής Ψυχολογίας προβλέπει την ύπαρξη αρνητικής σχέσης μεταξύ άγχους και ακαδημαϊκής αυτοαντίληψης των μαθητών (Εικόνα 11.6(β)). Σε καθεμία από αυτές τις περιπτώσεις η θεωρία μπορεί να ελεγχθεί με τον καθορισμό της συσχέτισης των δύο μεταβλητών.
- **Πρόβλεψη:** Αν η σχέση μεταξύ δύο μεταβλητών είναι συστηματική, είναι πιθανόν να χρησιμοποιηθεί η μία μεταβλητή, για να γίνουν ακριβείς προβλέψεις των τιμών της άλλης. Για παράδειγμα, είναι γνωστή η σχέση ανάμεσα στους βαθμούς των Πανελλαδικών Εξετάσεων και τους βαθμούς των πρωτοετών φοιτητών. Ο μαθητής που έχει υψηλό βαθμό στις Πανελλαδικές αναμένεται να έχει υψηλό μέσο όρο στο τέλος του πρώτου έτους στο Πανεπιστήμιο. Αντίθετα, ο μαθητής που δεν απέδωσε ικανοποιητικά στις Πανελλαδικές αναμένεται να έχει δυσκολίες στο πρώτο έτος των σπουδών του. Οι υπεύθυνοι των Πανεπιστημίων σε πολλές χώρες

χρησιμοποιούν τον βαθμό του μαθητή στις τελικές εξετάσεις του Λυκείου για να προβλέψουν τη δυνατότητα επιτυχίας του στο Πανεπιστήμιο. Περισσότερα για την πρόβλεψη θα αναφερθούν στη συνέχεια του Κεφαλαίου.

11.2 Εισαγωγή στην παλινδρόμηση

Προηγουμένως, έχει παρουσιαστεί ο συντελεστής συσχέτισης *Pearson* ως η ποσότητα που περιγράφει τη γραμμική συσχέτιση μεταξύ δύο ποσοτικών μεταβλητών. Εδώ, θα εξεταστεί ένα πρόβλημα πρόβλεψης της τιμής μιας μεταβλητής με τη βοήθεια της τιμής μιας άλλης, όταν οι δύο μεταβλητές είναι ποσοτικές και η σχέση τους είναι γραμμική.

Ένας καθηγητής Μαθηματικών, προκειμένου να οδηγηθεί σε κατάλληλη διδακτική επιλογή, πρέπει να διαπιστώσει το επίπεδο των μαθητών του σε μια τάξη Α΄ Γυμνασίου, πριν ξεκινήσουν τα μαθήματα. Ο καθηγητής διαθέτει τα δημογραφικά στοιχεία των νέων μαθητών του και τα αποτελέσματα μιας έρευνας (Barkatsas, Gialamas, Karageorgos, & Kasimatis, 1998) που συνδέει το μέσο μορφωτικό επίπεδο του πατέρα με τη μέση βαθμολογία σε ένα τυποποιημένο διαγώνισμα για ένα τυχαίο δείγμα 29 τμημάτων Α΄ Γυμνασίου. Το διάγραμμα σκεδασμού (Εικόνα 11.11) περιγράφει αυτή τη σχέση που προκύπτει από τα στοιχεία της έρευνας. Διαπιστώνεται μια γραμμική και ισχυρά θετική συσχέτιση. Έχει σχεδιαστεί μια ευθεία γραμμή στο γράφημα που περνά ανάμεσα από τα σημεία- τμήματα.



Εικόνα 11.11. Συσχέτιση ετών εκπαίδευσης πατέρα και επίδοσης στα Μαθηματικά.

Η γραμμή αυτή διευκολύνει στην κατανόηση της σχέσης ανάμεσα στο μορφωτικό επίπεδο (σε έτη εκπαίδευσης) και την επίδοση στα Μαθηματικά και εκφράζει τη γενική ή κεντρική τάση τοποθέτησης των σημείων, όπως η μέση τιμή εκφράζει την κεντρική θέση μιας ομάδας τιμών. Η γραμμή αυτή θα μπορούσε να χρησιμοποιηθεί για πρόβλεψη. Η γραμμή εκφράζει μια συγκεκριμένη σχέση μεταξύ μορφωτικού επιπέδου του πατέρα και της αντίστοιχης επίδοσης στα Μαθηματικά. Έτσι ο καθηγητής θα μπορούσε να χρησιμοποιήσει τη γραμμή, για να προβλέψει ότι η τάξη του με μέση τιμή ετών εκπαίδευσης του πατέρα 9,5 αναμένεται να αντιστοιχεί σε μια επίδοση περίπου 6,7 που είναι κοντά στη μέση επίδοση όλων των τάξεων ($\bar{x} = 7,17$ και $s = 1,5$).

Ο στόχος της ανάλυσης παλινδρόμησης είναι να βρεθεί μία διαδικασία που να προσδιορίζει μία ευθεία γραμμή, η οποία εκφράζει κάθε φορά την καλύτερη προσαρμογή για οποιαδήποτε ομάδα δεδομένων. Όπως είναι αντιληπτό, δεν είναι απαραίτητο να σχεδιαστεί η ευθεία γραμμή, αφού είναι γνωστό από τα Μαθηματικά ότι σε κάθε ευθεία μπορεί να αντιστοιχηθεί μια απλή εξίσωση της μορφής $Y = bX + a$ που εκφράζει τη γραμμική σχέση μεταξύ των μεταβλητών X και Y . Τα a και b εκφράζουν τις σταθερές της εξίσωσης. Το b εκφράζει την κλίση της ευθείας, δηλαδή το πόσο μεταβάλλεται η μεταβλητή Y , όταν η μεταβλητή X μεταβάλλεται κατά μία μονάδα. Η σταθερά a εκφράζει τη διατομή της γραμμής, δηλαδή την τιμή της μεταβλητής Y , όταν $X = 0$.

Αν, για παράδειγμα, χρησιμοποιηθεί η εξίσωση $Y = 69 - 0,4X$, για να περιγράψει τη σχέση των μεταβλητών της Εικόνας 11.9, όπου Y είναι το ποσοστό μαθητών που εισάγονται στο Πανεπιστήμιο και X το ποσοστό μαθητών με χαμηλό οικογενειακό εισόδημα ενός Σχολείου, η διατομή 69 (69%) εκφράζει το ποσοστό των μαθητών του Σχολείου που εισάγονται σε Πανεπιστήμιο, αν το ποσοστό μαθητών χαμηλού εισοδήματος

σε αυτό είναι 0%. Η κλίση 0,4% εκφράζει το ποσοστό κατά το οποίο ελαττώνεται το ποσοστό μαθητών που εισάγεται στο Πανεπιστήμιο, όταν το ποσοστό των μαθητών του Σχολείου με χαμηλό εισόδημα αυξηθεί κατά μία ποσοστιαία μονάδα.

11.2.1 Η ευθεία των ελαχίστων τετραγώνων

Επειδή, για ένα σύνολο σημείων, τα οποία φαίνεται να εκτείνονται κατά μήκος μιας ευθείας γραμμής, θα μπορούσαν να χαραχθούν πολλές γραμμές που εκφράζουν αυτή τη γενική τάση των σημείων, έχουν αναπτυχθεί στατιστικές μέθοδοι που προσδιορίζουν μία συγκεκριμένη ευθεία που διέρχεται όσο το δυνατόν πλησιέστερα απ' όλα τα σημεία. Με λίγα λόγια, έχει τη βέλτιστη προσαρμογή σύμφωνα με ένα μαθηματικό κριτήριο. Η ευθεία που προσδιορίζεται με αυτές τις μεθόδους ονομάζεται *γραμμή παλινδρόμησης*. Μία τεχνική που συμπεριλαμβάνει μία από αυτές τις μεθόδους ονομάζεται *ανάλυση παλινδρόμησης*.

Για τον καθορισμό της έννοιας της καλής προσαρμογής μιας γραμμής πάνω στα σημεία ενός διαγράμματος σκεδασμού, χρειάζεται να ορισθεί η απόσταση ενός σημείου από τη γραμμή. Για κάθε τιμή X των δεδομένων με τη χρήση της εξίσωσης ή της ευθείας που εκφράζεται από την Y και συμβολίζεται με $\hat{Y}$, η απόσταση μεταξύ προβλεπόμενης και πραγματικής τιμής Y ορίζεται από τη διαφορά $e = Y - \hat{Y}$, που ονομάζεται *υπόλοιπο* (residual).

Στο παράδειγμα της Εικόνας 11.11 δίνεται γραφικά η απόσταση e που εκφράζει το σφάλμα (υπόλοιπο) που γίνεται στην πρόβλεψη μιας τιμής, όταν χρησιμοποιείται η εξίσωση αυτής της γραμμής. Επειδή οι τιμές του e μπορεί να είναι θετικές ή αρνητικές, ανάλογα με τη θέση του σημείου πάνω ή κάτω από τη γραμμή, χρησιμοποιείται το τετράγωνο e^2 της απόστασης e . Έτσι, το άθροισμα των τετραγώνων των σφαλμάτων e^2 εκφράζει κατάλληλα το συνολικό σφάλμα ανάμεσα στα δεδομένα και τη γραμμή:

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

Είναι λογικό να αναζητηθεί ως βέλτιστη λύση η γραμμή η οποία έχει το μικρότερο άθροισμα τετραγώνων των σφαλμάτων SS_e . Η γραμμή αυτή ονομάζεται *γραμμή των ελαχίστων τετραγώνων*. Η εξίσωση της ευθείας με την παραπάνω ιδιότητα δίνει για κάθε τιμή της X την καλύτερη πρόβλεψη $\hat{Y}$ και έχει τη μορφή:

$$\hat{y} = bx + a.$$

Το πρόβλημα είναι να βρεθούν οι συγκεκριμένες τιμές των συντελεστών a και b , ώστε να εκφράζεται η ιδιότητα της ευθείας των ελαχίστων τετραγώνων. Οι πολύπλοκοι υπολογισμοί που οδηγούν στο αποτέλεσμα δεν παρουσιάζονται εδώ. Τα αποτελέσματα όμως των υπολογισμών δίνονται παρακάτω:

$$b = \frac{SP}{SS_X}, \quad (11.6)$$

$$a = \bar{y} - b\bar{x}, \quad (11.7)$$

όπου

$$SP = \sum (x - \bar{x}) \cdot (y - \bar{y}) \text{ και}$$

$$SS_X = \sum (x - \bar{x})^2.$$

Ο υπολογιστικός τύπος που δίνει την κλίση b της ευθείας και διευκολύνει τους υπολογισμούς είναι:

$$b = \frac{N(\sum xy) - (\sum x)(\sum y)}{[N(\sum x^2)] - (\sum x)^2}. \quad (11.8)$$

Ως παράδειγμα, θα υπολογίσουμε τους συντελεστές a και b της εξίσωσης του προβλήματος που παρουσιάστηκε στην εισαγωγή της παραγράφου, τα δεδομένα του οποίου προαναφερθέντος προβλήματος δίνονται στον Πίνακα 11.3. Με τη χρήση του Τύπου στο (11.8) βρίσκουμε:

$$\begin{aligned} b &= \frac{N(\sum xy) - (\sum x)(\sum y)}{[N(\sum x^2)] - (\sum x)^2} = \frac{29(2133,14) - (292,15)(207,8)}{[29(2995,96)] - (292,15)^2} \\ &= \frac{61861,06 - 60708,77}{86882,84 - 85351,62} = \frac{1152,29}{1531,22} \\ &= 0,75 \end{aligned}$$

Στη συνέχεια, από τον Τύπο στο (11.7) υπολογίζεται η τιμή της διατομής a :

$$a = \bar{y} - b\bar{x} = \frac{207,80}{29} - 0,75 \cdot \frac{292,15}{29} = -0,39.$$

Η εξίσωση της ευθείας που εκφράζει καλύτερα τα δεδομένα είναι η εξής:

$$\hat{y} = 0,75x - 0,39.$$

x	x ²	y	y ²	xy	x	x ²	y	y ²	xy
9,35	87,42	7,52	56,55	70,31	10,67	113,85	7,27	52,85	77,57
9,11	82,99	7,57	57,30	68,96	9,17	84,09	6,00	36,00	55,02
9,74	94,87	8,68	75,34	84,54	8,41	70,73	5,00	25,00	42,05
9,38	87,98	6,88	47,33	64,53	10,71	114,70	7,18	51,55	76,90
8,82	77,79	5,44	29,59	47,98	8,57	73,44	8,50	72,25	72,85
11,38	129,50	6,68	44,62	76,02	11,43	130,64	9,50	90,25	108,59
13,24	175,30	10,35	107,12	137,03	7,11	50,55	4,00	16,00	28,44
12,32	151,78	9,28	86,12	114,33	8,55	73,10	5,94	35,28	50,79
10,58	111,94	8,04	64,64	85,06	8,75	76,56	5,31	28,20	46,46
10,43	108,78	7,04	49,56	73,43	11,88	141,13	8,21	67,40	97,53
8,70	75,69	7,30	53,29	63,51	10,00	100,00	6,04	36,48	60,40
11,72	137,36	7,14	50,98	83,68	9,17	84,09	6,33	40,07	58,05
10,56	111,51	9,23	85,19	97,47	10,83	117,29	8,42	70,90	91,19
11,14	124,10	7,00	49,00	77,98	10,00	100,00	5,04	25,40	50,40
10,43	108,78	6,91	47,75	72,07	10,67	113,85	7,27	52,85	77,57
9,35	87,42	7,52	56,55	70,31	9,17	84,09	6,00	36,00	55,02
Σύνολα	292,15	2995,96	207,80	1552,01	2133,14				

Πίνακας 11.3 Υπολογισμός εξίσωσης παλινδρόμησης ανάμεσα στη μέση τιμή ετών εκπαίδευσης του πατέρα (X) και στη μέση επίδοση στα Μαθηματικά (Y) για 29 τμήματα της Α' Γυμνασίου.

11.2.2 Η παλινδρόμηση και η ευθεία των ελαχίστων τετραγώνων με το SPSS

Προκειμένου να υπολογιστούν η κλίση και η διατομή της ευθείας παλινδρόμησης αλλά και να υπολογιστεί η προσαρμοστικότητα της στα δεδομένα, με τη βοήθεια του SPSS θα εκτελεστεί ως παράδειγμα η παλινδρόμηση της μέσης επίδοσης της τάξης πάνω στο μέσο μορφωτικό επίπεδο του πατέρα των μαθητών της τάξης αυτής. Τα ζεύγη των τιμών δίνονται στον Πίνακα 11.3 της προηγούμενης υποενότητας.

Από το βασικό μενού του SPSS επιλέγονται διαδοχικά **Analyze=>Regression=>Linear** και στη συνέχεια επιλέγουμε τη μεταβλητή «ED_YR_F» (μέση εκπαίδευση πατέρα των μαθητών της τάξης) και «grade_1» (μέσος βαθμός τάξης), στα πλαίσια **Independent(s)** και **Dependent**, αντίστοιχα. Στην Εικόνα 11.12 δίνεται το παράθυρο της **Linear Regression** με τις προαναφερόμενες επιλογές μας. Στους Πίνακες 11.4 έως 11.6 δίνονται τα αποτελέσματα της παλινδρόμησης.

Στον Πίνακα 11.4 εξετάζεται ο βαθμός προσαρμογής του υποδείγματος της γραμμής παλινδρόμησης. Η ποσότητα R είναι η απόλυτη τιμή του συντελεστή συσχέτισης r . Ο R -square εκφράζει μια πολύ δημοφιλή αναλογία, η οποία προαναφέρθηκε στην παρουσίαση του συντελεστή r . Πρόκειται για το ποσοστό της μεταβλητότητας της εξαρτημένης μεταβλητής που ερμηνεύεται από το υπόδειγμα.

Επειδή έχουμε μία μόνο μεταβλητή στο δεξί μέρος της εξίσωσης μπορεί να ειπωθεί ότι ο R^2 (47,3% στο παράδειγμά μας) είναι το ποσοστό της μεταβλητότητας της εξαρτημένης μεταβλητής (μέσος βαθμός τάξης), που εξηγείται από την ανεξάρτητη μεταβλητή (μέση εκπαίδευση του πατέρα των μαθητών της τάξης). Ο διορθωμένος συντελεστής R^2 (Adjusted R square) εκφράζει την εκτίμηση της προσαρμογής του υποδείγματος σε ένα άλλο τυχαίο δείγμα από τον ίδιο πληθυσμό. Επειδή η γραμμή παλινδρόμησης έχει υπολογιστεί με βάση τα δεδομένα του δείγματός μας, αναμένεται το υπόδειγμα να προσαρμόζεται κάπως καλύτερα στο δείγμα που χρησιμοποιήθηκε για τη δημιουργία του απ' ό,τι σε ένα άλλο δείγμα. Συνεπώς, η τιμή του διορθωμένου R^2 θα είναι πάντα μικρότερη από την τιμή του R^2 . Η ποσότητα *Std. error of the Estimate* εκφράζει τη μέση απόκλιση

από την εκτιμώμενη τιμή της εξαρτημένης μεταβλητής που προκύπτει από την εξίσωση παλινδρόμησης ή την τυπική απόκλιση $SEE = \sqrt{\frac{e^2}{N-2}}$ των υπολοίπων $e = Y - \hat{Y}$, των οποίων η μέση τιμή είναι ίση με 0. Όσο μικρότερη είναι η τυπική απόκλιση των υπολοίπων, τόσο καλύτερη η προσαρμογή του υποδείγματος.

Καθώς η μεταβλητότητα της επίδοσης της τάξης (sum of Squares, Total=63,022) μπορεί να διασπαστεί σε ένα μέρος που οφείλεται στο μορφωτικό επίπεδο του πατέρα (sum of Squares, Regression = 29,811) και σε ένα δεύτερο μέρος (sum of Squares, Residual = 33,210), το οποίο δεν εξηγείται από το μορφωτικό επίπεδο του πατέρα, είναι προφανές ότι για να υπολογίσουμε το ποσοστό της μεταβλητότητάς της που εξηγείται από την παλινδρόμηση θα πρέπει να εκτελέσουμε τη διαίρεση 29,811/63,022. Το αποτέλεσμα είναι 47,3%. Οι διαφορές μεταξύ των τάξεων που δεν εξηγούνται από το υπόδειγμα της παλινδρόμησης είναι το 100-47,3=52,7% της συνολικής μεταβλητότητας της μέσης επίδοσης. Τα τελευταία ποσοστά είναι ακριβώς ίδια με εκείνα που δίνει ο R^2 .

Το πηλίκο $F = (\text{mean Square, regression}) / (\text{mean Square, residual})$ ακολουθεί κατανομή F με 1 και $N-2$ βαθμούς ελευθερίας. Ο έλεγχος της μηδενικής υπόθεσης ότι η τιμή του ρ^2 είναι 0 στον πληθυσμό γίνεται με το πηλίκο F και είναι απολύτως ισοδύναμος με τον έλεγχο για την τιμή του συντελεστή ρ που είδαμε παραπάνω στη γραμμική συσχέτιση με την κατανομή t . Οι δύο έλεγχοι δίνουν την ίδια p . Στο παράδειγμά μας, απορρίπτεται ισχυρά η μηδενική υπόθεση, καθώς από τον Πίνακα 11.5 έχουμε $F(1,27)=24,237$; $p<0,001$.

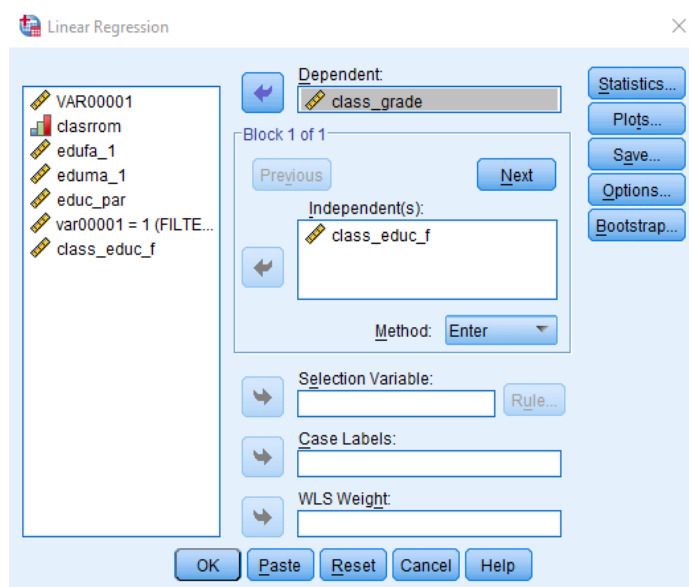
Στον Πίνακα 11.6 υπολογίζονται οι συντελεστές της εξίσωσης, οι οποίοι βρίσκονται στη στήλη B. Στο κελί (B, constant) βρίσκεται η σταθερά της εξίσωσης και στο κελί (B, CLASS_EDUC_f) βρίσκεται ο συντελεστής της μεταβλητής «CLASS_EDUC_F» ή κλίση της ευθείας παλινδρόμησης.

Η εξίσωση για την προβλεπόμενη μέση επίδοση της τάξης μπορεί να γραφτεί:

$$CLASS_GRADE = -0,399 + 0,751 \times CLASS_EDUC_F.$$

Οι τιμές των συντελεστών B του πίνακα είναι ικανοποιητικές εκτιμήσεις των συντελεστών της εξίσωσης, η οποία εξηγεί την επίδοση στον πληθυσμό των τάξεων και ως εκ τούτου για κάθε συντελεστή μπορεί να ελεγχθεί η μηδενική υπόθεση $\beta=0$.

Όμως η αναφορά στους ελέγχους θα γίνει παρακάτω στη γενικότερη μορφή της παλινδρόμησης που είναι η πολλαπλή παλινδρόμηση. Στη στήλη «Beta» δίνεται ο τυποποιημένος συντελεστής για την ανεξάρτητη μεταβλητή που εκφράζει την τιμή του B, αν η παλινδρόμηση εκτελεστεί με τυποποιημένες τις τιμές (αφαιρώντας τη μέση τιμή και διαιρώντας με την τυπική απόκλιση της μεταβλητής) των δύο μεταβλητών. Μπορεί να παρατηρηθεί ότι η τιμή του τυποποιημένου συντελεστή στην απλή παλινδρόμηση δεν είναι άλλη από την τιμή 0,688 του συντελεστή συσχέτισης *Pearson r*. Η τυποποίηση είναι απαραίτητη για να δοθεί το μέγεθος της επίδρασης, αφού η τιμή B μοιράζεται την ίδια μονάδα μέτρησης με την εξαρτημένη μεταβλητή και κατά συνέπεια δεν είναι κατάλληλη γι' αυτό.



Εικόνα 11.12 Επιλογή μεταβλητών στο παράθυρο της «Γραμμικής Παλινδρόμησης».

Model Summary				
Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	,688 ^a	,473	,454	1,10906
a. Predictors: (Constant), class_educ_f				

Υποσημείωση: R: Η απόλυτη τιμή του συντελεστή συσχέτισης, R Square: το τετράγωνο του συντελεστή (συντελεστής προσδιορισμού), Adjusted R square: Διορθωμένος συντελεστής προσδιορισμού, Std. Error of the Estimate: το τυπικό σφάλμα της εκτίμησης.

Πίνακας 11.4 Δείκτες προσαρμογής του υποδείγματος.

ANOVA ^a						
Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	29,811	1	29,811	24,237	,000 ^b
	Residual	33,210	27	1,230		
	Total	63,022	28			

a. Dependent Variable: class_grade

b. Predictors: (Constant), class_educ_f

Υποσημείωση: Sum of squares, Regression: άθροισμα τετραγώνων των αποκλίσεων των εκτιμώμενων τιμών, Sum of squares, Residual: άθροισμα τετραγώνων των υπολοίπων, df: βαθμοί ελευθερίας, Mean square regression: είναι το μέσο τετράγωνο των ποσοτήτων αποκλίσεων για regression και Mean Square Residual: η διακύμανση των υπολοίπων. F: είναι το πηλίκο των δύο τελευταίων και Sig.: η στατιστική του σημαντικότητας, προκειμένου να απορριφθεί η $H_0: \rho=0$.

Πίνακας 11.5 Ανάλυση διακύμανσης της στατιστικής σημαντικότητας του συντελεστή προσδιορισμού R^2 .

Coefficients ^a					
Model		Unstandardized Coefficients		Standardized Coefficients	Sig.
		B	Std. Error	Beta	
1	(Constant)	-,399	1,550		,799
	class_educ_f	,751	,153	,688	,000
a. Dependent Variable: class_grade					

Υποσημείωση: B: οι συντελεστές της εξίσωσης παλινδρόμησης, Sd. Error: το τυπικό σφάλμα (τυπική απόκλιση της δειγματοληπτικής κατανομής του συντελεστή B), Beta: ο τυποποιημένος συντελεστής B του συντελεστή της ανεξάρτητης μεταβλητής, t: Το πηλίκο B/ Sd. Error που ακολουθεί κατανομή student t και Sig.: η στατιστική σημαντικότητα του t για τον έλεγχο της μηδενικής $H_0: \beta=0$

Πίνακας 11.6 : Οι συντελεστές της εξίσωσης παλινδρόμησης και η σημαντικότητα των τιμών τους.

11.3 Προϋποθέσεις εγκυρότητας της παλινδρόμησης και του συντελεστή r Pearson και μη παραμετρικοί συντελεστές συσχέτισης

Στις προηγούμενες ενότητες του Κεφαλαίου συζητήθηκαν οι έννοιες της συσχέτισης και της απλής γραμμικής παλινδρόμησης με τις οποίες περιγράφεται η σχέση μεταξύ δύο μεταβλητών. Όμως, όπως συμβαίνει με τις περισσότερες μεθόδους ή τεχνικές, υπάρχουν προϋποθέσεις για την εγκυρότητά τους. Οι προϋποθέσεις της απλής παλινδρόμησης, αν και διαφέρουν ελάχιστα από αυτές της συσχέτισης με τον r , μελετώνται διεξοδικά στο Κεφάλαιο 12.

Οι προϋποθέσεις του συντελεστή *Pearson* είναι:

1. Οι μεταβλητές είναι ποσοτικές συνεχείς.
2. Η μορφή της σχέσης είναι γραμμική.
3. Οι μεταβλητές αναμένεται να κατανέμονται κατά προσέγγιση κανονικά.

Αν από την κατασκευή του διαγράμματος σκεδασμού μεταξύ δύο μεταβλητών και τα ιστογράμματα των κατανομών τους διαπιστωθεί η παραβίαση της γραμμικότητας ή της κανονικότητας, αλλά η σχέση φαίνεται να είναι τουλάχιστον *μονότονη*, υπάρχει επίλυση του προβλήματος. Μονότονη είναι μια σχέση όταν σε κάθε αύξηση της τιμής της μίας μεταβλητής αυξάνεται και η τιμή της άλλης (θετική συσχέτιση) ή σε κάθε αύξηση της τιμής της μίας μειώνεται η τιμή της άλλης (αρνητική συσχέτιση). Σε αυτή την περίπτωση οι αρχικές τιμές των περιπτώσεων σε κάθε μεταβλητή μετατρέπονται σε θέσεις. Για παράδειγμα, στον Πίνακα 11.7 δίνονται οι τιμές των ωρών μελέτης και αντίστοιχες θέσεις για πέντε μαθητές. Με αυτόν τον τρόπο μετατρέπονται οι κλίμακες διαστήματος ή πηλίκου σε τακτικές κλίμακες. Η διαδικασία αυτή μετατρέπει συνήθως μια μη γραμμική σχέση δύο μεταβλητών σε γραμμική και καθιστά εφικτή τη χρήση του συντελεστή *Pearson*. Ωστόσο, ο υπολογισμός του συντελεστή συσχέτισης *Spearman*, που παρουσιάζεται πιο κάτω, προτιμάται λόγω της ευκολίας στους υπολογισμούς. Μια άλλη περίπτωση που οδηγεί στον συντελεστή *Spearman* είναι η περίπτωση που οι τιμές της μίας ή και των δύο μεταβλητών είναι θέσεις ή απλώς μεταβλητές τακτικής κλίμακας (παραβίαση της 1^{ης} από τις προϋποθέσεις που αναφέρθηκαν).

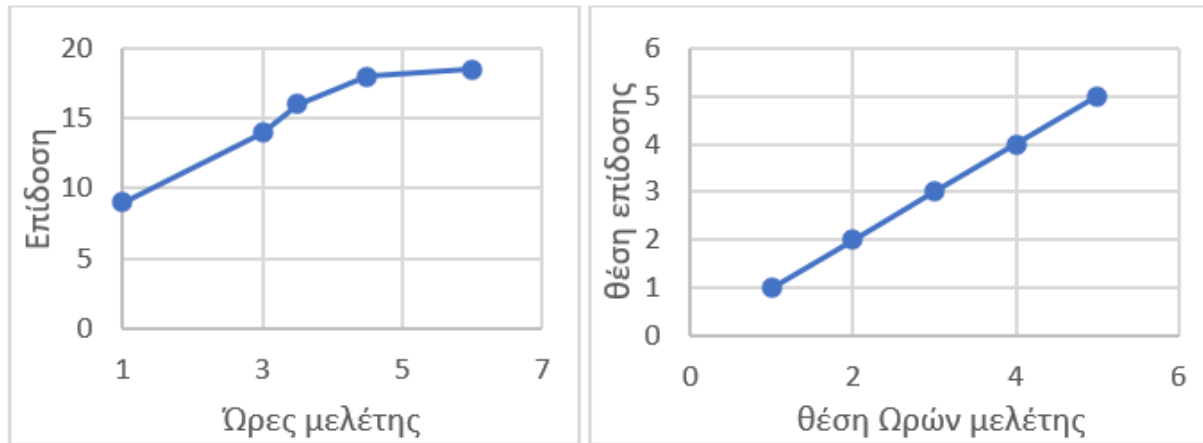
11.3.1 Ο Υπολογισμός του συντελεστή συσχέτισης *Spearman*

Προκειμένου να κατανοηθούν τα παραπάνω θα δώσουμε ένα παράδειγμα. Στον Πίνακα 11.7 παρουσιάζονται οι τιμές των μεταβλητών «Ώρες μελέτης ημερησίως» και «Επίδοση μαθητή» για ένα δείγμα $N = 5$ μαθητών. Παρουσιάζονται τόσο οι αρχικές τιμές όσο και οι θέσεις των ατόμων. Στην Εικόνα 11.13 παρουσιάζονται τα διαγράμματα σκεδασμού των αρχικών τιμών και των θέσεων των δύο μεταβλητών.

Από τον Πίνακα 11.7, εξετάζοντας τις αρχικές τιμές, παρατηρούμε μια ιδιαίτερα υψηλή συνάφεια ανάμεσα στις δύο μεταβλητές. Πράγματι, ο μαθητής M_1 έχει τις λιγότερες ώρες μελέτης και τον μικρότερο βαθμό επίδοσης. Ο μαθητής M_2 έχει τον δεύτερο λιγότερο χρόνο μελέτης και τον δεύτερο μικρότερο βαθμό επίδοσης κ.ο.κ. Δηλαδή, κάθε αύξηση των ωρών μελέτης συνοδεύεται από μία αύξηση της επίδοσης.

Μαθητής	Ώρες μελέτης		Επίδοση	
	Τιμή	Θέση	Τιμή	Θέση
M_1	1	1	9	1
M_2	3	2	14	2
M_3	3,5	3	16	3
M_4	4,5	4	18	4
M_5	6	5	18,5	5

Πίνακας 11.7 Μετατροπή τιμών σε θέσεις.



Εικόνα 11.13 Διάγραμμα αρχικών τιμών (α) και διάγραμμα θέσεων (β).

Παρατηρώντας το διάγραμμα σκεδασμού αυτών των αρχικών τιμών του Πίνακα 11.7 (Εικόνα 11.13 (α)), διαπιστώνουμε ότι τα σημεία ακολουθούν διαρκώς αυξητική πορεία, αλλά δεν βρίσκονται πάνω σε μια γραμμή. Όταν η σχέση απεικονιστεί στο διάγραμμα σκεδασμού των θέσεων του ίδιου Πίνακα (Εικόνα 11.13 (β)), διαπιστώνουμε ότι η σχέση είναι απόλυτα γραμμική. Συνεπώς, η αποκατάσταση της γραμμικότητας επιτρέπει τη χρήση του *Pearson r*. Αν η σχέση των δύο μεταβλητών είναι καμπυλόγραμμη, για παράδειγμα είχε τη μορφή U ή N, ο μετασχηματισμός σε θέσεις δεν έχει το ίδιο αποτέλεσμα και δεν χρησιμοποιούνται οι r και r_s , αλλά απαιτείται άλλη προσέγγιση.

Ο συντελεστής συσχέτισης θέσεων *Spearman*, που συμβολίζεται με r_s , δίνεται από τον Τύπο:

$$r_s = 1 - \frac{6\sum D^2}{n_p(n_p^2 - 1)}, \quad (11.9)$$

όπου

D^2 = η διαφορά ενός ζεύγους θέσεων στο τετράγωνο
και n_p = ο αριθμός των ζευγών ή ατόμων

Ως παράδειγμα εφαρμογής του τύπου 11.9 θα χρησιμοποιήσουμε τα δεδομένα του παραδείγματος του Πίνακα 11.8, αφού υπολογίσουμε τις θέσεις για κάθε μέτρηση του άγχους πριν και μετά την αεροβική άσκηση. Ο υπολογισμός της διαφοράς των θέσεων D και της διαφοράς στο τετράγωνο D^2 δίνονται επίσης στον Πίνακα 11.8.

Με τη βοήθεια του Τύπου στο (11.9) υπολογίζουμε τον συντελεστή *Spearman*:

$$\begin{aligned} r_s &= 1 - \frac{6\sum D^2}{n_p(n_p^2 - 1)} = \\ r_s &= 1 - \frac{6(67)}{9(9^2 - 1)} = \\ r_s &= 1 - \frac{402}{720} = 1 - 0,56 = 0,44 \end{aligned}$$

Είναι λογικό να υπάρχει θετικός συντελεστής συσχέτισης $r_s=0,44$ ανάμεσα σε δύο διαφορετικές μετρήσεις του άγχους των ατόμων που έλαβαν μέρος στο πείραμα.

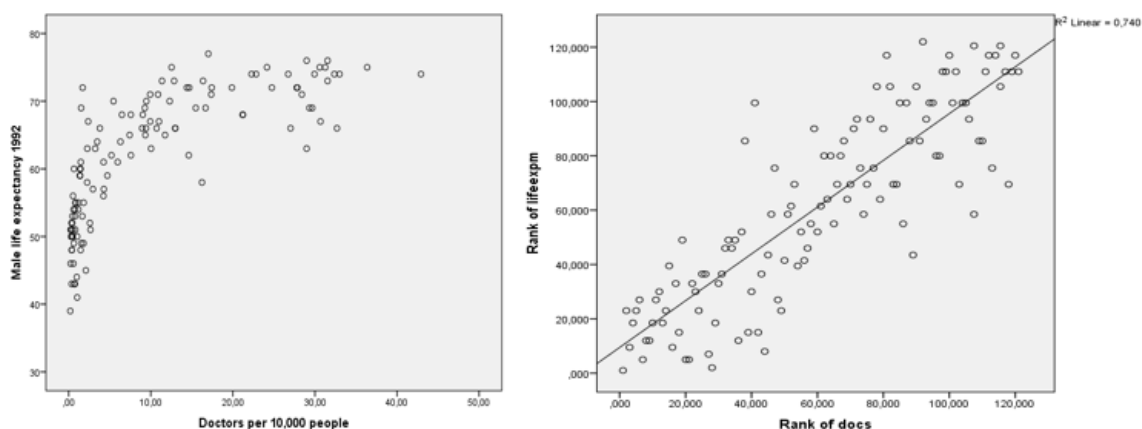
	Άγχος		Θέσεις		Διαφορά	
Άτομο	Πριν	Μετά	Πριν	Μετά	D	D^2
A ₁	17	10	8	6,5	1,5	2,25
A ₂	8	3	2,5	1	1,5	2,25
A ₃	18	16	9	9	0	0
A ₄	5	5	1	4,5	-3,5	12,25
A ₅	12	4	7	2,5	4,5	20,25
A ₆	8	5	2,5	4,5	-2	4
A ₇	10	4	5,5	2,5	3	9
A ₈	9	13	4	8	-4	16
A ₉	10	10	5,5	6,5	-1	1

Πίνακας 11.8. Υπολογισμός θέσεων με πρόσημο για το πρόγραμμα ελάττωσης του άγχους.

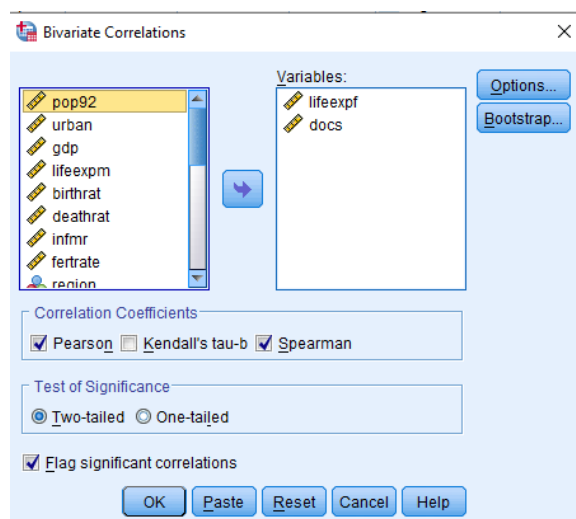
Όταν ελέγχουμε τη στατιστική σημαντικότητα του συντελεστή r_s , χρησιμοποιούμε τον Πίνακα Λ του παραρτήματος. Αυτός ο πίνακας περιέχει τις κρίσιμες τιμές για τον έλεγχο της σημαντικότητας του συντελεστή r_s . Οι κρίσιμες τιμές βρίσκονται με τη βοήθεια του αριθμού των ζευγών n_p , που βρίσκονται στην αριστερότερη στήλη του Πίνακα και του επιθυμητού επιπέδου σημαντικότητας που δίνεται στην πρώτη γραμμή του πίνακα. Η μηδενική υπόθεση, που ελέγχεται εδώ κατά αναλογία με τον συντελεστή *Pearson* r , είναι ότι η συσχέτιση των μεταβλητών στον πληθυσμό ρ_s είναι ίση με 0. Σημαντικός συντελεστής r_s σημαίνει ότι απορρίπτεται η μηδενική υπόθεση και γίνεται δεκτή η εναλλακτική, δηλαδή $\rho_s \neq 0$. Για να απορριφθεί η μηδενική υπόθεση, πρέπει η απόλυτη τιμή του r_s που υπολογίζεται από τα δεδομένα να είναι μεγαλύτερη από την κρίσιμη τιμή του πίνακα Λ.

Προκειμένου να ελεγχθεί η σημαντικότητα του $r_s = 0,44$ που υπολογίσαμε με τα δεδομένα του παραδείγματος του Πίνακα 11.8, για επίπεδο $\alpha = 0,05$ και $n_p = 9$, από τον Πίνακα Λ βρίσκουμε την κρίσιμη τιμή 0,7. Επειδή $|0,44| < 0,7$ δεν απορρίπτεται η μηδενική υπόθεση. Από το αρχείο «COUNTRY.SAV», στο οποίο καταγράφονται μια σειρά από δείκτες για ένα δείγμα 122 χωρών, θα μελετηθεί η συσχέτιση ανάμεσα στη μεταβλητή «DOCS» που εκφράζει τον αριθμό των γιατρών μια χώρας ανά 10.000 κατοίκους και τη μεταβλητή «LIFEEXPM» με το προσδόκιμο επιβίωσης των ανδρών σε έτη. Από το παρακάτω γράφημα της Εικόνας 11.14 είναι ξεκάθαρο ότι πρόκειται για μια μη γραμμική αλλά καμπυλόγραμμη σχέση ανάμεσα στις δύο μεταβλητές. Στην περίπτωση αυτή, ο συντελεστής *Pearson* r δεν είναι κατάλληλος αλλά μπορεί να χρησιμοποιηθεί ο *Spearman* r_s αφού η σχέση είναι μονότονη. Πράγματι, χρησιμοποιώντας τις θέσεις αντί των αρχικών τιμών φαίνεται ότι η σχέση είναι γραμμική. Για τον υπολογισμό του r_s από το βασικό μενού του SPSS επιλέγονται διαδοχικά **Analyze=>Correlate=>Bivariate**.

Στη συνέχεια, επιλέγουμε τις δύο μεταβλητές για τις οποίες επιθυμούμε τις συσχετίσεις, όπως φαίνεται στην Εικόνα 11.15, και μαρκάρουμε την επιλογή *Spearman*. Έπειτα πατάμε «κλικ» στο εικονίδιο **OK**.



Εικόνα 11.14 Διάγραμμα του αριθμού ιατρών απέναντι στο προσδόκιμο επιβίωσης των ανδρών με αρχικές τιμές (αριστερά) και με τις θέσεις των χωρών στις δύο μεταβλητές (δεξιά).



Εικόνα 11.15 Επιλογή του συντελεστή Spearman.

Επειδή δεν αφαιρέθηκε η επιλογή *Pearson*, στο παράθυρο *Bivariate Correlations* της Εικόνας 11.15 στο αποτέλεσμα του SPSS έχουμε δύο πίνακες, έναν για κάθε τύπο συντελεστή, όπως οι πίνακες αυτοί εμφανίζονται στην Εικόνα 11.16.

Correlations				
		lifeexpf Female life expectancy 1992	docs Doctors per 10,000 people	
lifeexpf Female life expectancy 1992	Pearson Correlation	1	,779**	
	Sig. (2-tailed)		,000	
	N	122	121	
docs Doctors per 10,000 people	Pearson Correlation	,779**	1	
	Sig. (2-tailed)	,000		
	N	121	121	
**. Correlation is significant at the 0.01 level (2-tailed).				
Correlations				
		lifeexpf Female life expectancy 1992	docs Doctors per 10,000 people	
Spearman's rho	lifeexpf Female life expectancy 1992	Correlation Coefficient	1,000	,881**
		Sig. (2-tailed)	.	,000
		N	122	121
	docs Doctors per 10,000 people	Correlation Coefficient	,881**	1,000
		Sig. (2-tailed)	,000	.
		N	121	121
**. Correlation is significant at the 0.01 level (2-tailed).				

Εικόνα 11.16. Οι πίνακες συσχέτισεων *Pearson* (άνω) και *Spearman* (κάτω).

