

UNIVERSITETI POLITEKNIK I TIRANES

METODA STATISTIKORE DHE ZBATIMET E TYRE

"Ndërtimi dhe analiza e një modeli probabilitaro - statistikor për studimin e efektit të ndotjes në gjendjen shëndetësore të banorëve në zona industriale"

M.Sc. Etleva Beliu (Llagami)

Udhëheqës Shkencor: Prof. Dr. Shpëtim Leka

Tirane 2015



REPUBLIKA E SHQIPERISE
UNIVERSITETI POLITEKNIK – TIRANE
FAKULTETI INXHINIERISE MATEMATIKE & INXHINIERISE FIZIKE
Departamenti i Inxhinierise matematike

DISERTACION
i paraqitur nga
M.Sc. Etleva Beliu (Llagami)
për marrjen e gradës shkencore
"DOKTOR"
në Inxhinieri Matematike

Tema: Metoda statistikore dhe zbatimet e tyre

"Ndërtimi dhe analiza e një modeli probabilitaro - statistikor për studimin e efektit të ndotjes në gjendjen shëndetësore të banorëve në zona industriale"

Udhëheqës Shkencor: Prof. Dr. Shpëtim Leka

Mbrohet më datë _____ para jurisë

1. _____ Kryetar

2. _____ Anëtar

3. _____ Anëtar

4. _____ Anëtar

5. _____ Anëtar

Abstrakti

Në këtë punim synohet që, me paraqitjen e metodave të ndryshme statistikore, analizën e detajuar të menyres së realizimit të tyre, t'i tregohet kujdes kushteve që ato kërkojnë për zbatim, dhe për më tepër, analizës së vlefshmërisë së tyre, apo si të thuash vlerësimin që modeli i ndërtuar në bazë të të dhënave të zgjedhjes, të mund të përdoret për tërë popullimin.

Kështu, gjatë studimit, synohet që krahas të përbashkëtave të metodave statistikore, të vërehen të vecantat e tyre, duke bërë të mundur zgjedhjen e metodës më të përshtatshme në një situatë konkrete. Të gjitha këto metoda janë përdorur për ndërtimin dhe analizën e një modeli probabilitaro- statistikor për studimin e efektit të ndotjes në gjendjen shëndetësore të banorëve në zona industriale.

Është paraqitur e përdorur metoda e njohur ANOVA, por jo vetëm si një metodë e krahasimit të mesatareve, por janë interpretuar vlerësimet intervalore me të cilat ajo është shoqëruar, dhe përdorimi i saj i gjere në përfundimet e analizave të tërë metodave të tjera. Metoda e Regresit Linear te Shumëfishtë është parë në tërë llojshmeritë e tij, jo vetëm të regresit tëzakonshëm, por edhe të atij Hirarkik dhe Stepwise, si metoda që i japin përgjigje pyetjeve të ndryshme që lindin gjatë modelimeve statistikore. Regresi Logjistik, si metoda më e mirë e analizës së një modeli te një ndryshoreje rasti kategorike, është përdorur, e madje është ndërtuar edhe modeli probabilitaro-statistikor për gjendjen e banorëve të zonës në pneumologji. MANOVA, si metodë e krahasimit të mesatareve të komponentëve të një vektori është aplikuar, e përfundimet e saj janë krahasuar me rezultatet me të mira te metodës Analizës se Komponentëve Kryesorë, në interpretimin të gjendjes së gjendjes së përgjithshme të banorëve të zonës.

Gershetimi i terësisë së këtyre metodave ndihmoi edhe në analizen komplekse të e problemeve banoreve në hematologji.

HYRJE	1
1 ANALIZA E VARIANCËS NJË PËRMASORE	2
1.1 TEORIA BAZË E ANOVA-S	2
1.2 MODELI NË NËNPOPULLIME TË PËRCAKTUARA	2
1.3 MODELI NË NËNPOPULLIME TË RASTËSISHME	5
1.4 KRITERE DHE VËREJTJE NË APLIKIMIN E ANOVA-S	6
1.5 TESTI LEVENE (METODA E HOMOGENITETIT TË VARIANCAVE)	7
1.5.1 VËSHTRIM I PËRGJITHSHËM I TESTIT	7
1.5.2 PËRCAKTIMI I TESTIT	7
1.6 KRAHASIMI I SHUMËFISHTË	8
1.6.1 NJË TRAJTIM I PËRGJITHSHËM I PROBLEMIT	8
1.6.2 METODA TUKEY	9
1.6.3 METODA SCHEFFE	12
1.6.4 KRAHASIMI I DY METODAVE	12
1.6.4.1 Zbatim 1	13
2 MODELI I REGRESIT LINEAR TË SHUMËFISHTË	15
2.1 QËLLIMI I REGRESIT TË SHUMËFISHTË	15
2.2 TIPET E NDRYSHME TË REGRESIT TË SHUMËFISHTË	15
2.3 TRAJTIMI MATEMATIK I REGRESIT LINEAR TË ZAKONSHËM, STANDART (GLM)	15
2.3.1 VLERËSIMI I PARAMETRAVE TË REGRESIT LINEAR TË SHUMËFISHTË (HAPI I)	15
2.3.2 INTERPRETIMI GJEOMETRIK I ZGJIDHJES ME METODËN E KATRORËVE MË TË VEGJËL	17
2.3.3 VLERËSIMI INTERVALOR I PARAMETRAVE TË REGRESIT (HAPI II)	19
2.3.4 KONTROLI I HIPOTEZAVE (HAPI III)	19
2.3.5 VËREJTJE	20
2.3.6 PËRFUNDIME	21
2.4 ZBATIM 1	21
2.4.1 VËREJTJE	23
2.5 KOEFICIENTI I KORRELACIONIT TË SHUMËFISHTË	24
2.5.1 PËRFUNDIME	25
2.5.2 ZBATIM 1 (VAZHDIM)	26
2.6 KOEFICIENTI I KORELACIONIT TË PJESSHËM	27
2.6.1 SHTRIMI I PROBLEMIT	27
2.6.1.1 Përfundime	27
2.6.2 VLERËSIMI DHE TESTIMI I HIPOTEZAVE PËR KOEFICIENTËT E KORELACIONIT TË SHUMËFISHTË DHE TË PJESSHËM	28
2.6.2.1 Vërejtje	30
2.7 REGRESI I SHUMËFISHTË STEPWISE	31
2.7.1 SHTRIMI MATEMATIK I REGRESIT TË SHUMËFISHTË STEPWISE	31
2.7.2 PROCEDURAT ITERATIVE	32
2.7.2.1 Përfundime	35
2.7.3 RREGULLAT E NDËRPRERJES	35
2.8 ZBATIM 2	37
2.9 REGRESI I SHUMËFISHTË HIRARKIK	38

Përmbajtja

2.9.1	ZBATIM 1(VAZHDIM)	39
2.9.2	ZBATIM 2:(VAZHDIM)	40
2.9.3	ZBATIM 3 (VAZHDIM)	42
3	ANALIZA E VARIANCËS DYPËRMASORE (METODA ANOVA DYPËRMASORE)	44
3.1	MODELI ME NENPOPULLIME TE FIKSUARA (MODEL I)	44
3.1.1	PËRFUNDIME	47
3.1.2	ZBATIM 1	47
3.2	MODELI ME NËNPOPULLIME TË RASTESISHME (MODEL II)	48
3.2.1	ZBATIM 2	50
3.2.2	VËREJTJE	51
3.3	SHTRIMI MATEMATIK I ANALIZËS SË KOVARIANCËS (ANCOVA)	52
4	ANALIZA E VARIANCËS SHUMËPËRMASORE MANOVA	55
4.1	ANALIZA E VARIANCES SHUMËPËRMASORE (MANOVA)	55
4.1.1	ZBATIMI 1	56
4.1.2	VLERËSIMI I PARAMETRAVE	57
4.1.3	TESTI I HIPOTEZËS SE LINEARITETIT	58
4.1.4	ZBATIMI 1 (VAZHDIM)	58
4.1.5	TESTI I DIFERENCES MES VLERËS SE MESATAREVE MIDIS DISA POPULLIMEVE:MANOVA	58
4.1.5.1	Vërejtje	59
4.1.5.2	Zbatimi i 4.1 (vazhdim)	60
4.2	ANALIZA E KOMPONENTËVE KRYESORË	61
4.2.1	VËREJTJE	63
4.2.2	ZBATIMI 2(VAZHDIM)	64
4.3	ANALIZA FAKTORIALE	65
4.3.1	ZBATIMI 3(VAZHDIM)	66
5	MODELI I REGRESIT LOGJISTIK TË SHUMËFISHTË	68
5.1	REGRESI LOGJISTIK DHE LLOJET E TIJ	68
5.1.1	MODELI I REGRESIT LOGJISTIK TË SHUMËFISHTË	68
5.1.2	MODELI I REGRESIT LOGJISTIK I THJESHTË ME NJË VARIABËL NË MODELIM	69
5.1.2.1	Përfundime	70
5.1.3	MODELI LOGJISTIK I THJESHTË ME PARAMETRA KATEGORIKË	71
5.1.3.1	Përfundime	71
5.1.3.2	Zbatim 1	72
5.1.4	REGRESIONI LOGJISTIK MULTINOMINAL	73
5.1.4.1	Përfundimisht:	73
5.2	VLERËSIMI I PARAMETRAVE	73
5.2.1	PROBLEMI I VLERËSIMIT INTERVALOR TË PARAMETRAVE:	76
5.2.2	ZBATIM 1 (VAZHDIM)	76
5.3	PËRPUTSHMËRIA E MODELIT LOGJISTIK	77
5.3.1	VËREJTJE	78
5.3.2	ZBATIM 1 (VAZHDIM)	79
5.4	MODELIMI I VLERAVE TË PARAMETRIT TË PARASHIKUAR	79

Përmbajtja

5.4.1	ZBATIM 1 (VAZHDIM)	80
5.5	FUNKSIONIT LIDHËS 'PROBABIT'	83
5.5.1	VËREJTJE:	83

6 NDËRTIMI E ANALIZA E NJË MODELI PROBABILITARO-STATISTIKOR PËR STUDIMIN E EFEKTIT TË NDOTJES

84

6.1	NDIKIMI I NIVELIT TE NDOTJES NE PROBLEMET SHENDETËSORE TË PËRGJITHSHME, DERMATOLOGJIKE DHE ATO PNEUMOLOGJIKE. (PROBLEMI 1)	86
6.2	NDËRTIMI I MODELIT STATISTIKORE PËR PROBLEMEVE PNEUMOLOGJIKE TË BANORËVE TË ZONËS INDUSTRIALE (PROBLEMI 2)	100
6.2.1	MODELI I ANALIZËS MBI KOMPONENTËT E ANALIZËS FAKTORIALE	100
6.2.2	PËRFUNDIME	103
6.3	ANALIZA STATISTIKORE E PROBLEMEVE HEMATOLOGJIKE TË BANORËVE TË ZONËS (PROBLEMI 3)	105
6.4	ANALIZA STATISTIKORE E PROBLEMEVE HEMATOLOGJIKE TË BANORËVE TË ZONËS (PROBLEMI 4)	115
6.4.1	BAZA E TË DHËNAVE	115
6.4.2	ANALIZA STATISTIKORE E PROBLEMIT	115
6.5	PËRFUNDIME	121

7 SHTOJCË - ANALIZA STATISTIKORE E TË DHËNAVE PËR PËRCAKTIMIN E FAKTORËVE QË NDIKOJNË NË MEMORJE

122

7.1	DISA FJALË PËR STUDIMIN	122
7.2	<i>CFARË E SHKAKTON KËTË RRIJTJE NË MEMORJE?</i>	123
7.3	METODA STATISTIKORE E NDJEKUR E INTERPRETIMI I REZULTATEVE	124
7.4	PËRFUNDIME	132

REFERENCAT

133

Tabela 1–1 Afishimet e Testit Levene	13
Tabela 1–2 Rezultatet e Metodës së Kontrastit	14
Tabela 2–1 Tabela e koeficientëve të modelit të regresit	22
Tabela 2–2 Afishimet e t-testeve dhe koeficienteve të regresit	23
Tabela 2–3 Tabela e vlefshmerisë së modelit për 'frutat'	26
Tabela 2–4 Tabela e vlefshmerisë së modelit për 'misri '	26
Tabela 2–5 Tabela e koeficientëve të korelacionit	38
Tabela 2–6 Afishimet e t-testeve dhe koeficienteve të regresit	39
Tabela 2–7 Tabela e regresit hirarkik	41
Tabela 2–8 Koeficientët e korelacionit të regresit hirarkik	41
Tabela 2–9 Koeficientët e korelacionit të pjesshëm	42
Tabela 2–10 Regresi stepwise	42
Tabela 3–1 Përbërësit e variacionit Modeli 1	46
Tabela 3–2 Testet sipas hipotezave	46
Tabela 3–3 Tabela e të dhënave deskriptive dhe analizës ANOVA dypërmasore	48
Tabela 3–4 Përbërësit e variacionit Modeli II	49
Tabela 3–5 Afishimet e përbërësve të variacionit	51
Tabela 3–6 ANOVA e modelit të qendëruar	54
Tabela 3–7 Përbërësit e analizës së kovariancës mes grupeve	54
Tabela 4–1 Modeli i Matrices sipas GLM dhe MANOVA-s	57
Tabela 4–2 ANOVA për 13 variablat e varur	60
Tabela 4–3 Vlerat e veta të komponentëve kryesorë	64
Tabela 4–4 Tabela e vlerave të veta	67
Tabela 5–1 Vlerësimi para modelimit	73
Tabela 5–2 Koeficientët e variablave në ekuacionin e modelimit	77
Tabela 5–3 Kontrolli i vlefshmërisë së modelit	79
Tabela 5–4 Informacione për modelimin	79
Tabela 5–5 Të dhënat e modelimit të regresit logjistik	82
Tabela 6–1 Matrica e kovariancës	86
Tabela 6–2 Kovarianca mes variablave	87
Tabela 6–3 Testi i Multivariancës	89
Tabela 6–4 ANOVA për 13 variablat e varur	90
Tabela 6–5 t testet për secilin prej 13 variablave të varur	91
Tabela 6–6 Analiza e ndikimit të variablave të pavarur tek variablat e varur në shqyrtim	94
Tabela 6–7 Vlerat e veta	94
Tabela 6–8 Koeficientët e variablave kanonik	94
Tabela 6–9 Analiza deskriptive e <i>fshatrave</i> në ' <i>superndikim</i> '	95
Tabela 6–10 Efekti në ' <i>superndikim</i> ' i variablave të pavarur	95
Tabela 6–11 Analiza e Komponentëve Kryesorë	97
Tabela 6–12 Afishimet për të tri faktorët	98
Tabela 6–13 Anova për komponentin pneumologjik	99
Tabela 6–14 Anova për komponentin e përgjithshëm	99
Tabela 6–15 Anova për komponentin e dermatologjik	99
Tabela 6–16 Afishimet e regresit linear	101
Tabela 6–17 Rregulli i vlerësimit të regresit të shumfishtë linear	102
Tabela 6–18 Koeficientët e modelit të regresit të shumfishtë	102
Tabela 6–19 Koeficientët e variablave në ekuacionin e modelimit	104
Tabela 6–20 Tabela e klasifikimit pas modelimit	105
Tabela 6–21 Rezultatet e analizës deskriptive	106
Tabela 6–22 Rezultatet e krahasimit të paramerave kryesor në gjak	106
Tabela 6–23 Kontrolli hipotezës për kufirin minimal	107
Tabela 6–24 Kontrolli hipotezës për minimumin e vleras normale	108
Tabela 6–25 Anova dypërmasore	108

Tabela

Tabela 6–26 Analiza deskriptive	109
Tabela 6–27 Anova njëpërmasore per komponentet e tjera të gjakut	109
Tabela 6–28 Regresi logjistik për indikatorët rasteve problematike	111
Tabela 6–29 Analiza dy përmasore	113
Tabela 6–30 Kontrolli i efektit nderthurës tek ezinofailet	114
Tabela 6–31 Analiza deskriptive e bimëve	117
Tabela 6–32 Rezultatet e testit Levin dhe te t-testit për mesataret	117
Tabela 6–33 Rezultatet e analizës së kontrastit	118
Tabela 6–34 Rezultatet e regresit hirarkik dhe regresit të shumëfishtë të zakonshëm	120
Tabela 7–1 Analiza Deskriptive e nrword	125
Tabela 7–2 Testi i Normalitetit të nrword	125
Tabela 7–3 Analiza Deskriptive dhe testi i normalitetit të sqrt nrword	126
Tabela 7–4 Rezultatet e stepwise regression	128
Tabela 7–5 Rezultatet e ANOVA dy përmasore	130
Tabela 7–6 Rezultatet e krahasimit të sqtrnrfaqëve sipas grupeve të nivelit të sportit	132

Figura 1–1 Vlerat Kritike – numri i mesatareve qe krahasen	11
Figura 2–1 Paraqitja gjeometrike	18
Figura 4–1 Interpretimi gjeometrik i metodës së komponentëve kryesor në rastit n=2	63
Figura 5–1 Treguesi 'kollë' sipas zonave	72
Figura 5–2 Grafiku i predicted_logit_probability	81
Figura 6–1 Pozicioni i zonë së studimit	85
Figura 6–2 Pozicioni i puseve tnaftës	85
Figura 6–3 Plotimi I vlerave të veta	97
Figura 6–4 Harta e bimësisë në zonën në studim	116
Figura 6–5 Harta e rrjetit hidrologjik	116
Figura 7–1 Histograma e numrit të fjalëve të memorizuara	124
Figura 7–2 Grafiku i ndryshimit të vlerave të <i>nrfjalëve</i> nga vlerat normale	126
Figura 7–3 Grafiku i ndryshimit të vlerave të <i>sqtrnrfjalëve</i> nga vlerat normale	128
Figura 7–4 Grafiku i mesatareve të <i>sqtrnrfjalëve</i> sipas <i>gjinise</i> dhe <i>grupit te nivelit</i> te shtrimeve fizike	130

HYRJE

Analiza matematikore e modeleve statistikore, të kuptuarit në thellësi të të gjithë specifikave të tyre, njohja e hipotezave me të cilat këto metoda shoqërohen, statistikave që ato përdorin për vlerësim, vecantive të tyre, është hapi i parë në procesin e ndërtimit të modeleve probabilitare-statistikore. Pa dyshim zhvillimi i teknologjisë kompjuterike dhe rritja e fuqisë dhe shpejtësisë së tyre përpunuese, ka rritur efektet e këtyre metodave në mënyrë dramatike

Këto metoda janë përdorur për ndërtimin e një modeli matematiko-statistik që analizon gjendjen shëndetësore të banorëve në një zonë industriale .

Konkretisht të dhënat janë marrë nga Instiruti i Shendetit Publik, dhe përfshijnë informacione përdorur gjendjen e përgjithshme shëndetësore, për problemet në pneumologji, të dhëna të fushes hematologjike , dermatologjike, etj.

Për ndërtimin e një modeli të saktë statistikor , u shfrytëzuan edhe të dhëna të tjera specifike për problemet e trajtuara.

1 Analiza e variancës një përmasore

1.1 Teoria bazë e ANOVA-s

ANOVA, (shkurtim i fjalëve Analizë e Variancës), është **iniciuar** rreth viteve 1920 nga statistikani anglez Sir Ronald Aylmer **Fisher**. Metoda synon të testojë hipotezën që disa variabla kanë të njëjtën mesatare në një apo më shumë popullime, apo jo.

Fisher bëri një varg arsyetimesh mbi variancën dhe arriti në konkluzione të rëndësishme. Bazë e këtyre arsyetimeve ishte ndarja e variancës së zgjedhjes të një popullimi, në dy komponente dhe, krahasimi i këtyre komponenteve **duke përdorur statistikën Fisher**. Këto janë dhe arsytet pse ANOVA njihet ndryshe me emrin ANOVA e Fisherit ose analiza e Fisherit.

Pa dyshim zhvillimi i teknologjisë kompjuterike dhe rritja e fuqisë dhe shpejtësisë së tyre përpunuese, bënë që efektet e kësaj metode statistikore të jenë dramatike.

Le të shqyrtojmë këto metodë statistikore.

Janë dhënë I nënpopullime të një popullimi. Shënohet μ_i mesatarja e nënpopullimit të i -të për $i=1,2,\dots,I$. Duke marrë zgjedhje të rastit të i nënpopullimeve, vlerësohen pikërisht këto mesatare. Pastaj, kërkohet të kontrollohet hipoteza:

$$H_0: \mu_1 = \mu_2 = \dots = \mu_I = \mu,$$

sipas së cilës, të gjitha nënpopullimet e popullimit kanë pritje matematike të njëjtë.

Kjo hipotezë ka *hipotezë alternative*, ka të paktën dy nënpopullime me pritje matematike të ndryshme.

Për të kontrolluar H_0 , supozohet se jemi duke shqyrtuar nënpopullime normale që kanë variancë të njëjtë, pra që nënpopullimet kanë ligj $N(\mu_i, \sigma^2)$

Hipotezës së mësipërme mund ti jepet forma :

$$H_0: \alpha_1 = \alpha_2 = \dots = \alpha_I = 0,$$

(efektet e nënpopullimeve janë 0)

Dihet se për popullimin në shqyrtim, $MS_T = \sum_{i=1}^I \sum_{j=1}^{J_i} (y_{ij} - \bar{y})^2 / (n-1)$ (1) është vlerësim

i pazhvendosur për σ^2 . Pikërisht vlerësimin e mësipërm të variancës, Fisher e ndan në dy komponente.

1.2 Modeli në nënpopullime të përcaktuara

Në modelin e nënpopullimeve të përcaktuara, shqyrtohen I nënpopullime të popullimit, të fiksuara nga eksperimentuesi. Në secilin prej nënpopullimeve të fiksuara merret një zgjedhje e rastit me vëllim J_i . Vlera e j -te e zgjedhjes së nënpopullimit të i -të shënohet y_{ij} .

Modeli ndertohet si me poshte:

Së pari: Shënohet me μ_i mesatarja e nënpopullimit të i -të për i nënpopullimet e popullimit, (nënpopullime këto të fiksuara nga eksperimentuesi). Ndaj, me supozimin se nënpopullimi i i -të ka shpërndarje $N(\mu_i, \sigma^2)$, cdo y_{ij} vlere e zgjedhjes te marre, mund të paraqitet ne formën: $y_{ij} = \mu_i + e_{ij}$, ku e_{ij} kanë shpërndarje $N(0, \sigma^2)$.

Së dyti: μ_i mund të shkruhen si shumë e μ (pritje matematike të popullimit) me α_i që tregon efektin e nënpopullimit të i -të. Përfundimisht vlerat e marra të zgjedhjes shkruhen në formën: $y_{ij} = \mu + \alpha_i + e_{ij}$

Ky model linear i fituar quhet *modeli linear i ANOVA-s* e, meqë nënpopullimet janë paracaktuar nga eksperimentuesi, justifikohet edhe emërtimi i këtij modeli, si **model në nënpopullime të përcaktuara**^[1].

Pikërisht, ky model linear i ANOVA-s përdoret për vlerësimin pikesor te parametrave μ dhe α_i . **Metoda e Katroreve më të Vegjël**, siguron në bazë të teoremës së Gaus –Markovit, një vlerësim pikesor të pazhvendosur dhe me dispersion minimal të parametrave μ dhe α_i , kur ato fitohen si kombinim linear i vlerave të vrojtuar y_{ij} .

Pra, ndërtohet $S(\mu, \alpha_i) = \sum_{i=1}^I \sum_{j=1}^{J_i} (y_{ij} - \mu - \alpha_i)^2$, dhe kërkohet të gjenden vlerat μ dhe α_i që minimizojnë këtë shumë.

$$\frac{\partial S}{\partial \mu} = - \sum_{i=1}^I \sum_{j=1}^{J_i} (y_{ij} - \mu - \alpha_i) = 0$$

$$\sum_{i=1}^I \sum_{j=1}^{J_i} y_{ij} = n\mu + \sum_i \sum_j \alpha_i = n\mu + \sum_i J_i \alpha_i$$

Atehere vlerësimi për μ (me supozimin se $\sum_i J_i \alpha_i = 0$) është

$$\hat{\mu} = \frac{\sum_i \sum_j y_{ij}}{n}$$

Ndërsa për të vlerësuar α_i , llogarisim derivatet e pjesshme sipas parametrat α_i

$$\frac{\partial S}{\partial \alpha_i} = - \sum_j (y_{ij} - \mu - \alpha_i) = 0$$

$$\sum_j y_{ij} = J_i \mu + J_i \alpha_i$$

$$\hat{\alpha}_i = \frac{\sum_j y_{ij}}{J_i} - \mu = \bar{y}_i - \bar{y}$$

Pra vlerësimet që merren me MKV janë

$$\hat{\mu} = \bar{y} \quad (*)$$

$$\hat{\alpha}_i = \bar{y}_i - \bar{y}$$

Kështu $\hat{\mu} + \hat{\alpha}_i = \bar{y} + (\bar{y}_i - \bar{y}) = \bar{y}_i$

Por dihet nga teorema Gaus –Markovit se, jo vetëm vlerësimet që merren për parametrat janë të pazhvendosur, por dhe se këto vlerësime sigurojnë një minimum të

variancës nga të gjithë vlerësuesit e pazhvendosur të μ dhe α_i , që fitohen si kombinim linear i vlerave të vrojtuar.

Kështu, shuma e katrorëve të gabimeve

$$SSB = \sum_{i=1}^I \sum_{j=1}^{J_i} (y_{ij} - \hat{\mu} - \hat{\alpha}_i)^2 = \sum_{i=1}^I \sum_{j=1}^{J_i} (y_{ij} - \bar{y}_{i.})^2$$

duke u pjesuar me numrin e gradëve të lirisë të saj, siguron një vlerësim të pazhvendosur minimal të σ^2 .

Një vlerësim i pazhvendosur për σ^2 , është dhe mesatarja e variancave të grupeve të zgjedhjes. Duke marrë

$$MS_B = \frac{\sum_i s_i^2}{I} = \sum_{i=1}^I \sum_{j=1}^{J_i} (y_{ij} - \bar{y}_{i.})^2 / I(J_i - 1) = \sum_{i=1}^I \sum_{j=1}^{J_i} (y_{ij} - \bar{y}_{i.})^2 / (n - I) \quad (2)$$

nëse hipoteza H_0 është e vërtetë dhe, duke ditur për më tepër s_i^2 -të, kanë shpërndarje $\chi^2(J_i-1)$, MS_B do të ketë shpërndarje $\chi^2(I(J_i-1))$, pra $\chi^2(n-I)$.

Tani, shuma e dyfishtë (1) transformohet në shumën (2) që, nëse hipoteza është e vërtetë, siguron një tjetër vlerësim të pazhvendosur për σ^2 .

Duke u nisur nga barazimi $y_{ij} - \bar{y} = (y_{ij} - \bar{y}_{i.}) + (\bar{y}_{i.} - \bar{y})$ dhe pasi ngrihen në katrorë dy anët e tij, shumëhet sipas i -ve dhe j -ve dhe merret

$$SST = \sum_{i=1}^I \sum_{j=1}^{J_i} (y_{ij} - \bar{y})^2 = \sum_{i=1}^I \sum_{j=1}^{J_i} (y_{ij} - \bar{y}_{i.})^2 + \sum_{i=1}^I \sum_{j=1}^{J_i} (\bar{y}_{i.} - \bar{y})^2 + 2 \sum_{i=1}^I \sum_{j=1}^{J_i} (y_{ij} - \bar{y}_{i.})(\bar{y}_{i.} - \bar{y})$$

Meqënëse termi i fundit i kësaj shume mund të transformohet në trajtën (2)

$\sum_{i=1}^I (\bar{y}_{i.} - \bar{y}) \sum_{j=1}^{J_i} (y_{ij} - \bar{y}_{i.})$, duket qartë që ky praktikisht është 0. Ndaj, shuma merr formën

$$\sum_{i=1}^I \sum_{j=1}^{J_i} (y_{ij} - \bar{y}_{i.})^2 + \sum_{i=1}^I J_i (\bar{y}_{i.} - \bar{y})^2 = SS_B + SS_M$$

Duke u nisur nga barazimi $SS_T = SS_B + SS_M$ dhe nga përfundimet (1) dhe (2), në bazë të Teoremës Koçran, SS_M dhe SS_B janë të pavarur dhe janë vlerësime për σ^2 , madje $SS_M/(I-1)$ ka shpërndarje $\chi^2(I-1)$.

Vërehet se SS_M është shuma e katrorëve të gabimeve të mesatareve të grupeve me mesataren e tërë zgjedhjes.

Kështu, MS_B / MS_M ka shpërndarje $F(I-1, n-I)$, si raport i dy statistikave të pavarura me shpërndarje χ^2 , raport ky, bazë në analizën ANOVA.

Në programet kompjuterike, kur japim komandën për vlerësim me metodën ANOVA, afishohen:

vlera g_0 e testit të mësipërm, pra vlera e raportit Fisher për zgjedhjen

p- vlera koresponduese e saj sipas tabelës së Fisherit (pra llogaritet probabiliteti që testi statistik të marrë vlerën g_0 ose vlera më ekstreme se g_0 , kur H_0 është e vërtetë).

Gjeometrikisht, p-vlera, paraqet dyfishin e sipërfaqes në të djathtë të vlerës së vrojtuar, nën grafikun $F(I-1, n-I)$. Kështu mund të përcaktojmë jo vetëm se hipoteza

është apo jo e vërtetë, por dhe nivelin e gabimit të llojit të parë me të cilin ne mund ta pranojmë këtë hipotezë nëse është e mundur.

1.3 Modeli në nënpopullime të rastësishme

Tabela e ANOVA-s shfrytëzohet edhe në një këndvështrim tjetër.

Për këtë, më parë ndërtohet një tjetër model, që quhet **modeli në nënpopullime të rastësishme**^[3].

Shihet popullimi $N(\mu, \sigma^2)$, dhe, lidhur me të, edhe tërë nënpopullimet e këtij popullimi të përcaktuar sipas një faktori. Me anë të zgjedhjes së marrë, ndërtohet një zgjedhje për popullimin e nënpopullimeve.

Për këtë, llogariten më parë mesataret m_i për cdo zgjedhje nga grupet e marrë sipas një faktori, dhe mendohen nënpopullime me shpërndarje $N(m_i, \sigma^2)$, për $i = 1, 2, \dots, I$, që tani shërbejnë si zgjedhje me vëllim 1 për popullimin e nënpopullimeve të popullimit të madh.

Me σ_a shënohet varianca e vlerave të kësaj zgjedhje nga μ .

Kështu, në bazë të modelit linear, elementet e kësaj zgjedhje mund të shkruhen në formën:

$m_i = \mu + a_i$, ku a_i e pavarura kanë shpërndarje normale me mesatare 0 dhe variancë σ_a^2 (10).

Përfundimisht, vlerave të zgjedhjes fillestare të popullimit i jipet një paraqitje tjetër në formën :

$y_{ij} = m_i + e_{ij} = \mu + a_i + e_{ij}$, ku a_i e pavarura kanë shpërndarje normale me mesatare 0 dhe variancë σ_a^2 , ndërsa e_{ij} , gjithashtu të pavarura, kanë shpërndarje $N(0, \sigma^2)$.

Në këtë model të ndërtuar, interesohemi për σ_a^2 që paraqet variacionin e pritjes matematike të nënpopullimeve të përcaktuara sipas një faktori, nga pritja matematike e popullimit μ .

Shtrohet hipoteza

$$H_0: \sigma_a^2 = 0,$$

(domethënë që nuk ka variacion mes nënpopullimeve sipas faktorit në studim), me hipotezë alternative

$$H_1: \sigma_a^2 \neq 0.$$

Duke përdorur Metodën e Katrorëve më të Vegjël, merret përsëri i njëjti vlerësim për pritjen matematike të popullimit, si në rastin e modelit në nënpopullime të përcaktuara,

$$\hat{\mu} = \sum \frac{m_i}{I} = \bar{y}$$

Nga teorema Kocran merret dhe një vlerësim i pazhvendosur për σ_a^2 , që jipet nga shuma e mëposhtme $\sum_i (m_i - \hat{\mu})^2 = \sum_i (m_i - \bar{y})^2$, pjesuar me $I-1$.

Por, dihet se m_i fitohen nga y_{ij} për i të fiksuar dhe $j=1, \dots, J$ sipas barazimit $y_{ij} = m_i + e_{ij}$, ku e_{ij} , kanë shpërndarje $N(0, \sigma^2)$.

Ndaj $SS_M / (I-1) = \sum_{i=1}^I J(\bar{y}_i - \bar{y})^2 / (I-1)$ është një mesatare e vlerësimeve të variacionit

të grupeve, pra përmban edhe variacionin brenda cdo grupi σ^2 dhe J herë variacionin nga grupi në grup, σ_a^2 , (me supozimin që vëllimi i zgjedhjeve në cdo grup është marrë i njëjtë), ndërkohë që dihet se vlera e kësaj statistike afishohet në tabelën ANOVA, dhe për më tepër se statistika ka shpërndarje $\chi^2(I-1)$

Duke bërë të njëjtën ndarje të katrorit të gabime të popullimit, si në rastin e modelit të parë, kemi :

$$SST = \sum_{i=1}^I \sum_{j=1}^{J_i} (y_{ij} - \bar{y}_i)^2 + \sum_{i=1}^I J_i (\bar{y}_i - \bar{y})^2 = SS_B + SS_M$$

$SS_B / (n-I)$ është një vlerësim i pazhvendosur për variacionin brenda grupeve që është σ^2 , me shpërndarje të njohur $\chi^2(n-I)$ (mund të shihet si mesatare e variacionit brenda cdo grupi, që e shënuam σ^2).

Për të vlerësuar σ_a^2 , bazohemi në vlerësimin:

$\hat{\sigma}_a^2 = (MS_B - MS_M) / J$ (ku J është vëllimi i njëjtë i zgjedhjes i marrë në cdo grup)

Kështu për testimin e hipotezës $H_0: \sigma_a^2 = 0$ do të shfrytëzojmë po tabelën ANOVA dhe rezultatin e afishuar të statistikës së Fisherit.

1.4 Kritere dhe vërejtje në aplikimin e ANOVA-s

1. Metoda ANOVA kërkon, para aplikimit të saj, të kontrollohen nëse plotësohen apo jo dy kushtet e sipërpërmendura, normalizimi i të dhënave si dhe kontrolli i barazimit të variancës mes nënpopullimeve në studim.

Metodat e përdorura për kontrollin e normalizimit të të dhënave është **metoda Kolmogorov-Smirnov-it** dhe **Shapiro-Wilk-it**.

Ndërsa për të kontrolluar kushtin e dytë të barazimit të variancave përdoret **testi Levene-t**, në format e tij të përpunuara (shih Testin Levene, 1.5).

2. Metoda e ANOVA-s përdoret shume vecanërisht kur sipas niveleve të faktorit në shqyrtim, bëhet ndarja e popullimit të madh në disa (më shumë se dy) nënpopullime.

Nëse ndarja do të ishte në dy nënpopullime atëherë trajtimi i problemit si me metodën e ANOVA-s do të jipte të njëjtin rezultat sikur ai të trajtohet me ane të t testit.

3. Lind pyetja në këto raste cila nga metodat është më e pershtatshme. Në kushtet e mosplotësimi të Kushtit të Homogjenitetit të Variances mes grupeve, metoda më mire është ajo e t-testit, sepse me modifikimin e bere ajo lejon përdorimin e saj në dy raste dhe kur plotësohet dhe kur nuk plotësohet kushti në fjalë.

Kur Kushti i Homogjenitetit të Variances plotësohet përdoret t-testi Student, ndërsa kur ky kusht nuk plotësohet përdoret t-testi Welch,

Ndërsa në kushtet e më shumë se dy nënpopullimeve, përdorimi i t-testit për cdo dyshe të studiuar s'do të kishte kuptim sepse për cdo t-test vendimi do të merret me një gabim të llojit të parë α , kështu do të rritej shumë gabimi në marrjen e vendimit që të gjitha nënpopullimet janë statistikisht me mesatare të njëjte.

4. F testi për hipotezën $\sigma_a^2 = 0$ në modelin e dytë të metodës në fjalë është një vlerë e përafëruar e testit në fjalë, kur J_i s'janë të barabarata. Megjithatë kjo pranohet si procedurë praktike, derisa testi i sakte do të ishte shumë i komplikuar.

1.5 Testi Levene (Metoda e Homogjenitetit të Variancave)

1.5.1 Vështrim i përgjithshëm i testit

Testi Levene përdoret në statistike për të kontrolluar barazimin apo jo të variancës mes dy ose më shumë grupeve.

Shume teste statistikore e kanë si kusht supozimin se dispersionet e popullimeve nga të cilat janë marrë zgjedhje të ndyshme, janë të barabarta. Procedurat tipike ku ky kusht kërkohet është Analiza e Variancës (ANOVA) dhe t- testi.

Testi i Levene përdoret zakonisht para krahasimit të mesatareve. Duke u nisur nga rezultatet e këtij testi vendoset menyra që do të përdoret për krahasimin e mesatareve, do të jetë ajo parametrike apo jo parametrike, metode kjo e fundit që është e lire nga kushti i barazimit të variancave.

Ai mund të përdoret edhe si test, thjesht për të përcaktuar se dy ose më shumë zgjedhje të marra nga i njëjti popullim kanë apo jo variancë të barabartë.

Testi njihet ndryshe me emrin Testi i Homogjenitetit të Variancës.

Testi kontrollon hipotezën H_0 që variancat janë të barabarta, me hipotezë alternative jo të gjitha variancat (dispersionet e popullimeve) janë të barabarta.

Nëqoftëse p-vlera e rezultatit të këtij testi është e vogël, më e vogël se një nivel rëndësie, (zakonisht 0.05), atëherë diferencat në variancat e zgjedhjeve kanë dhënë rezultate të tilla që nuk ndodhin nëse zgjedhjet e rastit do të merreshin nga popullime me variancë të njëjtë.

Përfundimisht, hipoteza, që variancat janë të barabarta, bie poshtë dhe arrihet në përfundim që ka diferencë mes variancave (dispersioneve) të popullimeve.

1.5.2 Përcaktimi i testit

Testi statistik W përcaktohet si më poshtë:

$$W = \frac{(N - k)}{(k - 1)} \frac{\sum_{i=1}^k N_i (Z_{i.} - Z_{..})^2}{\sum_{i=1}^k \sum_{j=1}^{N_i} (Z_{ij} - Z_{i.})^2},$$

ku W është rezultati i testit

k numri i grupeve që i nënshtrohen testit,

N numri total i rasteve në gjithë grupet,

N_i numri i rasteve në grupin i -te

$Z_{ij} = |Y_{ij} - \bar{Y}_i|$ ku Y_{ij} është vlera e variablit të matur për elementin j -te në grupin i -të ndërsa $\bar{Y}_i$ mesatarja e grupit të i -të.

ose $Z_{ij} = |Y_{ij} - \tilde{Y}_i|$ ku $\tilde{Y}_i$ është kuartili i dytë i grupit

Të dy vlerësimet përdoren.

Sipas vlerësimit të parë, shmangia e vlerave të vrojtura behet duke llogaritur distancën e tyre nga mesatarja e grupit. Mesataret e grupeve konsiderohen si vlerësime pikësore për mesataren e popullimit, të cilit zgjedhja i përket. Vlerësimi në fjalë përdoret kur të dhënat e zgjedhjes plotësojnë kushtin e normalizimit.

Ndërsa në formën e dytë, si bazë në llogaritje merren distancat nga kuartilet e dyta, median-at e vlerave të vrojtura të zgjedhjes të popullimit në studim. Metoda e dytë njihet si **testi Brown-Forsythe**. Ajo përdoret, dhe jep rezultate të mira, në rastet kur të dhënat e zgjedhjes nuk plotësojnë kushtin e normalizimit.

$Z_{..}$ është mesatarja e gjithë Z_{ij} pra

$$Z_{..} = \frac{1}{N} \sum_{i=1}^k \sum_{j=1}^{N_i} Z_{ij}$$

$Z_{i.}$ është mesatarja e grupit të i -te

$$Z_{i.} = \frac{1}{N_i} \sum_{j=1}^{N_i} Z_{ij}$$

Testi i përcaktuar W provohet se ka shpërndarje Fisher me parametra($k-1$, $N-k$).

Ndaj për marrjen e vendimit krahasohet vlera e vrojtuar me vlerën kritike të testit, që përcaktohet nga $F(k-1, N-k)$ për një nivel besimi të zgjedhur (0.05 ose 0.01).

Në përpunimin kompjuterik të të dhënave procesi i marrjes së vendimit bëhet sic është sqaruar në paragrafin e mësipërm.

Testi e Levene konsiderohet robust dhe i fuqishëm. Qëndrueshmëria e testit lidhet me aftësinë e testit për të arritur në konkluzionin që variancat s'janë të barabarta kur në fakt ato s'janë të tilla. Robuste lidhet me faktin që testi arrin në konkluzionin e gabuar që variancat s'jane të barabarta, kur zgjedhja s'ka shpërndarje normale dhe në të vërtetë ato janë të barabarta.

Punimi origjinal i Levene në fakt propozohet vetëm me mesataren. Ishte Brown dhe Fosythe që më 1974 e përmirësuan testin në fjale dhe përdoren medianen në këtë vlerësim.

Me ane të studimeve të kryera me metodën Monte Karlo, u tregua se përdorimi i mesatares vlen nëse të dhënat tregojnë për shpërndarje simetrike dhe plotësojnë kushtin e normalitetit. Ndërsa, kur të dhënat nuk plotësojnë kushtin e normalitetit, psh kur kemi shpërndarje Koshi, rezultat më të mira i siguron përdorimi i medianës.

1.6 Krahasimi i shumëfishtë

1.6.1 Një trajtim i përgjithshëm i problemit

Në problemin e trajtuar lind pyetja: Cili nga çiftet e mesatare të nënpopullimeve është i ndryshëm?

Një variant është të analizojmë të gjithë çiftet e mundshme të nënpopullimeve duke krahasuar pritjet matematike. Në rastin tonë kemi për të kontrolluar disa hipoteza

$H_0: \mu_i = \mu_j$, për $i \neq j$

Lind problemi i studimit të nënpopullimeve që i korespondojnë nivele të ndryshme të këtij faktori.

Për të përcaktuar se cili nga nivelet do të konsiderohet i ndryshëm, duhen dalluar mesataret statistikisht të ndryshme nga njëra-tjetra.

Matematikisht kemi të bëjmë me problemin e **krahasimit të shumëfishtë**, të shumëfishtë në kuptimin se kemi të bëjmë me disa krahasime të njëpasnjëshme. Për çdo krahasim të tillë, përgjigja e studimit të një hipoteze, jepet me një koeficient besimi të lartë. Pra, mundësia që kjo hipotezë të anulohet, ndërsa ajo është e vërtetë, është shumë e vogël. Por, nëse bëhen shumë krahasime të njëpasnjëshme, mundësia që të gjitha këto hipoteza të hidhen poshtë, ndërsa ato janë të vërteta, rritet ndjeshëm.

Shume matematikanë kanë paraqitur teknika të ndryshme për kontrollimin e këtij gabimi, në rastin e krahasimeve të shumëfishta. Teknikat që ato propozojnë, tentojnë një reduktim të gabimit të llojit të parë për secilin krahasim, në mënyrë të tillë që megjithëse gabimi i shumë krahasimeve të rrisë α -ën, ajo të mos kalojë vlerën e dëshiruar.

Por ka dhe metoda që synojnë rritjen e zonës të vlerave të lejuara për diferencën midis dy mesatareve që vlerësohen, në mënyrë që, ndërprerja e këtyre zonave pas të gjitha krahasimeve të nevojshme, të japë një gabim të llojit të parë, sa ai i dëshiruar. Një metodë e tillë është **metoda Tukey**, e përdorur në rastin kur për vlerësimin e nënpopullimeve përdoren zgjedhje nga nënpopullime që plotësojnë kushtin e barazimit të variancës, dhe **Scheffe**, që përdoret në rastin kur kushti në fjalë nuk plotësohet.

1.6.2 Metoda Tukey

Kjo metode ka marrë emrin e matematikanit dhe kimistit amerikan të shekullit XX, John Tukey^[3].

Ajo synon të dallojë mesataret e ndryshme nga njëra-tjetra duke i krahasuar ato dy e nga dy.

Le ta shohim metoden Tukey, duke trajtuar disa çështje të saj, supozimet ku bazohet kjo metode, testin që ndërton për vlerësim, statistikën në të cilën bazohet si dhe rradhën e krahasimit që propozon.

Metoda Tukey ka në bazë të saj *disa supozime*. Së pari, vëzhgimet janë të pavarura dhe, së dyti, mesataret, që studiohen, janë nga popullime me shpërndarje normale. Por kushti që e dallon përdorimin e metodes Tukey nga metodat e tjera është që vëllimi i zgjedhjeve të nënpopullimeve është i njëjtë.

Testi Tukey bazohet në testin e zakonshëm Studenti, ose Gosset, që përdoret për krahasimin e dy mesatareve.

Shembull i thjeshtë:

Duam të kontrollojmë hipotezën $H_0: \mu_1 = \mu_2$

me hipotezë alternative që pritjet matematike të dy popullimeve, janë të ndryshme.

Ky problem është i njëjtë me $H_0: \mu_1 - \mu_2 = 0$, me $H_1: \mu_1 - \mu_2 \neq 0$

Për vlerësimin e kësaj hipoteze sipas metodës Studenti, përdoret statistika $T = \frac{\bar{x} - \bar{y}}{\frac{s}{\sqrt{n}}}$

që ka shpërndarje Studenti me $2n-2$ shkallë lirie (n është vëllimi i zgjedhjeve që do të shqyrtohen dhe, ku në këtë formulë është supozuar të jetë i njëjtë. Në rast se ky vëllim nuk është i njëjtë, atëherë në vend të n në formulë do të punohet me mesataren e dy vëllimeve të zgjedhjeve). Pasi përcaktohet zona e vlerave të lejuara ($-t$, t), duke shfrytëzuar tabelën e shpërndarjes Studenti për α -ën, dhe $2n-2$ shkallë lirie, shihet nëse vlera e testit të marrë bie në zonën kritike. Në këtë rast themi se popullimet kanë mesatare të ndryshme.

Por, përdorimi i këtij testi në tërë krahasimet e cifteve të popullimeve, thamë, se do të sjellë rritjen e gabimit të llojit të parë.

Testi Tukey mundohet të zgjidhë këtë problem, duke modifikuar zonën e vlerave të lejuara të statistikës Studenti, në varësi të numrit të mesatareve që do të kontrollohen. Kuptohet se sa më i madh të jetë ky numër, aq më e madhe duhet marrë vlera kritike e

modifikuar t_T , duke berë që dhe zona e diferencës së lejuar të dy mesatareve të një krahasimi të vetëm të jetë më e madhe.

Për të kuptuar më mirë këtë ndryshim midis zonës së vlerave të lejuar të marra sipas metodës Studenti, dhe asaj Tukey, le ti kthehemi problemit të thjeshtë të mësipërm, problemit të krahasimit vetëm të dy mesatareve.

Per kontrollin e kesaj hipoteze, sipas Studentit ndërtohet statistika:

$$T = \frac{\bar{x} - \bar{y}}{\frac{s}{\sqrt{n}}}$$

Kurse, sipas Tukey, ndërtohet statistika $T_T = \frac{\bar{x} - \bar{y}}{\frac{s\sqrt{2}}{\sqrt{n}}}$, ku 2-shi përcaktohet nga numri i

mesatareve që po krahasohen.

Kuptohet që ky ndryshim në test do të rrisë zonën e vlerave të lejuar nga $(-t, t)$ sipas metodës Studenti, në $(-t\sqrt{2}, t\sqrt{2})$ sipas metodës Tukey.

P.sh, nëqoftëse. në cdo grup kemi marrë një zgjedhje me vëllim 13, ($n=13$), atëhere për dy grupet që krahasohen, do të arrihej në përfundim që kanë pritje matematike të ndryshme, nëse vlera e testit T është në $(-\infty, -2.06)$ ose në $(+2.06, +\infty)$. Por me modifikimin e bërë sipas Tukey, ky test do t'i takonte zonës kritike, pra nënpopullimet do të konsiderohen me pritje matematike të ndryshme, nëse vlera e testit në fjalë është në $(-\infty, -2.06*1.414) = (-\infty, -2.91)$ ose në $(+2.91, +\infty)$. Dhe nga tabela Studenti, kësaj vlere kritike i korespondon $\alpha = 0.025$.

Nga dy krahasimet që janë të nevojshme të bëhen në këtë rast, probabiliteti që të dy nënpopullimet të konsiderohen me pritje matematike të njëjtë, është $(1 - \alpha)^2 = (1 - 0.025)^2 = 0.95$, duke siguruar kështu përfundimisht një përgjigje me nivel besimi 0.05, sa ai i kërkuari.

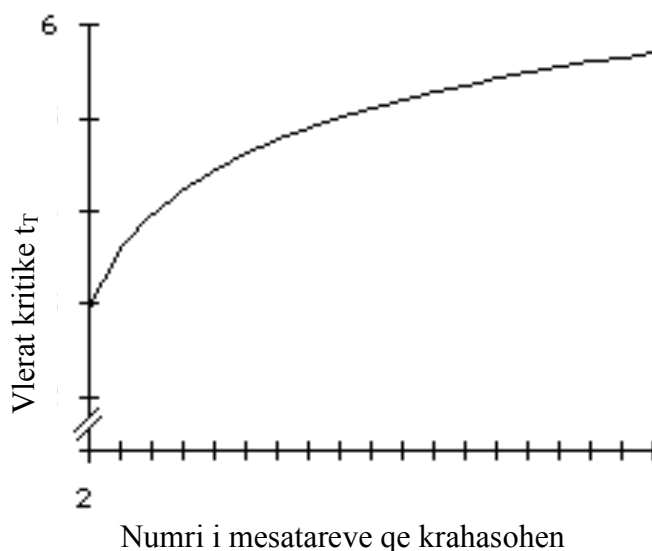
Duhet theksuar se, sipas kësaj metode, në rast se vëllimi i cdo zgjedhjeje të marrë që do të krahasohet është i njëjtë, koeficientin e besimit, e ka pikërisht $1 - \alpha$.

Nga ky sqarim shihet se, kur numri i mesatareve që do të kontrollojmë është 2, $t_T = t\sqrt{2}$, pra testi Tukey dhe Studenti janë identik.

Problemi i ndryshimit midis dy testeve do të shfaqet nëse krahasojmë më shumë mesatare.