Από το αποτέλεσμα προκύπτει ότι υπάρχει μια ισχυρή στατιστικά σημαντική θετική συσχέτιση ανάμεσα στον αριθμό των γιατρών, το προσδόκιμο και τα έτη επιβίωσης των ανδρών μιας χώρας. Παρατηρούμε ότι και ο *Pearson r* παρουσιάζει επίσης θετική σημαντική συσχέτιση, αλλά μικρότερου μεγέθους, τουτέστιν 0,779 έναντι του 0,881 που δίνεται από τον *Spearman*. Μια άλλη προσέγγιση, που θα μελετηθεί εκτενώς στο επόμενο Κεφάλαιο, θα ήταν η αποκατάσταση της κανονικότητας των τιμών της κατανομής της μεταβλητής «DOCS» υπολογίζοντας τον φυσικό λογάριθμό τους ($\ln(\text{DOCS})$). Με αυτόν τον μετασχηματισμό αποκαθίσταται

παράλληλα και η γραμμικότητα της σχέσης της $\text{Ln}(\text{DOCS})$ με την «LIFEEXPM» και, έτσι, ο συντελεστής *Pearson*, που είναι πλέον έγκυρος, δίνει μια τιμή $r=0,879$ σχεδόν ταυτόσημη με την 0,881 του *Spearman*.

11.4 Προβλήματα για εξάσκηση και ανατροφοδότηση

Άσκηση 1

Για ποια από τις παρακάτω περιπτώσεις θα προτεινάτε τη χρήση της γραμμικής παλινδρόμησης;

- Για να προβλεφθεί η επιλογή πολιτικού κόμματος από την ηλικία του ψηφοφόρου.
- Για να προβλεφθεί η επιλογή πολιτικού κόμματος από τον βαθμό αστικοποίησης της περιοχής διαμονής του ψηφοφόρου (αστική, ημιαστική, αγροτική).
- Για να προβλεφθεί το εισόδημα από το μορφωτικό επίπεδο (έτη εκπαίδευσης).
- Για να συγκριθεί η τηλεθέαση (ώρες ημερησίως) στις τρεις διαφορετικές περιοχές κατοικίας (αστική, ημιαστική, αγροτική).
- Για να προβλεφθεί το προσδόκιμο επιβίωσης των ανδρών από τον αριθμό των γιατρών ανά χίλιους κατοίκους μιας χώρας, χρησιμοποιώντας ένα τυχαίο δείγμα χωρών.

Άσκηση 2

Οι επιδόσεις στα μαθηματικά ενός τυχαίου δείγματος μαθητών Γυμνασίου μετρήθηκαν με ένα τεστ, δύο φορές με χρονική διαφορά τριών μηνών. Διαπιστώθηκε ότι η συσχέτιση ανάμεσα στις δύο μετρήσεις είναι γραμμική και υπολογίστηκε ο συντελεστής συσχέτισης r ανάμεσα στις δύο επιδόσεις (Επίδοση_1 και Επίδοση_2). Ποια από τις παρακάτω προτάσεις είναι λανθασμένη;

- Αν ο συντελεστής r έχει θετική τιμή, αυτό σημαίνει ότι οι μαθητές με τους μικρότερους βαθμούς στην πρώτη μέτρηση είναι εκείνοι με τους μικρότερους βαθμούς και στη δεύτερη μέτρηση.
- Αν ο συντελεστής r έχει θετική τιμή, αυτό σημαίνει ότι ο μέσος μαθητής του δείγματος βελτίωσε την επίδοσή του κατά τη δεύτερη μέτρηση.
- Υπάρχει σχεδόν βεβαιότητα ότι ο συντελεστής r έχει θετική τιμή.
- Αν ο συντελεστής r έχει αρνητική τιμή, αυτό σημαίνει ότι οι μαθητές με τους μεγαλύτερους βαθμούς στην πρώτη μέτρηση είναι εκείνοι με τους μικρότερους βαθμούς και στη δεύτερη μέτρηση.

Άσκηση 3

Χρησιμοποιώντας τα στοιχεία αναφορικά με τη σχέση ανάμεσα στα έτη εκπαίδευσης μεταξύ δύο συζύγων (μεταβλητές «husbeduc» και «wifeduc») που βρίσκονται στο αρχείο δεδομένων κοινωνικής έρευνας «gssnet.sav» (θα το βρείτε στη διεύθυνση: <http://www.cs.uni.edu/~jacobson/SCL/SPSS/data/>) να κατασκευαστεί το κατάλληλο γράφημα, προκειμένου να απαντηθούν τα παρακάτω:

- Ο τύπος των δεδομένων επιτρέπει τον υπολογισμό της συσχέτισης του *Pearson*.
- Η διαφανόμενη σχέση επιτρέπει τη χρησιμοποίηση της απλής γραμμικής παλινδρόμησης. Γιατί;
- Θα χαρακτηρίζατε τη σχέση θετική ή αρνητική; Γιατί;
- Ποια είναι η τιμή του συντελεστή συσχέτισης; Να χαρακτηριστεί ως προς το μέγεθός της και τη σημαντικότητά της.

Άσκηση 4

Από τα δεδομένα του θέματος 2 να εκτελεστεί στο SPSS η ανάλυση απλής παλινδρόμησης προκειμένου να μελετηθεί η πρόβλεψη της εκπαίδευσης του συζύγου από αυτή της συζύγου του.

- Να γραφεί η εξίσωση παλινδρόμησης.
- Ποιο ποσοστό της διακύμανσης της εκπαίδευσης του συζύγου ερμηνεύεται από την εκπαίδευση της συζύγου;
- Τι ακριβώς εκφράζει η τιμή της σταθεράς και τι η κλίση της ευθείας που αντιστοιχεί στην εξίσωση της παλινδρόμησης;

- δ. Ποια είναι η προβλεπόμενη τιμή για τα έτη εκπαίδευσης του συζύγου, αν η εκπαίδευση της συζύγου του είναι 12 χρόνια;
- ε. Αν η εκπαίδευση του συζύγου είναι 15 έτη και της συζύγου 10, ποια είναι η τιμή του υπολοίπου;

Άσκηση 5

Να χρησιμοποιηθεί το αρχείο «Math_attitudes» που περιέχει τα δεδομένα από πανελλαδική έρευνα που πραγματοποιήθηκε σε δείγμα μαθητών της Α΄ Γυμνασίου. Με την κατάλληλη επιλογή από το SPSS να δημιουργηθεί ο πίνακας συσχετίσεων μεταξύ των μεταβλητών «C» (αυτοπεποίθηση απέναντι στα Μαθηματικά), «U» (χρησιμότητα των Μαθηματικών), «AS» (στάση για την επιτυχία στα Μαθηματικά), «E» (κίνητρο αποτελεσματικότητας), «f» (υποστηρικτικός πατέρας), «m» (υποστηρικτική μητέρα), «t» (υποστηρικτικός δάσκαλος) και «grade» (επίδοση σε τεστ Μαθηματικών) που δόθηκε σε μαθητές της Α΄ Γυμνασίου. Οι σχέσεις μεταξύ των παραπάνω ποσοτικών κλιμάκων θα θεωρηθούν γραμμικές. Να απαντηθούν τα παρακάτω:

- α. Ποιος είναι ο πληθυσμός της έρευνας;
- β. Ποια είναι η μεταβλητή που ερμηνεύει καλύτερα τη μεταβλητότητα της «grade»; Η σχέση της μεταβλητής αυτής με τη «grade» είναι στατιστικά σημαντική; Η διεύθυνση της σχέσης είναι η αναμενόμενη σύμφωνα με την κοινή λογική; Πώς χαρακτηρίζεται η έντασή της;
- γ. Να συζητηθεί η πρόταση: «Η επίδοση των μαθητών στα Μαθηματικά εξαρτάται περισσότερο από την υποστήριξη των γονέων και του δάσκαλου παρά από την αυτοπεποίθηση και το κίνητρο αποτελεσματικότητάς των μαθητών απέναντι στα Μαθηματικά».
- δ. Τα ερωτήματα β) και γ) να διερευνηθούν ξεχωριστά για μαθητές και μαθήτριες και να σχολιαστούν τα αποτελέσματα συγκριτικά ως προς τα δύο φύλα των μαθητών.

Βιβλιογραφία

- Γναρδέλης, Χ. (2003). *Εφαρμοσμένη Στατιστική*. Αθήνα: Εκδ. Παπαζήση.
- Cohen, J. (1977). *Statistical power analysis for the behavioral sciences* (Revised). New York: Academic.
- Howell, C. D. (1997). *Statistical Methods in Psychology* (4th ed.). Duxbury Press, by Wadsworth Publishing Co. International Thomson Publishing Inc.
- Norusis, M. J. (2002). *SPSS 11.0 Guide to Data Analysis*. New York: Prentice-Hall.
- Ott, L., & Longnecker, M. (2001). *An Introduction to Statistical Methods and data Analysis* (5th ed.). Duxbury, Pacific Grove.
- Tabachnick, B. G., & Fidell, L. S. (1996). *Using multivariate statistics*. (3rd ed.). New York: Harper Collins.

Παράρτημα

$df = N - 2$	Επίπεδο σημαντικότητας για μονόπλευρο έλεγχο				
	0,05	0,25	0,01	0,005	0,0005
	Επίπεδο σημαντικότητας για αμφίπλευρο έλεγχο				
	0,1	0,05	0,02	0,01	0,001
1	0,988	0,997	0,9995	0,9999	0,999999
2	0,900	0,950	0,980	0,990	0,999
3	0,805	0,878	0,934	0,959	0,991
4	0,729	0,811	0,882	0,917	0,974
5	0,669	0,754	0,833	0,875	0,951
6	0,621	0,707	0,789	0,834	0,925
7	0,582	0,666	0,750	0,798	0,898
8	0,549	0,632	0,715	0,765	0,872
9	0,521	0,602	0,685	0,735	0,847
10	0,497	0,576	0,658	0,708	0,823
11	0,476	0,553	0,634	0,684	0,801
12	0,458	0,532	0,612	0,661	0,780
13	0,441	0,514	0,592	0,641	0,760
14	0,426	0,497	0,574	0,623	0,742
15	0,412	0,482	0,558	0,606	0,725
16	0,400	0,468	0,543	0,590	0,708
17	0,389	0,456	0,529	0,575	0,693
18	0,378	0,444	0,516	0,561	0,679
19	0,369	0,433	0,503	0,549	0,665
20	0,360	0,423	0,492	0,537	0,652
21	0,352	0,413	0,482	0,526	0,640
22	0,344	0,404	0,472	0,515	0,629
23	0,337	0,396	0,462	0,505	0,618
24	0,330	0,388	0,453	0,496	0,607
25	0,323	0,381	0,445	0,487	0,597
26	0,317	0,374	0,437	0,479	0,588
27	0,311	0,367	0,430	0,471	0,579
28	0,306	0,361	0,423	0,463	0,570
29	0,301	0,355	0,416	0,456	0,562
30	0,296	0,349	0,409	0,449	0,554
35	0,275	0,325	0,381	0,418	0,519
40	0,257	0,304	0,358	0,393	0,490
45	0,243	0,288	0,338	0,372	0,465
50	0,231	0,273	0,322	0,354	0,443
60	0,211	0,250	0,295	0,325	0,408
70	0,195	0,232	0,274	0,302	0,380
80	0,183	0,217	0,257	0,283	0,357
90	0,173	0,205	0,242	0,267	0,338
100	0,164	0,195	0,230	0,254	0,321

Πίνακας Ζ: Κρίσιμες τιμές του συντελεστή συσχέτισης Pearson.

<i>N</i>	Επίπεδο σημαντικότητας για μονόπλευρο έλεγχο			
	0,05	0,25	0,01	0,005
	Επίπεδο σημαντικότητας για αμφίπλευρο έλεγχο			
	0,1	0,05	0,02	0,01
5	0,900	*	*	*
6	0,829	0,886	0,943	*
7	0,714	0,786	0,893	*
8	0,643	0,738	0,833	0,881
9	0,600	0,683	0,783	0,833
10	0,564	0,648	0,745	0,794
11	0,523	0,623	0,736	0,818
12	0,497	0,591	0,703	0,780
13	0,475	0,566	0,673	0,745
14	0,457	0,545	0,646	0,716
15	0,441	0,525	0,623	0,689
16	0,425	0,507	0,601	0,666
17	0,412	0,490	0,582	0,645
18	0,399	0,476	0,564	0,625
19	0,388	0,462	0,549	0,608
20	0,377	0,450	0,534	0,591
21	0,368	0,438	0,521	0,576
22	0,359	0,428	0,508	0,562
23	0,351	0,418	0,496	0,549
24	0,343	0,409	0,485	0,537
25	0,336	0,400	0,475	0,526
26	0,329	0,392	0,465	0,515
27	0,323	0,385	0,456	0,505
28	0,317	0,377	0,448	0,496
29	0,311	0,370	0,440	0,487
30	0,305	0,364	0,432	0,478

Πίνακας Α: Κρίσιμες τιμές για τον συντελεστή συσχέτισης *Spearman*.

Κεφάλαιο 12

Πολλαπλή Γραμμική Παλινδρόμηση

Σύνοψη

Στο κεφάλαιο αυτό εξετάζεται το υπόδειγμα της πολλαπλής γραμμικής παλινδρόμησης για μία εξαρτημένη ποσοτική μεταβλητή και μία ομάδα ανεξάρτητων μεταβλητών που μπορεί να ανήκουν σε οποιονδήποτε τύπο μεταβλητής. Παρουσιάζονται και ελέγχονται οι προϋποθέσεις καταλληλότητας των κατανομών αυτών με τη βοήθεια των τυποποιημένων υπολοίπων, καθώς και η γραμμικότητα των σχέσεων μεταξύ της εξαρτημένης μεταβλητής και των ανεξάρτητων μεταβλητών. Επίσης, παρουσιάζεται και ερμηνεύεται πίνακας με δείκτες προσαρμογής του υποδείγματος και, πιο συγκεκριμένα, ο πίνακας των μερικών συντελεστών της εξίσωσης παλινδρόμησης με τη στατιστική τους σημαντικότητα. Ακόμα, παρουσιάζονται στατιστικές μέθοδοι ελάττωσης του πλήθους των ανεξάρτητων μεταβλητών. Τέλος, δίνονται και εξετάζονται δείκτες που προσδιορίζουν περιπτώσεις με ισχυρή επίδραση στους συντελεστές και στην προσαρμογή του υποδείγματος.

Προαπαιτούμενη γνώση

Ιδιότητες κανονικής κατανομής, μέτρα κεντρικής θέσης και διακύμανσης, έλεγχος υποθέσεων με τις κατανομές F και t , συντελεστής συσχέτισης Pearson r , απλή γραμμική παλινδρόμηση: Κεφάλαια 4, 6, 7, 8 και 9 του συγγράμματος.

12.1 Εισαγωγή στην Πολλαπλή παλινδρόμηση

12.1.1 Περιγραφή και έλεγχος υποδείγματος

Η Πολλαπλή Γραμμική Παλινδρόμηση (ΠΓΠ), επίσης γνωστή απλώς ως πολλαπλή παλινδρόμηση, αποτελεί μια στατιστική τεχνική που χρησιμοποιεί πολλές επεξηγηματικές (ερμηνευτικές ή ανεξάρτητες) μεταβλητές, για να προβλέψει το αποτέλεσμα μιας εξαρτημένης ποσοτικής μεταβλητής.

Στην ενότητα αυτή θα χρησιμοποιηθούν δεδομένα από την έρευνα των Λεονταρή και Γιαλαμά (2005), προκειμένου να ερμηνευτεί η μεταβλητότητα της κατάθλιψης (μεταβλητή «DEPR») σε εφήβους από τις μεταβλητές που δίνονται στον παρακάτω Πίνακα 12.1.

Όνομα μεταβλητής	Περιγραφή	Τύπος μεταβλητής
BEH	Έλεγχος Συμπεριφοράς των γονέων	Ποσοτική
PSY_F	Ψυχολογικός έλεγχος Πατέρα	Ποσοτική
PSY_M	Ψυχολογικός έλεγχος Μητέρας	Ποσοτική
MARKS	Γενικός βαθμός (0-20)	Ποσοτική
AGE	Ηλικία σε έτη	Ποσοτική
FEMALE	Φύλο μαθητή (0=Αγόρι/ 1=κορίτσι)	Κατηγορική (Ψευδο-μεταβλητή των κοριτσιών)

Πίνακας 12.1 Ερμηνευτικές μεταβλητές της κατάθλιψης.

Το υπόδειγμα ή η εξίσωση που θα χρησιμοποιηθεί στην πρόβλεψη της κατάθλιψης από τις ερμηνευτικές μεταβλητές του Πίνακα 12.1 είναι:

$$\begin{aligned} & \text{προβλεπόμενη τιμή της κατάθλιψης (DEPR)} = \\ & = B_0 + B_1 \text{BEH} + B_2 \text{PSY_F} + B_3 \text{PSY_M} + B_4 \text{MARKS} + B_5 \text{AGE} + B_6 \text{FEMALE} \end{aligned}$$

Αντί να έχουμε μία σταθερά και μία κλίση έχουμε μία σταθερά (B_0) και 6 συντελεστές B_1 έως B_6 , για τις έξι προβλεπτικές μεταβλητές, οι οποίοι ονομάζονται μερικοί συντελεστές παλινδρόμησης. Επειδή στις περισσότερες περιπτώσεις έχουμε ένα δείγμα με τη βοήθεια του οποίου εξάγονται συμπεράσματα για τον πληθυσμό από τον οποίο προέρχονται οι μερικοί συντελεστές, οι παραπάνω εξισώσεις αποτελούν εκτιμήσεις των άγνωστων παραμέτρων $\beta_0, \beta_1, \dots, \beta_6$ της εξίσωσης που ισχύει για τον πληθυσμό.

12.1.2 Κωδικοποίηση κατηγορικών ανεξάρτητων μεταβλητών

Οι κατηγορικές μεταβλητές, πριν εισέλθουν ως ανεξάρτητες στο δεξί μέρος της εξίσωσης στην πολλαπλή παλινδρόμηση, θα πρέπει να υποστούν τον παρακάτω μετασχηματισμό:

- Αν η κατηγορική μεταβλητή X έχει k κατηγορίες, δημιουργούνται ισάριθμες k διχοτομικές μεταβλητές (με τιμές 0 και 1) $X_1, X_2, \dots, X_k$, οι οποίες ονομάζονται και *ψευδομεταβλητές*. Κάθε τέτοια ψευδομεταβλητή αντιστοιχεί σε μια κατηγορία της X . Αν ένα άτομο ανήκει στην κατηγορία j της X , τότε θα έχει την τιμή 1 στην ψευδομεταβλητή X_j και την τιμή 0 σε όλες τις άλλες ψευδομεταβλητές.

•	•	•	•	•	•	•	•
•	•	•	•	•	•	•	•
•	•	•	•	•	•	•	•
•	•	•	•	•	•	•	•
•	•	•	•	•	•	•	•
•	•	•	•	•	•	•	•
•	•	•	•	•	•	•	•
•	•	•	•	•	•	•	•

Πίνακας 12.2 Κατηγορικής μεταβλητής X με k κατηγορίες και $X_1 \dots X_k$ ψευδομεταβλητές.

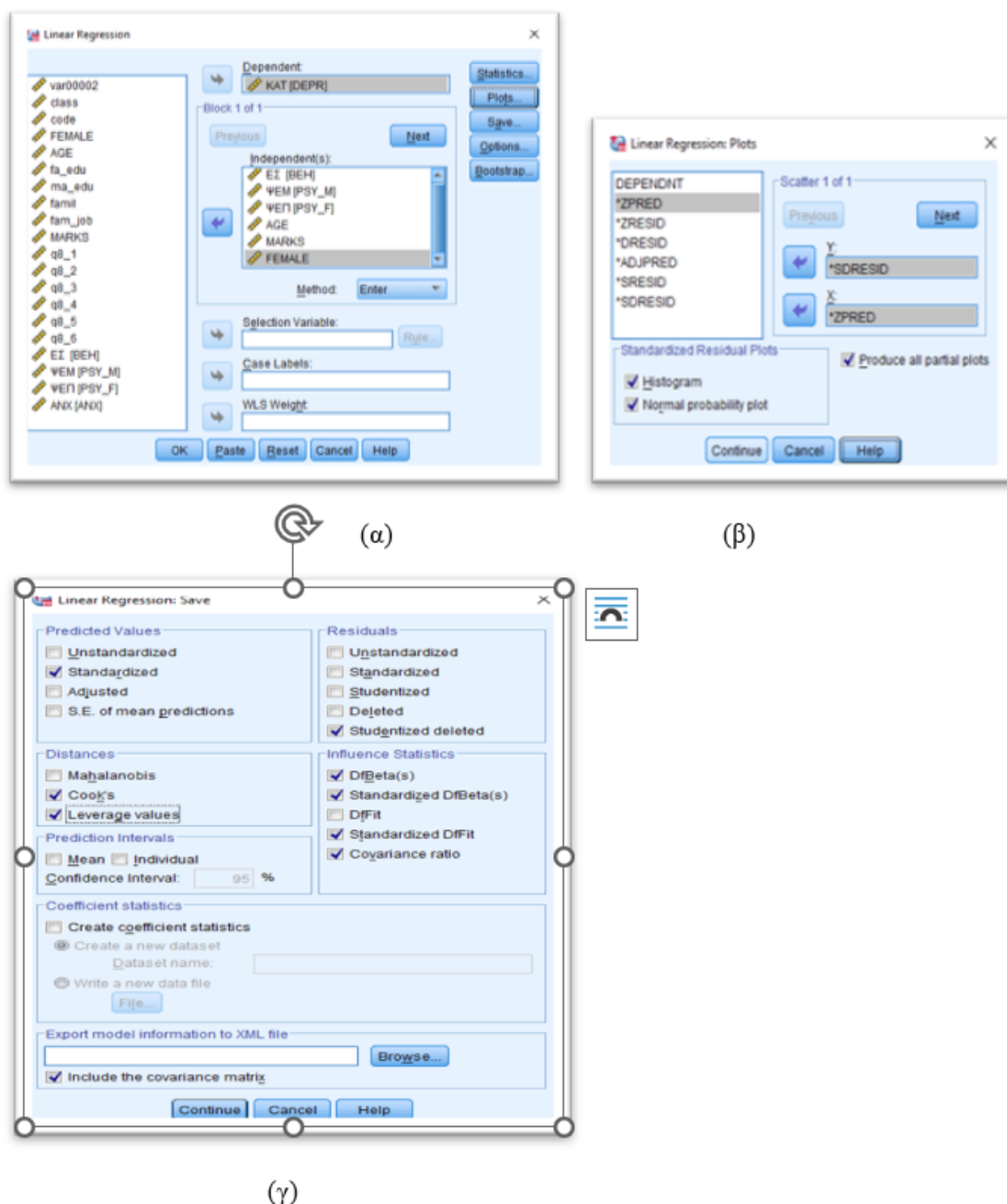
- Στην ανάλυση συμμετέχουν αντί της X οι $k-1$ ψευδομεταβλητές και παραμένει εκτός η λεγόμενη ψευδομεταβλητή της κατηγορίας αναφοράς.
- Το ποια θα επιλεγεί ως κατηγορία αναφοράς εξαρτάται από τους στόχους του ερευνητή. Αν στο παράδειγμά μας η βαθμολογία ήταν ανεξάρτητη κατηγορική με τιμές «Χαμηλός», «Μέτριος» και «Υψηλός» βαθμός, θα ήταν λογικό να εισαχθούν στο υπόδειγμα οι ψευδομεταβλητές που αντιστοιχούν στις δύο κατηγορίες «Μέτριος» και «Υψηλός» βαθμός, προκειμένου να συγκριθούν με τον «Χαμηλό» βαθμό ως προς το επίπεδο της κατάθλιψης των εφήβων. Στο παράδειγμα που μελετάται σε αυτή την Ενότητα, ό,τι αφορά το Φύλο αντιπροσωπεύεται με την ψευδομεταβλητή των κοριτσιών.

12.1.3 Προϋποθέσεις εγκυρότητας του υποδείγματος της πολλαπλής παλινδρόμησης

Οι βασικές προϋποθέσεις για την εγκυρότητα της πολλαπλής παλινδρόμησης είναι στην πραγματικότητα μία γενίκευση των προϋποθέσεων που ισχύουν στην απλή παλινδρόμηση:

1. Η κατανομή των τιμών της εξαρτημένης μεταβλητής για κάθε συνδυασμό τιμών των ανεξάρτητων μεταβλητών του υποδείγματος είναι κανονική.
2. Η διακύμανση των τιμών της εξαρτημένης μεταβλητής παραμένει σταθερή για κάθε συνδυασμό τιμών των ανεξάρτητων μεταβλητών του υποδείγματος.
3. Η σχέση καθεμίας ανεξάρτητης μεταβλητής με την εξαρτημένη είναι γραμμική και παραμένει γραμμική, ακόμα και όταν αφαιρεθεί η επίδραση των υπολοίπων ανεξάρτητων μεταβλητών από αυτή.
4. Οι παρατηρήσεις είναι ανεξάρτητες μεταξύ τους.
5. Δεν υπάρχει πολυσυγγραμμικότητα, δηλαδή ισχυρές συσχετίσεις μεταξύ των ανεξάρτητων μεταβλητών, οι οποίες αλλοιώνουν την κατεύθυνση, το μέγεθος και τη σημαντικότητα ορισμένων συντελεστών.

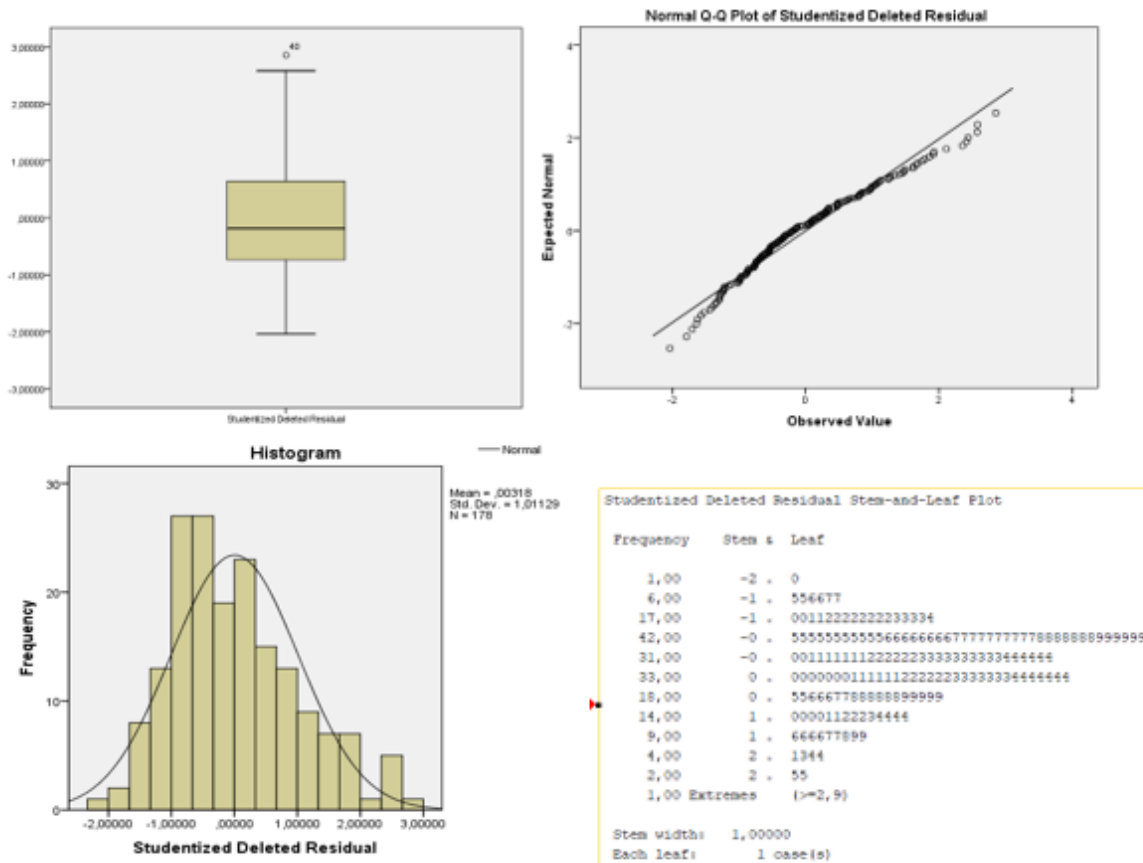
Στη διερεύνηση των περισσότερων από τις παραπάνω προϋποθέσεις χρησιμοποιούνται τα υπόλοιπα $e = Y - \hat{Y}$, που έχουν υποστεί όμως μια τυποποίηση. Στη βιβλιογραφία προτείνεται η χρήση των «τυποποιημένων κατά *student* διαγραμμένων» (SDR). Ένα τέτοιο υπόλοιπο για ένα άτομο είναι ένα τυποποιημένο κατά *student* υπόλοιπο που υπολογίζεται όταν το άτομο εξαιρεθεί από την παλινδρόμηση. Η εκτέλεση της ανάλυσης παλινδρόμησης της «DEPR» για τη δημιουργία των υπολοίπων, όπως άλλες επιλογές που θα χρησιμοποιηθούν στη συνέχεια, απεικονίζονται στα πλαίσια διαλόγου [Linear regression](#), [Linear regression: Plots](#) και [Linear regression: Save](#), που δίνονται στην Εικόνα 12.1.



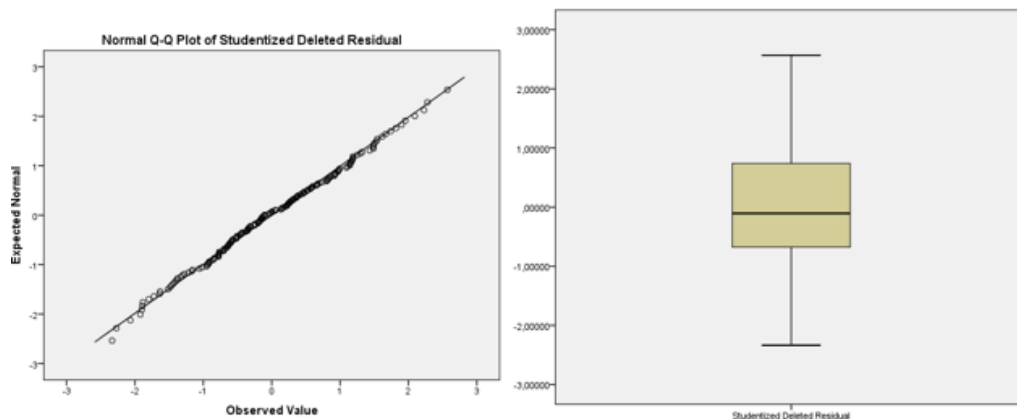
Εικόνα 12.1 (α), (β), (γ) Επιλογές αποθήκευσης υπολοίπων και προβλεπόμενων και κατασκευής γραφημάτων για τον έλεγχο των προϋποθέσεων της παλινδρόμησης.

1. Τα υπόλοιπα αυτά ακολουθούν κατανομή *student* όταν ισχύουν οι προϋποθέσεις της παλινδρόμησης με $N-p-2$ βαθμούς ελευθερίας, όπου N το μέγεθος του δείγματος και p ο αριθμός των ανεξάρτητων μεταβλητών. Για επαρκώς μεγάλα δείγματα ($N > 30$) η κατανομή των SDR είναι κατά προσέγγιση κανονική. Το απλούστερο γράφημα για τη διαπίστωση της

κανονικότητας είναι το ιστόγραμμα των SDR ή το φυλλογράφημα ή το θηκόγραμμα και, τέλος, το «Q-Q plot» το οποίο είναι ένα πιο εξειδικευμένο διαγνωστικό εργαλείο κανονικότητας. Το τελευταίο είναι ένα διάγραμμα σκεδασμού των SDR απέναντι στις αναμενόμενες τιμές της τυπικής κανονικής κατανομής. Αν όλα τα σημεία βρίσκονται πάνω σε μία ευθεία γραμμή τότε η κατανομή των SDR είναι κανονική. Από τα τέσσερα γραφήματα αναδείχτηκε η μονοκόρυφη θετικής ασυμμετρίας κατανομή των SDR με μικρό σχετικά βαθμό απόκλισης από την κανονική κατανομή (Εικόνα 12.2). Ο μετασχηματισμός της εξαρτημένης μεταβλητής, χρησιμοποιώντας την τετραγωνική ρίζα των τιμών της και η εκτέλεση εκ νέου της πολλαπλής παλινδρόμησης, δημιουργεί υπόλοιπα SDR που κατανέμονται κατά προσέγγιση κανονικά και επιπλέον δεν υπάρχουν άτομα με ακραίες τιμές (Εικόνα 12.3).

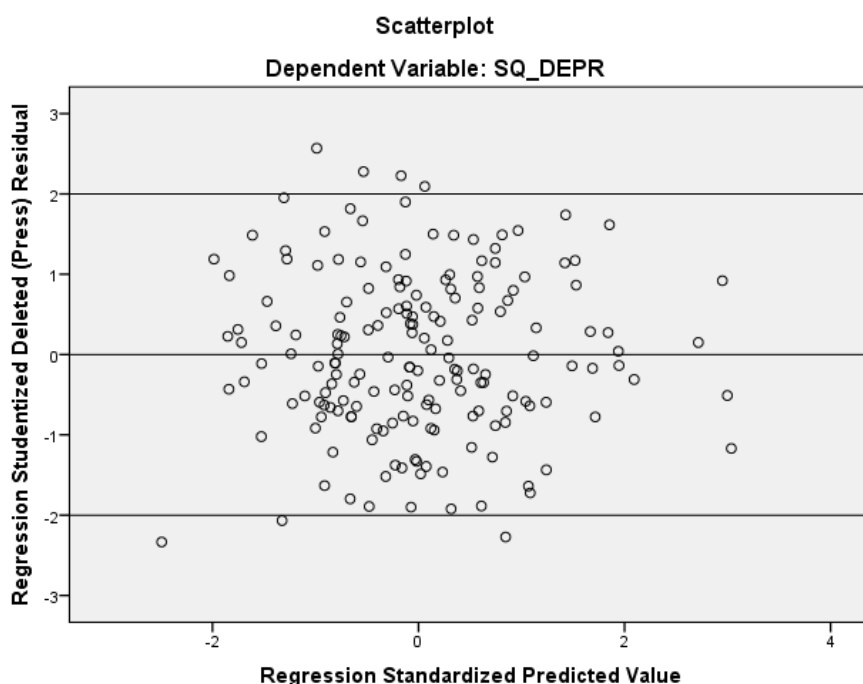


Εικόνα 12.2 Γραφήματα για τον έλεγχο της κανονικότητας των υπολοίπων της παλινδρόμησης της «DEPR».



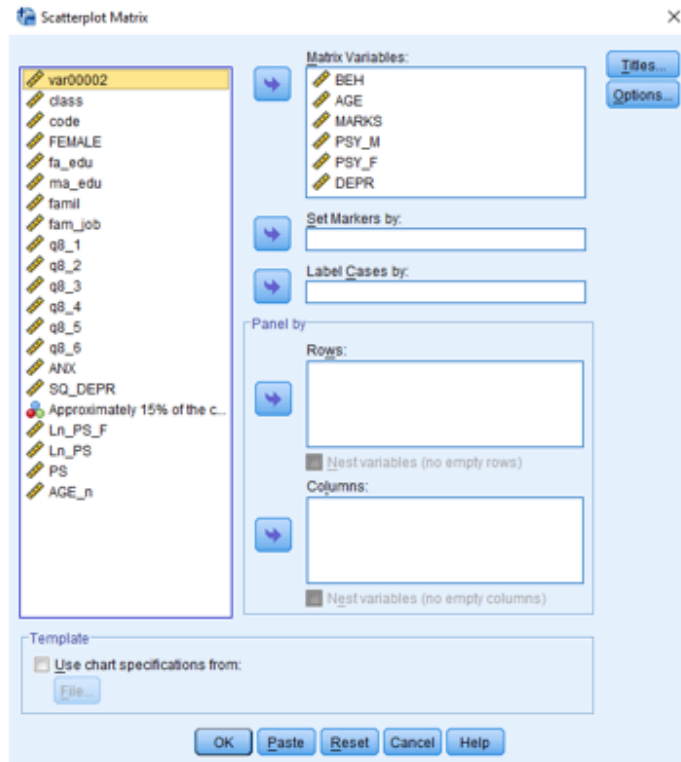
Εικόνα 12.3 Γραφήματα ελέγχου κανονικότητας μετά τον μετασχηματισμό της τετραγωνικής ρίζας της εξαρτημένης μεταβλητής του υποδείγματος της «DEPR».

2. Το διάγραμμα σκεδασμού των SDR πάνω στις προβλεπόμενες τιμές της «DEPR» αναδεικνύει την ομοσκεδαστικότητα. Αν τα SDR συγκεντρώνονται λίγο ως πολύ γύρω από μία οριζόντια γραμμή που διέρχεται από το μηδέν και σχηματίζουν μία ζώνη ίδιου περίπου πλάτους για κάθε προβλεπόμενη τιμή της «DEPR», τότε φαίνεται να ικανοποιείται η σταθερότητα διακύμανσης των SDR σε κάθε προβλεπόμενη τιμή της «DEPR». Από την Εικόνα 12.4 μπορεί να διαπιστωθεί η κατά προσέγγιση σταθερότητα της διακύμανσης αλλά και η σχετικά μικρότερη διακύμανση για τα άτομα με μεγάλες τιμές της «DEPR». Ο μετασχηματισμός της εξαρτημένης μεταβλητής με τη χρήση κάποιας από τις συναρτήσεις $\sqrt{X}$, $\log_{10} X$, ή $\frac{1}{X}$ συνήθως αποκαθιστά τη σταθερότητα της διακύμανσης.

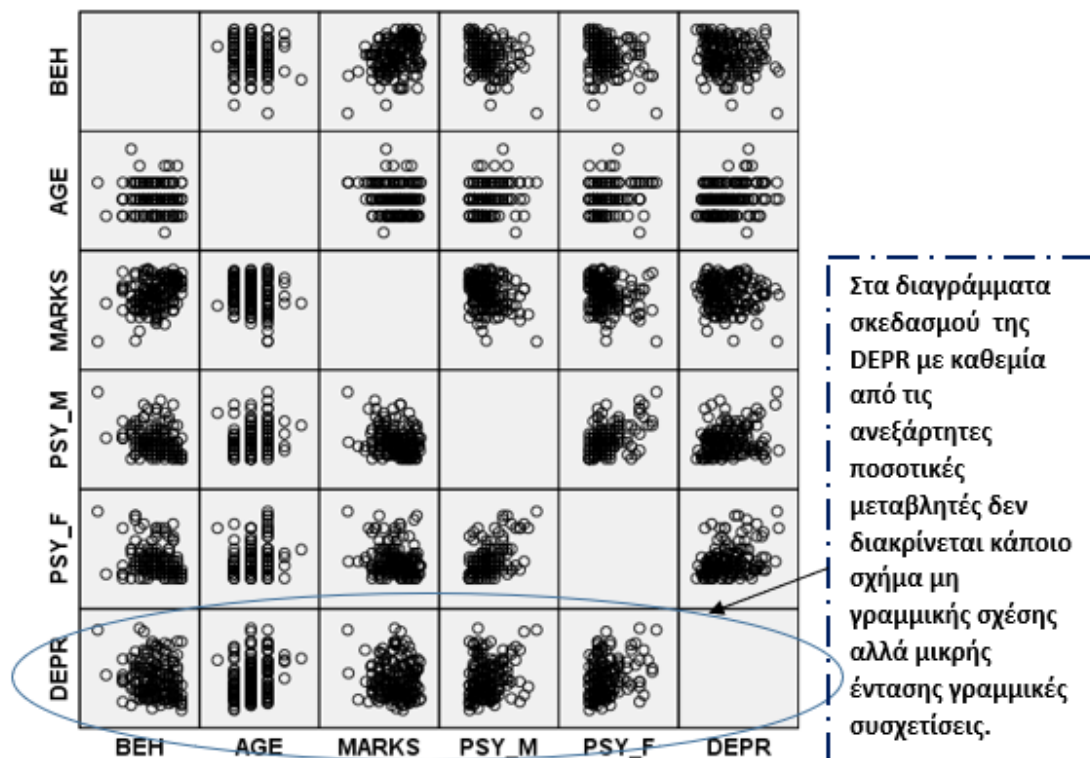


Εικόνα 12.4 Διάγραμμα σκεδασμού των υπολοίπων SDR πάνω στις προβλεπόμενες τυποποιημένες τιμές από το υπόδειγμα.

3. Η γραμμικότητα της σχέσης εξαρτημένης με καθεμία ανεξάρτητη μεταβλητή θα πρέπει να διαπιστωθεί πριν από την εκτέλεση της πολλαπλής παλινδρόμησης χρησιμοποιώντας διαγράμματα σκεδασμού. Για την ευκολότερη δημιουργία όλων των διαγραμμάτων σκεδασμού μεταξύ των ποσοτικών μεταβλητών οι οποίες συμμετέχουν στο υπόδειγμα, θα επιλέξουμε **Graphs=>=>Legacy Dialogs=>Scatter/Dot...=>Matrix Scatter** και έπειτα πατάμε **Define** και εισάγουμε τις ποσοτικές μεταβλητές «DEPR», «BEH», «PSY_F, PSY_M, AGE, MARKS» στο πλαίσιο **Matrix Variables** και πατάμε **OK** (Εικόνα 12.5 α). Στην Εικόνα 12.5 (α) και (β) δίνεται το παράθυρο **Scatterplot Matrix** και ο Πίνακας των διαγραμμάτων σκεδασμού. Η μελέτη των διαγραμμάτων σκεδασμού της εξαρτημένης με κάθε ανεξάρτητη μεταβλητή θα αναδείξει τις ανεξάρτητες μεταβλητές για τις οποίες η σχέση με την εξαρτημένη αποκλίνει της γραμμικής. Ο μετασχηματισμός των τιμών της ανεξάρτητης με τις συναρτήσεις $\sqrt{X}$, $\log_{10} X$, ή $\frac{1}{X}$, σε αντιστοιχία με την ένταση του προβλήματος (από τη μικρότερη προς τη μεγαλύτερη) θα μπορούσε να αποκαταστήσει τη γραμμικότητα.



(α)



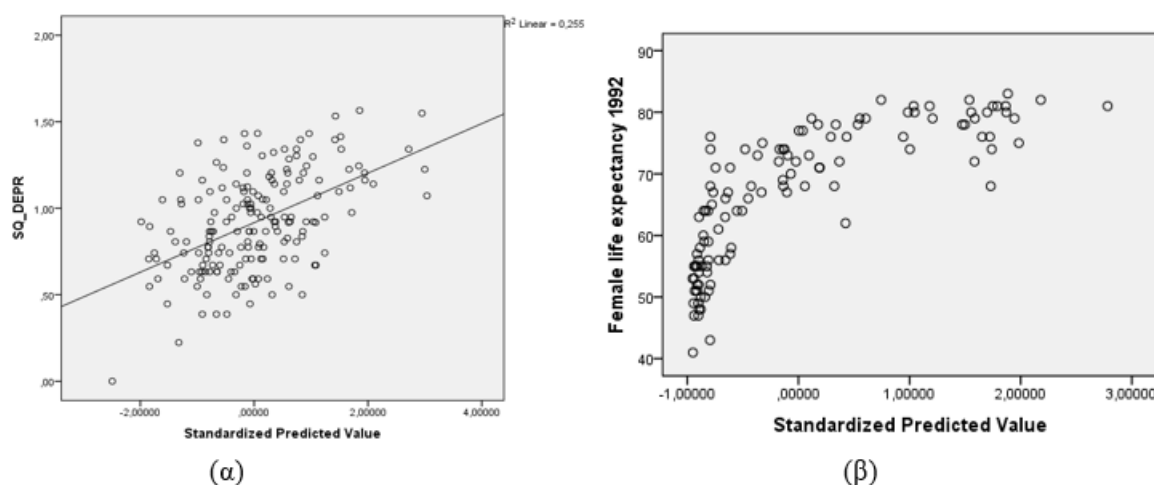
(β)

Εικόνα 12.5 Επιλογές για την κατασκευή πολλαπλών διαγραμμάτων σκεδασμού (α) και πίνακας πολλαπλών διαγραμμάτων σκεδασμού (β).