Në tabela të vecanta të paraqitura, mund të lexohen vlerat e t_T për një nivel besimi të dhënë α dhe, për një numër të përcaktuar të zgjedhjes së nënpopullimeve n , që marrin pjesë në këtë krahasim të shumëfishtë. Në grafikun e paraqitur mund të lexohen këto vlera t_T për $\alpha = 0.05$ në varësi të numrit të mesatareve që do të krahasojmë.

Figura 1-1 Vlerat Kritike – numri i mesatareve që krahasohen



Tukey kujdeset edhe për reduktimin e numrit të krahasimeve të mesatareve të popullimeve. Psh, në rastin e shqyrtimit të mesatareve (A,B,C,D), të cilat supozojmë se i kemi renditur sipas rendit $A > B > C > D$, sipas testit Tukey, nuk është patjetër e domosdoshme të bëhen të gjitha krahasimet e mundshme.

Për këtë sugjerohet kjo *radhë krahasimi*: Krahasohet më parë popullimi me mesataren më të madhe (A), me atë më të vogël (D). Nëqoftëse t_T (vlera e statistikës që Tukey propozon për vlerësim) është më e vogël se vlera t_T e shpërndarjes së propozuar, hipoteza H_0 pranohet si e vërtetë. Pra, mesataret e këtyre zgjedhjeve nuk janë statistikisht të ndryshme. Kuptohet, derisa nuk kanë ndryshim këto dy mesatare që kishin diferencën më të madhe, krahasimi i mesatareve të popullimeve të tjera është i panevojshëm. Procesi i krahasimit të çifteve ndërpritet.

Përfundimisht, për rradhën e krahasimit në testin Tukey, fillohet gjithnjë me krahasimin e mesatares më të madhe me atë më të vogël. Pastaj vazhdohet me krahasimin e kësaj mesatareje më të madhe, me mesataren e dytë më të ulët, derisa mesatarja më e lartë të jetë krahasuar me të gjithë mesataret e tjera, ose sic thamë, derisa të mos ketë ndryshim mes mesatareve në studim. Pas këtij krahasimi, i njëjti proces përsëritet me mesataren e dytë për nga madhësia: krahasohet me mesataren më të vogël, pastaj me atë të dytë më të ulët, e kështu me rradhë.

Procesi i krahasimit përfundon ose kur janë bërë të gjithë krahasimet e mundshme e ka ndryshim mes cdo çifti që krahasohet, ose sa gjendet çifti i pare i zgjedhjeve mesataret e të cilëve s'kanë ndryshim statistikor.

Si konkluzion, mund të themi se metoda Tukey është një metodë e krahasimit të shumëfishtë që ka një kontroll të dobët të nivelit të besimit, në kuptimin që në cdo krahasim çiftesh ajo nuk kontrollon drejtpërdrejt nivelin e besimit për këtë krahasim, por ndërhyr në mënyrë të tërthortë në të.

Me këtë metodë, në rast se vëllimi i cdo zgjedhje të marrë që do të krahasohet, është i njëjtë, koeficienti i besimit është pikërisht $1 - \alpha$. Në programet kompjuterike statistikore mund të përdoret lehtësisht kjo metodë e krahasimit të shumëfishtë

(Tukey). Kjo do të na afishojë gjithë intervalet e besimit të hipotezave $\mu_i = \mu_j$, ose $\mu_i - \mu_j = 0$, me koeficient besimi 0.95. Kështu, thjeshtë mund të identifikohet cili nga çiftet e nënpopullimeve, ka mesatare të ndryshme. Dy nënpopullime konsiderohen me mesatare të ndryshme nëse intervali i tyre i besimit nuk përmban 0.

1.6.3 Metoda Scheffe

Metoda Scheffe^[1] është metodë e krahasimit të shumëfishtë, por ky krahasim bëhet në një hap të vetëm. Kjo përben dhe thelbin e ndryshimit të saj me metodën e mësipërme Tukey, e cila duke krahasuar dyshe nënpopullimesh, zhvillohej në disa etapa.

Kjo është një ndër metodat bazë që përdoret në analizën e regresit linear për të realizuar krahasimin e shumëfishtë, dhe natyrisht dhe në ANOVA-n (si një rast i vecantë i regresit linear) Metoda ka marre emrin e statistikanit amerikan Henry Scheffe.

Ndërtohet një kontras C i përcaktuar nga barazimi

$$C = \sum_{i=1}^r c_i \mu_i$$

dhe shtrohet hipoteza $H_0: C = 0$ me hipotezë alternative $C \neq 0$

Vlerësimi i C behet sipas barazimit të mëposhtëm

$$\hat{C} = \sum_{i=1}^r c_i \bar{Y}_i$$

dhe varianca e së cilës është

$$s_{\hat{C}}^2 = \hat{\sigma}_e^2 \sum_{i=1}^r \frac{c_i^2}{n_i},$$

ku n_i është vëllimi i zgjedhjes së marrë nga nënpopullimi i i -të , dhe

$\hat{\sigma}_e^2$ varianca e vlerësuar (estimated)

Forma e intervalleve të besimit, me koeficient besimi $1-\alpha$, është:

$$\hat{C} \pm s_{\hat{C}} \sqrt{(r-1) F_{\alpha; r-1; N-r}}$$

ku N është vëllimi i zgjedhjes së gjithë popullimit, r numri i nënpopullimeve në të cilat është ndarë popullimi i madh.

Pas afishimit të këtij intervali, kontrollohet nëse ai e përmban apo jo zeron. Nëse nuk e përmban, atëherë hipoteza H_0 bie.

Metoda në fjalë konsiderohet shume rigoroze, sepse kontrolllet i ben per një gabim të rendit shumë më të vogël se α .

1.6.4 Krahasimi i dy metodave

Metoda Tukey përdoret në nënpopullime me vëllime të njëjta, ndërsa metoda Scheffe nuk vendos një kufizim të tillë.

Metoda Scheffe nuk e ka të domosdoshëm kushtin që $\sum_{i=1}^r c_i = 0$ por, nëse ky kusht plotësohet në vendosjen e konstanteve, atëherë themi se po përdoren koeficientë që

plotesojnë **kushtin e kontrastit**. Ndërsa metoda Tukey e ka domosdoshmeri qe gjate krahasimit të cifteve të plotësohet kushti i kontrastit.

Së fundmi, metodat në fjalë bazohen në statistika me shpërndarje të ndryshme. Metoda Tukey bazohet ne statistiken me shperndarje Studenti, ndersa ajo Scheffe ne shperndarjen Fisher. Është e kuptueshme që rezultatet e dy metodave në fjalë të jenë të ndryshme.

Përfundimisht, në rastin e krahasimit të mesatareve të tri nënpopullimeve me vëllim të njëjtë zgjedhje, interval më të ngushtë jep metoda Tukey, ndërsa ne rastet e 4 e me shumë nënpopullimeve interval më të ngushtë e siguron metoda Scheffe (bazuar kjo nga punimi i O'Neil dhe Wetherill më 1971)^[3]

Metoda Scheffe coi në trajtimin e një problemi tjetër, atij të Analizes së Kontrastit^[2]. Është pikërisht kjo analizë ajo që përcakton tërsisht modelin e ANOVA-s me anë e vlerave të këtyre koeficientëve.

Kontrastin si të thuash mes nënpopullimeve ajo e përkthen në bashkësi numrash, që janë pikërisht koeficientët e kontrastit (cij quhen edhe koeficientë të kontrastit).

Në varësi të rastit që shqyrtohet, koeficientët e kontrastit mund të jenë të planifikuar ose jo. Të përcaktuar janë në rastet kur eksperimenti kryhet për të përcaktuar një hipoteze me koeficientë të paracaktuar. Pra, cdo gjë është mirë e planifikuar. Ndërsa rasti tjetër është kur këto koeficientë varen nga rezultatet e vrojtimit që është kryer. Metoda e Kontrastit, në këtë rast quhet Post Hoc.

Nëse koeficientët nuk e plotesojnë kushtin e kontrastit, atëhere këto koeficientët tregojnë koeficientin e pjesshëm te korelacionit mes variablit në shqyrtim e variablit të varur. Pra përcaktojnë kontributin e variablit ne shqyrtim në variablin e varur.

1.6.4.1 Zbatim 1

Një perdorim i metodes Scheffe me koeficiente qe nuk plotesojne kushtin e kontrastit, është bërë në trajtimin e **problemit të ndikimit te SO₂ tek prodhueshmëria e bimëve** (shih problemin e trajtuar tek Kapitulli 6).

Në rastin në fjalë, pas aplikimit të metodës Leven, në krahasimin e variancës se nënpopullimeve tek tri komponentët (patate, perime, fasule) u morën vlera signifigative (shih Tabelen 1-1). Pra, në bazë të zgjedhjes së marre, u tregua per hedhje poshtë të hipotezës së barazimit të variances, duke na penguar në përdorimin e ANOVA-s .

Atëhere u përdor metoda e t-testit Welch,që tregoi ndikim negativ të zonës së ndotur, pra që tri komponentët e vektorit bimë kanë prodhueshmeri me të larte tek zona e kontrollit se tek zona industriale, e ndotur. U shtrua problemi i përcaktimit se në c'shcallë ishte ky ndikim.

Tabela 1–1 Afishimet e Testit Levene

Levene's Test for Equality of Variances	
F	Sig.

grure	Equal variances assumed	.234	.629
miser	Equal variances assumed	2.447	.120
patate	Equal variances assumed	31.638	.000
perime	Equal variances assumed	21.773	.000
fasule	Equal variances assumed	4.941	.028
jonxhe	Equal variances assumed	16.393	.000
fruta	Equal variances assumed	1.313	.254

Konkluzioni i lartpërmendur (1.6.4 Krahاسimi i dy metodave) tregoi se metoda e Kontrastit mund të perdoret edhe në krahاسimin e dy nënpopullimeve me ane të t-testit. Koeficientët e Kontrastit që nuk plotësojnë kushtin e kontrastit, përcaktojnë tërësisht koeficientet ^[2] me të cilën lidhen mesataret e dy nënpopullimeve.

Pas shkrimit të kodit në SPSS:

ONEWAY patate perime fasule BY area
/CONTRAST=1 -1.3
/MISSING ANALYSIS.

Përdorimi i Metodes së Kontrastit gjatë t-testit, dha afishimet e mëposhtme (Tabela 1-2):

Tabela 1–2 Rezultatet e Metodës së Kontrastit

Contrast Tests

		Contrast	Value of Contrast	Std. Error	t	df	Sig. (2-tailed)
patate	Assume equal variances		67.8788 ^a	17.45542	3.889	153	.000
	Does not assume equal variances		67.8788 ^a	23.76892	2.856	47.916	.006
perime	Assume equal variances		26.1727 ^a	14.10142	1.856	153	.065
	Does not assume equal variances		26.1727 ^a	17.38772	1.505	58.402	.138
fasule	Assume equal variances		12.4687 ^a	5.55062	2.246	153	.026
	Does not assume equal variances		12.4687 ^a	6.30780	1.977	71.498	.052

a. The sum of the contrast coefficients is not zero.

Rezultatet e afishuara treguan se prodhimi në zonën e kontrollit të tri bimëve në fjalë është 130% e prodhimit të tek zonës së ndotur.

2 Modeli i regresit linear të shumëfishtë

2.1 Qëllimi i regresit të shumëfishtë

Qëllimi i metodës së regresit të shumëfishtë është analiza e lidhjes mes variablave të pavarur dhe një variabli të varur. Variablat e varur mund të jenë metrik, pra ordinal apo i vazhdueshëm, apo dischotomous, pra variabla nominal me dy vlera. Variabli i varur duhet të jetë metrik.

Metoda e regresit kontrollon nëse një lidhje e tillë ekziston, e nëqoftëse ajo ekziston, synon të përdoret informacioni ekzistues për variablat e pavarur, që të përmirësohet saktësia në parashikimin e vlerës së variablit të varur.

2.2 Tipet e ndryshme të regresit të shumëfishtë

Njihën tri tipe të regresit të shumëfishtë, secila e dezignuar në SPSS për tu përgjigjur pyetjeve të ndryshme:

Regresi i shumëfishtë standart që përdoret për të vlerësuar lidhjen mes një bashkësie variablash të pavarura dhe një variabli të varur.

Regresi statistic (ose stepwise) ose që përdoret për të gjetur nënbashkësinë e variablave të pavarur (Stepwise) që kanë lidhje më të fortë me variablin e varur.

Regresi hirarkik që përdoret për të kontrolluar lidhjen mes një bashkësie variablash të pavarur dhe një variabli të varur, pasi më parë është kontrolluar për efektin e disa variablave të pavarur në variablin e varur.

2.3 Trajtimi matematik i regresit linear të zakonshëm, standart (GLM)

Le të paraqesim trajtimin matematik të këtij problemi të regresit linear të shumëfishtë. Supozojmë se variabli y varet nga p variabla “jo statistikisht të pavarur” $x_1, x_2, \dots, x_p$, që quhen variabla të pavarur. Problemi i regresit të shumëfishtë ka të bëjë me ndërtimin e një modeli, ku y shprehet si funksion i këtyre variablave në formën

$$y = f(x_1, x_2, \dots, x_p) + \epsilon, \text{ ku } \epsilon \text{ ka shpërndarje } N(0, \sigma^2),$$

Ndërsa, kur modelin e kërkojmë në formën

$y = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_p x_p + \epsilon$, ku ϵ ka shpërndarje $N(0, \sigma^2)$, matematikisht kemi të bëjmë me modelin polinomial të regresit të shumëfishtë.

Në ndërtimin e këtij modeli duhen parë këto probleme:

2.3.1 Vlerësimi i parametrave të regresit linear të shumëfishtë (Hapi I)

Përcaktimi i zgjedhjes në shqyrtim mund të bëhet në dy mënyra :

Së pari, për një vlerë të fiksuar të variablave $x_1, x_2, \dots, x_p$ bëjmë disa matje dhe marrim vlera të ndryshme të y_i , e më pas, fiksojmë një tjetër p -she vlerash të $x_1, x_2, \dots, x_p$ e masim disa vlera të Y , duke marrë një zgjedhje të nënpopullimit në fjalë. Këtë proces

e vazhdojmë derisa të kemi marrë n vëzhgime. Sipas kësaj mënyre, vetëm Y shihet si variabël i rastit.

Së dyti, marrim n individë nga popullimi në fjalë dhe shikojmë, përcaktojmë, p+1 variablat për të cilën jemi të interesuar.

Kështu për të ndërtuar modelin e formës

$$y_i = \beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2} + \dots + \beta_p x_{ip} + \varepsilon_i \quad \text{për } i=1..n$$

ne duhet të vlerësojmë $\beta_0, \beta_1, \beta_2, \dots, \beta_p$ që janë parametrat e këtij modeli dhe $\varepsilon_1, \varepsilon_2, \dots, \varepsilon_n$, n variabla të pavarur $N(0, \sigma^2)$.

Modelin matematik që kërkojmë të ndërtojmë, mund ta paraqesim edhe në formë vektoriale^[4].

Le të shënojmë me $\beta = (\beta_0, \dots, \beta_p)'$ vektorin parametrik $p + 1$ dhe $y = (Y_0, \dots, Y_n)'$ vektorin i vëzhgimeve, ndërsa $e = (e_0, \dots, e_n)'$ vektori i gabimeve dhe matrica e modelit $n * (p + 1)$

$$\begin{bmatrix} 1 & x_{11} & \dots & x_{1n} \\ 1 & x_{21} & \dots & x_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & x_{n1} & \dots & x_{pn} \end{bmatrix}$$

Kështu modeli matematik ka formën e këtij ekuacioni vektorial:

$$Y = X'\beta + e$$

ku e është një matricë $N(0, \sigma^2 I)$

Për vlerësimin e parametrave $\beta_0, \beta_1, \beta_2, \dots, \beta_p$, përdorim metodën e katrorëve më të vegjël, pra ndërtojmë shumën e mëposhtme:

$$S(\beta) = \sum_i (y_i - \beta_0 - \beta_1 x_{i1} - \beta_2 x_{i2} - \dots - \beta_p x_{ip})^2$$

Duke përdorur shënimet e mësipërme, matrica ku bazohet vlerësimi mund të paraqitet në formë matricore $S = (y - X'\beta)'(y - X'\beta)$.

$$\begin{aligned} \frac{\partial S}{\partial \beta} &= 0 = -X(Y - X'\beta) + (Y - X'\beta)'(-x^T) \\ 0 &= -XY + XX^T\beta - Y^T X^T + (X^T\beta)^T X^T \\ XX^T\beta &= XY \end{aligned}$$

Kështu vektori i koeficientave të regresit $b = (b_0, b_1, \dots, b_p)'$ mund të llogaritet duke zgjidhur ekuacionin normal: $(XX')\beta = Xy$ dhe përcaktojmë vlerat e parametrave β duke kërkuar minimizimin e kësaj shume. Zgjidhja e marë është

$$b = (XX')^{-1}(Xy)$$

Ky vlerësim i parametrave β është një vlerësim i pazhvendosur, sepse:

$$E(\hat{\beta}) = E(b) = E((XX^T)^{-1}(XY)) = (XX^T)^{-1}E(XY) = (XX^T)^{-1}XE(Y) = \\ = (XX^T)^{-1}XX^T\beta = \beta$$

Nga teorema e Gaus-Markovit dimë, se vlera e parashikuar $\tilde{y}$ ka variancën minimale, nga të gjithë vlerësuesit e tjerë linearë të vlerave të X

$S = \sum_i (y_i - \tilde{y})^2$ na shërben si një vlerësim për σ^2 . Madje, një vlerësim të pazhvendosur të σ^2 , na e siguron dhe $MS_R = \sum_i (y_i - \tilde{y})^2 / (n-p-1)$ ose në trajtë vektoriale vlerësimi i

pazhvendosur për variancën është $s^2 = (\mathbf{y} - \mathbf{X}'\mathbf{b})'(\mathbf{y} - \mathbf{X}'\mathbf{b}) / (\mathbf{n} - \mathbf{p} - 1)$

dhe matrica e kovariancës është $cov(b) = s^2(XX')^{-1}$

Paraqitet edhe një formë tjetër e modelit të regresit linear të shumëfishtë, forma e qendërzuar. Modeli i qendërzuar jepet nga:

$$y_i = \beta'_0 + \beta_1(x_{1i} - \bar{x}_1) + \dots + \beta_p(x_{pi} - \bar{x}_p) + e_i, \quad 1, \dots, n,$$

ku

$$\bar{x}_j = \frac{1}{n} \sum_{k=1}^n x_{jk}, \quad j = 1, \dots, p, \quad \text{dhe} \quad \beta'_0 = \beta_0 + \beta_1\bar{x}_1 + \dots + \beta_p\bar{x}_p$$

Vlerësuesit me Metoden e Katrorëve më të Vegjël të $\beta_1, \dots, \beta_p$, janë të njëjta si më sipër dhe emërohen $b_1, \dots, b_p$, ndërkohë që vlerësuesi me këtë metodë i β'_0 është $b'_0 = \bar{y}$. Avantazhi i këtij modeli është se vlerësuesit me Metoden e Katrorëve më të Vegjël $b_1, \dots, b_p$ janë të pakorreluar me b'_0 . Kjo do të thotë që është më e lehtë të gjesh intervale besimi, duke u bazuar në vlerat e parashikuara $\hat{y} = \bar{y} + b_1(x_1 - \bar{x}_1) + \dots + b_p(x_p - \bar{x}_p)$, siç do të shohim më poshtë.

Modeli i qendërzuar mund të shkruhet në formë matricore si më poshtë.

Le të jetë $\mathbf{A}$ matrica e shumes të katrorëve $p \times p$ dhe prodhimit të devijimeve me elementët i, j $a_{ij} = \sum_{k=1}^n (x_{ik} - \bar{x}_i)(x_{jk} - \bar{x}_j), i, j = 1, \dots, p. (*)$

Le të jetë $\mathbf{g}$ vektori $p \times 1$ me elementin e i -të $g_i = \sum_{k=1}^n (y_k - \bar{y})(x_{ik} - \bar{x}_i), i = 1, \dots, p.$

Atëherë vektori i katrorëve më të vegjël

$$\bar{\mathbf{b}} = (b_1, \dots, b_p)' = \mathbf{A}^{-1}\mathbf{g}$$

dhe për më tepër

$$Cov(\bar{\mathbf{b}}) = \sigma^2 \mathbf{A}^{-1}, \text{ dhe } cov(\bar{y}, b_i) = 0, \text{ për } i = 1, \dots, p$$

2.3.2 Interpretimi gjeometrik i zgjidhjes me Metodën e Katrorëve më të Vegjël

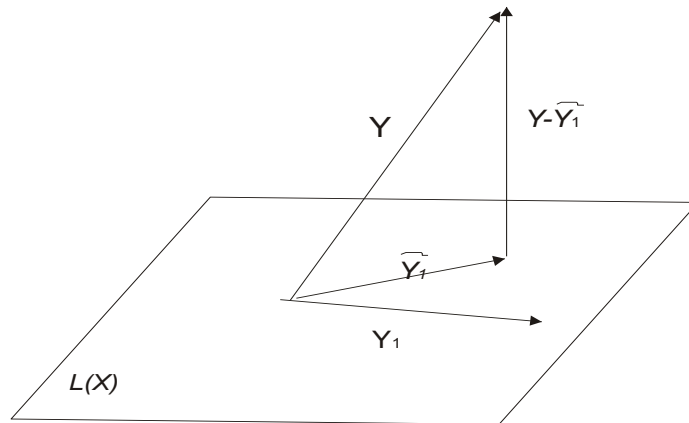
Në thelb të kësaj metode, qëndron përcaktimi i β , që të

$$\min_{\beta \in R^p} S(\beta) = \min_{\beta \in R^p} (\underline{Y} - \mathbf{X}'\beta)'(\underline{Y} - \mathbf{X}'\beta) = \min_{\beta \in R^p} (\underline{Y} - \mathbf{Y}_1)'(\underline{Y} - \mathbf{Y}_1)$$

ku $\underline{Y}$ janë shënuar vlerat e matura të Y dhe me $\mathbf{Y}_1 = \mathbf{X}^T \beta$,

Ky minimum kërkohet për çdo $\beta \in \mathbb{R}^p$, ndërsa $Y_1 = X^T \beta$ në fakt, nuk mbushin të gjithë $\mathbb{R}^n$. Ato janë vetëm mbështjellset e shtyllave të X . Për këtë bashkësi mbështjelljesh, përdorim shënimin $L(x)$.

Figura 2-1 Paraqitja gjeometrike



$(Y - Y_1)$, gjeometrisht na paraqet vektorin diferencë të dy vektoreve Y dhe Y_1 .

Me këto simbolika të futura, përcaktimi i β kthehet në përcaktimin e vektorit diferencë $(Y - Y_1)$ që ka gjatësi më të vogël, dhe Y_1 për këtë vektor është projekcioni i Y në hiperplanin $L(x)$.

Siç e dimë ky vektor ka gjatësi më të vogël nëse ai është pingul me hiperplanin e Y_1 , pra me $L(x)$, që, siç e sqaruam më lart, është pjesë e $\mathbb{R}^n$, ku ndodhen tërë $X \beta$.

Kështu në bazë të njohurive gjeometike, për vektorin ku arrihet minimumi i dëshiruar (po e shënojmë me $Y - \tilde{Y}_1$), nga kushti i ortogonalitetit, mund të shkuajmë:

$$X(Y - \tilde{Y}_1) = 0$$

$$XY - X\tilde{Y}_1 = 0$$

$$XY = X\tilde{Y}_1 = X(X' \beta)$$

Vlerësimi b për parametrat β do të merret nga zgjidhja e ekuacionit matricor $XY = X(X' \beta)$

Ndoshta ky vëzhgim gjeometrik i problemit mund të duket, i vështirë kur flitet për ndërtimin e një modeli të shumëfishtë. Por, në rastin kur $X \in \mathbb{R}^2$, pra në rastin e modelit linear të thjeshtë, ai ka paraqitjen gjeometrike si në figurën e mësipërme.

2.3.3 Vlerësimi intervalor i parametrave te regresit (Hapi II)

Programet e plota të posaçme kompjuterike, përveç vlerësimit pikësor të parametrave $\beta_0, \dots, \beta_p$ me metoden e të katrorëve më të vegjël, japin informacion edhe për intervalet e besimit dhe për të kontrollin e hipotezave rreth këtyre parametrave.

Programet llogarisin gabimet standarde të koeficientëve. Për çdo β_i , gabimi standard i koeficientit, $se(b_i)$ është vlerësimi i devijimit standard të vlerësimit b_i të parametrat $\beta_i, i = 1, \dots, p$. Meqënëse β është marrë si kombinim linear i y_i , që supozohet se kanë shpërndarje normale, atëherë mund të shfrytëzojmë faktin që $(\beta_i - b_i)/se(b_i)$ ka shpërndarje Studenti me $n-p-1$ shkallë lirie.

Kështu që intervali i besimit me koeficient besimi $(1 - \alpha) = \gamma$ të β_i është:

$[L1, L2]_{\beta_i} = [b_i - se(b_i)t_{1-\alpha/2}(v_R); b_i + se(b_i)t_{1-\alpha/2}(v_R)]$ për $i = 1, \dots, p$

2.3.4 Kontrolli i hipotezave (Hapi III)

Kontrolli i hipotezave për koeficientët $\beta_1, \dots, \beta_p$ ndahet në tri kategori: mund të testohen të gjithë koeficientët $\beta_1 = \beta_2 = \dots = \beta_p = 0$, ose mund të testohet që secili koeficient $\beta_k = 0, k = 1, \dots, p$, ose mund të testojmë që çdo nënbashkësi me m koeficientë është e barabartë me zero. Të gjitha këto hipoteza mund të interpretohen nëpërmjet variabla të pavarur si më poshtë.

Së pari kontrollojmë hipotezën $H_0: \beta_1 = \beta_2 = \dots = \beta_p = 0$

(pra praktikisht që variablat $x_1, x_2, \dots, x_p$ s'ndikojnë në vlerësimin e y)

Kjo hipotezë s'është gjë tjetër veçse hipoteza "variablat e pavarur $X_1, \dots, X_p$ nuk e përmirësojnë parashikimin e Y më shumë se $\hat{y} = \bar{y}$ ",

me hipotezë alternative "jo të gjithë koeficientët janë të barabartë me zero"

Nëse nuk hidhet poshtë kjo hipotezë do të thotë që, pranohet $\bar{y}$ si parashikuesi më i mirë i Y , dhe nuk ia vlen ndërtimi i modelit të regresit linear.

Për kontrollin e kësaj hipoteze përdorim statistikën

$$\frac{MS_D}{MS_R} = \frac{\sum (\bar{y} - \tilde{y})^2 / p}{\sum (y_i - \tilde{y})^2 / (n - p - 1)} \text{ me shpërndarje } F(p, n-p-1)$$

Për të spjeguar këtë, mjafton të kujtojmë se SS_R / σ^2 ka shpërndarje $\chi^2 (n-p-1)$.

dhe $SS_T / \sigma^2 = \sum (y_i - \bar{y})^2 / \sigma^2$ ka shpërndarje $\chi^2 (n-1)$.

Si përfundim $SS_D / \sigma^2 = \sum (\tilde{y} - \bar{y})^2 / \sigma^2$ do të ketë shpërndarje $\chi^2 (p)$.

Nëse hipoteza H_0 pranohet si e vërtetë, kjo do të thotë se për vlerësimin e Y mund të përdorim $\bar{y}$. Nëse kjo hipotezë nuk pranohet si e vërtetë, si e tillë pranohet hipoteza alternative, Nëse hipoteza H_0 hidhet poshtë atëherë do të pranohet hipoteza alternative si e vërtetë, pra që jo të gjithë parametrat janë 0 ose njëloj, "disa nga variablat e pavarur përmirësojnë parashikimin e Y më mirë se $\hat{y} = \bar{y}$ ".

Ndaj kalohet në etapën tjetër.

Etapa e dytë: Shtrohen hipotezën $H_0: \beta_k=0$, për $k=1..p$, me hipotezë alternative $H_1: \beta_k \neq 0$, Hipotezat $H_0: \beta_k = 0$ për $k = 1, \dots, p$ mund të shihet si hipoteza ku “variabli X_k nuk përmirëson parashikimin e Y krahasuar me atë që fitohet nga efekti në Y i $p - 1$ variablave të tjerë” Gjeometrikisht mendohet që hiperplani që po ndërtohet është paralel me këtë bosht.

Të dhënat për kontrollin e hipotezës së shtruar përcaktohen nga raporti F që jepet në kolonën e fundit të tabelës së analizës së variancës, dhe është,

$$F = \frac{MS_D}{MS_R}.$$

Kur është e vërtetë H_0 , F ka njëshpërndarje F me $v_D = p$ dhe $v_R = n - p - 1$ shkallë lirie. p -vlëra është sipërfaqja në të djathtë të F të vrojtuar, nën grafikun e shpërndarjes së $F(v_D, v_R)$.

Si test statistikor për kontrollin e hipotezës merret:

$$F = \frac{b_k^2}{[se(b_k)]^2}$$

e cila kur H_0 është e vërtetë, ka një shpërndarje F me 1 dhe $v_R = n - p - 1$ shkallë lirie. p -vlëra është hapësira në të djathtë të F në grafikun e shpërndarjes $F(1, v_R)$.

Disa programe kompjuterike e japin këtë vlerë të F për çdo b_k .

Programe të tjerë (SPSS) japin vlerën e statistikës ekuivalente

$$t = \frac{b_k}{se(b_k)}$$

e cila, nëse H_0 është e vërtetë ka shpërndarjen Studenti me $v_R = n - p - 1$ shkallë lirie. p -vlëra është dyfishi i hapësirës në të djathtë të $|t|$ në grafikun e $S(v_R)$.

Etapa e tretë: Hipoteza e fundit, që një nënbashkësi prej m koeficientësh është zero është më e vështirë të provohet. Pa bërë ndonjë kufizim të madh, le të mendojmë që kjo nënbashkësi konsiston në m koeficientët e parë $\beta_1, \dots, \beta_m$. Atëherë të provosh $H_0: \beta_1 = \dots = \beta_m = 0$ është e njëjtë sikur të provosh që “ m variablat $X_1, \dots, X_m$ nuk përmirësojnë parashikimin e Y nga ajo që fitohet nga regresi i Y nga $X_{m+1}, \dots, X_p$ ”.

Fillimisht, ekzekutrohet programi për të llogaritur regresin e Y me $X_{m+1}, \dots, X_p$. Nga analiza e tabelës së variancës marrim shumën e katrorëve të mbetjes prej SS_R' . Më pas përdoret programi për të llogaritur regresin e Y me $X_1, \dots, X_m, \dots, X_p$, dhe nxjerrim shumën e katrorëve të mbetjes dhe katrorin mesatar prej respektivisht SS_R dhe MS_R . Atëherë testi statistikor për H_0 është

$$F = \frac{(SS_R' - SS_R)/m}{MS_R}$$

e cila kur H_0 është e vërtetë ka shpërndarje F me m dhe $v_R = n - p - 1$ shkallë lirie. p -vlëra është hapësira në të djathtë të F së vrojtuar nën grafikun e shpërndarjes $F(m, v_R)$.

2.3.5 Vërejtje

1. Varianca e $\hat{y}$ prej $x_1, \dots, x_p$ është

$$V(\hat{y}) = \sigma^2 \left(\frac{1}{n} + d'A^{-1}d \right)$$

ku elementët e A janë përcaktuar në seksionin 2.3.1 nga (*), dhe

$d = (x_1 - \bar{x}_1, \dots, x_p - \bar{x}_p)'$, me $\bar{x}_i = \frac{1}{n} \sum_{k=1}^n x_{ik}$.

Një interval besimi me koeficient besimi $(1 - \alpha)$, për mesataren e vlerës së vërtetë të Y nëse $x_1, \dots, x_p$ njihet, është

$$\hat{y} \pm \left[s^2 \left(\frac{1}{n} + d' A^{-1} d \right) \right]^{\frac{1}{2}} t_{1-\left(\frac{\alpha}{2}\right)}(n - p - 1)$$

Një interval besimi me koeficient besimi $(1 - \alpha)$ për një vlerë të vetme të parashikuar të Y nëse $x_1, \dots, x_p$ njihet, është

$$\hat{y} \pm \left[s^2 \left(1 + \frac{1}{n} + d' A^{-1} d \right) \right]^{\frac{1}{2}} t_{1-\left(\frac{\alpha}{2}\right)}(n - p - 1)$$

Disa programe shfaqin elementët e A^{-1} për të lehtësuar llogaritjen e ekuacioneve të mësipërme.

2. Nëse testojmë Hipotezat $H_0: \beta_k = 0$ për disa vlera të k me të njëjtin nivel rëndësie α , atëherë niveli i rëndësisë së përbashkët nuk është patjetër α . Për të shmangur këtë problem mund të ndërtojmë intervalet e besimit të shumëfishtë për të gjitha $\beta_k, k = 1, \dots, p$, në mënyrë që koeficienti i besimit të përbashkët të jetë $1 - \alpha$. Intervali i besimit të shumëfishtë me koeficient $(1 - \alpha)$ për β_k është

$$b_k \pm (se(b_k)[pF_{1-\alpha}(p, n - p - 1)]^{1/2}$$

Hipoteza $H_0: \beta_k = \beta_k^{(0)}$ bie poshtë me nivel rëndësie α , nëse $\beta_k^{(0)}$ nuk është në këtë interval.

2.3.6 Përfundime

- Në regresin e shumëfishtë standart të gjithë variablat e pavarur futen njëkohësisht në model.
- Vlerësimi i R and R^2 përcaktojnë fortësinë e lidhjes mes variablave të pavarur dhe variablit të varur. Testi Fisher përdoret për të përcaktuar nëse kjo lidhje e përcaktuar nga zgjedhja mund të përgjithsohet, apo jo, për të gjithë popullimin.
- Testi t përdoret për të vlerësuar lidhjen individuale mes çdo variabli të pavarur dhe variablit të varur.

2.4 Zbatim 1

Nëse i referohemi zbatimit "**Efekti i SO₂ në prodhimet bujqësore**", vihet re përdorimi i kësaj metode studimi të të dhënave. Në këtë zbatim kërkohet të kontrollohet nëse ndotja industriale ndikon apo jo në prodhueshmërinë e produkteve bujqësore. Pra a ndikon rritja e produkteve në zonen industriale apo jo industriale, në sasinë e produktit të marrë?

Për këtë qëllim janë kontrolluar në të dy zonat, sasia e prodhimit të të njëjtave produkte bujqësore. Pas përdorimit të testeve të përshtatshme, është vënë re se për produktet 'grurë', 'misër' dhe 'fruta', nuk ka dallim midis zonave. Por, sasia e produkteve bujqësore dihet që varet nga përkujdesja ndaj sipërfaqeve të

shfrytëzueshme për të mbjella dhe kushtet abiotike e biotike. Për të studiuar efektin e këtyre faktorëve në këtë problem para aplikimit të metodës në fjalë është kontrolluar :

A është bashkësia e të dhënave e përshtatshme për këtë aplikim statistikor?

Kontrollohet më parë që të mos ketë probleme me mungesën e të dhënave, probleme me outliers, plotësohen kushtet e aplikimit të kësaj metode dhe më pas, që ndarja e të dhënave që shfrytëzohen, të konfirmojë për besueshmeri të rezultateve.

Variablat 'distance', që përcakton distancën në metër të individit nga pusi, 'mbetje urbane' që kodohet sipas shkallës së trajtimit të mbetjës urbane, në rastin tonë më i lartë kodi më i ulët niveli i trajtimit të kësaj lloji mbetje; dhe kujdesi ndaj 'ujrave të zeza', që është koduar në mënyrë të tillë që rritja e kodit tregon përkujdesje të ulët të trajtimit të tyre, janë konsideruar pas studimit, si variablat që duhen marrë në shqyrtim.

Metoda e regresit të shumëfishtë kërkon që variabli i varur 'bimet' të jetë metrik, dhe variablat e pavarur, të jenë metrikë ose kategorikë, që përcaktojnë parametrat në fjalë. 'ujrat e zeza' dhe 'sistemi i mbeturinave' janë variabla ordinal të përcaktuara sipas niveleve, pra në një farë mënyre mund të mendohet që këto variabla mund të masin parametrat në fjalë. Mendohet që variablat e këtij lloji mund të trajtohen si variabla metrikë. Por, ka statisticienë që nuk janë dakord me këtë përdorim. Ndaj, në fund të interpretimit duhet bërë një vërejtje, nëse këto lloj variablash përdoren.

'distanca nga pusi' është variabël metrik ndaj ai përdoret në metodën e regresit linear. Minimumi i raportit të vëllimit të zgjedhjes që shfrytëzohet me variablat e pavarur që marrin pjesë në studim, duhet të jetë minimumi 5 me 1. Por 155 rastet e vrojtues dhe 3 variabla të pavarur bëjnë që raporti i kësaj analize të jetë 51.7 me 1, duke plotësuar kështu edhe kushtin e preferuar 15 me 1.

Tabela 2–1 Tabela e koeficientëve të modelit të regresit

Coefficients

Model	Unstandardized Coefficients		Standardized Coefficients	t	Sig.
	B	Std. Error	Beta		
1 (Constant)	731.418	293.331		2.493	.014
area	-.76.385	116.533	-.157	-.655	.513
sistemi i ujrave të zeza	-.215.809	56.983	-.298	-3.787	.000
terheqja e mbeturinave	-.7.930	21.784	-.034	-.364	.716
distanca shtëpi-pusne meter	-.053	.183	-.065	-.288	.774

a. Dependent Variable: fruta

Tabela 2–2 Afishimet e t-testeve dhe koeficienteve të regresit

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	929.111	287.930		3.227	.002
	area	-20.027	114.388	-.043	-.175	.861
	sistemi i ujrave të zeza	-150.558	55.933	-.216	-2.692	.008
	terheqja e mbeturinave	18.431	21.383	.083	.862	.390
	distanca shtepi-pusneti në meter	-.022	.180	-.027	-.120	.905

a. Dependent Variable: miser

Ndikimi i faktorit abiotik *ujrat e zeza* duhet marrë parasysh, sepse tek *gruri* (t-test -3.415 dhe sig. 0.01), *misri* (t-test -2.692 dhe sig. 0.08) dhe *frutat* (t-test -3.787 dhe sig. < 0.01). Vlerat negative të koeficientëve përkatësisht -436, -150.5 dhe -215 në modelin e regresit linear përcaktojnë drejtimin e lidhjes. Duke verejtur kodin e variablit *ujra të zeza*, ndërsa kodi rritet nuk ka asnjë proces sanitar në trajtimin e ujrave të zeza, arrihet në përfundimin se një sistem më i mirë i ujrave të zeza ndimon në prodhueshmëri me të lartë tek disa bime.

Hipotezat zero që nuk kanë lidhje gjithë variablat e pavarur me variablin e varur, *prodhueshmëria e bimës*, bie poshtë. Në tabele (model summary) janë paraqitur dhe vlerat e statistikes R^2 -change. Informacioni i marrë nga variablat e shtuar, redukton parashikimin në prodhueshmërinë e bimëve (7.6% për *grurin*, 5.4% për *misrin*, and 6.7% për *frutat*).

2.4.1 Vërejtje

Në modelin e regresit, β_i përcakton shkallën e ndryshimit të Y në varësi të X_i , kur vlerat e $X_j, j = 1, \dots, p, j \neq i$ janë mbajtur të pandryshueshme. Megjithatë, këto koeficientë mund të mos ketë kuptim t'i krahasojmë në madhësi për shkak të diferencave në njësitë e matjes së $X_1, \dots, X_p$. Kjo vështirësi mund të kalohet duke përdorur *variabla të pavarur të standartizuar*^[5]. Kështu, për $j = 1, \dots, p$, le të jenë $Z_j = X_j/s_j$, ku $s_j^2 = \sum_{i=1}^n (x_{ji} - \bar{x}_j)^2 / (n - 1)$.

Modeli i regresit të shumëfishtë në termat e Z_j jepet tashmë prej ekuacionit

$$y_i = \gamma_0 + \gamma_1 z_{1i} + \dots + \gamma_p z_{pi} + e_i, \quad i = 1, \dots, n$$

ku $\gamma_k, k = 0 \dots, p$, janë parametra të panjohur dhe e_i janë terma gabimi të pavarura $N(0, \sigma^2)$. Vlerësuesit e katrorëve më të vegjël c_k të γ_k dhe testet e hipotezave rreth γ_k ndjekin teorinë e mësipërme me Z_j dhe γ_k që zëvendësojnë x_j dhe β_k respektivisht. Avantazhi i këtij standartizimi është që $\gamma_1, \dots, \gamma_p$ janë tani madhësitë që përcaktojnë ndryshimin, të matura jo më në shkallë të ndryshme. Kjo na lejon të arrijmë në konkluzionin rreth shkallës së kontributit të variablave të pavarur $Z_1, \dots, Z_p$ (ose në

mënyrë të njëjtë, $X_1, \dots, X_p$). Pra, një vlerë e lartë e c_j sjell një shkallë të lartë të kontributit të Z_j (ose X_j), $j = 1, \dots, p$.

2.5 Koeficienti i korrelacionit të shumëfishtë

Do diskutojmë teorikisht modelin e regresit linear të shumëfishtë. Teoria pranon që $p + 1$ variablat $Y, X_1, \dots, X_p$ janë të gjithë variabla të rastësishëm, dhe për më tepër ata kanë një shpërndarje normale të shumëpërmasore. Do të tregohet këtu që mesatarja e shpërndarjes së kushtëzuar të Y , në se $X_1 = x_1, \dots, X_p = x_p$ njihet, është funksion i regresit të shumëfishtë linear të $\beta_0 + \beta_1 x_1 + \dots + \beta_p x_p$. Kjo motivon modelin e regresit të shumëfishtë linear, në të cilin varianca e gabimit σ^2 është funksion i variancës σ_y^2 të Y dhe i një madhësie që quhet koeficient i korrelacionit të shumëfishtë.

Shqyrtojmë shpërndarjen normale të shumëfishtë të $Y, X_1, \dots, X_p$ me mesatare $\mu_y, \mu_1, \dots, \mu_p$ dhe varianca $\sigma_y^2, \sigma_1^2, \dots, \sigma_p^2$, respektivisht. Shënojmë kovariancën e Y me X_i me σ_{yi} , dhe kovariancën e X_i me X_j me σ_{ij} për $i, j = 1, \dots, p$. Gjithashtu përcaktojmë koeficientët e korelacionit të popullimit:

$$\rho_{yx_i} = \frac{\sigma_{yi}}{\sigma_y \sigma_i} \quad \text{dhe} \quad \rho_{x_i x_j} = \frac{\sigma_{ij}}{\sigma_i \sigma_j}$$

Për vlera të dhëna të $X_1 = x_1, \dots, X_p = x_p$, ekziston një nënpopullim i vlerave korresponduese të Y . Shpërndarja e saj, shpërndarja e kushtëzuar e Y në se $X_1 = x_1, \dots, X_p = x_p$, është gjithashtu normale me mesatare

$$\mu_{y, x_1 \dots x_p} = \mu_y + \beta_1(x_1 - \mu_1) + \dots + \beta_p(x_p - \mu_p)$$

Kjo quhet dhe pritja matematike me kusht e Y , kur dihet se $X_1 = x_1, \dots, X_p = x_p$, ose e regresit të Y ndaj $X_1, \dots, X_p$. Vlerat $\beta_1, \dots, \beta_p$ quhen koeficientë të regresit të pjesshëm, dhe janë funksion i variancave dhe kovariancave të mësipërme. Varianca e kësaj shpërndarje të kushtëzuar është

$$\sigma^2 = \sigma_y^2(1 - \rho_{y, x_1 \dots x_p}^2) \quad (*)$$

ku $\rho_{y, x_1 \dots x_p}$, rrënja katrore positive e $\rho_{y, x_1 \dots x_p}^2$, quhet koeficienti i korrelacionit të shumëfishtë të Y ndaj $X_1, \dots, X_p$.

Nëse përcaktojmë variablin e rastësishëm $e = Y - \mu_{y, x_1 \dots x_p}$, atëherë shpërndarja me kusht e e -së, nëse dihet se $X_1 = x_1, \dots, X_p = x_p$, është $N(0, \sigma^2)$. Bazuar në këtë, mund të shkruajmë

$$Y = \beta_0 + \beta_1 x_1 + \dots + \beta_p x_p + e$$

ku

$$\beta_0 = \mu_y - \beta_1 \mu_1 - \dots - \beta_p \mu_p$$

dhe e ka shpërndarje $N(0, \sigma^2)$. Vihet re se ekuacioni i mëparshëm ka të njëjtën formë me modelin e regresit të shumëfishtë linear të trajtuar më parë.

2.5.1 Përfundime

1. Koeficienti i korelacionit të shumëfishtë $\rho_{y,x_1,\dots,x_p}$ përcakton masën e lidhjes lineare ndërmjet Y dhe bashkësisë së variablave $\{X_1, \dots, X_p\}$. Koeficienti i korelacionit të shumëfishtë merr vlera $0 \leq \rho_{y,x_1,\dots,x_p} \leq 1$. Një vlerë zero tregon që Y është i pavarur nga bashkësia $\{X_1, \dots, X_p\}$, ndërsa një vlerë 1 tregon varësi perfekte lineare, pra që, Y mund të shprehet saktësisht si një kombinim linear i $X_1, \dots, X_p$.

Për të analizuar fortësinë e modelit të regresit linear të shumëfishtë, përdoret zakonisht ky rregull. Nëse koeficienti i korelacionit është më e vogël se 0.2, lidhja e variablave sipas modelit të regresit konsiderohet shumë e dobët; nëse koeficienti i korelacionit është më i madh se 0.2 dhe më i vogël se 0.4, atëherë thuhet se lidhja sipas modelit të regresit është e dobët, për vlerën e koeficientit nga 0.4 në 0.6 modeli quhet i moderuar. Modeli i regresit linear quhet i fortë dhe shumë i fortë përkatësisht në rastet kur ky koeficient i korelacionit, i llogaritur si mësipërm, ka vlera përkatësisht nga 0.6-0.8 dhe nga 0.8 në 1.

2. Duke zgjidhur ekuacionin për koeficientin e korelacionit të shumëfishtë, marrim:

$$\rho_{y,x_1,\dots,x_p}^2 = \frac{\sigma_y^2 - \sigma^2}{\sigma_y^2} \quad (*)$$

Katrori i koeficientit të korelacionit të shumëfishtë përcakton pjesën e variancës së Y që “shpjegohet” nga relacionin e regresit me $X_1, \dots, X_p$.

3. Koeficienti i korelacionit të shumëfishtë është jonegativ, me përkufizim. Kështu në rast të shpërndarjes normale dypërmasore, ku $(p = 1)$, kemi

$$\rho_{y,x_1} = \rho_{x_1,y} = |\rho_{x_1y}|$$

ku ρ_{x_1y} është koeficienti i korelacionit të thjeshtë ndërmjet X_1 dhe Y .

4. Kur $p = 2$, arsyetimet e mësipërme mund të shkruhen:

$$\mu_{y,x_1x_2} = \mu_y + \beta_1(x_1 - \mu_1) + \beta_2(x_2 - \mu_2)$$

Grafiku i këtij ekuacioni është një plan (që quhet plani i regresit të Y sipas X_1 dhe X_2) në hapësirën e përcaktuar nga boshtet koordinative x_1, x_2 , dhe nga μ_{y,x_1x_2} . Grafiku i ekuacionit për $p > 2$ paraqet një hiperplan në hapësirën $p + 1$ dimensionale të përcaktuar nga $x_1, \dots, x_p$ dhe nga vlera $\mu_{y,x_1,\dots,x_p}$.

5. Korelacioni i shumëfishtë është maksimumi i një korelacioni të thjeshtë mes Y dhe kombinimit linear të $X_1, \dots, X_p$. Për më tepër, $\mu_{y,x_1,\dots,x_p}$ është kombinimi linear që siguron këtë korelacion maksimal. Lidhja ndërmjet koeficientit të korelacionit të shumëfishtë dhe parametrave të regresit $\beta_1, \dots, \beta_p$ do të diskutohet më vonë.

6. Koeficienti i korelacionit të shumëfishtë nuk varet nga shkalla e matjes së ndryshoreve të rastit të $Y, X_1, \dots$, ose X_p .

7. Koeficienti i korelacionit të shumëfishtë mund të shfrytëzohet edhe në një formë tjetër. Kështu, $(1 - \rho_{y,x_1,\dots,x_p}^2)^{1/2}$ është pjesa e i devijimit standard të Y që mbetet “e pa shpjeguar” nga ndryshimet e $X_1, X_2, \dots, X_p$. Për shembull, një koeficient i korrelimit të shumëfishtë 0.9, në dukje shumë i lartë, lë akoma 44% të devijimit standard të Y të pa shpjeguar.

2.5.2 Zbatim 1 (vazhdim)

Nëse i referohemi përsëri "**Efekti I SO₂ në prodhimet bujqësore**", pasi u vu re se për produktet '*grurë*', '*misër*' dhe '*fruta*', nuk kishte dallim midis zonave, u kontrollua nëse sasia e produkteve bujqësore varet nga kushtet abiotike e biotike. Variablat '*distance*', që përcakton distancën në metër te individit nga pusi, '*mbetje urbane*' që kodohet sipas shkallës së trajtimit të mbetjeve urbane, në rastin tonë më i lartë kodi më i ulët niveli i trajtimit të kesaj lloji mbetje; dhe kujdesi ndaj '*ujrave të zeza*', që është koduar në mënyrë të tillë që rritja e kodit tregon përkujdesje të ulët të trajtimit të tyre, u trajtuan me metodën e regresit linear të shumëfishtë dhe u mor modeli I regresi te përcaktuar sipas tabelës 2-2.

Tabela 2-3 Tabela e vlefshmerisë së modelit për 'frutat'

Model Summary

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Change Statistics				
					R Square Change	F Change	df1	df2	Sig. F Change
1	.301 ^a	.091	.067	213.48478	.091	3.743	4	150	.006

a. Predictors: (Constant), distanca shtepi-pusne meter, sistemi i ujrave te zeza, terheqja e mbeturinave, area

Tabela 2-4 Tabela e vlefshmerisë së modelit për 'misri '

Model Summary

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Change Statistics				
					R Square Change	F Change	df1	df2	Sig. F Change
1	.232 ^a	.054	.028	209.55396	.054	2.126	4	150	.050

a. Predictors: (Constant), distanca shtepi-pusne meter, sistemi i ujrave te zeza, terheqja e mbeturinave, area

Afishimet, tabela 2.3 dhe 2.4, tregojnë se kemi të bëjme me modele regresi që duhen marrë parasysh, (p-vlerat për të dy modelet janë 0.006 dhe 0.05) . Por nëse përqëndrohemi tek vlerat e afishuara, vleresimet për R janë 0.35 dhe 0.24, të cilat në baze të vërejtjes 1, flasin për një model të dobët regresi. Pikerisht modelet e shqyrtuara justifikojnë 11% dhe 5.5% të ndryshimeve në prodhshmërinë e bimeve në fjalë.

2.6 Koeficienti i korelacionit të pjesshëm

Këtu diskutohet një tjetër koeficient korelacioni, koeficienti i korelacionit të pjesshëm. Ky përdoret për të matur lidhjen lineare ndërmjet çfarëdo dy prej $p + 1$ variablave $Y, X_1, \dots, X_p$, pas heqjes së 'efektit' të një nënbashkësie, jo boshe, të $p - 1$ variablave të mbetur. Në veçanti, ai përcakton varësinë ndërmjet Y dhe një variabli të pavarur X_m , pas heqjes nga Y , të varësisë lineare të një nënbashkësie të k prej $p - 1$ variablave të pavarur të mbetur $X_i, i = 1, \dots, p, i \neq m$. Kjo varësi lineare e Y ndaj nënbashkësisë së k variablave quhet 'efekt' i nënbashkësisë në fjalë, tek ndryshorja e rastit Y . Teoria e koeficientit të korelacionimit të pjesshëm lidhet me shpërndarjen me kusht.

2.6.1 Shtrimi i problemit

Le të jenë l dhe h , dy prej variablave çfarëdo nga $Y, X_1, \dots, X_p$, dhe le të jetë c një nënbashkësi jo boshe e $p - 1$ variablave të mbetur. Përcaktohen ndryshoret e rastit $Z_1 = l - \mu_{l,c}$ ku $\mu_{l,c}$ pritja me kusht e l , për c e dhënë dhe, $Z_2 = h - \mu_{h,c}$, ku $\mu_{h,c}$ pritja me kusht e h , për c e dhënë. Duhet patur parasysh se Z_1 dhe Z_2 janë ndryshore rasti, sepse ato janë funksione të ndryshoreve të rastit në c . Atëherë koeficienti i korelacionit të pjesshëm të l dhe h , për c e dhënë është:

$$\rho_{lh,c} = \rho_{Z_1 Z_2}$$

ku $\rho_{Z_1 Z_2}$ është koeficienti i zakonshëm i korelacionit ndërmjet Z_1 dhe Z_2 .

Shfaqen zakonisht dy raste të veçanta.

Një rast është kur $l = Y, h = X_m, m = 1, \dots, p$, dhe c është nënbashkësia e $p - 1$ variablave të pavarur të mbetur. Ky korelacion i pjesshëm do shënohet me $\rho_{Y, X_m, c}$.

Rasti i dytë është kur $l = Y, h = X_m$, dhe c është nënbashkësi e k variablave të parë të pavarur $\{X_1, X_2, \dots, X_k\}$, ku $l \leq k < m \leq p$. Koeficienti i korelacionit të pjesshëm në këtë rast, do të shënohet me $\rho_{Y, X_m, X_1 \dots X_k}$. Në përgjithësi në se janë k variabla në c , atëherë koeficienti i korelacionit të pjesshëm thuhet se është i rendit k .

2.6.1.1 Përfundime

1. Koeficienti i korelacionimit të pjesshëm $\rho_{lh,c}$ është percaktues I lidhjes lineare ndërmjet l dhe h , kur vlerat e variablave të paraqitura c s' ndryshojnë. Ajo merr vlera nga -1 në 1; një vlerë zero tregon që lështë e pavarur nga h , kur c mbahet e pandryshuar.

2. Përdoret shumë një identitet ndërmjet koeficientëve të korelacionimit të pjesshëm dhe të shumëfishtë për bashkësinë e variablave $Y, X_1, \dots, X_{k-1}, X_k, k = 2, \dots, p$, që është

$$1 - \rho_{Y, X_1 \dots X_k}^2 = (1 - \rho_{Y, X_1 \dots X_{k-1}}^2)(1 - \rho_{Y, X_k, X_1 \dots X_{k-1}}^2)$$

Që do të thotë se

$$V(Y|X_1, \dots, X_k) = V(Y|X_1, \dots, X_{k-1})(1 - \rho_{Y, X_k, X_1 \dots X_{k-1}}^2)$$

ku $V(Y|X_1, \dots, X_i)$ është variance (dispersioni) me kusht e Y , për $X_1, \dots, X_i$, $i = 1, \dots, p$ të dhënë. Kështu

$$\rho_{y x_k, x_1 \dots x_{k-1}}^2 = \frac{V(Y|X_1, \dots, X_{k-1}) - V(Y|X_1, \dots, X_k)}{V(Y|X_1, \dots, X_{k-1})}$$

Kështu që katrori i koeficientit të korelacionit të pjeshëm është në proporcional me variancën e Y “që shpjegohet” nga shtesa e X_k tek bashkësia $\{X_1, \dots, X_{k-1}\}$.

3. Kemi gjithashtu

$$\rho_{y x_m, c} = \beta_m \left[\frac{V(X_m|c)}{V(Y|c)} \right]^{1/2}, \quad m = 1, \dots, p$$

ku c është tërë nënbashkësia e $p - 1$ variabla të mbetur dhe $V(X_m|c)$ është varianca me kusht e X_m për variabla të dhëna c . Kështu të provosh që $\beta_m = 0$ është e njëjtë sikur të provosh që $\rho_{y x_m, c} = 0$.

4. Korelacioni i pjeshëm mund të llogaritet nga një proces rekrusiv si më poshtë. Në se l, h dhe d paraqesin çdo tre variabla të marrë në bashkësinë $\{Y, X_1, \dots, X_p\}$, atëherë të gjitha korelacionet e pjeshme të rendit të parë jepen prej

$$\rho_{lh, d} = \frac{\rho_{lh} - \rho_{ld}\rho_{hd}}{\sqrt{(1 - \rho_{ld}^2)}\sqrt{(1 - \rho_{hd}^2)}}$$

ku afishimi në të djathtë paraqet korelacionet e thjeshta. Hap pas hapi, duke përdorur formulën rekursive

$$\rho_{lh, cd} = \frac{\rho_{lh, c} - \rho_{ld, c}\rho_{hd, c}}{\sqrt{(1 - \rho_{ld, c}^2)}\sqrt{(1 - \rho_{hd, c}^2)}}$$

ku c është çdo nënbashkësi e variabla të mbetur, mund të fitojmë koeficientët e korelacionit të pjeshëm të çdo rendi.

2.6.2 Vlerësimi dhe testimi i hipotezave për koeficientët e korelacionit të shumëfishtë dhe të pjeshëm

Shqyrtohet tani problemi i llogaritjes dhe interpretimit i koeficientëve të korelacionit nisur nga shfrytëzimi i bazës së të dhënave. Me qenë se teorikisht kërkohet që të gjitha $p + 1$ variablat të jenë variabla të rastit, supozohet se vlerat e vrojtuar $(y_1, x_{11}, \dots, x_{p1}), \dots, (y_n, x_{1n}, \dots, x_{pn})$ janë fituar nga zgjedhja e rastësishme e n individëve nga popullimi normal shumëpërmasor. Për çdo individ, të gjithë $p + 1$ variablat janë matur njëkohësisht. Vlerësuesit e mesatare, variancave dhe kovariancave të këtij popullimi merren nga vlerësimi i mesatare, variancave dhe kovariancave të zgjedhjes. Lehtësisht softet statistikore bejne llogaritjen e mesatares, variancës dhe kovariancës së vektorit të të dhënave, përmes analizës deskriptive. Madje ato llogarisin edhe matricën e kovariancës midis $p + 1$ variabla.

U vu re se që shpërndarja me kusht e Y , për $X_1 = x_1, \dots, X_p = x_p$ të dhënë, çon në modelin e regresit linear të shumëfishtë. Kështu rastin tonë kemi:

$$y_i = \beta_0 + \beta_1 x_{1i} + \dots + \beta_p x_{pi} + e_i, \quad i = 1, \dots, n$$

me e_i të pavarur $N(0, \sigma^2)$. Me qenë se ky ekuacion është identik me ekuacionin e regresit të diskutuar në seksionin e mëparshëm, atëherë vlerësuesit për $\beta_0, \beta_1, \dots, \beta_p$,

dhe σ^2 dhe hipotezat e testit të diskutuara në atë seksion janë të vlefshme këtu. Ato që mbeten për t'u vlerësuar janë koeficientët e korrelacionit të pjesshëm dhe të shumëfishtë.

Diskutojmë tani vlerësimin e koeficientëve të korrelacionit të shumëfishtë dhe atij të pjesshëm. Vlerësuesi i koeficientit të korrelimit të shumëfishtë, i shënuar si $r_{y,x_1,\dots,p}$, zakonisht afishohet nga programet statistikore me shënimin R i *shumefishte* ose Koeficient i korrelacionit të shumëfishtë. Ai mund të llogaritet gjithashtu në analizën e tabelës së variancës, nga :

$$r_{y,x_1\dots x_p} = + \sqrt{\frac{SS_D}{SS_T}}$$

Ai mund të përcaktohet edhe si rrënja katrore, pozitive, e koeficientit të përcaktimit R^2 . Ky vlerësim pikësor për koeficientin e korrelacionit të shumëfishtë është gjithnjë jo negativ.

Për çfarë vlen vlerësimi për koeficientin e korrelacionit?

Së pari, ai përcakton masën e varësisë lineare të Y nga të gjithë variablat e pavarur. Për të kontrolluar hipotezën bazë se, nuk ka varësi lineare, pra që, $H_0: \rho_{y,x_1\dots x_p} = 0$, mund të përdorim të rezultatet e F – *testit*, me qenë se hipoteza e shtruar është ekuivalente me $H_0: \beta_1 = \dots = \beta_p = 0$. Mund të përdoret gjithashtu edhe F – *testi* ekuivalent:

$$F = \frac{n - p - 1}{p} \frac{r_{y,x_1\dots x_p}^2}{1 - r_{y,x_1\dots x_p}^2}$$

p-vlera është sipërfaqja në të djathtë të F nëndensitetin probabilitar të shpërndarjes $F(p, n - p - 1)$.

Së dyti, thamë se katrori i këtij koeficienti të vlerësuar përcakton “pjesën e variacionit të Y që shpjegohet nga regresit linear i Y ndaj $X_1, \dots, X_p$ ”.

Vëzhgimet që mungojnë, në rastin shumëpërmasor

Në analizat që janë me një variabël (p.sh. testet t), trajtimi më arsyeshëm i rastit kur vlera mungon te X (variabli që do analizohet) është heqja e këtij vrojtimi nga zgjedhja. Situata është e ndryshme kur kemi të bëjme me analiza që janë të natyrës shumëpërmasore, p.sh. kur p variablat $X_1, \dots, X_p$ vëzhgohen në çdo vrojtim. Në këtë situatë në se një prej variablave u mungon një vlerë përshembull X_1 , nuk është e nevojshme të heqim këtë vrojtim nga analiza dhe të humbim informacionin për $X_2, \dots, X_p$. Me që ekuacioni i analizës së regresit të shumëfishtë sikurse edhe procedurat multivariate janë funksione të vektorit mesatar μ dhe të matricës së kovariancës Σ , mund ta lëmë këtë rast në zgjedhje dhe të përdorim informacion që kemi në llogaritjen e vlerësimit të vektorit mesatar $\bar{x}$ dhe matricës së kovariancës S .

Tani të diskutojmë metoda të ndryshme për vlerësimin e μ dhe Σ (apo njëlloj të matricës së korrelimit R), kur mungojnë disa vlera. Le të jenë n_i numri i rasteve kur X_i është pa mungesa, n_{ij} numri i rasteve kur të dyja X_i dhe X_j janë pa mungesa, dhe n_c numri i rasteve kur të gjitha $X_1, \dots, X_p$ janë pa mungesa ($n_i, n_{ij}, n_c \leq n$ - vëllimi I zgjedhjes, $i, j = 1, \dots, p, i \neq j$). Atëherë disa metoda për të përcaktuar $\bar{x}$ dhe S (ose $\hat{R}$) janë si më poshtë:

Metoda 1

Vlerësimi i $\bar{x}$ dhe S duke përdorur të gjitha rastet e n_c . Kjo quhet metoda e eliminimit të rastit.

Metoda 2

Përdorimi i n_i vëzhgime e llogaritet $\bar{x}_i$. Zëvendësojmë $\bar{x}_i$ për çdo observim të munguar në X_i . Më pas përdoret vërtetimi i plotësuar i n vëzhgimeve, për të marrë $\bar{x}$ dhe S . Kjo quhet metoda e zëvendësimit mesatar.

Metoda 3

Përdoren n_i vëzhgimet e llogaritet $\bar{x}_i$ dhe s_i^2 , dhe përdoren i n_{ij} observimeve për të fituar s_{ij} . Këto statistika përcaktojnë $\bar{x}$ dhe S .

Metoda 4

Përdoren n_i vëzhgimet e llogaritet $\bar{x}_i$ dhe s_i^2 , dhe përdoren i n_{ij} vëzhgimet e llogaritet r_{ij} . Më pas s_{ij} llogaritet si $s_{ij} = r_{ij}(s_i)(s_j)$, që është i vetmi ndryshim nga metoda 3. Metodat 3 dhe 4 quhen dhe metodat e fshirjes në çifte.

Metoda 5

Përdoren të gjitha n_c vëzhgimet e llogaritet ekuacionin e regresit të një variabli me të gjithë variablat e tjerë. Për shembull, le të jetë ky ekuacion regresi $X_1 = f(X_2, \dots, X_p)$. Më pas një vlerë e munguar e X_1 për rastin e j -të vlerësohet $\hat{x}_{1j} = f(x_{2j}, \dots, x_{pj})$. Ekuacione të ngjashme mund të gjenden për $X_2, \dots, X_p$. Tërësia e vërtetimeve që rezultojnë tashmë e plotësuar, përdoret për vlerësimin e $\bar{x}$ dhe S .

Metoda 6

Përdoret metoda 5 me ca ndryshime. Për shembull X_1 mund të parashikohet vetëm nga një variabël i vetëm $X_2, \dots, X_p$ që është më i koreluar me X_1 , ose X_1 mund të parashikohet nga nënbashkësia e $X_2, \dots, X_p$. Metodat 5 dhe 6 quhen metodat e zëvendësimit të regresit.

Disavantazhi i parë i përdorimit të një prej këtyre metodave është që vetitë statistikore të tyre nuk njihen përveçse në raste të veçanta. Për më tepër këto metoda shpesh prodhojnë rezultate të shmangura. Rekomandimet tregojnë se vërtetimi dhe/ose variablat duhet të eliminohen, sipas rastit. Kështu vërtetimet dhe variablat duhet të eliminohen për të maksimizuar numrin e vërtetimeve të plota. Kështu, nëse një vërtetim ka shumë variabla që mungojnë, duhet të eliminohet gjithë vërtetimi. Nga ana tjetër, nëse një variabël mungon në shumë vërtetime, ky variabël eliminohet. Më pas përdoren procedurat e zakonshme të katrorëve më të vegjël ose procedurat shumëpërmasore.

2.6.2.1 Vërejtje

1. Shumica e paketave kompjuterike kanë si opsion 'default' metodën e eliminimit të vërtetimit.
2. Në disa programe ekziston si opsion edhe fshirja në çift. Ky opsion mund të përdoret kur një vërtetues ka shumë variabla që i mungojnë pak vlera, dhe kur metoda e eliminimit të vërtetimit do të reduktojë numrin e vërtetimeve shumë. Përdoruesi duhet të ketë parasysh që mund të lindin dhe situata jo të vërteta në këtë rast nga llogaritjet

(si p.sh. shumë katrorësh ose vlera testesh F negative). Në rastet kur përdoret opsioni nuk aplikohet në mënyrë rigorozë teoria e nderthurjes statistikore.

2.7 Regresi i Shumëfishtë Stepwise

Regresioni i shumëfishtë shërben për të gjetur nënbashkësinë më efektive të variablove përcaktues, në parashikimin e variablit të varur.

Variablat shtohen në ekuacionin e regresit duke përdorur kriterin e maksimizimit të R^2 . Kur shtimi i variabave nuk sjell ndonjë ndryshim statistikor të rëndësishëm në R^2 , atëherë analiza e regresit të shumëfishtë Stepwise ndalon.

2.7.1 Shtrimi matematik i Regresi të Shumëfishtë Stepwise

Në shumë raste, gjatë aplikimit të regresit, vrojtuesi nuk ka informacion të mjaftueshëm rreth shkallës së rëndësisë së variablove të pavarur $X_1, X_2, \dots, X_p$ në parashikimin e variablit të varur Y . Kontrolli i $H_0: \beta_i = 0$ për çdo variabël X_i , $i = 1, \dots, p$, nuk na siguron këtë informacion.

Meqë statistika që përcakton 'efektivitetin' e një bashkësie variablash të pavarur në parashikim, është koeficienti i korelacionit të shumëfishtë, një zgjidhje për problemin e mësipërm është llogaritja e koeficientit të regresit të Y ndaj të gjitha nënbashkësive të mundshme të variablove të pavarur dhe, më pas, të zgjidhet nënbashkësia më e mirë, duke shfrytëzuar këto koeficientë.

Për çdo nënbashkësi me k , elementë, për $k = 1, \dots, p$, përcaktohet nënbashkësia S_k që përmban koeficientët më të mëdhenj të korelacionit të shumëfishtë. Për regresin e Y ndaj nënbashkësisë S_1 , kontrollohet hipoteza që përfshirja e $p - 1$ variablove të mbetur nuk përmirëson parashikimin e Y , me hipotezë alternative: verëhet përmirësim. Për kontrollin e kësaj hipoteze përdoret testi statistikor i kontrollit të koeficientëve

$$F = \frac{(SS_R' - SS_R)/m}{MS_R} \quad \text{pra} \quad \text{testi:}$$

Nëse hipoteza hidhet poshtë, atëherë për regresin e Y në nënbashkësinë S_2 , kontrollohet nëse përfshirja e $p - 2$ variablove të mbetur nuk përmirëson parashikimin e Y . Përsëritet testi në mënyrë të njëpasnjëshme, deri sa për nënbashkësite S_m , $1 \leq m \leq p$, hipotezat që $p - m$ variablat e mbetur nuk përmirësojnë parashikimin e Y , pranohen. Atëherë S_m është nënbashkësia më e mirë e variablove në parashikimin e Y me qenë se: (a) ajo siguron koeficientin e korelacionit të shumëfishtë më të madh' ndërmjet të gjithë nënbashkësive me vëllim m ; dhe (b) përfshirja e $p - m$ variablove të mbetur nuk përmirëson ndjeshëm parashikimin e Y . Nëse kjo zgjidhje nuk është e vetme, natyra e problemit mund të sugjerojë nënbashkësinë më të përshtatshme.

Kur numri i variablove të pavarur është i madh, aplikimi i kësaj metode s'është dicka praktike, madje edhe nëse kemi kompjuter me shpejtësi të madhe, përcaktimi i nënbashkësisë më të mirë duke përdorur këtë procedurë është i vështirë. Për shembull kur $p = 5$, atëherë janë $5+10+10+5+1=31$ ekuacione regresi, dhe kur $p = 10$, janë $2(10+45+120+210)+252+1=1023$ ekuacione regresi. Në përgjithësi, janë të nevojshëm $2^p - 1$ ekuacione regresi. Koha dhe shpenzimet e nevojshme na shtyjnë të gjejmë një rrugë tjetër për këtë problem.

Një zgjidhje është teknika e regresit hap pas hapi, në të cilën variablat e pavarur $X_1, X_2, \dots, X_p$ futen njëri pas tjetrit në ekuacion, sipas disa kriterëve të parapërcaktuar. Sapo një variabël futet në ekuacion, ai mund të ndërrohet me një variabël që është në modelin regresit, ose mund të hiqet fare nga ekuacioni. Bashkësia e kriterëve që përcakton se si një variabël futet, ndërrohet ose hiqet quhet procedura stepwise. Diskutohen katër procedura stepwise.

Nga një procedurë stepwise, fitohet një listë e renditur parashikuesish. Për shembull në se $p = 5$ kjo listë mund të jetë X_2, X_5, X_1, X_4 dhe X_3 . Për të përcaktuar ekuacionin më të mirë nga kjo listë zgjidhen $m \leq p$ variablat e parë kështu që (i) ata parashikojnë Y sa më mirë që të jetë e mundur, dhe (ii) m është sa më e vogël që të jetë e mundur. Me fjalë të tjera një nënbashkësi variablash të varur zgjidhet nga lista e renditur, e tillë që të ketë aftësi parashikuese të lartë. Në shembullin e mësipërm mundet që nënbashkësitë të konsistojnë vetëm në X_2 dhe X_5 dhe, se regresi me X_2 dhe X_5 është pothuaj aq i mirë sa regresi në X_2, X_5, X_1, X_4 dhe X_3 . Procedura për përcaktimin e m njihet më emrin procedura Iterative. Mëposhtë paraqiten tri rregulla të tilla ndërprerje.

2.7.2 Procedurat iterative

Le të supozojmë se kemi një bashkësi variablash të pavarur $X_1, X_2, \dots, X_p$ të cilët janë të gjithë për tu shqyrtuar për parashikimin e Y . Le të supozojmë gjithashtu dhe që kemi një zgjedhje të rastit me vëllim n . Procedura ndërprerjes standarde, që do diskutohet fillimisht, konsiston në një rregull për të futur variablat dhe një rregull për heqjen e variablave (ndërrimi i variablave nuk përfshihet në procedurën e iterative standard). Siç do të shohim procedurat e tjera rrjedhin nga kjo procedurë.

1. Procedura standarde iterative (Metoda F).

Variablat futen ose hiqen duke përdorur një test statistikor, t testin, i cili është prezantuar për kontrollin e hipotezave, koeficienti i korrelacionit të pjesshëm janë zero. Programet kompjuterike paraqesin katrorët e kesaj statistike, e cila ka shpërndarje F (me shkallë lirie të përcaktuar më poshtë).

Më pas supozojmë se një nënbashkësi c që përmban k variabla është futur në ekuacion, $k = 0, 1, \dots, p - 1$. Më pas testi F teston nëse një variabël X (jo nga c) përmirëson në mënyrë të ndjeshme parashikimin e Y në krahasim me atë të variablave në c , p.sh. $H_0: \rho_{yx.c} = 0$, duke përdorur statistikën e testit F

$$F_{yx.c} = \frac{r_{yx.c}^2(n - k - 2)}{1 - r_{yx.c}^2}$$

e cila nën hipotezën zero ka shpërndarje $F(1, n - k - 2)$.

Në mënyrë të ngjashme merret edhe një test statistikor po me shpërndarje Fisher, që teston procesin e heqjes, pra nëse një variabël X në një nënbashkësi c që përmban k variabla nuk nevojitet më. Pra nënbashkësia e zvogëluar c' që përmban $k' = k - 1$ variabla, parashikon Y po aq mirë sa nënbashkësia c . Kështu, ne testojmë $H_0: \rho_{yx.c'} = 0$ duke përdorur testin statistikor me shpërndarje Fisher $F(1, n - k' - 2)$:

$$F_{yx.c'} = \frac{r_{yx.c'}^2(n - k' - 2)}{1 - r_{yx.c'}^2}$$

Siç do të shohim rregulli Iterativ sipas procedurës standard, përfshin zgjedhjen e një vlere minimale F për përfshirjen e variablave. Disa paketa programuese marrin një vlerë të paracaktuar 4.0. Për heqjen e variablave, zgjidhet gjithashtu një minimum i testit F për të hequr (më i vogël se minimumi F për përfshirjen; disa paketa programesh marrin minimumin $F = 3.9$).

Hapat e procedurës standard më të detajuara janë:

Hapi 0:

Testohet $H_0: \rho_{yx_i} = 0, i = 1, \dots, p$. (Koeficienti i korelacionit të thjeshtë është koeficient i korelacionit të pjesshëm me $k = 0$ dhe $c =$ bashkësi boshe.) Për këtë llogariten koeficientët i korelacionit të thjeshtë r_{yx_i} dhe vlerat e F testit, F_{yx_i} , për $i = 1, \dots, p$. Testi statistikor jepet nga

$$F_{yx_i} = \frac{r_{yx_i}^2(n-2)}{1-r_{yx_i}^2}$$

që përfitohet nga (*) me $k = 0$. Ai ka shpërndarje Fisher, F me 1 dhe $n - 2$ shkallë lirie.

Hapi 1

Variabli X_{i_1} që ka vlerën e testit F më të madh, (ose ndryshe, katrorët e koeficientëve të korrelacionit me Y më të mëdhenj) zgjidhet si parashikuesi më i mirë i Y . Nga ekuacioni i katrorëve më të vegjël, bëhet analiza e tabelës së variancës, dhe llogaritet koeficienti i korrelacionit të shumëfishtë $r_{y.x_{i_1}} = |r_{yx_{i_1}}|$. Në këtë hap, testi F , që shërben për të hequr X_{i_1} , është e njëjtë me vlerën e F -testit për të futur për X_{i_1} . Më pas llogariten koeficientët e korelacionit të pjesshëm $r_{yx_i.x_{i_1}}$ dhe F -testet për përfshirjen e çdo njerit prej tyre

$$F_{yx_i.x_{i_1}} = \frac{r_{yx_i.x_{i_1}}^2(n-3)}{1-r_{yx_i.x_{i_1}}^2}$$

për $i = 1, \dots, p, i \neq i_1$, për çdo variabël që nuk përfshihet në ekuacionin e regresit. Kjo statistikë F ka 1 dhe $n - 3$ shkallë lirie, dhe kontrollon hipotezat $H_0: \rho_{yx_i.x_{i_1}} = 0$ për $i = 1, \dots, p, i \neq i_1$. Nëse të gjitha F për të përfshirjen e variablave të tjerë janë më të vegjël se minimum i F për përfshirjen e variablit, atëherë ekzekutohet Hapi S. Në rast kundërt vazhdojmë me Hapin 2.

Hapi 2

Variabli X_{i_2} që ka vlerën e testit F për përfshirjen e variablit më të madh (që do të thotë, katrorin e korelacionit të pjesshëm më të madh me Y për X_{i_1} të dhënë) zgjidhet si parashikuesi më i mirë i Y nëse, më parë, ka qenë zgjedhur X_{i_1} . Llogariten nga ekuacioni i katrorëve më të vegjël, analiza e tabelës së variancave, dhe koeficienti i korelacionit të shumëfishtë $r_{y.x_{i_1}x_{i_2}}$. Gjithashtu, llogariten $F_{yx_{i_1}.x_{i_2}}$ dhe $F_{yx_{i_2}.x_{i_1}}$. Këto statistika me 1 dhe $n - 3$ shkallë lirie përcaktohen si më poshtë:

$$F_{yx_{i_1}.x_{i_2}} = \frac{r_{yx_{i_1}.x_{i_2}}^2(n-3)}{1-r_{yx_{i_1}.x_{i_2}}^2} \quad \text{dhe} \quad F_{yx_{i_2}.x_{i_1}} = \frac{r_{yx_{i_2}.x_{i_1}}^2(n-3)}{1-r_{yx_{i_2}.x_{i_1}}^2}$$

Ato testojnë respektivisht hipotezat zero $H_0: \rho_{yx_{i_1}.x_{i_2}} = 0$ dhe $H_0: \rho_{yx_{i_2}.x_{i_1}} = 0$. Së fundmi, llogariten koeficienti i korelacionit të pjesshëm $r_{yx_{i_1}.x_{i_2}}$, dhe vlera e testit F me 1 dhe $n - 4$ shkallë lirie për $i = 1, \dots, p, i \neq i_1, i \neq i_2$ për përfshirjen e x_{i_2}

$$F_{yx_{i_1}.x_{i_2}} = \frac{r_{yx_{i_1}.x_{i_2}}^2(n-4)}{1-r_{yx_{i_1}.x_{i_2}}^2}$$

të cilat bëjnë të mundur kontrollin e hipotezës $H_0: \rho_{yx_{i_1}.x_{i_2}} = 0$. Nëse të gjitha F për përfshirjen e x_{i_2} janë më të vogla se minimumi i vlerës së statistikës F për përfshirjen x_{i_1} , atëherë ekzekutohet Hapi S. Në të kundërt, kalohet në Hapin 3.

Hapi 3

(a) Le të shënojmë me L bashkësinë e l variablave të pavarur që janë përfshirë në ekuacionin e regresit. Nëse ndonjë nga vlerat e statistikës F për të hequr variablin e përfshirë në L , është më i vogël se minimumi i paracaktuar i F për të hequr, atëherë variabli me F -vlere për të hequr më të vogël nuk përfshihet më në L , dhe ekzekutohet hapi (3b) ku l zëvendësohet me $l - 1$. Nëse të gjitha vlerat e F testit për të përfshirë variablat që tashmë s'janë në L , janë më të vogla se minimumi i F të paracaktuar të përfshirjes, atëherë ekzekutohet Hapi S. Në të kundërt variabli me vlerën e F përfshirjes më të madh merret dhe, i shtohet L -së kështu që l zëvendësohet me $l + 1$.

(b) llogariten për Y dhe variablat në L , ekuacioni i katrorëve më të vegjël, tabela e analizës së variancave dhe koeficienti i korrelimit të shumëfishtë $r_{y.l}$. Gjithashtu llogaritet vlerat e testeve F për të hequr parametra, $F_{yx_{i_j}.(l-1)}$ ndërmjet Y dhe X_{i_j} në L nëse $l - 1$ variabla mbesin në L . Kjo ka 1 dhe $n - l - 1$ shkallë lirie dhe testohet hipoteza $H_0: \rho_{yx_{i_j}.(l-1)} = 0$. Së fundmi llogariten koeficienti i korrelacionit të pjesshëm $r_{yx_{i_l}.l}$ dhe testi F për përfshirjen $F_{yx_{i_l}.l}$ për Y dhe X_{i_l} jo në L për variabla të dhënë në L . Kjo statistikë ka 1 dhe $n - l - 2$ shkallë lirie dhe testohet $H_0: \rho_{yx_{i_l}.l} = 0$ për X_{i_l} jo L , $i = 1, \dots, p$.
Hapat 4,5,...

Përsëritet Hapi 3 në mënyrë në çdo iteracion. Hapi S ekzekutohet (a) kur vlerat e testit F për përfshirjen e variablave të , që s'janë në L , janë më të vogla se minimumi i statistikës F për përfshirjen e variablit, (b) kur të gjithë variablat janë përfshirë dhe vlera e testit F për të hequr për variablat e përfshirë është më e madhe se F për heqjen e variablit, vlerë kjo e paracaktuar në fillim, ose (c) kur numri i variablave të përfshirë në regresin e shumëfishtë është $n - 2$.

Hapi S

Zakonisht, me kërkesë të përdoruesit, jepet një tabelë përmbledhëse. Për çdo hap afishohen: numri i hapit, variabli i përfshirë apo i hequr, vlera e F testeve për përfshirjen dhe heqjen e variablit në fjalë, dhe koeficienti i korrelacionit të shumëfishtë të Y me variablat e përfshirë në këtë model.

2. Procedura Standarde me Zëvendësim (FSWAP)

Kjo procedurë përdor të njëjtat rregulla për heqjen dhe shtimin e variablave si procedura e mësipërme iterative, përveç që në çdo hap, mund të bëhen ndërrime ndërmjet variablave në ekuacion. Kjo procedurë është si të thuash më e mira, sipas variantit që dëshirojmë, dhe ekuacionit të regresit të marrë sipas metodës më sipërme. Duke përdorur këtë procedurë, një variabël në ekuacion ndërrohet me një variabël që nuk është në ekuacion, nëse ndërrimi rrit koeficientin e korelacionit të shumëfishtë (madje edhe nëse nuk siguron ndonjë rritje të ndjeshme). Në një hap të caktuar, nëse janë k variabla në nënbashkësinë c të cilët janë përfshirë në ekuacion, procedura iterative (i) heq variabla nga c duke përdorur testin F për heqje, (ii) ndërrimin e një variabli në c me një variabël jo në c , dhe (iii) shtimin e një variabli në c duke përdorur testin F përfshirës.

3. Metoda e Korelacionit të Shumëfishtë (Metoda R)

Kjo procedurë përdor të njëjtin rregull për përfshirjen e variablave në ekuacionin e shumëfishtë, por përdor një rregull tjetër për heqjen e variablave. Në një iteracion të caktuar, ajo përdor rregullin R^2 , sipas të cilit një variabël hiqet nga modeli i regresit linear, në se siguron një katror të koeficientit të korelacionit të shumëfishtë më i madh. Është e mundur të arrihet rritje e R^2 , meqenëse R^2 nuk është vetëm një funksion i madhësive të përcaktuara (siç janë, n dhe s_y^2), por edhe të dy madhësive që ndryshojnë (s^2 mesatares së katrorit të mbetjes dhe p numri i variablave në ekuacion). Kështu është e mundur për këto sasi të ndryshojnë, në atë mënyrë që një variabël më pak të japë një R^2 më të madhe. Kjo procedurë përfundimisht konsiston sipas rradhës (i) heqjen e një variabli, duke përdorur rregullin R^2 , dhe (ii) shtimin e një variabli, duke përdorur testin F të përfshirjes.

4. Metoda e Korelacionit të Shumëfishtë me Ndërrim (RSWAP)

Kjo procedurë është e njëjtë me metodën R, por përdor edhe rregullin i ndërrimit. Rendi i kësaj procedure është: (i) heqja e variablit duke përdorur rregullin R^2 , (ii) ndërrimi i dy variablave për të rritur R^2 , dhe (iii) shtimi i një variabli duke përdorur testin F të përfshirjes.

2.7.2.1 Përfundime

Shumica e paketave të programeve të regresit stepwise ka një opsion për ndërhyrjen në procesin e gjetjes së variablave më me interes në ekuacionin e regresit. Për çdo variabël përdoruesi specifikon, me anë të një parametri që fut, *le ta quajmë nivel detyrimit*, një instruksion që përcakton nëse variabli do të futet pavarësisht nga vlera e testit F të përfshirjes. Kështu përdoruesi ka kontroll mbi variablat (në kundërshtim me zgjedhjen statistikore përshkruar më sipër), dhe në këtë mënyrë variablat që kanë vlerë më të madhe në parashikim mund të futen të parët. Procedura stepwise përfshin më pas variablat që kanë një nivel rëndësie të vogël.

2.7.3 Rregullat e Ndërprerjes

Do paraqesim tri rregulla për përcaktimin e numrit të variablave (predictors) për gjetjen e ekuacionit më të mirë të regresit. Rregulli standard, i cili është i pranishëm në shumicën e paketave të programeve statistikore, kontrollon numrin e variablave që futen në ekuacion me anë të një parametri që përcaktohet nga përdoruesi dhe, që

quhet *minimum vlerë e testit F të përfshirjes*. Siç thamë më sipër ky parametër është njëvlerë minimum e statistikes Fisher që i korespondon nivelit të rëndësisë α , $F_{1-\alpha}(1, v)$. α rekomandohet të jetë 0.15, me gjithë se shumë përdorues vendosin α në 0.05.

1. Rregulli Standard i Ndërprerjes

Rregulli Standard I Ndërprerjes për të fituar bashkësinë më të mirë të parashikuesve H mund përcaktohet lehtë nga tabela përmbledhëse e printuar në Hapin S. Në një modelin iterativ, vlera e testit F të përfshirjes për çdo variabël të futur krahasohet me vlerën e F për përfshirje të përcaktuar më parë.

Më të detajuara hapat ne procedurat kompjuterike janë:

- (a) Në hapin 1, futet një variabël X_{i_1} . Nëse vlera e testit F përfshirës është jo e konsiderueshme, d.m.th. F përfshirjes $<$ *minimum i F të përfshirjes se paracaktuar*, atëherë përfshirja e variablit në fjale në regres është e pa kuptimtë, dhe përdoruesi duhet të kërkojë metoda të tjera për të analizuar të dhënat. Në të kundërt, $H = \{X_{i_1}\}$
- (b) Në Hapin 2, futet një variabël X_{i_2} . Nëse vlera e testit F të përfshirjes së tij $<$ *minimum i F të përfshirjes së paracaktuar*, atëherë H mbetet vetëm me variablin X_{i_1} , dhe regresi më i mirë është ai i Hapi 1. Në të kundërt, $H = \{X_{i_1}, X_{i_2}\}$.
- (c) Për çdo hap të mëtejshëm, në se futet një variabël, krahasohet vlera e testit F përfshirës me *minimumin e F përfshirës*. Në se F përfshirëse e variablit është e konsiderueshme, shtohet variabli në H dhe shkohet në hapin tjetër. Nëse jo, ndërpritet shtimi i H dhe merret, si ekuacioni më i mirë i regresit, ai i marrë në hapin e mëparshëm.

2. Rregulli i Ndërprerjes Bazuar në Ndryshimet në R^2

Kjo procedurë kërkon që përdoruesi të paracaktojë me kujdes vlerën '*minimale të testit F të përfshirjes*' dhe vlerën '*minimale të F për të hequr variablin e përfshirë*'. Vlera '*minimale e testit F të përfshirjes*' mund të zgjidhet në mënyrë të tillë që çdo variabli që parashikohet të jetë i nevojshëm në parashikimin t'i japim shanse 50-50. Kështu, vlera '*minimale e testit F përfshirjes*' duhet të jetë $F_{0.50}(1, n - p - 1)$. Në anën tjetër, vlera '*minimale e testit F për të hequr variablin e përfshirë*' mund të zgjidhet në mënyrë të tillë që, të kemi shanse të vogla që një variabël i përfshirë, të hiqet. Kështu zgjidhet vlera '*minimale të F për të hequr variablin e përfshirë*' shumë i vogël, p.sh. 0.01. Më pas vëmendja përqëndrohet tek variablat që janë në tabelën përmbledhëse në hapin e fundit. Le të jetë L bashkësia e l variablave, $l \leq p$ dhe, le të jetë $r_{y,l}$ koeficienti i korelacionit të shumëfishtë ndërmjet Y dhe variablave në L . (Nëse një apo më shumë variabla hiqen, është e nevojshme të llogaritet serisht $r_{y,l}$). Le të jetë H bashkësia e h variablave në ekuacionin e regresit në një hap të ndërmjetëm. Procedura bën testimin e hipotezës $H_0: \rho_{y,h} = \rho_{y,l}$ duke përdorur

$$F = \frac{n - l - 1}{l - h} \frac{r_{y,l}^2 - r_{y,h}^2}{1 - r_{y,l}^2}$$

Nëse H_0 është e vërtetë, testi F ka shpërndarje Fisher me $l - h$ dhe $n - l - 1$ shkallë lirie. Kryhet ky test në mënyrë të njëpasnjëshme, për çdo iteracion derisa fitohet vlera e parë, e testit F josisignifikative. Supozojmë se kjo ndodh p.sh. në

Iteracion 3, ku Y është shprehur me një ekuacion regresi të një bashkësi H të h variablave. Më pas në se në Iteracion 4 duke dashur të përfshijmë variabla te rinj në ekuacion regresi, nëse vlerat e F testit janë josignifikative atëherë ndalohej së shtuari H -në dhe marrim si ekuacionin më të mirë të regresit, atë të fituar në Iteracion 3. Nëse testi që kryhet synon të heq variablat ekzistues, atëherë variabli hiqet, kryejmë kontrollin e mësipërm për Iteracion 4. Në se testi përfshirës është signifkativ, atëherë H është bashkësia e variablave e fituar në Iteracion 4, dhe marrim si regresin më të mirë atë të Iteracion 4. Në se vlera e F testit është josignifikative, përcaktojmë H me $h - 1$ variabla si bashkësinë e variablave që marrin pjesë në këtë ekuacion regresi.

3. Rregulli i Ndërprerjes Bazuar në Mesataren e Katrorit të Gabimit të Pakushtëzuar

Bendel dhe Afifi^[6] (1976) rekomandojnë një rregull ndalimi alternativ. Ata propozojnë të kontrollohet hipotezat që gabimi i pakushtëzuar i katrorit mesatar (UMSE)^[7] nuk ulet nga një hap në tjetrin. Kjo madhësi përcaktohet si $UMSE = E(Y - \hat{Y})^2$, ku vlera e parashikuar e Y , $\hat{Y}$, fitohet nga tërësia e vlerave $Y, X_1, \dots, X_p$ të cilët mendohet se kanë shpërndarje normale shumëpërmasore. Vlerësimi i UMSE në një hap të dhënë behet nga statistika:

$$\widehat{UMSE}(q, n) = \frac{n^2 - n - 2}{n(n - q - 2)} MS_R = \frac{(n-1)(n^2 - n - 2)(1 - r_{y.c}^2)s_y^2}{n(n - q - 1)(n - q - 2)},$$

ku q është numri i variablave në ekuacion dhe MS_R është mbetja e katrorit mesatar në hapin e dhënë.

2.8 Zbatim 2

Në punimin me teme: **Cfarë ndikon në memorje**, ku synohet të vlerësohet ndikimi në memorjen e adoleshentëve të disa variablave si, mënyra e të ngrënit, konsumimi apo jo i alkolit, sasia e orëve të fiskulturës, (aktiviteteve sportive) dhe gjinia. Përcaktues për memorjen, në studimin në fjalë, është marrë variabli 'nrwords', që përcakton numrin e fjalëve që mund të memorizojë një adoleshent pas leximit të 40 fjaleve për 5 minuta. Studimi i ndikimit të 5 variablave në fjale, në variablin tone, u mendua të behej me ane të modelit të regresit linear të shumëfishtë. Por për ndertimin e këtij modeli, një ndër kushtet që duhet të plotësohet është dhe ai që, variablat që marrin pjesë në model të jenë dy e nga dy të pavarur. Për të kontrolluar këtë u shkruajt dhe u ekzekutua kodi i mëposhtëm në SPSS :

PARTIAL CORR

/VARIABLES=sportlevel duhanpojo alkolpojo gjinia BY nrtotalIfjaleve
 /SIGNIFICANCE=TWOTAIL
 /STATISTICS=DESCRIPTIVES
 /MISSING=LISTWISE.

Afishimet paraqiten në tabelen 2-5.

Tabela 2-5 Tabela e koeficientëve të korelacionit

Correlations

Control Variables			sportlevel	duhan po/jo	alkol po/jo	gjinia f/m
nr total I fjaleve	sportlevel	Correlation	1.000	.090	.091	.272
		Significance (2-tailed)	.	.121	.116	.000
		df	0	295	295	295
	duhan po/jo	Correlation	.090	1.000	.444	.325
		Significance (2-tailed)	.121	.	.000	.000
		df	295	0	295	295
	alkol po/jo	Correlation	.091	.444	1.000	.263
		Significance (2-tailed)	.116	.000	.	.000
		df	295	295	0	295
	gjinia f/m	Correlation	.272	.325	.263	1.000
		Significance (2-tailed)	.000	.000	.000	.
		df	295	295	295	0

Koeficientat e korelacinit të afishuara mes çdo dy variablave në shqyrtim tregojnë se çdo dy prej variablave "alkol", "duhan", e "sport level", mund të konsiderohen të pavarur linearisht. Meqënëse ky kusht plotësohet mund të flasim për mundësinë e ndërtimit të një modeli regresi me këta variabla. Ndërsa vërehet që ka njëfarë korelacioni të dobët të gjinisë me variablat e tjerë në shqyrtim.

2.9 Regresi i Shumëfishtë Hirarkik

Metoda në fjalë shihet në programet kompjuterike si regresi i shumëfishtë hirarkik. Ajo përdoret zakonisht në rastet, kur gjatë studimit diskriptiv të të dhënave, vërehet varesi e variablit në studim nga një variabël tjetër i pavarur, por ne nuk jemi të interesuar aq shumë për këtë varësi. Interesi ynë është përqëndruar më shumë në kontrollin e varësisë së variablit në studim nga variabla të tjerë.

Në regresin e shumëfishtë hirarkik futja e variablave të pavarur bëhet në dy faza. Në fazën e parë futen variablat e pavarur që ne duhet t'i kontrollojmë më parë në lidhjen e regresit. Ndërsa në fazën e dytë futen ato variabla të pavarur që ne duhet t'i shqyrtojmë, pasi varësia e parë lineare merret parasysh. Vlerësimi i ndikimit të variablave, me të cilën punohet në fazën e dytë, bëhet me anë të testit statistikor që ndertohet duke pasur parasysh R^2 e marrë nga faza e parë.

2.9.1 Zbatim 1(vazhdim)

Nëse i referohemi sërish zbatimit "Efekti I SO_2 në prodhimet bujqësore", vihet re përdorimi i kësaj metode studimi të të dhënave. Në këtë zbatim u vu re se ndotja industriale ndikon në prodhueshmërinë e produkteve bujqësore jonxhes, patates, perimes dhe fasules.

Mesatarja e prodhimit të jonxhes është statistikisht më e lartë në zonën e ndotur, derisa diferenca e studiuar e vlerës së komponentit në fjalë në zonën e ndotur nga zona e kontrollit është negative.

Mesatarja e patates, perimes dhe fasules është statistikisht më e ulët në zonën e ndotur.

Por, sasia e produkteve bujqësore dihet që mund të varet nga kushtet abiotike të lartëpërmendura. Për të testuar nëse këto tri variablat e lartë përmendura ndikojnë në parashikimin e sasisë së prodhimit të patates, fasules, perimeve dhe jonxhes, është përdorur metoda e regresit hirarkik të shumëfishtë⁽⁵⁾.

Si variabli i pavarur i listuar më parë, the independent variable listed first, u përdor variabli zone, dhe si variabla të shtuar, the added variables, distance, ujrat e zeza dhe mbetjet urbane. Përgjithsisht interesi nuk qëndron tek varësia e variablave të kontrollit 'area' ndaj variablave të varura, ndaj jemi në kushtet e përdorimit të regresit hirarkik.

Afishimet paraqesin rezultatet e t-testeve, nëse gjatë ndertimit të hiperplaneve të regresit nuk I marrim parasysh faktorët abiotik, dhe tregojnë efekt të faktorit 'sistem i ujrave të zeza'

Tabela 2–6 Afishimet e t-testeve dhe koeficientëve të regresit

Coefficients^a

Model	Unstandardized Coefficients		Standardized Coefficients	t	Sig.
	B	Std. Error	Beta		
1 (Constant)	-.301.111	1619.402		-.186	.853
area	2405.556	915.482	.208	2.628	.009
2 (Constant)	-9952.987	7017.577		-1.418	.158
area	5846.296	2787.910	.505	2.097	.038
sistemi i ujrave të zeza	267.888	1363.239	.016	.197	.044
terheqja e mbeturinave	1375.448	521.163	.248	2.639	.009
distanca shtëpi-pusnetë	3.503	4.381	.180	.800	.425
meter					

a. Dependent Variable: jonxhe

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	229.364	22.768		10.074	.000
	area	-59.364	12.871	-.349	-4.612	.000
2	(Constant)	-14.914	98.412		-.152	.880
	area	43.286	39.097	.255	1.107	.270
	sistemi i ujrave të zeza	5.775	19.118	.023	.302	.063
	terheqja e mbeturinave	10.801	7.309	.133	1.478	.142
	distanca shtepi-pusne meter	.164	.061	.572	2.662	.009

a. Dependent Variable: perime

Faktoret abiotik sistemi i ujrave të zeza dhe, ne menyre jo shumë te sigurt dhe ai i mbetjeve urbane, kanë ndikim ne produktivitetin e bimeve. Sa më shumë kujdes të tregohet në to nga ana sanitare aq më pozitivisht ato ndikojne tek bimët. Informacioni i marrë nga variablat e shtuar, redukton parashikimin në ndryshueshmërinë e prodhueshmerisë se bimëve, vetëm si rezultat i zonës, me 20% per perimet, 4.3% for për jonxhën.

2.9.2 Zbatim 2:(vazhdim)

Në punimin me teme: **Cfarë ndikon në memorje**, pas analizës deskriptive të të dhënave është vënë re ndikim në variablin '*nrwords*,' dhe në përmirësimin e tij për sigurimin e kushtit të normalitetit '*sqrtnrwords*', ndikim i variablit '*gjinia*', ndikim i cili për ne nuk paraqet interes. Ndaj per të percaktuar efektin apo jo ne variablin në shqyrtim të kater variablave të tjerë te pavarur, që marrin pjesë në sdtudimin në fjale, efektin e duhanit alkolit te të ngrënit shëndetshëm dhe të sportit është kryer regresi linear hirarkik. Kjo metodë synon, pas heqjes nga *sqrword* e efektit të gjinise, të kontrollojë për efekt të alkolit, duhanit të të ngrënit shëndetshëm dhe të sportit.

Tabela 2–7 Tabela e regresit hirarkik

Model Summary

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Change Statistics				
					R Square Change	F Change	df1	df2	Sig. F Change
1	.186 ^a	.035	.031	.67475	.035	10.583	1	296	.001
2	.647 ^b	.418	.408	.52735	.384	48.150	4	292	.000

a. Predictors: (Constant), gjinia f/m

b. Predictors: (Constant), gjinia f/m, hani shendetshem po/jo, alkool po/jo, sa here bejne sport ne jave, duhan po/jo

Tek tabela 2-7 vlerësimi i R^2 -change tregon se, ky model justifikon 39% ndryshimet ne rrënjën katrore të numrit të fjalëve të memorizuara. Kontolli i hipotezës $R=0$, pra nuk ka lidhje lineare mes modelit të propozuar dhe faktoreve të tjere të testuar bie poshte. Kjo, sepse vlera e testit të kontrollit të hipotezës në fjalë është signifikativ, vlere e testit F jep $p=0.000002$. Madje $R=0.612$ tregon se lidhja ne fjale, e shprehur nga modeli i regresit, është e fortë.

Tabela 2–8 Koefficientët e korelacionit të regresit hirarkik

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	3.335	.118		28.184	.000
	gjinia f/m	.260	.080	.186	3.253	.001
2	(Constant)	3.267	.108		30.143	.000
	gjinia f/m	.068	.070	.049	.972	.332
	sa here bejne sport ne jave	.223	.018	.597	12.470	.000
	duhan po/jo	-.148	.075	-.104	-1.972	.050
	alkool po/jo	-.146	.072	-.104	-2.035	.043
	hani shendetshem po/jo	.042	.070	.028	.607	.544

a. Dependent Variable: sqrtword

Kryerja e metodes se regresit linear hirarkik konkludon për efekt të sportit, duhan dhe alkoolit (lexo p- vlerat në tabelen 2.8). Ndërsa variabli i mënyrës së të ngrënit, nuk sjell ndonjë ndryshim nëse nuk merret parasysh në modelin e regrest linear.