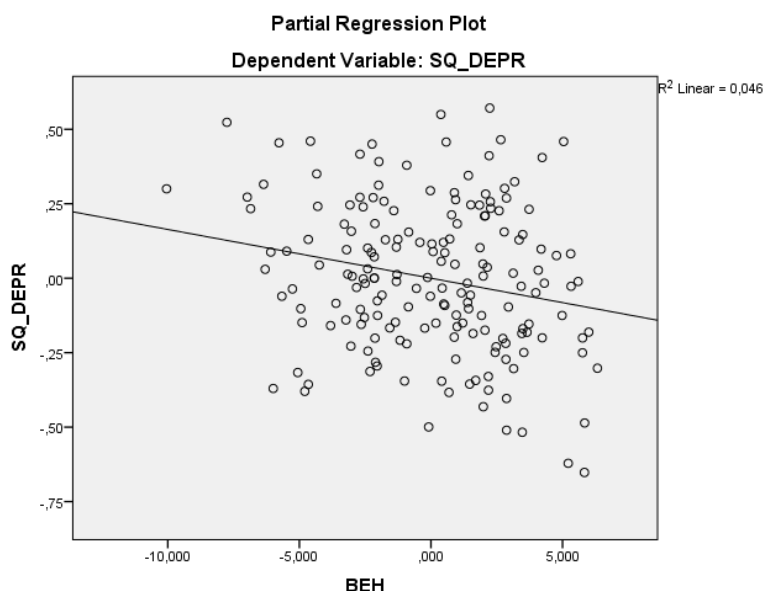
Ένα άλλο γράφημα στο οποίο αναμένεται γραμμικότητα είναι το διάγραμμα σκεδασμού της εξαρτημένης μεταβλητής πάνω στην προβλεπόμενη τιμή της εξαρτημένης. Στην Εικόνα 12.6 (α), που δημιουργήθηκε από

τις προβλεπόμενες και παρατηρούμενες τιμές της μεταβλητής «DEPR», δεν δημιουργείται αμφιβολία για τη γραμμικότητα της σχέσης (αντίθετα, στο διάγραμμα 12.6 (β) στο οποίο απεικονίζονται τα υπόλοιπα μιας άλλης παλινδρόμησης για την πρόβλεψη της αναμενόμενης επιβίωσης των γυναικών, υπάρχει σαφής ένδειξη ότι τα δεδομένα δεν εκφράζονται ικανοποιητικά από μια γραμμική σχέση αλλά από μια καμπυλόγραμμη).

Τέλος, η γραμμικότητα πρέπει να επιβεβαιώνεται και στα διαγράμματα μερικής παλινδρόμησης. Στην Εικόνα 12.7 παρουσιάζεται το διάγραμμα μερικής παλινδρόμησης της μεταβλητής «BEH». Στον οριζόντιο άξονα βρίσκονται τα υπόλοιπα της παλινδρόμησης της «SQ_DEPR» πάνω σε όλες τις ανεξάρτητες μεταβλητές εκτός της «BEH». Στον κάθετο άξονα βρίσκονται τα υπόλοιπα της παλινδρόμησης της «BEH» πάνω σε όλες τις άλλες ανεξάρτητες μεταβλητές. Με τον υπολογισμό των υπολοίπων στην πραγματικότητα αφαιρούνται οι επιδράσεις όλων των υπόλοιπων ανεξάρτητων μεταβλητών τόσο από την εξαρτημένη όσο και από την ανεξάρτητη μεταβλητή που βρίσκεται στο γράφημα. Αν υφίσταται γραμμική μερική συσχέτιση, τότε το γράφημα μερικής παλινδρόμησης είναι γραμμικό. Στο γράφημα που φαίνεται στην Εικόνα 12.7 δεν παρατηρείται απόκλιση από τη γραμμικότητα.



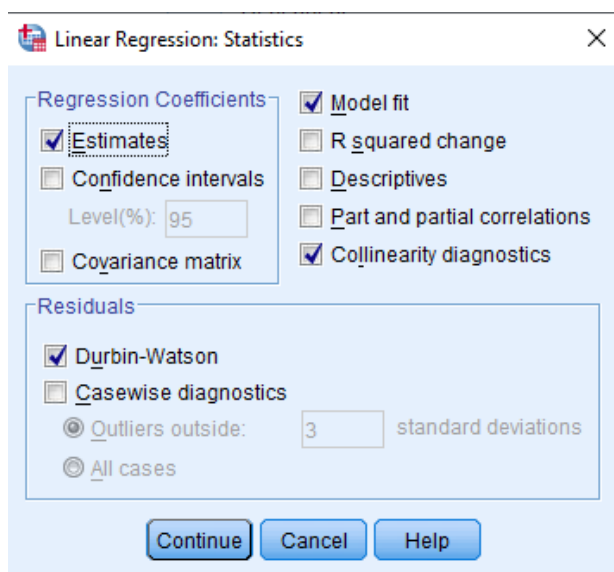
Εικόνα 12.6 Προβλεπόμενες και παρατηρούμενες τιμές της εξαρτημένης μεταβλητής «SQ_DEPR» του τρέχοντος υποδείγματος (α) και του προσδόκιμου επιβίωσης γυναικών από άλλη ανάλυση (β).



Εικόνα 12.7 Γράφημα μερικής παλινδρόμησης για «BEH».

4. Η υπόθεση της ανεξαρτησίας των παρατηρήσεων αναφέρεται στη μη παρουσία ισχυρής συσχέτισης μεταξύ των παρατηρήσεων. Το πρόβλημα μπορεί να προκύψει όταν τα δεδομένα συλλέγονται με μια χρονική σειρά. Για παράδειγμα, αν παρατηρούμε τον χρόνο που απαιτείται να κατασκευαστεί ένα προϊόν μετά την εισαγωγή νέων μηχανημάτων σε μια εταιρεία, καθώς οι υπάλληλοι εξασκούνται περισσότερο στη χρήση των μηχανημάτων, είναι πιθανόν τα προϊόντα στο ξεκίνημα να κατασκευάστηκαν σε περισσότερο χρόνο από εκείνα που κατασκευάστηκαν αργότερα. Συνεπώς, διαδοχικά προϊόντα συνδέονται περισσότερο από αυτό που αναμένεται, αν οι παρατηρήσεις είναι ανεξάρτητες. Ο έλεγχος *Durbin-Watson* μπορεί να διαγνώσει την παρουσία ισχυρής συσχέτισης μεταξύ γειτονικών παρατηρήσεων μελετώντας τη συσχέτιση διαδοχικών υπολοίπων. Τιμές του στατιστικού *Durbin-Watson* κοντά στο 0 δείχνουν ισχυρή θετική συσχέτιση, ενώ τιμές κοντά στο 4 ισχυρή αρνητική. Τιμές του στατιστικού *Durbin-Watson* μεταξύ του 1,5 και 2,5 δεν δημιουργούν ανησυχία. Για να βρεθεί η τιμή του στατιστικού *Durbin-Watson* από την επιλογή **Statistics** στο παράθυρο της **Linear Regression** επιλέγουμε **Durbin-Watson**, όπως φαίνεται στην Εικόνα 12.8 και στην Εικόνα 12.9 βλέπουμε το αποτέλεσμα.
5. Η πολυσυγγραμμικότητα αναφέρεται σε μια κατάσταση κατά την οποία μία ανεξάρτητη μεταβλητή προβλέπεται σε υψηλό βαθμό από άλλες ανεξάρτητες μεταβλητές. Η ποσότητα $1 - R^2$ ονομάζεται *ανοχή* (tolerance) και εκφράζει το ποσοστό της ανεξάρτητης μεταβλητής X , το οποίο δεν εξηγείται από τις υπόλοιπες ανεξάρτητες μεταβλητές, όπου R^2 είναι ο συντελεστής προσδιορισμού από την παλινδρόμηση της X πάνω στις υπόλοιπες ανεξάρτητες μεταβλητές. Οι τιμές της ανοχής κυμαίνονται από το 0 έως το 1. Όσο πιο κοντά στο 0 είναι η τιμή της ανοχής, τόσο μεγαλύτερο είναι το πρόβλημα της πολυσυγγραμμικότητας στα δεδομένα. Αν η ανοχή μιας ανεξάρτητης μεταβλητής είναι $<0,1$ ή ισοδύναμα ο δείκτης πληθωρισμού διακύμανσης VIF (variance inflation factor) ($VIF=1/\text{ανοχή}$) είναι μεγαλύτερος του 10, η πολυσυγγραμμικότητα πιθανόν να επηρεάζει τους συντελεστές της παλινδρόμησης. Είναι πιθανόν σε μια τέτοια περίπτωση το γενικό τεστ F να είναι σημαντικό, απορρίπτοντας την υπόθεση ότι όλοι οι συντελεστές είναι 0, αλλά στον πίνακα των συντελεστών κανένας να μην είναι σημαντικός σύμφωνα με τον έλεγχο t . Επίσης μπορεί να έχουμε συντελεστές παλινδρόμησης με εσφαλμένο πρόσημο. Στην περίπτωση αυτή, πρέπει βρεθεί ποιες μεταβλητές συνδέονται ισχυρά μεταξύ τους (π.χ. με τη βοήθεια του VIF) και να απομακρυνθούν κάποιες απ' αυτές από το υπόδειγμα.

Προκειμένου να πάρουμε τους διαγνωστικούς δείκτες της πολυσυγγραμμικότητας, από την επιλογή **Statistics** στο παράθυρο της **Linear Regression** επιλέγουμε **Collinearity Diagnostics**, όπως φαίνεται στην Εικόνα 12.8 και στον Πίνακα που φαίνεται στην Εικόνα 12.11, ο οποίος είναι ο συνήθης Πίνακας των συντελεστών του υποδείγματος στον οποίο δίνονται οι τιμές *Tolerance* και *VIF*.



Εικόνα 12.8 Επιλογές ελέγχου πολυσυγγραμμικότητας και ανεξαρτησίας παρατηρήσεων.

Model Summary ^b					
Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Durbin-Watson
1	,506 ^a	,256	,230	,24855	1,998

a. Predictors: (Constant), FEMALE, PSY_F, AGE_n, MARKS, BEH, PSY_M
b. Dependent Variable: SQ_DEPR

Η τιμή 1,998 του στατιστικού Durbin-Watson δεν δηλώνει κάποια μορφή συσχέτισης μεταξύ των υπολοίπων της παλινδρόμησης. Δεν αμφισβητείται η ανεξαρτησία των παρατηρήσεων.

Εικόνα 12.9 Δείκτες προσαρμογής και η τιμή του στατιστικού Durbin-Watson.

ANOVA ^a						
Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	3,642	6	,607	9,826	,000 ^b
	Residual	10,564	171	,062		
	Total	14,206	177			

a. Dependent Variable: SQ_DEPR

b. Predictors: (Constant), FEMALE, PSY_F, AGE_n, MARKS, BEH, PSY_M

Εικόνα 12.10 Ανάλυση διακύμανσης για τον έλεγχο της σημαντικότητας του R^2 .

Coefficients ^a								
Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.	Collinearity Statistics	
		B	Std. Error	Beta			Tolerance	VIF
1	(Constant)	-,300	,424		-,706	,481		
	BEH	-,016	,006	-,219	-2,869	,005	,747	1,338
	PSY_M	,006	,007	,071	,871	,385	,661	1,513
	PSY_F	,015	,007	,177	2,226	,027	,691	1,446
	MARKS	,003	,010	,026	,358	,720	,848	1,179
	AGE_n	,074	,025	,207	2,981	,003	,904	1,107
	FEMALE	,202	,040	,358	5,045	,000	,862	1,160

a. Dependent Variable: SQ_DEPR

Για καμία από τις ανεξάρτητες μεταβλητές δεν παρατηρείται τιμή VIF μεγαλύτερη του 10. Συνεπώς, η πολυσυγγραμμικότητα δεν αποτελεί πρόβλημα για τα δεδομένα μας.

Εικόνα 12.11 Συντελεστές του υποδείγματος με τιμές VIF για τον έλεγχο πολυσυγγραμμικότητας.

12.1.4 Η προσαρμογή του υποδείγματος

Μετά τη διαπίστωση της ικανοποίησης των προϋποθέσεων εγκυρότητας του υποδείγματος θα πρέπει να εξετασθεί κατά πόσο είναι καλή η προσαρμογή του. Στον Πίνακα που φαίνεται στην Εικόνα 12.9 η ποσότητα «R Square» μας πληροφορεί πως το 25,6% της μεταβλητότητας στις τιμές της κατάθλιψης ερμηνεύεται από τις έξι ανεξάρτητες μεταβλητές.

Ο συντελεστής R εκφράζει τη συσχέτιση ανάμεσα στην παρατηρούμενη τιμή της εξαρτημένης μεταβλητής και την προβλεπόμενη τιμή της από το υπόδειγμα της παλινδρόμησης. Οι τιμές του R βρίσκονται στο διάστημα $[0,1]$, με την τιμή 1 να αντιστοιχεί στην τέλεια πρόβλεψη της μεταβλητής από το υπόδειγμα και την τιμή 0 να εκφράζει την πλήρη έλλειψη γραμμικής σχέσης ανάμεσα στις ανεξάρτητες μεταβλητές και την εξαρτημένη μεταβλητή. Η τιμή $R=0,506$ εκφράζει μέτριο βαθμό προβλεπτικής ικανότητας του υποδείγματος.

Η ανάλυση διακύμανσης, που φαίνεται στον Πίνακα της εικόνας 12.10 της προηγούμενης ενότητας, ελέγχει τη μηδενική υπόθεση ότι ο συντελεστής R^2 έχει τιμή 0 στον πληθυσμό ή ισοδύναμα ότι όλοι οι μερικοί συντελεστές παλινδρόμησης είναι 0 στον πληθυσμό. Ο έλεγχος βασίζεται στο πηλίκο F του μέσου τετραγώνου της παλινδρόμησης προς το μέσο τετράγωνο των υπολοίπων. Το πηλίκο αυτό στον Πίνακα της εικόνας 12.10 είναι 9,826 και, επειδή το παρατηρούμενο επίπεδο σημαντικότητας p είναι μικρότερο από 0,0005, απορρίπτεται η μηδενική υπόθεση, ότι δεν υπάρχει γραμμική σχέση της κατάθλιψης με τις 6 ανεξάρτητες μεταβλητές. Τούλάχιστον, ένας μερικός συντελεστής παλινδρόμησης διαφέρει από το 0 στον πληθυσμό.

12.1.5 Η σημαντικότητα και η ερμηνεία των μερικών συντελεστών

Οι εκτιμώμενες τιμές των συντελεστών των ανεξάρτητων μεταβλητών, οι οποίες δίνονται στη στήλη Β του Πίνακα 12.3, οδηγούν στην παρακάτω εξίσωση παλινδρόμησης:

$$SQ_DPER = -0,3 - 0,016 \times BEH + 0,006 \times PSYM + 0,015 \times PSYF + 0,003 \times MARKS + 0,074 \times AGE_n + 0,202 \times FEMALE,$$

όπου

SQ_DPER είναι η εκτιμώμενη τιμή της μεταβλητής «SQ_DEPR».

Η ερμηνεία του μερικού συντελεστή B μιας ανεξάρτητης μεταβλητής X είναι η εξής: Όταν η τιμή της ανεξάρτητης μεταβλητής X αυξάνει κατά 1, ενώ οι τιμές των υπολοίπων ανεξάρτητων μεταβλητών παραμένουν σταθερές, τότε η τιμή της εξαρτημένης μεταβλητής μεταβάλλεται κατά B . Θετική τιμή του συντελεστή B ισοδυναμεί με αύξηση της εξαρτημένης μεταβλητής, όταν αυξάνει η τιμή της ανεξάρτητης. Αντίστοιχα, αρνητική τιμή του B εκφράζει τη μείωση της εξαρτημένης μεταβλητής, όταν η ανεξάρτητη μεταβλητή αυξάνει.

Για αύξηση κατά έναν χρόνο της ηλικίας («AGE_n») προβλέπεται αύξηση B_j της μεταβλητής «SQ_DEPR», που εκφράζει το επίπεδο κατάθλιψης. Προκειμένου να γίνει κατανοητή η επίδραση της ηλικίας στις αρχικές τιμές της μεταβλητής «DEPR», δηλαδή η μεταβολή G_j της «DEPR» σε αύξηση κατά έναν χρόνο της ηλικίας, θα πρέπει να χρησιμοποιήσουμε τον μετασχηματισμό $G_j = \frac{\Delta Y}{\Delta X_j} = 2 \times B_j \times (B_0 + B_1 X_1 + \dots + B_k X_k)$. Είναι φανερό ότι η τιμή G_j εξαρτάται από την τιμή της X_j και αυτό οφείλεται στο γεγονός ότι η συνάρτηση της τετραγωνικής ρίζας δεν είναι γραμμική. Ως παράδειγμα, θα υπολογίσουμε τη μεταβολή της «DEPR», όταν από την ηλικία των 15,2 έχουμε μεταβολή κατά ένα έτος, ενώ οι άλλες ανεξάρτητες μεταβλητές έχουν ως τιμές τις αντίστοιχες μέσες τιμές που βρίσκουμε από τα δεδομένα και δίνονται στον παρακάτω Πίνακα. Στην περίπτωση του φύλου, η τιμή 1 δηλώνει ότι η προβλεπόμενη τιμή ισχύει για τα κορίτσια. Χρησιμοποιώντας την αρχική εξίσωση παλινδρόμησης υπολογίζουμε:

$$\sqrt{Y} = (B_0 + B_1 X_1 + \dots + B_k X_k) = (-0,013 + 0,058 \times 15,12 + 0,004 \times 12,07 + 0,019 \times 11,99 - 0,017 \times 18,55 + 0,201 \times 1) = 1,027.$$

Ο υπολογισμός του συντελεστή G_{AGE_n} , $G_{AGE_n} = 2 \times 1,027 \times 0,058 = 0,119$, εκφράζει την αύξηση κατά 0,119 της «DEPR» σε αύξηση κατά 1 χρόνο της ηλικίας των μαθητριών που βρίσκονται σε ηλικία περίπου 15 ετών (15,2 ετών).

	B_j	Μέσες τιμές	G_j
(Constant)	-,013		
AGE_n	,058	15,12	0,119
PSY_M	,004	12,07	0,008
PSY_F	,019	11,99	0,038
BEH	-,017	18,55	-0,034
FEMALE	,201	1	0,413

Πίνακας 12.3 Υπολογισμός μερικών συντελεστών παλινδρόμησης για την «DEPR».

Άλλοι μετασχηματισμοί που στηρίζονται στον φυσικό λογάριθμο της εξαρτημένης ή/και ανεξάρτητων μεταβλητών, οι οποίοι αποτελούν και τη δημοφιλέστερη επιλογή, έχουν την εξής ερμηνεία:

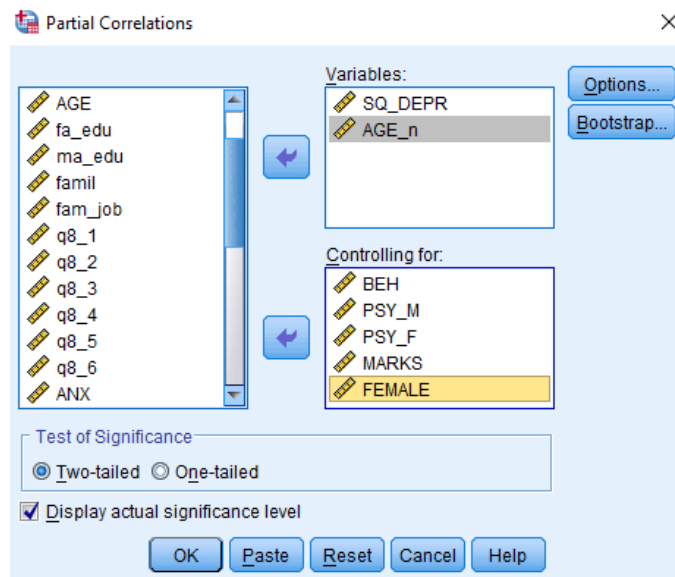
- Φυσικός λογάριθμος μίας ή περισσότερων ανεξάρτητων μεταβλητών: $Y = (B_0 + B_1 \ln(X_1) + \dots + B_k X_k)$. Η ερμηνεία του συντελεστή B_1 παλινδρόμησης είναι ότι 1% μεταβολή της X_1 , οδηγεί σε μεταβολή $B_1 \times 0,01$ της Y .
- Φυσικός λογάριθμος της εξαρτημένης $\ln Y = (B_0 + B_1 X_1 + \dots + B_k X_k) + u$. Η τιμή του συντελεστή δηλώνει ότι αύξηση της X_j κατά μία μονάδα οδηγεί σε μεταβολή της Y κατά $(\exp(B_j) - 1) \times 100\%$.
- Φυσικός λογάριθμος της εξαρτημένης μεταβλητής και μιας ή περισσότερων ανεξάρτητων μεταβλητών $\ln Y = (B_0 + B_1 \ln(X_1) + \dots + B_k X_k) + u$. Η τιμή του συντελεστή δηλώνει ότι αύξηση της X_1 κατά μία μονάδα οδηγεί σε μεταβολή της Y κατά $B_1\%$.

Η στατιστική σημαντικότητα των μερικών συντελεστών της παραπάνω εξίσωσης παλινδρόμησης, δηλαδή η απόφαση για την απόρριψη της μηδενικής υπόθεσης ότι η τιμή του συντελεστή είναι 0 στον πληθυσμό κρίνεται από την τιμή του στατιστικού t και το παρατηρούμενο επίπεδο σημαντικότητας p , που δίνονται αντίστοιχα στις στήλες «t» και «Sig.» του Πίνακα που φαίνεται στην Εικόνα 12.14.

Το πηλίκο $t = \frac{B}{Std. Error}$, όταν ισχύει η μηδενική υπόθεση, ακολουθεί κατανομή t -student με $N-k-1$ βαθμούς ελευθερίας, οι οποίοι στο παράδειγμά μας είναι $BE=178-6-1=171$. Από τον Πίνακα που φαίνεται στην Εικόνα 12.14, προκύπτει, για παράδειγμα, σημαντικότητα ($p<0,05$) των συντελεστών για τις μεταβλητές «BEH» ($B=-0,16$, $Beta=(-0,219)$, $t(171)=-2,87$, $p=0,005$), «PSY_F» ($B=0,015$, $Beta=-0,177$, $t(171)=2,23$, $p=0,027$), «AGE_n» ($B=0,074$, $Beta=-0,177$, $t(171)=2,98$, $p=0,003$) και «FEMALE» ($B=0,202$, $Beta=0,358$, $t(171)=5,05$, $p<0,001$). Οι τιμές των συντελεστών δεν βρέθηκαν σημαντικές για τις μεταβλητές «PSY_M» και «MARKS».

Οι τυποποιημένοι μερικοί συντελεστές παλινδρόμησης $Beta$ είναι οι μερικοί συντελεστές παλινδρόμησης, όταν η παλινδρόμηση εκτελεστεί με τις τυποποιημένες μεταβλητές ($\mu=0$, $\sigma=1$) στη θέση των αρχικών. Οι συντελεστές $Beta$ μπορούν να ερμηνευτούν ως η μεταβολή, με μονάδα την τυπική απόκλιση, που προκύπτει στην εξαρτημένη μεταβλητή από την αύξηση κατά μία τυπική απόκλιση της ανεξάρτητης μεταβλητής X_j , όταν οι άλλες ανεξάρτητες μεταβλητές παραμένουν σταθερές. Για παράδειγμα, διαπιστώνεται ότι σε αύξηση κατά μια τυπική απόκλιση της «PSY_F» έχουμε αύξηση της «SQ_DEPR» κατά 0,177 της τυπικής της απόκλισης. Οι συντελεστές αυτοί αποτελούν ασφαλή τρόπο για τη σύγκριση μεταξύ των ανεξάρτητων μεταβλητών ως προς την ένταση της σχέσης τους με την εξαρτημένη. Η χρήση των συντελεστών B για τη σύγκριση των ανεξάρτητων μεταβλητών δεν είναι κατάλληλη, αφού οι συντελεστές αυτοί εξαρτώνται από τη μονάδα μέτρησης των μεταβλητών αυτών. Για παράδειγμα, αν η ηλικία «AGE_n» δοθεί σε μήνες αντί για έτη, ο συντελεστής B θα γίνει 0,003 από 0,074 και θα είναι ίσος με τον συντελεστή της «PSY_M», ο οποίος δεν είναι στατιστικά σημαντικός και έχει $Beta=0,071$ πολύ μικρότερο από το 0,207 της μεταβλητής «AGE_n».

Ο υπολογισμός του μεγέθους της γραμμικής σχέσης μιας ανεξάρτητης μεταβλητής X_j με την εξαρτημένη Y , κρατώντας σταθερές τις επιδράσεις των άλλων ανεξάρτητων μεταβλητών, μπορεί να γίνει υπολογίζοντας τον συντελεστή μερικής (partial) συσχέτισης pr_j . Ο συντελεστής υπολογίζεται όταν οι επιδράσεις των υπολοίπων μεταβλητών αφαιρεθούν. Προκειμένου να υπολογιστεί ο συντελεστής μερικής συσχέτισης ανάμεσα στην «SQ_DEPR» και την «AGE_n» πρέπει να προβλεφθούν με δύο ξεχωριστές αναλύσεις παλινδρόμησης, με ανεξάρτητες τις «BEH», «PSY_M», «PSY_F», «MARKS» και «FEMALE». Ο συντελεστής μερικής συσχέτισης είναι ο συντελεστής συσχέτισης *Pearson* ανάμεσα στις δύο μεταβλητές υπολοίπων που προκύπτουν από τις δύο αναλύσεις. Στο SPSS μπορούμε να πάρουμε την τιμή του συντελεστή, αν επιλέξουμε [Analyze=>Correlate=>Partial](#) και εισάγουμε στο πλαίσιο [Variables](#) τις «SQ_DEPR» και «AGE_n» και έπειτα στο πλαίσιο [Controlling for](#) εισάγουμε τις «BEH», «PSY_M», «PSY_F», «MARKS» και «FEMALE», όπως παρουσιάζεται στην Εικόνα 12.12, και, τέλος, θα πατήσουμε [OK](#).



Εικόνα 12.12 Επιλογές για τον υπολογισμό της μερικής συσχέτισης.

Correlations				
Control Variables			SQ_DEPR	AGE_n
BEH & PSY_M & PSY_F & MARKS & FEMALE	SQ_DEPR	Correlation	1,000	,222
		Significance (2-tailed)	.	,003
		df	0	171
	AGE_n	Correlation	,222	1,000
		Significance (2-tailed)	,003	.
		df	171	0

Η τιμή του συντελεστή μερικής συσχέτισης είναι στατιστικά σημαντική με $p=0,003$ και μικρότερη από την τιμή του $r=0,24$.

Εικόνα 12.13 Συντελεστής μερικής συσχέτισης μεταξύ «SQ_DEPR» και «AGE_n».

Στον Πίνακα που φαίνεται στην Εικόνα 12.13 η τιμή του συντελεστή μερικής συσχέτισης είναι 0,22 με $p=0,003$. Η τιμή p στον έλεγχο του μερικού συντελεστή B της «AGE_n», από τον Πίνακα στην Εικόνα 12.14, είναι επίσης 0,003. Η ισότητα αυτή των παρατηρούμενων επιπέδων σημαντικότητας δεν είναι τυχαία αλλά οφείλεται στην ισοδυναμία των δύο ελέγχων. Όταν ελέγχεται η μηδενική υπόθεση ότι ο μερικός συντελεστής B_j της X_j είναι 0, συγχρόνως ελέγχεται και η υπόθεση ότι ο συντελεστής μερικής συσχέτισης μεταξύ X_j και Y είναι 0.

Ένας άλλος τύπος συντελεστή συσχέτισης ανάμεσα στην εξαρτημένη μεταβλητή Y και σε μία ανεξάρτητη μεταβλητή X_j είναι ο ημιμερικός (part) συντελεστής συσχέτισης sr_j . Η διαφορά αυτού του συντελεστή από τον pr_j έγκειται στο ότι υπολογισμός του γίνεται από τον συντελεστή *Pearson* ανάμεσα στην Y και στα υπόλοιπα της παλινδρόμησης της X_j πάνω στις υπόλοιπες ανεξάρτητες μεταβλητές. Το τετράγωνο του sr_j^2 είναι το πλέον χρήσιμο μέτρο της σπουδαιότητας της X_j , στον βαθμό που εκφράζει την ατομική συνεισφορά της στη συνολική διακύμανση της Y . Οι συντελεστές που παρουσιάστηκαν παραπάνω, καθώς και ο συντελεστής r , μπορούν να ληφθούν στο SPSS, αν από την επιλογή **Statistics** στο παράθυρο της **Linear Regression** επιλεγεί **Part and partial correlations** (βλ. Εικ. 12.8). Στον Πίνακα που φαίνεται στην Εικόνα 12.14 παρουσιάζονται οι τρεις τύποι συσχετίσεων.

Coefficients ^a								
Model	Unstandardized Coefficients		Standardized Coefficients	t	Sig.	Correlations		
	B	Std. Error	Beta			Zero-order	Partial	Part
1	(Constant)	-.300	,424		-.706	,481		
	BEH	-.016	,006	-.219	-2,869	,005	-,162	-,214
	PSY_M	,006	,007	,071	,871	,385	,247	,066
	PSY_F	,015	,007	,177	2,226	,027	,304	,168
	MARKS	,003	,010	,026	,358	,720	-,076	,027
	AGE_n	,074	,025	,207	2,981	,003	,271	,222
	FEMALE	,202	,040	,358	5,045	,000	,267	,360

a. Dependent Variable: SQ_DEPR

a. Dependent Variable: SQ_DEPR

Ο ημιμερικός συντελεστής (στήλη Part) συσχέτισης μεταξύ της Ηλικίας «AGE_n» και της «SQ_DEPR» είναι 0,197 μικρότερος από τον συντελεστή μερικής συσχέτισης.

Εικόνα 12.14. Συντελεστές: Pearson r (zer-order), μερικής pr_j (Partial) και ημιμερικής (Part) συσχέτισης sr_j .

12.1.6 Η κατασκευή ενός υποδείγματος (παλινδρόμηση πάνω σε όλα τα υποσύνολα των ανεξάρτητων μεταβλητών και Στατιστική Παλινδρόμηση)

Δεν υπάρχει αμφιβολία ότι υπάρχουν μεταβλητές που θα μπορούσαν να βρίσκονται στο υπόδειγμα με σημαντικούς μερικούς συντελεστές και να ενισχύουν την προβλεπτική του δυνατότητα. Επίσης, όπως είδαμε, κάποιες από τις μεταβλητές που συμπεριλάβαμε στο υπόδειγμα δεν έχουν σημαντική δυνατότητα στην πρόβλεψη της κατάθλιψης των εφήβων. Ο στόχος είναι να δημιουργηθεί ένα πιο λιτό υπόδειγμα αφαιρώντας μη χρήσιμες ανεξάρτητες μεταβλητές. Όταν συμπεριληφθούν στο υπόδειγμα μη σημαντικές μεταβλητές, αυξάνονται οι τιμές των τυπικών σφαλμάτων των συντελεστών χωρίς να βελτιώνεται η προβλεπτική του ικανότητα (R^2). Αν αφαιρεθούν από το υπόδειγμα σημαντικές προβλεπτικές μεταβλητές, τότε υπάρχει μεροληψία στην εκτίμηση των παραμέτρων του.

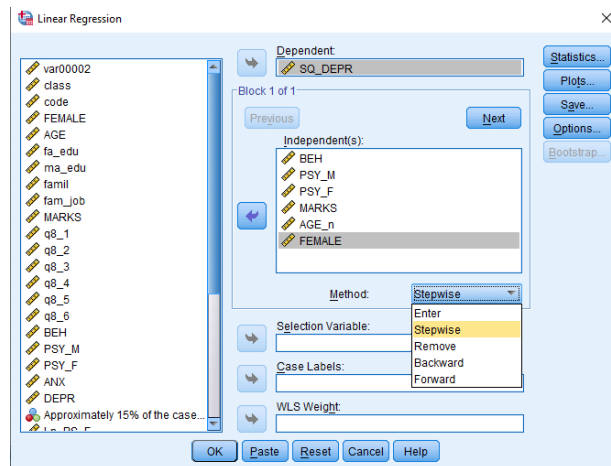
Παρακάτω δίνονται ορισμένες τεχνικές που οδηγούν στην επιλογή μιας ομάδας ανεξάρτητων μεταβλητών, οι οποίες προβλέπουν ικανοποιητικά την εξαρτημένη: η *παλινδρόμηση όλων των υποσυνόλων*, η *προς τα εμπρός επιλογή*, η *προς τα πίσω εξάλειψη* και η *σταδιακή παλινδρόμηση*.

1. Μια τεχνική που δεν υποστηρίζεται από το SPSS είναι η εκτέλεση της παλινδρόμησης για κάθε υποσύνολο ανεξάρτητων μεταβλητών. Στο παράδειγμά μας, με 6 ανεξάρτητες μεταβλητές θα έπρεπε να εκτελεστούν $2^6 = 64$ παλινδρομήσεις υποδειγμάτων με όλα τα υποσύνολα από το σύνολο των 6 αρχικών ανεξάρτητων μεταβλητών. Δηλαδή, 1 υπόδειγμα μόνο με σταθερά, 6 υποδείγματα με 1 μεταβλητή, 15 υποδείγματα με 2 μεταβλητές, 20 υποδείγματα με 3 μεταβλητές, 15 υποδείγματα με 4 ανεξάρτητες, 6 υποδείγματα με 5 μεταβλητές και, τέλος, 1 υπόδειγμα με 6 μεταβλητές. Στη συνέχεια, η χρήση ενός κριτηρίου, όπως ο συντελεστής προσαρμογής R^2 , μπορεί να χρησιμοποιηθεί για την επιλογή του υποσυνόλου προβλεπτικών μεταβλητών. Το μειονέκτημα της μεθόδου αυτής, πέρα από το γεγονός ότι είναι χρονοβόρα όταν ο αριθμός των ανεξάρτητων μεταβλητών είναι σημαντικός, είναι ότι υπάρχει ο κίνδυνος να επιλεγούν εκείνα τα υποδείγματα που προσαρμόζονται καλύτερα σε ιδιάζοντα δεδομένα που υπάρχουν μόνο στο σύνολο δεδομένων που αναλύουμε.
2. Προς τα εμπρός επιλογή (Forward). Το υπόδειγμα στο πρώτο βήμα αυτής της μεθόδου είναι εκείνο με μόνο σταθερό όρο την εξίσωση. Σε κάθε βήμα προστίθεται μία ανεξάρτητη μεταβλητή από αυτές που είναι εκτός υποδείγματος με το κριτήριο της μέγιστης αύξησης του R^2 με την προϋπόθεση ότι η μεταβολή ΔR^2 είναι στατιστικά σημαντική (σε επίπεδο σημαντικότητας $\alpha=0,05$). Το πρόγραμμα σταματά να εισάγει νέες μεταβλητές όταν δεν έχουν μείνει μεταβλητές εκτός υποδείγματος, οι οποίες με την είσοδο τους σε αυτό θα προκαλούσαν σημαντική αύξηση στον R^2 . Παρακάτω, στη σταδιακή επιλογή, δίνονται περισσότερες λεπτομέρειες και παράδειγμα με το υπόδειγμα της «SQ_DEPR» που μελετάται σε αυτό το Κεφάλαιο καθώς η *σταδιακή επιλογή* αποτελεί σύνθεση των μεθόδων της *προς τα εμπρός επιλογής* και της *προς τα πίσω εξάλειψης*.
3. Προς τα πίσω εξάλειψη (Backward). Στην τεχνική αυτή όλες οι ανεξάρτητες μεταβλητές εισάγονται στην εξίσωση του υποδείγματος και στη συνέχεια αφαιρούνται διαδοχικά. Η μεταβλητή με τη μικρότερη ημιμερική συσχέτιση με την εξαρτημένη μεταβλητή είναι αυτή που προκαλεί τη μικρότερη μεταβολή ΔR^2 στον R^2 μεταξύ των μεταβλητών για τις οποίες η ΔR^2 δεν είναι στατιστικά σημαντική

($p > 0,1$), αν απομακρυνθούν από το υπόδειγμα. Αφού αφαιρεθεί η πρώτη μεταβλητή, για τις μεταβλητές που παραμένουν στην εξίσωση ακολουθούμε την ίδια διαδικασία. Η διαδικασία σταματά όταν δεν υπάρχουν μεταβλητές στην εξίσωση που να ικανοποιούν τα κριτήρια απομάκρυνσης.

4. Σταδιακή επιλογή (stepwise). Αυτή είναι η συχνότερα χρησιμοποιούμενη μέθοδος. Σε κάθε βήμα εισάγεται μία ανεξάρτητη μεταβλητή από εκείνες που δεν υπάρχουν στην εξίσωση και η οποία βελτιώνει περισσότερο από όλες τον συντελεστή προσδιορισμού R^2 με την προϋπόθεση ότι η μεταβολή ΔR^2 είναι στατιστικά σημαντική (η τιμή του F $p < 0,05$). Μετά την εισαγωγή της νέας μεταβλητής στο υπόδειγμα, εκτιμώνται εκ νέου οι συντελεστές των μεταβλητών του και από τις μεταβλητές που υπάρχουν ήδη στην εξίσωση παλινδρόμησης αφαιρούνται αυτές που συνοδεύονται από μη σημαντική ΔR^2 ή ισοδύναμα μη σημαντικό μερικό συντελεστή B ή ($p > 0,1$). Η μέθοδος τερματίζεται όταν δεν μπορούν να συμπεριληφθούν ή να αφαιρεθούν άλλες μεταβλητές στο υπόδειγμα. Η σταδιακή επιλογή στο SPSS για το υπόδειγμα της «SQ_DEPR» πραγματοποιείται αν από τη λίστα [Method](#) στο κεντρικό παράθυρο της [Linear Regression](#) επιλεγεί [Stepwise](#) (Εικόνα 12.15).

- Στο πρώτο βήμα εισάγεται η μεταβλητή με τη μεγαλύτερη σε απόλυτη τιμή στατιστικά σημαντική συσχέτιση με την εξαρτημένη μεταβλητή. Αν δεν υπάρχει σημαντική συσχέτιση με καμία από τις ανεξάρτητες μεταβλητές, το υπόδειγμα δεν έχει καμία προβλεπτική αξία (δεν απορρίπτεται η $R^2 = 0$). Στην περίπτωσή μας, η πρώτη μεταβλητή που εισάγεται είναι η «PSY_F», αφού είναι αυτή με την υψηλότερη σημαντική συσχέτιση ($r = 0,304$, $p < 0,001$) μαζί με την «SQ_DEPR» από όλες τις ανεξάρτητες μεταβλητές.
- Στο δεύτερο βήμα εισάγεται η μεταβλητή που από τις στατιστικά σημαντικές μεταβολές προκαλεί τη μεγαλύτερη βελτίωση στον συντελεστή R^2 . Ο έλεγχος της υπόθεσης ότι μεταβολή ΔR^2 είναι 0, όταν μια μεταβλητή εισέλθει στο υπόδειγμα, είναι απολύτως ισοδύναμος με τον έλεγχο της υπόθεσης ότι ο μερικός συντελεστής B είναι 0. Στην Εικόνα 12.16 παρατηρούμε ότι μεταξύ των μεταβλητών «BEH», «PSY_M», «MARKS,AGE_n» και «FEMALE», που βρίσκονται εκτός υποδείγματος μετά την ολοκλήρωση του βήματος 1, η μεταβλητή με τη μεγαλύτερη και σημαντική τιμή Beta (Beta=0,274, $t=3,97$, $p < 0,001$) είναι η «FEMALE». Στη στήλη «Tolerance» (Εικόνα 12.17) δίνεται ο δείκτης ανοχής, ο οποίος εκφράζει την πολυσυγγραμμικότητα της μεταβλητής, αν εισαχθεί στο υπόδειγμα. Μεταβλητές με πολύ μικρή τιμή ανοχής ($< 0,1$) δεν εισάγονται στο υπόδειγμα επειδή προκαλούν υπολογιστικά προβλήματα. Για τη μεταβλητή «FEMALE» η ανοχή είναι 1,0 και δεν εμποδίζει την είσοδό της στο υπόδειγμα. Μετά την είσοδο της «FEMALE» στο υπόδειγμα, για την οριστικοποίηση της σύνθεσης του υποδείγματος κατά το 2^ο βήμα, ελέγχεται η σημαντικότητα των μερικών συντελεστών των «PSY_F», «FEMALE» (Εικόνα 12.16), όπου διαπιστώνεται ότι για καμία μεταβλητή δεν προκύπτει $p > 0,1$. Συνεπώς, και οι δύο μεταβλητές παραμένουν και οριστικοποιείται η σύνθεση του υποδείγματος κατά το 2^ο βήμα.
- Με αυτόν τον τρόπο συνεχίζεται η επιλογή μεταβλητών που συμβάλλουν σημαντικά στη βελτίωση του συντελεστή R^2 του υποδείγματος. Στο τελευταίο 4^ο βήμα δεν υπάρχει μεταβλητή που μπορεί να συμπεριληφθεί στο υπόδειγμα, αφού καμία από τις «PSY_M» και «MARKS» δεν έχει σημαντική μερική συσχέτιση με την εξαρτημένη μεταβλητή, όταν η σχέση τους ελέγχεται από τις «PSY_F», «FEMALE», «AGE_n» και «BEH» (Εικόνα 12.17), οι οποίες βρίσκονται ήδη στο υπόδειγμα. Συνεπώς, η σταδιακή εισαγωγή περατώνεται και στην τελική σύνθεση του υποδείγματος βρίσκονται οι ανεξάρτητες μεταβλητές «PSY_F», «FEMALE», «AGE_n» και «BEH».



Εικόνα 12.15 Επιλογή μεθόδου κατασκευής λιτότερου υποδείγματος.