Koefficientet negativ pranë variablit duhan e alkool tregojnë efektin negativ te dy këtyre faktorëve, dhe efekt pozitiv të sportit në memorje.

Për të përcaktuar se cilët nga variablat e pavarur në fjalë luajnë rol me të madh në *sqrtnrwords* kontrollohen vlerat e koeficientëve të korelacionit të pjesshëm, tabela 2.9 Vlerat tregojnë për ndikim të sportit në masen $0.6^2=0.36$, pra 36% e ndryshimeve në variablin e shqyrtuar shpjegohen nga variabli '*sa herë bëjnë sport në javë*'. Ndërsa ndryshimet në duhan dhe alkool shpjegojnë vetëm 4% të ndryshimeve të '*sqrtnrword*'

Tabela 2–9 Koeficientët e korelacionit të pjesshëm

Excluded Variables^a

Model	Beta	In	t	Sig.	Partial Correlation	Collinearity Statistics
						Tolerance
1	sa here bejne sport ne jave	.629 ^b	13.085	.000	.606	.897
	duhan po/jo	-.240 ^b	-4.141	.000	-.234	.921
	alkol po/jo	-.224 ^b	-3.906	.000	-.222	.950
	hani shendetshem po/jo	.072 ^b	1.251	.212	.073	.993

a. Dependent Variable: *sqrtnrword*

b. Predictors in the Model: (Constant), gjinia f/m

2.9.3 Zbatim 3 (vazhdim)

Në shembullin e mësipërm për të parë se cili nga 4 variablat e pavarura sporti, alkoli, te ushqyerit dhe duhani, ndikojnë më shumë në parashikimin e '*sqrtnrword*', është përdorur regresi linear stepwise.

Afishimet e mëposhtme tregojnë se cilët janë variablat kryesorë që ndikojnë në këtë model regresi.

Tabela 2–10 Regresi stepwise

Variables Entered/Removed^a

Model	Variables Entered	Variables Removed	Method
1	sa here bejne sport ne jave		Stepwise (Criteria: Probability-of-F-to-enter <= .050, Probability-of-F-to-remove >= .100).

2			Stepwise (Criteria: Probability-of- F-to-enter <= .050, Probability-of- F-to-remove >= .100).
	duhan po/jo	.	

a. Dependent Variable: sqrtword

Model Summary

Model	R	R Square	Adjusted Square	Std. Error of the Estimate
1	.624 ^a	.389	.387	.53684
2	.639 ^b	.409	.405	.52887

a. Predictors: (Constant), saherebejnepornejave

b. Predictors: (Constant), saherebejnepornejave, duhan po/jo

Metoda e regresit stepwise tregon se nga tere faktoret e konsideruar te pavarur ndikim tek *srtnword* ka *numri i oreve te fiskultures* dhe *konsumojne apo jo duhan*.

Faktori *kosumim alkoli, mënyra e të ushqyerit* janë eliminuar nga teresia e faktoreve te pavarur te futur ne shqyrtim. Modeli justifikon rreth 40% te ndryshimeve te variablit *sqrtword*. Madje vetëm variabli sport justifikon 38.96% të ndryshimeve në *sqrtwords*.

Nese shikohet koeficienti para tij ne modelin e regresit te shumefishte linear stepwise, eshte 0.23 modeli i ndertuar vetem ne baze te faktorit sa here bejne sport ne jave kalimi nga dy here fiskulture ne jave ne tri here fiskulture do te ne jave do te sjell mesatarisht 13.69 nga 11.7 fjale të memorizuara.

3 Analiza e Variancës Dypërmasore (Metoda ANOVA Dypërmasore)

3.1 Modeli me nenpopullime te fiksuara (Modeli I)

Me këtë metode, vëzhgohet ndikimi i dy faktorëve të ndryshëm A e B , mbi një variabël, Y . Kjo metodë i shikon nenpopullimet të përcaktuara nga nivelet e secilit prej dy faktorëve, si dhe nga nderthurja e tyre. Faktorët A e B , në fakt janë variabla të tipit kategorik. Vlerat e këtyre dy variablave të pavarur përcaktojnë nivelet e faktoreve në shqyrtim dhe, kombinimet e niveleve të ndryshme të tyre, përcaktojnë modelin që po shqyrtojmë. Kështu, nëse A ka I nivele (variabli A ka I vlera) dhe B ka J nivele, (variabli B ka J vlera), atëherë, IJ janë kombinimet e niveleve të dy faktorëve. Kombinimi i nivelit të i –të të faktorit A , me atë të j -të të faktorit B përcaktojnë atë që mund të quhet *qelizë ij*.

Në Analizën Dypërmasore dallohen rastet kur numri i vëzhgimeve të variablit të rastit Y , për çdo qelizë të jetë i njëjtë dhe kur ai s'është i tillë. Këto vëzhgime përcaktojnë zgjedhjen e nenpopullimit të përcaktuar nga nderthurja e nivelit të i -të të A -së, me nivelin e j -të të B -së.

Ndaj ndërtohet modeli dy përmasor, ku çdo vlerë e n.r. në shqyrtim paraqitet si kombinim i mesatares, i nivelit të ndikimit të faktorëve në shqyrtim, të nderthurjes së tyre dhe të një termi gabimi.

Shënojmë y_{ijk} vlerën e k -të të Y në qelizën e ij -të, dhe ndërtohet modeli për shqyrtimin e efekteve të dy faktorëve të fiksuara

$$Y_{ijk} = \mu + \alpha_i + \beta_j + (\alpha\beta)_{ij} + e_{ijk}$$

ku $i=1..I$; $j=1..J$; $k=1..K$

μ është mesatarja, α_i është efekti i faktorit A , dhe β_j efekti i faktorit B , $(\alpha\beta)_{ij}$ paraqet nderthurjen e dy faktorëve ij -të. Gabimet e rastit e_{ijk} janë variabla të rastit të pavarur, dhe kanë shpërndarje $N(0, \sigma_e^2)$.

Në rastin kur $(\alpha\beta)_{ij}=0$ për çdo i dhe j , modelet A dhe B quhen aditive.

Le të kalojmë tani në vlerësimin e parametrave që përcaktojnë efektin e faktorëve në shqyrtim.

Vlerësimet e këtyre parametrave me Metodën e Katrorëve më të Vegjël nuk janë të vetëm, ndaj vihen edhe disa kushte të tjera. Supozohet se :

$$\sum_{i=1}^I \alpha_i = 0$$
$$\sum_{j=1}^J \beta_j = 0$$

$$\sum_{i=1}^I (\alpha\beta)_{ij} = 0 \text{ për } j=1..J$$

$$\sum_{j=1}^J (\alpha\beta)_{ij} = 0 \text{ për } i=1..I$$

Si zakonisht, duke përdorur Metodën e Katrorëve më të Vegjël, përcaktohen vlerësime të pazhvendosura të parametrave të këtij modeli, μ , β_j , α_i . Shuma e marrë është kjo:

$$S = \sum_i \sum_j (y_{ijk} - \mu - \alpha_i - \beta_j - (\alpha\beta)_{ij})^2$$

dhe, gjejmë vlerësuesit e parametrave:

$$\begin{aligned}\hat{\mu} &= \bar{y} \\ \hat{\alpha}_i &= \bar{y}_{i..} - \bar{y} \\ \hat{\beta}_j &= \bar{y}_{.j.} - \bar{y} \\ (\alpha\beta)_{ij} &= \bar{y}_{ij.} - \bar{y}_{i..} - \bar{y}_{.j.} + \bar{y} \\ &\text{për } i=1..I \text{ dhe } j=1..J\end{aligned}$$

duke vënë në dukje se $\bar{y}$ është mesatarja e vlerave të y_{ijk} ,

$\bar{y}_{.j.}$ është mesatarja e kollonës së j -të

$\bar{y}_{i..}$ është mesatarja e rreshtit

Nga teorema e Gaus Markovit këto vlerësime të pavarura të parametrave sigurojnë një minimum të variancës nga të gjithë vlerësimet që fitohen si kombinim linear i vlerave të vrojtuar. Kështu shuma S është bazë për vlerësimin e pazhvendosur të σ^2_e . Kjo shumë, mund të ndahet në katër shuma, sipas burimeve të variacionit, A,B,AB dhe Mb(Regresit).

$$\begin{aligned}SKT &= \sum_i \sum_j \sum_k (y_{ijk} - \bar{y})^2 = JK \sum_i (\bar{y}_{i..} - \bar{y})^2 + IK \sum_j (\bar{y}_{.j.} - \bar{y})^2 + K \sum_j \sum_i (\bar{y}_{ij.} - \bar{y}_{i..} - \bar{y}_{.j.} + \bar{y})^2 + \\ &\sum_i \sum_j \sum_k (y_{ijk} - \bar{y}_{ij.})^2 = SK_A + SK_B + SK_{AB} + SK_R\end{aligned}$$

Anova Dy Përmasore shqyrton katër komponentët e variancës^[9], duke ndertuar statistika te pazhvendosura per çdo burim te variacionit, si tabela 3.1.

Për të përcaktuar ndryshimin midis grupeve sipas një faktori vlerësojmë shumën e katrorëve të gabimeve sipas secilit faktor S_{KA} dhe S_{KB} dhe kështu mesatarja e

$$\text{gabimit për çdo faktor do të jetë } M_A = \frac{S_{KA}}{I-1} \text{ dhe } M_B = \frac{S_{KB}}{J-1}$$

Për të kuptuar nëse ka ndryshime midis grupeve të ndryshme sipas secilit prej dy faktorëve, duhet që ato të krahasohen me variancën e gabimit σ^2

Tabela 3–1 Përbërësit e variacionit Modeli 1

Burimet e variacionit	Shuma e katrorëve të gabimeve	Mesatarja e katroreve
Faktori A	$S_{KA} = JK \sum_i (\bar{y}_{i..} - \bar{y})^2$	$M_{KA} = \frac{S_{KA}}{I - 1}$
Faktori B	$S_{KB} = IK \sum_j (\bar{y}_{.j.} - \bar{y})^2$	$M_{KB} = \frac{S_{KB}}{J - 1}$
Ndërthurja e dy faktorëve	$S_{KAB} = K \sum_j \sum_i (\bar{y}_{ij.} - \bar{y}_{i..} - \bar{y}_{.j.} + \bar{y})^2$	$M_{KAB} = \frac{S_{KAB}}{(I - 1)(J - 1)}$
Mbetjet	$S_{KR} = \sum_i \sum_j \sum_k (y_{ijk} - \bar{y}_{ij.})^2$	$M_{KR} = \frac{S_{KR}}{IJ(K - 1)}$
Totali	$S_{KT} = \sum_i \sum_j \sum_k (y_{ijk} - \bar{y})^2$	

Ajo supozohet e njëjtë për të gjithë grupet. Ndaj, për vlerësimin e saj përdoret statistika

$$\hat{\sigma}^2 = \frac{\sum_{i=1}^I \sum_{j=1}^J K_{ij} s^2_{ij.}}{IJ \sum (K_{ij} - 1)} = \frac{\sum \sum \sum (y_{ijk} - \bar{y}_{ij.})^2}{IJ \sum (K_{ij} - 1)} \equiv M_{KR} \text{ që na siguron një vlerësim të } \sigma^2$$

panjohur të σ^2
 Meqë y_{ijk} janë një zgjedhje nga një popullim i madh më shpërndarje $N(\mu, \sigma^2)$ atëherë
 $\frac{(IJK - 1)S_{KT}}{\sigma^2} \sim \chi^2(IJ(K - 1))$

Ndërtohen testet statistikore më shpërndarje të përcaktuar, Tabela 3.2, dhe në bazë të tyre kontrollon testet statistike të paraqitura në Tabelën 3.2.

Tabela 3–2 Testet sipas hipotezave

$H_0: (\alpha\beta)_{ij}=0$	$H_0: \alpha_i=0$	$H_0: \beta_j=0$
$F = \frac{M_{KAB}}{M_{KR}}$	$F = \frac{M_{KA}}{M_{KR}}$	$F = \frac{M_{KB}}{M_{KR}}$
$((I-1)(J-1); IJ(K-1))$	$((I-1); IJ(K-1))$	$((J-1); IJ(K-1))$

Në bazë të p-vlerave, pra të sipërfaqeve në të djathtë nën kurbat e këtyre shpërndarjeve, diskutohen :

Në fillim hipoteza $H_0: (\alpha\beta)_{ij}=0$,

- Nëse p-vlera është shumë, shumë e vogël, atëherë nuk kemi ndonjë ndërthurje midis dy faktorëve.

- *Vazhdohet*, duke shfrytëzuar p-vlerat e testeve perkatës, për t'i dhënë përgjigje pyetjeve nëse ka efekt secili prej dy faktoreve A dhe B. Modeli I i Anova-s i jep përgjigje pyetjes, a janë të njëjta mesataret e nenpopullimeve të caktuara sipas secilit faktor, apo jo.
- Nëse p-vlera është e madhe, atëherë shumës së katrorëve të gabimeve të mbetjeteve, duhet ti shtojmë dhe këtë shumë tashmë të konsiderueshme të ndërthurjes midis dy niveleve, duke marrë kështu një vlerësim tjetër të σ^2 : $M_{SP} = S_{KP} / \nu_P$ ku $S_{KP} = S_{KR} + S_{KAB}$ me një numër shkallësh lirie $\nu_P = \nu_R + \nu_{AB}$.
- Nderhyrja në sintaksën e programeve statistikore bën të mundur dhe kontrollin e hipotezës për ndërthurjen apo jo të dy faktorëve në studim.

3.1.1 Përfundime

Duhet vënë në dukje se nuk ka ndonjë vlerë kufi të përcaktuar të p-vlerës, për të thënë se kur duhet bërë dhe kur jo ky modifikim, çdo gjë varet nga personi që studion problemin.

Megjithatë ky modifikim keshillohet të bëhet në rastet kur p-vlera prane ndërthurjes së dy faktoreve, $P > 0.5$. Mosmarrja parasysh e këtij konkluzioni, mund të çojë në përfundime të gabuara në analizën e dy testeve të tjera.

Punohet më pas për shqyrtimin e ndikimit apo jo njëri nga njëri të faktorëve A dhe më pas B, duke shfrytëzuar përkatësisht $F = M_{SA} / M_{SP}$ për testimin e hipotezës $H_0: \alpha_i = 0$ dhe $F = M_{SB} / M_{SP}$ të hipotezës $H_0: \beta_j = 0$.

3.1.2 Zbatim 1

Në një studim, "**Cilët faktorë ndikojnë në memorje**", është kontrolluar me parë ndikimi i shumë faktorëve që mendohej të ndikonin në ndryshoren e rastit, që përcakton sasinë e fjalëve të memorizuara, (n.r. që për të plotësuar kushtin e normalitetit është modifikuar në *sqrtnrword*). Faktori *gjini* si dhe *numri i orëve që merren me sport në javë*, janë përcaktuar të dy si faktorë ndikues në *sqrtnrword*.

Për të vërtetuar si ndikojnë këto dy faktore, në përputhje dhe me problemin që trajton studimi, ndikimin tek memorja e adoleshentëve, është bërë krijimi i një variabli të ri kategorik, *nivsportlevel* (faktor me $J=3$ nivele). Me metodën e ANOVA-s Dypërmasore, modeli me nenpopullime të fiksuara, është kontrolluar ndikimi i faktorit *gjini* ($I=2$ nivele) dhe *sportlevel*. Tabela 3.3 paraqit të dhënat e kësaj analize.

Afishimet e para të tabelës 3.3 përmban përvec mesatare në çdo 6 kombinime të dy faktorëve të përcaktuar dhe vëllimin e zgjedhjes për çdo $2 \times 3 = 6$ qelizat e krijuara.

Ndërsa, në pjesën tjetër të tabelës 3.3 lexojmë përfundimet e kësaj metode. Vlera e statistikës $F=1.03$ dhe e p-vlerës korresponduese .902, tregojnë nuk mund të flasim për ndërthurje mes faktoreve^[10]. Madje dhe vlera e variancës që përcakton kjo ndërthurje është tej mase e vogël, siç lexohet në tabelë vetëm 0.01.

Ndërsa afishimet prane faktoreve tregojnë, për ndikim *sportlevel* ($F=42.256$ dhe e p-vlerës 7.62 E-17), dhe për mos ndikim të faktorit *gjini*, ($F=.333$ dhe p-vlera .565).

Tabela 3–3 Tabela e të dhënave deskriptive dhe analizës ANOVA dypërmasore

Descriptive Statistics

Dependent Variable: sqrtword

gjinia f/m sportlevel		Mean	Std. Deviation	N
femer	0-2 here sport	3.4465	.65658	137
	3- 4 here sport	4.0243	.59823	35
	me shume se 4 here sport	4.3390	.60750	7
	Total	3.5944	.69653	179
mashkull	0-2 here sport	3.4594	.47177	59
	3- 4 here sport	4.1089	.52193	34
	me shume se 4 here sport	4.4163	.52522	26
	Total	3.8540	.64051	119
Total	0-2 here sport	3.4504	.60572	196
	3- 4 here sport	4.0660	.55943	69
	me shume se 4 here sport	4.3999	.53454	33
	Total	3.6981	.68555	298

Dependent Variable: sqrtword

Source	Type III Sum of Squares	df	Mean Square	F	Sig.	Partial Eta Squared
Corrected Model	37.779 ^a	5	7.556	21.672	.000	.271
Intercept	2148.261	1	2148.261	6161.623	.000	.955
gjinia	.116	1	.116	.333	.565	.001
sportlevel	29.465	2	14.733	42.256	.0000000	.224
gjinia * sportlevel	.072	2	.036	.103	.902	.001
Error	101.806	292	.349			
Total	4215.000	298				
Corrected Total	139.585	297				

a. R Squared = .271 (Adjusted R Squared = .258)

3.2 Modeli me nënpopullime të rastesishme (Modeli II)

Tabela Anova-s mund të shfrytëzohet edhe në një kendveshtrim tjetër, sipas metodës me nenpopullime të rastësishme.

Keshtu, popullimi $N(\mu, \sigma^2)$ ndahet në nenpopullime të përcaktuara sipas dy faktoreve në shqyrtim, si rezultat i kombinimit të tyre.

Zgjedhja e marre siguron, zgjedhje për çdo nenpopullim me shpërndarje $N(m_{ij}, \sigma_{ij}^2)$ ku $i=1..I; j=1..J$.

Nga zgjedhjet e këtyre grupeve, llogariten me parë mesataret m_{ij} për çdo grup, të caktuar nga kombinimi i dy faktoreve të shqyrtuar A dhe B.

Vlerat e zgjedhjes fillestare (si zgjedhje e nenpopullimit të grupit ij) mund të shkruhen në formën:

$y_{ijk} = m_{ij} + e_{ijk}$ ku e_{ijk} të pavarura me shpërndarje $N(0, \sigma^2)$

(Në fakt vëllimi i zgjedhjes për çdo grup është i ndryshëm, pra k ndryshojnë, kështu simbolika me e përshtatshme për k do të ishte k_{ij} . Por, nëse paraqes rastin kur vëllimi në çdo grup të përcaktuar nga dy faktorët është i njëjti, nuk kam ndonjë ndryshim, vetëm na thjeshton simboliken)

Nenpopullimet me shpërndarje $N(m_{ij}, \sigma_{ij}^2)$, shërbejnë tani si zgjedhje për popullimin e madh, të parë si popullim i nenpopullimeve. Ndaj elementet e kësaj zgjedhje shkruhen në formën:

$m_{ij} = \mu + a_i + b_j + (ab)_{ij}$,

ku $a_i \sim N(0, \sigma_{i.}^2)$ të pavarura, $b_j \sim N(0, \sigma_{.j}^2)$ të pavarura dhe të pavarura edhe $(ab)_{ij} \sim N(0, \sigma_{ij}^2)$

Perfundimisht, ndërtohet modeli $Y_{ijk} = \mu + a_i + b_j + (ab)_{ij} + e_{ijk}$ ku $i=1..I; j=1..J; k=1..K$, ku μ mesatarja e gjithë popullimit, a_i, b_j dhe $(ab)_{ij}$ janë të pavarura dhe kanë shpërndarje me mesatare 0 e variance perkatesisht. $\sigma_a, \sigma_b, \sigma_{ab}$ dhe $e_{ijk} \sim N(0, \sigma^2 + \sigma_{ij}^2)$.

Shtrohet problemi i vlerësimit të këtyre parametrave. Duke përdorur të njëjten metodë vlerësohet

$$\hat{\sigma}_a^2 = (M_{KA} - M_{KAB})/JK, \quad \hat{\sigma}_b^2 = (M_{KB} - M_{KAB})/IK \text{ dhe } \hat{\sigma}_{ab}^2 = (M_{KAB} - M_{KR})/K$$

Për të sqaruar ku ndryshojnë dy metodat në ndarjen në katër komponente të variancës, është paraqitur tabela 3.4.

Tabela 3–4 Përbërësit e variacionit Modeli II

Burimet e variacionit	Perberesit sipas modelit I	Perberesit sipas modelit II
Faktori A	$\sigma^2 + \frac{JK \sum \alpha_i^2}{I-1}$	$\sigma^2 + K\sigma_{ab}^2 + JK\sigma_a^2$
Faktori B	$\sigma^2 + \frac{IK \sum \beta_j^2}{J-1}$	$\sigma^2 + K\sigma_{ab}^2 + IK\sigma_b^2$
Ndërthurja e dy faktorëve	$\sigma^2 + \frac{JK \sum \sum (\alpha\beta)_{ij}^2}{(I-1)(J-1)}$	$\sigma^2 + K\sigma_{ab}^2$
Mbetjet	σ^2	σ^2

Ndaj testet statistikore sipas modelit II do te jene:

Ho: $(\sigma_{ab})^2=0$	Ho: $\sigma_a^2=0$	Ho: $\sigma_b=0$
$\frac{M_{KAB}}{M_{KR}}$	$\frac{M_{KA}}{M_{KAB}}$	$\frac{M_{KB}}{M_{KAB}}$
F=	F=	F=
$((I-1)(J-1); IJ(K-1))$	$((I-1); (I-1)(J-1))$	$((J-1); (I-1)(J-1))$

Në fillim kontrollohet hipoteza Ho: $(\alpha\beta)_{ij}=0$,

- Nëse p-vlera është shumë, shumë e vogël, atehere nuk kemi ndonjë ndërthurje midis dy faktorëve.

Modeli II i Anova-s *vazhdon*, duke shfrytëzuar p-vlerat e testeve perkatës, dhe i jep përgjigje pyetjes, a ka apo jo variacion midis nenpopullimeve te caktuara sipas secilit faktor.

- Nëse p-vlera është e madhe, atëhere shumës së katrorëve të gabimeve të mbetjeve, duhet ti shtojmë dhe këtë shumë tashmë të konsiderueshme të ndërthurjes midis dy niveleve, duke marrë kështu një vlerësim tjetër të σ^2 :

$$M_{SP}=S_{KP}/v_P \text{ ku } S_{KP}=S_{KR}+S_{KAB} \text{ me një numër shkallësh lirie } v_P=v_R+v_{AB}.$$

Vazhdohet me shqyrtimin e ndikimit apo jo të faktorëve A dhe B, duke shfrytëzuar përkatësisht $F=M_{SA}/M_{SP}$ për testimin e Ho: $\alpha_i=0$ dhe $F=M_{SB}/M_{SP}$ të Ho: $\beta_j=0$.

Shihet se statistikatat F ndryshojne nga ato te modelit te pare, e kjo ne rastet e nje nderthurje te dy faktoreve me njeri tjetrin, na ben qe vleresimi i marre me modelin II te ANOVA-s te jete ndoshta ndryshe nga ai sipas modelit I.

3.2.1 Zbatim 2

Diskutimi i problemit të ndikimit të ndotjes në pneumologji me Anoven Dy Përmasore

Për problemin pneumologjik kemi zgjedhur si tregues përcaktues ndryshorjen e rastit nivkoll

Në studimin e problemit të pneumologjisë, është vënë re se faktori zonë ka ndikuar në nivelin e kollës dhe, është shtruar problemi i ndikimit apo jo mbi këtë ndryshore rasti të kohëqëndrimit në zonën industriale (kohëqëndrim, faktori A me 4 nivele)[8] dhe të shkallës së duhanpirjes (nivduhan, faktori B me tri nivele).

Ky problem u trajtua me metodën e Anova-s dy përmasore.

Tabela e Anovës dy permasore shfaq $F= 3.54$ dhe $P= 0.039$ pranë variablilit *kohëqëndrim*, që dëshmon se roli i kohëqëndrimit në këto zona, per ndryshoren nivkoll është domethënës.

Ndersa prane variablilit shkalla e duhan pirjes p-vlera =0.056, qe dëshmon per mosndikim te duhanit ne ndryshoren e rastit ne shqyrtim.

Sipas *Modelit II*, modeli i ANOVA-s Dypërmasore të rastwsishme, marrim këto vlerësime për përbërësit e variancës, Tabela 3-5:

Tabela 3-5 Afishimet e përbërësve të variacionit

Përbërësi i variacionit	Vlera
niveli kohëqëndrimit	0.10873
niveli i duhan pirjes	0.08524
niveli kohëqëndrimit* niveli i duhan pirjes	-0.06260
mbetjet(gabimet)	0.50283

Meqë varianca e një komponenti nuk mund të jetë më e vogël se 0, atëherë vleresimi për variancën nën efektin e ndërthurjes së dy faktorëve merret 0.

Pra, jemi në rastin kur dy faktorët konsiderohen aditive, dhe kjo metode për shqyrtimin e variancës së dy faktorëve afishon pranë faktorit nivelkohëqendrim $F=8.93$ dhe $P=0.043$ dhe pranë faktorit nivel i duhanpirjes $F=7.44$ dhe $P=0.026$, dhe tregojnë tani **ndikimin e dy faktorëve në shqyrtim në ndryshoren nivkollë**. Kështu meqë faktorët nuk ndërthuren, në zgjidhjen e këtij problemi mund të jipet edhe përgjigja se të dy faktorët ndikojnë në n.r në shqyrtim, pra në këtë tregues të problemeve pneumologjike.

Pra, mund të themi se fakti që modelin e ANOVA-s e shikojmë sipas **modelit I apo II, na con në rezultate tërësisht të ndryshme**.

Por, kontrolli i hipotezës së ndërthurjes apo jo të dy faktorëve në shqyrtim shfaq, pranë vlerës së statistikës $F=2.3367$, p vlerën 0.077, tregon se ndonëse nuk ka një ndërthurje të qarte mes dy faktoreve, keto faktore *nuk mund të konsiderohen si aditiv*. Kjo na bën të mos ndikohemi shumë nga vlerat e afishuara, por të nisemi nga vlerësimi tjetër. Në bazë të sqarimeve të bera, modifikohen vlerat e emeruesve të statistikave.

Fillimisht modifikohet shuma $S_{KP}=S_{KR}+S_{KAB}=59.83$ dhe me pas statistika në emeruesin e dy statistikave vleresuese: $M_{SP}=S_{KP}/v_P=0.4913$ që shfaq $F=2.3367$ dhe $P=0.089$.

Përfundimisht, një vëzhgim më i kujdesshëm i Anovas Dypërmasore, tregon se në bazë të kësaj metode nuk mund të themi se ndikimi i kohëqëndrimit në këtë zonë industriale është domethënës mbi këtë tregues të problemeve pneumologjike.

Po kështu modifikohet dhe vlerësimi për ndikimin apo jo të faktorit tjetër, shkallës së konsumimit të duhanit. Duke u nisur nga vlerësimi i mësipërm i M_{KP} , merret $F=1.729$ dhe $P=0.18$, që do të thote se nuk rezulton ndikim vetëm i këtij faktori në ndryshoren e rastit që shqyrtohet.

Nëse do të rishikonim rrezultatet e ti grumbullonim ato, duke përfillur ndërthurjen mes dy grupeve atëherë, do të fitonim të njëjtat F-vlera si në rastin e parë, pra të njëjtat përgjigje për ndikimin e faktorëve në shqyrtim.

Pra, shtojmë se sa të ndryshme mund të jenë përfundimet varet, dhe nga qëndrimi që mbahet ndaj ndërthurjes së dy faktorëve.

3.2.2 Vërejtje

SPSS shoqërohet me tri editor, editori i të dhënave, i output-ve dhe editori i sintaksës.

Editori i të dhënave përmban variablat, llojin e tyre e bën të mundur kryerjen e analizës së ndryshme statistikore.

Editori vier i outputeve, krijon lehtësi për afishimin e tabelave nga komanda të ndryshme të ekzekutuara nga editor i të dhënave dhe ai i sintaksës.

Editori i sintaksës hapet File/Data/Editor/New/ Syntax, dhe bën të mundur ruajtjen e sintaksës së gatshme për bashkësi të dhënash të ndryshme. Kështu editor i vecante mund të ekzekutohet në bazën e të dhënave të reja, duke ndryshuar vetëm emertimet e variablave.

Afishimet për përcaktimin e ndërthurjes mes dy faktorëve arrihen pasi në këtë editor të fundit shkruhet kodi i posaçëm.

3.3 Shtrimi matematik i Analizës së Kovariancës (ANCOVA)

Bazë e analizës së kovariancës është gërshetimi i metodës ANOVA me një analizë të thjeshtë të regresit linear.

Kuptohet, se ajo përdoret në rastet kur ndryshorja e rastit. nuk varet vetëm nga nivelet e një faktori, por dhe nga një ndryshore rasti tjetër, e vazhdueshme.

Pra, supozohet se kemi p.sh. I nivele të një popullimi sipas një faktori. Shënohen y_{ij} vëzhgimet e marra në eksperimentin e j-të

të nivelit të i-të. Supozohet gjithashtu se, y_{ij} kanë shpërndarje $N(\mu_i, \sigma^2)$.

Në të tilla raste modeli i ndërtuar matematik gjatë analizës ANOVA është $y_{ij} = \mu + \alpha_i + e_{ij}$

ku e_{ij} kanë shpërndarje $N(0, \sigma^2)$

μ pritja matematike e popullimit,

α_i ndikimi i nivelit të i-të dhe, siç e dihet, për të siguruar një vlerësim të këtyre parametrave me metodën e katrorëve më të vegjël, vihet kushti $\sum_i J_i \alpha_i = 0$

Por, ndërkohë krahas vrojtimeve të vlerave y_{ij} në vëzhgim e i-të të nivelit të j-të, bëhet një tjetër vëzhgim i variablit x, që supozohet të jetë linearisht i varur me y, vëzhgime që shënohen me x_{ij} .

Vlerësimi sipas ANOVA-s do të bëhej më i saktë nëse y_{ij} do të zbritej ajo pjesë që përcaktohet nga regresi linear me x_{ij} .

Kështu nga vlerat e marra y_{ij} , fitohen vlera të reja $y_{ij}^* = y_{ij} - \beta(x_{ij} - \bar{x})$

Këto, quhen vlera të reduktuar, sepse i është hequr efekti i ndikimit tek to i x_{ij} , duke u supozuar ky një ndikim linear.

Modeli matematik ANOVA, tani ndërtohet për këto vlera të reduktuara y_{ij}^* .

$$y_{ij}^* = \mu + \alpha_i + e_{ij} \text{ ku :}$$

e_{ij} kanë shpërndarje $N(0, \sigma^2)$

μ pritja matematike e popullimit të madh,

α_i ndikimi i nivelit të i-të.

Ky model matematik i vlerave të reduktuar nuk është gjë tjetër vetëm se modeli matematik :

$$y_{ij} = \mu + \beta(x_{ij} - \bar{x}) + \alpha_i + e_{ij}$$

i vlerave të y_{ij} , që mund ta quajmë dhe model të analizës së kovariancës.

Përfundimisht, në këtë model, $\beta(x_{ij} - \bar{x})$ është ndikimi në y_{ij} i variablit x , i supozuar si një ndikim linear dhe α_i , ndikimi i nivelit të i -të në sajë të faktorit në shqyrtim.

Kështu, meqë modeli i analizës së kovariancës kthehet në një model të përgjithshëm linear, duke përdorur metodën e katrorëve më të vegjël, përcaktohen vlerësime të pazhvendosura të parametrave të këtij modeli, μ , β , α_i .

Pas ndertimit të shumës $S = \sum_i \sum_j (y_{ij} - \beta(x_{ij} - \bar{x}) - \mu - \alpha_i)^2$

kërkohet minimumi i saj, duke përcaktuar kështu këto vlerësime për parametrat:

$$\begin{aligned} \frac{\partial S}{\partial \mu} &= - \sum_i \sum_j (y_{ij} - \beta(x_{ij} - \bar{x}) - \mu - \alpha_i) = 0 \\ &\quad - n\bar{y} - n\mu = 0 \\ \hat{\mu} &= \bar{y} \end{aligned}$$

$$\begin{aligned} \frac{\partial S}{\partial \beta} &= - \sum_i \sum_j (x_{ij} - \bar{x})(y_{ij} - \beta(x_{ij} - \bar{x}) - \mu - \alpha_i) = 0 \\ \sum_i \sum_j (x_{ij} - \bar{x})y_{ij} - \beta \sum_i \sum_j (x_{ij} - \bar{x})^2 - (\mu + \alpha_i) \sum_i \sum_j (x_{ij} - \bar{x}) &= 0 \\ \text{nga} \dots \text{ku} \end{aligned}$$

$$\hat{\beta} = \frac{\sum_i \sum_j (x_{ij} - \bar{x})(y_{ij} - \bar{y}_i)}{\sum_{i,j} (x_{ij} - \bar{x})^2} = \frac{S_{xy}}{S_{xx}}$$

$$\frac{\partial S}{\partial \alpha_i} = - \sum_j (y_{ij} - \beta(x_{ij} - \bar{x}) - \mu - \alpha_i) = 0$$

$$\hat{\alpha}_i = \frac{\sum_j y_{ij}}{n} - \hat{\beta} \sum_j (x_{ij} - \bar{x})$$

$$\hat{\alpha}_i + \hat{\mu} = \bar{y}_i - \hat{\beta}(\bar{x}_i - \bar{x})$$

Dihet që këto vlerësime të parametrave të modelit linear të ndërtuar, sigurojnë dhe një minimum të variancës nga të gjithë vlerësuesit linear të pazhvendosur të këtyre parametrave.

Ndaj,

$$\begin{aligned} S &= \sum_i \sum_j (y_{ij} - \hat{\beta}(x_{ij} - \bar{x}) - \hat{\mu} - \hat{\alpha}_i)^2 = \sum_i \sum_j (y_{ij} - \frac{S_{xy}}{S_{xx}}(x_{ij} - \bar{x}) - \bar{y}_i - \frac{S_{xy}}{S_{xx}}(\bar{x}_i - \bar{x}))^2 = \\ &= \sum_i \sum_j (y_{ij} - \bar{y}_i - \frac{S_{xy}}{S_{xx}}(x_{ij} - \bar{x} - \bar{x}_i + \bar{x}))^2 = \sum_i \sum_j (y_{ij} - \bar{y}_i - \frac{S_{xy}}{S_{xx}}(x_{ij} - \bar{x}_i))^2 = S_{yy} - \frac{S_{xy}^2}{S_{xx}} \end{aligned}$$

është një vlerësim i zhvendosur i variancës së gabimit.

Ndërsa, $MS = \frac{1}{n - I - 1} (S_{yy} - \frac{S_{xy}^2}{S_{xx}})$ është vlerësim i pazhvendosur për σ^2

Në këtë mënyrë është ndërtuar një vijë e regresionit linear

$$\hat{y} = \bar{y}_i + \frac{S_{xy}}{S_{xx}}(x - \bar{x}_i) \quad \text{për } i=1..I.$$

Për këtë arsye, këto quhen vija të regresionit linear brenda grupit.

Pas dhënies së komandës së vlerësimit me ANOVA të modelit të modifikuar do të afishohen këto vlera:

Tabela 3–6 ANOVA e modelit të qendëruar

Midis niveleve	$SS_M = \sum_i J(\bar{y}_i - \bar{y}^*)^2$	$SS_M/(I-1)$
Brenda niveleve	$SS_B = \sum_i \sum_j (y_{ij}^* - \bar{y}_i^*)^2$	$SS_B/(n-I-1)$
Totali	$SS_T = \sum_i \sum_j (y_{ij}^* - \bar{y}^*)^2$	$SS_T/(n-1)$

Kështu në rastin tonë, tabela 3-6 paraqet vlerat e analizës së variancës për vlerat e modifikuara. Ndërsa, për vlerat reale të vrojtuar, nuk është gjë tjetër vecse një analizë kovariance.

Tabela 3–7 Përbërësit e analizës së kovariancës mes grupeve

Midis niveleve	$SS_M = \sum_i J(\bar{x}_i - \bar{x})(\bar{y}_i - \bar{y})$	$SS_M/(I-1)$
Brenda niveleve	$SS_B = \sum_i \sum_j (x_{ij} - \bar{x})(y_{ij} - \bar{y}_i)$	$SS_B/(n-I-1)$
Totali	$SS_T = \sum_i \sum_j (x_{ij} - \bar{x})(y_{ij} - \bar{y})$	$SS_T/(n-1)$

4 Analiza e Variancës Shumëpërmase MANOVA

Kjo metodë është si të thuash një metodë ANOVA, por në rastin kur kemi shumë variabla të varur. Nga pikpamja e përgjithshme ANOVA teston për diferencën midis mesatareve të dy ose më shumë grupeve, në këtë këndvështrim metoda MANOVA teston për diferencën në dy ose me shume mesatare vektorësh. Duke qenë një problem i komplikuar, problemi në fjalë, për zgjidhjen e tij janë dhënë mënyra të ndryshme. Në këtë kapitull do të paraqesim disa variante për zgjidhjen e tij, dhe do të diskutohet për prioritetet e secilës metodë, duke konkluduar kështu dhe për rastet kur secila nga metodat e trajtuara do të ishte mirë të përdorej.

4.1 Analiza e Variances Shumëpërmase (MANOVA)

Shtrimi teorik i metodës MANOVA përcakton te veçantën e kësaj metode. Kjo metodë bazohet në një model regresi linear ndaj shpesh njihet me emrin MANOVA e variantit GLM.

Nga n individë, pra nga një zgjedhje me vëllim n , janë vëzhguar vlerat e p variablave të rastit, dhe vëzhgimet janë shënuar si më poshtë:

Variablat	Individët					
	1	2	.	.	.	n
1	Y_{11}	Y_{12}				Y_{1n}
.	.	.	.	.	.	.
.	.	.	.	.	.	.
.	.	.	.	.	.	.
p	Y_{p1}	Y_{p2}	.	.	.	Y_{pn}

Pra, vektori $Y_i = (y_{i1}, y_{i2}, \dots, y_{in})$ paraqet vlerat koresponduese të n vëzhgimeve në tërë individet të variablilit të i -të. Modeli linear për çdo komponente Y_i , $i=1, 2, \dots, p$, që mund të shihet si një model me një variabël, përcaktohet nga barazimi vektorial:

$$E(Y_i) = X' \beta_i, \quad \text{cov}(Y_i) = \sigma_{ij} I$$

ku $(X')^{n \times m}$ është matrica e desinjuar e rankut $r \leq m < n$, σ_{ii} është dispersioni I variablilit të i -të dhe

$\beta_i = (\beta_{i1}, \dots, \beta_{im})'$ është vektori me m komponentë, komponentët e të cilit janë parametra të panjohur të variablilit të i -të.

Varesia e variablave shprehet nga

$$\text{cov}(Y_i, Y_j) = \sigma_{ij} I \quad i, j = 1, 2, \dots, p$$

ku σ_{ij} është kovarianca mes variablilit të i -të dhe të j -të. Është supozuar se $p \leq n-r$ dhe $m < n$

Modeli mund të shkruhet në formë matricore

$$Y = X'\beta + e$$

ku vlerat e vrojtura përcaktojnë matricën

$$\begin{bmatrix} y_{11} & y_{21} & \cdot & y_{p1} \\ y_{12} & y_{22} & \cdot & y_{p2} \\ \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot \\ y_{1n} & y_{2n} & \cdot & y_{pn} \end{bmatrix}$$

(X') $n \times m$ është matrica e desinjuar e rankut r , dhe

$$\begin{bmatrix} \beta_{11} & \beta_{21} & \cdot & \beta_{p1} \\ \beta_{12} & \beta_{22} & \cdot & \beta_{p2} \\ \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot \\ \beta_{1n} & \beta_{2n} & \cdot & \beta_{pn} \end{bmatrix}$$

është matrica e parametrave të panjohur.

Së fundmi, e $n \times p$ është matricë, rreshtat e së cilës janë një zgjedhje e rastit n përmase që formohet nga shpërndarja p - përmase $N(0, \Sigma)$ ku $\Sigma^{p \times p}$ është matrica e kovariancës dhe $0^{p \times 1}$ është vektori 0.

Ekuacioni I marë nuk është gjë tjetër veç se mënyra formale e shprehjes së modelit të regresit të shumëfishtë linear.

4.1.1 Zbatimi 1

Është përdorur metoda MANOVA, Modeli i Regresit Linear të saj (GLM), në trajtimin e **Ndikimit të nivelit të ndotjes në problemet shëndetësore të përgjithshme, në problemet dermatologjike dhe ato pneumatike**. Kështu, në problemin në fjalë janë marë të dhëna nga 171 individë, banues të tri zonave, dy të konsideruara industriale, Marinëz dhe Sheqishtë-Belinë, dhe, një e konsideruar e paster Kurjan. Të intervistuarve i janë bërë dhe 13 pyetje të ndryshme, për gjendjen shëndetësore të të intervistuarve.

Në këtë situatë $n=171$, $p=13$ dhe Y_i është vektori i 171 vëzhgimeve, përgjigjeve të pyetjes së i -të, për $i=1,2,\dots,13$. Për këtë (pyetje), variabël të varur të i -të, modeli do të shkruhej në formën:

$$E(Y_i) = X_i'\beta_i$$

pra $Y_{ij} = \mu_i + \alpha_{i1}X_{j1} + \alpha_{i2}X_{j2} + e_{ij}$, për $i=1,2,\dots,13$ dhe $j=1,2,\dots,171$.

Kështu $\beta_i = (\mu_i, \alpha_{i1}, \alpha_{i2})'$ është vektori me 3 parametra i faktorit të i -të. X_{j1} dhe X_{j2} janë indekset e variablit '*fshatrat*'

Ishte statisticieni David P. Nichols⁽¹²⁾, që më 1997, pasi bëri një analizë të dallimeve mes MANOVA-n dhe Metoden e Regresit Linear të Shumefishtë të një variabli kategorik, tregoi se Metoda e Regresit Linear të Shumefishtë nuk siguron një zgjidhje të vetme të këtij ekuacioni matricor. Një ndër zgjidhjet më të thjeshta mund të kërkohej duke u bazuar në një matricë njësi. Ndërsa në MANOVA, modeli e Regresit Linear të Shumefishtë kërkohej të jetë kanonik, ndaj modeli I matrices X merr formën si në tabelën 4-1. Kjo, sepse kollona e fundit duhet shprehur si kombinim linear i kollonave para saj, pra $A_3 = C - A_1 - A_2$

Tabela 4–1 Modeli i Matrices sipas GLM dhe MANOVA-s

Nivelet	Matrica GLM				Matrica e MANOVA-s		
	C	A1	A2	A3	C	A1	A2
1	1	1	0	0	1	1	0
2	1	0	1	0	1	0	1
3	1	0	0	1	1	-1	-1

Në rastin tonë $m=r=3$ modeli mund të shkruhet në këtë formë matricore:

$$\begin{bmatrix} y_{11} & y_{21} & \cdot & y_{13,1} \\ y_{12} & y_{22} & \cdot & y_{13,2} \\ \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot \\ y_{1,63} & y_{2,63} & \cdot & y_{13,63} \\ y_{1,64} & \cdot & \cdot & y_{13,64} \\ \cdot & \cdot & \cdot & \cdot \\ y_{1,116} & \cdot & \cdot & y_{13,116} \\ y_{1,117} & \cdot & \cdot & y_{13,117} \\ \cdot & \cdot & \cdot & \cdot \\ y_{1,1712} & \cdot & \cdot & y_{13,171} \end{bmatrix} = \begin{bmatrix} 1 & 1 & 0 \\ 1 & 1 & 0 \\ \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot \\ 1 & 1 & 0 \\ 1 & 0 & 1 \\ \cdot & \cdot & \cdot \\ 1 & 0 & 1 \\ 1 & -1 & -1 \\ \cdot & \cdot & \cdot \\ 1 & -1 & -1 \end{bmatrix} \begin{bmatrix} \mu_1 & \mu_2 & \cdot & \mu_{13} \\ \alpha_{1,1} & \alpha_{2,1} & \cdot & \alpha_{13,1} \\ \alpha_{1,2} & \alpha_{2,2} & \cdot & \alpha_{13,2} \end{bmatrix} + e$$

4.1.2 Vlerësimi i parametrave

Vlerësuesi për matricen e parametrave të panjohur përcaktohet me vlerësimet që merren me Metoden e Katroreve më të Vegjel, duke shfrytëzuar vlerat e vëzhguara Y_i pra të variablit të i -të.

Kështu zgjidhja e ekuacionit

$$(XX') \beta_i = XY_i$$

na siguron vlerësuesit që na duhen të $\beta_i = (\beta_{i1}, \dots, \beta_{ip})$, e që mund të shihen si një model linear një permasor për Y_i .

Vlerësuesi i pazhvendosur për σ_{ij} nga teoria e modelit me një variabël përfitohet si më poshtë:

$$\frac{R_0(i,j)}{n-r} = \frac{Y_i' Y_i - (Y_i' X') \hat{\beta}_i}{n-r}$$

ku $r = \text{rank } X'$. Vlerësuesi i pazhvendosur për σ_i është

$$\frac{R_0(i,i)}{n-r} = \frac{Y_i' Y_i - (Y_i' X') \hat{\beta}_i}{n-r}$$

për $i,j = 1,2,\dots,p$. $R_0(i,j)$ quhen shumatat e katroreve të mbetjeve dhe $R_0(i,j)$ është shuma e prodhimit të mbetjeve $i,j = 1,2,\dots,p$. Matrica

$$\begin{bmatrix} R_0(1,1) & R_0(1,2) & \dots & R_0(1,p) \\ R_0(2,1) & R_0(2,2) & \dots & R_0(2,p) \\ \vdots & \vdots & \ddots & \vdots \\ R_0(p,1) & R_0(p,2) & \dots & R_0(p,p) \end{bmatrix}$$

quhet matrica e shumës së katrore të mbetjeve.

4.1.3 Testi I Hipotezës së Linearitetit

Testi i hipotezës së Modelit të Regresit Linear është trajtuar në 1965 nga Rao. Ai, për secilin nga variablat e varur, kryente ANOVA-n një-përmase. Kështu ai analizonte hipotezën :

$$H_0: H' \beta_i = \xi_i \quad \text{për } i=1,2,\dots,p$$

Këto hipoteza në trajtë matricore mund të shkruhen në formën

$$H' \beta = [\xi_1 \ \xi_2 \ \dots \ \xi_p]$$

ku matrica H' $s^* \times m$, me rank $s \leq k$ dhe, vektori ξ_i $s^* \times 1$ janë të dhënë. Kështu kemi një model të përcaktuar, nga i cili mund të sigurojme një vlerësim për β_i , β_i^* , dhe matricë produkt dhe shumë katrorësh të mbetjeve R_1 . Matrica $R_1 - R_0$ quhet matrica produkt dhe shumë katrorësh që përcaktuan nga devijimin nga hipoteza. Analiza e variacionit, duke patur parasysh zberthimin e $R_1 = R_0 + R_1 - R_0$ kthehet, në problemin e krahasimit të R_0 me $R_1 - R_0$.

Më 1965 Rao tregoi se hipoteza e mësipërme e trajtës matricore merr formën e gjetjes së p rrënjëve $\lambda_1, \lambda_2, \dots, \lambda_p$ të ekuacionit

$$|R_0 - \lambda R_1| = 0$$

Testi I kontrollit të hipotezës zero kthehet kështu në teste të funksioneve të λ_i .

Një mënyrë është të gjejmë minimumin e rrënjëve, sepse kështu arrihet maksimumi i devijimeve nga hipoteza zero, dhe teston për të.

Ndërsa sipas kriterit Willks kriteri Λ përcaktohet nga

$$\Lambda = \lambda_1 \lambda_2 \dots \lambda_p = \frac{|R_0|}{|R_1|}$$

4.1.4 Zbatimi 1 (vazhdim)

Në rastin tone testohet përmes metodes ANOVA njëpërmase hipotezat H_0 marrin formën:

$$H_0: \alpha_{i1} = \alpha_{i2} = 0$$

për $i=1,2,\dots,13$.

4.1.5 Testi i differences mes vlerës së mesatareve midis disa popullimeve: MANOVA

Le të shënojmë me $n_1, n_2, \dots, n_k$ vëzhgimet e kryera për k popullime, dhe të shënojmë me $\bar{Y}_{ir}$ mesataret e zgjedhjeve të p variablave nga zgjedhja e r -të, për $r = 1, 2, \dots, k$.

Ndërsa $S_{ij}^{(r)}$ elementi I matricës së shumës së katrorëve me n_r-1 shkallë lirie që merret nga zgjedhja e r-të..

Dhe së fundmi, le të jenë $\bar{Y}_1, \bar{Y}_2, \dots, \bar{Y}_p$ vlerat mesatare dhe, (S_{ij}) matrica e shumës së katrorëve të fituar nga kombinimi I të gjithë zgjedhjeve në një të vetëm. Duke vepruar si tek ANOVA një përmase përcaktojmë:

$$B_{ij} = \sum_{r=1}^r n_r \bar{Y}_{ir} \bar{Y}_{jr} - \bar{Y}_i \bar{Y}_j$$

si shumë të produkteve midis popullimeve dhe

$$\ddot{S}_{ij} = \sum_{r=1}^k S_{ij}^{(r)}$$

shumën e produkteve brenda popullimit të i-të , për $i=1,2,\dots,p$. Këto madhësi afishohen në tabelën e MANOVA-s . Kriteri Λ , në të cilën kjo metodë bazohet për të vlerësuar barazimin e vlerave të mesatareve është:

$$\Lambda = \frac{|W|}{|W+B|}$$

ku $|W|$ dhe $|W+B|$ janë përcaktorët e matricave (W_{ij}) dhe $(B_{ij}) + (W_{ij})$. Statistika Λ ka shpërndarje të vështirësi në përcaktim, ndaj në praktikë përdoret përafrimi F :

$$F = \left(\frac{ms - 2\lambda}{2r} \right) \left(\frac{1 - \Lambda^{1/s}}{\Lambda^{1/s}} \right)$$

$$\text{me } m = n - 1 - \frac{1}{2}(p + k) \quad s = \sqrt{\frac{p^2(k-1)^2 - 4}{p^2 + (k-1)^2 - 5}}$$

Kështu për përcaktimin e përqindjeve përdoret shpërndarja $F(2r, (ms-2\lambda))$ shkallë lirie

Idea e përafrit të shpërndarjes së vërtetë me shpërndarjen e Fisherit është bërë nga Rao, dhe testi në fjalë njihet me emrin e tij. Por ky test përdoret atëherë kur variabli I pavarur që studjohet, plotëson kushtin e normalitetit dhe, natyrisht, atë të Homogjenitetit të Variancës , që kontrollohet me anë të testit Leven.

Ndërsa Bartlett e përafroi këtë shpërndarje me shpërndarjen X^2 .

Sipas Bartlett transformimi i mëposhtëm i statistikës Λ realizon pikerisht , dhenien e përgjigjes se testit, duke përdorur shpërndarjen X^2 . Transformimi që ai propozoi është:

$$-(n - 1 - \frac{1}{2}(p + k)) \ln \Lambda$$

dhe statistika e formuar ka shpërndarje X^2 me $p(k-1)$ shkallë lirie.

Idea e përafrit të shpërndarjes në shqyrtim me shpërndarjen X^2 me $p(k-1)$ shkallë lirie është bërë nga Bartlett që më 1947. Ky modifikim I statistikës përdoret veçanërisht nëse nuk jemi të treguar të kujdesshem për testin e normalizimit të variablit të pavarur Y.

4.1.5.1 Vërejtje

Programet kompjuterike të cilat kryejnë metodën MANOVA, të programuar e modeluar nga Philip Burns nga Universiteti I NorthWest-it, bëjnë Analizën e Kovariancës Shumëpërmase dhe Analizë të Variancës Njëpërmase për çdo

variabël, si dhe analizën e regresit. Për çdo efekt të testuar afishohet një faqe output-i. Pra afishohen tri output-e

- Testet e Variancës së Shumëfishtë. Këto teste përfshijnë kriterin llambda Λ Wilk-i që përafrohet duke shfrytëzuar përafrimin e Rao-s me shpërndarjen F , si dhe teste të tjerë të përmendura në paragrafin e mësipërm.
- Testet ANOVA një-përmasore për çdo variabël, të cilat bëjnë të mundur përcaktimin e ndikimit apo jo të variablit të pavarur në secilin prej variablave të varur në studim.
- t-testet të cilat kanë vlerë nëse shihet ndikim I variablit të pavarur tek variablat e varur, për të kuptuar ç'nivele të variablit të pavarur ndikojnë tek këto variabla.

4.1.5.2 Zbatimi i 4.1 (vazhdim)

Duke vërtetuar me kujdes tri output-et e MANOVA-s (shih problemin Ndikimi I nivelit të ndotjes në problemet shendetësore të përgjithshme, në problemet dermatologjike dhe ato pneumatike), arrihet në dënimin e disa konkluzioneve.

Nga Analiza e Variancës së Shumëfishtë, meqënëse jemi në rastin kur nuk i është kushtuar kujdes normalizimit të të dhënave, por ato janë konsideruar normale nisur nga vlerat e analizës deskriptive të variablave, lexohet vlera e statistikës Fisher nga afishimi I Pillai's Trace, dhe kjo është 4.654 dhe i korespondon p- vlera 1.2 E-12. Kjo tregon se ka efekt variabli I pavarur, në rastin tonë 'fshatrat', në vlerat e gjithë variablave të varura, kur ato të 13 shihen si një grup I vetëm. Por ky ndikim mund të jetë në një nëngrup të këtyre 13 variablave të varur.

Nga Testet ANOVA një-përmasore për çdo variabël, Tabela 4.1, vërehet ndikim I variablit 'fshatrat' në secilin prej 13 variablave të varur, të cilat mund të shikohen të përcaktuara në tabelën në fjalë.

Për të përcaktuar se midis cilave vlerave të variablit të pavarur 'fshatrat' vërehet ndryshim, janë kryer t-testet për secilën nga dy çiftet e zonave në shqyrtim, për çdo variabël të varur.

Tabela 4-2 ANOVA për 13 variablat e varur

Source	Dependent Variable	Type III Sum of Squares	df	Mean Square	F	Sig.	Partial Eta Squared
fshatrat	sy te perlotur	4.706	2	2.353	10.392	.000	.110
	kruarje hundesh	1.918	2	.959	5.610	.004	.063
	tharje gryke	10.481	2	5.240	27.313	.000	.245
	gelbaze	4.852	2	2.426	11.020	.000	.116
	kolle	2.674	2	1.337	5.935	.003	.066
	lunga ne lekure	.905	2	.452	4.755	.010	.054
	kruarje ne lekure	2.357	2	1.179	9.113	.000	.098
	shperthim puçrrash	1.404	2	.702	6.373	.002	.071
	apo flluskash						
	skuqje nxehtesi ne	14.098	2	7.049	44.434	.000	.346
	fytyre						

pezmatime te lekures	3.556	2	1.778	12.302	.000	.128
marramendje	6.952	2	3.476	16.507	.000	.164
nauze, te vjella	5.380	2	2.690	12.306	.000	.128
dhimbje koke te forta	13.601	2	6.800	45.622	.000	.352

. Duke u perqendruar në vlerat e t-testeve te afishuara mund të arrihet në perfundimin që

- o Ka ndryshim signifikativ midis fshatrave Kurjan dhe Sheqishtë- Beline dhe Kurjan e Marinez, për çdo pyetje te bëre. Vlera e t-testit rezulton pozitiv, ndaj themi se vlerat e këtyre variablave janë më të larta në Kurjan. Duke vërejtur mënyrën e kodimit të përgjigjeve në fjalë, konkludohet se ka ndikim negativ të zonës industriale në tërë parametrat e përcaktuara si përcaktues të gjendjes së përgjithshme, të gjendjes pneumatike dhe asaj dermatologjike.
- o Nuk vihet re ndryshim midis dy fshatrave që ndodhen në zonat industriale. Mesataret e nënpullimeve përkatëse konsiderohen të njehta.

4.2 Analiza e Komponentëve Kryesorë

Le të jenë $X_1, X_2, \dots, X_p$ p variabla të rastit që kanë shpërndarje shumë përmase, jo dosmosdoshmërisht normale me pritje matematike $\mu^{p \times 1} = (\mu_1, \mu_2, \dots, \mu_p)'$ dhe matricë kovariance . Zakonishti tregohet shume kujdes në aplikime studimit të kesaj matrice kovariance, pra studimit të lidhjes mes variablave $X_1, X_2, \dots, X_p$. Studimi I kësaj lidhje ndryshe, njihet me emrin, studimi I *struktures së variancës*. Struktura e variancës mund të përcaktohet ose nga llogaritja e kovariancës mes variablave, ose e koorelacionit midis tyre.

Qëllimi i këtij aplikacioni është të gjenden, nëse është e mundur, disa funksione $Y_1, Y_2, \dots, Y_q$, ku çdo Y_i është kombinim linear i $X_1, X_2, \dots, X_p$ ku ($q < p$), i cili afersisht gjeneron strukturën mes $X_1, X_2, \dots, X_p$. Për me tepër të gjendet strukturë e re varësie, që kjo të mund të sigurojë të njëjtën sasi informacioni si variablat origjinalë.

Më poshtë paraqitet një metodë që njihet me emrin *analiza e komponentëve kryesore*, që tregon pikërisht parimet e punës së kësaj analize.

Metoda synon të gjejë p kombinime lineare

$$Y_1 = \sum_{j=1}^p \alpha_{1j} X_j, \dots, Y_p = \sum_{j=1}^p \alpha_{pj} X_j$$

të tillë që

$$\text{cov}(Y_i, Y_j) = 0 \text{ për } i, j = 1, \dots, p \text{ nese } i \neq j,$$

$$D(Y_1) \geq D(Y_2) \geq \dots \geq D(Y_p)$$

dhe

$$\sum_{j=1}^p D(Y_j) = \sum_{j=1}^p \sigma_{jj}$$

Nga barazimet e mësipërme duket sikur variablat e rinj janë të pakooreluar dhe të radhitur sipas dispersioneve, kovariancës së tyre. Për më tepër *totali I variancës mbetet I njëjtë* pas transformimit, në këtë mënyrë, nënbashkësia e q variablave të rinj të parë, Y_i për $i=1, 2, \dots, q$ mund të shpjegojë shumicën e variancës totale dhe për më tepër të sigurojë një përshkrim të strukturës së varësisë mes variablave origjinalë. Për përcaktimin e koeficientëve α_{ij} , për $i, j = 1, \dots, p$ përdoret metoda e komponentëve kryesorë.

Shpërndarja e variablave $X_1, X_2, \dots, X_p$ përcaktohet tërësisht nga μ dhe Σ . Me supozimin që $\mu = (0, 0, \dots, 0)'$, struktura e varësisë mes variablave që jipet nga Σ , shpjegon tërësisht shpërndarjen e $X_1, X_2, \dots, X_p$.

Duke ditur që Σ është e njohur, le ta kërkojmë $Y_1 = \alpha_{11}X_1 + \dots + \alpha_{1p}X_p$. Duhet të gjejmë $\alpha_{11}, \dots, \alpha_{1p}$, të tillë që

$$D(Y_1) = \sum_{i=1}^p \sum_{j=1}^p \alpha_{1i} \alpha_{1j} \sigma_{ij}$$

te ketë vlerë maksimale dhe të plotësohet kushtit që vendoset për të siguruar unicitetin e zgjidhjes, $\sum_{j=1}^p \alpha_{1j}^2 = 1$

Zgjidhja $\alpha_1 = (\alpha_{11}, \alpha_{12}, \dots, \alpha_{1p})'$ quhet vektorë të vetë dhe shoqërohet me me të madhin vlerë të vetë të Σ . Kjo vlerë e vektorit të vetë është e barabartë me variancën $D(Y_1)$ kombinimi linear $Y_1 = \alpha_{11}X_1 + \dots + \alpha_{1p}X_p$ quhet edhe komponenti I parë kryesor I komponentëve të $X_1, X_2, \dots, X_p$, dhe ai shpjegon $100 \cdot D(Y_1)/V$ përqind të variancës totale.

Pastaj, kërkojmë të përcaktojmë $Y_2 = \alpha_{21}X_1 + \dots + \alpha_{2p}X_p$, dhe të vlerësojmë $\alpha_{21}, \alpha_{22}, \dots, \alpha_{2p}$.

Përcaktimi I këtyre koeficientëve bëhet I tillë që të maksimizojë

$$D(Y_2) = \sum_{i=1}^p \sum_{j=1}^p \alpha_{2i} \alpha_{2j} \sigma_{ij}$$

dhe $\sum_{j=1}^p (\alpha_{2j})^2 = 1$ (*) si dhe $\text{cov}(Y_1, Y_2) = \sum_{i=1}^p \sum_{j=1}^p \alpha_{1i} \alpha_{2j} \sigma_{ij} = 0$ (**)

Kushti (*) si guron unicitetin e zgjidhjes ndërsa (**) siguron që Y_1 dhe Y_2 të jenë të pakooreluar. Zgjidhja $\alpha_2 = (\alpha_{21}, \alpha_{22}, \dots, \alpha_{2p})'$ është vektor I vetë dhe shoqërohet me vlerën e vetë të dytën për nga madhësia e të Σ . Kjo vlerë e vetë është e barabartë me variancën $D(Y_2)$. Kombinimi linear $Y_2 = \alpha_{21}X_1 + \dots + \alpha_{2p}X_p$ është dhe komponenti I dytë kryesor i komponentëve të $X_1, X_2, \dots, X_p$, dhe, të dy komponentët e parë kryesorë shpjegojnë $100 \cdot (D(Y_1) + D(Y_2))/V$ përqind të variancës totale.

Pasi $Y_1, Y_2, \dots, Y_{q-1}$, për $q=2, 3, \dots, p$ të jenë marre në këtë mënyrë fitohet $Y_q = \alpha_{q1}X_1 + \dots + \alpha_{qp}X_p$ I tillë që

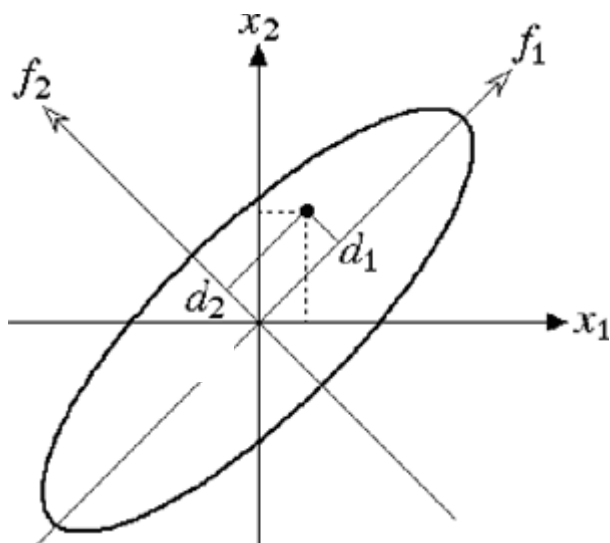
$$D(Y_q) = \sum_{i=1}^p \sum_{j=1}^p \alpha_{qi} \alpha_{qj} \sigma_{ij}$$

dhe që plotëson kushtin $\sum_{j=1}^p \alpha_{qj}^2 = 1$ dhe $\text{cov}(Y_m, Y_q) = \sum_{i=1}^p \sum_{j=1}^p \alpha_{qi} \alpha_{mj} \sigma_{ij} = 0$

për $m=1, \dots, q-1$. Kështu fitojmë vektorin e vetë $\alpha_q = (\alpha_{q1}, \alpha_{q2}, \dots, \alpha_{qp})'$ që I korespondon vlera e vetë e q-të më e madhe $D(Y_q)$. Kështu Y_q është komponenti i q-të kryesor, dhe variablat $Y_1, Y_2, \dots, Y_q$ shpjegojnë $100 \cdot (D(Y_1) + D(Y_2) + \dots + D(Y_q))/V$ përqind të variancës totale.

Një interpretim gjeometrik I problemit në fjale mund të shihet në figuren, për $p=2$. Çdo variabël $X_1, X_2, \dots, X_p$, paraqitet gjeometrikisht nga një bosht koordinativ me origjinë në $\mu^{p \times 1} = (\mu_1, \mu_2, \dots, \mu_p)'$. Këto p boshte në hapësirën p dimensionale, bëjnë të mundur konceptimin e

$X = (x_1, x_2, \dots, x_p)$ si një pikë koordinatat e të cilës janë $X_1=x_1, X_2=x_2, \dots, X_p=x_p$.

Figura 4-1 Interpretimi gjeometrik i metodës së komponentëve kryesor në rastit $n=2$ 

Problemi i analizës së komponenteve, sipas kësaj paraqitje gjeometrike, mund të ishte i njëjtë si të rrotulluarit e këtyre boshteve që variabli Y_1 i paraqitur nga i pari bosht i ri të ketë dispersion maksimal. Variabli i dytë Y_2 i paraqitur në boshtin tjetër të ri është i pakoreluar me Y_1 dhe ka një dispersion maksimal. Në mënyrë të ngjashme, variabli Y_q paraqet boshtin e ri të q -të që është i pakoreluar me $Y_1, Y_2, \dots, Y_{q-1}$ dhe ka maksimumin e dispersionit, pa e thyer kushtin e lartpërmendur. Kështu nëse $f(x)$ do të ishte një densitet normal i vektoreve të rastit $X = (x_1, x_2, \dots, x_p)'$, atëherë inekuacioni $f(x) \leq c$, ku c është një konstante, përcakton si të thuash në hapësirën p -dimensionale një elipsoid. Drejtimi I boshteve kryesore të këtij elipsoidi përcakton dhe drejtimin e komponenteve të rinj Y_i .

I parë në hapësirën dy dimensionale të përcaktuar nga X_1 dhe X_2 me origjinë në (μ_1, μ_2)

elipsoidi është elipsoid i zakonshëm. Komponentja e parë është në drejtimin e boshtit të madh të elipsit, ndërsa komponentja e dytë kryesore është në drejtim të boshtit të vogël të elipsit.

Pra, problemi i përcaktimit të komponentëve kryesorë të rinj bazohet në matricën e kovariancës Σ . Nëse kjo matricë është e panjohur atëherë në vend të saj do të përdoret matrica e kovariancës së zgjedhjes.

4.2.1 Vërejtje

1. Kryerja e metodës MANOVA në programet kompjuterike, si metode me shumë variabla bëhet nëpërmjet analizës 'multivariate'. Dhe kjo ka arsyen e saj, theksuam se në këto metode kemi shumë variabla të varur dhe një variabël në dukje I pavarur me to. Por, sipas Tabachnik dhe Fidell (2007)^[14] është e pakuptimtë përdorimi I MANOVA-s nëse të gjithë variablat e varur masin të njëjtën gjë. Ato duhet të mendohet se I përkasin, të paktën dy grupeve të ndryshme.

2. Në dallim me procedurat e tjera të analizës shumë përmase, të cilat kanë si parim reduktimin e numrit të variablave të nevojshëm për tu matur, procedura në fjalë

kërkon që të gjithë p variablat të maten, sepse ato të gjithë do të ndikojnë në llogaritjen e vlerës së komponenteve kryesore. Pra, është e pakuptimtë të perdoret nëse qëllimi i studimit është se cilët nga variablat e duhet të hiqen, që të marr më pak variabla në studim.

3. Homogjeniteti i variancës me grupeve është shumë i rëndësishëm. Siç dihet kontrolli në fjalë bëhet përmes testit Levene. Metoda interesohet që varianca mes grupeve të jetë e njëjtë, por, duhet theksuar, se ky kusht do të konsiderohet që nuk plotësohet kur p -vlera e afishuar të jetë <0.005 , jo si zakonisht <0.05 .^[13]

4. Në qoftëse $X_1, X_2, \dots, X_p$ kanë shpërndarje normale shumëpërmase, atëhere komponentët kryesorë $Y_1, Y_2, \dots, Y_q$, duke qenë se janë çift e çift te pakooreluar, do të jone çift e çift normale.

5. Vërehet se varianca totale e komponentëve kryesorë të rinj është sa varianca mes variablave të varur. Por, duke patur parasysh se për çdo variabël të varur X_i mund të behet normalizimi i tij, varianca totale V do të jetë sa numri i variablave të varur X_i .

6. Njehsimi i koorelacionit mes X_i dhe Y_j jipet nga $\frac{\alpha_{ij}(D(Y_j))^{1/2}}{\sigma(X_i)}$, ndaj për të vlerësuar

kontributin e $X_1, X_2, \dots, X_p$ në Y_j , përdoren pikërisht $\frac{\alpha_{ij}}{\sigma(X_i)}$. Po nëse vlerësimi behet duke u bazuar tek kovariancat atëhere përdoren vetëm α_{ij} për vlerësimin e komponentëve X_i që ndikojnë në përcaktimin e komponentit të ri kryesor. Kur ky koeficientë është më i lartë, atëhere më i lartë është kontributi i variablit te varur.

4.2.2 Zbatimi 2(vazhdim)

Po duhet theksuar se vlera me e madhe e metodës MANOVA verëhet në mënyrën tjetër që mund të kryhet kjo metode, në analizën e komponenteve kryesore. Përcaktimi i ketyre komponentëve kryesor është një përcaktim i i variablave të rinj që janë kombinim linear i 13 variablave te varur ne studim, e shpesh për këtë arsye njihen me emrin përcaktimi i supervariablave. Këto variabla të rinj, supervariabla^[15], të pavarur synojnë të paraqesin të njëjtat matje apo informacione të marra nga 13 variablat në studim.

MANOVA për analizën e komponentëve kryesorë do te ekzekutohet pas shkrimit të kodit në editorin e sintaksës. Editori është emertuar ' *pyetje*' dhe permban kodin e meposhtëm:

MANOVA p23 p24 p25 gelbaze kolle p28 p29 p30 p31 p32 p33 p34 p35 BY fshatrat(1,3)

/DISCRIM=STAN RAW CORR

/PRINT =SIGNIF(MULTIV,UNIV,EIGEN,DIMENR)/DESIGN.

Ekzekutimi i kësaj analize siguron gjetjen e dy kombinimeve lineare të 13 variablave në studim. Shihet se keto dy supervariabla të krijuar, ruajnë totalin e variances.

Tabela 4-3 Vlerat e veta të komponentëve kryesorë

Root No	Eigenvalue.	Pct.	Cum. Pct.	Canon Cor.
1	1.2537	100.00000	100.00000	.74585
2	.00000	.00000	100.00000	.00000

Por, afishimet e Tabeles 4.3, tregojnë se supervariabli i parë shpjegon $0.74585^2 = 0.556\%$ të variancës totale. Ndërsa supervariabli I dytë i korespondon një vlerë e vete shumë e vogël $2.3E-5$ dhe, për rrjedhojë dhe sasia e variancës që shpjegon supervariabli I dytë, siç shihet edhe nga afishimet e tabelës 1, është e papërfillshme (<0.005). Kështu studimi përqëndrohet vetëm në supervariablin e parë me vlerë të vetë 1.25375.

superndikimi = $-0.31412 \cdot p_{23} - 0.23080 \cdot p_{24} - 0.50926 \cdot p_{25} - 0.32348 \cdot p_{26} - 0.3739 \cdot p_{27} - 0.21248 \cdot p_{28} - 0.29417 \cdot p_{29} - 0.24600 \cdot p_{30} - 0.64955 \cdot p_{31} - 0.34177 \cdot p_{32} - 0.39591 \cdot p_{33} - 0.34184 \cdot p_{34} - 0.65818 \cdot p_{35}$

Në rastin në fjalë problemi i studimit të ndikimit të 'fshatra'-ve tek 13 variablat e varur reduktohet në studimin e 'fshatra'-ve tek variabli kanonik 'superndikim'.

4.3 Analiza Faktoriale

Në seksioni 4.2 u tregua se teknika e komponenteve kryesore ishte një metode që bazohej në varësinë mes $X_1, X_2, \dots, X_p$ p variabla të rastit që kanë shpërndarje shumë përmase, me pritje matematike $\mu^{p \times 1} = (\mu_1, \mu_2, \dots, \mu_p)'$ dhe matricë kovariance $\Sigma^{p \times p}$.

Komponentët kryesorë u kërkuar si p kombinime lineare

$$Y_1 = \sum_{j=1}^p \alpha_{1j} X_j, \dots, Y_p = \sum_{j=1}^p \alpha_{pj} X_j \quad (*)$$

të variablave të pakoreluar $X_1, X_2, \dots, X_p$, dhe rradhiten sipas dispersionit të tyre

$$D(Y_1) \geq D(Y_2) \geq \dots \geq D(Y_p)$$

dhe

$$\sum_{j=1}^p D(Y_j) = \sum_{j=1}^p \sigma_{jj} = V$$

Por, për ekuacionet (*) si ekuacione lineare, mund të shtrohet dhe problemi I anasjelltë, pra që variablat e varur të shprehin si kombinim linear I Komponentëve kryesorë. Pra

$$X_1 = \sum_{j=1}^p \beta_{1j} Y_j, \dots, X_p = \sum_{j=1}^p \beta_{pj} Y_j$$

për disa konstante β_{ij} , $i, j = 1, 2, \dots, p$. Mund të tregohet⁽⁵⁾ që $\beta_{ij} = \alpha_{ji}$, ndaj $X_1 = \sum_{j=1}^p \alpha_{j1} Y_j$,

$$\dots, X_p = \sum_{j=1}^p \alpha_{jp} Y_j$$

Nga këto ekuacione Modeli I Komponentëve Kryesorë që plotëson kushtin

$$D(Y_2) = \sum_{i=1}^p \sum_{j=1}^p \alpha_{2i} \alpha_{2j} \sigma_{ij} = \sum_{i=1}^p \sum_{j=1}^p \alpha_{2i} \alpha_{2j} \sigma_{ij}$$

$$\text{tani merr formën } \sigma_{ij} = \sum_{k=1}^p \alpha_{ki} D(Y_k) \alpha_{kj}$$

$$\sigma_{ii} = \sum_{k=1}^p \alpha_{ki}^2 D(Y_k) \text{ pwr } i, j = 1, 2, \dots, p$$

Këto ekuacione paraqesin faktorizimin e variancës dhe kovariancës së variablave të varur si funksion të α_{ij} dhe të Variances së Variablave Kryesorë.

Analiza faktoriale është një disiplinë e tërë, por qëllimi I këtij seksioni nuk është të trajtojë problemin në tërësi. Qëllimi është të shohim se si kjo metode na ndihmon në përcaktimin e Faktorëve Kryesorë. Programet statistikore kompjuterike kanë paketa të gatshme për kryerjen e kësaj metode, ndonëse në dukje ajo është shumë e vështirë.

Vlerësimi I parë I kësaj metode është bërë që më 1967 nga Morrisin^[16], I cili sygjeroi që për përcaktimin e Faktorëve Kryesorë, të punohej me teknikën e rrotullimit. Paketat e programeve kompjuterike specifikojnë:

- numrin e Faktorëve Kryesorë që mjafton të studiohen
- formën e matricës , që përcakton këto faktorë kryesorë
- vlerëson numrin e hapave të nevojshëm që ky vlerësim të bëhet I vlefshëm.

Nëse shtrohet problemii numrit të faktorëve kryesorë që është i nevojshëm të merren në shqyrtim sygjerohet^[17]

- të merren faktorë që vlerat e veta t'i kenë >1
- këto komponentë në tërësi justifikojnë së bashku rreth 70-75% të variancës totale.

4.3.1 Zbatimi 3(vazhdim)

Nëse shembullin e ndikimit të 'fshatra'-ve tek 13 pyetjet që përcaktonin 13 variabla kategorikë, e shohim sipas Analizës së Komponentëve Kryesorë, duhet të përmëndim që 13 pyetjet ndahen në tri kategori, në pyetje për gjendjen e përgjithshme të individit, për gjendjen pneumatike dhe dermatologjike. Kryerja e kësaj analize afishon tabelën e vlerave të veta, Tabela 4.4, që në bazë të verejtjes së mësipërme tregon se numri I faktorëve kryesorë duhet të jetë 3, e këto tri faktorë së bashku justifikojnë 58.9% të variancës totale.

Tabela 4-4 Tabela e vlerave të veta

Total Variance Explained

Comp.	Initial Eigenvalues			Extraction Sums of Squared Loadings			Rotation Sums of Squared Loadings		
	Total	% of Variance	Cumulative %	Total	% of Variance	Cumulative %	Total	% of Variance	Cumulative %
1	4.313	33.174	33.174	4.313	33.174	33.174	3.091	23.776	23.776
2	2.103	16.179	49.353	2.103	16.179	49.353	2.614	20.111	43.887
3	1.231	9.466	58.820	1.231	9.466	58.820	1.941	14.933	58.820
4	.988	7.602	66.421						
5	.885	6.804	73.225						
6	.793	6.097	79.322						
7	.618	4.754	84.076						
8	.493	3.790	87.866						
9	.450	3.463	91.329						
10	.368	2.828	94.157						
11	.311	2.390	96.547						
12	.264	2.035	98.582						
13	.184	1.418	100.000						

Duke u mbeshtetur ne vlerat e afishuara nga kjo metodë për përcaktimin e faktorëve arrihen të ndertohen tri faktorët ne editorin e sintaksës, si mëposhtë.

Komp_dermatologjik= .324* p23+ .460*p24+.312*p25+.142* gelbaze-.062* kolle +.834*p28 +.773*p29+.872* p30+.020* p31+.692* p32-.168* p33+.161* p34+.256*p35

Komp_pergjith= .542* p23+ .273*p24+.614*p25+.098* gelbaze+.063* kolle +.026*p28 +.165*p29+.104* p30+.437* p31+.108* p32+.830* p33+.576* p34+.770*p35

Komp_pneumologjik= 122* p23+.178*p24+.405*p25+.730* gelbaze+.783* kolle -.084*p28 +.034*p29-.068* p30+.600* p31+.310* p32-.104* p33+.275* p34+.167*p35

Theksojme se faktoret e ndertuar janë kombinim linear i 13 variablave të varur, por nisur nga vlera me e lartë e koeficientëve prane pyetjeve të karakterit dermatologjik në faktorin e parë, të pyetjeve të karakterit të përgjithshëm në faktorin e dytë dhe atij pneumologjik në të tretin është bërë dhe emërtimi I faktorëve. Të tri faktoret jane përberës të konsiderueshëm te variacionit të përgjithshëm, ndaj emertohen *komponente. (shih problemin 1, Kapitulli 6)*

Pastaj studiohet problemi I ndikimit të 'fshata'-ve në këto tri faktorë kryesore.

5 Modeli i regresit logjistik të shumëfishtë

5.1 Regresi logjistik dhe llojet e tij

Në seksionin 2 është trajtuar metoda e regresit të shumëfishtë linear (GLM). Por, lidhjet mund të jenë më të komplikuar se

$$y_i = \beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2} + \cdots \beta_p x_{ip} + \varepsilon_i \quad \text{për } i=1..n$$

ku duhet të vlerësojmë $\beta_0, \beta_1, \beta_2, \dots, \beta_p$ që janë parametrat e këtij modeli dhe $\varepsilon_1, \varepsilon_2, \dots, \varepsilon_n$, n variabla të pavarur $N(0, \sigma^2)$. Në këtë seksion do të përqëndrohemi në modelin e regresionit logjistik, duke trajtuar problemin e llojshmërisë së tij në përputhje me problemin e shtruar, vlerësimet që ato sigurojnë dhe, më pas, mënyren se si ky model kontrollohet për vlefshmerinë e tij.

5.1.1 Modeli i regresit logjistik të shumëfishtë

Një model regresi përmban tri komponentë: komponentin e rastit, komponentin sistematik dhe funksionin lidhës.

Komponentët e rastit janë $Y_1, Y_2, \dots, Y_n$ që mendohet se janë ndryshore rasti të pavarura, me shpërndarje nga familja eksponenciale. Pra, shpërndarja e Y_i nuk është identike, por i takon të njëjtës familje eksponenciale.

Komponenti sistematik është modeli. Ai është funksion i parametrave.

$$y_i = \beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2} + \cdots \beta_p x_{ip}$$

Në modelin e regresit linear, është i lidhur me mesataren e Y_i .

Dhe së fundmi, *funksioni lidhës* i dy komponentëve, që teorikisht paraqitet:

$$g(\mu_i) = \beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2} + \cdots \beta_p x_{ip}, \text{ ku } \mu_i = E(Y_i).$$

Dimë se një ndër kushtet që duhet të plotësoheshin në këtë model ishte, Y_i të kishin shpërndarje normale. Por, meqënëse familja normale është pjesë e asaj eksponenciale atëherë ky lloj modelimi u përdor edhe kur Y_i kishin shpërndarje nga e njëjta familje, jo vetëm normale, por dhe Puasoni e Binomiale. Në rastet kur komponenti i rastit merr një numër të kufizuar vlerash, propozohet të përdoret metoda e modelimit të regresit logjistik.

Në këtë model, kur variablat e rastit $Y_1, Y_2, \dots, Y_n$, janë të pavarura dhe kanë shpërndarje Bernuli (π_i), lidhja mes x_i dhe π_i kërkohet në formën

$$\ln\left(\frac{\pi_i}{1-\pi_i}\right) = \beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2} + \cdots \beta_p x_{ip} = \text{logit}(\pi_i) \quad (1)$$

Ana e majte paraqet $\ln$ e mundësive të suksesit për Y_i (që njihet me emrin *logit*)^[18]. Modeli, siç shihet, propozon që $\ln$ e mundësive të suksesit të jetë funksion linear i x -it.

Densiteti probabilitar i vektorit të rastit, në këtë rast, mund të shkruhet në formën:

$$p^y(1-p)^{1-y} = (1-p)\exp\{y \ln(p/(1-p))\}$$

Termi $\ln(\frac{p}{1-p})$ është parameter natyror i familjes eksponenciale, dhe siç shihet si funksion lidhës është përdorur $g(p) = \ln(\frac{p}{1-p})$.

Ekuacioni (1) mund të shkruhet si:

$$p(x_i) = \frac{e^{\beta_0 + \beta_1 x_{i1} + \dots + \beta_n x_{in}}}{1 + e^{\beta_0 + \beta_1 x_{i1} + \dots + \beta_n x_{in}}} \quad (2)$$

ose, në mënyrë më të përgjithshme, si:

$$p(x) = \frac{e^{\beta_0 + \beta_1 x_1 + \dots + \beta_n x_n}}{1 + e^{\beta_0 + \beta_1 x_1 + \dots + \beta_n x_n}} \quad (3)$$

Shihet se $0 < p(x) < 1$ e në rastin në fjale, kjo përputhet me faktin që, ai është probabiliteti.

5.1.2 Modeli i regresit logjistik i thjeshtë me një variabël në modelim

Si lloj i vecante i regresit logjistik mund të konsiderohet edhe regresi logjistik i thjeshtë, që është rasti i regresit logjistik të shumëfishtë, por me një variabël në modelim. Kjo është forma më e thjeshtë e regresit logjistik^[19].

Ai merr rëndësi edhe teorike, për të kuptuarit më mirë të komponenteve të këtij regresit dhe interpretimin korrekt të tyre. Pra edhe për lehtësi në studim e $p(x)$, komponentin sistematik mund ta paraqesim, duke e kufizuar në 1 numrin e variablave që marrin pjesë në modelim.

Atehere (3) do merrte formën

$$p(x) = \frac{e^{\beta_0 + \beta x}}{1 + e^{\beta_0 + \beta x}}$$

dhe derivati mund të shkruhej:

$$\frac{dp(x)}{dx} = \beta p(x)(1 - p(x)) \quad (4)$$

Meqë termi $p(x)(1 - p(x)) > 0$, atehere shenja e derivatit varet nga shenja e β . Nqs $\beta > 0$ atëhere derivati i $p(x)$ është rigorozisht pozitiv; nqs $\beta < 0$ atëhere derivati $p(x)$ është rigorozisht negativ. Pra, nqs $\beta > 0$ atëhere $p(x)$ është funksion rigorozisht rritës; nqs $\beta < 0$ atëhere $p(x)$ është funksion rigorozisht zbritës; nqs $\beta = 0$ atëhere $p(x) = \frac{e^{\beta_0}}{1 + e^{\beta_0}}$ për çdo x . Në këtë rast si në modelin e regresit linear, nuk ka lidhje mes p dhe x .

Parametrat β_0 dhe β kanë pothuajse kuptim të ngjashëm me ato të regresit linear të zakonshëm.

Nëse $x=0$, po të shihet ekuacioni (1) kuptohet se, β_0 paraqet $\ln$ e raportit të mundësisë së suksesit në $x=0$. Nëse duke zëvendësuar tek (1) me x dhe $x+1$ kemi se për çdo x është e vërtetë:

$$\ln \frac{p(x+1)}{1-p(x+1)} - \ln \frac{p(x)}{1-p(x)} = \beta_0 + \beta(x+1) - \beta_0 - \beta(x) = \beta$$

Kështu β ndryshimi i $\ln$ së mundësive të suksesit që ndodh nëse x rritet me 1 njësi. (*) Duke marrë eksponencialin e të dy anëve të këtij ekuacioni, fitohet

$$e^\beta = \frac{p(x+1)/(1-p(x+1))}{p(x)/(1-p(x))} \quad (5)$$

Ana e djathtë paraqet raportin e mundësive dhe krahason mundësitë e suksesit në $x+1$ me ato të suksesit në x ^[21]. Përfundimisht,

$$\frac{p(x+1)}{1-p(x+1)} = e^\beta \frac{p(x)}{1-p(x)} \quad (6)$$

Pra, e^β është faktori që lidh ndryshimet e mundësisë së suksesit, që i korespondojnë rritjes me një njësi të x -it.

Ekuacioni (3) sygjeron një tjetër mundësi të modelimit të probabilitetit të Bernulit $p(x)$ si funksion i x -it. Dihet se densiteti probabilitar i n.r. *logistic(0,1)* përcaktohet nga barazimi

$$F(w) = \frac{e^w}{(1+e^w)}$$

Kështu, barazimi (3) mund të shkruhet në formën:

$$p(x) = F(\beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_p x_p)$$

Interpretimi i dhënë tek (*) është një interpretim mese i saktë, por nëse e shohim në këndvështrimin që $p(x)$ paraqet një kurbë, duke u bazuar në ekuacionin (4), themi se nëse ndërtohet tangentja me kurbaturën ^[20], në një pikë x , koeficienti këndor i tangentës do të jetë $\beta * p(x)(1-p(x))$.

Kështu për $p(x)=1/2$, ky koeficient e ka vlerën $\beta * \frac{1}{2} * \frac{1}{2} = \beta/4$.

Shihet se $p(x)=1/2$, arrihet në $x = -\beta_0/\beta$.

Kjo vlerë e x konsiderohet edhe si niveli i efektit të medianës. Në studimet mjekësore kjo quhet ndryshe doza e rrezikshme shkatërruese ^[23].

Duke u nisur nga perafrimi linear i kurbes, Agresti ^[22] sygjeron të vërohet se ç'ndodh kur kemi një ndryshim të x me $\frac{1}{\beta}$.

Këtij ndryshimi i korespondon një ndryshim afërsisht sa $\frac{1}{\beta} * \frac{\beta}{4} = 1/4$ i $p(x)$. Pra pikat ku $p(x)=0.25$ dhe 0.75 (në të vërtetë kjo ndodh në 0.27 dhe 0.73) dhe në 0.50 , janë me distancë $\frac{1}{\beta}$ nga njëra tjetra. (**)

Kur $p(x)=0.9$ ose 0.1 , koeficienti këndor do të jetë 0.09β .

5.1.2.1 Përfundime

Interpretimi i koeficientëve në rastin logjistik të thjeshtë me një variabël në modelim, është:

- kur x rritet me një njësi, β paraqet ndryshimin e $\ln$ së mundësive të suksesit, ndërsa

- e^{β} është faktori që lidh ndryshimet e mundësisë së suksesit.

5.1.3 Modeli logjistik i thjeshtë me parametra kategorikë

Regresi logjistik, e ndërton modelin edhe nëse variabli që duhet të modelojmë është kategorik.

Për të thjeshtuar paraqitjen e problemit, supozohet se duam të punojmë me një parametër të tillë të vetëm, me I kategori, pra në rreshtin e i -të të tabelës përmbledhese të të dhënave do të ishte y_i , numri i ndodhjes së vlerës përkatëse të variablit që duam të modelojmë, gjatë n_i rasteve që shfaqet kjo vlerë e variablit kategorik. Y_i mund të shihet si variabël binomial me parametër p_i . Modeli logjistik i tij me një faktor është:

$$\ln \frac{p_i}{1-p_i} = \beta_0 + \beta_i \quad (7)$$

Sa më e lartë β_i aq më e lartë vlera e p_i ,

Po ana e djathtë e modelit të mësipërm është modeli ANOVA njëpërmasor. Si në metodën në fjalë modeli ka aq parametra sa nivele ka faktori i shqyrtuar. Por, këtu një nga nivelet e faktorit është reduktuar. Pra faktori i shqyrtuar me I kategori ka tani $i-I$ kategori jo të reduktuara dhe kategorine e I te reduktuar. Një ndër parametrat merret kështu 0. Nëse $\beta_I = 0$ atëherë β_0 është logit në rreshtin I (zëvendësoj tek (7) indeksin $i = I$ dhe fitoj: $\text{logit}(p_I) = \beta_0$)

Ndaj β_i është diferenca mes *logit* të rreshtit të i dhe I. Kështu β_i është $\ln$ e raportit të mundësive, për këtë çift radhësh.

Nëse për faktorin në fjale ekziston një I që p_i të jetë pozitive, pra të ekzistojë $\beta_i \neq 0$, atëherë modeli është i kënaqshëm^[25]. Ndërsa nëse të gjitha $\beta_i = 0$, atëherë themi se X lidhet me Y vetëm një term konstant.

5.1.3.1 Përfundime

- Regresioni logjistik përdoret në situata kur duam të parashikojmë prezencën ose mungesën e një karakteristike në subjektet, individët, duke u bazuar në vlerat e një bashkësie variablash të studiara mbi to.
- Kjo metodë pranon, që variablat në bazë të së cilës ndërtohet modeli, të jenë variabla kategorikë dhe të vazhdueshëm.
- Nëse tërë variablat që marrin pjesë në modelim janë kategorikë, metoda është shumë efektive.
- Nëse ndonjë nga variablat e modelit është i vazhdueshëm, ai mund të përcaktohet 'covariate', dhe modeli bëhet duke punuar me variablin si në rastin e Analizës së Kovariancës (seksioni 3.3)
- Koeficientët e kësaj metode përdoren për të vlerësuar raportet e mundësive për secilin nga variablat e pavarur në model, nëse individi, subjekti i takon një grupi të përcaktuar nga vlera të caktuara të variablave të tjerë, që marrin pjesë në modelim.

5.1.3.2 Zbatim 1

Në problemin 3 Kapitulli 6, **Trajtimi i problemeve pneumologjike**, nga analiza e treguesve të përzgjedhur në këtë fushë u vu re se treguesi kryesor është treguesi 'kollë'.

Kërkohej, në fakt, të kontrollohet nëse ndikon apo jo *zona* në mënyrë të dukshme në këtë parametër pneumatik.

Për të analizuar problemin në fjalë janë shfrytëzuar të dhënat e 426 subjekteve të marra rastësisht nga tri fshatrat e sipërpërmendura, 196 subjekte janë nga zona e pastër, zona e kontrollit, pjesa tjetër 230 janë marrë nga zona industriale, e konsideruar e ndotur.

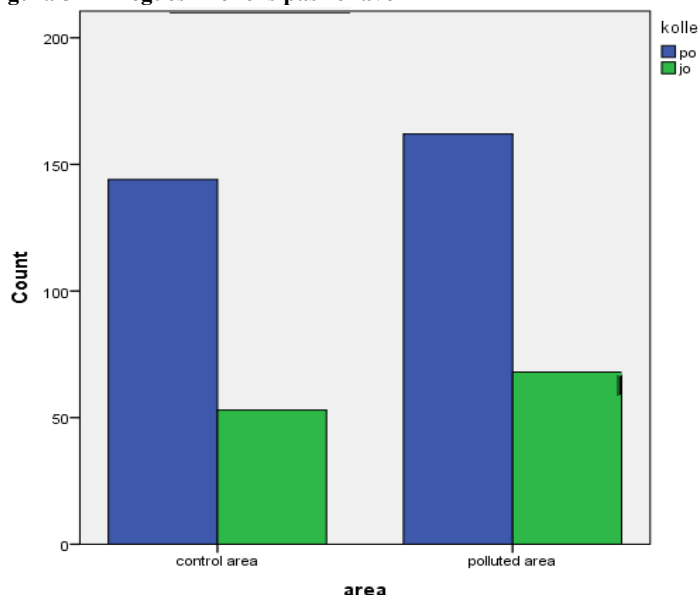
Në bazën tonë të të dhënave, variabli *kollë* nga transformimi i rekordeve, (shih Problemin 3), është i koduar, 0 nëse personi nuk ka kollë dhe 1 kur ka kollë. Pra është variabël binar.

Grafiket e paraqitur, Figura 5. 1, tregojnë për të dhënat e shfaqjes apo jo të 'kollës' në dy zonat. Grafiket kanë vetëm karakter informues, në kuptimin që, shihet se numri i rasteve me kollë është më i lartë në zonën e ndotur, por aty është më i lartë dhe numri i rasteve pakollë (vëllimi i zgjedhjes në zonën e ndotur është më i lartë). Nga llogaritjet mundësitë e kollës në zonën e pastër janë 36.7%, ndërsa në zonën e ndotur janë 42%. Ndaj, shtrohet problemi, a është kjo gjë statistikisht e vërtetë.

Për të kontrolluar këtë, për shkak të natyrës së variablit 'kollë', është përdorur metoda e regresit logjistik binar me variabla kategorikë.

Variabli i varur është konsideruar variabli '*kollë*', variabël binar; ndërsa variablat me të cilën është ndërtuar parashikimi është variabli '*zone*' (*area*), '*konsum_duhan_në_shtëpi*' dhe variabli '*konsum_duhan*', i trajtuar më parë si ndikues në '*kollë*'^[26].

Figura 5–1 Treguesi 'kollë' sipas zonave



Vëllimi i të zgjedhjes për shkak të mungesës së vlerave është 219, nga të cilët 112 nuk konsumojnë duhan pjesa tjetër konsumon. Kështu, nëse për një person të zonës do të

na duhej të konkludonim konsumon apo jo duhan, pa bërë asnjë modelim do të thonim jo. Sipas tabelës së klasifikimit ky pohim do të ishte i vërtetë në 51.1% të rasteve.

Tabela 5–1 Vlerësimi para modelimit

Classification Table^{a,b}

Observed	Predicted		
	kolle		Percentage Correct
	jo	po	
kolle jo	112	0	100.0
po	107	0	.0
Overall Percentage			51.1

a. Constant is included in the model.

b. The cut value is .500

Vlera $\exp(b)=0.955$ tregon raportin individëve e me kolle ndaj pakolle, ndryshe $1-0.955=0.053$ mund të interpretohetse 5% janë me shume rastet pa kolle se me kolle. Të gjitha këto vlerësime bëhen para modelimit. (në bllokun 0 të output-eve)

5.1.4 Regresioni logjistik multinominal

Regresi multinominal logjistik përdoret në shumë situata, kur duhet të klasifikohen subjektet bazuar në vlerat e një bashkësie variablash. Ky tip regresi është i ngjashëm me regresin logjistik, por është më i përgjithshëm, sepse nuk kërkon që patjetër variabli i varur të ketë vetëm dy kategori.

5.1.4.1 Përfundimisht:

- Regresi multinominal logjistik përdoret kur duhet të klasifikohen subjektet bazuar në vlerat e një bashkësie variablash.
- Variablat e pavarur që marrin pjesë në model mund të jenë kategorik ose të vazhdueshëm.
- Nëse ato janë kategorik, konsiderohen si faktorë gjatë përpunimit të të dhënave, ndërsa nëse janë të vazhdueshëm, si 'covariates'
- Ky model ndërtohet, nëse pranohet se mundësitë e raportit të dy kategorive të variablit, që do të modelohet, nuk varen nga kategorite e tjera të variablit.

5.2 Vlerësimi i parametrave

Në modelin e regresit linear, ndërtohej modeli

$$y_i = \beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2} + \cdots \beta_p x_{ip} + \varepsilon_i$$

e për vlerësimin e parametrave përdoret metoda e katrorëve më të vegjël. Ndërsa në rastin në fjalë kur $Y_i \sim \text{Bernuli}(p_i)$, ato nuk kanë më lidhje direkte me

$\beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2} + \dots + \beta_p x_{ip}$, kështu nuk mund të përdoret më metoda e katrorëve më të vegjël.

Metoda që përdoret për **vlerësimin pikësor** është Metoda e Përngjasisë Maksimale. Ndërtohet funksioni I Përngjasisë Maksimale^[25]

$$L(\beta_0, \dots, \beta_p) = \prod_{i=1}^n p(x_i)^{y_i} (1 - p(x_i))^{1-y_i} = \prod_{i=1}^n F_i^{y_i} (1 - F_i)^{1-y_i}.$$

Për të kontrolluar ekstremumin e tij, maksimumin, përdorim funksionin ndihmës të Funksionit të Përngjasisë

$$\ln(L(\beta_0, \dots, \beta_p)) = \sum_{i=1}^n \left\{ \ln(1 - F_i) + y_i \ln\left(\frac{F_i}{1 - F_i}\right) \right\}$$

dhe gjej derivatet e pjesshme.

Derivimi sipas β_0 siguron:

$$\begin{aligned} \frac{\partial \ln(1 - F_i)}{\partial \beta_0} &= -\frac{f_i}{1 - F_i} = -\frac{F_i f_i}{F_i(1 - F_i)} \\ \frac{\partial \ln\left(\frac{F_i}{1 - F_i}\right)}{\partial \beta_0} &= \frac{f_i}{F_i(1 - F_i)} \end{aligned}$$

Pra,

$$\frac{\partial \ln(L(\beta_0, \dots, \beta_p))}{\partial \beta_0} = \sum_{i=1}^n (y_i - F_i) \frac{f_i}{F_i(1 - F_i)} \quad (5)$$

Ndërsa derivimi sipas variablae të tjerë do të japë:

$$\frac{\partial \ln(L(\beta_0, \dots, \beta_p))}{\partial \beta_i} = \sum_{i=1}^n (y_i - F_i) \frac{f_i}{F_i(1 - F_i)} x_i \quad (6)$$

Por, për funksionin logjistik, meqë $F(w) = \frac{e^w}{(1 + e^w)}$

atëhere $\frac{f_i}{F_i(1 - F_i)} = 1$

Ndaj, (5) dhe (6) marrin formë më të thjeshtë.

Vlerësimi I parametrave të modelit mund të behet duke barazuar me 0 ekuacionet (5) dhe (6) dhe, duke vlerësuar me anë të këtyre ekuacioneve β_0 dhe β_i . Ekuacionet e formuara janë jo linear dhe duhen zgjidhur numerikisht.

Për regresionin logjistik, funksioni ndihmes i Përngjasisë Maksimale është rigorozisht i myset, kështu nëse ekuacionet në fjalë kanë zgjidhje. Kjo zgjidhje është e vetme, dhe është një vlerësim më Metodën e Përngjasisë Maksimale. Megjithatë qëllon, që në raste të vecanta të dhënash, ekuacionet nuk kanë zgjidhje. Maksimumi arrihet në këto raste kur parametri x shkon në $\mp\infty$. Kjo lidhet me faktin se, gjatë modelimit logjistik, u supozua se $0 < p(x) < 1$, por në një bashkësi të dhënash të caktuara, maksimumi I Funksionit të Përngjasisë arrihet në limit, me $p(x)=0$ ose 1. Në këto raste thuhet se probabiliteti i të dhënave konvergjon në 0, nëse modeli logjistik është I vërtetë.