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	,616	,074		8,354	,000
	PSY_F	,025	,006	,304	4,236	,000
2	(Constant)	,535	,074		7,253	,000
	PSY_F	,026	,006	,310	4,492	,000
	FEMALE	,155	,039	,274	3,971	,000
3	(Constant)	-,578	,367		-1,575	,117
	PSY_F	,021	,006	,255	3,669	,000
	FEMALE	,158	,038	,280	4,157	,000
	AGE_n	,077	,025	,215	3,094	,002
4	(Constant)	-,214	,377		-,567	,571
	PSY_F	,017	,006	,206	2,954	,004
	FEMALE	,203	,040	,360	5,099	,000
	AGE_n	,076	,024	,212	3,114	,002
	BEH	-,017	,005	-,226	-3,111	,002

a. Dependent Variable: SQ_DEPR

Στο πρώτο Βήμα εισάγεται η «PSY_F» και στο τελευταίο βήμα η «BEH». Όπως είναι φανερό, σε κάθε βήμα εισάγεται μια νέα μεταβλητή, αλλά σε κανένα βήμα δεν αφαιρείται καμία από εκείνες που βρίσκονταν στο υπόδειγμα κατά το προηγούμενο βήμα. Οι μεταβλητές ενός βήματος υπάρχουν πάντα και στο επόμενο βήμα.

Εικόνα 12.16 Σταδιακές μεταβολές και τελικό υπόδειγμα με τη μέθοδο «Stepwise».

Model		Beta In	t	Sig.	Partial Correlation	Collinearity Statistics Tolerance
1	BEH	-,097 ^b	-1,315	,190	-,099	,947
	PSY_M	,118 ^b	1,400	,163	,105	,716
	MARKS	-,014 ^b	-,189	,850	-,014	,958
	AGE_n	,207 ^b	2,843	,005	,210	,936
	FEMALE	,274 ^b	3,971	,000	,288	1,000
2	BEH	-,230 ^c	-3,091	,002	-,228	,822
	PSY_M	,147 ^c	1,811	,072	,136	,711
	MARKS	-,063 ^c	-,886	,377	-,067	,936
	AGE_n	,215 ^c	3,094	,002	,228	,935
3	BEH	-,226 ^d	-3,111	,002	-,230	,821
	PSY_M	,115 ^d	1,435	,153	,108	,698
	MARKS	-,035 ^d	-,493	,623	-,037	,913
4	PSY_M	,068 ^e	,841	,402	,064	,668
	MARKS	,019 ^e	,269	,788	,021	,857

a. Dependent Variable: SQ_DEPR
b. Predictors in the Model: (Constant), PSY_F
c. Predictors in the Model: (Constant), PSY_F, FEMALE
d. Predictors in the Model: (Constant), PSY_F, FEMALE, AGE_n
e. Predictors in the Model: (Constant), PSY_F, FEMALE, AGE_n, BEH

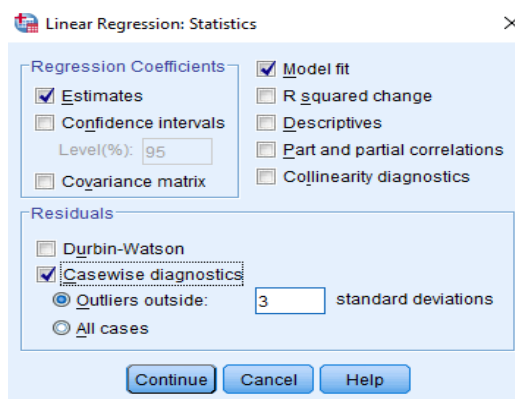
Στο δεύτερο Βήμα μετά την εισαγωγή της «PSY_F», η επιλεγόμενη μεταβλητή είναι η «FEMALE». Η τελευταία είναι αυτή με τη μεγαλύτερη στατιστικά σημαντική μερική συσχέτιση με την εξαρτημένη «SQ_DEPR» όταν η σχέση τους ελέγχεται ως προς την «PSY_F» η οποία βρίσκεται ήδη στο υπόδειγμα. Στο τελευταίο 4^ο βήμα δεν υπάρχει μεταβλητή που μπορεί να συμπεριληφθεί στο υπόδειγμα, αφού καμία από τις «PSY_M» και «MARKS» δεν έχει σημαντική μερική συσχέτιση με την εξαρτημένη μεταβλητή, όταν η σχέση τους ελέγχεται από τις «PSY_F», «FEMALE», «AGE_n» και «BEH» που βρίσκονται ήδη στο υπόδειγμα. Έτσι, η σταδιακή εισαγωγή τελειώνει.

Εικόνα 12.17 Μεταβλητές εκτός υποδείγματος και των κριτηρίων εισαγωγής στο υπόδειγμα κατά τα διαδοχικά στάδια κατασκευής του με τη μέθοδο «Stepwise».

12.1.7 Περιπτώσεις με ισχυρή επίδραση στην προσαρμογή και στους συντελεστές του υποδείγματος

Για τον προσδιορισμό περιπτώσεων με ισχυρή με επίδραση στους δείκτες προσαρμογής, όπως ο συντελεστής προσδιορισμού R^2 , στο τυπικό σφάλμα της εκτίμησης SEE αλλά και στις τιμές των μερικών συντελεστών παλινδρόμησης και στα τυπικά σφάλματα των συντελεστών αυτών έχει προταθεί μια σειρά από δείκτες:

- Τα τυποποιημένων κατά student υπόλοιπα, αν έχουν απόλυτη τιμή μεγαλύτερη από το 3 υποδεικνύουν περιπτώσεις με πιθανή ισχυρή επίδραση στην προσαρμογή αλλά και στους συντελεστές του υποδείγματος. Από την επιλογή Save στο παράθυρο της Linear Regression μαρκάρουμε το Casewise diagnostics και στο Outliers δηλώνουμε 3 τυπικές αποκλίσεις (Εικόνα 12.18). Επειδή η πρώτη ανάλυση στο παράδειγμά μας δεν έδωσε αποτέλεσμα, καθώς κανένα τυποποιημένο υπόλοιπο δεν υπερβαίνει απολύτως το 3, στην Εικόνα 12.19 δίνεται το αποτέλεσμα για απόλυτες τιμές μεγαλύτερες του 2.



Εικόνα 12.18 Επιλογή δημιουργίας πίνακα με στοιχεία ακραίων τυποποιημένων υπολοίπων.

Case Number	Std. Residual	SQ_DEPR	Predicted Value	Residual
2	-2,243	,00	,5519	-,55189
4	-2,119	,22	,7451	-,52147
6	-2,108	,50	1,0187	-,51869
26	2,023	1,20	,7065	,49771
40	2,651	1,38	,7260	,65238
109	2,255	1,40	,8416	,55483
112	2,004	1,43	,9397	,49300
179	2,320	1,43	,8609	,57092

a. Dependent Variable: SQ_DEPR

Case number: αριθμός γραμμής της περίπτωσης.
 Std. Residual: Τυποποιημένο υπόλοιπο.
 SQ_DEPR: Παρατηρούμενες τιμές της εξαρτημένης μεταβλητής.
 Predicted Value: προβλεπόμενη τιμή της εξαρτημένης μεταβλητής από το υπόδειγμα.
 Residual: Το απλό υπόλοιπο (η διαφορά SQ_DEPR-Predicted Value).

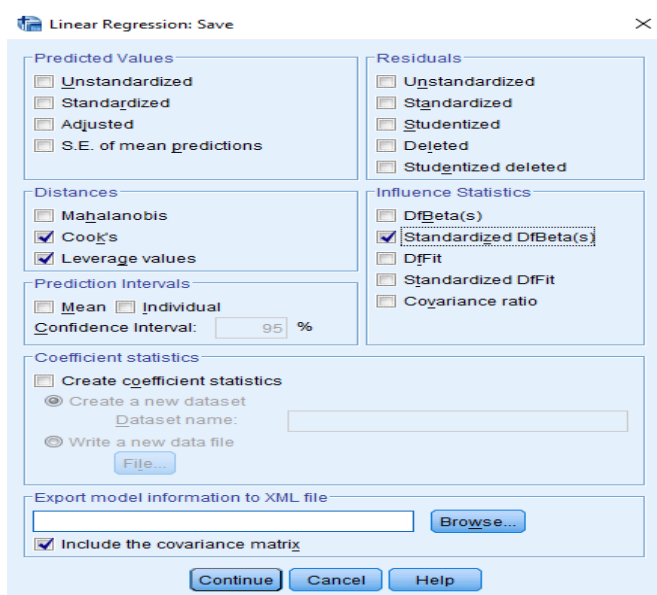
Εικόνα 12.19 Στοιχεία περιπτώσεων με ακραίες τιμές τυποποιημένων υπολοίπων.

Από τα αποτελέσματα που παρουσιάζονται στην Εικόνα 12.19 διαπιστώνεται ότι οι περιπτώσεις στις θέσεις 2, 4, 6, 26, 40, 109, 112, 179 είναι τιμές μιας πιθανής επίδρασης στη διαμόρφωση των συντελεστών της εξίσωσης, η οποία θα διερευνηθεί παραιτέρω με τους υπολοίπους δείκτες που ακολουθούν:

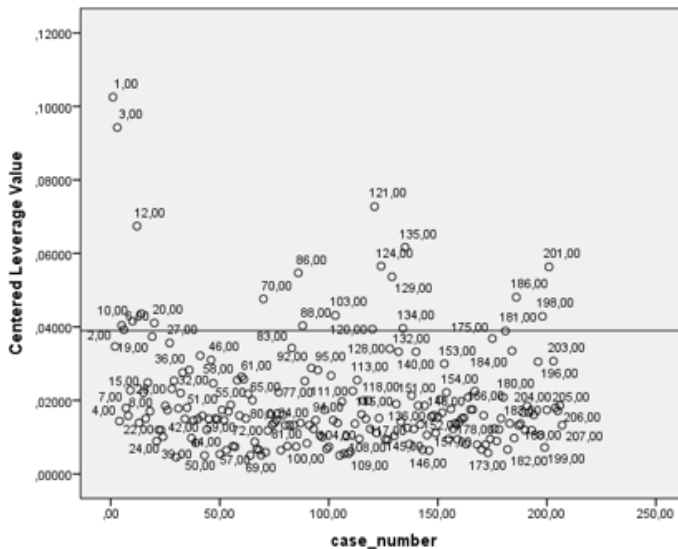
- Απόσταση Leverage. Εκφράζει ασυνήθιστους συνδυασμούς τιμών για τις ανεξάρτητες τιμές των μεταβλητών που χρησιμοποιούνται στο υπόδειγμα. Συγκεκριμένα, είναι μια απόσταση που δείχνει πόσο απέχουν οι τιμές μιας περίπτωσης από τις μέσες τιμές των ανεξάρτητων μεταβλητών. Ένας πρακτικός κανόνας είναι να θεωρούνται τιμές με πιθανή ισχυρή επίδραση αυτές που έχουν μια τιμή Leverage μεγαλύτερη από $\frac{2p}{N}$, όπου p είναι ο αριθμός των ανεξάρτητων μεταβλητών και N το μέγεθος του δείγματος. Στο υπόδειγμα της «SQ_DEPR» με τις τέσσερις ανεξάρτητες μεταβλητές η οριακή αυτή τιμή είναι $\frac{2 \times 4}{101} = 0,0398$. Για τον υπολογισμό των αποστάσεων Leverage στο SPSS από την επιλογή Save στο παράθυρο της Linear Regression μαρκάρουμε Leverage values και στη συνέχεια βγαίνουμε από την επιλογή Save πατώντας Continue (Εικόνα 12.20). Ένας απλός τρόπος, για να ταυτοποιηθούν γραφικά περιπτώσεις με την ισχυρότερη επίδραση, είναι να εξετάσουμε ένα διάγραμμα σκεδασμού με τον αριθμό γραμμής της περίπτωσης (μεταβλητή «CASE_NUMBER» στον οριζόντιο άξονα) και την απόσταση Leverage ως (μεταβλητή «LEV_1» στον κάθετο άξονα). Για τη δημιουργία της μεταβλητής «CASE_NUMBER» θα χρησιμοποιηθεί η συνάρτηση \$CASENUM από την εντολή Transform=>Compute Variable (βλ. Κεφ. 2). Στην Εικόνα 12.21 που παρουσιάζεται το αντίστοιχο γράφημα διαπιστώνεται ότι οι παρατηρήσεις με αριθμούς γραμμής 1, 3, 12 και 121,

μεταξύ μιας δεκάδας περιπτώσεων που ξεπερνούν το όριο 0,0398, είναι αυτές που ξεχωρίζουν λόγω της μεγάλης τους απόστασης από τις μέσες τιμές των ανεξάρτητων μεταβλητών.

- Αλλαγή στον μερικό συντελεστή B. Η τιμή DFBeta, που αποθηκεύεται στο αρχείο, είναι η μεταβολή που προκαλείται στην τιμή του μερικού συντελεστή παλινδρόμησης μιας ανεξάρτητης μεταβλητής όταν εκτελεστεί η ανάλυση χωρίς αυτήν την περίπτωση. Για κάθε ανεξάρτητη μεταβλητή το SPSS αποθηκεύει μια μεταβλητή. Με όνομα DFB0 είναι η μεταβολή στη σταθερά της εξίσωσης, με όνομα DFB1 είναι οι μεταβολές για την πρώτη ανεξάρτητη μεταβλητή που βρίσκεται στο υπόδειγμα κ.λπ. Στην πράξη, χρησιμοποιούνται οι τυποποιημένες τιμές για κάθε συντελεστή με όνομα SDB0, SDB1, ..SDB4 τις οποίες αποθηκεύσαμε (Εικόνα 12.20) και για τις οποίες δημιουργήσαμε ενδεικτικά, κατά αναλογία με τις αποστάσεις Leverage, το διάγραμμα σκεδασμού για τις τιμές SDF2 της μεταβλητής «PSY_F». Πρέπει να δοθεί προσοχή σε περιπτώσεις με απόλυτη τιμή μεγαλύτερη από $\frac{2}{\sqrt{N}}$, δηλαδή $\frac{2}{\sqrt{201}} = 0,141$ για τα δεδομένα που μελετώνται. Από το γράφημα (Εικόνα 12.22) μπορεί να διαπιστωθεί ότι οι μαθητές που βρίσκονται στις σειρές 6 και 121 παρουσιάζουν τη μεγαλύτερη επίδραση στην τιμή του συντελεστή αφού όταν αφαιρεθούν από την ανάλυση η τιμή του θα αυξηθεί.
- Η απόσταση Cook παίρνει υπόψη της τόσο την απόσταση Leverage όσο και το τυποποιημένο κατά Student υπόλοιπο και είναι ένα μέτρο ανάδειξης της επίδρασης μιας περίπτωσης σε όλους τους μερικούς συντελεστές της παλινδρόμησης όταν εκτελεστεί η ανάλυση χωρίς αυτήν την περίπτωση. Με την απόσταση αυτή δίνεται η επίδραση μιας περίπτωσης στη συνολική προσαρμογή του υποδείγματος, όπως η τελευταία εκφράζεται από το τυπικό σφάλμα της εκτίμησης SEE. Τιμές της απόστασης Cook μεγαλύτερες από 1 χρειάζονται παραιτέρω διερεύνηση. Στο διάγραμμα σκεδασμού (Εικόνα 12.23) παρουσιάζονται οι περιπτώσεις με συντεταγμένες τον αριθμό γραμμής και την τιμή της απόστασης Cook. Η κατασκευή στο SPSS έγινε σύμφωνα με την κατασκευή της Εικόνας 12.20. Αν και σε καμία περίπτωση δεν υπάρχει απόσταση Cook μεγαλύτερη ή έστω κοντά στην τιμή 1, οι περιπτώσεις που βρίσκονται στις γραμμές 1, 2, 6 και 121 θα προκαλούσαν ελάττωση στο τυπικό σφάλμα της εκτίμησης που έχει την τιμή SEE 0,246. Επιλέγοντας από το Data=>select cases (βλ. Κεφ. 3) περιπτώσεις με τιμή μεταβλητής COOK_1< 0,04, αυτές θα εξαιρεθούν από τις επόμενες αναλύσεις. Εκτελώντας ξανά την παλινδρόμηση προκύπτει μια μικρότερη τιμή, 0,239, για το τυπικό σφάλμα της εκτίμησης, αλλά η διαφορά είναι μικρή και η απόφασή μας είναι να κρατήσουμε το αποτέλεσμα της αρχικής μας ανάλυσης.

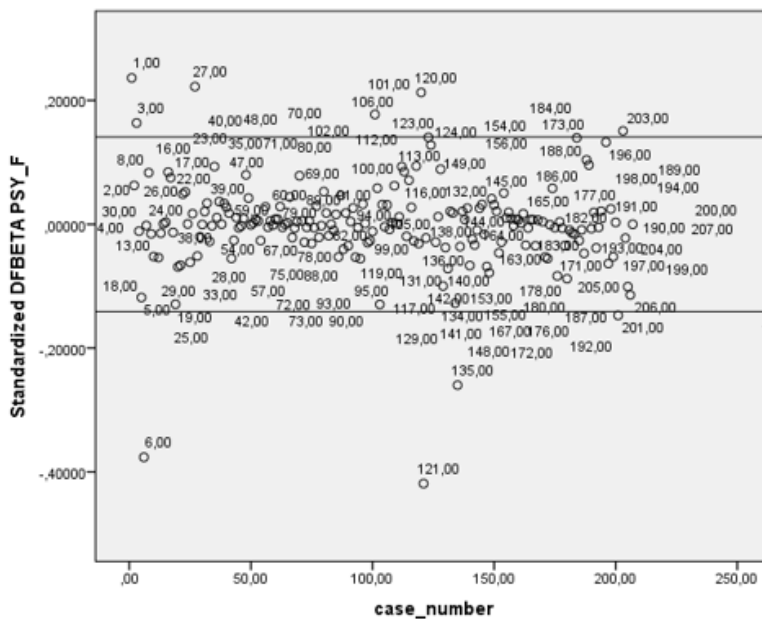


Εικόνα 12.20 Επιλογή δεικτών επίδρασης των περιπτώσεων στις παραμέτρους του υποδείγματος.



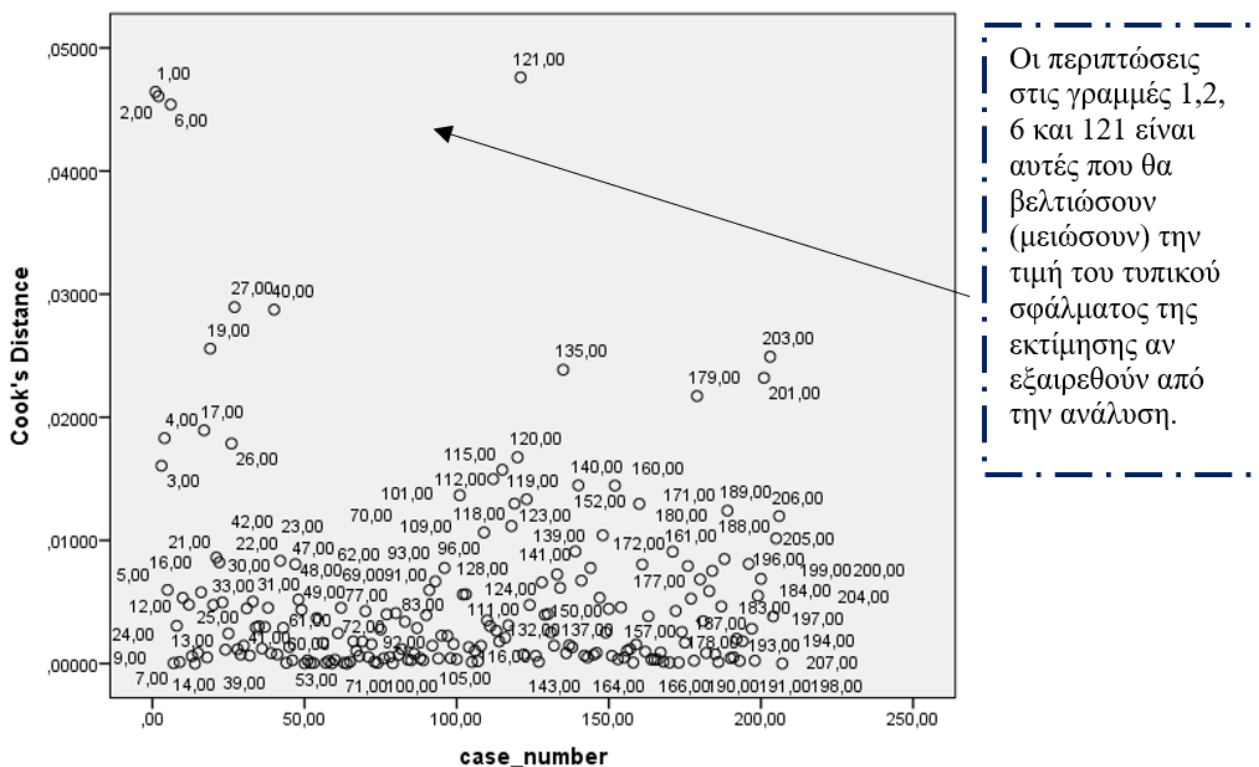
Στον επεξεργαστή γραφήματος πατώντας στο εικονίδιο  δημιουργούμε μια γραμμή που τέμνει τον άξονα της «LEV_1» στην τιμή 0,0398. Σε κάθε περίπτωση που σημειώνεται με κύκλο δίνεται και ο αριθμός γραμμής του αρχείου στην οποία βρίσκεται.

Εικόνα 12.21 Διάγραμμα σκεδασμού «CASE_NUMBER» με τη «LEV_1».



Στον επεξεργαστή γραφήματος πατώντας στο εικονίδιο  δημιουργούμε δύο γραμμές που τέμνουν τον άξονα της «SDB_2_1» με τιμή 0,141 και -0,141. Σε κάθε περίπτωση που σημειώνεται με κύκλο δίνεται και ο αριθμός γραμμής του αρχείου στην οποία βρίσκεται. Αν ο μαθητής της γραμμής 121 αφαιρεθεί από την ανάλυση, θα προκληθεί αύξηση του συντελεστή της «PSY_F» κατά 0,002 και η τιμή του από 0,017 θα αλλάξει σε 0,019.

Εικόνα 12.22 Διάγραμμα σκεδασμού «CASE_NUMBER» με τη «SDB2_1».



Εικόνα 12.23 Διάγραμμα σκεδασμού «CASE_NUMBER» με την «COOK_I».

12.1.8. Η διασταυρούμενη εγκυρότητα του υποδείγματος

Ένα τελικό πρόβλημα που πρέπει να επιλυθεί είναι η διαδικασία ελέγχου της εγκυρότητας του υποδείγματος της Ανάλυσης Παλινδρόμησης. Με αυτήν τη διαδικασία θα στηρίξουμε τη γενίκευση των συμπερασμάτων μας στον πληθυσμό και θα αποδυναμώσουμε την υπόθεση ότι αυτά ισχύουν μόνο στο συγκεκριμένο δείγμα. Μια συνήθης πρακτική είναι η σύγκριση του R^2 με το προσαρμοσμένο R^2 (adjusted R square). Μικρή διαφορά ανάμεσα στα δύο οδηγεί στο συμπέρασμα ότι το υπόδειγμα δεν είναι υπερ-προσαρμοσμένο στο δείγμα και υπάρχει επαρκής εγκυρότητα και κατάλληλος αριθμός παρατηρήσεων ανά μεταβλητή. Μια δεύτερη προσέγγιση, που είναι και προτεινόμενη, είναι η διαίρεση του αρχικού δείγματος σε δύο περίπου ισομεγέθη δείγματα και η εκτέλεση της ανάλυσης σε ένα από αυτά, το οποίο ονομάζεται *δείγμα διαλογής* (screening sample). Στη συνέχεια, θα χρησιμοποιήσουμε τους μερικούς συντελεστές παλινδρόμησης $B_0, B_1, \dots, B_k$ της ανάλυσης από το δείγμα διαλογής, για να εκτιμήσουμε τις προβλεπόμενες τιμές $\hat{Y}$ των ατόμων του δεύτερου δείγματος διακρίβωσης (calibration sample) από τα δεδομένα των ανεξάρτητων μεταβλητών αυτού του δείγματος. Στη συνέχεια, το τετράγωνο της συσχέτισης των $\hat{Y}$ με τις παρατηρούμενες τιμές Y του δείγματος διακρίβωσης θα δώσει την τιμή R^2 της προσαρμογής του. Η μικρή διαφορά ανάμεσα στον R^2 του δείγματος διαλογής και στον R^2 του δείγματος διακρίβωσης αντιστοιχεί στην επιθυμητή εγκυρότητα του υποδείγματος.

12.2 Προβλήματα για εξάσκηση και ανατροφοδότηση

Άσκηση 1

Μια μελέτη για την αντιλαμβανόμενη ποιότητα ζωής (Z) σε έναν μεγάλο αριθμό πόλεων (μονάδα ανάλυσης: η πόλη) έδωσε την παρακάτω εξίσωση χρησιμοποιώντας ως ερμηνευτικές μεταβλητές τη μέση θερμοκρασία (Θ), τη διάμεσο εισόδημα (Ε), την κατά κεφαλήν δαπάνη σε κοινωνικές υπηρεσίες (Κ) και την πυκνότητα πληθυσμού (Π):

$$\hat{Z} = 5,40 - 0,02 \cdot \Theta + 0,05 \cdot E + 0,003 \cdot K - 0,01 \cdot \Pi$$

- Να ερμηνευτούν οι συντελεστές της εξίσωσης.
- Ποια είναι η προβλεπόμενη τιμή αναφορικά με την ποιότητα ζωής σε μια πόλη με μέση θερμοκρασία 20° C, διάμεσο εισόδημα 10.000 ευρώ, κατά κεφαλήν δαπάνη 300 ευρώ για κοινωνικές υπηρεσίες και πυκνότητα 180 κατοίκους στο οικοδομικό τετράγωνο;
- Πόσο διαφέρουν ως προς την κλίμακα της αντιλαμβανόμενης ζωής δύο πόλεις, αν διαφέρουν μόνο ως προς το εισόδημα (8000 και 10500 αντίστοιχα);
- Αν για τους συντελεστές των ερμηνευτικών μεταβλητών της εξίσωσης τα τυπικά σφάλματα ήταν {0,015, 0,011 0,002 και 0,001} αντίστοιχα, ποια είναι η μεταβλητή η οποία είναι πιθανότερο να εξαιρεθεί, αν θα θέλατε να βελτιώσετε το υπόδειγμα και ποια θα εμφανίζει την υψηλότερη τιμή beta;

Άσκηση 2

Δίνονται οι δύο πίνακες συσχετίσεων μεταξύ τριών μεταβλητών από διαφορετικές μελέτες. Ο συντελεστής προσδιορισμού R^2 , που προκύπτει από την Ανάλυση Πολλαπλής Παλινδρόμησης της Y πάνω στις $X1$ και $X2$, είναι 0,995 και 0,803 αντίστοιχα.

Μελέτη A: Correlations			
	X1	X2	Y
X1 Pearson Correlation	1	,029	,487**
Sig. (2-tailed)		,687	,000
N	200	200	200
X2 Pearson Correlation	,029	1	,884**
Sig. (2-tailed)	,687		,000
N	200	200	200
Y Pearson Correlation	,487**	,884**	1
Sig. (2-tailed)	,000	,000	
N	200	200	200
**. Correlation is significant at the 0.01 level (2-tailed).			

Μελέτη B: Correlations			
	X1	X2	Y
X1 Pearson Correlation	1	,533**	,571**
Sig. (2-tailed)		,000	,000
N	200	200	200
X2 Pearson Correlation	,533**	1	,889**
Sig. (2-tailed)	,000		,000
N	200	200	200
Y Pearson Correlation	,571**	,889**	1
Sig. (2-tailed)	,000	,000	
N	200	200	200
**. Correlation is significant at the 0.01 level (2-tailed).			

- Γιατί στη Β μελέτη το R^2 είναι μικρότερο από εκείνο της Α, ενώ από τον πίνακα των συσχετίσεων προκύπτει ότι τόσο η $X1$ όσο και η $X2$ συνδέονται ισχυρότερα με την Y στη Β μελέτη;
- Ποιας ανεξάρτητης μεταβλητής και σε ποια μελέτη ο τυποποιημένος συντελεστής (Beta) αναμένεται να επηρεαστεί δραματικά σε σχέση με την απλή συσχέτιση της μεταβλητής με την εξαρτημένη μεταβλητή του υποδείγματος;

Άσκηση 3

Από το αρχείο «STATISTICS GRADE & SATS.sav» να επιλεγεί ένα τυχαίο δείγμα που θα περιλαμβάνει το 70% των περιπτώσεων του αρχείου και θα αποθηκευτεί σε ξεχωριστό αρχείο με όνομα «STAT_1.sav». Το υπόλοιπο 30% των περιπτώσεων θα αποθηκευτεί σε ένα δεύτερο αρχείο με το όνομα «STAT_2» το οποίο θα χρησιμοποιηθεί στο επόμενο πρόβλημα 5.

Άσκηση 4

Να εκτελεστεί μια Πολλαπλή Παλινδρόμηση στο SPSS με τα δεδομένα του «STAT_1», με την οποία προβλέπεται η επίδοση σε ένα εισαγωγικό μάθημα Στατιστικής (μεταβλητή «βαθμός Στατιστικής», Y) στα πλαίσια του προγράμματος σπουδών ενός παιδαγωγικού τμήματος. Οι ερμηνευτικές μεταβλητές είναι η «Στάση απέναντι στη στατιστική» X1, η οποία εκφράζει την κλίμακα του γνωστικού τομέα του SATS κατά την έναρξη του Μαθήματος, η «Εμπλοκή στη Μαθησιακή διαδικασία» X2, η «Επίδοση στα Μαθηματικά του Λυκείου» X3, και η «Στατιστική Εμπειρία» X4.

- Να γίνουν οι κατάλληλες επιλογές στο SPSS για να δημιουργούν τα γραφήματα που θα χρησιμοποιηθούν για τον έλεγχο των προϋποθέσεων της κανονικότητας και της ομοσκεδαστικότητας των υπολοίπων αλλά και της γραμμικότητας των προβλεπόμενων τιμών απέναντι στις παρατηρούμενες. Επιπλέον, να γίνουν επιλογές στο SPSS προκειμένου να εμφανιστούν τα κατάλληλα στοιχεία στους πίνακες παλινδρόμησης, έτσι ώστε να γίνουν οι έλεγχοι της πολυσυγγραμμικότητας και της ανεξαρτησίας των υπολοίπων.
- Ποιο είναι το συμπέρασμά σας σχετικά με την ικανοποίηση αυτών των προϋποθέσεων στη συγκεκριμένη περίπτωση; Αν χρειάζεται να γίνουν κάποιοι μετασχηματισμοί στις τιμές των μεταβλητών, ποιες μεταβλητές θα επιλέγατε και γιατί;
- Αν στο δεύτερο ερώτημα έχετε απαντήσει θετικά, μεταβείτε στο τέταρτο ερώτημα για να γίνουν οι απαραίτητες αλλαγές και να εκτελέσετε ξανά την ανάλυση, ελέγχοντας εκ νέου την ικανοποίηση των προϋποθέσεων.
- Από τα αποτελέσματα της παλινδρόμησης που κάνατε αποδεκτή, να γραφτεί η εξίσωση που προβλέπει τον βαθμό στη Στατιστική από τις ανεξάρτητες μεταβλητές: «Στάση απέναντι στη Στατιστική», «Μαθησιακή Εμπλοκή», «Επίδοση στα Μαθηματικά του Λυκείου» και «Εμπειρία στη Στατιστική».
- Μπορεί να απορριφτεί η μηδενική υπόθεση ότι όλοι οι μερικοί συντελεστές παλινδρόμησης είναι 0 στον φοιτητικό πληθυσμό από τον οποίο προέρχεται το δείγμα; Πού στηρίζεται το συμπέρασμά σας; Αν η απάντησή σας είναι θετική, να δοθεί και να αξιολογηθεί ως προς το μέγεθός του, το ποσοστό της μεταβλητότητας των βαθμών στη Στατιστική το οποίο ερμηνεύεται από τις ανεξάρτητες μεταβλητές του υποδείγματος.
- Από τις ανεξάρτητες μεταβλητές του υποδείγματος ποια είναι αυτή που προβλέπει καλύτερα τον βαθμό στη Στατιστική; Να εξηγηθεί η απάντησή σας.

Άσκηση 5

Από τη μελέτη των πινάκων που προκύπτουν από την παλινδρόμηση του προηγούμενου προβλήματος να απαντηθούν τα παρακάτω:

- Φαίνεται να υπάρχει λόγος για αφαίρεση κάποιων ανεξάρτητων μεταβλητών από το υπόδειγμα έτσι ώστε να προκύψει ένα λιτότερο;
- Ανεξάρτητα από την απάντησή σας, χρησιμοποιήστε τη μέθοδο «stepwise» για την κατασκευή ενός λιτότερου υποδείγματος. Το αποτέλεσμα συμφωνεί με την απάντησή σας στο πρώτο ερώτημα;
- Να γίνουν κατάλληλες επιλογές στο SPSS, ώστε να υπολογιστούν οι αποστάσεις *Cook*, οι τυποποιημένες τιμές *DFBeta* και να αποθηκευτούν στο αρχείο. Επίσης, να επιλεγεί από το «Statistics» η ταυτοποίηση των περιπτώσεων με ακραία υπόλοιπα. Στη συνέχεια, χρησιμοποιώντας κατάλληλα γραφήματα για τις παραπάνω ποσότητες, να διερευνήσετε την παρουσία περιπτώσεων με σοβαρή επίδραση στην προσαρμογή και τους συντελεστές του υποδείγματος. Αν διαπιστωθεί η παρουσία τέτοιων περιπτώσεων, να εκτελεστεί ξανά η ανάλυση μετά την απομάκρυνσή τους.

- Οι μερικοί συντελεστές παλινδρόμησης B του υποδείγματος που επιλέχτηκε ως καταλληλότερο με τα δεδομένα του «STAT_1» από τις προηγούμενες επιλογές, θα χρησιμοποιηθούν για να υπολογιστούν οι προβλεπόμενες τιμές για τις περιπτώσεις του αρχείου «STAT_2» και θα αποθηκευτούν σε αυτό ως μεταβλητή με το όνομα «Y_PRED». Να χρησιμοποιηθεί κατάλληλα η μεταβλητή αυτή, έτσι ώστε να αξιολογηθεί η εγκυρότητα του υποδείγματος που κατασκευάστηκε με τα δεδομένα του δείγματος του αρχείου «STAT_1». Ποιο είναι το συμπέρασμά σας;

Βιβλιογραφία

- Γναρδέλης, Χ. (2003). *Εφαρμοσμένη Στατιστική*. Αθήνα: Εκδ. Παπαζήση.
- Howell, C. D. (1997). *Statistical Methods in Psychology* (4th ed.). Duxbury Press, by Wadsworth Publishing Co. International Thomson Publishing Inc.
- Norusis, M. J. (2002). *SPSS 11.0 Guide to Data Analysis*. New York: Prentice-Hall.
- Ott, L., & Longnecker, M. (2001). *An Introduction to Statistical Methods and data Analysis* (5th ed.). Duxbury, Pacific Grove.
- Tabachnick, B. G., & Fidell, L. S. (1996). *Using multivariate statistics* (3rd ed.). New York: Harper Collins.

Κεφάλαιο 13

Λογαριθμική Παλινδρόμηση

Σύνοψη

Στο παρόν κεφάλαιο, παρουσιάζεται αρχικά μια εισαγωγή στη λογαριθμική παλινδρόμηση και περιγράφονται οι διαφορές και οι ομοιότητες με την πολλαπλή γραμμική παλινδρόμηση. Στη συνέχεια, παρουσιάζεται η λογαριθμική παλινδρόμηση στο περιβάλλον του SPSS και εξετάζεται το υπόδειγμά της, ακολουθώντας συγκεκριμένο παράδειγμα από τον χώρο της εκπαιδευτικής έρευνας.

Προαπαιτούμενη γνώση

Συσχέτιση, Απλή και Πολλαπλή Γραμμική Παλινδρόμηση: Κεφάλαια 11 και 12 του βιβλίου.

13.1 Εισαγωγή στη λογαριθμική παλινδρόμηση

Η ανάλυση παλινδρόμησης που μελετήθηκε στο Κεφάλαιο 12 αποτελεί μια μέθοδο για τη διερεύνηση των συναρτησιακών σχέσεων μεταξύ των μεταβλητών. Η σχέση εκφράζεται με τη μορφή μιας εξίσωσης ή ενός υποδείγματος που συνδέει την εξαρτημένη μεταβλητή με μία ή περισσότερες επεξηγηματικές ή προγνωστικές μεταβλητές. Οι περισσότερες μεταβλητές σε αυτό το μοντέλο έχουν ποσοτικό χαρακτήρα. Η εκτίμηση παραμέτρων στο υπόδειγμα της παλινδρόμησης βασίζεται σε τέσσερις βασικές παραδοχές. Πρώτον, η εξαρτημένη μεταβλητή σχετίζεται γραμμικά με τις επεξηγηματικές μεταβλητές. Δεύτερον, τα σφάλματα του υποδείγματος (υπόλοιπα) είναι ανεξάρτητα και κατανέμονται με τον ίδιο τρόπο ακολουθώντας την κανονική κατανομή, με μέση τιμή το μηδέν και κοινή διακύμανση. Τρίτον, οι επεξηγηματικές μεταβλητές μετρούνται χωρίς σφάλματα. Η τελευταία παραδοχή αναφέρεται στο ότι όλες οι παρατηρήσεις είναι εξίσου αξιόπιστες.

Σε περίπτωση που η εξαρτημένη μεταβλητή του υποδείγματος είναι κατηγορική, τότε η πιθανότητα μιας περίπτωσης να ανήκει στις διάφορες κατηγορίες της, μπορεί να μοντελοποιηθεί χρησιμοποιώντας το ίδιο μοντέλο. Παρ' όλα αυτά, υπάρχουν πολλοί περιορισμοί όσον αφορά τις παραδοχές αυτού του υποδείγματος πολλαπλής παλινδρόμησης. Πρώτον, δεδομένου ότι το εύρος της πιθανότητας είναι μεταξύ 0 και 1, η συνάρτηση της δεξιάς πλευράς της εξίσωσης στα υποδείγματα πολλαπλής παλινδρόμησης δεν έχει όρια. Δεύτερον, ο όρος σφάλματος του μοντέλου μπορεί να πάρει μόνο περιορισμένες τιμές και η διακύμανση των σφαλμάτων δεν είναι σταθερή αλλά εξαρτάται από την πιθανότητα της τιμής της εξαρτημένης μεταβλητής.

Σε γενικές γραμμές, η γραμμική πολλαπλή παλινδρόμηση χρησιμοποιείται όταν η εξαρτημένη μεταβλητή είναι ποσοτική, ενώ για κατηγορικές εξαρτημένες μεταβλητές και, πιο συγκεκριμένα, για διχοτομικές εξαρτημένες μεταβλητές, είναι προτιμότερο να χρησιμοποιούνται εναλλακτικά υποδείγματα. Διχοτομικές εξαρτημένες μεταβλητές έχουμε, για παράδειγμα, στις παρακάτω περιπτώσεις:

- Ένας οικονομολόγος μπορεί να ενδιαφέρεται να προσδιορίσει την πιθανότητα αποτυχίας μιας βιομηχανίας που βασίζεται στη γεωργία, δεδομένου ενός αριθμού χρηματοοικονομικών δεικτών και του μεγέθους της επιχείρησης (μεγάλη ή μικρή επιχείρηση).
- Μια ερευνητική ομάδα μελετά την επίδραση της αυτο-αποτελεσματικότητας του εκπαιδευτικού, του συνολικού χρόνου εξοικείωσης και τη συμμετοχή του σε σεμινάρια σχετικά με τις ΝΤ (Νέες Τεχνολογίες), πάνω στην πιθανότητα χρήσης των ΝΤ στη διδασκαλία.

Η διακρίνουσα ανάλυση θα μπορούσε να χρησιμοποιηθεί στη διερεύνηση των παραπάνω ερωτημάτων. Όμως, επειδή οι ανεξάρτητες μεταβλητές είναι μια σύνθεση ποσοτικών και κατηγορικών μεταβλητών, η παραδοχή πολυμεταβλητής κανονικότητας μπορεί να μην επαληθευτεί. Στην περίπτωση μια διχοτομικής κατηγορικής μεταβλητής, προτιμώμενες τεχνικές είναι η *λογαριθμική παλινδρόμηση* ή *παλινδρόμηση πιθανότητας* (logistic or probit), οι οποίες δεν απαιτούν καμία παραδοχή για τις κατανομές των ανεξάρτητων μεταβλητών. Οι τεχνικές αυτές έχουν τη σημαντική ιδιότητα να παράγουν προβλεπόμενες τιμές, δηλαδή τις πιθανότητες να συμβεί κάτι μεταξύ του 0 και του 1. Επιπλέον, είναι απαλλαγμένες από τις παραδοχές της πολλαπλής παλινδρόμησης, οι οποίες, όπως αναφέραμε προηγουμένως, δεν ικανοποιούνται.

Στις παρακάτω ενότητες παρουσιάζεται η προσαρμογή του υποδείγματος της λογαριθμικής παλινδρόμησης. Για την εφαρμογή αυτής της ανάλυσης με τη βοήθεια του SPSS, θα χρησιμοποιηθούν τα δεδομένα του αρχείου «Logit regression example barriers.sav» με τις απαντήσεις 134 εκπαιδευτικών-νηπιαγωγών από την έρευνα των Νικολοπούλου και Γιαλαμά (2013) προκειμένου να αναδειχθούν οι παράγοντες που συνδέονται με την πιθανότητα χρήσης των NT στη διδασκαλία. Η εξαρτημένη μεταβλητή του υποδείγματος που θα μελετηθεί είναι η διχοτομική (Ναι/Όχι) και δηλώνει την αξιοποίηση ή μη των υπολογιστών στη διδασκαλία. Οι ανεξάρτητες μεταβλητές του υποδείγματος δίνονται στον Πίνακα 13.1.

Οι Ενότητες του Κεφαλαίου που ακολουθούν δομούνται ως εξής: στην Ενότητα 13.2 δίνεται η διαδικασία εκτέλεσης της Λογαριθμικής Παλινδρόμησης (ΛΠ) στο SPSS. Στις επόμενες Ενότητες 13.3 έως 13.6 παρουσιάζονται οι ορισμοί, οι έννοιες και οι μέθοδοι στα διάφορα βήματα της τεχνικής της ΛΠ και δίνονται παραδείγματα αλλά και η ερμηνεία των αποτελεσμάτων χρησιμοποιώντας τους πίνακες αποτελεσμάτων που προκύπτουν από την εκτέλεση της ανάλυσης στην Ενότητα 13.2.