Ashtu si për vlerësimin e parametrave përdoret Metoda e Përngjasisë Maksimale, edhe për të vlerësuar **variancën** përdoret, si në rastin një përmasor përdoret

Informacioni I Fisherit, por në formë më të përgjithshur, në atë të matricës së informacionit. Kjo matricë në rastin e dy parametrave ka formën:

$$\begin{pmatrix} -\frac{\partial^2 y}{\partial \beta_0^2} \ln L(\beta_0, \beta_1)/y & -\frac{\partial^2 y}{\partial \beta_0 \partial \beta_1} \ln L(\beta_0, \beta_1)/y \\ -\frac{\partial^2 y}{\partial \beta_0 \partial \beta_1} \ln L(\beta_0, \beta_1)/y & -\frac{\partial^2 y}{\partial \beta_1^2} \ln L(\beta_0, \beta_1)/y \end{pmatrix} \quad (7)$$

Duke punuar me ekuacionet (5) dhe (6) e regresit logjistik , matrica e informacionit (7) merr formen:

$$\begin{pmatrix} \sum_j n_j F_j (1 - F_j) & \sum_j x_j n_j F_j (1 - F_j) \\ \sum_j x_j n_j F_j (1 - F_j) & \sum_j x_j^2 n_j F_j (1 - F_j) \end{pmatrix} \quad (8)$$

dhe varianca e parametrave β_0, β_1 vlerësohen duke përdorur këtë matricë.^[21]

Për llogaritjen e variancës në bazë të numrit të informacionit^[22] $I(\tilde{\beta})$ përdoret formula,

$\text{Var}(\tilde{\beta}) \approx [1/I(\tilde{\beta})]$. Këtu në rastin shumë përmasor mund të përdorim, në vend të inversit të informacionit të Fisherit, inversin e matricës së informacionit. Duke patur parasysh formën e matrisës inverse dypërmasore, kjo matricë e anasjelltë do të ketë në diagonalen kryesore elemente diagonales kryesore të (7). Kështu, në këtë diagonale do të jenë vlerësuesit e $[se(\tilde{\beta}_0)]^2$ dhe $[se(\tilde{\beta}_1)]^2$ ku $se(\tilde{\beta})$ është devijimi standart i $\tilde{\beta}$.

Ashtu si në modelin e regresit linear , kontrollohen hipotezat $H_0: \beta=0$, pra a ka lidhje variabli që kërkojmë të modelojmë me variablat që marrin pjesë në modelim.

Për të kontrolluar këto hipoteza përdoret testi Wald. Ky test ka marrë emrin e statisticienit hungarez Abraham Wald. Sipas këtij testi parametrat vlerësohen nga vlerësimet e zgjedhjes.

Testi që përdoret për kontrollin e hipotezës në fjalë është :

$$\frac{\tilde{\beta}}{se(\beta)} \quad (9)$$

Ky test nëse është e vërtetë H_0 ka shpërndarje afërsisht normale, vecanërisht kur vëllimi i zgjedhjes është i madh.

Si zakonisht, nëse vlera e vrojtuar e testit bie në zonën kritike , atëhere hipoteza bazë bie poshtë.

Ndryshe për vlerësimin e ndikimit apo jo të parametrat në modelim mund të përdoret testi^[11]

$$-2 \ln \lambda(y^*) = 2[\ln L(\tilde{\beta}_0; \tilde{\beta}_1 | y^*) - \ln L(\tilde{\beta}_0, 0 | y^*)] \quad (10)$$

ku $\tilde{\beta}_0$ është vlerësimi për β_0 kur $\beta = 0$

Tregohet se testi në fjalë ka shpërndarje X_1^2 .

Kështu nëse vlera e vrojtuar e testit në fjalë plotëson kushtin $-2\ln\lambda \geq \chi_{1,nkr}^2$, atehere variabli në fjalë nuk ndikon në modelim.

5.2.1 Problemi i vlerësimit intervalor të parametrave:

Për modelin me vetëm një variabel sipas modelimit,

$$\text{logit}[p(x)] = \alpha + \beta x,$$

shqyrtoam kontrollin e hipotezës $H_0 : \beta = 0$. dhe testet me të cilat ky kontroll bëhet.

Kontrolli i saj bëhet, sic u përmend më sipër, me anë të **Testit Wald**-it që përdor funksionin logaritmik të përngjasisë maksimale për $\hat{\beta}$, me test: $Z = \hat{\beta}/SE$ ose **katrorin** e tij; kur H_0 është e vërtetë z^2 asimptotikisht ka shpërndarje χ_1^2 .

Metoda e **maksimizimit të raportit të funksionit të përngjasisë maksimale për dy vlera**, e kthen problemin, duke punuar me funksionin logaritmik, në maksimizimin e diferencës ndërmjet logaritmit të funksionit të përngjasisë maksimale në $\hat{\beta}$ dhe në $\beta = 0$. E ky test (-2LL), gjithashtu ka asimptotikisht një shpërndarje χ_1^2 .

Për zgjedhje me vëllim të madh, të tre testet zakonisht japin rezultate të ngjashme. Raporti Përngjasia Maksimale (-2LL) preferohet më shumë në krahasim me testin Wald. Merren me shumë informacione nëse përdoret kjo metode. Kur $|\beta|$ është relativisht e madhe, testi Wald nuk është aq i fuqishëm sa test (-2LL), dhe mund edhe të çojë në përfundime jo normale [Hauck dhe Donner (1977)].

Intervalet e besueshmërisë japin më shumë informacion se sa kontrolli i hipotezave. Një interval për β sigurohet duke u pershtatur statistiken që do të duhej për vlerësimin e testit $H_0 : \beta = \beta_0$. Intervali besimit është bashkësia e β_0 për të cilën testit me shpërndarje hi-katror merr vlerë jo më e madhe se $\chi_1^2(\alpha) = z_{\alpha/2}^2$. Sipas metodës Wald, kjo do të interpretohej $[(\hat{\beta} - \beta_0)/SE]^2 \leq z_{\alpha/2}^2$. Përfundimisht intervali besimit është $\hat{\beta} \pm z_{\alpha/2} * SE$.

Por, karakteristika të tjera mund të kenë rëndësi më të madhe se β , p.sh.përcaktimi i vlerës së $p(x)$ për x të ndryshme. Për $x = x_0$ të fiksuar, $\text{logit}[\hat{p}(x_0)] = \hat{\alpha} + \hat{\beta}x_0$ ka një SE që jepet nga rrënja katrore e llogaritur nga vlerat e vlerësuar të parametrave:

$$\text{var}(\hat{\alpha} + \hat{\beta}x_0) = \text{var}(\hat{\alpha}) + x_0^2 \text{var}(\hat{\beta}) + 2x_0 \text{cov}(\hat{\alpha}, \hat{\beta})$$

Një interval besueshmërie 95% për $\text{logit}[p(x_0)]$ është $(\hat{\alpha} + \hat{\beta}x_0) \pm 1.96SE$. Por, këto intervale vlejne nëse vëllimi i zgjedhjes është i madh. Duke zëvendësuar çdo skaj në transformimin e anasjelltë,

$$p(x_0) = \exp(\text{logit}) / [1 + \exp(\text{logit})] \quad (11)$$

merret një interval korespondues $p(x_0)$.

Çdo metodë e paraqitur mund të sigurojë interval besueshmërie dhe në teste me vëllim të vogël.

5.2.2 Zbatim 1 (vazhdim)

Tabela 5-2 tregon rezultatet e hipotezës se faktori, me anë të së cilit modeli po ndërtohet, mund të mos marrë pjesë në modelim. Sic përmendëm, kjo bëhet me ane të

testit Wald, vlerat e vrojuara të të cilit paraqiten në tabelë. Ndërsa p_vlerat përkatëse, ndihmojnë në marrjen e vendimit.

Rezultatet e regresit logjistik, tabela 5.4, tregojnë se faktorë të rëndësishëm në këtë modelim janë 'area' dhe 'duhan_shtëpi'.

Koeficienti $\beta = 1.028$ pranë 'areas', tregon se kur 'area' rritet me një njësi, pra kur kalohet në zonën e ndotur, dhe vlerat e variablove të tjerë që marrin pjesë në modelim nuk ndryshojnë, mundësitë e kollës rriten, ndaj mundësive për kollë në zonën e kontrollit. Mundësitë për kollë në zonën e ndotur janë $\exp(1.028)=2.798266$ e mundësive për kollë në zonën e kontrollit.

Në të njëjtin përfundim arrihet edhe duke parë vlerën pranë 'duhan_shtëpi' 0.122, faktorit tjetër domethënës në këtë modelim, rritja e tij, pra konsumi i duhanit, rrit mundësinë e shfaqjes së kollës ndaj rasteve pakollë. Vlera e kësaj rritje mund të vlerësohet me $\exp(0.122)=1.13$, Kështu, në zonën e ndotur mundësia që të shfaqet ky tregues i rëndësishëm i problemeve pneumologjike, është 13% më e lartë se në zonën e kontrollit.

Tabela 5–2 Koeficientët e variablove në ekuacionin e modelimit

Variables in the Equation

		B	S.E.	Wald	df	Sig.	Exp(B)	95% C.I. for EXP(B)	
								Lower	Upper
Step 1 ^a	area	1.028	.326	9.967	1	.002	2.796	1.477	5.293
	pini_duhan	.147	.489	.090	1	.764	1.158	.444	3.018
	duhan_shtepi	.122	.037	10.602	1	.001	1.130	1.050	1.216
	Constant	-2.267	.596	14.491	1	.000	.104		

a. Variable(s) entered on step 1: area, pini_duhan, duhan_shtepi.

5.3 Përputhshmëria e modelit logjistik

Në praktikë, nuk ka garanci që një model i caktuar i regresionit logjik përputhet mirë me të dhënat.

Në rastin e modelit të regresit linear shumëpërmasor, zgjidhja e këtij problemi është relativisht e lehtë. Ndërsa këtu nderlikohet.

Në këto raste problemi synohet të zgjidhet duke analizuar përputhshmërinë e modelit logjistik.

Kur variablat pjesëmarrës në modelim janë vetëm kategorikë, përdoret kjo metodë për vlerësimin e mungesës së përputhshmërisë : krahasohen denduritë e vlerave të modeluara, me denduritë e matura, për $y = 0$ dhe $y = 1$. Në çdo konfigurim të x , mund të shumëzohet probabiliteti i vlerësuar i dy rezultateve Y, me numrin e subjekteve në atë bashkësi, duke marrë një vlerësim për densitetin e parashikuar. Këto quhen *vlera të përputhshmerise*.

Test i përputhshmërisë krahason vlerat e vëzhguara, dhe vlerat e përputhshmerisë se parashikuar, duke përdorur testin Pirson X^2 . Për një numër të caktuar konfigurimesh, testi I Pirson X^2 ka shpërndarje Hi-katror $X^2(1)$.

Arsyeja për përdorimin e kësaj metode, testit të përshtatshmërisë, kur variablat në modelim janë kategorik , kuptohet se lidhet me faktin, se, nëse do të kishim ndërtuar tabelen e kontigjences për variablat e vazhdueshëm, ajo mund të përmbante vlera

shumë të vogla. Atehere mund të ishte e pamundur krahasimi me anë të testit Hi-katror.

Për më tepër, me rritjen e numrit të variablave që sigurojnë modelimin, grupimi i njëkohshëm i vlerave, për çdo kombinim vlerash të variablave nv model, mund të prodhojë një tabelë pasigurie, me një numër të madh qelizash, shumica e të cilave kanë vlerë të vogël.

Sipas **Hosmer dhe Lemeshow** (1980), problemi në këto raste zgjidhet kështu:

Vlerat e vëzhguara dhe të përputhshmërisë, ndahen në grupe sipas probabilitetit të suksesit të vlerësuar, duke përdorur të dhënat origjinale. Rruga që përdoret zakonisht, është ndarja në grupe (p.sh. 10) me vëllim të njëjtë. Çifti i parë i vlerave të vëzhguara dhe vlerave koresponduese të përputhshmërisë, i referohet $n/10$ vëzhgimeve, që kanë probabilitetin e më të lartë të vlerësuar. Çifti tjetër i referohet $n/10$ vëzhgimeve që sigurojnë, probabilitetet e dyta të vlerësuar, e kështu me rradhë. Çdo grup ka një numër subjektsh të vëzhguar dhe, një vlerë të përputhshmërisë. Vlera e përputhshmërisë për një rezultat është shuma e probabiliteteve të vlerësuar, për këtë rezultat për të gjitha vëzhgimet në atë grup.

Ata propozuan një test Pirson, duke krahasuar vlerat e vëzhguara dhe të përputhshmërisë për ndarjet e lartëpërmendura.

Testin e ndertuan kështu: Shënojmë me y_{ij} rezultatin binary për vëzhgimin j në grupin i të ndarjes, $i = 1, \dots, n_i$. Le të shënojmë $\hat{p}_{ij}$ probabilitetin e përputhshmërisë korespondues për modelin në fjalëme të dhënat e pagrupuara. Testi është

$$\sum_{i=1}^g \frac{(\sum_j y_{ij} - \sum_j \hat{p}_{ij})^2}{(\sum_j \hat{p}_{ij})[1 - (\sum_j \hat{p}_{ij})/n_i]}$$

Kur shumë vëzhgime kanë të njëjtin probabilitet të vlerësuar, ka një arbitraritet në formimin e grupeve, dhe software të ndryshëm mund të raportojnë deri diku vlera të ndryshme.

Testi Hosmer-Lemeshow nuk është shumë i fuqishëm (Hosmer, 1997)^[27] për të dedektuar raste të veçanta mungese përputhshmërie. Mënyra e krahasimit të modelit të propozuar, me vlerat reale është më e përdorshme shkencërisht. (*)

Nëse p-vlera e afishuar e testit është e madhe, kjo tregon se përshtatshmëria e të dhënave me modelin është e mirë.

5.3.1 Vërejtje

- Në modelin e regresit linear të shumëfishtë, shih 2.2.1" Për cfarë vlen vlerësimi i koeficientit të korelacionit", është treguar se për të kontrolluar nëse vlen apo jo modeli i propozuar, kontrollohet hipoteza $H_0: \rho_{y, x_1 \dots x_n} = 0$. Ndaj kryhet testi ANOVA, e kuptohet se kur hipoteza bie poshte, pra për p-vlera shumë të vogla, konkludohet se *modeli është i vlefshëm*.
- Në rastin e regresit logjistik ndodh ndryshe. Aty, me anë të testit me shpërndarje X^2 , kontrollohet hipoteza 'modeli logjistik është i vlefshëm'. Ndaj

p-vlera te vogla, tregojnë se testi bie poshte pra, nuk kemi model logjistik të vlefshëm.

5.3.2 Zbatim 1 (vazhdim)

Për modelin e mësipërm është bërë kontrolli i përshtatshmërisë së modelit. Kërkohet të kontrollohet hipoteza:

'ky model logjistik është i vlefshëm',

me hipoteze alternative:

'modeli në fjalë nuk ia vlen të shqyrtohet'.

Afishimet e Tabelës 5-3 japin përgjigjen e hipotezës . Në tabelen e afishuar lexohen vlera e testit X^2 , dhe p-vlera e tij.

Tabela 5-3 Kontrolli i vlefshmërisë së modelit

Hosmer and Lemeshow Test

Step	Chi-square	df	Sig.
1	6.205	4	.184

P-vlera 0.184, jo signifikative, tregon se modeli është i vlefshëm.

Pas modelimit afishohet dhe Tabela 5-4 që përmban vlerat e Cox & Snell R Square dhe Nagelkerke R Square. Këto vlera të R^2 s'janë të njëjta, se ato janë llogaritur në shkallë të ndryshme. Vlera e Cox & Snell R Square, sipas statisticienëve luhatet deri në 0.7, ndërsa Nagelkerke R Square I merr vlerat si zakonisht deri në 1, ndaj zakonisht në interpretim i referohemi këtij vlerësimi. Ky numer jep njëfarë informacioni për perafrimin e modelit.

Tabela 5-4 Informacione për modelimin

Model Summary

Step	-2 Log likelihood	Cox & Snell R Square	Nagelkerke R Square
1	266.107 ^a	.157	.209

a. Estimation terminated at iteration number 5 because parameter estimates changed by less than .001.

5.4 Modelimi i vlerave të parametrave të parashikuar

Qëllimi i modelit logjistik është të arrijme, që duke njohur vlerat e variablave që marrin pjesë në modelim të mund të parashikojmë vlerën e variablës të varur. Në programet statistikore, krahas përfundimeve të mësipërme, në tabelat e data-ve afishohen edhe dy variabla të rinj, vlera e parashikuar e probabilitetit dhe 'grup ku bën pjesë', që në fakt është vlera e variablës që ne synojmë të parashikojmë përmes modelimit. Duke lexuar këtë vlerë mund të konkludojmë, ku bën pjesë individi.

Le të shpjegojmë mënyrën se si arrihet të behet kjo. Procesi kalon në disa etapa:

- Duke shfrytëzuar koeficientët e pastandardizuar të tabelës dhe konstanten e fituar nga analiza e regresit logjistik binar, arrihet të modelohet variabli *predicted_logit_probability*.
- Së dyti, radhiten vlerat e variablit në fjalë.
Krijohet një variabël i ri index dhe me ndihmën e tij dhe të vlerave të radhitura të *predicted_logit_probability*, nëse bëhet paraqitja grafike, qoftë dhe e tipit scatter, mund të shihet qartë se ky variabël i ri i modeluar me ndihmën e koeficientëve të tabelës ka formë S, pra është një vijë sigmoidale^[28].
- Janë pikërisht vlerat e këtij variabli *predicted_logit_probability* që transformohen në vlerat e variablit *predic_probability*.
- Etapa e katër mbyll procesin e parashikimit, bën të mundur që me anë të një transformimi, vlerat e *predicted_logit_probability*, të kthehen në vlerat e duhura që përcaktojnë se ku bën pjesë individ. Zbatimi i mëposhtëm paraqet saktë këto procese.

5.4.1 Zbatim 1 (vazhdim)

Në zbatimin e paraqitur mësipër, do të tregohet si mund të arrihet, punë që sot bëhet nga softet statistikore, të gjenerohen vlerat e parashikuara.

Së pari, me ndihmën e sintaxes, duke shfrytëzuar konstanten e koeficientët e pastandardizuar të fituar nga analiza e regresit logjistik ndërtoj funksionin

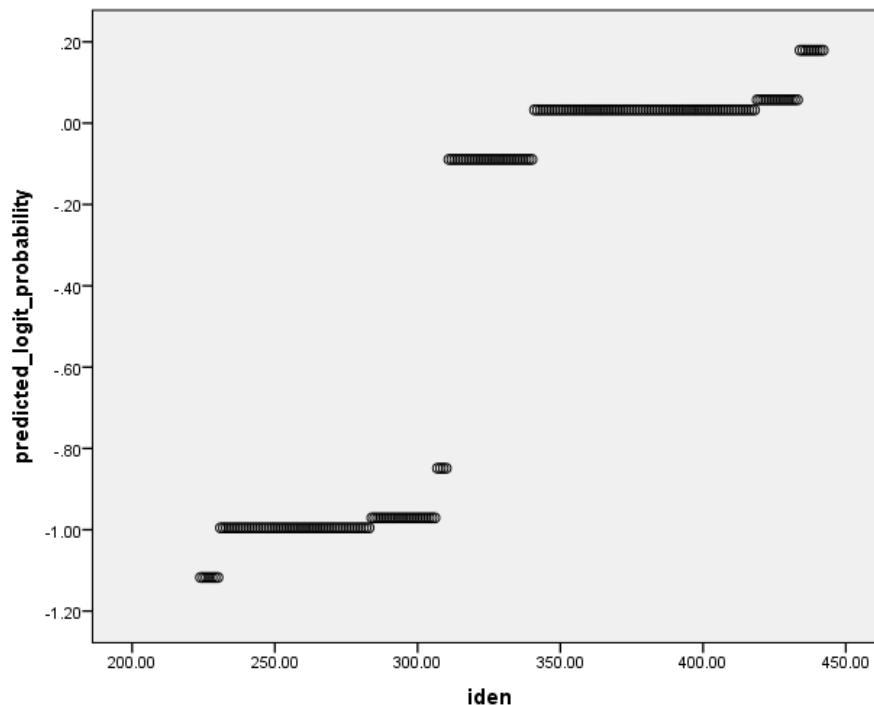
```
Compute predicted_logit_probability=1.02810746*area+.14654518*pini_duhan  
+.12196304*duhan_shtepi -2.26733415
```

Duhet patur parasysh se këto nuk janë vlerat e parashikuara të probabiliteteve, këto janë vlerat e logit të tyre, pra

$$F(w) = \frac{e^w}{(1+e^w)}$$

Së dyti, radhis vlerat e variablit *predicted_logit_probability*, dhe me ndihmën e variablit index, të krijuar nga ne bëj paraqitjen e tyre me një grafik scatter që është paraqitur në Figura 5-2

Figura 5-2 Grafiku i predicted_logit_probability



Sa më shumë grafiku ka formën e S aq më i mirë del parashikimi. Në fakt në rastin tonë të modelimit, raporti i numrit të personave me kollë (të koduar 1), me ato pakollë, (të koduar 0) ishte afërsisht i njëjtë. Nëse do të ishte shumë i ndryshëm atëhere parashikimi që në fillim pa bërë asnjë modelim do të ishte shumë më i lartë, dhe kur të përmirësohet nga modeli kuptohet do të sigurohej, së fundmi, vlera shumë më të afërta me ato të matura.

Së treti, duke shfrytëzuar formulën, me ndihmën e sintaksës arrihet të krijohet variabli *predictit_probabilitet*.

```
compute predictit_probabilitet=2.7182818281828** predicted_logit_probability
/(1+2.7182818281828)** predicted_logit_probability
```

Së fundmi, behet transformimi i *predictit_probabilitet* në variabël binar. Transformimi më i llogjikshëm është, që në rastet që ky variabël merr vlera më të vogla 0.5, atëhere vlera parashikohet 0, ndërsa në rastet që ajo është më e madhe se 0.5, vlera e parashikuar e variablit të varur të jetë 1.

Shpesh herë propozohet se si vlerë kufiri përcaktuese, ndarëse, të merret 0.45. Në tabelën 5-5 paraqitet një pjesë e vlerave origjinale të variablit '*kolle*'. Në kollonën e dytë, janë paraqitur një pjesë e vlerave të variablave të sqaruar mësipër, *predicted_logit_probability*, *predictit_probabilitet* dhe *predicted value*.

Tabela 5-5 Të dhënat e modelimit të regresit logjistik

Vl.k.tkolle	Index	Pred_logit_pr	Pred_prob	Pred_value
.00	6	.03	.39653	.00
.00	7	.03	.39653	.00
1.00	8	.03	.39653	.00
1.00	9	.03	.39653	.00
.00	10	.03	.39653	.00
1.00	11	.03	.39653	.00
1.00	12	.03	.39653	.00
1.00	13	.03	.39653	.00
.00	14	.03	.39653	.00
.00	15	.03	.39653	.00
1.00	16	.06	.85568	1.00
1.00	17	.06	.85568	1.00
1.00	18	.06	.85568	1.00
1.00	19	.06	.85568	1.00
.00	20	.06	.85568	1.00
1.00	21	.06	.85568	1.00
1.00	22	.06	.85568	1.00
1.00	23	.06	.85568	1.00
1.00	24	.06	.85568	1.00
1.00	25	.06	.85568	1.00
1.00	26	.06	.85568	1.00
1.00	27	.06	.85568	1.00
1.00	28	.06	.85568	1.00
1.00	29	.06	.85568	1.00
1.00	30	.06	.85568	1.00
.00	31	-1.00	.24721	.00
1.00	32	-1.00	.24721	.00
.00	33	-1.00	.24721	.00
.00	34	-1.00	.24721	.00
.00	35	-1.00	.24721	.00
.00	36	-1.00	.24721	.00
.00	37	-1.00	.24721	.00
.00	38	-1.00	.24721	.00
1.00	39	-1.00	.24721	.00
1.00	40	-1.00	.24721	.00
1.00	41	-1.00	.24721	.00
1.00	42	-1.00	.24721	.00
1.00	43	-1.00	.24721	.00
.00	44	-1.00	.24721	.00
.00	45	-1.00	.24721	.00
.00	46	1.00	.24721	.00
.00	47	1.00	.24721	.00
.00	48	1.00	.24721	.00
.00	49	1.00	.24721	.00

.00	50	1.00	.24721	.00
.00	51	1.00	.24721	.00
1.00	52	1.00	.24721	.00
.00	53	1.00	.24721	.00
.00	54	1.00	.24721	.00

Po te shihet me kujdes ka nje numër të madh vlerash 1 që parashikohen në rregull, ndërsa pa modelin në fjalë ato nuk parashikohen dot. asnjëra prej tyre. Megjithatë duhet thënë se më parë parashikoheshin të gjithë rastet pa kollë.

5.5 Funksionit lidhës 'Probit'

Modeli i ndërtuar ka formën $p(x)=F(x)$. Duke shfrytëzuar klasën e gjërë të shpërndarjes normale me shumëllojshmëritë e pritjes matematike dhe dispersioneve, mund të gjendet një i tillë që shkallën e rritjes ta ketë të njëjtë me atë të shpërndarjes normale të normuar. Por, me ndihmen e funksionit shpërndarës të shpërndarjes normale θ , lidhja e mësipërme mund të shkruhet në formën:

$$p(x) = \theta(\beta_0 + \beta x)$$

Meqë θ është funksion monoton rritës, funksioni i anasjelltë θ^{-1} ekziston^[29], ndaj barazimi (6) merr formën:

$$\theta^{-1}(p(x)) = \beta_0 + \beta x \quad (7)$$

Në këtë rast funksioni lidhës është θ^{-1} i cili pasqyron zonën]0;1[të probabiliteteve në $(-\infty; +\infty)$, zona ku marrin vlera variablat pjesëmarrës të modelit të regresit. Modeli (7) quhet model *probit*.

5.5.1 Vërejtje:

- Softet e posacme kompjuterike ofrojnë mundësinë e përdoruesit në zgjedhjen e funksionit lidhes. Pra, mund të zgjidhet funksioni '*logit*' i trajtuar më parë, ose funksioni në fjalë, '*probit*'.
- Softet ofrojnë dhe mundësinë e funksioneve të tjera lidhës.
- Funksionet e tjera nga '*logit*', zakonisht, përdoren nëse kushti i koeficientëve të njëjtë, për të gjitha kategoritë e variablit të varur, nuk plotësohet.

6 Ndërtimi e analiza e një modeli probabilitaro-statistikor për studimin e efektit të ndotjes

Oksidet e sulfurit(SO_x) janë molekula të squfurit dhe oksigjenit. Dioksidi i sulfurit është forma më e zakonshme që gjendet në shtresat e ulta të atmosferës. Ai, në sasi të 1,000 deri në 3,000 mikron për metër kub($\mu\text{g}/\text{m}^3$), është pa arome dhe shije. Ndërsa në përqëndrimet 10,000 $\mu\text{g}/\text{m}^3$, ka një aromë të fortë jo të këndshme⁽¹⁾.

Ky gaz ndodhet në sasira më të medha afer zonave industriale, e vecanërisht atyre naftëmbajtëse. Në vendin tonë, fusha e Kucovës dhe Marinëz janë zonat më njohura naftëmbajtëse. Thuajse çdo familje ka brenda truallit të vet, një pus naftë ende në shfrytëzim madje në të shumicën e rasteve edhe dy puse. Banorët ankohen, se këto puse rrjedhin pa pushim dhe ç'ka është më kryesorja ata janë me avari. Amortizimi i grupeve të grumbullimit të naftës, si rezultat i korrozionit ka sjellë rredhjen në bazën e rezervuarve, në kolektor, saraqineska, e këto rrjedhje pa dyshim shkaktojnë ndotje të mjedisit në tërësi

Burimi i këtij gazi është ndotja jo vetëm nga zjerrja por dhe, nga përpunimi që i bëhet nënprodukteve të tij. Mbetjet e naftës derdhen në rezervuar dhe ndotin ujin e vaditjes duke e nxjerrë jashtë përdorimit, ujë që banorët e këtij fshati dhe të fshatrave fqinjë e përdorin për vaditjen e bimëve bujqësore gjatë verës. Gjithashtu dhe shkalla e lartë e amortizimit të tubacioneve pus- grup dhe grup-impian, dekantimi si rezultat i kohës së gjatë të shfrytëzimit, korrozioni i shkaktuar nga nafta me sasi squfuri, sjellin përsëri rrjedhje të sasirave të konsiderueshme, përzierje naftë-ujë në mjedis

Nisur nga problemet në fjalë, jeta e banorëve rrezikohet seriozisht. Për këtë nga Instituti i Shëndetit Publik dhe Ministria e Mjedisit merr të dhëna për treguesit kryesor të gjendjes së njerëzve dhe këto të dhëna analizohen, me qëllim marrjen e masave të nevojshme nga organet kompetente për zbatimin e kushteve teknike, në shfrytëzimin e puseve të naftës, që ndodhen pranë shtëpive, si psh, përdorimi i filtrave pastrues dhe monitorizimi i ndotjes. Në këtë punim janë shfrytëzuar të dhënat e marra nga këto institucione, për pesë fshatra të zonës Belinë-Sheqishtë, Zharrë, Marinëz, Patos e Kurjan, ne jug të Shqipërisë, në Fier, Figura 1. Fshatrat ndodhen pranë njëri tjetrit.

Në baze të të dhënave të Institutit të Gjeoshkencës u bë paraqitja e vendodhjes së puseve në zonë, Figura 2. Nga vëzhgimi i hollësishëm i pozicionimit të tyre kuptohet, pse Kurjan konsiderohet zonë e paster ndërsa fshatrat e tjerë jo. Ato janë te pasura me puse e janë zona të përpunimit të naftës.

Figura 6-1 Pozicioni i zonë së studimit

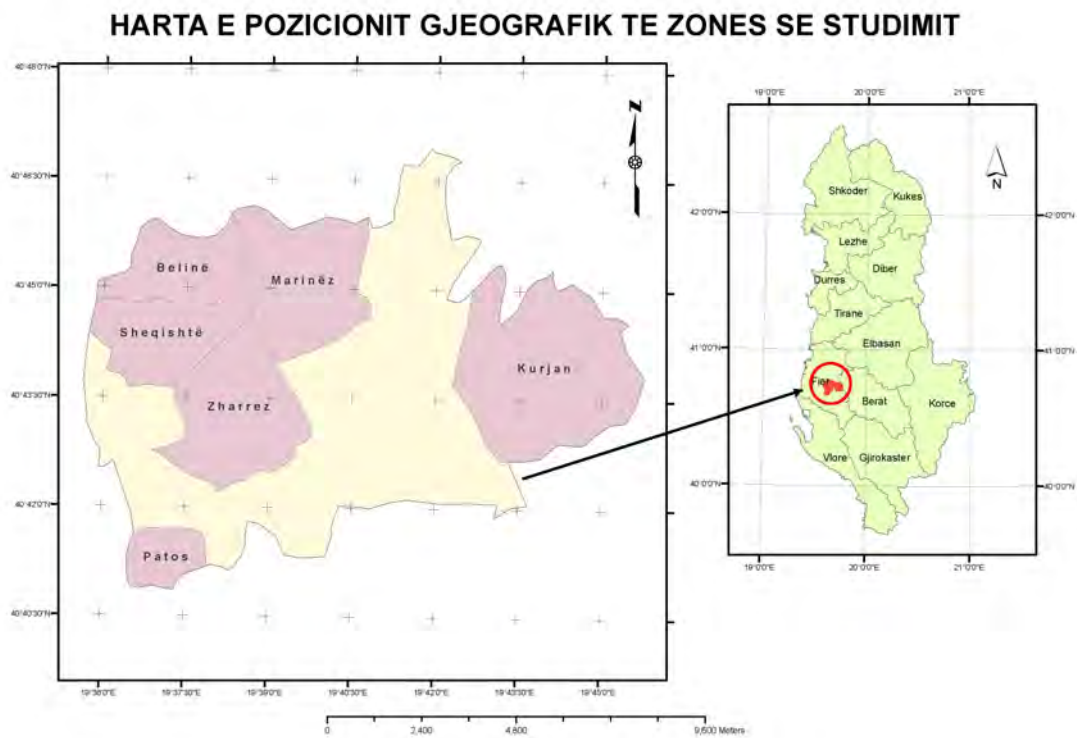
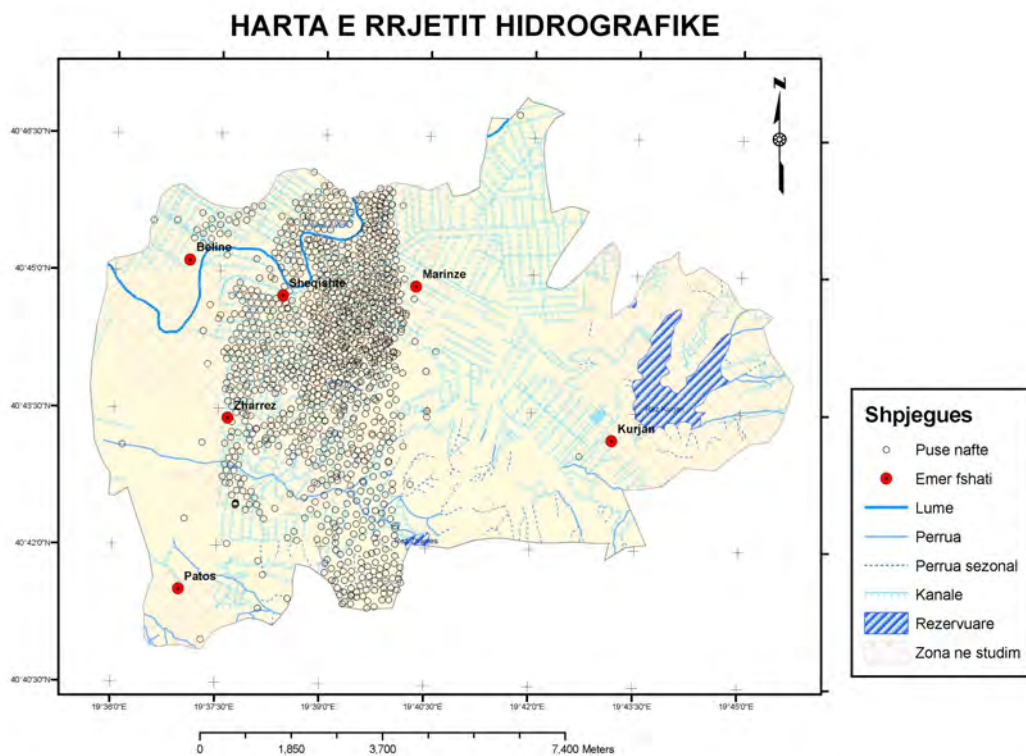


Figura 6-1 Pozicioni i puseve tnaftës



6.1 Ndikimi i nivelit të ndotjes në problemet shëndetësore të përgjithshme, dermatologjike dhe ato pneumologjike. (Problemi 1)

Jane marre të dhëna nga 171 individë, banues të tri zonave, dy të konsideruar industriale, Marinëz dhe Sheqishtë-Belinë, dhe një e konsideruar e pastër, Kurjan. Të intervistuarve i janë bërë dhe 13 pyetje të ndryshme. Pyetjet ndahen në tri kategori: pyetje të gjendjes së përgjithshme shëndetësore (3 pyetje), pyetje për problemet dermatologjike (5 pyetje) si dhe pyetje për problemet pneumologjike (5 pyetje).

Pyetje të gjendjes së përgjithshme janë: marrje mendje (p33), nauze apo të vjella (p34), dhimbje koke (p35).

Pyetje për problemet dermatologjike janë: lunga në lëkurë (p28), kruajtje në lëkurë (p29), shpërthim puçrash (p30), skuqje të lëkurë (p31), dhe pezmatim të lëkurës (p32).

Pyetjet për problemet pneumologjike janë: sy të përllotur(p23), kruarje hundësh(p24), tharje gryke(p25), gëlbbaze dhe kollë.

Nga kodimi i pyetjeve të të intervistuesve janë përcaktuar disa variabla, që siç shihet, duhet të konsiderohen si variabla të varur. Pra, variabla të varur janë 13.

Shtohet problemii ndikimit në këto variabla të fshatit, pra të faktit që kemi të bëjme me të dhëna të marra nga një fshat industrial apo një fshat joindustrial.

Në fillim do të trajtohen të tri llojet e pyetjeve si të dhëna të problemit 'shenja, efekte', më pas do të shihen si variabla të pavarur në analizën MANOVA.

Për përdorimin e metodës MANOVA kërkohet

Tabela 6-1 Matrica e kovariancës

Box's Test of Equality of Covariance Matrices ^a	
Box's M	.000
F	.000
df1	91
df2	35214.891
Sig.	1.000

Tests the null hypothesis that the observed covariance matrices of the dependent variables are equal across groups.

a. Design: Intercept + fshatrat

plotësimin e disa kushte. Një kusht, tej mase i rëndësishëm, që duhet të plotësohet, është kushti I varincës së barabarte mes tri grupeve të të dhënave të krijuara në rastin tonë nga ndarja e të 13 të dhënave të ndryshme, sipas fshatrave ku individi është banor.

Në rastin tonë ku kusht plotësohet. Kjo mund të shihet nga rezultati i vlerës së statistikës Fisher dhe p- vlerës koresponduese. Me gjithatë duhet theksuar se hipoteza në fjalë është e vertetë dhe do të pranohet si e tillë edhe për p-vlera >0.005. Konkretisht, në rastin tonë, pasi të dhenat ndahen në bazë të fshatit, banor i të cilit është individi, varesia e 13 variablave shprehet nga

$$\text{cov}(\mathbf{Y}_i, \mathbf{Y}_j) = \sigma_{ij} \mathbf{I} \quad i, j = 1, 2, \dots, p$$

ku σ_{ij} është kovarianca mes variablit të i-të dhe të j-të, e llogaritur janë paraqitur në Tabela 6-2.

Tabela 6-2 Kovarianca mes variablave

Inter-Item Covariance Matrix														
fshatrat	sy t	kruarje	Tharj	gelbaz	koll	Lung	kruarj	Shperthi	skuqje	pezmatim	marramendj	nauze	dhimbj	
	perlotu	hundesh	e gryke	e	e	ne lekure	ne lekure	Pucrrash	i ne fytire	e te lekures	e	te vjella	e koke	te forta
Vlosh sy te	.207	.005	.058	.037	.000	-.005	-.005	-.005	-.005	-.005	.078	.037	.111	
perlotur														
kruarje	.005	.088	.014	.012	-	.015	.015	-.002	-.002	-.002	.037	.012	.021	
hundesh					.005									
tharje gryke	.058	.014	.157	.025	.038	-.003	-.003	-.003	.013	-.003	.058	.073	.090	
gelbaze	.037	.012	.025	.166	.082	-.003	-.003	-.003	-.003	-.003	.021	.054	.021	
kolle	.000	-.005	.038	.082	.176	-.004	-.004	.013	-.004	.013	.000	.050	.032	
lunga ne	-.005	.015	-.003	-.003	-	.016	.016	.000	.000	.000	.012	-.003	-.005	
lekure					.004									
kruarje ne	-.005	.015	-.003	-.003	-	.016	.016	.000	.000	.000	.012	-.003	-.005	
lekure					.004									
shperthim														
pucrrash	-.005	-.002	-.003	-.003	.013	.000	.000	.016	.000	.016	.012	-.003	.012	
apo														
fluskash														
skuqje														
nxehtesi ne	-.005	-.002	.013	-.003	-	.000	.000	.000	.016	.000	.012	-.003	-.005	
fytire					.004									
pezmatime	-.005	-.002	-.003	-.003	.013	.000	.000	.016	.000	.016	.012	-.003	.012	
te lekures														
marramendj	.078	.037	.058	.021	.000	.012	.012	.012	.012	.012	.207	.037	.078	
e														
nauze, te	.037	.012	.073	.054	.050	-.003	-.003	-.003	-.003	-.003	.037	.166	.085	
vjella														
dhimbje	.111	.021	.090	.021	.032	-.005	-.005	.012	-.005	.012	.078	.085	.207	
koke te forta														
Sheqishte sy te	.238	.062	.020	.059	-	.044	.079	.058	.042	.043	.020	.085	.008	
-Beline perlotur					.007									

Ndërtimi&analiza e një modeli probabilitar-statistikor për studimin e efektit të ndotjes

	kruarje	.062	.220	.095	.029	-	.022	.049	.048	.012	.107	-.037	.042	.042
	hundesh				.022									
	tharje gryke	.020	.095	.212	.055	.051	.050	.022	.043	.034	.076	.024	.098	.074
	gelbaze	.059	.029	.055	.252	.067	.019	.004	-.002	.050	.029	-.040	.034	-.021
	kolle	-.007	-.022	.051	.067	.254	-.025	.005	-.024	.096	-.003	-.006	.001	.007
	lunga ne					-								
	lekure	.044	.022	.050	.019	.025	.142	.088	.116	-.047	.060	-.044	.016	.022
	kruarje ne													
	lekure	.079	.049	.022	.004	.005	.088	.196	.116	.008	.087	-.016	.018	.034
	shperthim					-								
	pucrrash apo	.058	.048	.043	-.002	.024	.116	.116	.165	-.032	.086	-.033	.032	.027
	fluskash													
	skuqje													
	nxehtesi ne	.042	.012	.034	.050	.096	-.047	.008	-.032	.242	.030	.015	.039	.024
	fytyre													
	pezmatime					-								
	te lekures	.043	.107	.076	.029	.003	.060	.087	.086	.030	.220	-.075	.023	.042
	marramendj					-								
	e	.020	-.037	.024	-.040	.006	-.044	-.016	-.033	.015	-.075	.212	.022	.055
	nauze, te													
	vjella	.085	.042	.098	.034	.001	.016	.018	.032	.039	.023	.022	.249	.019
	dhimbje													
	koke te forta	.008	.042	.074	-.021	.007	.022	.034	.027	.024	.042	.055	.019	.115
Marinez	sy te	.238	.062	.020	.059	-	.044	.079	.058	.042	.043	.020	.085	.008
	perlotur					.007								
	kruarje					-								
	hundesh	.062	.220	.095	.029	.022	.022	.049	.048	.012	.107	-.037	.042	.042
	tharje gryke	.020	.095	.212	.055	.051	.050	.022	.043	.034	.076	.024	.098	.074
	gelbaze	.059	.029	.055	.252	.067	.019	.004	-.002	.050	.029	-.040	.034	-.021
	kolle	-.007	-.022	.051	.067	.254	-.025	.005	-.024	.096	-.003	-.006	.001	.007
	lunga ne					-								
	lekure	.044	.022	.050	.019	.025	.142	.088	.116	-.047	.060	-.044	.016	.022
	kruarje ne													
	lekure	.079	.049	.022	.004	.005	.088	.196	.116	.008	.087	-.016	.018	.034
	shperthim					-								
	pucrrash apo	.058	.048	.043	-.002	.024	.116	.116	.165	-.032	.086	-.033	.032	.027
	fluskash													

skuqje													
nxehtesi ne	.042	.012	.034	.050	.096	-.047	.008	-.032	.242	.030	.015	.039	.024
fytyre													
pezmatime	.043	.107	.076	.029	-	.060	.087	.086	.030	.220	-.075	.023	.042
te lekures					.003								
marramendj	.020	-.037	.024	-.040	-	-.044	-.016	-.033	.015	-.075	.212	.022	.055
e					.006								
nauze, te	.085	.042	.098	.034	.001	.016	.018	.032	.039	.023	.022	.249	.019
vjella													
dhimbje	.008	.042	.074	-.021	.007	.022	.034	.027	.024	.042	.055	.019	.115
koke te forta													

Afishimet e marra nga metoda e Multivariancës janë paraqitur në Tabelën 6-3.

Vlera e statistikës Fisher, në rastet ku nuk i është kushtuar kujdes normalizimit të të dhënave, lexohet nga afishimi i Pillai's Trace, është 4.654 dhe i korrespondon p- vlera 1.2 E-12. Kjo tregon se ka efekt variabli i pavarur , ne rastin tonë '*fshatrat*', në vlerat e gjithë variablave të varura, kur ato të 13 shihen si një grup i vetëm. Madje, p-vlera tej mase e vogël, percakton dhe shkallën e lartë të ndryshimit të 'efekteve' sipas fshatrave. Duhet të theksohet se në rastin tone efekti është signifkativ sipas të gjitha modeleve të ofruara, madje edhe sipas testit Roy^[30]. Ky test përdoret pasi të jenë plotësuar shumë kushte. Psh, kërkon jo thjeshtë kontrollin e kufijve të disa parametrave që na japin informacion për normalizimin, por patjetër plotësim të plotë të kushteve të normalizimit.

Tabela 6-3 Testi i Multivariancës

Multivariate Tests ^a							
Effect		Value	F	Hypothesis df	Error df	Sig.	Partial Eta Squared
Intercept	Pillai's Trace	.992	1447.439 ^b	13.000	156.000	.000	.992
	Wilks' Lambda	.008	1447.439 ^b	13.000	156.000	.000	.992
	Hotelling's Trace	120.620	1447.439 ^b	13.000	156.000	.000	.992
	Roy's Largest Root	120.620	1447.439 ^b	13.000	156.000	.000	.992
fshatrat	Pillai's Trace	.556	4.654	26.000	314.000	.000	.278
	Wilks' Lambda	.444	6.015 ^b	26.000	312.000	.000	.334
	Hotelling's Trace	1.254	7.474	26.000	310.000	.000	.385
	Roy's Largest Root	1.254	15.141 ^c	13.000	157.000	.000	.556

a. Design: Intercept + fshatrat

b. Exact statistic

c. The statistic is an upper bound on F that yields a lower bound on the significance level.

Madje në kollonën e fundit lexohet, se sipas testit me të dobët Pillai, 30% e ndryshimeve në gjendjen shëndetësore të njerëzve shpjegohen si rezultat I tri nivele të ndryshme të variablilit 'fshatrat'. Ndërsa sipas testit të fuqishëm Roy, kushtet e të cilit në rastin tonë plotësohen, jemi në kushtet e variancës së barabartë mes grupeve, më shumë se 50% e ndryshimeve në gjendjen e përgjithshme varen nga zona e banimit.

Konkluzioni në të cilin është arritur, duhet të interpretohet drejt, pra ka një ndikim të variablilit 'fshatrat' tek ndonjë nënbashkësi e variablave të trajtuar deri tani si një grup. Procesi i MANOVA-s vazhdon me etapën e dytë, atë të kryerjes së disa testeve 'univariate' ndaj secilit prej 13 pyetjeve, pra kryhem ANOVA për çdo pyetje. Afishimet tregojnë për ndikim të variablilit *fshatrat* tek secili nga 13 pyetjet e studiuar. (Tabela 6-4)

Shihet se ndikimi më i vogël I '*fshatrat*' është tek *kruarja e hundëve* dhe *kolle*. Në këtë përfundim arrihet duke u mbështetur tek vlerat e F-testit dhe afishimet koresponduese të p-vlerave.

Duhet theksuar se nuk vërehet ndikim shumë i ndjeshëm i *fshatrat* nga ku është marrë zgjedhja në '*lunga në lëkurë*'.

Tabela 6-4 ANOVA për 13 variablat e varur

Source	Dependent Variable	Type III	df	Mean Square	F	Sig.	Partial Eta Squared
		Sum of Squares					
fshatrat	sy te perlotur	4.706	2	2.353	10.392	.000	.110
	kruarje hundesh	1.918	2	.959	5.610	.004	.063
	tharje gryke	10.481	2	5.240	27.313	.000	.245
	gelbaze	4.852	2	2.426	11.020	.000	.116
	kolle	2.674	2	1.337	5.935	.003	.066
	lunga ne lekure	.905	2	.452	4.755	.010	.054
	kruarje ne lekure	2.357	2	1.179	9.113	.000	.098
	shperthim pucrrash apo filuskash	1.404	2	.702	6.373	.002	.071
	skuqje nxehtesi ne fytire	14.098	2	7.049	44.434	.000	.346
	pezmatime te lekures	3.556	2	1.778	12.302	.000	.128
	marramendje	6.952	2	3.476	16.507	.000	.164
	nauze, te vjella	5.380	2	2.690	12.306	.000	.128
	dhimbje koke te forta	13.601	2	6.800	45.622	.000	.352

Tabela 6–5 t testet për secilin prej 13 variablave të varur

Dependent Variable	(I) fshatrat	(J) fshatrat	Mean Difference (I-J)	Std. Error	Sig.
sy te perlotur	Kurjan	Sheqishte-Beline	.3439*	.08825	.000
		Marinez	.3439*	.08825	.000
	Sheqishte-Beline	Kurjan	-.3439*	.08825	.000
		Marinez	.0000	.09158	1.000
	Marinez	Kurjan	-.3439*	.08825	.000
		Sheqishte-Beline	.0000	.09158	1.000
kruarje hundesh	Kurjan	Sheqishte-Beline	.2196*	.07668	.005
		Marinez	.2196*	.07668	.005
	Sheqishte-Beline	Kurjan	-.2196*	.07668	.005
		Marinez	.0000	.07958	1.000
	Marinez	Kurjan	-.2196*	.07668	.005
		Sheqishte-Beline	.0000	.07958	1.000
tharje gryke	Kurjan	Sheqishte-Beline	.5132*	.08123	.000
		Marinez	.5132*	.08123	.000
	Sheqishte-Beline	Kurjan	-.5132*	.08123	.000
		Marinez	.0000	.08430	1.000
	Marinez	Kurjan	-.5132*	.08123	.000
		Sheqishte-Beline	.0000	.08430	1.000
gelbaze	Kurjan	Sheqishte-Beline	.3492*	.08701	.000
		Marinez	.3492*	.08701	.000
	Sheqishte-Beline	Kurjan	-.3492*	.08701	.000
		Marinez	.0000	.09030	1.000
	Marinez	Kurjan	-.3492*	.08701	.000
		Sheqishte-Beline	.0000	.09030	1.000
kolle	Kurjan	Sheqishte-Beline	.2593*	.08803	.004
		Marinez	.2593*	.08803	.004
	Sheqishte-Beline	Kurjan	-.2593*	.08803	.004
		Marinez	.0000	.09135	1.000
	Marinez	Kurjan	-.2593*	.08803	.004
		Sheqishte-Beline	.0000	.09135	1.000
lunga ne lekure	Kurjan	Sheqishte-Beline	.1508*	.05720	.009
		Marinez	.1508*	.05720	.009
	Sheqishte-Beline	Kurjan	-.1508*	.05720	.009
		Marinez	.0000	.05936	1.000

Ndërtimi&analiza e një modeli probabilitar-statistikor për studimin e efektit të ndotjes

	Marinez	Kurjan	-.1508*	.05720	.009
		Sheqishte-Beline	.0000	.05936	1.000
	Kurjan	Sheqishte-Beline	.2434*	.06669	.000
		Marinez	.2434*	.06669	.000
kruarje ne lekure	Sheqishte-Beline	Kurjan	-.2434*	.06669	.000
		Marinez	.0000	.06921	1.000
	Marinez	Kurjan	-.2434*	.06669	.000
		Sheqishte-Beline	.0000	.06921	1.000
	Kurjan	Sheqishte-Beline	.1878*	.06154	.003
		Marinez	.1878*	.06154	.003
shperthim puçrrash apo fluskash	Sheqishte-Beline	Kurjan	-.1878*	.06154	.003
		Marinez	.0000	.06387	1.000
	Marinez	Kurjan	-.1878*	.06154	.003
		Sheqishte-Beline	.0000	.06387	1.000
	Kurjan	Sheqishte-Beline	.5952*	.07386	.000
		Marinez	.5952*	.07386	.000
skuqe nxehtesi ne fytire	Sheqishte-Beline	Kurjan	-.5952*	.07386	.000
		Marinez	.0000	.07665	1.000
	Marinez	Kurjan	-.5952*	.07386	.000
		Sheqishte-Beline	.0000	.07665	1.000
	Kurjan	Sheqishte-Beline	.2989*	.07050	.000
		Marinez	.2989*	.07050	.000
pezmatime te lekures	Sheqishte-Beline	Kurjan	-.2989*	.07050	.000
		Marinez	.0000	.07316	1.000
	Marinez	Kurjan	-.2989*	.07050	.000
		Sheqishte-Beline	.0000	.07316	1.000
	Kurjan	Sheqishte-Beline	.4180*	.08510	.000
		Marinez	.4180*	.08510	.000
marramendje	Sheqishte-Beline	Kurjan	-.4180*	.08510	.000
		Marinez	.0000	.08831	1.000
	Marinez	Kurjan	-.4180*	.08510	.000
		Sheqishte-Beline	.0000	.08831	1.000
	Kurjan	Sheqishte-Beline	.3677*	.08671	.000
		Marinez	.3677*	.08671	.000
nauze, te vjella	Sheqishte-Beline	Kurjan	-.3677*	.08671	.000
		Marinez	.0000	.08998	1.000
	Marinez	Kurjan	-.3677*	.08671	.000
		Sheqishte-Beline	.0000	.08998	1.000
	Kurjan	Sheqishte-Beline	.5847*	.07160	.000

		Marinez	.5847*	.07160	.000
dhimbje koke te forta	Sheqishte-Beline	Kurjan	-.5847*	.07160	.000
		Marinez	.0000	.07430	1.000
	Marinez	Kurjan	-.5847*	.07160	.000
		Sheqishte-Beline	.0000	.07430	1.000
	Kurjan	Sheqishte-Beline	.5847*	.07160	.000
		Marinez	.5847*	.07160	.000

Duke u përqendruar në vlerat e t-testeve të afishuara mund të arrihet në përfundimin :

- Ka ndryshim statistikativ midis fshatrave Kurjan dhe Sheqishtë- Beline dhe Kurjan e Marinez për çdo pyetje të bërë. Mesatarja e variablave, të përcaktuar nga përgjigjet e çdo pyetje të bërë, është statistikisht më e lartë në Kurjan. Duke vërejtur mënyrën e kodimit të përgjigjeve në fjalë, konkludohet se ka ndikim negativ të zonës industriale në tërë parametrat e përcaktuara si përcaktues të gjendjes së përgjithshme, të gjendjes pneumologjike dhe asaj dermatologjike.
- Nuk vihet re ndryshim midis dy fshatrave që ndodhen në zonat industriale.

Por metoda e paraqitur deri tani nuk është gjë tjetër veçse metoda me e zakonshme e Manova-s.

Vlera me e madhe e metodës MANOVA verëhet në mënyrën tjetër që mund të kryhet kjo metode, në **analizën e komponenteve kryesore**. Përcaktimi i ketyre komponentëve kryesorë është një përcaktim i variablave të rinj që janë kombinim linear i 13 variablave të varur në studim, e shpesh për këtë arsye njihen me emrin përcaktimi i supervariablave. Këto variabla të rinj, supervariabla^[15], të pavarur synojnë të paraqesin të njëjtat matje apo informacione të marra nga 13 variablat në studim.

MANOVA për analizën e komponentëve kryesorë do të ekzekutohet pas shkrimit të kodit në editorin e sintaksës. Editori është emërtuar ' *pyetje*' dhe përmban kodin e mëposhtëm:

MANOVA p23 p24 p25 gelbaze kolle p28 p29 p30 p31 p32 p33 p34 p35 BY fshatrat(1,3)

/DISCRIM=STAN RAW CORR

/PRINT =SIGNIF(MULTIV,UNIV,EIGEN,DIMENR)/DESIGN.

Pjesa e parë e afishimeve, Tabela 6-4, tregon për ndikim të variablit '*fshatrat*' në tërë keto variabla të varur. Këtë informacion e morëm edhe nga ekzekutimi i metodës MANOVA sipas variantit të parë. Afishimet e ndikimit verëhet se janë të njëjta me ato të metodës së parë. Madje po të krahasohen me kujdes, vlerat e testit Pillais është e njëjtë , .55629, dhe po ashtu vlerat e testeve të tjera.

Tabela 6-6 Analiza e ndikimit të variablit të pavarur tek variablat e varur në shqyrtim

EFFECT .. fshatrat
Multivariate Tests of Significance (S = 2, M = 5 , N = 77)

Test Name	Value	Approx. F	Hypoth. DF	Error DF	Sig. of F
Pillais	.55629	4.65353		26.00	314.00
.000					
Hotellings	1.25375	7.47425		26.00	310.00
.000					
Wilks	.44371	6.01498		26.00	312.00
.000					
Roys	.55629				
Note.. F statistic for WILKS' Lambda is exact.					

Ekzekutimi i kësaj analize siguron gjetjen e dy kombinimeve lineare të 13 variablave në studim. Kombinimi i dy këtyre supervariablave është signifkativ. Dihet se keto dy supervariabla të krijuar ruajnë totalin e variances, dhe janë ortogonalë. Por, supervariabli i parë shpjegon gjithë variancën e tri variablave të pavarur, që fshihen në 13 variablat e varur.

Tabela 6-7 Vlerat e veta

Root No.	Eigenvalue	Pct.	Cum. Pct.	Canon Cor.
1	1.25375		100.00000	100.00000
.74585				
2	.00000		.00000	100.00000
.00000				

Por afishimet e Tabeles 6-7 tregojnë se $0.74585^2 = 55.6\%$ të variancës supervariablit shpjegohet nga ndryshimet e vlerave të variablit 'fshati'.

Në tabelen 6-8 afishohen vlerat përcaktuese të supervariablit në fjalë, koeficientët e ketij supervariabli.

Tabela 6-8 Koeficientet e variablit kanonik

Correlations between DEPENDENT and canonical variables
Canonical Variable

Variable	1
p23	-.31412
p24	-.23080
p25	-.50926
gelbaze	-.32348
kolle	-.23739
p28	-.21248
p29	-.29417
p30	-.24600
p31	-.64955
p32	-.34177
p33	-.39591

p34 -.34184
p35 -.65818

Pra, këtë variabël kryesor e përcaktojmë si kombinim linear të koeficientëve të afishuar dhe të 13 variablave në fjalë. E llogarisim këtë variabël dhe e emërtojmë '*superndikim*', duke shkruar kodin përkatës në editorin e sintaksës.

compute super_ndikim= -.20647* p23 +.00410*p24-.5260*p25+.54939* gelbaze
 -.10457* kolle +.90458*p28 +.00185*p29
 +.19729* p30+1.58608* p31+.28266* p32 +.64448*
 p33+.16127* p34+1.20061* p35

Tabela 6–9 Analiza deskriptive e *fshatrave* në '*superndikim*'

Descriptive Statistics

Dependent Variable: super_ndikim

fshatrat	Mean	Std. Deviation	N
Vlosh	8.8516	.70669	63
Sheqishte-Beline	6.7938	1.01302	54
Marinez	6.7938	1.01302	54
Total	7.5519	1.34653	171

Kryhet metoda e ANOVA-s në variablin e ri, për të përcaktuar ndikimin në të fshatrave. Afishimet e vlerave të statistikës descriptive të variablit të ri '*superndikim*', në të tri nënpopullimet e përcaktuara nga vlerat e variablit '*fshatrat*' janë paraqitur në Tabelen 6-9. Mesataret e nënpopullimeve të *superndikimit* janë të njëjta, më të vogla, për dy fshatrat që ndodhen në zonën industriale, dhe po kështu dhe devijimet standarte. Duke patur parasysh kodimin e 13 variablave (sa më I lartë kodi aq më e mirë gjendja e përgjithshme), themi se ka ndikim negativ ndotja, efekt negativ në gjendjen e përgjithshme të individit.

Afishimet e ANOVA-s për '*superndikim*'-in që shpjegon maksimumin e variancës mes individeve, Tabela 6-10, tregon ndikimin e *fshatra*-ve në variablin e ri. Madje tabela në fjalë tregon se 55 % e ndryshimeve të *superndikimi*-t vijnë si rezultat I ndikimit të *fshatra*-ve (vlera e statistikës F është 101.28 ndërsa p-vlera <.001)

Tabela 6–10 Efekti në '*superndikimi*' i variablit të pavarur

Tests of Between-Subjects Effects

Dependent Variable: super_ndikim

Source	Type III Sum of Squares	df	Mean Square	F	Sig.	Partial Eta Squared
Corrected Model	168.492 ^a	2	84.246	101.282	.000	.547
Intercept	9516.397	1	9516.397	11440.809	.000	.986
fshatrat	168.492	2	84.246	101.282	.000	.547

Error	139.741	168	.832			
Total	10060.576	171				
Corrected Total	308.234	170				

a. R Squared = .547 (Adjusted R Squared = .541)

Vihet re se:

- I njëjti problem i trajtuar me metoden e komponentëve kryesorë, tregon se ndotja ndikon në 55% të gjendjes së përgjithshme (të shprehur nga variabli 'super_ndikim'), ndërsa sipas metodes MANOVA(GLM), ndotja spjegonte 27.8-50% te variances së 13 komponentëve të gjendjes së përgjithshme.
- Zgjidhja e problemit me MANOVA (GLM) është me e ndërlikuar e ka një vëllim të madh analizash, ndërsa 'analiza e komponenteve kryesore' e redukton problemin në ndikimin në një variabel të vetëm, për të cilin mund të vlerësohen qartë dhe vlerat mesatare të tij në tri nivelet e variablit 'fshatrat'.

Problemi në fjalë mund të trajtohet edhe me anë të *Metodës se Analizës Faktoriale*. Siç është përmendur në vërejtjen 1, metoda MANOVA, do të ishte e pakuptimtë të përdorej nëse tërë faktoret, variablat e varur do të ishin të ndryshëm. Këto variabla u tha se duhet t'i perkisnin të paktën dy grupeve të ndryshme. Në rastin tonë ato I perkasin tri grupeve të ndryshme. Tri pyetje kanë të bëjnë me gjendjen e përgjithshme të individit, 5 me problemet e tij dermatologjike dhe 5 me problemet pneumologjike. Shkrimi i kodit të mëposhtëm përdoret për të kryer këtë metodë:

```

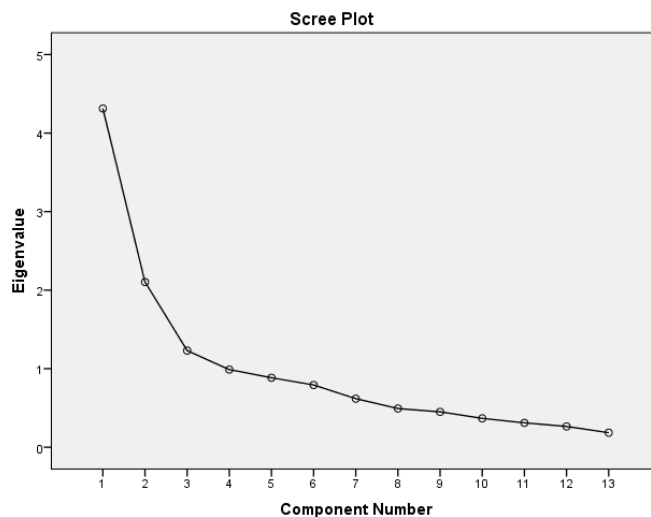
FACTOR
/VARIABLES p23 p24 p25 gelbaze kolle p28 p29 p30 p31 p32 p33 p34
p35
/MISSING LISTWISE
/ANALYSIS p23 p24 p25 gelbaze kolle p28 p29 p30 p31 p32 p33 p34 p35
/PRINT INITIAL SIG EXTRACTION ROTATION
/PLOT EIGEN
/CRITERIA FACTORS(3) ITERATE(25)
/EXTRACTION PC
/CRITERIA ITERATE(25)
/ROTATION VARIMAX
/METHOD=CORRELATION.

```

Afishimet e mëposhtme Figura 2, tregojnë se merren tri komponente kryesore. Kjo sepse, sic shihet edhe nga grafiku, vlerat e veta janë më të larta për 3-4 komponentët e parë. Madje kjo verëhet edhe në afishimet e tabelës së Analizës së Komponentëve Kryesorë.

Figura 6–3 Plotimi I vlerave të veta

a. Rotation converged in 4 iterations.



Pesë komponente kryesore shpjegojnë më shumë se 73% të variancës së përgjithshme. Në rastin tonë ne po përqëndrohemi në tri komponentët e parë kryesorë që shpjegojnë rreth 60% të kësaj variance të përgjithshme, sepse siç shihet vetëm tri vlerat e veta të para janë >1.

Në këtë rast do të krijoheshin tri variabla të rinj, variabla që janë 'faktore të 13 variablave'.

Tabela 6–11 Analiza e Komponentëve Kryesorë

Total Variance Explained									
Comp.	Initial Eigenvalues			Extraction Sums of Squared Loadings			Rotation Sums of Squared Loadings		
	Total	% of Variance	Cumulative %	Total	% of Variance	Cumulative %	Total	% of Variance	Cumulative %
1	4.313	33.174	33.174	4.313	33.174	33.174	3.091	23.776	23.776
2	2.103	16.179	49.353	2.103	16.179	49.353	2.614	20.111	43.887
3	1.231	9.466	58.820	1.231	9.466	58.820	1.941	14.933	58.820
4	.988	7.602	66.421						
5	.885	6.804	73.225						
6	.793	6.097	79.322						
7	.618	4.754	84.076						
8	.493	3.790	87.866						
9	.450	3.463	91.329						
10	.368	2.828	94.157						
11	.311	2.390	96.547						

12	.264	2.035	98.582					
13	.184	1.418	100.000					

Analiza faktoriale realizon krijimin e 3 variablove të rinj, e në përputhje me tri grupimet e 13 variablove të varur këto variabla po i emërtojmë *efektifaktorial1*, *efektifaktorial2*, *efektifaktorial3*.

Përqëndrimi në vlerat e afishuara të koeficienteve prane 13 variablove për tre komponentët e rinj *efektifaktorial*, Tabela 6-12, tregon se komponenti i parë ka koeficientë më të lartë pranë ndikimit në lekure, ndaj ndryshe *efektifaktorial1* mund të emërtohet *komponent_dermatologjik*.

Ky variabël justifikon 33% të ndryshimeve të variancës totale. Komponenti I dyte I ka koeficientët më të lartë pranë treguesve të gjendjes së përgjithshme (*komponent_pergjithshem*), dhe të dy komponentet e rinj të parë sëbashku justifikojnë rreth 50% të variancës së përgjithshme. Ndersa i treti është më i koreluar me pyetje që kanë të bëjnë me fushën pneumologjike, ndaj *efektifaktorial3* mund të emërtohet *komponent_pneumologjik*.

Ndikimi i 'fshat'-it vërehet në të tri komponentet, dhe ai është signifkativ në të tri komponentet, por ai përcakton 44% të ndryshimit të variablit përfaqësues të gjendjes së përgjithshme dhe 35% të variablit të gjendjes së përgjithshme pneumologjike, ku variabli 'kolle' zë vendin kryesor, me koeficientin më të lartë. Ndikimi I këtij variabli justifikon vetëm 22% të ndryshimeve në komponentin dermatologjik, që përcaktohet nga lunga, puçra dhe kruarje në lëkurë.

Nëse krahasojmë përfundimet me metoden MANOVA (Analiza e Komponenteve Kryesor) dhe Analizën Faktoriale vërehet se ndikimet e fshatit, ndotjes, në variablat e varur janë shumë larg 55% të fituar me metoden MANOVA. Por në këtë përfundim mund të na çojë vetëm kjo metodë, se kur krijon komponentët përbërës kjo metodë bazohet në idene e maksimizimit të variancës.

Tabela 6–12 Afishimet për të tri faktorët

Rotated Component Matrix ^a			
	Component		
	1	2	3
sy të perlotur	.324	.542	.122
kruarje hundesh	.460	.273	.178
tharje gryke	.312	.614	.405
gelbaze	.142	.098	.730
kolle	-.062	.063	.783
lunga në lekure	.834	.026	-.084
kruarje në lekure	.773	.165	.034
shperthim puçrash apo flluskash	.872	.104	-.068
skuqje nxehtesi në fytyrë	.020	.437	.600
pezmatime të lekures	.692	.108	.310

Marramendje	-.168	.830	-.104
nauze, te vjella	.161	.576	.275
dhimbje koke te forta	.256	.770	.167

Extraction Method: Principal Component Analysis.

Rotation Method: Varimax With Kaiser Normalization.

Tabela 6–13 Anova për komponentin pneumologjik

Tests of Between-Subjects Effects

Dependent Variable: Komp_pneumologjik

Source	Type III Sum of Squares	df	Mean Square	F	Sig.	Partial Eta Squared	Noncent. Parameter	Observed Power ^b
Corrected Model	70.889 ^a	2	35.444	45.764	0	0.353	91.528	1
Intercept	4715.665	1	4715.665	6088.652	0	0.973	6088.652	1
fshatrat	70.889	2	35.444	45.764	0	0.353	91.528	1
Error	130.116	168	0.775					
Total	5026.33	171						
Corrected Total	201.005	170						

a. R Squared = .353 (Adjusted R Squared = .345)

b. Computed using alpha = .05

Tabela 6–14 Anova për komponentin e përgjithshëm

Tests of Between-Subjects Effects

Dependent Variable: Komp_pergjith

Source	Type III Sum of Squares	df	Mean Square	F	Sig.	Partial Eta Squared	Noncent. Parameter	Observed Power ^b
Corrected Model	155.513 ^a	2	77.757	67.716	0	0.446	135.431	1
Intercept	8269.91	1	8269.91	7201.985	0	0.977	7201.985	1
fshatrat	155.513	2	77.757	67.716	0	0.446	135.431	1
Error	192.911	168	1.148					
Total	8828.331	171						
Corrected Total	348.425	170						

a. R Squared = .446 (Adjusted R Squared = .440)

b. Computed using alpha = .05

Tabela 6–15 Anova për komponentin e dermatologjik

Tests of Between-Subjects Effects

Dependent Variable: Komp_dermatologjik

Source	Type III Sum of Squares	df	Mean Square	F	Sig.	Partial Eta Squared	Noncent. Parameter	Observed Power ^b
Corrected Model	61.298 ^a	2	30.649	24.372	.000	.225	48.745	1.000
Intercept	11116.743	1	11116.743	8840.170	.000	.981	8840.170	1.000
fshatrat	61.298	2	30.649	24.372	.000	.225	48.745	1.000
Error	211.264	168	1.258					
Total	11568.856	171						
Corrected Total	272.562	170						

a. R Squared = .225 (Adjusted R Squared = .216)

b. Computed using alpha = .05

6.2 Ndërtimi I modelit statistikore për problemeve pneumologjike të banorëve të zonës industriale (Problemi 2)

Në pjesën e parë të studimit u përqëndruam te shqetësimet e banorëve të zonës së ndotur. Në etapën e fundit të këtij studimi, pasi përdorëm metodën e Analizës së Komponenteve, u arrit të ndërtohej një 'supervariabël' i problemeve pneumologjike, i cili u emërtua 'komponent pneumatik'. Në këtë pjesë të tretë duam të studiojmë problemin pneumatik, dhe të mund ti pergjigjemi pyetjeve: A ndikon ndotja në këtë probleme? A ndikojne në këto probleme faktorë apo variabla të tjerë ?

6.2.1 Modeli i analizës mbi komponentët e Analizës Faktoriale

Në trajtimin e problemit të ndikimit apo jo të zonës së pneumologji, variabël bazë u përdor variabli 'komponent pneumatik'. Kujtojmë se variabli në fjalë është llogaritur me anë të metodës Analiza Faktoriale sipas formulës

$$\text{compute Komp_pneumologjik} = .122 * p23 + .178 * p24 + .405 * p25 + .730 * \text{gelbaze} + .783 * \text{kolle} - .084 * p28 + .034 * p29 - .068 * p30 + .600 * p31 + .310 * p32 - .104 * p33 + .275 * p34 + .167 * p35$$

Kujtojmë, se ndonëse ai është kombinim i përgjigjeve të 13 pyetjeve, koeficientet me të lartë, kur ai modelohet i ka pranë 'kollës' e 'gëlbasës' e 'kruarjes në grykë', ndërsa koeficientët e tjerë janë më të vegjel se 0.4. Për zgjidhjen më të mirë të problemit, u llogarit 'komponenti pneumatik specifik', komponent ky vetëm I faktorëve pneumatik.

$$\text{compute Komp_pneumologjik_specifik} = .122 * p23 + .178 * p24 + .405 * p25 + .730 * \text{gelbaze} + .783 * \text{kolle}$$

Duke parë mënyrën se si ai është krijuar, themi se, ky variabël është i vazhdueshëm.

Por, kontrolli i problemit të ndikimit të zonës në këtë variabël, është i pakuptimtë, pa studimin e faktorëve të tjerë që mund të ndikojnë në të. Ndaj, nëse kërkojmë të analizojmë ndikimin apo jo në të disa variabla të tjerë, duhet përdorur metoda e regresit linear të shumëfishtë.

Nga analiza deskriptive e të dhënave u vu re 14.3% e të intervistuarve konsumojnë duhan, pjesa tjetër nuk konsumon.

Ndërtimi i modelit linear me variablat 'konsum duhani' dhe 'area' japin afishimet e mëposhtme, Tabela 16.

Rezultatet e afishuara tregojnë se modeli i regresit të ndërtuar, është i vlefshëm. Hipoteza se dy variablat e modelimit nuk kanë lidhje me 'komponent pneumatik specifik', bie poshtë. Vlera e statistikës Fisher 46.465 dhe p-vlera shumë e vogël, Tabela 6-16, tregojnë pikërisht këtë.

E modeli i propozuar, duke u bazuar në vërejtjes që njihet me emrin 'rule of thumb', Tabela 16 siguron një korelacion të lartë pozitiv me $R=0.599$. Ai justifikon rreth 36% të ndryshimeve të variablit në shqyrtim, 'komponent pneumatik specifik', Tabela 6-16.

Tabela 6–16 Afishimet e regresit linear

Model Summary

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Change Statistics				
					R Square Change	F Change	df1	df2	Sig. F Change
1	.598 ^a	.358	.350	.8751505531	.358	46.465	2	167	.000

a. Predictors: (Constant), a pini duhan, area

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.	Collinearity Statistics	
		B	Std. Error	Beta			Tolerance	VIF
1	(Constant)	6.749	.515		13.100	.000		
	area	-1.322	.139	-.588	-9.475	.000	1.000	1.000
	a pini duhan	.374	.237	.098	1.581	.116	1.000	1.000

a. Dependent Variable: Komp_pneumologjik

Ndërsa, kontrolli i hipotezave për domosdoshmërinë e variabla në modelim, me anë të t-testit, tregojnë se variabli 'area' ka ndikim të rëndësishëm në modelin e propozuar.

Modelimi me regres të shumëfishtë i 'komponentin pneomologjik specifik', duke përfshirë dhe variabla të tjerë si, 'konsumoni duhan në shtëpi', 'vitesh banorë në zonë', variabël i fundit mund të konsiderohet i vazhdueshëm, e rrit efikasitetin e modelit. Modeli justifikon tani 43% të ndryshimeve të 'komponentin pneomologjik specifik'

Dihet se minimumi i raportit të vëllimit të zgjedhjes që shfrytëzohet me variablat e pavarur duhet të jetë 5 me 1, ndërsa vlera e raportit të preferuar është 15 me 1. Në rastin në fjalë, ky raport më i lartë se 20.

Tabela 6–17 Rregulli i vlerësimit të regresit të shumëfishtë linear

Rregulli i njohur i vlerësimit të modelit të regresit^[11]

Koeficienti i korelacionit	Interpretimi
.80 në 1.00 (-.80 në -1.00)	Të koreluara shumë fort pozitivisht (negativisht)
.60 në .80 (-.60 në -.80)	Të koreluara fort pozitivisht (negativisht)
.40 në .60 (-.40 në -.60)	Mjaftueshëm të koreluara pozitivisht (negativisht)
.20 në .40 (-.30 në -.50)	Të koreluara dobët pozitivisht (negativisht)
.00 në .20 (.00 në -.20)	Aspak të koreluara

Tabela 6–18 Koeficientët e modelit të regresit të shumëfishtë

Model Summary

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Change Statistics				
					R Square Change	F Change	df1	df2	Sig. F Change
1	.656 ^a	.430	.408	.8356693197	.430	20.1584	4	107	.000

a. Predictors: (Constant), a pini duhan njeri ne shtepine tuaj, a pini duhan, area, prej sa vitesh banoni ne kete zone

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.	Collinearity Statistics	
		B	Std. Error	Beta			Tolerance	VIF
1	(Constant)	6.343	.943		6.726	.000		
	area	-1.222	.185	-.562	-6.619	.000	.740	1.350
	a pini duhan	.005	.401	.001	.012	.991	.908	1.101

prej sa vitesh banoni ne kete zone	.001	.005	.011	.134	.894	.730	1.36 9
a pi duhan njeri ne shtepine tuaj	.524	.200	.202	2.613	.010	.896	1.11 7

a. Dependent Variable: Komp_pneumologjik_specifik

Tabela e afishimeve të testeve, për vlerën apo jo të variablave ne modelim, tregon, se nga kater variablat e shqyrtuar, ndikim ne modelim ka *zona* në të cilen individi është banor, pra fakti që ai jeton në zonën e ndotur apo jo, dhe a *konsumon njeri duhan në shtëpi*.

Vlera e koeficientit është negative prane '*area*' , duke treguar se kur indeksi zonë bëhet 2, pra kur zona ndotet, vlere e variablit '*komponent_pneumologjik_specifik*' zvogelohet. Duke u perqëndruar tek mënyra e kodimit të të dhënave në bazën origjinale , pergjigjet 2 tregojne se nuk ka shenja te indikatorit pneumologjik në fjalë, ndersa 1 ka shenja të indikatorit. Kështu, variabli '*komponent pneumologjik specifik*' zvogelohet, do të thote se ka me shume vlere 1 në pergjigjet për indikuesit pneumatik. Pra, se ka me shume raste te shfaqjes se indikuesve të problemeve pneumologjike, kur kalohet në zonen e ndotur.

Ndërsa, vlere e koeficientit pranë komponentes tjetër signifikative '*konsumi duhan në shtëpi*', është pozitive. Kodimi i variablit '*konsumi duhan në shtëpi*', me 2 ne rast mos konsumi dhe 1 në rast konsumi, tregon se konsumi i duhanit në shtëpi (zvoglimi me nje njësi e variablit '*konsumi duhan në shtëpi*') shoqerohet me zvogëlim të '*komponent pneumologjik specifik*'. Bazuar në arsyetimet e mësipërme në me shume raste te shfaqjes se indikuesve të problemeve pneumologjike, ndërsa kemi konsum duhani në shtëpi dhe vlerat e faktoreve të tjerë nuk ndryshojnë.

Duke patur parasysh perfundimet tek 2.4, vlere e koeficientëve të modelit të regresit linear të ndërtuar, tregon shkallën e ndryshimit të '*komponent pneumologjik specifik*' në varësi të variablit, pranë të cilit ai ndodhet, kur komponentet e tjerë ne modelim mbahen konstantë.

S' mund të konkludohet bazuar në madhesinë e këtyre koeficientëve në vlerë absolute, për kontribut më të madh në modelim të një variabli se variabli tjetër.

6.2.2 Përfundime

- Ka më shume raste te shfaqjes se indikuesve të problemeve pneumologjike, kur kalohet në zonen e ndotur.
- Ka me shume raste te shfaqjes se indikuesve të problemeve pneumologjike, ndërsa kemi konsum duhani në shtëpi dhe vlerat e faktoreve të tjerë nuk ndryshojnë.

Analiza e problemeve pneumatike të banorëve, në bazë të një indikatorit të vetëm

Me interes është kontrolli i ndikimit në, **indikatorin më të rëndësishëm në pneumologji, 'kolle'**. Për të analizuar problemin në fjalë janë shfrytëzuar të dhënat e 426 subjekteve të marra rastësisht nga pesë fshatrat e sipërpërmendura, 196 subjekte janë nga zona e pastër, zona e kontrollit, pjesa tjetër 230 janë marrë nga zona e ndotur.

Për kontrollin e faktorëve, variabla, që ndikojnë në 'kolle', nisur nga përfundimi I mësipër është vënë re se ka një ndikim të *zonës* dhe *konsumit* të duhanit në shtëpi, kur kalohet në zonën e ndotur. Nga studimet e më parshme^[15], është vënë re dhe ndikimi I *'konsumit vetiak të duhanit'*. Kështu, që tek Kapitulli 5 , *Zbatim*, është paraqitur hollësisht, zbatimi i regresit logjistik, kur në modelim përdoret faktori *'zone, 'konsum_duhami' dhe 'konsum_duhami_ne_shtëpi'*.

Në bazën tonë të të dhënave, variabli *kolle* është variabel kategorik. Ai pasi i është bërë një transformim rekordesh, është i koduar, 0 nëse personi nuk ka kolle dhe 1 kur ka kolle. Kështu ai është bërë një variabël binar, që është modeluar me anë të modelit logjistik, për të kontrolluar ndikim apo jo të variablave të lartpërmendur.

Variabel i varur është konsideruar variabli *'kolle'*, variabël binar; ndërsa variablat me të cilën është ndërtuar parashikimi është variabli *'zone', area, 'konsum_duhami_ne_shtëpi'* dhe variabli *'konsum_duhami'*, i trajtuar më parë si ndikues në *'kolle'*^[14].

Modeli I komentuar gjërësisht tek Kapitulli 5 , nga kontrolli statistik ka dalë I vlefshëm. Janë komentuar koeficientët e modelimit, të paraqitura edhe tek Tabela 6-19. Por, duhet patur parasysh që me këto koeficientë, mund të modelohet lehtësisht variabli, *parashikim_kolle'*.

Tabela 6–19 Koeficientët e variablave në ekuacionin e modelimit

Variables in the Equation

	B	S.E.	Wald	df	Sig.	Exp(B)	95% C.I. for EXP(B)	
							Lower	Upper
Step 1 ^a								
area	1.028	.326	9.967	1	.002	2.796	1.477	5.293
pini_duhami	.147	.489	.090	1	.764	1.158	.444	3.018
duhami_shtëpi	.122	.037	10.602	1	.001	1.130	1.050	1.216
Constant	-2.267	.596	14.491	1	.000	.104		

a. Variable(s) entered on step 1: area, pini_duhami, duhami_shtëpi.

Në output-et e afishuara merren edhe informacione se c'do të arrijë modeli I fituar. Në tabelën e kontigjencës sipas testit Hosmer and Lemeshow, tregohet se behet automatikisht një ndarje e të dhënave në gjashtë grupe, e për secilin grup tregohet numri konkret i vlerave të vërtetuara të grupit dhe, numri i vlerave të grupit që siguron modeli. Krahasimi i këtyre vlerave tregon se për vlera të pritshme të ulta modeli nuk është aq i mirë.

Tabela 6–20 Tabela e klasifikimit pas modelimit

Classification Table^a

Observed	Predicted		
	Kolle		Percentage Correct
	jo	Po	
kolle jo	58	54	51.8
po	37	70	65.4
Overall Percentage			58.4

a. The cut value is .500

Tabela e klasifikimit tregon se modeli percakton saktësisht 51.8% të atyre që s'kane kolle, por 65.40% të atyre që kane kolle, ndërsa më parë, para modelimit, nuk arrihej të parashikohej asnjë nga keto persona. Kështu përgjindja totale e parashikimit me atë përfundim që jipej pa modelim është rritur me 18%

6.3 Analiza statistikore e problemeve hematologjike të banorëve të zonës (Problemi 3)

Ajri rrotull kësaj fushe, ku nafta është nxjerrë e përpunuar që me 1928 është i ndotur. Niveli i SO₂ është statistikisht i lartë (33 part per billion). Të dhënat e analizave të gjakut, përmbajnë rezultatet e analizave në gjak të banorëve të zonave në fjalë. Në këto rezultate vecohen eritrocitet, leukocitet dhe hemoglobina. Por, aty përmbahen edhe të dhëna për parametrat e tjerë limfocite, monocite, granulocite, që përfshijnë neutrofilet dhe eosinofile, që përbëjnë formulën e leukocideve.

Për secilin person janë marrë të dhëna të përgjithshme si mosha, gjinia, sa kohë kanë që banojnë në zonë, historiku i ndonjë sëmundje në gjak. Database është ndërtuar nga Instituti i Shëndetësisë dhe i Shëndetit Publik, duke patur parasysh listën e tërë faktorëve që mund të shkaktojnë probleme hematologjike.

Rrezultatet e të dhënave në gjak për banorët e zonës së ndotur, u krahasuan me rezultatet e të dhënave koresponduese të zonës së kontrollit. Analiza statistikore, e sqaruar hollësisht në seksionin 1.4, synonte të kontrollojë nëse këto mesatare janë statistikisht të njëjta, apo jo. Gjatë kryerjes së procedurave të testit Leven, e më pas t_{testit}, Studenti apo Welch, në varësi të rezultateve të testit të parë, afishohen dhe vlerat mesatare të parametrave në fjalë Tabela 6-21.

Tabela 6–21 Rezultatet e analizës deskriptive

Group Statistics					
	area	N	Mean	Std. Deviation	Std. Error Mean
eritrocide	1.00	56	4718750.00	511578.884	68362.604
	2.00	84	4575588.24	1034865.505	177477.970
leukocide	1.00	56	7866.07	1768.130	236.276
	2.00	84	12029.41	29198.054	5007.425
hemoglobine	1.00	55	13.33	1.480	.200
	2.00	83	12.07	2.535	.441

Si rezultat fuqisë së madhe kompjuterike të përpunimit të të dhënave, rezultatet e testeve shoqërohen edhe me afishimin e intervaleve të besimit të tyre, Tabela 6-22. Kështu, duke njohur edhe kufijtë normal ku duhet të lëvizin këto vlera, mund të lindin dyshime nëse këto kufij lëvizin jashtë normalitetit.

Tabela 6–22 Rezultatet e krahasimit të paramerave kryesor në gjak

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
eritrocide	Equal var. assumed	3.115	.081	.876	88	.383	143161.765	163447.949	-181656.688	467980.217
	Equal var. not assum.			.753	42.952	.456	143161.765	190189.052	-240403.509	526727.039
leukocide	Equal var. assumed	5.126	.026	-1.068	88	.289	-4163.340	3899.244	-11912.269	3585.588
	Equal var. not assum.			-.831	33.147	.412	-4163.340	5012.996	-14360.640	6033.959
hemoglobine	Equal var. assumed	.625	.431	2.931	86	.004	1.253	.427	.403	2.102

Equal var. not assum.			2.586	45.301	.013	1.253	.484	.277	2.228
--------------------------	--	--	-------	--------	------	-------	------	------	-------

Nga analiza e tri parametrave eritrocite, leukocyte dhe hemoglobinës u morrën këto afishime:

- Tabela e mësipërme tregon se, për eritrocitet plotësohet kushti i homogjenitetit të variancës së popullimit, dhe me pas, vërehet se s'ka dallim statistikor midis mesatares eritrocideve të grupit të ndotur dhe atij të kontrollit.
- p- vlera .412 tregon se niveli i leukocideve është i njëjtë në dy grupet në studim, nga ana statistikore, por nuk plotësohet kushti i homogjenitetit të variancës.
- p-vlera prane hemoglobinës .004 , dhe vlerat e afishuara të mesatares së diferencës, tregojnë se kjo vlerë është më e ulët në zonën e kontrollit.

Niveli mesatar i eritrocideve në të dy grupet e analizuara duket se është nën kufijtë normal. Duke patur parasysh dhe problemet për të cilat zona përmendej, u shtrua problemi ka apo jo shenja të anemisë në gjak. Sipas specialistëve kufiri përcaktues për të ishte 3500000. Pra u shtrua hipoteza:

hipoteza zero (ka shenja të anemisë)

$$H_0 : \mu = 3500000$$

hipoteza alternative (nuk ka shenja të anemisë)

$$H_a : \mu > 3500000$$

Testi, output-et mund të lexohen në Tabelën 6-23, vlera e statistikës t që jep një p-vlerë 8.427E-30, tregojnë se nuk ka shenja të anemisë.

Tabela 6-23 Kontrolli hipotezës për kufirin minimal

One-Sample Test

	Test Value = 3500000					
	t	df	Sig. (2-tailed)	Mean Difference	95% Confidence Interval of the Difference	
					Lower	Upper
eritrocide	15.886	136	.000	1128971.963	988070.29	1269873.64

Por, testi i hipotezës tjetër me vlerën 4500000 që sipas specialistëve është kufiri më i ulët i lejuar i këtij parametri

hipoteza zero

$$H_0 : \mu = 4500000$$

hipoteza alternative

$$H_a : \mu > 4500000$$

tregon se niveli I eritrocideve është në kufijtë minimum, të vlerave normale me nivel besimi 0.1 ($P=0.07$), Tabela 6-24.

Tabela 6–24 Kontrolli hipotezes për minimumin e vleras normale

One-Sample Test						
	Test Value = 4500000					
					95% Confidence Interval of the Difference	
	t	df	Sig. (2-tailed)	Mean Difference	Loëer	Upper
eritrocide	1.815	136	.072	128971.963	-11929.71	269873.64

Sic u përmend edhe më sipër, në vlerat e komponentit të tretë në gjak është vënë re një diferencë e vogël, po signifikative midis dy zonave në shqyrtim. Vlera e t-testit dhe p-vlera tregojnë influencën e zonës, në këte komponente të gjakut. Nisur nga përvoja e specialistëve, duke shfrytëzuar të dhënat e përgjithshme, u kontrollua efekti i dy faktoreve 'gjinia' dhe 'area'. Rezultatet e paraqitura në Tabelën 6-25 tregojnë ndikimin e dy faktoreve në shqyrtim. p-vlerat 0.000012 pranë 'zone'-s, dhe tej mase më e ulët, pranë 'gjinisë' tregojnë për efektin e dy faktoreve në fjalë përqindja e femrave në të dhënat Data Analizagjak. U vure re se përqindja e femrave në zonën industriale ishte (61.2%), ndërsa vlera e kësaj përqindje në zonën e kontrollit ishte (48%).

Tabela 6–25 Anova dypërmasore

Tests of Between-Subjects Effects					
Dependent Variable: hemoglobine					
Source	Type III Sum of Squares	Df	Mean Square	F	Sig.
Intercept	68859.115	1	68859.115	32842.165	.000
area	41.261	1	41.261	19.679	.000
gjinia	115.873	1	115.873	55.265	.000
area * gjinia	.354	1	.354	.169	.681
Error	853.344	407	2.097		
Total	70295.400	411			
Corrected Total	1007.894	410			

a. R Squared = .153 (Adjusted R Squared = .147)

Një analizë statistikore e detajuar u bë për komponentët e tjerë të gjakut, limfocite, monocite dhe granulocite (neutrofile dhe esonofile).

Studimi i tabelës së rezultateve, Tabela 6-26, tregon se, mesataret e gjithë parametrave të formulës së leukocideve janë me të ulta në zonën industriale se ato të nivelit mesatar .

Tabela 6–26 Analiza deskriptive

Groups	Formula e leukocideve	Statistika deskriptive			
		Mesatarja matematike	Devijimi standart	Q1	Q3
Zona kontrollit	Limfocide	33.46	1.2000	25.75	41.00
	Monocide	5.920	0.226	5.0000	7.0000
	Granulocise neutrofile eosinofile	59.03 2.9800	1.56 0.147	53.95 2.00	64.11 4.00
Zona industriale	Limfocide	24.04	2.60	0.37	37.75
	Monocide	4.149	3.3.7	0.070	7.000
	Granulocise neutrofile eosinofile	57.04 2.1150	1.36 0.2490	50.50 0.5200	64.75 3.000

Rezultatet e metodës Anova një-përmasore, Tabela 6-27, për *limfocitet*, *neutrofilet* dhe *eosinofilet* ($P= 0.002$, $F=9.98$) , ndotjes, në këto komponente të gjakut.

Tabela 6–27 Anova njëpërmasore per komponentet e tjera të gjakut

Tests of Between-Subjects Effects

Dependent Variable: monocide

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Intercept	17.271	1	17.271	2.551	.114
area	346.804	2	173.402	25.617	.000
Error	568.608	84	6.769		
Total	1481.180	87			
Corrected Total	915.411	86			

a. R Squared = .379 (Adjusted R Squared = .364)

Tests of Between-Subjects Effects

Dependent Variable: limfocide

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
--------	-------------------------	----	-------------	---	------

Intercept	583.342	1	583.342	2.433	.122
area	11900.009	2	5950.005	24.821	.000
Error	20855.661	87	239.720		
Total	53230.089	90			
Corrected Total	32755.670	89			

a. R Squared = .363 (Adjusted R Squared = .349)

Faktoret që mund të ndikojnë në shfaqjen e problemeve hematologjike për banorët, janë listuar nga Instituti i Shendetit Publik dhe në bazë tyre është hartuar dhe një pyetsor, me informacione të detajuara për 'moshë'n, 'histori te tilla sëmundjesh në familje', 'koheqëndrimi në zonë'.

Qëllimi është analiza dhe modelimi i një modeli statistiko-probabilitar që të shqyrtojë efektin e tyre në cdo komponente gjaku, duke synuar që këto parametra të modelojnë komponentin e shqyrtuar.

Për këtë u kthyen variablat në shqyrtim 'mosha' dhe koheqëndrimi në zonë ' në variabla kategorik, me vlera të përcaktuara si më poshtë:

$$kohëqëndrim_kategorik = \begin{cases} 0, & nqs\ personi\ ka\ më\ pak\ se\ 20\ vjet\ banorë \\ 1, & nqs\ personi\ është\ 20 - 40\ vjet\ banorë \\ 2, & nqs\ personi\ ka\ me\ shume\ se\ 40\ vjet\ banore \end{cases}$$

dhe

$$mosha_binar = \begin{cases} 0, & nqs\ person\ nën\ 25\ vjeç \\ 1, & nqs\ person\ mbi\ 25\ vjeç \end{cases}$$

Meqë komponentet e formulës së leukocideve janë të gjithë n. r. të vazhdueshme, dhe kërkohet të shqyrtohet efekti i n.r të tipit kategorik, një metodë e ndjekur vazhdimisht në studimin e problemeve të tilla përfshin kthimin e vlerave të parametrave të formulës së leukocidare në n.r kategorike. Pra nga parametrat e formulës së leukocideve, janë krijuar kater variabla binar yi, për $1 \leq i \leq 4$. Ato përcaktohen 1 nqs vlera e parametratit perkatës është me e vogël se kuartili i pare, dhe 0 në rastet e tjera. Është përdorur kuartili i pare si vlerë përcaktuese e funksioneve në fjalë. Kështu testohet nëse kapja e këtyre vlerave nën kuartil, ndikohet nga faktorët në shqyrtim, e për me tepër cilët faktorë janë ato që rrisin rezikun , riskun, e vlerave te tilla, jo të mira.

Me anë të regresit logjistik është shqyrtuar ndikimi i 'mosha_ binar' dhe 'koheqëndrimi _kategorik' në,' zonë' në komponentët e formulës së leukocideve dhe rezultatet janë paraqitur në Tabelën 6-28.

Tabela 6–28 Regresi logjistik për indikatorët rasteve problematike

Variables in the Equation		B	S.E.	Wald	df	Sig.	Exp(B)
Step 1 ^a	mosha_binar	-.124	.125	5.264	1	.035.	.674
	kateg_qendrim	-.642	.271	5.598	2	.018	.526
	Constant	.638	.636	1.008	1	.315	1.893

a. Variable(s) entered on step 1: kategqendrimi, mosha_binar

Tabela 6-28 tregon për ndikim të dy variablave (p-vlerat janë .025. dhe .018) mosha_binar dhe kateg_qendrim . Koeficientat negativ do të thotë se ulja e parametrit mosha_binar, pra kur flitet për banorë nën 25 vjeç , do të rrisë logit e probabiliteti, pra mundësitë e shfaqjeve të parametrave të ulta janë më të mëdha se mundëstë e shfaqjes së vlerave normale.

Po kështu, mund të përqëndrohemi edhe në faktorin tjetër të rëndësishëm të këtij modeli, ulja e kategorisë së qëndrimit, pra banorët që s'janë rezidentë për një kohë të gjatë, kemi një rritje të mundësive për kapjen e vlerave të komponentëve të gjakut nën normalen.

Pra, në moshat nën 25 vjeç ose tek banorët me një kohë jo shumë të madhe qëndrimi mundësitë, që shënjat e parametrave të ulta të formulës leukucidare të shfaqen, janë më të mëdha.

Gjithë parametrat e formulës së leukucideve janë variabla të pavarur të vazhdueshëm. Mosha_binar (dy nivele) and kohëqëndrimi (tre nivele), mund të shihen si faktore. Kështu mund të zbatohet edhe procedura e Anova-Dypërmasore për të kontrolluar se *kohe qendrimi* dhe *mosha_binar* kanë ndikim në parametrat e shqyrtuara, ose jo dhe të kontrollohet ndërthurja e faktoreve.

Kontrolli i vlerave të afishuara (Tabela 6-29) tregon :

Tek *monocidet*: Kemi ndikim të *kohëqëndrimit* p_value 0.010, por jo ndikim të '*mosha-binar*' Vlera e statistikes F=.257 dhe p_value .774

Vërehet se efekti ndërthurës i dy faktorëve në shqyrtim është josignifikativ (ky efekt paraqitet më qartë në Tabelën 29, p_value 0.324, në afishimeve për monocidet)

Tek *ezinofailet* ndikimi i *kohe_qëndrimit* është i pranueshëm me nivel besimi 0.1.

P_value e testit ANOVA për këtë variabël është .073 . Dhe, gjithashtu i qartë është dhe ndikimi i rolit ndërthurës i dy faktorëve në shqyrtim (p_vlera .018).

Tek *limfocidet* ndikimi i kohë_qëndrimit, shihet se pas kontrollit të metodës ANOVA

. Ne tabelen 6-29 p_vlera e afishuar është 0.004. Por, ndryshe nga parametrat e tjerë të formulës së leukocideve këtu, pra për këtë bazë

të dhënash, nuk vërehet ndërthurja mes dy faktorëve në shqyrtim.

Kështu, përfundimisht, ekzistojnë te paktën dy grupe kohë_qëndrimi ku vërehet ndryshimi I komponentëve të formules së leukocideve. Për përcaktimin e grupit të 'kohëqëndrimit' me ndryshim të dallueshem të komponentit në gjak është përdorur metoda Tukey (variabli kategorik kohëqëndrim' është me tri kategori ndaj kjo metodë jep vlera më të mira).

Përdorimi I post hoc Tukey, për komponentin *monocite*, tregon se deri në 20 vjetët e para të kohëqëndrimit niveli I uljes së monociteve është i ndryshëm (p=0.04), por me pas ai përmirësohet dhe është jo i ndryshëm pas 40 vjet kohëqëndrimi (p=0.21).

Sasia e *limfociteve* të banorëve për tri nivelet e kohëqëndrimit është respektivisht: 24.3 (për banorët që janë 1-20 vjet në këtë zonë); 23.1 (për banorët që janë 21-40 vjet në këtë zonë); 22.3 (për banorët që janë mbi 40 vjet në këtë zonë).

Për përcaktimin e grupeve me vlerë mesatare të ndryshme, të parametrave të gjakut, të përdoren Post Hoc e ANOVA së komponentes së gjakut ndaj *kateg_kohëqëndrimi*, do të thotë ta ndërtohet modeli pa marrë parasysh rolin ndërthurës së *moshës_binar* me *kateg_kohëqëndrim*.

Modelimi do të ishte më i saktë, për *ezinofailet* nëse kjo ndërthurje do të merret parasysh. Ndaj zgjidhjen e marrim duke ndërhyrë në sintaksën e SPSS, e cila mundëson ekzekutimin e komandave që analizojnë modelin, duke patur parasysh këtë faktor të ndërthurjes. Sintaksa e shkruar në SPSS është bërë si më poshtë:

DATASET ACTIVATE DataStudimi I analizes se gjakut zona industriale.

UNIANOVA ezinofile BY mosha_binar kategqendrimi

/METHOD=SSTYPE(3)

/INTERCEPT=INCLUDE

/POSTHOC=mosha_binar kategqendrimi(SCHEFFE)

*/PLOT=PROFILE(kategqendrimi*mosha_binar)*

/EMMEANS=TABLES(mosha_binar) COMPARE ADJ(LSD)

/EMMEANS=TABLES(kategqendrimi) COMPARE ADJ(LSD)

*/EMMEANS=TABLES(mosha_binar*kategqendrimi) COMPARE ADJ(mosha_binar)*

/PRINT=ETASQ HOMOGENEITY DESCRIPTIVE

/CRITERIA=ALPHA(.05)

*/DESIGN=mosha_binar kategqendrimi mosha_binar*kategqendrimi.*

Rezultatet e *ezinofaileve*, Tabela 6-30, tregojnë :

- Ne vitet e para të kohëqëndrimit vlera e limfocideve mesatare është përkatësisht .003 dhe .0035 për dy nivelet e moshave, pra është nën kufijtë minimal që është më pak se 0.35 %.
- Por tabela 6-30, e afishuar nga ndërhyrja në sintakse, duke shfaqur një diferencë negative, tregon se efekti i ndryshimit në limfocide vërehet shumë

mes grup moshës së re e tjetrës, në nivelin e parë të kohëqëndrimit. Vlera negative mesatare tregon madje, se vlera mesatare është me e ulët tek te rinjtë.

- Kjo diferencë mes grupmoshave nuk është më signifkative në nivelet e tjera të kohëqëndrimit.

Fillimisht mund të konkludohet, rendesia e përdorimit të efektit ndërthurës mes dy vektoreve në shqyrtim.