Όνομα μεταβλητής	Περιγραφή	Τύπος μεταβλητής
lack of support	Έλλειμμα υποστήριξης	ποσοτική
lack of confidence	Έλλειμμα εμπιστοσύνης	ποσοτική
lack of equipment	Έλλειμμα τεχνολογικών πόρων	ποσοτική
class conditions	Δυσμενείς συνθήκες τάξης	ποσοτική
confidence with technology	Αυτοπεποίθηση με NT	ποσοτική
ncomp_01	Παρουσία Υπολογιστών στην τάξη (0=Όχι / 1=Ναι)	Κατηγορική
yearsexp	Εμπειρία τους με τους υπολογιστές σε έτη (1 less than 1 year 2 [1-2] years/ 3 [3-5] years/ 4 more than 5 years)	Κατηγορική με διάταξη
certificA	Πιστοποίηση στις βασικές υπολογιστικές δεξιότητες (0=Όχι / 1=Ναι)	Κατηγορική
yearsteach	Διδακτική τους εμπειρία (1 [1-5] years/ 2 [6-10] years/ 3 [11-15] years/ 4 [16-20] years/ 5 more than 20 years)	Κατηγορική με διάταξη

Πίνακας 13.1 Ερμηνευτικές μεταβλητές της χρήσης NT στη διδασκαλία.

13.1.1 Η λογαριθμική παλινδρόμηση στο SPSS

Για να εκτελέσουμε τη λογαριθμική παλινδρόμηση, από του μενού επιλέγουμε διαδοχικά **Analyze=>Regression=>Binary Logistic** και εμφανίζεται (Εικόνα 13.1) το πλαίσιο διαλόγου **Logistic Regression**. Στην περιοχή **Dependent** θα μεταφέρουμε τη δίτιμη μεταβλητή (με τιμές: Ναι και Όχι) που αντιστοιχεί στην αξιοποίηση των υπολογιστών από τους εκπαιδευτικούς κατά την εκπαιδευτική διαδικασία. Στην περιοχή **Covariates** θα τοποθετήσουμε τις ανεξάρτητες μεταβλητές: «lack_of_support», «lack_of_confidence», «lack_of_equipment», «class_conditions», «confidence_with_technology», «ncomp_01», «yearsexp», «certificA» και «yearsteach». Οι μεταβλητές που τοποθετούνται στην περιοχή **Covariates** μπορεί να είναι ποιοτικές (που θα πρέπει να μετατραπούν σε αντίστοιχες ψευδομεταβλητές) ή και ποσοτικές. Στην περίπτωση που θέλουμε να εισαγάγουμε στο μοντέλο και αλληλεπίδραση μεταξύ μεταβλητών, επιλέγουμε από τη λίστα των μεταβλητών αριστερά, ταυτόχρονα δύο ή περισσότερες μεταβλητές (με την ταυτόχρονη χρήση του **ctrl** από το πληκτρολόγιο και του αριστερού **κλικ** του ποντικιού) και πατάμε το πλήκτρο **>a*b>** που φαίνεται κάτω από το βελάκι μετάβασης των μεταβλητών στην περιοχή **Covariates**. Στην περίπτωση του συγκεκριμένου παραδείγματος δεν διερευνούμε αλληλεπιδράσεις μεταβλητών.

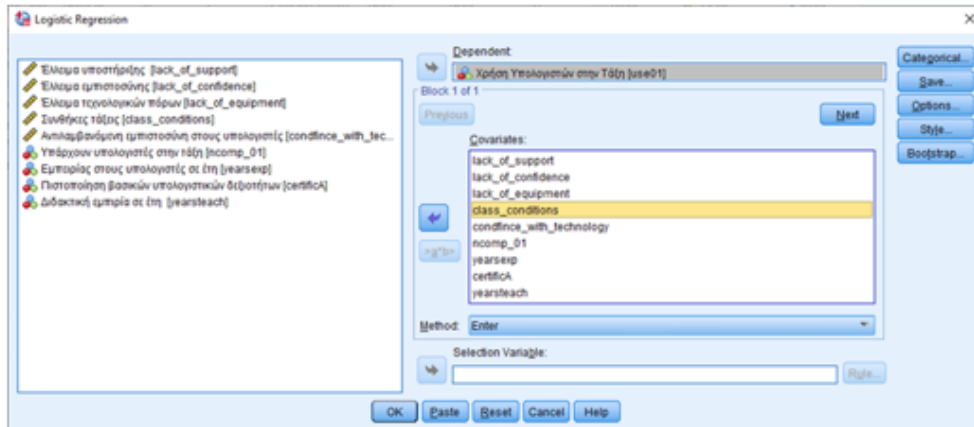
Για μια τυπική λογαριθμική παλινδρόμηση θα πρέπει να αγνοήσουμε τα κουμπιά [Previous](#) και [Next](#), αφού αυτά αντιστοιχούν στην ιεραρχική λογαριθμική παλινδρόμηση. Επίσης, η προεπιλεγμένη μέθοδος είναι το [Enter](#), όπου οι ανεξάρτητες μεταβλητές εισέρχονται στο μοντέλο όλες μαζί σαν μπλοκ. Τέλος, μέσω της επιλογής [Selection Variable](#) μπορούμε να ορίσουμε, τοποθετώντας στο συγκεκριμένο πεδίο μια μεταβλητή, σε ποιο υποσύνολο των δεδομένων θέλουμε να βασιστεί η ανάλυση.

Στην περίπτωση που στο υπόδειγμα εισέρχονται κατηγορικές μεταβλητές με περισσότερες από δύο κατηγορίες ($v > 2$) χρειάζεται να προχωρήσουμε στον ορισμό *ψευδομεταβλητών* (dummy variables) και ταυτόχρονα στον ορισμό της κατηγορίας αναφοράς, έτσι ώστε να είναι εφικτή η αξιοποίηση αυτών στο υπόδειγμα (μοντέλο). Κάθε ψευδομεταβλητή παίρνει την τιμή 1, όταν το υποκείμενο υπέδειξε αυτή την κατηγορία, και 0 σε οποιαδήποτε άλλη περίπτωση. Στο μοντέλο θα εισαχθούν τόσες ψευδομεταβλητές όσες είναι και ο αριθμός των κατηγοριών της μεταβλητής μείον 1 ($v-1$). Η ψευδομεταβλητή που δεν εισάγεται στο μοντέλο αποτελεί το *επίπεδο αναφοράς* (reference category). Η κωδικοποίηση των κατηγορικών μεταβλητών (με ή χωρίς διάταξη) ακολουθεί την ίδια διαδικασία με εκείνη που έχει συζητηθεί στο Κεφάλαιο 12 της Πολλαπλής Παλινδρόμησης με τη χρήση ψευδομεταβλητών. Στην περίπτωση της ΛΠ, όμως, δεν χρειάζεται κανένας χειρισμός προετοιμασίας στο SPSS από τον χρήστη, αφού το πρόγραμμα δημιουργεί αυτομάτως τις ψευδομεταβλητές και ο χρήστης απλώς επιλέγει την κατηγορία αναφοράς.

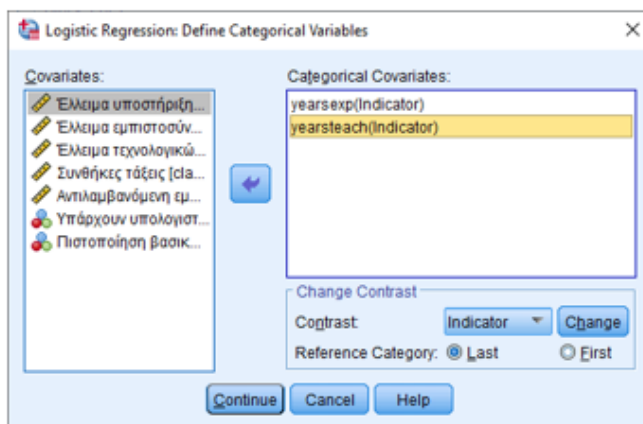
Για τον ορισμό των ψευδομεταβλητών στο πλαίσιο διαλόγου [Logistic Regression](#) (Εικόνα 13.1 (α)) επιλέγουμε [Categorical](#) και έπειτα στο πλαίσιο διαλόγου [Logistic Regression: Define Categorical Variables](#) (Εικόνα 13.1(β)) που εμφανίζεται, περνάμε διαδοχικά τις δύο μεταβλητές «yearsexp» και «yearsteach» στην περιοχή [Categorical Covariates](#). Η προεπιλεγμένη επιλογή [Reference category Last](#) υποδεικνύει ότι θα χρησιμοποιηθεί για τη δημιουργία των ψευδομεταβλητών ως κατηγορία αναφοράς η τελευταία κατηγορία (με τη μεγαλύτερη αριθμητική τιμή) κάθε μεταβλητής.

Στην περίπτωση αυτή, η σχετική πιθανότητα πραγματοποίησης του γεγονότος που αντιστοιχεί στην κατηγορία της μεταβλητής συγκρίνεται με τη σχετική πιθανότητα κατηγορίας αναφοράς. Περισσότερες λεπτομέρειες αναφορικά με τις διαθέσιμες επιλογές ορισμού κατηγορίας αναφοράς θα βρείτε στα βιβλία των Γναρδέλλη (2013) και Field (2018).

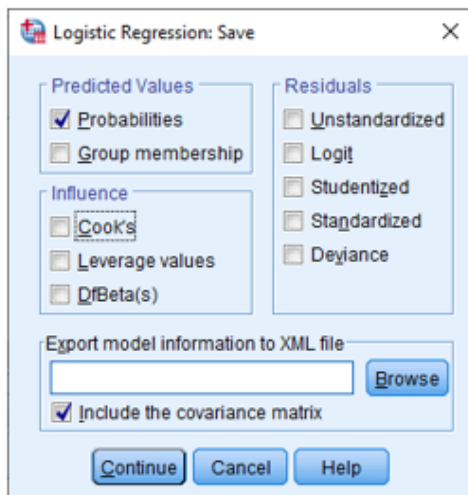
Με την επιλογή [Save](#) του πλαισίου διαλόγου [Logistic Regression](#) μπορούμε να ζητήσουμε στο αντίστοιχο πλαίσιο διαλόγου (Εικόνα 13.1(γ)) να αποθηκευτούν στο αρχείο δεδομένων μια σειρά καινούριων μεταβλητών αλλά και διαγνωστικών στατιστικών μέτρων που μας παρέχει η ανάλυση. Οι μεταβλητές αυτές είναι εκτιμώμενες τιμές (Predicted Values), τα υπόλοιπα της παλινδρόμησης (Residuals), καθώς και κάποια στατιστικά (Influence) που αξιοποιούνται για τον εντοπισμό παρατηρήσεων με ισχυρή επίδραση στη διαμόρφωση του τελικού μοντέλου (Γναρδέλλης, 2013). Από τα παραπάνω εμείς θα ζητήσουμε μόνο την εξαγωγή των [Probabilities](#), δηλαδή των εκτιμώμενων πιθανοτήτων πραγματοποίησης του γεγονότος (αξιοποίηση των υπολογιστών από τους εκπαιδευτικούς κατά τη διδασκαλία) για κάθε υποκείμενο της έρευνας (εκπαιδευτικοί).



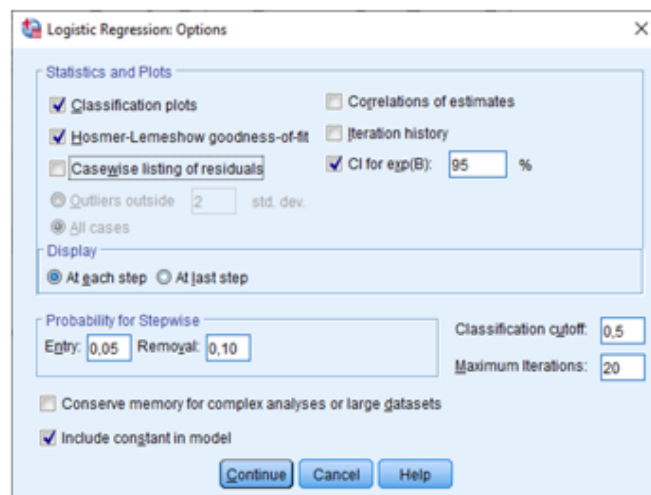
(α)



(β)



(γ)



(δ)

Εικόνα 13.1 Πλαίσιο διαλόγου «Logistic Regression» (α) «Logistic Regression», (β) «Logistic Regression: Define categorical Variables», γ) «Logistic Regression: Save», δ) «Logistic Regression: Options».

Επιπλέον, στην επιλογή **Options** μπορούμε να επιλέξουμε στο αντίστοιχο πλαίσιο διαλόγου (Εικόνα 13.1 (δ)) εξαγωγή στατιστικών μέτρων και γραφικές απεικονίσεις οι οποίες θα υποστηρίξουν τα αποτελέσματα της ανάλυσης αλλά και της ερμηνείας του μοντέλου. Τα κυριότερα από αυτά, τα οποία θα ζητήσουμε στην παρούσα ανάλυση, είναι:

- το διάγραμμα ταξινόμησης των παρατηρήσεων (Classification plot),
- το κριτήριο καλής προσαρμογής των Hosmer-Lemeshow (goodness-of-fit),
- τα 95% διαστήματα εμπιστοσύνης των «exp(B)» (CI for exp(B): 95%) και
- τα υπόλοιπα του υποδείγματος (Casewise listing of residuals).

Επίσης, μεταξύ άλλων μπορούμε να επιλέξουμε κριτήρια για να συμπεριληφθεί (**Entry**) ή όχι (**Removal**) μια μεταβλητή στο μοντέλο, καθώς επίσης και αν το μοντέλο θα περιλαμβάνει σταθερό όρο (**Include constant in model**). Τέλος, στην επιλογή **Classification Cutoff** μπορούμε να ορίσουμε την τιμή της πιθανότητας βάσει της οποίας θα γίνει η επαναταξινόμηση των παρατηρήσεων από το πρόγραμμα. Η χρησιμότητα των επιλογών θα εξηγηθεί στις παρακάτω ενότητες. Από τους αρχικούς πίνακες αποτελεσμάτων (Εικόνα 13.2) ο πίνακας (α) παρουσιάζει τον αριθμό των έγκυρων παρατηρήσεων της ανάλυσης, καθώς και τον αριθμό των παρατηρήσεων με τουλάχιστον μία ελλείπουσα τιμή (missing value) στις μεταβλητές της ανάλυσης, ενώ ο πίνακας (β) δείχνει την κωδικοποίηση των κατηγοριών της εξαρτημένης μεταβλητής. Στην προκειμένη περίπτωση η κωδικοποίηση είναι 0 για το όχι και 1 για το ναι (δηλαδή, με την τιμή 1 υποδεικνύεται η πραγματοποίηση του γεγονότος). Στον επόμενο πίνακα (γ) (Categorical Variables Codings) δίνεται η κωδικοποίηση των κατηγορικών μεταβλητών με βάση το σχήμα κωδικοποίησης για τη δημιουργία ψευδομεταβλητών. Υπενθυμίζουμε ότι έχουμε ορίσει ως κατηγορία αναφοράς τη μεγαλύτερη αριθμητική τιμή της μεταβλητής και η κατηγορία αναφοράς έχει κωδικό μηδέν σε όλες τις άλλες κατηγορίες.

Case Processing Summary			
Unweighted Cases ^a		N	Percent
Selected Cases	Included in Analysis	134	100,0
	Missing Cases	0	,0
	Total	134	100,0
Unselected Cases		0	,0
Total		134	100,0
a. If weight is in effect, see classification table for the total number of cases.			

Dependent Variable Encoding	
Original Value	Internal Value
0 no	0
1 yes	1

Categorical Variables Codings						
		Frequency	Parameter coding			
			(1)	(2)	(3)	(4)
yearsteach	1 [1-5] years	22	1,000	,000	,000	,000
	2 [6-10] years	34	,000	1,000	,000	,000
	3 [11-15] years	16	,000	,000	1,000	,000
	4 [16-20] years	19	,000	,000	,000	1,000
	5 more than 20 years	43	,000	,000	,000	,000
yearsexp	1 less than 1 year	15	1,000	,000	,000	
	2 [1-2] years	17	,000	1,000	,000	
	3 [3-5] years	29	,000	,000	1,000	
	4 more than 5 years	73	,000	,000	,000	

Εικόνα 13.2 Αρχικοί πίνακες αποτελεσμάτων της λογαριθμικής παλινδρόμησης.

13.1.2 Σχετική πιθανότητα και η συνάρτηση «logit»

Η ΛΠ βασίζεται στους λογάριθμους και στον λόγο πιθανοτήτων. *Σχετική πιθανότητα* (Odds ratio) είναι η πιθανότητα να συμβεί ένα γεγονός προς την πιθανότητα να μην συμβεί. Αν p είναι η πιθανότητα να συμβεί ένα γεγονός ($Y=1$) και $1-p$ η πιθανότητα να μην συμβεί ($Y=0$), τότε:

$$\text{σχετική πιθανότητα (OR)} = \frac{p}{1-p}.$$

Αντί να χρησιμοποιηθεί ένα γραμμικό υπόδειγμα για να εξεταστεί η εξάρτηση της πιθανότητας εμφάνισης μιας ασθένειας από τις επεξηγηματικές μεταβλητές, χρησιμοποιείται ο λογαριθμικός μετασχηματισμός, ο οποίος ορίζεται ως εξής:

$$\ln\left(\frac{p}{1-p}\right) = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_k x_k, \quad (13.1)$$

όπου

η έκφραση $\ln\left(\frac{p}{1-p}\right)$ ονομάζεται *logit* της εξαρτημένης Y .

$\beta_0, \beta_1, \beta_2, \dots$ είναι ο σταθερός όρος αλλά και οι συντελεστές παλινδρόμησης των ανεξάρτητων μεταβλητών $x_1, x_2, \dots, x_k$.

Ένα από τα προβλήματα είναι ότι δεν υπάρχει ταύτιση των πεδίων τιμών του αριστερού και του δεξιού μέρους του γραμμικού υποδείγματος. Ενώ, δηλαδή, οι τιμές που παίρνει η παράσταση $\beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_k X_k$, στο δεξί μέρος της (1), ανήκουν σε όλο τον χώρο των πραγματικών αριθμών, ούτε η Y , που παίρνει τιμές 0 και 1, αλλά ούτε και η πιθανότητα εμφάνισης ενός γεγονότος, με πεδίο τιμών $[0,1]$, ικανοποιούν την ισότητα των δύο μερών. Το πρόβλημα αυτό αντιμετωπίζεται με τον λογάριθμο, αφού μπορεί να πάρει τιμές σε όλο το εύρος των πραγματικών αριθμών. Έτσι, με το *logit* μπορούμε να οδηγηθούμε στη σχετική πιθανότητα OR και από αυτήν στην πιθανότητα p να συμβεί το υπό μελέτη γεγονός:

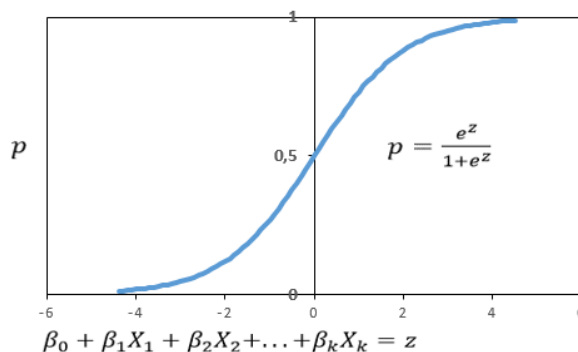
Αν για πρακτικούς λόγους ορίζουμε $\beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_k X_k = z$, από τον ορισμό του φυσικού λογαρίθμου και την (1) έχουμε:

$$\begin{aligned} \ln\left(\frac{p}{1-p}\right) &= z \Rightarrow \\ \frac{p}{1-p} &= e^z \Rightarrow \\ p &= e^z - pe^z \Rightarrow p = \frac{e^z}{1+e^z}. \end{aligned}$$

Δηλαδή:

$$p = \frac{\exp(\beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_k X_k)}{1 + \exp(\beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_k X_k)} \quad (13.2)$$

Η εξίσωση στο 13.2 δεν εκφράζει γραμμική σχέση, αλλά απεικονίζεται από μια σιγμοειδή καμπύλη (Εικόνα 13.2).



Εικόνα 13.3 Απεικόνιση της συνάρτησης $p = \frac{e^z}{1+e^z}$

13.1.3 Οι μερικοί συντελεστές και η ερμηνεία τους

Από την εξίσωση στο 13.1 προκύπτει ότι καθένας από τους συντελεστές $\beta_1, \beta_2, \dots, \beta_k$ εκφράζει τη μεταβολή του λογάριθμου της σχετικής πιθανότητας, για μία μονάδα αύξησης της αντίστοιχης ανεξάρτητης μεταβλητής X_i , εφόσον οι τιμές των υπολοίπων ανεξάρτητων μεταβλητών παραμένουν σταθερές.

Από τον ορισμό των *logit* για τα πηλικά σχετικών πιθανοτήτων OR και OR_1 για τις τιμές x_i και $x_i + 1$ της μεταβλητής X_i αντίστοιχα, όταν οι άλλες μεταβλητές παραμένουν σταθερές, ισχύει:

$$\begin{aligned} OR &= \frac{p}{1-p} = e^{a+\beta_1 x_1 + \beta_2 x_2 + \dots + \beta_i x_i + \dots + \beta_k x_k} = e^a \cdot e^{\beta_1 x_1} \cdot e^{\beta_2 x_2} \cdot \dots \cdot e^{\beta_i x_i} \cdot \dots \cdot e^{\beta_k x_k} \\ OR_1 &= \frac{p_1}{1-p_1} = e^{a+\beta_1 x_1 + \beta_2 x_2 + \dots + \beta_i (x_i + 1) + \dots + \beta_k x_k} = e^a \cdot e^{\beta_1 x_1} \cdot e^{\beta_2 x_2} \cdot \dots \cdot e^{\beta_i (x_i + 1)} \cdot \dots \cdot e^{\beta_k x_k} \\ OR_1 &= \frac{p_1}{1-p_1} = e^a \cdot e^{\beta_1 x_1} \cdot e^{\beta_2 x_2} \cdot \dots \cdot e^{\beta_i x_i} \cdot e^{\beta_i} \cdot \dots \cdot e^{\beta_k x_k} = e^{\beta_i} \cdot \left(\frac{p}{1-p} \right) \\ OR_1 &= e^{\beta_i} \cdot OR \end{aligned} \quad (13.3)$$

Άρα, η ποσότητα e^{β_i} , δηλαδή ο αντιλογάριθμος του β_i , είναι ο παράγοντας επί τον οποίο πολλαπλασιάζεται η σχετική πιθανότητα πραγματοποίησης του γεγονότος, όταν η ανεξάρτητη μεταβλητή X_i αυξηθεί κατά μία μονάδα και εφόσον οι υπόλοιπες μεταβλητές παραμένουν σταθερές (Γναρδέλλης, 2013).

Λαμβάνοντας υπόψη τις ιδιότητες της εκθετικής συνάρτησης με βάση το e , αν $\beta_i > 0$, τότε ο $e^{\beta_i} > 1$, δηλαδή η σχετική πιθανότητα αυξάνεται, όταν αυξηθεί κατά μία (ή περισσότερες) μονάδα η τιμή της X_i και οι υπόλοιπες μεταβλητές παραμένουν σταθερές. Αν $\beta_i < 0$, τότε ο $e^{\beta_i} < 1$, δηλαδή η σχετική πιθανότητα μειώνεται και αν $\beta_i = 0$, τότε ο $e^{\beta_i} = 1$, δηλαδή η σχετική πιθανότητα παραμένει ίδια.

Σε ένα λογαριθμικό υπόδειγμα πρόβλεψης της εμφάνισης στεφανιαίας νόσου, για παράδειγμα, αν $\beta_i = 0,0953$ στην εξίσωση για την ανεξάρτητη ποσοτική μεταβλητή της διαστολικής πίεσης, τότε σε μια αύξηση κατά μία μονάδα της τιμής της διαστολικής πίεσης έχουμε ($e^{\beta_i} = 1,10$) αύξηση κατά 10% της σχετικής πιθανότητας εμφάνισης στεφανιαίας νόσου.

13.1.4 Η εκτίμηση και η σημαντικότητα των μερικών συντελεστών και οι δείκτες καλής προσαρμογής του υποδείγματος

Στη λογαριθμική παλινδρόμηση, όπως και στην πολλαπλή, οι συντελεστές της εξίσωσης είναι οι εκτιμήσεις των πληθυσμιακών συντελεστών με τη βοήθεια του δείγματος που αναλύεται. Ενώ στην πολλαπλή παλινδρόμηση η εκτίμηση γίνεται με τη μέθοδο των ελαχίστων τετραγώνων, στη ΛΠ χρησιμοποιείται η *εκτίμηση μεγίστης πιθανοφάνειας*. *Πιθανοφάνεια* ονομάζεται η πιθανότητα να ληφθούν οι συγκεκριμένες παρατηρούμενες τιμές της εξαρτημένης μεταβλητής στο δείγμα για ένα συγκεκριμένο συνδυασμό τιμών των συντελεστών $a, \beta_1, \beta_2, \dots, \beta_k$.

Η πιθανοφάνεια εκφράζεται από τη συνάρτηση πιθανοφάνειας:

$$L(a, \beta_1, \beta_2, \dots, \beta_k) = \prod_{i=1}^N p_i^{y_i} (1 - p_i)^{1-y_i}, \quad (13.4)$$

όπου

y_i παίρνει τις τιμές 0 ή 1 της εξαρτημένης μεταβλητής Y , για τις παρατηρήσεις $i=1, 2, \dots, N$ και p_i η πιθανότητα επιτυχίας (να ληφθεί η τιμή 1) για την παρατήρηση i .

Οι $k+1$ τιμές $\hat{a}, \hat{\beta}_1, \hat{\beta}_2, \dots, \hat{\beta}_k$, που μεγιστοποιούν τη συνάρτηση πιθανοφάνειας L και, πιο συγκεκριμένα, τον λογάριθμο $\ln(L)$, αποτελούν τις εκτιμήσεις των πληθυσμιακών συντελεστών της ΛΠ. Η μεγιστοποίηση της συνάρτησης πιθανοφάνειας, για τον προσδιορισμό των συντελεστών του υποδείγματος γίνεται με τη χρήση του μηδενισμού των μερικών παραγώγων της $\ln(L)$ ως προς $a, \beta_1, \beta_2, \dots, \beta_k$ και για την επίλυση του συστήματος των $k+1$ εξισώσεων με $k+1$ αγνώστους χρησιμοποιούνται ειδικοί επαναληπτικοί αλγόριθμοι διαδοχικών

προσεγγίσεων (iterative) ένας εκ των οποίων χρησιμοποιείται από το SPSS. Οι αλγόριθμοι αυτοί εκτός των εκτιμήσεων των συντελεστών $\alpha, \beta_1, \beta_2, \dots, \beta_k$ παρέχουν και τα αντίστοιχα τυπικά σφάλματα.

Μετά την εκτίμηση των συντελεστών και τον προσδιορισμό της $\ln(L)$ ή LL μπορεί να μελετηθεί η στατιστική σημαντικότητα του υποδείγματος, δηλαδή να ελεγχθεί η μηδενική υπόθεση $H_0: \beta_1 = \beta_2 = \dots = \beta_k = 0$ κατά αναλογία με τον γενικό έλεγχο F της πολλαπλής παλινδρόμησης.

Για τον έλεγχο της υπόθεσης συγκρίνονται οι δύο πιθανοφάνειες λογάριθμου LL_1 και LL_0 που αντιστοιχούν στο υπόδειγμα με όλες τις ανεξάρτητες μεταβλητές (πλήρες) και σε αυτό που περιλαμβάνει μόνο τον σταθερό όρο α (μειωμένο υπόδειγμα). Η ποσότητα $\chi^2 = 2[LL_1 - LL_0]$ ακολουθεί τη χ^2 κατανομή με k βαθμούς ελευθερίας, όπου k είναι η διαφορά των παραμέτρων που εκτιμά κάθε υπόδειγμα: το πλήρες υπόδειγμα εκτιμά $k+1$ συντελεστές, ενώ αυτό με τον σταθερό παράγοντα έναν. Αν η ποσότητα $\chi^2 = 2[LL_1 - LL_0]$ ξεπερνά την κρίσιμη τιμή $\chi^2_{\alpha, k}$ της χ^2 που ορίζεται από το επίπεδο σημαντικότητας $\alpha=0,05$, τότε απορρίπτεται η παραπάνω μηδενική και υπάρχει τουλάχιστον ένας συντελεστής ανεξάρτητης μεταβλητής που διαφέρει από το 0. Στο SPSS δίνεται το παρατηρούμενο επίπεδο σημαντικότητας p για την τιμή χ^2 που υπολογίζεται από τα δεδομένα.

Ο έλεγχος *Wald* χρησιμοποιείται για τη σημαντικότητα του κάθε συντελεστή και γίνεται με τη βοήθεια του στατιστικού $W = \frac{\hat{\beta}_i^2}{\text{var}(\hat{\beta}_i)}$ το οποίο ακολουθεί ασυμπτωτικά τη χ^2 κατανομή με 1 βαθμό ελευθερίας. Το στατιστικό $\sqrt{W} = \frac{\hat{\beta}_i}{SE(\hat{\beta}_i)}$ ακολουθεί ασυμπτωτικά την κατανομή z . Ο έλεγχος *Wald* δίνεται στον πίνακα των συντελεστών της εξίσωσης (Peduzzi et al., 1996). Η επισκόπησή του οδηγεί στην ταυτοποίηση των ανεξάρτητων μεταβλητών στις οποίες οφείλεται η σημαντικότητα του γενικού ελέγχου της διαφοράς των LL , που είδαμε παραπάνω.

Σε αναλογία με την πολλαπλή παλινδρόμηση της ΛΠ δίνεται η τιμή R^2 των *Cox* και *Snell*, το οποίο υπολογίζεται από την παράσταση $R^2 = 1 - \left[\frac{L_0}{L_1}\right]^{2/N}$, όπου N το μέγεθος του δείγματος και αποτελεί έναν τρόπο έκφρασης του μεγέθους της επίδρασης του υποδείγματος. Ο συντελεστής του *Nagelkerke* ισούται με $\tilde{R}^2 = \frac{R^2}{R_{\max}^2}$, όπου $R_{\max}^2 = 1 - [L_0]^{2/N}$ και η τιμή του ανήκει στο διάστημα $(0,1)$ διορθώνοντας το πρόβλημα που υπάρχει στον συντελεστή R^2 των *Cox* και *Snell*. Συνεπώς, είναι προτιμότερο να αναφερόμαστε στην R^2 τιμή του *Nagelkerke*. Αυτοί οι δύο συντελεστές αναφέρονται ως *ψευδο- R^2* , επειδή ακριβώς οι τιμές τους δεν εκφράζουν το ποσοστό μεταβλητότητας της εξαρτημένης μεταβλητής που ερμηνεύεται από το υπόδειγμα, όπως συμβαίνει με την εκτίμηση ελαχίστων τετραγώνων στην πολλαπλή παλινδρόμηση. Κυριολεκτικά, αυτοί οι συντελεστές R^2 αντιπροσωπεύουν την αναλογική μείωση στην απόλυτη τιμή του μέτρου λογαριθμικής πιθανοφάνειας. Οι παραπάνω συντελεστές δίνουν τιμές χαμηλότερες της τιμής του συντελεστή προσδιορισμού R^2 της πολλαπλής παλινδρόμησης.

Ένας δεύτερος τρόπος αξιολόγησης της προβλεπτικής ικανότητας είναι η επισκόπηση του πίνακα ταξινόμησης (classification table). Στις γραμμές του βρίσκονται οι παρατηρούμενες τιμές της εξαρτημένης μεταβλητής Y , ενώ στις στήλες βρίσκονται οι προβλεπόμενες τιμές της ίδιας μεταβλητής. Αν η πιθανότητα να συμβεί το γεγονός είναι μεγαλύτερη από 0,5, τότε η προβλεπόμενη τιμή της Y είναι ίση με 1 διαφορετικά είναι ίση με 0.

Η δοκιμή καλής προσαρμογής *Hosmer-Lemeshow* είναι ένα από τα πιο κοινά εργαλεία που χρησιμοποιούνται στη λογαριθμική παλινδρόμηση. Μετά την εκτέλεση της λογαριθμικής παλινδρόμησης, ταξινομούνται οι παρατηρήσεις κατά αύξουσα σειρά των εκτιμώμενων πιθανοτήτων του να συμβεί το γεγονός. Οι παρατηρήσεις στη συνέχεια χωρίζονται σε περίπου δέκα ομάδες με βάση την εκτίμηση αυτής της πιθανότητας. Στη συνέχεια, πραγματοποιείται σύγκριση μεταξύ του αριθμού των παρατηρήσεων για τις οποίες πράγματι συμβαίνει το γεγονός (παρατηρούμενες) σε κάθε ομάδα και του αριθμού των παρατηρήσεων που προβλέπονται από το υπόδειγμα λογαριθμικής παλινδρόμησης (προβλεπόμενες).

Το στατιστικό καλής προσαρμογής *Hosmer-Lemeshow* υπολογίζεται ως στατιστικό χ^2 που προκύπτει από τον πίνακα $(2 \times g)$ των παρατηρούμενων και αναμενόμενων συχνοτήτων, όπου g είναι ο αριθμός των ομάδων. Το στατιστικό είναι:

$$\chi_{HL}^2 = \sum_{i=1}^g \frac{(O_i - N_i \hat{\pi}_i)^2}{N_i \hat{\pi}_i (1 - \hat{\pi}_i)} \sim \chi_{g-2}^2, \quad (13.5)$$

όπου

O_i = ο παρατηρούμενος αριθμός γεγονότων στην ομάδα i ,

N_i = ο αριθμός των παρατηρήσεων στην ομάδα i και

$\widehat{\pi}_i$ = η μέση εκτιμώμενη πιθανότητα γεγονότος στην ομάδα i .

Η τιμή που υπολογίζεται για κάποιο υπόδειγμα δεν πρέπει να είναι στατιστικά σημαντική για να γίνει αποδεκτή η καλή προσαρμογή του. Ένα γραφικός τρόπος παρουσίασης της προσαρμογής του υποδείγματος δίνεται από έναν τύπο διαγράμματος με τίτλο «observed groups and predicted probabilities» στο SPSS, όπου στον άξονα των x βρίσκονται τιμές στο διάστημα $[0 \text{ έως } 1]$ και στον άξονα των y δίνεται η συχνότητα των πιθανοτήτων. Κάθε περίπτωση απεικονίζεται στον άξονα x με την προβλεπόμενη από το υπόδειγμα πιθανότητα να συμβεί το γεγονός και συμβολίζεται με την παρατηρούμενη τιμή του στην Y . Συνεπώς, με το γράφημα αυτό μπορούν να εντοπιστούν οι περιπτώσεις για τις οποίες γίνεται κακή πρόβλεψη από το υπόδειγμα.

Στο SPSS αρχικά δίνεται η προσαρμογή του μειωμένου υποδείγματος, δηλαδή εκείνου που στην εξίσωση περιλαμβάνει μόνο τον σταθερό όρο. Για το παράδειγμα, στην πιθανότητα χρήσης των ΝΤ στη διδασκαλία, στον Πίνακα 13.3(β) φαίνεται ότι η τιμή $B=0,716$ είναι στατιστικά σημαντική ($W(1)=15,13$, $p<0,001$) και από την τιμή $\text{Exp}(B)=2,045$ συμπεραίνεται ότι η σχετική πιθανότητα να συμβεί το γεγονός είναι 2,045, καθώς και η πιθανότητα να συμβεί το γεγονός από την εξίσωση στο 13.2 είναι $p=2,04/(1+2,04)=0,67$. Ο πίνακας ταξινόμησης (α) (Εικόνα 13.4) δείχνει 100% ορθή ταξινόμηση για τις περιπτώσεις με παρατηρούμενη τιμή $Y=1$, αλλά 0% ορθές ταξινομήσεις για τις περιπτώσεις με παρατηρούμενη τιμή $Y=0$. Ο πίνακας (γ) (Εικόνα 13.4) (Variables not in the Equation) αναφέρεται στη σημαντικότητα καθεμίας από τις ανεξάρτητες μεταβλητές, αν είχε συμπεριληφθεί μόνη της στο μοντέλο μαζί με τον σταθερό όρο. Η τιμή του στατιστικού (score) σύμφωνα με τον οποίο πραγματοποιείται ο έλεγχος, αναφέρεται στη βαρύτητα που έχει κάθε ανεξάρτητη μεταβλητή στην πρόγνωση των τιμών της εξαρτημένης (Γναρδέλλης, 2013). Για μεγάλα δείγματα το στατιστικό αυτό ακολουθεί τη χ^2 κατανομή με βαθμούς ελευθερίας, τον αριθμό των συντελεστών τους οποίους ελέγχει. Η παρουσίαση ξεφεύγει από τους στόχους του παρόντος βιβλίου. Έτσι, λαμβάνοντας υπόψη τον έλεγχο αυτό, διαφαίνεται σημαντική εξήγηση της αξιοποίησης των υπολογιστών κατά τη διδασκαλία από τις μεταβλητές «ncomp_01», «certificA» και «confidence_with technology», με τη σειρά παρουσίασης να βασίζεται στα scores, όπως παρουσιάζονται στον Πίνακα (γ) της εικόνας 13.4.

Classification Table^{a,b}

Observed			Predicted		
			use01		Percentage Correct
			0 no	1 yes	
Step 0	use01	0 no	0	44	,0
		1 yes	0	90	100,0
	Overall Percentage				67,2

a. Constant is included in the model.

b. The cut value is ,500

(α)

Variables in the Equation

	B	S.E.	Wald	df	Sig.	Exp(B)
Step 0 Constant	,716	,184	15,134	1	,000	2,045

(β)

Variables not in the Equation

		Score	df	Sig.
Step 0	Variables			
	lack_of_support	,710	1	,399
	lack_of_confidence	,388	1	,533
	lack_of_equipment	2,489	1	,115
	class_conditions	2,231	1	,135
	condfince_with_technology	4,645	1	,031
	ncomp_01	36,929	1	,000
	yearsexp	1,169	3	,760
	yearsexp(1)	,393	1	,531
	yearsexp(2)	,614	1	,433
	yearsexp(3)	,054	1	,815
	certificA	10,165	1	,001
	yearsteach	5,723	4	,221
	yearsteach(1)	3,516	1	,061
	yearsteach(2)	,242	1	,623
	yearsteach(3)	,021	1	,886
	yearsteach(4)	2,917	1	,088
Overall Statistics		53,586	14	,000

(γ)

Εικόνα 13.4 Η προσαρμογή του μειωμένου υποδείγματος και η σημαντικότητα των μεταβλητών που είναι εκτός υποδείγματος.

Όλοι οι πίνακες που ακολουθούν, (α) έως (στ) (Εικόνα 13.5), αναφέρονται στην αξιολόγηση αλλά και στη μορφή του τελικού υποδείγματος. Η αξιολόγηση της προσαρμογής του υποδείγματος στα δεδομένα του δείγματος γίνεται με τον λόγο των μεγίστων τιμών της συνάρτησης πιθανοφάνειας (likelihood ratio statistic) για το πλήρες μοντέλο και το υπόδειγμα που περιλαμβάνει μόνο τον σταθερό όρο. Στον πίνακα (α) (Εικόνα 13.5) με τίτλο «Omnibus Tests of Model Coefficients» πραγματοποιείται ο έλεγχος του πλήρους υποδείγματος,

κατά πόσο δηλαδή οι μεταβλητές που θεωρήσαμε ως ανεξάρτητες εξηγούν σημαντικά καλύτερα τη μεταβλητότητα της εξαρτημένης μεταβλητής συγκριτικά με το μειωμένο υπόδειγμα. Η τιμή του στατιστικού $\chi^2 = 2[LL_1 - LL_0] = -2LL_0 - (-2LL_1)$ (τρίτη γραμμή του Πίνακα (α) της εικόνας 13.5) είναι 62,494=169,647-107,153 με 14 βαθμούς ελευθερίας και είναι στατιστικά σημαντική ($p < 0,001$).

Στον πίνακα (β) (Εικόνα 13.5) με τίτλο «Model Summary» δίνεται η τιμή της συνάρτησης λογαριθμο-πιθανοφάνειας (-2Log likelihood = 107,153) για το πλήρες υπόδειγμα της αντίστοιχης συνάρτησης λογαριθμο-πιθανοφάνειας, ενώ για το μειωμένο υπόδειγμα η αντίστοιχη τιμή είναι 169,647 και εμφανίζεται, αν επιλέξουμε στο πλαίσιο διαλόγου **Options** την επιλογή **Iteration history**.

Στον πίνακα (β) (Εικόνα 13.5) δίνονται επίσης οι συντελεστές προσδιορισμού των *Cox & Snell* (0,373=37,3%) και *Nagelkerke* (0,519=51,9%). Ένα καλό υπόδειγμα θα έχει τυπικά έναν συντελεστή *Nagelkerke* $R^2 > 0,2$ ενώ η τιμή $> 0,5$ υποδηλώνει ένα εξαιρετικό υπόδειγμα. Επίσης, στον Πίνακα (β) (Εικόνα 13.5) δίνεται ο έλεγχος καλής προσαρμογής *Hosmer and Lemeshow* από τον οποίο αναδεικνύεται η ικανοποιητική προσαρμογή του πλήρους υποδείγματος του παραδείγματός μας ($\chi^2_{HL}(8) = 12,49, p = 0,131$).

Ο Πίνακας των αποτελεσμάτων που παρουσιάζει ενδιαφέρον είναι ο Πίνακας ταξινόμησης του τελικού υποδείγματος (Εικόνα 13.5 και πίνακας (ε)). Όπως φαίνεται, το μοντέλο κατάφερε να εκτιμήσει τις τιμές της εξαρτημένης μεταβλητής (πρακτικά να τοποθετήσει τα υποκείμενα στις δύο τιμές (ναι ή όχι) στην αξιοποίηση των υπολογιστών κατά την εκπαιδευτική διαδικασία) που συμφωνούν με το 85,1% των πραγματικών τιμών της. Συγκεκριμένα ταίριαξε ορθά το 75% των *όχι* και 90% των *ναι*.

Ο Πίνακας (στ) (Εικόνα 13.5) που είναι από τους σημαντικότερους και έχει τίτλο «Variables in the Equation» παρουσιάζει τους συντελεστές του τελικού μοντέλου (B και e^B) μαζί με τους αντίστοιχους επαγωγικούς ελέγχους, καθώς και τα 95% διαστήματα εμπιστοσύνης των συντελεστών (e^B). Σύμφωνα με το κριτήριο *Wald* φαίνεται στο μοντέλο να υπάρχουν τρεις μόνο μεταβλητές που έχουν σημαντική επίδραση στην εξαρτημένη. Τη μεγαλύτερη επίδραση φαίνεται να την παρέχει η μεταβλητή «ncomp_01» (παρουσία ή μη υπολογιστών στην τάξη του εκπαιδευτικού), η οποία έχει έναν συντελεστή $e^B = 19,58$. Δηλαδή, η σχετική πιθανότητα των εκπαιδευτικών που έχουν υπολογιστές στην τάξη τους να χρησιμοποιήσουν υπολογιστές κατά την εκπαιδευτική διαδικασία είναι 19,58 φορές μεγαλύτερη από τη σχετική συχνότητα των συναδέλφων που δήλωσαν ότι δεν διαθέτουν υπολογιστή στην τάξη τους. Ακολουθεί η επίδραση της μεταβλητής «certificA» (αν έχουν ή όχι οι εκπαιδευτικοί πιστοποίηση στις βασικές υπολογιστικές δεξιότητες), η οποία έχει συντελεστή $e^B = 4,901$. Αυτό σημαίνει ότι η σχετική πιθανότητα των εκπαιδευτικών που έχουν σχετική πιστοποίηση να χρησιμοποιήσουν υπολογιστές στην εκπαιδευτική διαδικασία είναι 4,9 φορές μεγαλύτερη από τους άλλους που δήλωσαν ότι δεν έχουν τη σχετική πιστοποίηση. Τέλος, η μεταβλητή «confidence_with_technology» (αντιλαμβανόμενη εμπιστοσύνη των εκπαιδευτικών στους υπολογιστές) έχει συντελεστή $e^B = 3,492$. Αυτό σημαίνει ότι όσο μεγαλύτερη είναι η αυτοπεποίθηση του εκπαιδευτικού απέναντι στη χρήση των υπολογιστών, τόσο πιθανότερη είναι και η χρήση NT κατά τη διδασκαλία. Η αύξηση της τιμής της αυτοπεποίθησης κατά μία μονάδα της κλίμακας μέτρησής της, οδηγεί σε αύξηση της σχετικής πιθανότητας να χρησιμοποιηθούν οι NT κατά 249%.

Τέλος, το γράφημα «observed groups and predicted probabilities» (Εικόνα 13.6) αποτελεί μια απεικόνιση των παρατηρούμενων και των προβλεπόμενων τιμών της εξαρτημένης μεταβλητής. Με την απεικόνιση αυτή μπορούμε να εκτιμήσουμε σύντομα την προσαρμογή του μοντέλου στα παρατηρούμενα δεδομένα. Ο διαχωρισμός γίνεται λαμβάνοντας υπόψη την πιθανότητα 0,5. Σε τέλεια περίπτωση, κάτω από το 0,5 θα βρίσκαμε όλες τις περιπτώσεις με παρατηρούμενη τιμή «No ή 0», ενώ από 0,5 και πάνω όλες τις περιπτώσεις με παρατηρούμενη τιμή «Yes ή 1». Μια εικόνα καλής προσαρμογής, όπως συμβαίνει και στο παράδειγμά μας, παρουσιάζει τη συντριπτική πλειονότητα των τιμών στις αντίστοιχες πλευρές, όπως σημειώνονται στον οριζόντιο άξονα (group). Πράγματι, μόνο για 8 περιπτώσεις για τις οποίες παρατηρήθηκε η χρήση NT στη διδασκαλία δεν γίνεται η ορθή πρόβλεψη από το υπόδειγμα. Ανάλογο είναι το συμπέρασμά μας για την παρουσία περιπτώσεων στις οποίες δεν παρατηρήθηκε η χρήση NT στη διδασκαλία, στη δεξιά πλευρά του άξονα των πιθανοτήτων.

Omnibus Tests of Model Coefficients

		Chi-square	df	Sig.
Step 1	Step	62,494	14	,000
	Block	62,494	14	,000
	Model	62,494	14	,000

(α)

Model Summary

Step	-2 Log likelihood	Cox & Snell R Square	Nagelkerke R Square
1	107,153 ^a	,373	,519

a. Estimation terminated at iteration number 6 because parameter estimates changed by less than ,001.

(β)

Contingency Table for Hosmer and Lemeshow Test

		use01 = 0 no		use01 = 1 yes		Total
		Observed	Expected	Observed	Expected	
Step 1	1	12	11,605	1	1,395	13
	2	11	9,735	2	3,265	13
	3	9	8,023	4	4,977	13
	4	4	5,498	9	7,502	13
	5	2	3,721	11	9,279	13
	6	2	2,562	11	10,438	13
	7	0	1,358	13	11,642	13
	8	3	,767	10	12,233	13
	9	1	,445	12	12,555	13
	10	0	,285	17	16,715	17

Hosmer and Lemeshow Test

Step	Chi-square	df	Sig.
1	12,493	8	,131

(γ)

(δ)

Classification Table^a

Observed		Predicted			
		use01		Percentage Correct	
		0 no	1 yes		
Step 1	use01	0 no	33	11	75,0
		1 yes	9	81	90,0
Overall Percentage					85,1

a. The cut value is ,500

(ε)

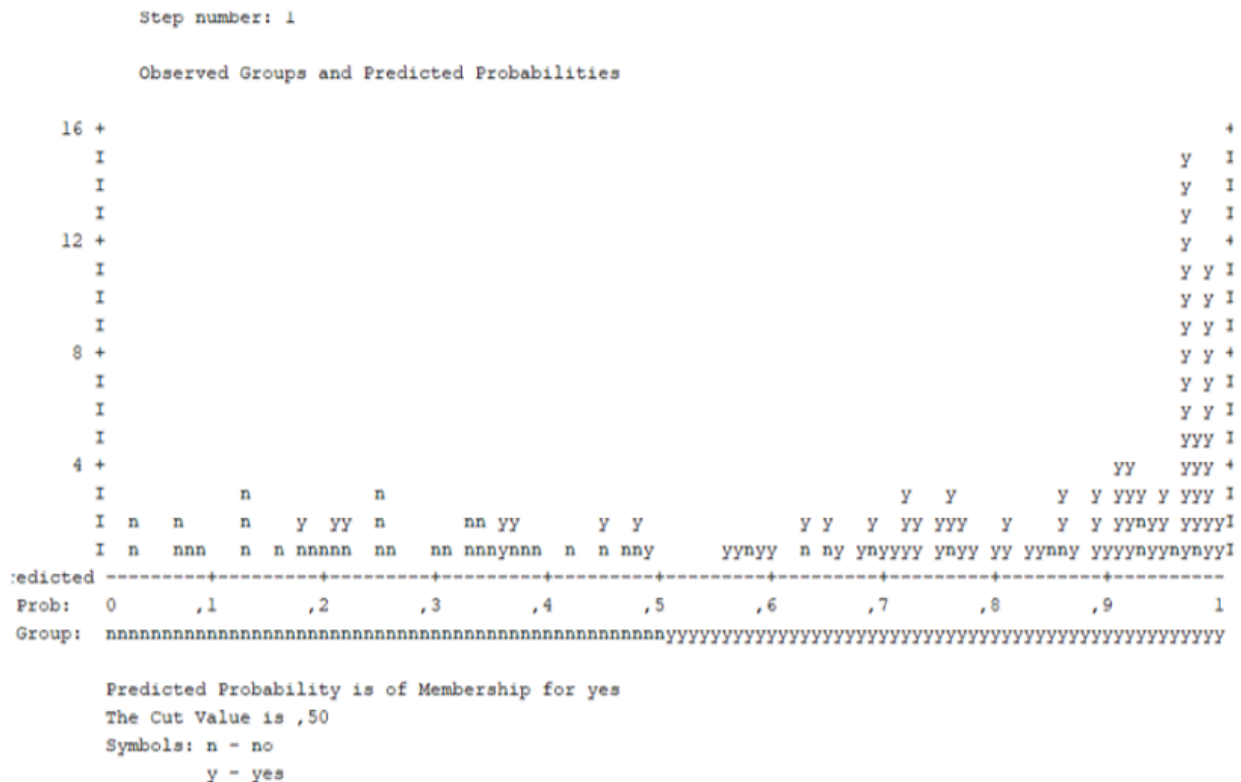
Variables in the Equation

		B	S.E.	Wald	df	Sig.	Exp(B)	95% C.I. for EXP(B)	
								Lower	Upper
Step 1 ^a	lack_of_support	,346	,652	,281	1	,596	1,413	,394	5,072
	lack_of_confidence	-,501	,451	1,234	1	,267	,606	,250	1,466
	lack_of_equipment	-,611	,447	1,868	1	,172	,543	,226	1,303
	class_conditions	,404	,477	,716	1	,397	1,498	,588	3,818
	condfince_with_technology	1,250	,553	5,109	1	,024	3,492	1,181	10,327
	ncomp_01	2,975	,579	26,351	1	,000	19,580	6,289	60,962
	yearsexp			1,280	3	,734			
	yearsexp(1)	,126	1,103	,013	1	,909	1,134	,131	9,851
	yearsexp(2)	-,522	,872	,359	1	,549	,593	,107	3,275
	yearsexp(3)	,450	,675	,443	1	,506	1,568	,417	5,892
	certificA	1,589	,589	7,292	1	,007	4,901	1,546	15,533
	yearsteach			2,402	4	,662			
	yearsteach(1)	-,198	,916	,047	1	,829	,821	,136	4,940
	yearsteach(2)	,295	,753	,154	1	,695	1,344	,307	5,876
	yearsteach(3)	,319	,919	,120	1	,729	1,375	,227	8,322
	yearsteach(4)	1,224	,890	1,891	1	,169	3,401	,594	19,467
	Constant	-4,026	2,452	2,695	1	,101	,018		

a. Variable(s) entered on step 1: lack_of_support, lack_of_confidence, lack_of_equipment, class_conditions, condfince_with_technology, ncomp_01, yearsexp, certificA, yearsteach.

(στ)

Εικόνα 13.5. Η προσαρμογή για το πλήρες υπόδειγμα.



Εικόνα 13.6 Γράφημα «observed groups and predicted probabilities».

13.1.5 Προϋποθέσεις της ΛΠ

Παρά το γεγονός ότι στη ΛΠ και στην εκτίμηση με τη μέθοδο μέγιστης πιθανοφάνειας δεν απαιτούνται οι προϋποθέσεις της πολυμεταβλητής κανονικότητας των ανεξάρτητων μεταβλητών, της κανονικότητας και της ομοσκεδαστικότητας των υπολοίπων που συνοδεύουν τις εκτιμήσεις με τη μέθοδο ελαχίστων τετραγώνων (OLS), υπάρχουν όμως ορισμένες προϋποθέσεις που πρέπει να ελέγχονται για να διασφαλιστεί η εγκυρότητα της:

- Αρχικά, μία από τις κύριες παραδοχές της ΛΠ είναι η κατάλληλη δομή της εξαρτημένης μεταβλητής. Η δυαδική λογαριθμική παλινδρόμηση απαιτεί η εξαρτημένη μεταβλητή να είναι διχοτομική. (Με 1 συμβολίζεται η παρουσία και με 0 η απουσία του γεγονότος που μελετάται).
- Η ΛΠ απαιτεί οι παρατηρήσεις να είναι ανεξάρτητες μεταξύ τους. Με άλλα λόγια, οι παρατηρήσεις δεν πρέπει να προέρχονται από επαναλαμβανόμενες μετρήσεις ή δείγματα στα οποία οι παρατηρήσεις να προκύπτουν κατά συστάδες.
- Η ΛΠ απαιτεί να υπάρχει μικρή ή καθόλου πολυσυγγραμμικότητα μεταξύ των ανεξάρτητων μεταβλητών. Αυτό σημαίνει ότι οι ανεξάρτητες μεταβλητές δεν πρέπει να συσχετίζονται ισχυρά μεταξύ τους. Η προϋπόθεση αυτή μπορεί να ελεγχθεί εκτελώντας μια πολλαπλή παλινδρόμηση με εξαρτημένη μεταβλητή τη διχοτομική εξαρτημένη της ΛΠ ή οποιαδήποτε άλλη και ανεξάρτητες μεταβλητές τις ίδιες μεταβλητές που βρίσκονται ως ανεξάρτητες στο υπόδειγμα της ΛΠ. Για τη διερεύνηση της προϋπόθεσης αυτής ο αναγνώστης μπορεί να ακολουθήσει ακριβώς τις οδηγίες που δόθηκαν στο Κεφάλαιο 12. Από το αποτέλεσμα αυτής της διερεύνησης αναφορικά με το παράδειγμα του Κεφαλαίου δεν προέκυψαν τιμές του VIF που να αντιστοιχούν σε υψηλό βαθμό πολυσυγγραμμικότητας στα δεδομένα.
- Θα πρέπει να ισχύει η γραμμικότητα στη σχέση μεταξύ κάθε ανεξάρτητης μεταβλητής και των σχετικών πιθανοτήτων των λογάρισμων. Αν και η ΛΠ δεν απαιτεί οι εξαρτημένες και οι ανεξάρτητες μεταβλητές να συσχετίζονται γραμμικά, απαιτείται οι ανεξάρτητες μεταβλητές να σχετίζονται γραμμικά με τις πιθανότητες καταγραφής. Είναι εύκολος ο έλεγχος με τη βοήθεια

ενός διαγράμματος σκεδασμού για κάθε ανεξάρτητη μεταβλητή στον άξονα των x. Οι σχετικές πιθανότητες λογάριθμων $\ln(OR)$, που θα βρίσκονται στον άξονα των y, μπορούν να προκύψουν εύκολα από τις εκτιμώμενες πιθανότητες που υπολογίζονται για κάθε περίπτωση από το SPSS στην επιλογή **Transform=>Compute Variable** (βλ. Κεφ.2) με τον τύπο $\ln(PRE_1 / (1 - PRE_1))$.

- Τέλος, η ΛΠ απαιτεί συνήθως μεγάλο μέγεθος δείγματος. Μια γενική οδηγία είναι ότι τουλάχιστον 10 περιπτώσεις για κάθε ανεξάρτητη μεταβλητή στο μοντέλο πρέπει να αντιστοιχούν στη συχνότητα της λιγότερο συχνής τιμής της εξαρτημένης μεταβλητής (Peduzzi et al., 1996). Για παράδειγμα, εάν έχουμε 5 ανεξάρτητες μεταβλητές και η αναμενόμενη πιθανότητα της λιγότερο συχνής τιμής της εξαρτημένης είναι 0,10, τότε θα χρειαστούμε ένα ελάχιστο μέγεθος δείγματος $N=500$ ($10 \times 5 / 0,10$). Για το παράδειγμα του κεφαλαίου, στο υπόδειγμα με 9 μεταβλητές και με 32,8% συχνότητα της μη χρήσης NT κατά τη διδασκαλία, οι απαιτήσεις του υποδείγματος σε δείγμα είναι $N=274$. Με δεδομένο ότι το υπόδειγμα στηρίχτηκε σε ένα δείγμα 134 εκπαιδευτικών, η παραπάνω διαπίστωση αποτελεί έναν περιορισμό αναφορικά με την εγκυρότητα των συμπερασμάτων μας.

Παρουσίαση του αποτελέσματος της λογαριθμικής παλινδρόμησης

Για την παρουσίαση του αποτελέσματος της λογαριθμικής παλινδρόμησης μπορούμε να γράψουμε συνοδευτικά τα εξής:

Εκτελέστηκε η Λογαριθμική Παλινδρόμηση προκειμένου να διερευνηθεί η επίδραση συγκεκριμένων παραγόντων (Έλλειμμα υποστήριξης, Έλλειμμα εμπιστοσύνης, Έλλειμμα τεχνολογικών πόρων, Συνθήκες τάξεις, Αντιλαμβανόμενη εμπιστοσύνη των εκπαιδευτικών στους υπολογιστές, αν έχουν ή όχι οι εκπαιδευτικοί στην τάξη τους υπολογιστές, την όποια εμπειρία των εκπαιδευτικών με τους υπολογιστές, αν έχουν ή όχι πιστοποίηση στις βασικές υπολογιστικές δεξιότητες, και την όποια διδακτική εμπειρία των εκπαιδευτικών) πάνω στην αξιοποίηση ή όχι των υπολογιστών από τους εκπαιδευτικούς κατά την εκπαιδευτική διαδικασία. Οι προϋποθέσεις εγκυρότητας της μεθόδου ικανοποιούνται με εξαίρεση το μέγεθος του δείγματος που είναι αρκετά μικρότερο από το απαιτούμενο. Το λογαριθμικό υπόδειγμα παλινδρόμησης βρέθηκε στατιστικά σημαντικό, $\chi^2(14) = 62,494, p < .001$. Ο έλεγχος *Hosmer and Lemeshow* έδειξε αποδεκτή προσαρμογή και η τιμή του συντελεστή προσδιορισμού R^2 -*Nagelkerke* = 51% δείχνει εξαιρετική προσαρμογή. Αναφορικά με τη διακριτική του ικανότητα βρέθηκε ότι το υπόδειγμα ταξινομεί με συνέπεια το 85,1% των περιπτώσεων. Από τα αποτελέσματα προέκυψε ότι η σχετική πιθανότητα των εκπαιδευτικών που έχουν υπολογιστές στην τάξη τους να χρησιμοποιήσουν υπολογιστές κατά την εκπαιδευτική διαδικασία είναι 19,58 φορές μεγαλύτερη από τους άλλους που δήλωσαν ότι δεν έχουν υπολογιστή στην τάξη τους. Επίσης, η σχετική πιθανότητα των εκπαιδευτικών που έχουν σχετική πιστοποίηση να χρησιμοποιήσουν υπολογιστές στην εκπαιδευτική διαδικασία είναι 4,9 φορές μεγαλύτερη από τους άλλους που δήλωσαν ότι δεν έχουν τη σχετική πιστοποίηση. Τέλος, οι εκπαιδευτικοί που έχουν μεγάλη εμπιστοσύνη στη χρήση των υπολογιστών είναι πιο πιθανό να χρησιμοποιήσουν υπολογιστές στην τάξη σε σχέση με τους άλλους που δήλωσαν ότι δεν έχουν μεγάλη εμπιστοσύνη.

13.2 Προβλήματα για εξάσκηση και ανατροφοδότηση

Άσκηση 1

Πολλές έρευνες υποστηρίζουν ότι η υπέρταση συνδέεται με αυξημένο κίνδυνο θανάτου από καρδιαγγειακή νόσο. Σε μια μεγάλη έρευνα σχετικά με αυτή τη σύνδεση εξετάστηκαν 3338 άνδρες με υψηλή πίεση και 2676 άνδρες με χαμηλή πίεση. Κατά την περίοδο της έρευνας 21 άνδρες από τη ομάδα χαμηλής πίεσης και 55 άνδρες από τη ομάδα υψηλής πίεσης έχασαν τη ζωή τους από καρδιαγγειακή νόσο.

- α. Να βρεθεί το ποσοστό των ανδρών που πέθαναν από καρδιαγγειακή νόσο στην ομάδα υψηλής πίεσης. Έπειτα να υπολογιστεί η σχετική πιθανότητα (odds).
- β. Να γίνει το ίδιο με την ομάδα χαμηλής αρτηριακής πίεσης.
- γ. Να υπολογιστεί το πηλίκιο των σχετικών πιθανοτήτων με τη σχετική πιθανότητα της ομάδας υψηλής πίεσης στον αριθμητή. Περιγράψτε με λέξεις το αποτέλεσμα.

Άσκηση 2

Στην εργασία (Nikolopoulou et al., 2021) χρησιμοποιήθηκε λογαριθμική παλινδρόμηση με στόχο να εντοπιστούν οι παράγοντες που εξηγούν την πρόθεση (ναι ή όχι) 920 εκπαιδευτικών να υιοθετήσουν την κινητή μάθηση. Στον Πίνακα που ακολουθεί παρουσιάζονται οι μεταβλητές που συμπεριελήφθησαν στο τελικό μοντέλο. Οι μεταβλητές αυτές είναι:

- A. Επίπεδο εκπαίδευσης «School level (SL)», ποιοτική μεταβλητή με κατηγορίες: 1. Δημοτικό «elementary school» η οποία έχει οριστεί ως κατηγορία αναφοράς, 2. Νηπιαγωγείο «early childhood», 3. Γυμνάσιο «high school», 4. Γενικό Λύκειο «general lyceum» και 5. Επαγγελματικό Λύκειο «vocational lyceum».
- B. Εκπαιδευτική εμπειρία «Years of teaching experience», ποσοτική μεταβλητή.
- Γ. Οφέλη από χρήση της κινητής μάθησης «F2: Benefits», ποσοτική μεταβλητή.
- Δ. Προτιμήσεις αναφορικά με τη χρήση της κινητής μάθησης «F3: Preferences», ποσοτική μεταβλητή.
- Ε. Εξωτερικές επιδράσεις στην εκπαιδευτική χρήση των κινητών συσκευών «F4: External Influences», ποσοτική μεταβλητή.

Λαμβάνοντας υπόψη τον πίνακα με τις μεταβλητές που συμπεριελήφθησαν στο τελικό μοντέλο, να παρουσιάσετε την ερμηνεία των σημαντικών παραγόντων στην εξήγηση της πρόθεσης των εκπαιδευτικών να υιοθετήσουν την κινητή μάθηση.

Table 7 Logistic regression model of mobile use in class (yes/no): Variables retained after forward selection method with Wald criterion

	B	SE	Wald	df	Sig.	Exp(B)	95% C.I. for EXP(B)	
							Lower	Upper
School level (SL) (elementary school=ref)			19.633	4	.001			
SL: early childhood	-.080	.245	0.106	1	.744	.923	.571	1.493
SL: high school	-.815	.200	16.611	1	.000	.443	.299	.655
SL: general lyceum	-.481	.211	5.197	1	.023	.618	.409	.935
SL: vocational lyceum	-.475	.234	4.103	1	.043	.622	.393	.985
Years of teaching experience	-.159	.050	10.102	1	.001	.853	.773	.941
F2: Benefits	.841	.171	24.290	1	.000	2.318	1.659	3.238
F3: Preferences	.393	.151	6.733	1	.009	1.482	1.101	1.994
F4: External Influences	.246	.094	6.908	1	.009	1.279	1.065	1.536
Constant	-3.779	.544	48.247	1	.000	.023		

Άσκηση 3

Από τη διερεύνηση των παραγόντων [agecat (ηλικιακή κατηγορία), sex(φύλο), educ (έτη εκπαίδευσης), log_inc (λογάριθμος του εισοδήματος σε δολ.)] που συνδέονται με τη χρήση του διαδικτύου δόθηκαν οι παρακάτω 5 Πίνακες αποτελεσμάτων.

Πίνακας 1. Categorical Variables Codings

		Frequency	Parameter coding			
			(1)	(2)	(3)	(4)
agecat Age category	18-29	193	,000	,000	,000	,000
	30-39	244	1,000	,000	,000	,000
	40-49	241	,000	1,000	,000	,000
	50-59	146	,000	,000	1,000	,000
	60-89	77	,000	,000	,000	1,000
sex Respondent's sex	Male	445	,000			
	Female	456	1,000			

Πίνακας 2. Variables in the Equation

	B	S.E.	Wald	df	Sig.	Exp(B)
Step 0 Constant	,338	,068	25,066	1	,000	1,403

Πίνακας 3. Omnibus Tests of Model Coefficients

	Chi-square	df	Sig.
Step 1 Step	203,822	7	,000
Block	203,822	7	,000
Model	203,822	7	,000

Πίνακας 4. Classification Table^a

Observed		Predicted		
		usenet Use Internet?		Percentage Correct
		No	Yes	
Step 1 usenet Use Internet?	No	225	150	60,0
	Yes	111	415	78,9
Overall Percentage				71,0

a. The cut value is 500

Πίνακας 5. Variables in the Equation

	B	S.E.	Wald	df	Sig.	Exp(B)
Step 1 ^a agecat			24,051	4	,000	
agecat(1)	-,016	,222	,005	1	,942	,984
agecat(2)	-,301	,224	1,808	1	,179	,740
agecat(3)	-,653	,255	6,575	1	,010	,520
agecat(4)	-1,346	,327	16,917	1	,000	,260
educ	,364	,036	103,202	1	,000	1,440
sex(1)	,361	,161	5,038	1	,025	1,434
log_inc	,520	,183	8,102	1	,004	1,681
Constant	-6,709	,829	65,468	1	,000	,001

a. Variable(s) entered on step 1: agecat, educ, sex, log_inc.

A. Από τα αποτελέσματα των Πινάκων του προηγούμενου θέματος γράψτε μια αναφορά συμπεριλαμβάνοντας:

- Στατιστική σημαντικότητα της προσαρμογής.
- Το αποτέλεσμα του πίνακα «Classification Table».
- Αναφορά στις μεταβλητές που ερμηνεύουν σημαντικά την πιθανότητα χρήσης διαδικτύου.
- Για τις κατηγορικές μεταβλητές να δοθεί το ποσοστό αύξησης ή μείωσης της σχετικής πιθανότητας χρήσης διαδικτύου κάθε κατηγορίας συγκριτικά με την κατηγορία αναφοράς.

B. Ποια είναι η εκτιμώμενη πιθανότητα χρήσης διαδικτύου άνδρα 25 ετών με 16 χρόνια εκπαίδευσης και μισθό 22.000 δολάρια ετησίως ($\ln(22000)=10$);

Άσκηση 4

Στο αρχείο «Postgraduate studies.sav» υπάρχουν πληροφορίες για το αποτέλεσμα της επιλογής 280 φοιτητών που ζήτησαν να συμμετάσχουν στις Μεταπτυχιακές Σπουδές (ΜΣ) του τμήματος ΑΑ.

Οι μεταβλητές στο αρχείο είναι:

- Αποτέλεσμα για τη συμμετοχή στις ΜΣ «Participate» (τιμές: Ναι, Όχι).
- Βαθμολογία στις εξετάσεις για την εισαγωγή στο πρόγραμμα ΜΣ «Exam_scores», ποσοτική μεταβλητή.
- Μέσος όρος βαθμολογίας από τις προπτυχιακές σπουδές «Grade_point_average», ποσοτική μεταβλητή.
- Βαθμολογία συνέντευξης «interview_level», ποιοτική μεταβλητή διάταξης με τέσσερις τιμές.

- A. Να εκτελεστεί η λογαριθμική παλινδρόμηση με εξαρτημένη μεταβλητή την «Participate» και ανεξάρτητες μεταβλητές «Exam_scores», «Grade_point_average» και «interview_level».
- B. Αφού γράψετε τα συμπεράσματα της ανάλυσης, υποδείξτε ποια μεταβλητή εξηγεί περισσότερο την εισαγωγή των αιτούντων φοιτητών στο πρόγραμμα ΜΣ.

Βιβλιογραφία

- Γναρδέλλης, Χ. (2013). *Ανάλυση δεδομένων με το IBM SPSS Statistics 21*. Αθήνα: Εκδ. Παπαζήση.
- Field, A. (2018). *Discovering statistics using IBM SPSS statistics* (5th ed.). London: SAGE.
- Nikolopoulou, K., & Gialamas, V. (2015). Barriers to the integration of computers in early childhood settings: Teachers' perceptions. *Education and Information Technologies*, 20, 285-301. <https://doi.org/10.1007/s10639-013-9281-9>
- Nikolopoulou, K., Gialamas, V., Lavidas, K., & Komis, V. (2021). Teachers' readiness to adopt mobile learning in classrooms: A study in Greece. *Technology, Knowledge and Learning*, 26, 53-77. <https://doi.org/10.1007/s10758-020-09453-7>
- Peduzzi, P., Concato, J., Kemper, E., Holford, T. R., & Feinstein, A. R. (1996). A simulation study of the number of events per variable in logistic regression analysis. *Journal of Clinical Epidemiology*, 49(12), 1373-1379. [https://doi.org/10.1016/S0895-4356\(96\)00236-3](https://doi.org/10.1016/S0895-4356(96)00236-3)

Κεφάλαιο 14

Διερευνητική Παραγοντική Ανάλυση

Σύνοψη

Στο κεφάλαιο αυτό γίνεται μια εισαγωγή στη διερευνητική παραγοντική ανάλυση (Factor analysis), καθώς επίσης παρουσιάζονται και τα βήματα που απαιτούνται για την ανάλυση αυτή στο περιβάλλον του SPSS. Τέλος, παρουσιάζονται ζητήματα εγκυρότητας και αξιοπιστίας που θα πρέπει να λαμβάνονται υπόψη κατά τη διερεύνηση της παραγοντικής δομής μιας κλίμακας.

Προαπαιτούμενη γνώση

Συσχέτιση και Απλή Γραμμική και Πολλαπλή Γραμμική Παλινδρόμηση: Κεφάλαια 11 και 12 του συγγράμματος.

14.1 Εισαγωγή στην Παραγοντική Ανάλυση

Η συσχέτιση και ο συντελεστής συσχέτισης (Pearson r) χρησιμοποιείται για την περιγραφή της σχέσης ανάμεσα σε δύο ποσοτικές μεταβλητές. Πολλές φορές οι μεταβλητές προέρχονται από παρατήρηση και δεν είναι αποτέλεσμα κάποιου χειρισμού εκ μέρους του ερευνητή. Για παράδειγμα, αν κάποιος ερευνητής ενδιαφέρεται για τη σχέση ανάμεσα στο άγχος του μαθητή και την επίδοσή του, θα μπορούσε να καταγράψει τους βαθμούς που έδωσε ο δάσκαλος μιας ομάδας μαθητών και ταυτόχρονα να μετρήσει το επίπεδο του άγχους τους με τη βοήθεια ενός ερωτηματολογίου ή κλίμακας. Εδώ, ο ερευνητής δεν χειρίζεται ούτε την επίδοση ούτε το άγχος των μαθητών, αλλά παρατηρεί ή μετράει αυτό που συμβαίνει εκ του φυσικού. Όπως είναι κατανοητό, στην περίπτωση του άγχους η έννοια είναι αφηρημένη και δεν μπορεί να υπάρχει μία μόνο ερώτηση με την οποία να γίνεται η μέτρηση ώστε να προκύπτει μια τιμή που αντιστοιχεί στο επίπεδο του αντιλαμβανόμενου άγχους, όπως για παράδειγμα συμβαίνει με τη μέτρηση του βάρους. Το άγχος είναι μια *υποκείμενη μεταβλητή* (latent variable) (μη παρατηρούμενη), και για να μετρηθεί χρειάζεται να καταγραφούν οι απαντήσεις των ερωτώμενων σε διάφορες δηλώσεις. Χρειάζεται, δηλαδή, να πραγματοποιηθεί μέτρηση σε διάφορες πτυχές που προκαλούνται από το άγχος. Επομένως, το άγχος αναμένεται να επιδρά (συσχέτιση) στις απαντήσεις-*παρατηρούμενες μεταβλητές* (manifest variables) που επιλέγουν οι μαθητές, στις δηλώσεις του ερωτηματολογίου που αντιστοιχούν σε αυτό.

Η παραγοντική ανάλυση είναι μια συλλογή μεθόδων που χρησιμοποιούνται για να εξεταστεί πώς οι υποκείμενες κατασκευές (latent variables ή constructs) επηρεάζουν τις απαντήσεις σε έναν αριθμό μετρούμενων μεταβλητών (manifest variables).

Μεταξύ των αποτελεσμάτων της παραγοντικής ανάλυσης βρίσκεται και η εκτίμηση του επιπέδου του άγχους με τη μορφή μίας (ή μερικών) ποσοτικής μεταβλητής η οποία μπορεί να χρησιμοποιηθεί για τη διερεύνηση της προαναφερθείσας σχέσης ανάμεσα στο άγχος του μαθητή και στην επίδοσή του με τον υπολογισμό του συντελεστή συσχέτισης r . Κυρίως υπάρχουν δύο είδη παραγοντικής ανάλυσης: η Διερευνητική και η Επιβεβαιωτική.

Η Διερευνητική Παραγοντική Ανάλυση (ΔΠΑ) (Exploratory Factor Analysis, EFA) προσπαθεί να ανακαλύψει τη φύση των κατασκευών που επηρεάζουν ένα σύνολο απαντήσεων. Για παράδειγμα, με τη βοήθεια της EFA οδηγούμαστε στο συμπέρασμα ότι οι έξι δηλώσεις που αναφέρονται στη νοημοσύνη ανήκουν σε μία δομή με δύο υποκείμενες κατασκευές. Επίσης, συμπεραίνουμε ότι οι δηλώσεις (items) q1, q4 και q6 αντανakλούν κυρίως τη «Χωρική νοημοσύνη», ενώ οι άλλες τρεις δηλώσεις q2, q3 και q5 συνδέονται κυρίως με τη «Λεκτική νοημοσύνη».

Η Επιβεβαιωτική Παραγοντική Ανάλυση (ΕΠΑ) (Confirmatory Factor Analysis, CFA) ελέγχει εάν ένα καθορισμένο σύνολο θεωρητικών κατασκευών επηρεάζει τις απαντήσεις σύμφωνα με έναν προβλεπόμενο τρόπο. Εδώ, δηλαδή, μία προϋπάρχουσα δομή επιβεβαιώνεται ή όχι. Για παράδειγμα, με τη βοήθεια της CFA βρίσκουμε ότι το υπόδειγμα που θέλει, τη δομή του ερωτηματολογίου να έχει δύο πτυχές (παράγοντες ή διαστάσεις), με τις q1, q4 και q6 να συνδέονται αποκλειστικά με την «Χωρική νοημοσύνη», ενώ τις δηλώσεις q2, q3 και q5 να συνδέονται αποκλειστικά με τη «Λεκτική νοημοσύνη», έχει καλή προσαρμογή με τα δεδομένα.

Και οι δύο τύποι Παραγοντικής Ανάλυσης βασίζονται στο Υπόδειγμα του Κοινού Παράγοντα, που απεικονίζεται στην Εικόνα 14.1. Αυτό το υπόδειγμα προτείνει ότι κάθε παρατηρούμενη μεταβλητή («Y1» έως «Y6») επηρεάζεται μερικώς από υποκείμενους κοινούς παράγοντες (F1 και F2) και εν μέρει από υποκείμενους μοναδικούς παράγοντες (u1 έως u6). Η δύναμη της σχέσης μεταξύ κάθε παράγοντα και κάθε μεταβλητής ποικίλλει, έτσι ώστε ένας συγκεκριμένος παράγοντας επηρεάζει ορισμένες μεταβλητές περισσότερο από άλλες. Στο παράδειγμα που φαίνεται στην Εικόνα 14.1, οι «Y1», «Y2» και «Y3» δέχονται τη σημαντική επίδραση του F1 (συνεχής γραμμή διαδρομών) σε αντίθεση με τις «Y4», «Y5» και «Y6», στις οποίες ο F1 ασκεί ασήμαντη επίδραση (διακεκομμένη γραμμή διαδρομών). Το υπόδειγμα αυτό χρησιμοποιείται από τα λογισμικά, όπως το SPSS ή το AMOS για (EFA) και (CFA) αντίστοιχα.

Το σύστημα εξισώσεων (14.1) γραμμικής παλινδρόμησης που επιλύεται για το παράδειγμα της Εικόνας 14.1 είναι το παρακάτω:

$$Y_j = b_{j1} \cdot F_1 + b_{j2} \cdot F_2 + U_j \quad (14.1)$$

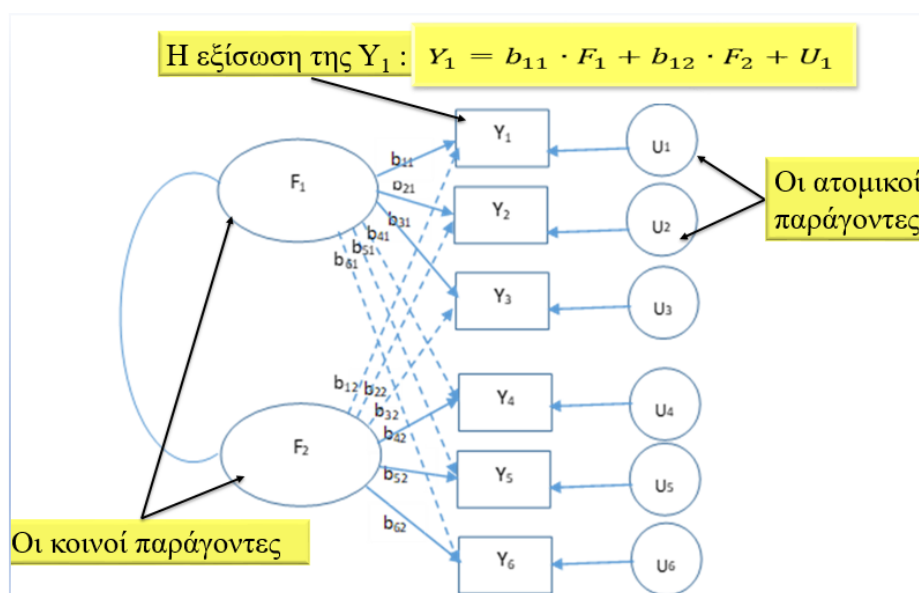
όπου

Y_j οι τυποποιημένες παρατηρούμενες μεταβλητές

F_i οι κοινόι παράγοντες

U_j ο μοναδικός ή ατομικός παράγοντας της Y_j

b_{ji} η φόρτιση (συσχέτιση) του παράγοντα της Y_j πάνω στον παράγοντα F_i



Εικόνα 14.1 Υπόδειγμα διερευνητικής παραγοντικής ανάλυσης με 6 παρατηρούμενες μεταβλητές και 2 κοινούς παράγοντες.

Οι παραγοντικές αναλύσεις πραγματοποιούνται εξετάζοντας το πρότυπο συσχετίσεων (ή συνδιακυμάνσεων) μεταξύ των παρατηρούμενων μεταβλητών (Πίνακας 14.1). Παρατηρούμενες μεταβλητές που έχουν μεγάλη συσχέτιση (είτε θετική είτε αρνητική) επηρεάζονται από τους ίδιους παράγοντες, ενώ εκείνοι που δεν συνδέονται ή συνδέονται ασθενώς μεταξύ τους επηρεάζονται από διαφορετικούς παράγοντες. Στον Πίνακα 14.1 οι παρατηρούμενες μεταβλητές «Y1» έως «Y3» έχουν ισχυρότερες σχέσεις μεταξύ τους, επειδή επηρεάζονται από τον παράγοντα F1. Αντίθετα, οι παραπάνω μεταβλητές έχουν ασθενέστερες σχέσεις με τις «Y4» έως «Y6», οι οποίες επηρεάζονται από τον παράγοντα F2. Παρόμοια, οι παρατηρούμενες μεταβλητές «Y4» έως «Y6» έχουν ισχυρότερες σχέσεις μεταξύ τους, επειδή επηρεάζονται από τον παράγοντα F2.

	Y1	Y2	Y3	Y4	Y5	Y6
Y1	1					
Y2	,326**	1				
Y3	,449**	,417**	1			
Y4	,342**	,228**	,328**	1		
Y5	,309**	,159*	,287**	,719**	1	
Y6	,317**	,195*	,347**	,714**	,685**	1

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

* . Correlation is significant at the 0.05 level (2-tailed).

Πίνακας 14.1 Συσχετίσεις μεταξύ των 6 παρατηρούμενων μεταβλητών της παραγοντικής ανάλυσης της Εικόνας 14.1.

14.1.1 Οι στόχοι της Διερευνητικής Παραγοντικής Ανάλυσης

Οι πρωταρχικοί στόχοι μιας ΔΠΑ είναι να προσδιοριστεί:

1. Ο αριθμός των κοινών παραγόντων που επηρεάζουν ένα σύνολο παρατηρούμενων μεταβλητών. Πρακτικά, πώς ομαδοποιούνται οι παρατηρούμενες μεταβλητές.
2. Η ένταση της σχέσης μεταξύ κάθε παράγοντα και κάθε παρατηρούμενης μεταβλητής.

Κάποιες συνηθισμένες χρήσεις της ΔΠΑ είναι:

1. Προσδιορισμός της φύσης των κατασκευών που βρίσκονται κάτω από τις απαντήσεις των συμμετεχόντων σε μια συγκεκριμένη περιοχή περιεχομένου.
2. Προσδιορισμός του τρόπου ομαδοποίησης των δηλώσεων που χρησιμοποιούνται σε μια κλίμακα (σύνολο δηλώσεων του ερωτηματολογίου που μετρούν την ίδια έννοια).
3. Ανάδειξη των διαστάσεων μιας κλίμακας μέτρησης. Οι ερευνητές συχνά επιθυμούν να αναπτύξουν κλίμακες που αντιστοιχούν σε ένα μόνο χαρακτηριστικό.
4. Υπολογισμός των παραγοντικών τιμών που αντιπροσωπεύουν τιμές των υποκείμενων κατασκευών (εννοιών) για χρήση σε άλλες αναλύσεις.

14.1.2 Η εκτέλεση της ΔΠΑ

Στην παρουσίαση της παραγοντικής ανάλυσης θα χρησιμοποιηθούν τα δεδομένα από τη διερεύνηση των στάσεων μαθητών απέναντι στη χρήση υπολογιστικού περιβάλλοντος κατά τη διδασκαλία των Μαθηματικών (Barkatsas et al., 2009). Το αρχείο των δεδομένων που θα χρησιμοποιήσουμε είναι το «Barkatsas et al.sav». Συγκεκριμένα, θα διερευνηθεί η παραγοντική δομή των στάσεων απέναντι στα Μαθηματικά με τη βοήθεια 12 δηλώσεων (q1-q4, q9-q16), στις οποίες απάντησαν 343 μαθητές Γ΄ Γυμνασίου και Α΄ Λυκείου με μία 5-βαθμών κλίμακα (1= σχεδόν ποτέ έως και 5 = σχεδόν πάντα).

Υπάρχουν ορισμένα βασικά βήματα για την εκτέλεση μιας ΔΠΑ:

1. Έλεγχος προϋποθέσεων εγκυρότητας:
 - Το πηλίκο [μέγεθος δείγματος (N) / αριθμός μεταβλητών (p)] πρέπει να είναι μεγαλύτερο του 5 και στην περίπτωση που οι μεταβλητές αποκλίνουν σοβαρά από την κανονική κατανομή πρέπει να ξεπερνά το 10 (Watkins, 2018). Στα δεδομένα του παραδείγματος η προϋπόθεση αυτή ικανοποιείται με N=343 και p=12.
 - Πριν εκτελεστούν τα παραπάνω βήματα πρέπει να ελεγχθούν οι μεταβλητές, μία προς μία, για πιθανή παρουσία ακραίων τιμών και σοβαρών αποκλίσεων από την κανονικότητα. Αν παρουσιαστούν προβλήματα μπορεί να χρειαστεί μετασχηματισμός μιας ή περισσότερων μεταβλητών (π.χ. λογάριθμος) ή/και η εξαίρεση ορισμένων ακραίων τιμών. Στο παράδειγμα των «στάσεων απέναντι στα Μαθηματικά» επιλέγουμε [Analyze=> Descriptive Statistics=> Descriptive](#) και, μαρκάροντας τις 12 παρατηρούμενες μεταβλητές (βλ. Κεφ. 4) που πρόκειται να χρησιμοποιηθούν στην παραγοντική ανάλυση, δημιουργείται ο Πίνακας (Εικόνα 14.2)

βασικών περιγραφικών στατιστικών. Τα πηλικά $|Skewness|/SE$ και $|Kurtosis|/SE$ παίρνουν τιμές μεγαλύτερες του 3 σε ορισμένες μεταβλητές, όπως οι «q16», «q10», «q2», «q1» για τη λοξότητα και «q2», «q3», «q14», «q15» για την κύρτωση. Ανάλογες τιμές αποτελούν πολύ συχνό φαινόμενο στη μελέτη των στάσεων στις οποίες χρησιμοποιούνται κλίμακες τύπου Likert.

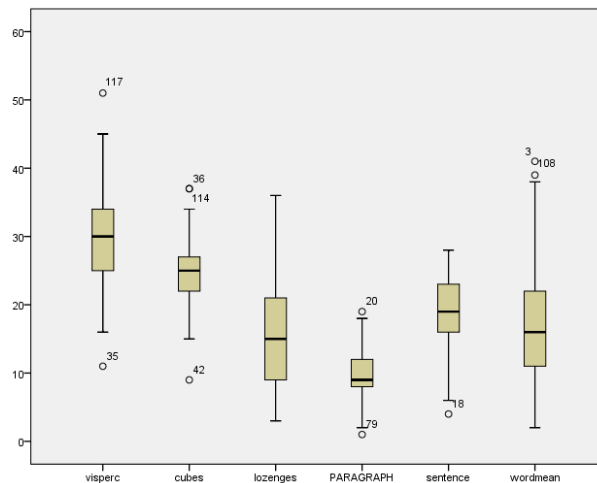
Επειδή η υπόθεση της πολυμεταβλητής και μονομεταβλητής κανονικότητας δεν έχει νόημα στην περίπτωση των διατάξιμων μεταβλητών, αλλά μόνο για εκείνη των συνεχών μεταβλητών, στη σύγχρονη βιβλιογραφία προτείνεται να γίνει η χρήση του Πίνακα πολυχωρικών (polychoric) συσχετίσεων αντί του Πίνακα συσχετίσεων *Pearson r*. Πολυχωρική ονομάζεται η συσχέτιση ανάμεσα σε δύο συνεχείς υποκείμενες μεταβλητές, η καθεμία από τις οποίες ακολουθεί την κανονική κατανομή και συνδέεται με μια τακτική παρατηρούμενη μεταβλητή. Στην επιβεβαιωτική αλλά και διερευνητική παραγοντική ανάλυση οι προτεινόμενες μέθοδοι εξαγωγής των παραγόντων για τακτικές μεταβλητές ή συνεχείς μεταβλητές με ισχυρή απόκλιση από την κανονικότητα είναι ο WLSMV (weighted least square mean and variance adjusted) και ULS (Unweighted least squares), σε αντίθεση με την περίπτωση κανονικών συνεχών μεταβλητών που κλασικά χρησιμοποιείται η μέθοδος της ML (Maximum likelihood) (Suh, 2015). Από τον Πίνακα 14.2 και το μέγεθος των τιμών της λοξότητας και της κύρτωσης δεν προκύπτει ισχυρή απόκλιση από την κανονική κατανομή ($|Skewness| < 2$ και $|kurtosis| < 3$) για καμία από τις δηλώσεις (Kim, 2013).

Descriptive Statistics							
	N	Mean	Std. Deviation	Skewness		Kurtosis	
	Statistic	Statistic	Statistic	Statistic	Std. Error	Statistic	Std. Error
q1	343	3,57	1,257	-,527	,132	-,908	,263
q2	343	3,64	1,330	-,574	,132	-,1011	,263
q3	342	3,22	1,309	-,238	,132	-,1159	,263
q4	343	3,08	1,369	-,129	,132	-,1292	,263
q9	340	3,13	1,129	-,195	,132	-,679	,264
q10	342	3,71	1,065	-,684	,132	-,468	,263
q11	342	3,35	,992	-,398	,132	-,107	,263
q12	338	3,13	1,161	-,210	,133	-,728	,265
q13	340	3,41	1,234	-,488	,132	-,753	,264
q14	340	3,45	1,224	-,553	,132	-,657	,264
q15	341	3,07	1,322	-,199	,132	-,1065	,263
q16	342	3,71	1,289	-,738	,132	-,564	,263
Valid N (listwise)	332						

Στις στήλες «Skewness/statistic» και «Kurtosis/Statistic» δίνονται οι τιμές της λοξότητας και της κύρτωσης και στη στήλη «Std.Error» το τυπικό σφάλμα τους.

Εικόνα 14.2 Περιγραφικά στατιστικά των κατανομών των 12 δηλώσεων.

Εναλλακτικά, η απόκλιση από την κανονικότητα σε συνεχείς μεταβλητές μπορεί να ελεγχθεί και γραφικά ή χρησιμοποιώντας αντίστοιχα τεστ (βλ. Κεφ. 6). Στο παράδειγμά μας, η επισκόπηση του θηκόγραμματος (Εικόνα 14.3) των τιμών των 6 παρατηρούμενων μεταβλητών που χρησιμοποιούνται για τον Πίνακα συσχετίσεων (Πίνακας 14.1) δεν αναδεικνύει σοβαρή απόκλιση από την κανονική κατανομή. Η οριζόντια γραμμή (αντιστοιχεί στη διάμεσο των παρατηρήσεων) στα θηκογράμματα βρίσκεται σε όλες τις περιπτώσεις, σχεδόν στη μέση του ορθογωνίου.

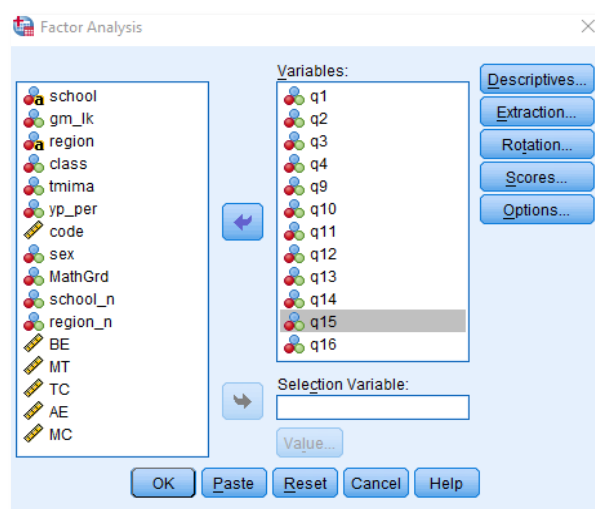


Εικόνα 14.3 Θηκόγραμμα των 6 μεταβλητών του παραδείγματος.

2. Έλεγχος για ελλείπουσες τιμές:

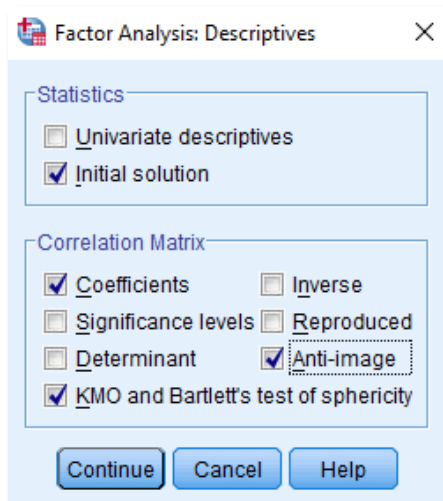
- Η εξαίρεση των περιπτώσεων με τουλάχιστον μία ελλείπουσα τιμή (listwise exclusion) είναι μια πιθανή επιλογή αρκεί να μην οδηγεί σε δραματική μείωση του δείγματος με το οποίο θα εκτελεστεί η ΔΠΑ. Αν δεν μπορεί να χρησιμοποιηθεί η παραπάνω προσέγγιση, ο υπολογισμός των συσχετίσεων με όλα τα διαθέσιμα δεδομένα (pairwise exclusion) είναι η επόμενη επιλογή. Μια σύγχρονη προσέγγιση είναι η εκτίμηση των ελλειπουσών τιμών με τη βοήθεια των συσχετίσεων όλων των μεταβλητών που εμπλέκονται στην Παραγοντική Ανάλυση. Παράδειγμα τέτοιας προσέγγισης είναι η υποκατάσταση των ελλειπουσών τιμών με τη βοήθεια του αλγόριθμου EM που υπάρχει στο πακέτο «Missing Values Analysis» του SPSS. Η παρουσίαση της τελευταίας ξεφεύγει από τους στόχους του παρόντος. Στα δεδομένα του Πίνακα (Εικόνα 14.2) παρατηρούμε ότι οι μαθητές με απαντήσεις σε όλες τις δηλώσεις του ερωτηματολογίου είναι 332 (valid N(listwise)) και η απώλεια $343-332=11$ περιπτώσεων κρίνεται αμελητέα και συνεπώς θα εξαιρεθούν οι 11 αυτές περιπτώσεις από την ανάλυση. Για την εξαίρεση αυτή επιλέγεται η επιλογή **listwise** (Εικόνα 14.6).

Για την εκτέλεση της παραγοντικής ανάλυσης από το μενού επιλέγουμε διαδοχικά **Analyze=> Dimension Reduction=> Factor** και εμφανίζεται το πλαίσιο διαλόγου **Factor analysis** (Εικόνα 14.4), όπου τοποθετούμε στο πεδίο **Variables** όλες τις μεταβλητές που θέλουμε να συμπεριλάβουμε στην παραγοντική ανάλυση.



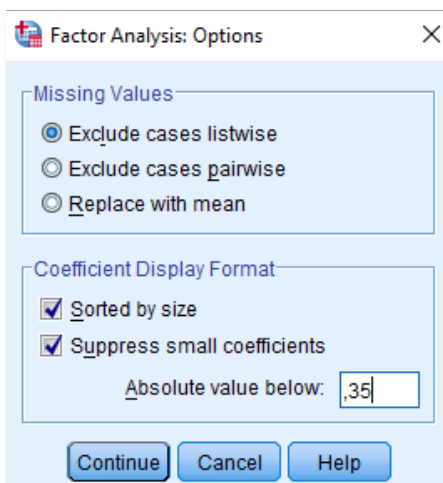
Εικόνα 14.4 Πλαίσιο διαλόγου «Factor Analysis».

Για τη δημιουργία του Πίνακα συσχετίσεων (pearson r) των παρατηρούμενων μεταβλητών και την καταλληλότητα τους για την παραγοντική ανάλυση επιλέγουμε την επιλογή **Descriptives** και στο αντίστοιχο πλαίσιο διαλόγου (Εικόνα 14.5) δηλώνουμε τη δημιουργία του Πίνακα συσχετίσεων (επιλογή **Coefficients**), τον υπολογισμό του Kaiser–Meyer–Olkin (KMO) και τον έλεγχο σφαιρικότητας του Πίνακα συσχετίσεων (επιλογή **KMO and Bartlett's test of sphericity**), καθώς επίσης και την επιλογή **Anti-image**.



Εικόνα 14.5 Πλαίσιο διαλόγου «Factor Analysis: Descriptives».

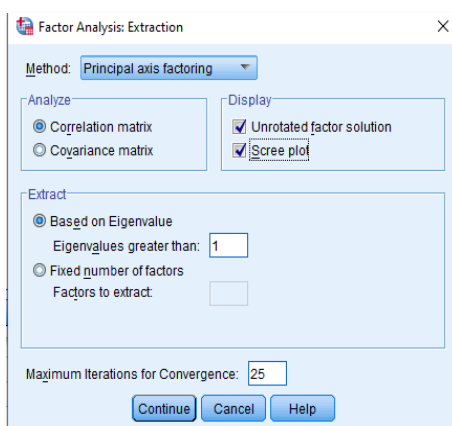
Για τον τρόπο επιλογής των περιπτώσεων, λαμβάνοντας υπόψη τις ελλείπουσες τιμές, επιλέγουμε την επιλογή **Options** και στο πλαίσιο διαλόγου **Factor Analysis Options** (Εικόνα 14.6) επιλέγουμε την επιλογή **Exclude cases listwise**. Επιπλέον, σε αυτό το πλαίσιο διαλόγου μπορούμε να επιλέξουμε την ταξινομημένη παρουσίαση των παρατηρούμενων μεταβλητών σε φθίνουσα σειρά ως προς τις φορτίσεις τους με τους παράγοντες (επιλογή **Sorted by size**). Τέλος, μπορούμε να ζητήσουμε να μην παρουσιαστούν φορτίσεις (επιλογή **Absolute value below**) για τις παρατηρούμενες μεταβλητές μέχρι κάποια τιμή.



Εικόνα 14.6 Πλαίσιο διαλόγου «Factor Analysis Options».

- Για τον προσδιορισμό του αριθμού των παραγόντων που θα συμπεριληφθούν στην τελική παραγοντική δομή αλλά και στην επιλογή της μεθόδου εκτίμησης ή εξαγωγής των παραγόντων για την παραγοντική ανάλυση επιλέγουμε την επιλογή **Extraction**. Στο πλαίσιο διαλόγου **Factor Analysis Extraction** που εμφανίζεται (Εικόνα 14.7) δηλώνουμε ως μέθοδο εκτίμησης των παραγόντων (επιλογή **Method**) την **Principal axis factoring**. Στο ίδιο πλαίσιο διαλόγου στο πεδίο **Fixed number of factor, Factors to extract** ορίζουμε και τον αριθμό των παραγόντων που θα ερμηνευτούν και θα πάρουν μέρος στην εξίσωση 14.1, που είδαμε παραπάνω. Τέλος, για

τον γραφικό προσδιορισμό του αριθμού των παραγόντων που τελικά θα κρατήσουμε επιλέγουμε **Scree plot**.



Εικόνα 14.7 Πλαίσιο διαλόγου «Factor Analysis Extraction».

3. Καταλληλότητα του πίνακα συσχετίσεων των παρατηρούμενων μεταβλητών:

- Τιμές του KMO που είναι μεγαλύτερες του 0,6 δηλώνουν καταλληλότητα των συσχετίσεων του δείγματος για την Παραγοντική Ανάλυση (Watkins, 2018). Η μηδενική υπόθεση στον έλεγχο *Bartlett* είναι ότι όλες οι 66 συσχετίσεις μεταξύ των 12 μεταβλητών ισούνται με μηδέν. Στην περίπτωσή μας απορρίπτεται, ($\chi^2(66)=1841,7$; $p<0,05$). Η πολύ υψηλή τιμή KMO και η σημαντικότητα του ελέγχου *Bartlett* επιβεβαιώνουν την καταλληλότητα της παραγοντικής ανάλυσης για το συγκεκριμένο δείγμα (Εικόνα 14.8). Η καταλληλότητα καθεμίας μεταβλητής ξεχωριστά αποτυπώνεται στη διαγώνιο του Πίνακα «anti-image matrices» και συγκεκριμένα στο μέρος του Πίνακα «Anti-image correlation» (Εικόνα 14.8). Τιμές στη διαγώνιο μεγαλύτερες από 0,60 αναδεικνύουν την καταλληλότητα των μεταβλητών στις οποίες αντιστοιχούν. Επειδή όλες οι τιμές του πίνακα είναι υψηλότερες από 0,9, δεν χρειάζεται να εξαιρεθεί καμία από τις 12 μεταβλητές.

KMO and Bartlett's Test		
Kaiser-Meyer-Olkin Measure of Sampling Adequacy.		,916
Bartlett's Test of Sphericity	Approx. Chi-Square	1841,662
	df	66
	Sig.	,000

Anti-image Matrices										
	q1	q2	q3	q4	q9	q10	q11	q12		
Anti-image Covariance	q1	.483	-.132	-.074	-.043	-.007	-.060	-.050	.021	
	q2	-.132	.558	-.094	-.031	.003	-.065	.015	-.043	
	q3	-.074	-.094	.513	-.176	.012	-.029	-.083	.060	
	q4	-.043	-.031	-.176	.615	.007	-.063	-.032	.021	
	q9	-.007	.003	.012	.007	.435	-.156	-.103	-.083	
	q10	-.060	-.065	-.029	-.063	-.156	.401	-.058	-.072	
	q11	-.050	.015	-.083	-.032	-.103	-.068	.483	-.116	
	q12	.021	-.043	.050	-.021	-.083	-.072	-.116	.564	
	q13	-.007	-.027	-.064	.008	.014	-.010	-.046	-.044	
	q14	-.050	-.039	-.033	.002	.010	-.059	.021	-.032	
	q15	-.059	.000	.005	-.018	-.102	.052	.012	-.030	
	q16	-.060	-.041	.029	-.034	.005	-.021	.041	-.021	
Anti-image Correlation	q1	.942*	-.255	-.148	-.080	-.015	-.137	-.104	.039	
	q2	-.255	.941*	-.176	-.052	.006	-.137	.028	-.077	
	q3	-.148	-.176	.910*	-.313	.024	-.064	-.166	.093	
	q4	-.080	-.052	-.313	.931*	.014	-.126	-.060	-.036	
	q9	-.015	.006	.024	.014	.900*	-.373	-.224	-.167	
	q10	-.137	-.137	-.064	-.126	-.373	.910*	-.155	-.151	
	q11	-.104	.028	-.166	-.060	-.224	-.155	.929*	-.223	
	q12	.039	-.077	.093	-.036	-.167	-.151	-.223	.939*	
	q13	-.015	-.057	-.141	.017	.033	-.024	-.105	-.092	
	q14	-.086	-.063	-.055	.003	.018	-.113	.037	-.052	
	q15	-.134	-.001	.011	-.037	-.243	.129	.028	-.064	
	q16	-.102	-.063	.047	-.051	.008	-.038	.069	-.033	

a. Measures of Sampling Adequacy(MSA)

Τόσο η τιμή KMO >> 0,6 όσο και ο έλεγχος Bartlett ($p<0,001$) αναδεικνύουν την καταλληλότητα των δεδομένων για ΔΠΑ. Το ίδιο διαπιστώνεται για την καταλληλότητα του δείγματος για καθεμία μεταβλητή αφού οι τιμές στην διαγώνιο του πίνακα Anti-image-Correlation είναι όλες μεγαλύτερες του >0.6

Υποσημείωση: Εξαιτίας του μεγέθους του, παρουσιάζεται μόνο ένα μέρος του πίνακα «ant-image matrices»

Εικόνα 14.8 Καταλληλότητα των συσχετίσεων για το σύνολό τους και καταλληλότητα κάθε μεταβλητής ξεχωριστά.

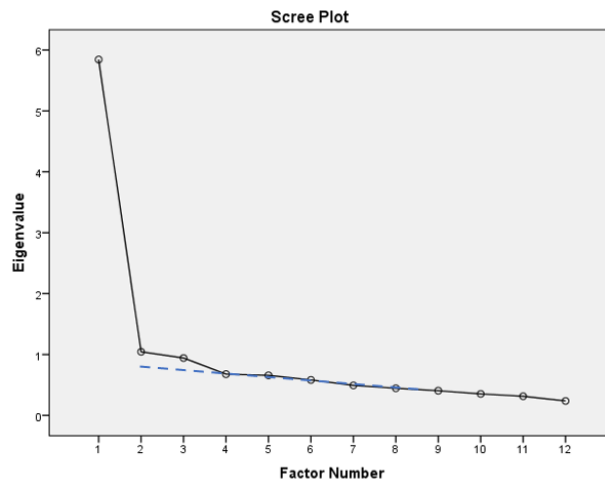
4. Επιλογή του αριθμού των παραγόντων που θα συμπεριληφθούν στην τελική δομή: Μερικές φορές υπάρχει μια συγκεκριμένη υπόθεση η οποία καθορίζει τον αριθμό των παραγόντων που θα συμπεριληφθούν στην τελική δομή, ενώ άλλες φορές απλώς στο τελικό υπόδειγμα λαμβάνεται υπόψη όσο το δυνατόν περισσότερη διακύμανση από εκείνη που υπάρχει μεταξύ των παρατηρούμενων μεταβλητών με όσο το δυνατόν λιγότερους παράγοντες. Εάν έχουμε k μεταβλητές, τότε μπορούμε να εξάγουμε το πολύ k παράγοντες. Υπάρχει ένας αριθμός μεθόδων για τον προσδιορισμό του βέλτιστου αριθμού παραγόντων εξετάζοντας τα δεδομένα. Στο σύγγραμμα αυτό θα παρουσιαστούν δύο τρόποι που προσφέρονται μέσα από τη διαδικασία εκτέλεσης της παραγοντικής ανάλυσης στο SPSS.
- α. Το κριτήριο Kaiser δηλώνει ότι πρέπει να χρησιμοποιηθεί ένας αριθμός παραγόντων ίσος με τον αριθμό των ιδιοτιμών του πίνακα συσχέτισης που είναι μεγαλύτερες από ένα. Από τον Πίνακα (Εικόνα 14.9) προκύπτει η παρουσία 2 παραγόντων με ποσοστό ερμηνευμένης μεταβλητότητας 48,7% και 8,7% αντίστοιχα.

Total Variance Explained						
Factor	Initial Eigenvalues			Extraction Sums of Squared Loadings		
	Total	% of Variance	Cumulative %	Total	% of Variance	Cumulative %
1	5,843	48,692	48,692	5,366	44,718	44,718
2	1,045	8,711	57,404	,594	4,947	49,665
3	,943	7,856	65,260			
4	,677	5,641	70,901			
5	,658	5,482	76,383			
6	,583	4,855	81,238			
7	,493	4,112	85,349			
8	,446	3,719	89,068			
9	,406	3,381	92,450			
10	,354	2,946	95,396			
11	,315	2,624	98,021			
12	,238	1,979	100,000			

Extraction Method: Principal Axis Factoring.

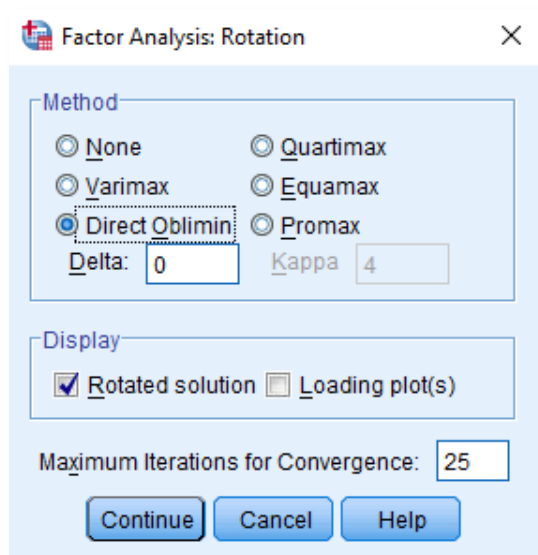
Εικόνα 14.9 Ιδιοτιμές και ποσοστό ερμηνευμένης μεταβλητότητας.

- β. Το «Scree test» δηλώνει ότι πρέπει να απεικονιστούν οι ιδιοτιμές του πίνακα συσχέτισης σε φθίνουσα σειρά και, στη συνέχεια, εξάγεται και ερμηνεύεται ένας αριθμός παραγόντων ίσος με τον αριθμό των ιδιοτιμών που εμφανίζονται από την πρώτη έως και αυτή που συναντάται πριν την τελευταία σημαντική πτώση του μεγέθους της ιδιοτιμής. Στην Εικόνα 14.10 το διάγραμμα των ιδιοτιμών (scree plot) δείχνει μια ισχυρή μείωση της ιδιοτιμής από τον 1^ο στον 2^ο παράγοντα και μια μικρού μεγέθους ελάττωση μετά τον 3^ο παράγοντα. Από τον 4^ο παράγοντα και μετά δεν παρατηρείται διακριτή ελάττωση αλλά τα σημεία φαίνεται να βρίσκονται πάνω σε μία ευθεία γραμμή (διακεκομμένη μπλε γραμμή στην Εικόνα 14.10). Το συμπέρασμα είναι ότι θα μπορούσε να αναδειχτεί μια δομή με έναν ισχυρό γενικό παράγοντα των στάσεων αλλά πιθανόν και μια δομή τριών παραγόντων στον βαθμό που οι παράγοντες είναι ερμηνεύσιμοι, δηλαδή εκφράζουν κάποιες προβλεπόμενες από τη θεωρία πτυχές των στάσεων.



Εικόνα 14.10 «Scree plot» για τον προσδιορισμό του αριθμού των εξαγόμενων παραγόντων.

5. Εξαγωγή των παραγόντων που το πλήθος τους έχει καθοριστεί στο προηγούμενο βήμα: Υπάρχει ένας αριθμός διαφορετικών μεθόδων εξαγωγής, συμπεριλαμβανομένης της μέγιστης πιθανοφάνειας (maximum likelihood (ML) και της εξαγωγής κύριου άξονα (principal axis factoring (PAF)). Η καλύτερη μέθοδος, που απαιτεί όμως μεγάλο δείγμα, είναι γενικά η εξαγωγή μέγιστης πιθανοφάνειας, εκτός και αν παρατηρείται σοβαρή απόκλιση πολυμεταβλητής κανονικότητας στις παρατηρούμενες μεταβλητές. Στο Fabrigar et al. (1999) αναφέρεται ότι οι μέθοδοι κύριων παραγόντων (PCA, PAF, ULS) έχουν το πλεονέκτημα ότι δεν απαιτούν παραδοχές αναφορικά με τις κατανομές των μεταβλητών. Η χρήση της μεθόδου εξαγωγής κύριου άξονα (PAF) προτείνεται στη βιβλιογραφία (Watkins, 2018).
6. Περιστροφή των παραγόντων για εύρεση μιας τελικής εκτίμησης των φορτίσεων και παραγοντικών συσχετίσεων: Για οποιοδήποτε δεδομένο σύνολο συσχετίσεων και αριθμό παραγόντων υπάρχει στην πραγματικότητα ένας μεγάλος αριθμός τρόπων με τους οποίους μπορούμε να προσδιορίσουμε τους παράγοντες μας λαμβάνοντας υπόψη τις συσχετίσεις μεταξύ των παρατηρούμενων μεταβλητών. Περιστρέφοντας τους παράγοντες μας προσπαθούμε να βρούμε μια λύση συντελεστών των παραγόντων στις εξισώσεις (φορτίσεις) που είναι ίση με αυτό που ελήφθη στην αρχική εξαγωγή, αλλά να έχει την απλούστερη ερμηνεία. Υπάρχουν πολλοί διαφορετικοί τύποι περιστροφής, αλλά όλοι προσπαθούν να κάνουν τους παράγοντες μας να σχετίζονται ισχυρά με ένα μικρό υποσύνολο των παρατηρούμενων μεταβλητών και να συνδέονται ασθενώς με τις υπόλοιπες μεταβλητές. Οι δύο μεγάλες κατηγορίες περιστροφών είναι: οι ορθογώνιες περιστροφές, οι οποίες παράγουν ασύνδετους παράγοντες (συσχέτιση περίπου 0), και οι πλάγιες περιστροφές, οι οποίες παράγουν σχετιζόμενους παράγοντες (Watkins, 2018). Η ενδεδειγμένη πρακτική είναι να επιτρέπονται οι συσχετίσεις μεταξύ των παραγόντων, όπως συμβαίνει συνήθως στις κατασκευές που εκφράζονται από αυτούς. Οι πτυχές των στάσεων που πιθανώς εκφράζονται από τα δεδομένα του παραδείγματος, που παρουσιάζεται εδώ, είναι αναμενόμενο να συνδέονται ισχυρά μεταξύ τους, όπως αναδεικνύεται και από την υπάρχουσα βιβλιογραφία. Μια καλή πρακτική για το είδος της περιστροφής που τελικά θα πρέπει να χρησιμοποιηθεί είναι να ελέγχουμε αρχικά το αποτέλεσμα της ανάλυσης με πλάγια περιστροφή και αν εξάγει παράγοντες με μικρή συσχέτιση μεταξύ τους ($|r| < 0,3$), τότε να επαναλαμβάνεται η ανάλυση με τη χρήση ορθογώνιας περιστροφής. Η καλύτερη ορθογώνια περιστροφή θεωρείται ευρέως ότι είναι η Varimax. Οι πλάγιες περιστροφές είναι λιγότερο διακριτές ως προς τη δημοφιλία, με τις δύο πιο συχνά χρησιμοποιούμενες που διαθέτει το SPSS να είναι η Direct Oblimin και η Promax. Η δεύτερη φαίνεται να μειώνει τον χρόνο επεξεργασίας σε μεγάλους πίνακες δεδομένων. Από το πλαίσιο διαλόγου της επιλογής **Rotation** (Εικόνα 14.11) επιλέγεται η **Direct Oblimin**, αφήνοντας την τιμή **Delta** στην προεπιλεγμένη τιμή 0.



Εικόνα 14.11 Πλαίσιο διαλόγου «Factor Analysis:Rotation».

7. Εκτέλεση της ΔΠΑ και ερμηνεία της δομής των παραγόντων: Πριν εκτελέσουμε την τελική ΔΠΑ, σύμφωνα με τις επιλογές που έγιναν στα προηγούμενα βήματα, προκειμένου να διευκολύνουμε τη μελέτη του Πίνακα των φορτίσεων που θα προκύψει από την ανάλυση, στο πλαίσιο **Options** ενεργοποιούμε τις επιλογές **Sorted by size** και **Suppress small coefficients** (Εικόνα 14.6). Η τιμή των φορτίσεων που επιθυμούμε να εμφανίζεται εξαρτάται από το μέγεθος του δείγματος της ανάλυσης. Σε μεγάλα δείγματα, όπως αυτό του παραδείγματος, τιμές μεγαλύτερες του 0,30 είναι ικανοποιητικές, ενώ σε μικρά δείγματα τιμές μεγαλύτερες του 0,5 θα ήταν πιο κατάλληλες.

Από την εκτέλεση της ανάλυσης προκύπτουν ορισμένοι πίνακες που περιγράφονται παρακάτω:

- Ο Πίνακας «Communalities» (Πίνακας (α), Εικόνα 14.12) στην πρώτη στήλη (initial) του οποίου παρουσιάζεται το ποσοστό της διακύμανσης μιας μεταβλητής, το οποίο εξηγείται από τις υπόλοιπες παρατηρούμενες μεταβλητές της ΔΠΑ. Στη στήλη «extraction» του Πίνακα δίνεται το ποσοστό της διακύμανσης της μεταβλητής που εξηγείται από τους παράγοντες που τελικά επιλέχθηκαν για την ανάλυση. Οι χαμηλές τιμές (π.χ. $< 0,2$) της ερμηνευμένης μεταβλητότητας μιας μεταβλητής προτείνουν την εξαίρεσή της από την ανάλυση. Στο παράδειγμά μας, δεν βρέθηκαν τόσο χαμηλές τιμές και συνεπώς όλες οι μεταβλητές συμμετέχουν στην ανάλυση.
- Ο Πίνακας (β) (Εικόνα 14.12) με τίτλο «Factor Matrix» παρουσιάζει τις φορτίσεις των παραγόντων πάνω στις μεταβλητές πριν την περιστροφή των παραγόντων. Καθώς έχουν εξαιρεθεί οι μικρές φορτίσεις ($< 0,30$), που αποτελεί ένα όριο το οποίο προτείνεται από τη βιβλιογραφία (Floyd & Widaman, 1995; Field, 2013), είναι αδύνατη η ερμηνεία τριών παραγόντων σε αυτόν τον Πίνακα, αφού όλες οι μεταβλητές φαίνεται να φορτίζονται από έναν και μοναδικό παράγοντα. Ο Πίνακας (α) (Εικόνα 14.13) με όνομα «Pattern Matrix», ο οποίος δίνει τις φορτίσεις μετά την περιστροφή, είναι ο σημαντικότερος και είναι αυτός που τελικά ερμηνεύεται και παρουσιάζεται στις αναφορές με τα αποτελέσματα μιας ΔΠΑ. Προσοχή στο ότι βασιζόμαστε σε αυτόν όταν έχουμε επιλέξει μια πλάγια περιστροφή. Στην περίπτωση της ορθογώνιας περιστροφής βασιζόμαστε στον Πίνακα «Rotated Component Matrix».
- Η φόρτιση μπορεί να ερμηνευθεί ως τυποποιημένος συντελεστής παλινδρόμησης, που παλινδρομεί τις μεταβλητές πάνω στους παράγοντες. Κάθε μεταβλητή θα σχετίζεται γραμμικά και ισχυρά με έναν από τους παράγοντες και ασθενώς με τους άλλους. Η απουσία των ασθενών φορτίσεων διευκολύνει ακόμα περισσότερο την ερμηνεία των παραγόντων. Αν όμως οι φορτίσεις μιας μεταβλητής είναι μικρότερες από 0,3 σε όλους τους παράγοντες (δεν φαίνεται καμιά φόρτιση στη γραμμή του Πίνακα που αντιστοιχεί σε αυτήν) η μεταβλητή αυτή θα πρέπει να εξαιρεθεί από την ανάλυση και η οποία θα εκτελεστεί εκ νέου. Στο παράδειγμά μας

(Πίνακας 14.5 (α)) δεν παρουσιάζεται τέτοια μεταβλητή, αφού η μικρότερη εμφανιζόμενη φόρτιση είναι 0,35 για την q14 «Στα Μαθηματικά επιβραβεύσαι για τις προσπάθειες που καταβάλλεις» στον 2^ο παράγοντα. Επίσης, μια φόρτιση θεωρείται σταθερή, αν ξεπερνά την τιμή 0,4 (Guadagnoli and Velicer, 1988). Οι παράγοντες που θα ερμηνευτούν πρέπει να διαθέτουν τρεις τουλάχιστον φορτίσεις μεγαλύτερες από 0,4 και αυτές να ανήκουν σε μεταβλητές οι οποίες δεν φορτίζουν και σε άλλους παράγοντες. Από τον Πίνακα (α) (Εικόνα 14.13) είναι σαφές ότι για καθέναν παράγοντα αντιστοιχούν τουλάχιστον 3 μεταβλητές, οι οποίες εμφανίζουν φόρτιση μεγαλύτερη από 0,4 μόνο στον συγκεκριμένο παράγοντα. Οι ισχυρότερες φορτίσεις (τις περισσότερες φορές όλες όσες εμφανίζονται στον Πίνακα) είναι αυτές που ερμηνεύουν και δίνουν όνομα στον παράγοντα. Αυτό σημαίνει ότι η αναγνώριση της κατασκευής βασίζεται στο περιεχόμενο των δηλώσεων-μεταβλητών με τις ισχυρές φορτίσεις στον παράγοντα. Τις περισσότερες φορές, όταν οι ισχυρές φορτίσεις ανήκουν λίγο έως πολύ στις μεταβλητές οι οποίες είναι προκαθορισμένο (από προηγούμενη έρευνα ή από την εγκυρότητα περιεχομένου) να ανήκουν σε μια συγκεκριμένη κατασκευή, η διαδικασία αυτή απλοποιείται.

- Στον Πίνακα (β) (Εικόνα 14.13) παρουσιάζεται το αποτέλεσμα της παραγοντικής ανάλυσης του Πίνακα των πολυχωρικών συσχετίσεων, που όπως αναφέρθηκε και προηγουμένως αποτελεί την κατάλληλη προσέγγιση των συσχετίσεων μεταξύ τακτικών μεταβλητών. Η εκτέλεση της ανάλυσης γίνεται με το πρόσθετο δωρεάν λογισμικό R Factor V2.0 (Basto and Pereira, 2012), το οποίο ενσωματώνεται στο SPSS στο μενού [Analyze=> Dimension Reduction=> R Factor V2.0](#). Από τη συγκριτική μελέτη των Πινάκων (α) και (β) (Εικόνα 14.13) δεν προκύπτουν αξιοσημείωτες διαφορές ούτε ως προς την ταξινόμηση των δηλώσεων στους 3 παράγοντες ούτε ως προς το μέγεθος των φορτίσεων για κάθε δήλωση.
- Είναι λογικό όλες οι ισχυρές φορτίσεις ενός παράγοντα να έχουν το ίδιο (συνήθως θετικό) πρόσημο στον βαθμό που πριν την ανάλυση έγιναν οι αντιστροφές στις αρνητικά διατυπωμένες δηλώσεις (αν υπάρχουν τέτοιες), έτσι ώστε σε όλες πλέον τις μεταβλητές οι υψηλές τιμές να εκφράζουν θετικές στάσεις. Επειδή στο SPSS κατά την επεξεργασία δίνεται αυθαίρετα το πρόσημο του παράγοντα θα πρέπει να δοθεί προσοχή στην ερμηνεία. Αν, για παράδειγμα, όλες οι φορτίσεις είναι αρνητικές, όπως συμβαίνει με τον τρίτο παράγοντα του παραδείγματός μας, μπορούμε να αντιστρέψουμε τις τιμές των αρνητικά διατυπωμένων δηλώσεων (βλ. Κεφ. 2) και έτσι να υποδείξουμε το θετικό νόημα, δηλαδή ότι οι υψηλές τιμές του παράγοντα εκφράζουν «θετικές στάσεις», όπως και οι υπόλοιποι παράγοντες. Μετά την παραπάνω αλλαγή, θα πρέπει να προσέξουμε τα εξής: Στον Πίνακα συσχετίσεων, που αποτελεί σημαντικό επίσης μέρος των αποτελεσμάτων, θα πρέπει να αλλάξει το πρόσημο της συσχέτισης του παράγοντα αυτού με όλους τους παράγοντες οι οποίοι εμφανίζουν θετικές φορτίσεις. Στο παράδειγμά μας, αυτό θα γίνει για τη μεταβλητή του τρίτου παράγοντα.
- Η ερμηνεία των παραγόντων του παραδείγματος είναι απλή επειδή οι δηλώσεις ομαδοποιούνται, όπως προβλέπεται από τους δημιουργούς του ερωτηματολογίου: η ομαδοποίηση των «q1», «q2», «q3» και «q4» στον πρώτο παράγοντα μας επιτρέπει να δηλώσουμε ότι ο παράγοντας εκφράζει την πτυχή της «Συμπεριφορικής Εμπλοκής» και για τον ίδιο λόγο ο δεύτερος παράγοντας εκφράζει τη «Συναισθηματική Εμπλοκή», ενώ ο τρίτος παράγοντας εκφράζει την «Εμπιστοσύνη απέναντι στα Μαθηματικά».
- Στον Πίνακα των συσχετίσεων «Factor correlation Matrix» (Πίνακας (α), Εικόνα 14.14) παρατηρούμε ότι υπάρχουν υψηλές (κατά απόλυτη τιμή) θετικές και αρνητικές συσχετίσεις μεταξύ των τριών παραγόντων. Οι αρνητικές συσχετίσεις οφείλονται σε αυτό που εξηγήθηκε παραπάνω και μετά τη διόρθωσή τους, όπως προαναφέρθηκε, προκύπτει ο Πίνακας (β). Η παρουσία των υψηλών θετικών συσχετίσεων μεταξύ των τριών παραγόντων αφενός αναδεικνύει τη στενή σχέση του επιπέδου θετικότητας της στάσης των μαθητών ως προς τις τρεις πτυχές των στάσεων και αφετέρου την ορθότητα της επιλογής μας για πλάγια περιστροφή των παραγόντων. Επιπλέον, οι ισχυρές σχέσεις υποδεικνύουν και τη δυνατότητα χρήσης ενός γενικού παράγοντα των στάσεων απέναντι στα Μαθηματικά. Αξίζει να προσέξουμε την εντυπωσιακή εγγύτητα των παραγοντικών συσχετίσεων του Πίνακα (β) με εκείνες του Πίνακα (γ) (Εικόνα 14.14), που προκύπτουν από την παραγοντική ανάλυση των πολυχωρικών συσχετίσεων με το R Factor V2.0

Communalities		
	Initial	Extraction
q1	,517	,571
q2	,442	,485
q3	,487	,599
q4	,385	,435
q9	,565	,668
q10	,599	,682
q11	,517	,574
q12	,436	,497
q13	,600	,664
q14	,306	,320
q15	,596	,755
q16	,265	,287

Extraction Method: Principal Axis Factoring.

(α)

Factor Matrix ^a			
	Factor		
	1	2	3
q1	,730		
q2	,662		
q3	,676		,347
q4	,598		
q9	,728		-,344
q10	,767		
q11	,707		
q12	,645		
q13	,742	,337	
q14	,547		
q15	,723	,471	
q16	,474		

Extraction Method: Principal Axis Factoring.

a. 3 factors extracted. 13 iterations required.

(β)

Εικόνα 14.12 (α) Ποσοστό ερμηνευμένης μεταβλητότητας μιας παρατηρούμενης μεταβλητής από τις υπόλοιπες (Initial) και από τους παράγοντες (Extraction) και (β) ο πίνακας φορτίσεων χωρίς περιστροφή «Factor Matrix».

Pattern Matrix ^a			
	Factor		
	1	2	3
q3 Όταν κάνω λάθη στα μαθηματικά, κρυγάζομαι μέχρι να τα διορθώσω	,763		
q4 Όταν δεν μπορώ να λύσω ένα πρόβλημα, συνεχίζω την προσπάθεια επίλυσης του, χρησιμοποιώντας διαφορετικές ιδέες	,592		
q2 Προσπαθώ να απαντάω στις ερωτήσεις που θέτει ο/η καθηγητής/τρια στα μαθηματικά	,539		
q1 Συγκεντρώνομαι πολύ στα μαθηματικά	,527		
q15 Είναι απολαυστικό να μαθαίνει κάποιος μαθηματικά		,875	
q13 Ενδιαφέρομαι να μαθαίνω καινούρια πράγματα στα μαθηματικά		,703	
q16 Αισθάνομαι ικανοί/ση όταν λύνω μαθηματικά προβλήματα		,490	
q14 Στα μαθηματικά επιβραβεύεσαι για τις προσπάθειες που καταβάλλεις		,350	
q9 Έχω μαθηματικό μυαλό			-,828
q10 Μπορώ να παίρνω καλούς βαθμούς στα μαθηματικά			-,680
q12 Έχω αυτοπεποίθηση στα μαθηματικά			-,662
q11 Ξέρω ότι μπορώ να αντιμετωπίσω σε δυσκολίες στα μαθηματικά			-,655

Extraction Method: Principal Axis Factoring.
Rotation Method: Oblimin with Kaiser Normalization.
a. Rotation converged in 11 iterations.

(α)

Sorted Pattern Matrix			
	F1	F2	F3
q3	0,815		
q4	0,627		
q1	0,591		
q2	0,543		
q15		0,981	
q13		0,687	
q16		0,518	
q14		0,324	
q9			-0,917
q10			-0,731
q12			-0,656
q11			-0,649

(β)

Εικόνα 14.13 Ο πίνακας φορτίσεων μετά την περιστροφή των παραγόντων «Pattern Matrix» και ο πίνακας από την ανάλυση των πολυχωρικών συσχετίσεων με περιστροφή «Sorted Pattern Matrix».

Correlation matrix of rotated factors			
	F1	F2	F3
F1	1,000	,602	-,673
F2	,602	1,000	-,639
F3	-,673	-,639	1,000

(α)

Factor Correlation Matrix			
Factor	1	2	3
1	1,000	,598	-,670
2	,598	1,000	-,641
3	-,670	-,641	1,000

Extraction Method: Principal Axis Factoring.
Rotation Method: Oblimin with Kaiser Normalization.

(β)

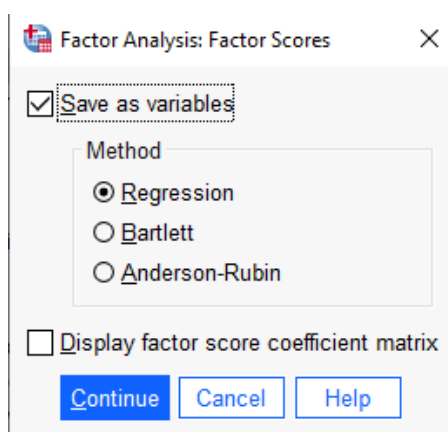
Factor Correlation Matrix			
Factor	1	2	3
1	1,000	,598	,670
2	,598	1,000	,641
3	,670	,641	1,000

Extraction Method: Principal Axis Factoring.
Rotation Method: Oblimin with Kaiser Normalization.

(γ)

Εικόνα 14.14 Πίνακες συσχετίσεων μεταξύ των παραγόντων από την PAF πριν (α) και μετά (β) τη διόρθωση των προσήμων και ο πίνακας από την ανάλυση των πολυχωρικών συσχετίσεων (γ).

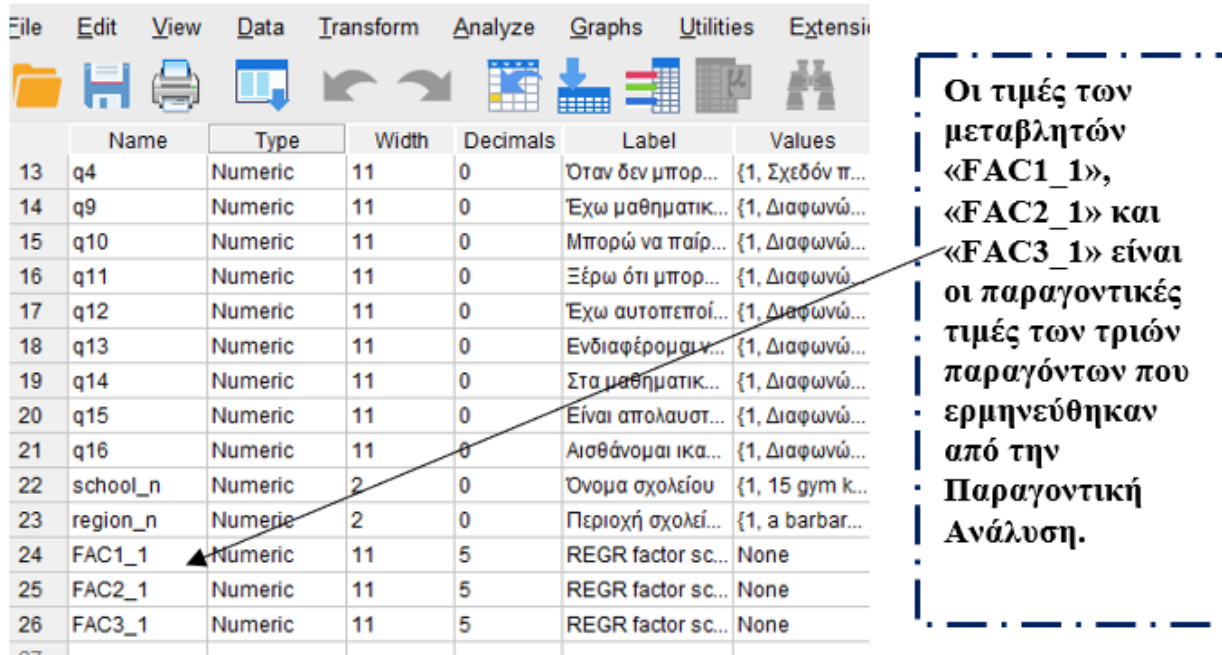
8. Μετά την οριστικοποίηση της δομής και την ερμηνεία των παραγόντων είναι πολλές φορές αναγκαίο να εκτιμηθούν οι παραγοντικές τιμές για κάθε περίπτωση του δείγματος, έτσι ώστε να μπορούν να χρησιμοποιηθούν οι παράγοντες ως μεταβλητές για περαιτέρω ανάλυση. Η τιμή που αντιστοιχεί σε μία περίπτωση για έναν δεδομένο παράγοντα είναι ο γραμμικός συνδυασμός όλων των μεταβλητών, σταθμισμένος με την αντίστοιχη συντελεστική φόρτιση.
 - Μεταξύ των τριών μεθόδων, που χρησιμοποιούνται, η καθεμία έχει τα πλεονεκτήματα και τα μειονεκτήματά της. Η μέθοδος παλινδρόμησης *Regression* μεγιστοποιεί τη συσχέτιση μεταξύ των παραγοντικών τιμών και του υποκείμενου παράγοντα και προφανώς μεγιστοποιεί την εγκυρότητα, αλλά οι βαθμολογίες μπορεί να είναι κάπως μεροληπτικές. Αυτό σημαίνει ότι ακόμα και αν έχει επιλεγεί μια ορθογώνια περιστροφή των παραγόντων, οι παράγοντες που θα δημιουργηθούν στο αρχείο δεδομένων δεν θα έχουν μηδενική *Bartlett's* μεταξύ τους. Στη μέθοδο *Bartlett* οι παραγοντικές τιμές που δημιουργούνται έχουν υψηλή συσχέτιση με τον υποκείμενο παράγοντα και λιγότερο με τους άλλους και επιπλέον δεν είναι μεροληπτικές ως εκτιμήσεις των πραγματικών παραγοντικών τιμών. Η μέθοδος *Anderson-Rubin* είναι κατάλληλη για ορθογώνιες περιστροφές διότι σχεδόν πάντα οδηγεί σε μη σχετιζόμενες παραγοντικές τιμές, ενώ για πλάγια περιστροφή προτείνεται περισσότερο η μέθοδος *Regression*.
 - Για τον υπολογισμό των παραγοντικών τιμών επιλέγουμε από το κεντρικό πλαίσιο διαλόγου της ανάλυσης την επιλογή **Scores** και στο πλαίσιο διαλόγου που εμφανίζεται (Εικόνα 14.15) μπορούμε να επιλέξουμε τον τρόπο υπολογισμού από τους τρεις προτεινόμενους και έπειτα ενεργοποιούμε την επιλογή **Save as variables**.



Εικόνα 14.15 Πλαίσιο διαλόγου «Factor Analysis: Factor Scores».

- Το αποτέλεσμα της δημιουργίας των παραγοντικών τιμών φαίνεται στην Εικόνα 14.16. Συγκεκριμένα, δημιουργούνται τρεις καινούριες μεταβλητές που τοποθετούνται ως τελευταίες. Εναλλακτικά, για τον υπολογισμό των παραγοντικών τιμών θα μπορούσαμε να υπολογίσουμε το άθροισμα ή το μέσο όρο των απαντήσεων (τιμές 1-5) του ατόμου στις δηλώσεις που αντιστοιχούν στον παράγοντα. Ως καλή πρακτική προτείνεται ο μέσος όρος, αφού αυτός

παίρνει τιμές από 1 έως και 5 και επομένως η όποια παραγοντική τιμή μπορεί να ερμηνευτεί χρησιμοποιώντας τις δυνατές απαντήσεις των δηλώσεων που συγκροτούν τον παράγοντα.



	Name	Type	Width	Decimals	Label	Values
13	q4	Numeric	11	0	Όταν δεν μπορ...	{1, Σχεδόν π...
14	q9	Numeric	11	0	Έχω μαθηματικ...	{1, Διαφωνώ...
15	q10	Numeric	11	0	Μπορώ να παίρ...	{1, Διαφωνώ...
16	q11	Numeric	11	0	Ξέρω ότι μπορ...	{1, Διαφωνώ...
17	q12	Numeric	11	0	Έχω αυτοπεποί...	{1, Διαφωνώ...
18	q13	Numeric	11	0	Ενδιαφέρομαι...	{1, Διαφωνώ...
19	q14	Numeric	11	0	Στα μαθηματικ...	{1, Διαφωνώ...
20	q15	Numeric	11	0	Είναι απολαυστ...	{1, Διαφωνώ...
21	q16	Numeric	11	0	Αισθάνομαι ικα...	{1, Διαφωνώ...
22	school_n	Numeric	2	0	Όνομα σχολείου	{1, 15 gym k...
23	region_n	Numeric	2	0	Περιοχή σχολεί...	{1, a barbar...
24	FAC1_1	Numeric	11	5	REGR factor sc...	None
25	FAC2_1	Numeric	11	5	REGR factor sc...	None
26	FAC3_1	Numeric	11	5	REGR factor sc...	None

Εικόνα 14.16 Πλαίσιο διαλόγου «Variable view» με τις τρεις καινούριες μεταβλητές που περιλαμβάνουν τις υπολογισμένες παραγοντικές τιμές.

14.2 Παραγοντική εγκυρότητα και αξιοπιστία των παραγόντων

- Οι ψυχομετρικές ιδιότητες μιας κλίμακας δηλώσεων συνιστάται κυρίως στον έλεγχο της εγκυρότητας και της αξιοπιστίας της. Οι ιδιότητες αυτές της κλίμακας των δηλώσεων που χρησιμοποιήθηκε για τη μέτρηση των στάσεων μαθητών απέναντι στη χρήση υπολογιστικού περιβάλλοντος κατά τη διδασκαλία και Μαθηματικών, πρέπει να διερευνηθούν μετά την εξαγωγή και ταυτοποίηση των παραγόντων που παρουσιάστηκε στην Εικόνα 14.13.
- Οι δύο πλευρές της εγκυρότητας αναφέρονται ως συγκλίνουσα (convergent) και διακρίνουσα (discriminant) εγκυρότητα. Η συγκλίνουσα εγκυρότητα αναδεικνύει τον βαθμό της σχέσης ανάμεσα στις παρατηρούμενες μεταβλητές που προσδιορίζουν τον παράγοντα με τον ίδιο τον παράγοντα. Η ποσοτική έκφραση της συγκλίνουσας εγκυρότητας δίνεται τόσο από το μέγεθος των φορτίσεων του παράγοντα πάνω στις μεταβλητές που του αντιστοιχούν όσο και από ένα συνολικό μέτρο που εκφράζει τη μέση μεταβλητότητα της παρατηρούμενης μεταβλητής που ερμηνεύεται από τον παράγοντα. Το μέτρο αυτό αναφέρεται ως *Μέση Εξαγόμενη Μεταβλητότητα* (Average Variance Extracted, AVE) ενός παράγοντα F_i και υπολογίζεται ως εξής:

$$AVE(F_i) = \frac{\sum_{j=1}^{K_i} \lambda_{ij}^2}{K_i}, \quad (14.2)$$

όπου

i = αύξων αριθμός του παράγοντα,

j = αριθμός φόρτισης,

λ_{ij}^2 = το τετράγωνο της φόρτισης του παράγοντα i στη μεταβλητή j ,

K_i = το πλήθος των μεταβλητών που φορτίζονται στον παράγοντα F_i .

- Σύμφωνα με τους Hair et al. (2016) οι τιμές της AVE πάνω από 0,4 δηλώνουν οριακή συγκλίνουσα εγκυρότητα, ενώ τιμές μεγαλύτερες του 0,5 υποδεικνύουν ικανοποιητική συγκλίνουσα εγκυρότητα. Στο παράδειγμά μας, εκτελώντας τις πράξεις (χρησιμοποιώντας τις αντίστοιχες φορτίσεις), βρίσκουμε τη μέση εξαγόμενη διακύμανση για κάθε παράγοντα. Για παράδειγμα, για τον 1^ο παράγοντα:

$$AVE = (0,763^2 + 0,592^2 + 0,539^2 + 0,527^2)/4 = 0,38.$$

- Οι τιμές της μέσης εξαγόμενης διακύμανσης είναι οριακά ικανοποιητικές για τον 1^ο και 2^ο παράγοντα, αλλά ικανοποιητικές για τον 3^ο. Οι αντίστοιχες φορτίσεις από τον Πίνακα 14.5(β) της ανάλυσης με πολυχωρικές συσχετίσεις έδωσαν παρόμοιες εκτιμήσεις συγκλίνουσας εγκυρότητας (οι τιμές σε παρένθεση στη στήλη AVE του Πίνακα 14.8) και για τους 3 παράγοντες. Η διακρίνουσα εγκυρότητα αναφέρεται στον βαθμό στον οποίο ένας παράγοντας διαφέρει από κάποιον άλλο παράγοντα της ίδιας κλίμακας. Για τη διακρίνουσα εγκυρότητα έχουν προταθεί διάφορα κριτήρια. Το απλούστερο κριτήριο είναι να μην υπάρχουν φορτίσεις μεταξύ των παραγόντων μεγαλύτερες από 0,8 κατά απόλυτο τιμή. Επίσης η διακρίνουσα εγκυρότητα μπορεί να βρεθεί με τη χρήση του κριτηρίου Fornell-Larcker (Fornell & Larcker, 1981), σύμφωνα με το οποίο η τετραγωνική ρίζα της AVE ενός παράγοντα πρέπει να είναι μεγαλύτερη από όλες τις συσχετίσεις του παράγοντα με τους υπόλοιπους παράγοντες. Στην παρουσίασή μας θα χρησιμοποιηθεί το πηλίκιο *Heterotrait-Monotrait Ratio of Correlations* (HTMT) που προτείνουν οι Henseler et al. (2015) και το οποίο εκφράζει το ποσοστό της συσχέτισης των παρατηρούμενων μεταβλητών μεταξύ δύο παραγόντων προς τη συσχέτιση εντός των παραγόντων. Όσο μικρότερο είναι το ποσοστό τόσο πιο ικανοποιητική είναι η διακρίνουσα εγκυρότητα των δύο παραγόντων.

$$HTMT(F_i, F_j) = \frac{c}{\sqrt{A \cdot B}} \quad (14.3)$$

όπου

$C = H$ μέση τιμή των συσχετίσεων ανάμεσα στα ζεύγη των μεταβλητών, εκ των οποίων μία συνδέεται με τον παράγοντα F_i και η άλλη με τον F_j .

$A =$ η μέση τιμή των συσχετίσεων για όλα τα ζεύγη των μεταβλητών του F_i .

$B =$ η μέση τιμή των συσχετίσεων για όλα τα ζεύγη των μεταβλητών του F_j .

- Οι τιμές του πηλίκου HTMT κάτω από το 0,85 εκφράζουν ικανοποιητική διακρίνουσα εγκυρότητα (Henseler et al., 2015). Στο SPSS δεν υπάρχει υπολογισμός του πίνακα HTMT, συνεπώς οι υπολογισμοί του Πίνακα 14.8 έχουν γίνει εύκολα σε ένα φύλλο του Excel αντιγράφοντας τον πίνακα συσχετίσεων που δημιουργούνται από την ενεργοποίηση της επιλογής **coefficients** στο πλαίσιο διαλόγου **Descriptives** (Εικόνα 14.5). Από την επισκόπηση των τιμών στη στήλη HTMT του Πίνακα (Εικόνα 14.20) διαπιστώνεται η παρουσία ικανοποιητικής διακρίνουσας εγκυρότητας των τριών παραγόντων, αφού καμία τιμή δεν είναι μεγαλύτερη του 0,85.
- Η αξιοπιστία και συγκεκριμένα η πτυχή της που αναφέρεται στην εσωτερική συνοχή μιας κλίμακας εκφράζεται από τον συντελεστή α του Cronbach (Cronbach, 1951). Ο συντελεστής αυτός εκφράζει το ποσοστό της μεταβλητότητας μια μονοδιάστατης κλίμακας (ενός παράγοντα), ως άθροισμα των μεταβλητών που συμπεριλαμβάνει, που οφείλεται στην παρουσία συμμεταβολής (συσχέτισης) μεταξύ των μεταβλητών που την απαρτίζουν. Όσο μεγαλύτερο είναι αυτό το ποσοστό τόσο μεγαλύτερη είναι και η συνοχή της. Οι τιμές του συντελεστή ανήκουν στο διάστημα [0, 1]. Επιθυμητές τιμές του συντελεστή είναι αυτές που υπερβαίνουν την τιμή 0,7. Αναλυτικότερα, τιμές από 0,6 μέχρι 0,7 θεωρούνται οριακά αποδεκτές, και τιμές από 0,7 έως 0,9 ικανοποιητικές (Hair et al., 2016). Ωστόσο, τιμές του συντελεστή μεγαλύτερες του 0,9 και κυρίως μεγαλύτερες του 0,95 πρέπει να αξιολογούνται με προσοχή, αφού οι τιμές αυτές υποδεικνύουν ότι όλες οι πτυχές της έννοιας που θέλουμε να μετρήσουμε είναι σχεδόν ίδιες (Hair et al., 2016).

- Ο υπολογισμός του α γίνεται από τον τύπο:

$$\alpha = \frac{k(\overline{cov}/\overline{var})}{1 + (k - 1)(\overline{cov}/\overline{var})} \quad (14.4)$$

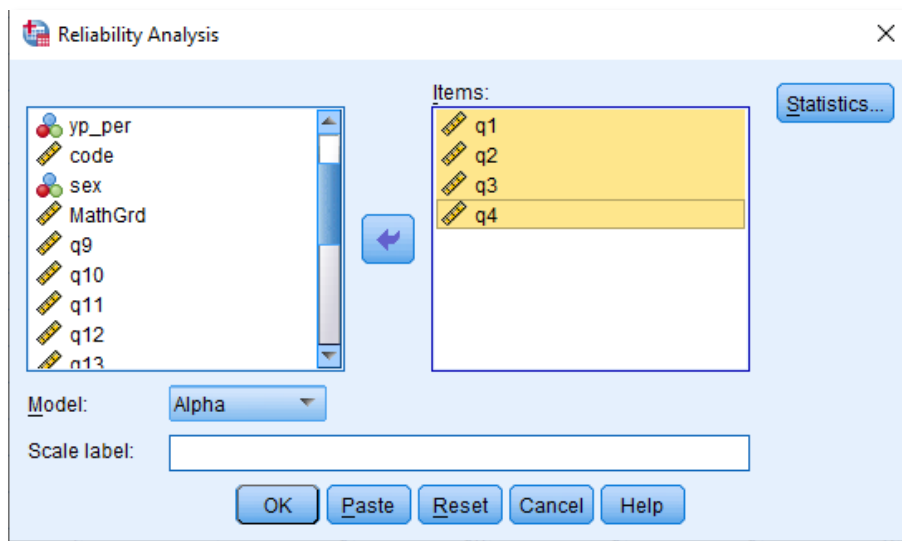
όπου

k = το πλήθος των μεταβλητών της κλίμακας.

$\overline{cov} = H$ μέση συνδιακύμανση των $\frac{k(k-1)}{2}$ ζευγών μεταβλητών της κλίμακας.

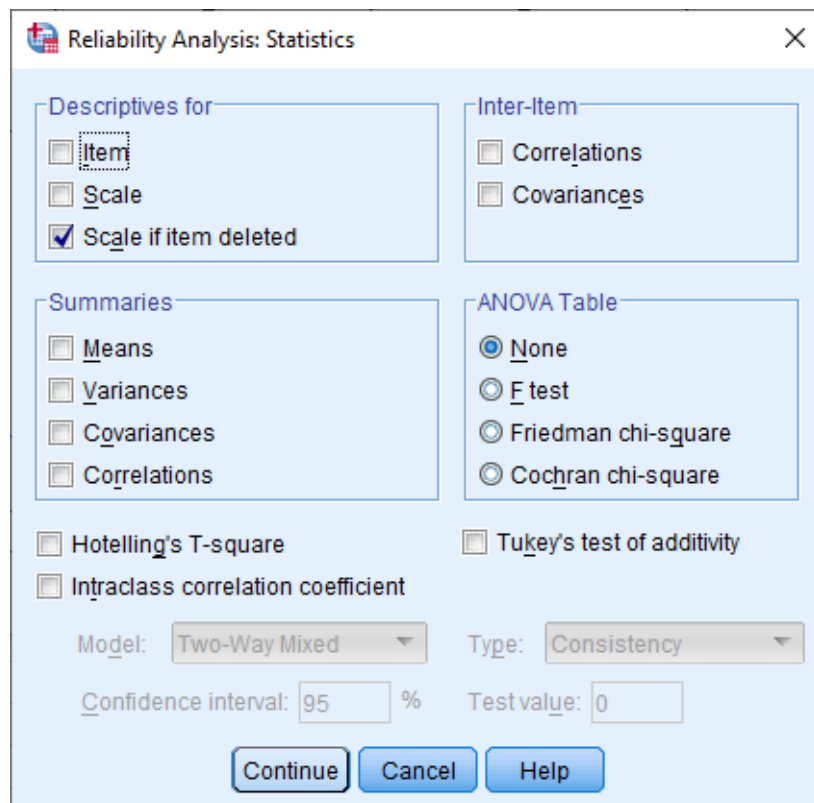
$\overline{var} = H$ μέση διακύμανση των μεταβλητών της κλίμακας.

- Προσοχή: σύμφωνα με τον τύπο υπολογισμού του συντελεστή Cronbach φαίνεται ότι όσο αυξάνεται ο αριθμός των δηλώσεων που συμπληρώνει μια κλίμακα υπάρχει περίπτωση (αν έχουν απαντηθεί παρόμοια) να αυξηθεί και ο συντελεστής χωρίς απαραίτητα να αυξάνεται και η πραγματική αξιοπιστία μέτρησης της κλίμακας. Αξίζει να τονίσουμε ότι, όταν ο συντελεστής αυτός υπερβαίνει το 0,7, δεν σημαίνει απαραίτητα ότι το σύνολο των δηλώσεων της κλίμακας υποδεικνύει μια μονοδιάστατη κλίμακα. Επίσης, σε πολλές περιπτώσεις μια κλίμακα μπορεί να μετρά επιμέρους διαστάσεις (υποκλίμακες). Για παράδειγμα, οι διαστάσεις της παιδαγωγικής αξίας μιας τεχνολογικής εφαρμογής και της ευχρηστίας της μπορεί να μην συνδέονται και επομένως η συνολική εκτίμηση της αξιοπιστίας όλης της κλίμακας δεν έχει νόημα, αλλά η εκτίμηση της αξιοπιστίας των δύο επιμέρους διατάσεων έχει νόημα.
- Ο υπολογισμός του συντελεστή α του Cronbach στο SPSS γίνεται επιλέγοντας από το μενού **Analyze=>Scale=>Reliability Analysis**. Στο πλαίσιο διαλόγου **Reliability Analysis** (Εικόνα 14.17) επιλέγουμε και τοποθετούμε στην περιοχή **Items** τις μεταβλητές που συνδέονται με τη μονοδιάστατη κλίμακα ή παράγοντα που προέκυψε από την παραγοντική ανάλυση.



Εικόνα 14.17 Πλαίσιο διαλόγου «Reliability Analysis».

- Στη συνέχεια, από την επιλογή **Statistics** (Εικόνα 14.18) ενεργοποιούμε μια ενδιαφέρουσα επιλογή, την **Scale if item deleted**, η οποία μπορεί να συμβάλει στο να ταυτοποιηθούν οι μεταβλητές που ευθύνονται για τη χαμηλή τιμή του α .



Εικόνα 14.18 Πλαίσιο διαλόγου «Reliability Analysis:Statistics».

- Το αποτέλεσμα των επιλογών δίνεται στον Πίνακα (Εικόνα 14.19). Από τον Πίνακα (α) (Εικόνα 14.19) διαπιστώνεται η ικανοποιητική τιμή 0,81 της αξιοπιστίας του παράγοντα. Στη στήλη «Cronbach's- a if item deleted» του Πίνακα (β) (Εικόνα 14.19) δίνεται η τιμή του συντελεστή, αν η μεταβλητή της γραμμής του Πίνακα που αντιστοιχεί εξαιρεθεί από τον υπολογισμό του. Στην περίπτωση μιας μη αποδεκτής τιμής (π.χ. $<0,7$) η εύρεση της μεγαλύτερης τιμής αυτής της στήλης, αν αυτή είναι μεγαλύτερη της αρχικής μη αποδεκτής τιμής, οδηγεί σε βελτίωση του συντελεστή. Στη συνέχεια, υπολογίζεται εκ νέου ο συντελεστής χωρίς τη μεταβλητή αυτή και η διαδικασία θα μπορούσε να συνεχιστεί έως ότου ο συντελεστής βρεθεί σε αποδεκτά όρια ή δεν μπορεί πλέον να βελτιωθεί.
- Στην περίπτωση του παράγοντα της «Συμπεριφορικής Εμπλοκής» που μελετάται στον πίνακα που φαίνεται στην εικόνα 14.20 καμιά τιμή δεν βελτιώνει την τιμή 0,81. Αντίθετα, οποιαδήποτε εξαίρεση μεταβλητής μειώνει την αξιοπιστία του παράγοντα και έτσι δεν πρέπει να προβούμε σε καμία εξαίρεση. Από τον Πίνακα 14.8 διαπιστώνεται η ικανοποιητική αξιοπιστία και των τριών παραγόντων του ερωτηματολογίου των στάσεων.
- Το αποτελέσματα της ανάλυσης των ψυχομετρικών ιδιοτήτων της κλίμακας των στάσεων των μαθητών απέναντι στη χρήση υπολογιστικού περιβάλλοντος κατά τη διδασκαλία των Μαθηματικών, όπως παρουσιάζονται στον Πίνακα (Εικόνα 14.20), επιβεβαιώνουν την παρουσία καλών ψυχομετρικών ιδιοτήτων και για τους τρεις παράγοντες των στάσεων. Έτσι, η μέτρηση που βασίζεται σε αυτή την κλίμακα (υπολογισμένες παραγοντικές τιμές) μπορεί να χρησιμοποιηθεί στη διερεύνηση ερευνητικών ερωτημάτων στα οποία εμπλέκονται οι συγκεκριμένες στάσεις. Στις διερευνήσεις αυτές οι παράγοντες θα εκπροσωπούνται από τις παραγοντικές τιμές που έχουν ήδη αποθηκευτεί στο αρχείο δεδομένων (Εικόνα 14.16).

Reliability Statistics	
Cronbach's Alpha	N of Items
,805	4

(α)

Item-Total Statistics				
	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation	Cronbach's Alpha if Item Deleted
q1	9,95	10,629	,643	,747
q2	9,89	10,379	,608	,763
q3	10,29	10,159	,661	,737
q4	10,46	10,401	,575	,779

(β)

Εικόνα 14.19 Η τιμή του α-Cronbach και διαγνωστικά στατιστικά για τη βελτίωσή του στην περίπτωση του παράγοντα της «Συμπεριφορικής Εμπλοκής».

Παράγοντας	AVE	HTMT	Cronbach's a
Συμπεριφορικής Εμπλοκής (1 ^{ος})	0,38 (0,43)	1 ^{ος} x2 ^{ος} =0,77	0,81
Συναισθηματική εμπλοκή(2 ^{ος})	0,41 (0,45)	1 ^{ος} x3 ^{ος} =0,79	0,78
Εμπιστοσύνη (3 ^{ος})	0,50 (0,56)	2 ^{ος} x3 ^{ος} =0,73	0,84

Εικόνα 14.20 Εγκυρότητα και αξιοπιστία της κλίμακας των στάσεων των μαθητών απέναντι στη χρήση υπολογιστικού περιβάλλοντος κατά τη διδασκαλία των Μαθηματικών. Σε παρένθεση βρίσκονται οι τιμές από την ανάλυση με πολυχωρικές συσχετίσεις.

Παρουσίαση των αποτελεσμάτων της παραγοντικής ανάλυσης

- Ενδεικτικά για την παρουσίαση των αποτελεσμάτων της παραγοντικής ανάλυσης γράφουμε τα εξής: Για την παραγοντική ανάλυση των 12 δηλώσεων αναφορικά με τις στάσεις των μαθητών απέναντι στη χρήση υπολογιστικού περιβάλλοντος κατά τη διδασκαλία των Μαθηματικών εφαρμόστηκαν η μέθοδος εκτίμησης των παραγόντων «εξαγωγή κύριου άξονα» (PAF) και η πλάγια περιστροφή. Το μέτρο KMO=0,916 επιβεβαίωσε τη δειγματική επάρκεια των δεδομένων για την ανάλυση αυτή. Το «scree plot» υπέδειξε μια παραγοντική δομή τριών παραγόντων που εξηγούν συνολικά το 65,2% της συνολικής διακύμανσης των απαντήσεων των μαθητών σε αυτές τις δηλώσεις. Σύμφωνα με τον Πίνακα (α) (που θα πρέπει να παρουσιαστεί συνοδευτικά) (Εικόνα 14.13) ο πρώτος παράγοντας εκφράζει την πτυχή της «Συμπεριφορικής Εμπλοκής» και σε αυτόν ομαδοποιούνται οι δηλώσεις «q1», «q2», «q3» και «q4». Ο δεύτερος παράγοντας εκφράζει τη «Συναισθηματική Εμπλοκή» και σε αυτόν ομαδοποιούνται οι δηλώσεις «q13», «q14», «q15» και «q16». Τέλος, ο τρίτος παράγοντας εκφράζει την «Εμπιστοσύνη απέναντι στα Μαθηματικά» και σε αυτόν ομαδοποιούνται οι δηλώσεις «q9», «q10», «q11» και «q12».
- Επιπλέον και όσον αφορά την εγκυρότητα των παραγόντων (συνοδευτικά θα πρέπει να παρουσιάσουμε τον Πίνακα (Εικόνα 14.20) ή τους αντίστοιχους δείκτες μέσα στο κείμενο), παρατηρείται οριακή συγκλίνουσα εγκυρότητα για τους δύο πρώτους παράγοντες (AVE1=0,38, AVE2=0,41) και ικανοποιητική εγκυρότητα για τον τρίτο παράγοντα (AVE3=0,50). Επιπρόσθετα, τόσο η διακρίνουσα εγκυρότητα (HTMT ration <0,85) όσο και η αξιοπιστία των παραγόντων (οι τρεις συντελεστές α είναι μεγαλύτεροι του 0,7) αξιολογούνται ως ικανοποιητικές.

14.3 Προβλήματα για εξάσκηση και ανατροφοδότηση

Άσκηση 1

Για τη μέτρηση της επαγγελματικής εξουθένωσης των εκπαιδευτικών χρησιμοποιήθηκε το εργαλείο MBI (Maslach Burnout Inventory) (Maslach et al., 1997). Το ερωτηματολόγιο αποτελείται από 22 στοιχεία, τα οποία κατανέμονται σε τρεις υποκλίμακες (παράγοντες ή διαστάσεις): α) συναισθηματική εξάντληση, β) αποπροσωποποίηση και γ) έλλειψη του αισθήματος της προσωπικής επίτευξης.

Διαστάσεις	Δηλώσεις
Συναισθηματική εξάντληση	q1, q2, q3, q6, q8, q13, q14, q16, q20
Αποπροσωποποίηση	q5, q10, q11, q15, q22
Έλλειψη του αισθήματος της προσωπικής επίτευξης	q4, q7, q9, q12, q17, q18, q19, q21

Στο αρχείο «JS_Burnout fa.sav» έχουμε τις απαντήσεις 158 εκπαιδευτικών πρωτοβάθμιας εκπαίδευσης του Νομού Α'. Δείγματα μόνο των δηλώσεων αναφέρονται στα labels των αντίστοιχων μεταβλητών.

- Να ελέγξετε τον πίνακα συσχετίσεων των μεταβλητών που αντιστοιχούν στις δηλώσεις.
- Χρησιμοποιώντας το «scree plot» να τεκμηριώσετε τη δομή των τριών παραγόντων.
- Να εκτελέσετε την παραγοντική ανάλυση με τη μέθοδο εκτίμησης των παραγόντων PAF για την εξαγωγή τριών παραγόντων.
- Να τεκμηριώσετε την ανάγκη για πλάγια περιστροφή της παραγοντικής δομής.
- Ποιες από τις δηλώσεις φορτίζουν περισσότερο σε άλλους παράγοντες από τους αναμενόμενους;
- Για την παραγοντική δομή που προκύπτει να υπολογίσετε τους δείκτες AVE και HTMT και α του Cronbach.
- Να γράψετε τα συμπεράσματά σας τόσο για την παραγοντική δομή που ταιριάζει με τα δεδομένα σας όσο και για τις ψυχομετρικές ιδιότητες της κλίμακας αυτής.

Άσκηση 2

Για τη μέτρηση των στάσεων των φοιτητών και μελλοντικών Φιλολόγων απέναντι στη διδασκαλία της Ιστορίας δημιουργήθηκε από τους ερευνητές (Michala et al, 2022) μια ομάδα 16 δηλώσεων και διανεμήθηκε σε 120 φοιτητές Φιλολογικού Τμήματος. Στο αρχείο «Attitudes towards teaching history fa.sav» έχουμε τις απαντήσεις 120 φοιτητών στις 16 δηλώσεις.

- Να ελέγξετε τον πίνακα συσχετίσεων των μεταβλητών που αντιστοιχούν στις δηλώσεις.
- Χρησιμοποιώντας το «scree plot» να τεκμηριώσετε τη δομή των τεσσάρων παραγόντων.
- Να εκτελέσετε την παραγοντική ανάλυση με τη μέθοδο εκτίμησης των παραγόντων PAF για την εξαγωγή τεσσάρων παραγόντων.
- Να τεκμηριώσετε την ανάγκη για πλάγια περιστροφή της παραγοντικής δομής.
- Για την παραγοντική δομή που προκύπτει να ερμηνεύσετε τους παράγοντες και να υπολογίσετε τους δείκτες AVE και HTMT και α του Cronbach.
- Να γράψετε τα συμπεράσματά σας τόσο για την παραγοντική δομή που ταιριάζει με τα δεδομένα σας όσο και για τις ψυχομετρικές ιδιότητες της κλίμακας αυτής.

Άσκηση 3

Για τη μέτρηση των αντιλήψεων των εκπαιδευτικών φιλολογικών μαθημάτων σχετικά με τη γνώση του περιεχομένου που διαθέτουν, δημιουργήθηκε από τους ερευνητές μια ομάδα 7 δηλώσεων και διανεμήθηκε σε 111 Φιλολόγους. Στο αρχείο «CK one dimension fa.sav» έχουμε τις απαντήσεις 111 Φιλολόγων στις 7 δηλώσεις (μεταβλητές «CK1» έως και «CK7») αναφορικά με τις αντιλήψεις τους για τη γνώση του περιεχομένου που διαθέτουν, καθώς και τις απαντήσεις τους για το φύλο και τη διδακτική τους εμπειρία (μεταβλητές «q1» και «q2» αντίστοιχα).

- Χρησιμοποιώντας το «scree plot» να τεκμηριώσετε τη δομή ενός παράγοντα (μονοδιάστατη δομή).
- Να εκτελέσετε την παραγοντική ανάλυση με τη μέθοδο εκτίμησης των παραγόντων PAF για την εξαγωγή του ενός παράγοντα.

- γ. Για την παραγοντική δομή που προκύπτει να ερμηνεύσετε τον παράγοντα και να υπολογίσετε τους δείκτες AVE και α του Cronbach.
- δ. Να γράψετε τα συμπεράσματά σας τόσο για την παραγοντική δομή που ταιριάζει με τα δεδομένα σας όσο και για τις ψυχομετρικές ιδιότητες της κλίμακας αυτής.
- ε. Να υπολογίσετε τις παραγοντικές τιμές με τη μέθοδο Regression.
- στ. Να ελέγξετε κατά πόσο διαφοροποιούνται οι αντιλήψεις των Φιλολόγων ανάλογα με το φύλο και ανάλογα με τη διδακτική εμπειρία (κύριες επιδράσεις). Υφίσταται αλληλεπίδραση των δύο αυτών παραγόντων (φύλο και διδακτική εμπειρία) πάνω στις αντιλήψεις τους; Γράψτε τα συμπεράσματά σας.

Βιβλιογραφία

- Barkatsas, A., Kasimatis, K., & Gialamas, V. (2009). *Learning secondary mathematics with technology: Exploring the complex interrelationship between students' attitudes, engagement, gender and achievement*. *Computers & Education*, 52(3), 562-570. <https://doi.org/10.1016/j.compedu.2008.11.001>
- Basto, M., & Pereira, J. M. (2012). An SPSS R-Menu for Ordinal Factor Analysis. *Journal of Statistical Software*, 46(4), 1–29. <https://doi.org/10.18637/jss.v046.i04>
- Beavers, A. S., Lounsbury, J. W., Richards, J. K., Huck, S. W., Skolits, G. J., & Esquivel, S. L. (2013). Practical considerations for using exploratory factor analysis in educational research. *Practical Assessment, Research, and Evaluation*, 18(1), 6. <https://doi.org/10.7275/qv2q-rk76>
- Cronbach, L. J. (1951). Coefficient alpha and the internal structure of tests. *Psychometrika*, 16, 297–334. <https://doi.org/10.1007/BF02310555>
- Fabrigar, L. R., Wegener, D. T., MacCallum, R. C., & Strahan, E. J. (1999). Evaluating the use of exploratory factor analysis in psychological research. *Psychological Methods*, 4(3), 272–299. <https://doi.org/10.1037/1082-989X.4.3.272>
- Floyd, F. J., & Widaman, K. F. (1995). Factor analysis in the development and refinement of clinical assessment instruments. *Psychological Assessment*, 7(3), 286–299. <https://doi.org/10.1037/1040-3590.7.3.286>
- Field, A. (2013). *Discovering statistics using IBM SPSS statistics* (5th ed.). London: SAGE.
- Fornell, C., & Larcker, D. F. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of Marketing Research*, 18(1), 39–50.
- Hair, J.F., Jr., Hult, G.T.M., Ringle, C., & Sarstedt, M. (2016), *A Primer on Partial Least Squares Structural Equation Modeling (PLS-SEM)*, Sage.
- Henseler, J., Ringle, C.M., & Sarstedt, M. (2015), A new criterion for assessing discriminant validity in variance-based structural equation modeling, *Journal of the Academy of Marketing Science*, 43(1), 1–21. <https://doi.org/10.1007/s11747-014-0403-8>
- Kim, H-Y. (2013). Statistical notes for clinical researchers: Assessing normal distribution (2) using skewness and kurtosis. *Restorative Dentistry and Endodontics*, 38, 52–54. <https://doi.org/10.5395/rde.2013.38.1.52>
- Maslach, C., Jackson, S. E., & Leiter, M. P. (1997). *Maslach burnout inventory*. Scarecrow Education.
- Suh, Y. (2015). The performance of maximum likelihood and weighted least square mean and variance adjusted estimators in testing differential item functioning with nonnormal trait distributions. *Structural Equation Modeling: A Multidisciplinary Journal*, 22(4), 568-580. <https://doi.org/10.1080/10705511.2014.937669>
- Watkins, M. W. (2018). Exploratory factor analysis: A guide to best practices. *Journal of Black Psychology*, 44(3), 219-246. <https://doi.org/10.1177/0095798418771807>

Το παρόν σύγγραμμα αποτελεί μια σύντομη περιγραφή των βασικότερων διαδικασιών στατιστικής ανάλυσης ποσοτικών και ποιοτικών δεδομένων. Σχεδιάστηκε κυρίως για φοιτητές ή και εκπαιδευτικούς που παρακολουθούν μαθήματα μεθοδολογίας εκπαιδευτικής έρευνας και για όλους όσους επιθυμούν να σχεδιάσουν και να υλοποιήσουν τη δική τους έρευνα. Η παρουσίαση των στατιστικών διαδικασιών ανάλυσης δεδομένων γίνεται παράλληλα με την παρουσίαση του περιβάλλοντος επεξεργασίας και στατιστικής ανάλυσης δεδομένων SPSS. Σε κάθε περίπτωση, γίνεται αρχικά μια εισαγωγή του θεωρητικού μέρους των στατιστικών διαδικασιών και στη συνέχεια, ακολουθώντας συγκεκριμένα παραδείγματα κυρίως της εκπαιδευτικής έρευνας, παρουσιάζονται στο περιβάλλον του SPSS λεπτομερώς η ανάλυση των δεδομένων, η ανάγνωση και η ερμηνεία των αποτελεσμάτων, καθώς και η διατύπωση των συμπερασμάτων. Αρχικά γίνεται μια περιγραφή των βασικών στατιστικών εννοιών που χρειάζονται για τις αναλύσεις με τη χρήση του SPSS, καθώς και των κύριων σταδίων της διεξαγωγής μιας ποσοτικής έρευνας. Στη συνέχεια, γίνεται μια σύντομη παρουσίαση των βασικών στατιστικών που περιγράφουν την κατανομή μιας μεταβλητής, παρουσιάζονται οι ερευνητικοί σχεδιασμοί ανεξάρτητων και εξαρτημένων δειγμάτων και οι παραμετρικοί και οι μη παραμετρικοί έλεγχοι υποθέσεων. Επιπρόσθετα, αφενός εξετάζονται το υπόδειγμα της πολλαπλής γραμμικής παλινδρόμησης και η λογαριθμική παλινδρόμηση και αφετέρου γίνεται μια εισαγωγή στη διερευνητική παραγοντική ανάλυση (Factor analysis), καθώς και στα βήματα που απαιτούνται για την ανάλυση αυτή στο περιβάλλον του SPSS. Τέλος, παρουσιάζονται ζητήματα εγκυρότητας και αξιοπιστίας που θα πρέπει να λαμβάνονται υπόψη κατά τη διερεύνηση της παραγοντικής δομής μιας κλίμακας.

Το παρόν σύγγραμμα δημιουργήθηκε στο πλαίσιο του Έργου ΚΑΛΛΙΠΟΣ+	
Χρηματοδότης	Υπουργείο Παιδείας και Θρησκευμάτων, Προγράμματα ΠΔΕ, ΕΠΑ 2020-2025
Φορέας υλοποίησης	ΕΛΚΕ ΕΜΠ
Φορέας λειτουργίας	ΣΕΑΒ/Παράρτημα ΕΜΠ/Μονάδα Εκδόσεων
Διάρκεια 2ης Φάσης	2020-2023
Σκοπός	Η δημιουργία ακαδημαϊκών ψηφιακών συγγραμμάτων ανοικτής πρόσβασης (περισσότερων από 700) • Προπτυχιακών και μεταπτυχιακών εγχειριδίων • Μονογραφιών • Μεταφράσεων ανοικτών textbooks • Βιβλιογραφικών Οδηγών
Επιστημονικά Υπεύθυνος	Νικόλαος Μήτρου, Καθηγητής ΣΗΜΜΥ ΕΜΠ
ISBN: 978-618-228-204-5	DOI: http://dx.doi.org/10.57713/kallipos-439

Το παρόν σύγγραμμα χρηματοδοτήθηκε από το Πρόγραμμα Δημοσίων Επενδύσεων του Υπουργείου Παιδείας.