Përfundimisht,nga krahasimi I përfundimeve të regresit logjistik me variablin parametrat e gjakut të kategorizuar, dhe modelit te fundit të analizës, mund të themi se kthimi i një variabli të vazhdueshem në një variabël kategorik shoqërohet me tre probleme;

- *Së pari*, problemi kontrollit të testeve të shumëfishta, e kthen në krahasim të cifteve të kuartileve.
- *Së dyti*, përdorimi i metodes kërkon plotësimin e kushtit të homogjenitetit brenda grupit të krijuar, duke zëvendësuar kushtin e homogjenitetit mes grupeve të variablit origjinal të vazhdueshëm. Statistikisht ky kusht, për ndryshoren e rastit të vazhdueshme, është i njëjtë me kushtin e homogjenitetit të të gjithë popullimit.
- *Së fundmi*, vështirësi paraqet punimi me këtë metode, sepse pengon shumë krahasimin e rezultateve te punimeve me njëri tjetrin, sepse këto pika përcaktuese, prerëse që përdoren për të karakterizuar ndryshoret e rastit të reja mund të jenë të ndryshme në studime të ndryshme.

Tabela 6–29 Analiza dy përmasore

Tests of Between-Subjects Effects

Dependent Variable: monocide

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	.004 ^a	6	.001	2.416	.072
Intercept	.020	1	.020	72.340	.000
mosha-binar	.000	2	7.131E-5	.257	.774
kategqendrimi	.004	3	.001	4.306	.010
mosha-binar * kategqendrimi	.000	1	.000	.997	.324
Error	.012	43	.000		
Total	.185	50			
Corrected Total	.016	49			

a. R Squared = .252 (Adjusted R Squared = .148)

Dependent Variable: limfocide

Ndërtimi&analiza e një modeli probabilitar-statistikor për studimin e efektit të ndotjes

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	.133 ^a	5	.027	3.553	.009
Intercept	.794	1	.794	106.105	.000
mosha_binar	.003	1	.003	.390	.536
kategqendrimi	.113	3	.038	5.041	.004
mosha_binar * kategqendrimi	.005	1	.005	.723	.400
Error	.329	44	.007		
Total	5.874	50			
Corrected Total	.462	49			

a. R Squared = .288 (Adjusted R Squared = .207)

Dependent Variable: ezinofile

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	.002 ^a	6	.000	2.512	.086
Intercept	.004	1	.004	38.640	.000
mosha-binar	.000	2	6.502E-5	.621	.542
kategqendrimi	.001	3	.000	3.714	.018
mosha-binar * kategqendrimi	.000	1	.000	2.386	.073
Error	.005	43	.000		
Total	.049	50			
Corrected Total	.006	49			

Tabela 6–30 Kontrolli i efektit ndërthurës tek ezinofailet

Pairwise Comparisons

Dependent Variable: ezinofile

			Mean Difference (I-J)	Std. Error	Sig. ^c	95% Confidence Interval for Difference ^c	
						Lower Bound	Upper Bound
kategqendrimi	(I) mosha_binar	(J) mosha_binar					
0-20	nen25	mbi 25	-.006 ^a	.004	.020	-.004	-.002
	mbi 25	nen25	.006 ^b	.004	.042	.002	.004
20-40	nen25	mbi 25	.004	.005	.457	-.007	.015
	mbi 25	nen25	-.004	.005	.457	-.015	.007
mbi 40	nen25	mbi 25	-.003	.004	.494	-.011	.005
	mbi 25	nen25	.003	.004	.494	-.005	.011

Based on estimated marginal means

6.4 Analiza statistikore e ndikimit të zones në bujqësi

6.4.1 Baza e të dhënave

Fushat e Patos Marinëz dhe Kucovës janë zonat më të rëndësishme industriale në vend, pra zona ku përqëndrimi i gazeve ndotës, si SO_2 , është i lartë. Por, këto zona shfrytëzohen edhe si zona bujqësore. Nga projekti ANFI I Ministrisë së Mjedisit, i financuar nga Banka Botërore, u morën informacione për mbulesën bimë në përgjithësi dhe shtrirjen e klasifikimit të pyjeve të vendit. Mbi bazë të këtyre informacioneve në programin ArcGIS 10 u realizuan hartat e mbulimit bujqësor të zonave.

Të dhënat janë marrë nga tre fshatra. Dy prej tyre Ballsh-Marinëz dhe Sheqishtë-Belinë janë në zonën industriale, ndërsa tjetri në zonën joindustriale të Kurjanit. Siç thamë tek shtrimi i problemit, Figura 6-2 të dy fshatrat e parë mund të konsiderohen njësoj të ekspozuar nga SO_2 , ndërsa i treti si zonë e pastër, sepse nuk ka përqëndrim të puseve në të. Kështu në studim Kurjani është konsideruar si zona e *kontrollit*, ndërsa Sheqishtë- Belinë dhe Marinëz është konsideruar si zonë e *ndotur*.

Në zonat e mbjella me *bimëve bujqësore* të këtyre fshatrave (Figura 6-4), janë marrë informacione për prodhueshmërinë e shtatë kulturave bujqësore, nga Kurjani (në zonën e *kontrollit*) dhe pjesa tjetër nga Sheqishtë- Belinë dhe Marinëz (në zonën e *ndotur*). Informacionet e marra lidhen me sasinë e prodhimit të grurit, misrit patates, perimeve, fasules, frutave dhe jonxhës. Është vëzhguar dhe distanca e shtëpive nga pusët e naftës në zonë, tipi i sistemit të ujërave të zeza të përdorura, menyrat e grumbullimit të mbeturinave dhe niveli i përkujdesjes së tyre për produktet bujqësore.

6.4.2 Analiza statistikore e problemit

Të dhënat për prodhueshmërinë e shtatë produkteve bujqësore janë analizuar dhe, nga analiza deskriptive e të dhënave të zgjedhjes janë marrë rezultatet e paraqitura në tabelën 6-31

Figura 6–4 Harta e bimësisë në zonën në studim

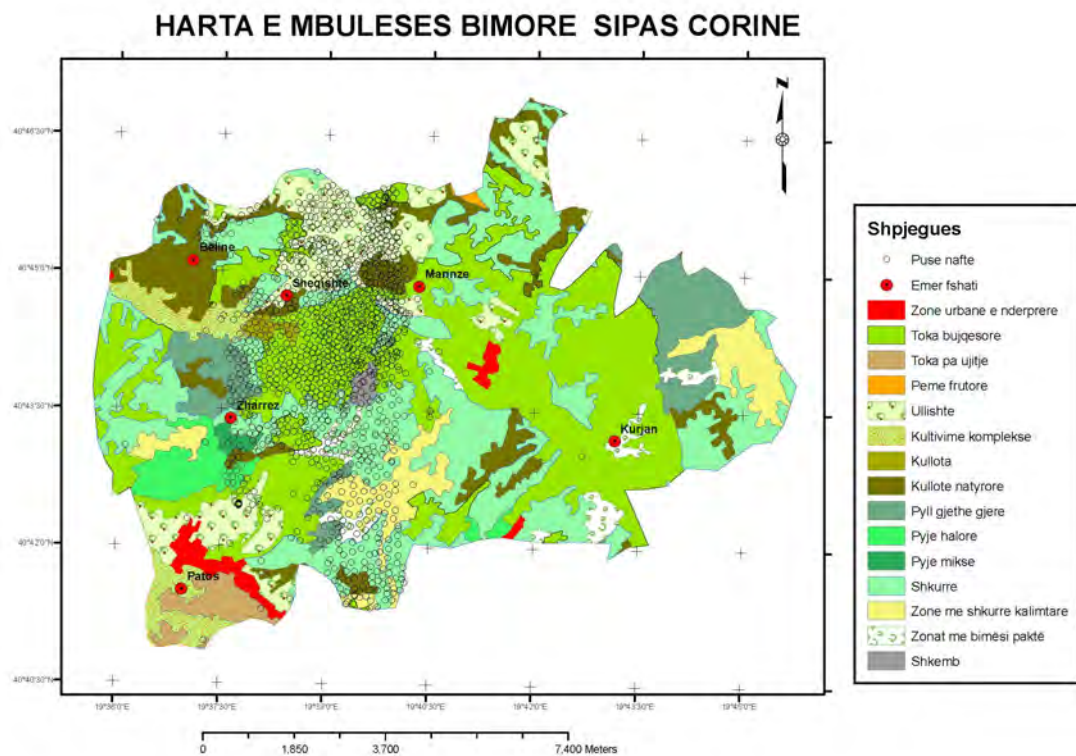


Figura 6–5 Harta e rrjetit hidrologjik

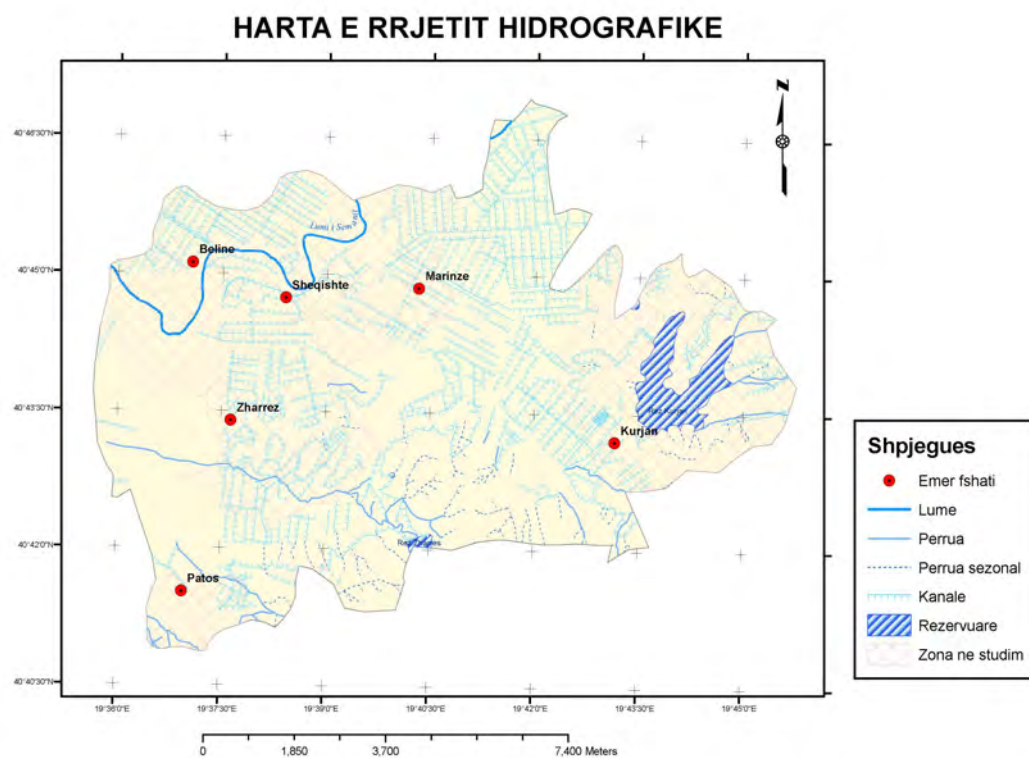


Tabela 6–31 Analiza deskriptive e bimëve

Descriptive Statistics				
Zona		N	Mean	Std. Deviation
Zona e kontrollit	grure	45	829.3333	331.24009
	miser	45	651.1111	179.16755
	patate	45	155.3333	156.05360
	perime	45	170.0000	108.39742
	jonxhe	45	2104.4444	1723.23105
	fasule	45	67.7778	37.15073
	fruta	45	184.4444	158.04999
Zona e ndotur	grure	110	812.9091	545.69263
	miser	110	630.9091	225.34192
	patate	110	67.2727	39.35044
	perime	110	110.6364	51.79843
	jonxhe	110	4510.0000	6030.83742
	fasule	110	42.5455	24.36106
	fruta	110	157.0000	242.22800

Duke pare tabelen 6-31, shihet se:

- *Mesatarja e prodhimit* të misrit, patates, perimes dhe fasules është më e ulët në zonën industriale.
- *Mesatarja e prodhimit* të frutave dhe jonxhës duket të jetë më e lartë në zonën e ndotur.

Kontrolli i kushtit të homogjenitetit të variaces, testi Levin, nuk jep te njejtin pergjigje per të gjithë komponentet e vektorit shtatë përmasor *Bimësi*, ndaj krahasimi I mesatareve te komponenteve te vektorit *Bimë* nuk është bërë me metodën ANOVA, por me anë të t- testit, sepse, sic dihet, kjo metodë bën të mundur vlerësimin e dy mesatareve pavarësisht plotësimit apo jo te kushtit në fjalë.

Në tabelen 6-32 janë paraqitur vlerat e *t*-testit Studenti apo *t*-testit Welch, në varësi të faktit qe komponenti I shqyrtuar e ploteson apo jo kushtin e homogjenitetit të variancës. Vlerat e statistikës F te testit Levin dhe p-vlerat koresponduese, tregojnë se varianca në zonat e ndotura është më e madhe për misrin (p-vlerë < 0.001), pataten (p-vlerë < 0.001), perimet (p-vlerë < 0.001), fasulen (p-vlerë < 0.001) dhe jonxhen (p-vlerë 0.028). Ky rezultat e ka dhe një farë justifikimi statistikor, sepse vëllimi I zgjedhjes ne zonën e ndotur eshte me i madh.

Tabela 6–32 Rezultatet e testit Levin dhe te t-testit për mesataret

Bimë	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
grure	.234	.629	.188	153	.851	16.424	87.35

miser	2.447	.120	.536	153	.593	20.202	37.70
patate	31.638	.000	3.737	46.306	.001	88.061	23.56
perime	21.773	.000	3.513	52.420	.001	59.364	16.90
jonxhe	4.941	.028	-3.820	142.163	.000	-2406	629.79
fasule	16.393	.000	4.202	60.051	.000	25.232	6.01
fruta	1.313	.254	.701	153	.485	27.444	39.16

Përdorimi *Independent-Simple T Test* në SPSS bën të mundur që, duke shfrytëzuar rezultatet e kësaj metode në varësi të përgjigjeve të testit Levin, të paraqitura në anën e djathtë të tabelës ³², të arrihet në këto përfundime:

- Ka ndryshim statistikor mes mesatareve të prodhimeve tek *patatja*, *fasulja* dhe *perimet*.
- Nuk ka ndryshim statistikor tek *frutat* (p-value 0 .701), *gruri* (p-value 0.188) dhe *misri* (p-value 0.536).
- Mesatarja e prodhimit të *jonxhes* është statistikisht më e lartë në zonën e ndotur, derisa diferenca e studiuar e vlerës së komponentit në fjalë në zonën e ndotur nga zona e kontrollit është negative.
- Mesatarja e *patates*, *perimes* dhe *fasules* është statistikisht më e ulët në zonën e ndotur.

Tabela 6–33 Rezultatet e analizës së kontrastit

Contrast Coefficients		
Contrast	area	polluted area
	controlled	
1	1	-1.3

Contrast Tests							
Bimë		Contrast	Value of Contrast	Std. Error	t	df	Sig. (2-tailed)
Patate	Does not assume 1 equal variances	1	67.8788 ^a	23.768	2.856	47.916	.006
Perime	Does not assume 1 equal variances	1	26.1727 ^a	17.387	1.505	58.402	.038
Fasule	Does not assume 1 equal variances	1	12.4687 ^a	6.307	1.977	71.498	.052

a. The sum of the contrast coefficients is not zero.

Më pas me anë të analizës së kontrastit mes dy grupeve për secilin nga tre komponentet që kanë prodhueshmëri më të ulët, pasi u përdorën koeficientët 1 dhe -1.3, u provua se për të tre bimët prodhueshmëria në zonën e kontrollit është statistikisht sa 1.3 e prodhueshmërisë në zonën e ndotur..

Programi SPSS mundëson llogaritjen e fuqisë së testit duke përdorur opsionin *the observed power of t-test*, me anë të efektit specifik të vëllimit të zgjedhjes, vëllimit të zgjedhjes, dhe p-vleres (significance). Duke analizuar këtë fuqi për tërë bimët arrihet në konkluzionin:

- Fuqia e testit për *perimet* është 0.96, për *fasulen*, *pataten* 0.99 dhe për *jonxhën* 0.788.
- Fuqia e vrojtuar e testit për *frutat* t-test është 0.11; e *grurit* t-test 0.053, dhe e *misrit* t-test është 0.083.

Kështu mund të konkludohet se po të kishim një vëllim zgjedhje më të madh, rezultati mund të ishte më signifikativ, duke patur parasysh ndikimin pozitiv të rritjes së vëllimit të zgjedhjes në fuqinë e testit. Për *grurin*, që sic lexohet nga tabela 6-32 e të dhënave deskriptive, ka devijim standart të madh, kjo mund të ndikonte në ngushtimin e intervaleve të besimit⁽³³⁾, duke çuar ndoshta përfundimisht në një vlerësim më të mirë të vlerave të popullimit.

Por duhet theksuar se fuqitë e ulta të testit në rastet e *grurit*, *misrit* dhe *frutave*, nuk do të thone se konkluzioni ynë, që nuk kishte ndryshim prodhueshmëria e ketyre produkteve në zonat e studiuara, nuk është i vertetë.

Problemi i fuqisë së testit ia vlen me shumë të studiohet nëse pergjigja e t-testit ka qenë signifikative, pra në rastin tonë ka treguar për ndikim të zonës në komponentin e studiuar bimë. Kështu, rezultatet e afishuara tregojnë një test më të dobët për *jonxhën*, sesa për *pataten*, *perimet* dhe *fasulen*.

Sasia e SO₂ që ndryshon në këto dy tipe zonash dhe influencon në produktivitetin e bimëve është faktor abiotik. Faktorët abiotik janë ato faktorë, fizik dhe kimik jojetësor, të cilët ndikojnë tek bimët. Por ka dhe një sërë faktorë të tjerë të tillë si zona ku luhet temperatura, lagështia e tokës, niveli i ndotjes, të cilët ndikojnë në këtë prodhueshmëri.

Temperatura është konsideruar e njëjtë në të dy sipërfaqet sepse zonat në fjalë janë të lokalizuara jo larg njëra tjetrës.

Duke shfrytëzuar Bazën e të Dhënave të Rrjetit Hidrologjik Kombëtar, të siguruar nga një projekt tjetër, Watershed, që ka studiuar basenët lumore dhe ka siguruar dixhitalizimin e të gjitha hartave topologjike të Shqipërisë, në GIS, u bë harta e rrjetit hidrologjik të zonave, Figura 6-5. Shqyrtimi i saj tregoi, se sistemi i kanalizimeve është i pasur në zonat në shqyrtim dhe nuk ka ndonjë dallim mes dy zonave për sa i përket problemit të lagështisë së tokës.

Niveli i ndotjes, në pamundesitë e matjes me aparaturë speciale të kontrollit të zonave, mendohet të varet nga distanca e shtëpive të intervistuarve nga pusët e naftës. Vlera e kësaj distance në metra përcaktojnë variablin numerik, *distanca*.

Për të studiuar edhe variabla të tjera të cilët ndikojnë në problemin e prodhueshmërisë së bimëve, janë përcaktuar edhe dy variabla të tjerë:

ujrat e zeza dhe *mbetjet urbane* u studiuan si variabla kategorik të koduara:

$$\begin{aligned} \text{ujrat e zeza} &= \begin{cases} 1 & \text{në rast kanalizimesh} \\ 2 & \text{në rastet e gropat septike} \\ 3 & \text{në rastet e përzierjes me ujë dhe rrjedhshëm} \end{cases} \\ \text{mbetjet urbane} &= \begin{cases} 1 & \text{në rast transporti të mbetjeve} \\ 2 & \text{në rast groposjesh të tyre në tokë} \\ 3 & \text{në rastin e djegjes së mbetjeve} \end{cases} \end{aligned}$$

Variablat *ujrat e zeza* dhe *mbetjet urbane* janë parë të lidhur me faktorët abiotik të tipit të dheut.

Të tri keto variabla janë trajtuar si variabla të pavarur. Të dhënat për to, si dhe të dhënat për prodhueshmerinë e bimëve janë analizuar me metodën e regresit për të identifikuar faktorë të tjerë abiotik që veprojnë tek bimët dhe mund të ndryshojnë ndikimin e SO₂ tek to.

Për të testuar nëse këto tri variabla ndikojnë në parashikimin e sasisë së prodhimit të *patates*, *fasules*, *perimeve* dhe *jonxhes* është përdorur metoda e regresit hirarkik të shumëfishtë⁽³⁴⁾. Si variabli i pavarur i listuar më parë, *the independent variable listed first*, u përdor variabli *zone*, dhe si variabla të shtuar, *the added variables*, *distance*, *ujrat e zeza* dhe *mbetjet urbane*.

Ndërsa për të studiuar efektin tek *gruri*, *misri* dhe *frutat* u përdor metoda e zakonshme e regresit linear, sepse nga kontrolli i bere ketu nuk u vu re ndikim i variablit *zone*.

Probabilitetet e statistikes F për të gjitha regresionet e studiara janë paraqitur në Tabelen 6-34, dhe tregojnë se ky model i regresit të shumëfishtë duhet marrë parasysh tek *perimet*, *jonxha*, *gruri*, *misri* dhe *frutat*, derisa vlerat e statistikes së siperpermendur kanë qenë jo-signifikative. Kështu hipotezat zero që nuk kanë lidhje gjithë variablat e pavarur me variablin e varur, *prodhueshmeria e bimes*, bie poshtë. Në tabele janë paraqitur dhe vlerat e statistikes *R²-change*. Informacioni i marrë nga variablat e shtuar, redukton parashikimin në prodhueshmerinë e bimëve (20% për *perimet*, 4.3% për *jonxhën*, 7.6% për *grurin*, 5.4% për *misrin*, and 6.7% për *frutat*).

Tabela 6-34 Rezultatet e regresit hirarkik dhe regresit të shumëfishtë të zakonshëm

Bimë	F-value	Sig.	R-Square
Perime	2.601	.05	0.199
Jonxhe	2.3	.07	0.043
Fasule	.144	.233	
Patate	1.591	.194	
Grurë	3.092	.018	0.076
Misër	2.126	.08	0.054
Fruta	3.743	.06	0.067

Duke analizuar tabelën e koeficientëve të regresit të shumëfishtë, shihet efekti i variablit *mbetje urbane*. Hipoteza zero $b=0$ ⁽³⁵⁾, nga afishimet e marra hidhet poshtë, duke treguar se nuk ka ndikim në komponentet e vektorit bimë. Por t-vlera 2.746 (sig. 0.009) për *jonxhën*, tregon për një farë lidhje të mundshme. Koeficientët pozitiv 10.801 dhe 1375 të modelit të regresit të linear si dhe drejtimi i kodimit të variablit *mbetje urbane*, tregon për një farë ndikimi pozitiv të trajtimit sa më të mirë të mbetjeve në rritjen e prodhimtarisë.

Ndikimi i faktorit tjetër abiotik *ujrat e zeza* duhet marrë parasysh, sepse tek *gruri* (t-test -3.415 dhe sig. 0.01), *misri* (t-test -2.692 dhe sig. 0.08) dhe *frutat* (t-test -3.787 dhe sig. < 0.01). Vlerat negative të koeficientëve përkatësisht -436, -150.5 dhe -215 në modelin e regresit linear përcaktojnë drejtimin e lidhjes. Duke verejtur kodin e

variablit *ujra të zeza*⁽³⁶⁾, ndërsa kodi rritet nuk ka asnjë proces sanitar në trajtimin e ujrave të zeza, arrihet në përfundimin se një sistem më i mirë i ujrave të zëza ndimon në prodhueshmëri me të lartë tek disa bime.

Analiza e rolit të faktorëve abiotik tek bimët, tregoi se efekti i dioksidit të squfurit reduktohet si rezultat i ndikimit të faktorëve të tjerë mbetje urbane dhe sistemi i ujrave të zeza.

6.5 Përfundime

Duke përmbledhur rezultatet në analizën e përfundimeve mund të themi:

- Testi Leven kontrollon barazimin e variancave në të gjitha nivelet e variablit të pavarur, të cilat përcaktohen nga variablat e pavarur. Afishimet e këtij testi (the spread-versus-level) përdoren për të kontrolluar testin dhe me pas duke patur parasysh kushtet perkatëse të secilës metode, aplikohet metoda e duhur në krahasimin e mesatareve.
- Llogaritja e fuqisë së testit na ndimon t'i tregojmë një kujdes më të madh të dhënave, por nuk tregon në rastet vecanerisht të një rezultati jo signifikativ, se vendimi i marre nuk është i drejtë. Ai ia vlen të studiohet në rastet kur përgjigja e testit është signifikative.
- Tipi i metodës së regresit të shumefishtë që duhet aplikuar, përcaktohet nga tipi i pyetjes të ciles ajo duhet të përgjigjet.

Studimi i problemit të ndikimit të dioksidit të sulfurit tek bimët tregoi se:

- Kur substanca në fjalë clirohet nga sipërfaqe të mëdha si zona industriale, nga kontiniere, puse ose fuçi, sasia e prodhimit tek disa bimë reduktohet mesatarisht 13%.
- Bimët sillen në mënyrë të ndryshme me dioksidin e squfurit, meqë kanë një ndryshime përse i përket mënyrës që ato e thithin gazin dhe aftësi të ndryshme në detoksifikimin e ndotësve dhe largimin e sasive të tepërta të squfurit. Konkretisht vetëm tek patatja, fasulja dhe tek perimet sasia e gazit të cliruar në këto zona ndikonte negativisht.
- Faktoret abiotik sistemi i ujrave të zeza dhe, në mënyrë jo shumë të sigurt dhe ai i mbetjeve urbane, kanë ndikim në produktivitetin e bimeve. Sa më shumë kujdes të tregohet në to nga ana sanitare, aq më pozitivisht ato ndikojnë tek bimët.

7 Shtojcë - Analiza statistikore e të dhënave për përcaktimin e faktorëve që ndikojnë në memorje

7.1 Disa fjalë për studimin

Këto ditë, gjithmonë e më shumë njerëz po zgjedhin palestrën si një mënyrë për të qenë të shëndetshëm në mungesë të aktivitetit fizik të përditshëm. Këtë e bëjnë për të “shpëtuar “ nga obeziteti dhe mbipesha, të dyja probleme të shpeshta. Gjithashtu, të gjithë e dinë, se nëse nuk ushtrohesh, muskujt dobësohen. Por, ajo që shumica e njerëzve nuk u kupton është që, dhe truri qendron në një gjendje më të mirë kur bën ushtrime fizike. Pra, ushtrimet fizike ndikojnë pozitivisht në të mësuarit dhe memorje.

Kjo është vënë re përmes një dekade studime në kafshë dhe njerëz.

Një nga provat e para për ndryshimet që sjellin ushtrimet në tru, është bërë që në vitet 1990. Dihej se truri kishte efekt në sjellje, por më parë askush nuk kishte pyetur për të kundërtën, e cila doli të ishte e vërtetë. Kur një grup shkencëtarësh të drejtuar nga Fred Gage, një neurolog nga Instituti Salk i Studime Biologjike, vendosi të krahasonte trurin e minjve të cilët kishin mundësi te pakufizuar për të lozur në një rrotë , ' minj vrapues' me ata të minjve' jovrapues'. Minjtë aktiv fizikisht kishin dyfishin e qelizave të reja nervore në hipokampus. Hipokampusi njihet si pjesa e trurit e përfshirë në aftësinë e të mësuarit dhe memorje. Eksperimentet treguan se minjtë e aftë fizikisht ishin më mirë në testet e memorjes. ^[37]

Sipas rezultateve, neurologia [Judy Cameron](#)^[37] e Universitetit të Pittsburgh, shtroi pyetjen nëse ushtrimet fizike do të ndikonin edhe tek majmunët. Në laboratorin e saj u trajnuan një grup majmunësh të një moshe jo të re, në pista vrapimi speciale një orë cdo ditë, pse ditë në javë për 5 muaj. Ndërkohë, një grup tjetër majmunësh nuk u përfshi në ushtrime. Ky eksperiment u përpoq të hidhte një vështrim te ndikimi i ushtrimeve mbi njerëzit, të cilët nuk janë si minjtë që vrapojnë vazhdimisht. Rezultatet treguan se majmunët në regjimin e caktuar, arriten të mësonin dyfish më shpejt se të tjerët.

Studimet mbi njerëzit, si ato të kryera nga Art Kramer, i cili vëzhgoi së si palestra mund të ndryshonte trurin gjatë plakjes, në Universitetin e Illinois, kanë treguar një lidhje të qartë mes të ushtruarit dhe përmirësimit të aftësive memorizuese gjatë jetës.

Sipas kërkimeve më të fundit, edhe njerëzit të cilët nuk kanë qenë shumë aktiv mund të përfitojnë nga aktivitetet , ndonëse në një moshe të thyer.

7.2 Cfarë e shkakton këtë rritje në memorje?

Ushtrimet fizike zmadhojnë volumin e gjakut në tru, duke rritur sasinë e oksigjenit të pompuar në çdo qelizë neuroniane.

Studimet kanë treguar që më shumë aktivitet fizik çon në një nivel më të lartë të molekulës BDNF^[38] (molekule kjo e njohur për vetinë neurotropike të trurit).

Një numër më i madh qelizash në rritje në hipokampus dhe një rritje në vëllim janë vënë re në subjekte që ushtrohen.^[37, 40]

Një studim i bërë në Irland^[38] zbuloi nëpërmjet analizave të gjakut se, pas ushtrimeve të rregullta, një nivel më i lartë BDNF, protein, gjendej në trup.

BDNF është një molekulë e familjes neurotropike. Kjo familje proteinash është e rëndësishme për mbijetesën , zhvillimin e funksionimin e neuroneve. Molekula e parë e kësaj familje NGF (Nerve Growth Factor) u zbulua nga fituesi i Çmimit Nobel, Levi Montalcini, dhe Hamburger më 1950. BDNF, pjesëtar i dytë i kësaj familje u zbulua në 1982. Më vonë janë zbuluar dhe të tjera proteina të ngjashme.^[39]

BDNF, më e studiuar e llojit të saj, është e rëndësishme për zhvillimin e trurit gjatë embriogenezës, fëmijërisë dhe rritjes. Kjo proteinë ndihmon mbijetesën e neuroneve duke ndikuar në rritjen, maturimin, dhe mirëmbajtjen e këtyre qelizave. Në tru, proteina e BDNF është aktive tek sinapses. Ajo rregullon plasticitetin sinaptik që është i rëndësishëm për procesin e të mësuarit dhe të memorizimit.

Duke patur parasysh nivelin e lartë të BDNF në gjakun e njerëzve me një jetesë të shëndetshme, shkencëtarët besojnë se aktiviteti fizik përmirëson memorjen dhe vëmendjen.^[3, 4]

Një studim i publikuar online 31 January 2014, nga Akademia Kombëtare Shkencore^[40] në zbuloi se ushtrimet rrisin volumin e hipokampusit të rriturit nga 55 në 80 vjeç, respektivisht me 2.12 në anën e majtë dhe 1.97 në të djathtë. Testet e memorjes mbi ta kanë treguar dhe një përmirësim të kësaj aftësie.^[41, 42]

Si konkluzion ushtrimet rrisin rrjedhjen e gjakut në tru dhe pompimin e oksigjenit. Në prezencën e një sasive të mjaftueshme oksigjeni më shumë neurone rriten në hapësirën e hipokampusit . Bazuar në kërkimet e mësipërme, me këtë studim është synuar të provohet se rezultatet e vëna re të rriturit janë të njëjtë edhe te të rinjtë dhe adoloshentët.

7.3 Metoda statistikore e ndjekur e interpretimi i rezultateve

Janë intervistuar rreth 300 persona të moshës 15-19 vjeçare nga shkolla të ndryshme të Shqipërisë; kryesisht nga Tirana e rrethinat e saj, nga Durrësi, Shkodra, Gjirokastra e Vlora. Para eksperimentit subjekteve iu shpjegua studimi e procedura. Personat që kishin probleme dhe çrregullime në memorje, humbje të mëdha në peshë, gjendje shëndetësore shumë të keqe nuk u përfshinë në studim.

Pyetsori përfshinte informacione demografike për moshën, gjininë dhe qytetin e lindjes, si dhe pyetje për mënyrën e jetesës. Pyetjet për mënyrën e jetesës ishin pyetje për konsumimin e alkolit, traditën e të ngrënit të ushqimeve shumë të rënda ose frutave.

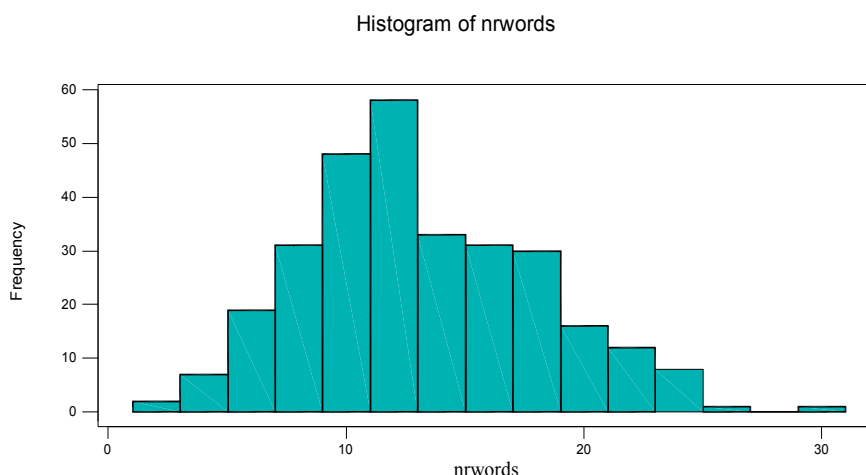
Për të përcaktuar memorjen e ndërmjetme (the intermediate term memory) është bërë një test për memorjen. Subjekteve i është dhënë një listë që përmban 40 fjalë të fushave të ndryshme, dhe ato janë lexuar për 5 minuta. Pastaj i është kërkuar të shkruanin fjalët që kanë mbajtur mend. Qëllimi i studimit ishte numri total i fjalëve të memorizuara.

Sasia e ushtrimeve fizike ishte e vështirë për tu matur, sepse ushtrime të ndryshme fizike ndryshojnë për nga koha që konsumohet për to, dhe për nga energjia. Kështu të luash tenis dhe të luash futboll nuk kërkon të njëjtin intesitet, por ato u maten njëllëj. Kjo do të thotë që *shpeshtësia në ushtrime* përmban sa 30 minuta ushtrime kryen gjatë një jave.

'*Të ngënit shëndetshëm*' ishte e vështirë gjithashtu të përcaktohej, kjo sepse personat konsumonin fruta psh, pas darke, së bashku me anëtarët e tjerë të familjes, duke i ndarë ato me njëri -tjetrin. Kështu, ishte e vështirë të rikujtohej dhe të përcaktohej sasia e konsumimit të frutave. Kjo është arsyeja pse ky variabël u konsiderua si variabël binar.

Duke shfrytëzuar këtë bazë të dhënash, u përcaktuan disa variabla. Variabli *nrffjalëve*, i cili përcaktohet nga numri i fjalëve të memorizuara për cdo person.

Figura 7-1 Histograma e numrit të fjalëve të memorizuara



U vu re se vlera më e vogël e tij ishte 1 dhe ajo më e larta 30. Më poshtë, Figura 7-1, paraqitet histogramën e këtij variabli me shkallë të zgjedhur 2 njësi. Derisa vëllimi i zgjedhjes ka qenë 300, pra i madh^[43], për kontrollin e normalitetit të të dhënave të këtij variabli është përdorur testi i normalitetit Shapiro-Wilky, i cili tregon për jo normalitet të të dhënave, tabela 7-1 dhe 7-2.

Tabela 7-1 Analiza Deskriptive e nrword

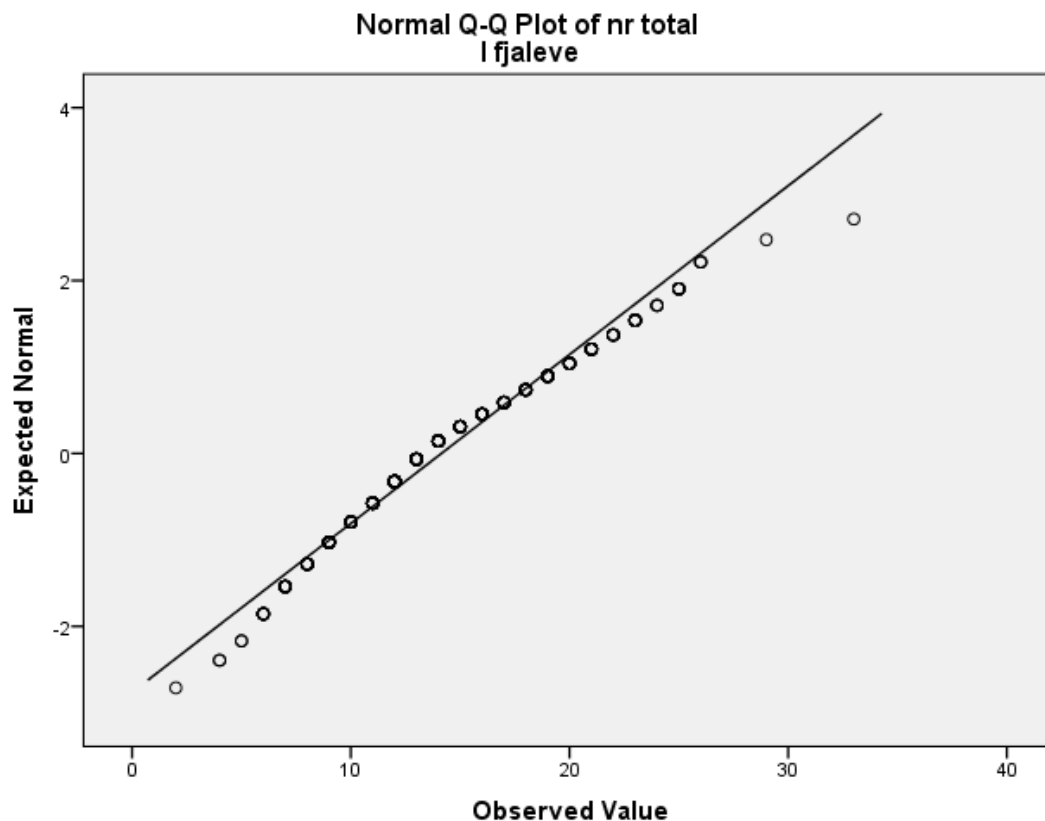
Descriptives		Statistic	Std. Error
nr total	Mean	14.144	.2966
I fjaleve	95% Confidence Interval for Lower Bound Mean	13.561	
	Upper Bound	14.728	
	5% Trimmed Mean	13.986	
	Median	13.000	
	Variance	26.218	
	Std. Deviation	5.1204	
	Minimum	2.0	
	Maximum	33.0	
	Range	31.0	
	Interquartile Range	7.0	
	Skewness	.529	.141
	Kurtosis	.083	.281

Tabela 7-2 Testi i Normalitetit të nrword

Tests of Normality						
	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
nr total	.109	298	.000	.976	298	.000
I fjaleve						

a. Lilliefors Significance Correction

Figura 7-2 Grafiku i ndryshimit të vlerave të *nr*fjalëve nga vlerat normale



Histograma tregon se kemi një pozitiv, ndaj punohet për normalizimin e të dhënave.

Transformimi me anë të funksionit *sqrt* të variablit në fjalë, krijon variablin e ri *sqrtnr*fjalëve, i cili përmbush kushtet e normalitetit. Vlera e Skewness .018 dhe e Kurtosis -.058 janë tashme shumë pranë 0. Rezultati i Shapiro-Wilky, Table 7-3, tregon se variabli i transformuar ka shpërndarje normale. Ndaj, në analizen e mëvonshme mund të përdorim këtë variabël të ri.

Tabela 7-3 Analiza Deskriptive dhe testi i normalitetit të *sqrt* nrword

Descriptives			
		Statistic	Std. Error
sqrtnword	Mean	3.6981	.03971
	95% Confidence Interval for Lower Bound Mean	3.6199	
	Upper Bound	3.7762	
	5% Trimmed Mean	3.6990	

Median	3.6056	
Variance	.470	
Std. Deviation	.68555	
Minimum	1.41	
Maximum	5.74	
Range	4.33	
Interquartile Range	.93	
Skewness	.018	.141
Kurtosis	-.058	.281

Tests of Normality

	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
sqrtnword	.074	298	.000	.993	298	.154

a. Lilliefors Significance Correction

Variablat si *gjinia* dhe ato që kanë të bëjnë me mënyrën e jetesës, si *konsumi i duhanit*, *konsumi i alkolit* dhe *numri i ushtrimeve fizike gjatë javës*, pas analizës së korelacionit mund të konsiderohen të pavarura.

Të studiosh rolin e *duhanit*, të *të ngënit*, dhe *shpeshtësisë së ushtrimeve në rrjfalëve*, është përdorur metoda e regresit stepwise^[44]. Rezultate e Tabelës 7-4 tregojnë se faktorë kryesorë në këtë modelim është *numri i ushtrimeve fizike gjatë javës* dhe *konsumi i duhanit*. Modeli në fjalë përcakton 40% të ndryshimeve të variablit në shqyrtim. Testi që $R=0$, pra nuk ka lidhje lineare mes modelit të propozuar dhe faktoreve të tjerë të testuar, bie poshte është signifikative me $p<0.001$ Madje $R=0.612$ tregon se lidhja në fjale është e fortë.

Figura 7-3 Grafiku i ndryshimit të vlerave të *sqrtnrfjalë* ve nga vlerat normale

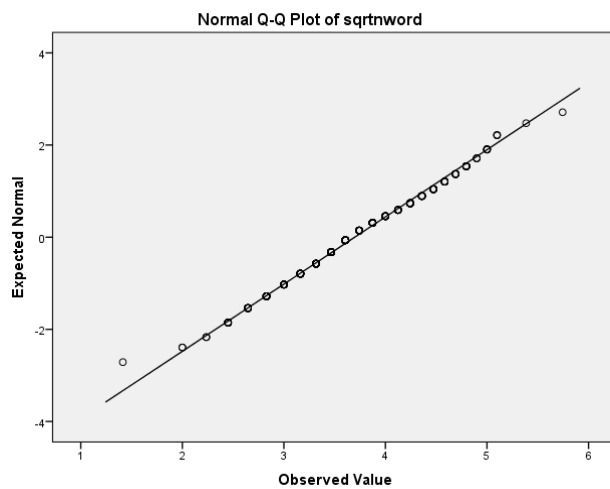


Tabela 7-4 Rezultatet e stepwise regression

Variables Entered/Removed ^a			
Model	Variables Entered	Variables Removed	Method
1	saherebejnes portnejave		Stepwise (Criteria: Probability- of-F-to-enter ≤ .050, Probability- of-F-to- remove ≥ .100).
2	duhan po/jo		Stepwise (Criteria: Probability- of-F-to-enter ≤ .050, Probability- of-F-to- remove ≥ .100).

a. Dependent Variable: sqtrnword

Model Summary

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	.624 ^a	.389	.387	.53684
2	.639 ^b	.409	.405	.52887

a. Predictors: (Constant), saherebejnesportnejave

b. Predictors: (Constant), saherebejnesportnejave, duhan po/jo

Mesatarja e *nrffalëve* kur subjektet nuk konsumojnë duhan është 13.270, kurse vlera mesatare e *nrffalëve* kur të personat konsumojnë duhan është më e ulët, 11.714. Vlerat mesatare janë përkatësisht 12.701 dhe 12.506 për personat që hanë shëndetshëm dhe jo. Por nga regresi stepwise këto variabla, *alkol* dhe *të ngrënit shëndetshëm* nuk janë marrë parasysh, pra nga baza e të dhënave nuk kemi ndikim të tyre në variablin *sqtrnrffalëve*.

Bazuar në të dhënat në fjalë fitet për ndikim të *alkolit* dhe *numri i ushtrimeve fizike gjatë javës*.

However the two last results of this investigations are in contrast with the research "Relationship of serum brain-derived neurotrophic factor (BDNF) and health-related lifestyle in healthy human subjects" by Ka Lok Chan, which reports higher BDNF level when the subjects eat healthy and when they smoke.

Por duhet thënë se këto dy rezultate bien në kundërshtim me rezultatet e punimit të Ka Lok Chan, "Lidhja e BDNF me mënyrën e jetesës së subjektit", në të cilën konkludohet se niveli i BDNF është më i lartë nëse subjektet ushqehen më shëndetshëm dhe kur konsumojnë duhan.

Studimi i rolit të duhanit në *nrffalëve*, tregoi se 212 subjekte që nuk konsumojnë duhan kanë një mesatare të *nrffalëve* prej 13.182. Ndërsa 85 njerëz të cilët konsumojnë duhan kanë një mesatare më të ulët të *nrffalëve*, prej 11.667 fjalësh. Koeficienti -.202 tregon për një rol negativ të konsumit të *duhanit*.

Ndërsa koeficienti tjetër .230 pranë *numri i ushtrimeve fizike* tregon që mesatarisht numri i fjalëve të memorizuara rritet nëse numri i orëve të ushtrimeve fizike rritet.

Nëqoftëse numri i ushtrime fizike rritet, nga dy herë në javë (sic është bërë zakonisht në shkollat tona) bëhet tre, atëherë numri i fjalëve të memorizuara ndryshon, nga 11.7 në 13.6.

Studimi u përqëndrua në 3 variabla: *gjini*, *numri i ushtrimeve fizike* dhe *sqtrnrffalë*.

Në këtë databazë për përcaktimin e gjinisë është përdorur kjo mënyrë: kodi 1=femra dhe kodi 2= meshkuj. Kështu nga 300 të intervistuar të të dy gjinive; 178 ishin femra (kodi 1) dhe 119 meshkuj (kodi 2).

Duke patur parasysh se shumica e adoleshentëve në Shqipëri bëjnë dy herë në javë ushtrime fizike, subjektet janë grupuar në tre grupe kryesore:

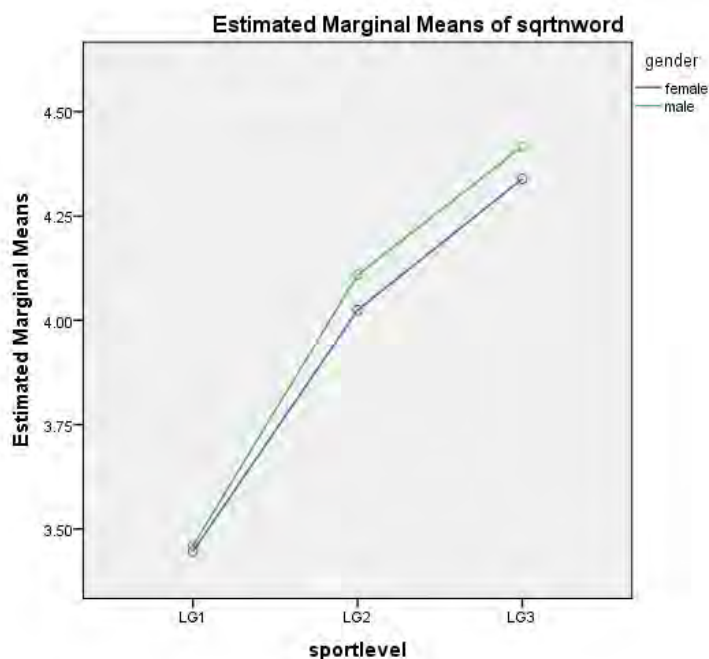
Grupi i nivelit 1 (LG1) 0-2 herë ushtrime fizike gjatë javës

Grupi i nivelit 2 (LG2) 3-4 herë ushtrime fizike në javë

Grupi i nivelit 3 (LG3) më shumë se 4 herë ushtrime fizike në javë.

ANOVA Dypërmasore^[9] është përdorur për të shqyrtuar ndikimin apo jo të variablave *gjini* me dy vlera , femra e meshkuj, dhe *grupi i nivelit* të ushtrimeve fizike, si variabel me tri vlera, në *sqrtnrfjalëve*. Grafiku i rezultateve, Figura 7-4, tregon se mesatarja e *sqrtnrfjalëve* është më e lartë në të tre grupet e niveleve për femrat se sa për meshkujt.

Figura 7-4 Grafiku i mesatareve të *sqrtnrfjalëve* sipas *gjinitë* dhe *grupit te nivelit* te shtrimeve fizike



Por rezultatet e Tabelës 7-5 (F .333 dhe Sig .565) tregojnë qartë se gjinia nuk sjell asnjë diferencë në *sqrtnrfjalëve* për cdo grup të nivelit të sportit.

Tabela 7-5 Rezultatet e ANOVA dy përmasore

Tests of Between-Subjects Effects

Dependent Variable: sqrtnword

Source	Type III Sum of Squares	df	Mean Square	F	Sig.	Partial Eta Squared
--------	-------------------------	----	-------------	---	------	---------------------

Corrected Model	37.779 ^a	5	7.556	21.672	.000	.271
Intercept	2148.261	1	2148.261	6161.623	.000	.955
gender	.116	1	.116	.333	.565	.001
sportlevel	29.465	2	14.733	42.256	.000	.224
gender * sportlevel	.072	2	.036	.103	.902	.001
Error	101.806	292	.349			
Total	4215.000	298				
Corrected Total	139.585	297				

a. R Squared = .271 (Adjusted R Squared = .258)

Duke u përqëndruar në afishimet pranë variablit nivel i sportit ($F=42.256$ dhe Sig. 7.62 E-17), kuptohet lehtë se *sqrtnrfjalëve* rritet ndersa numri i ditëve të ushtrimeve fizike rritet.

$F=.103$ and Sig.902 , tregojnë se nuk ka ndërhyrje mes dy faktorëve në fjalë.

Përqindjet e tre grupeve të ushtrimeve fizike sipas gjinise, tregojnë se femrat përbëjnë 73.91 % të LG1, ndërsa meshkujt 72.34% të LG3. Pra, shifrat se meshkujt kanë, në të memorizuarën e fjalëve, një mesatare më të lartë se femrat, lidhet me atë që djemtë bejnë më shumë ushtrime fizike se femrat.

Pas kësaj, nuk është e nevojshme të kontrollohet lidhja mes *gjinisë* e numrit të fjalëve të memorizuara, por vetëm ndikimi i *numrit të ditëve të ushtrimeve fizike* në *sqrtnrfjalë*.

Me metodën e krahasimit çift e çift, Tabela 7-6, është testuar se, kur ndryshimet mes dy niveleve të njëpasnjëshme të variablit *sportlevel* janë statistikisht të dukshme. Rritja, nëse përqëndrohesh në krahasimin e dy grupeve të njëpasnjëshme, është e ndjeshme kur bëhet kalimi nga Grupi i Nivelit 1 të sportit në Grupi i Nivelit 2(Sig 3.8 E-12). Ndërsa mes dy grupeve të fundit s' është statistikisht e ndjeshme me nivel rendesie $\alpha=0.01$.

Tabela 7–6 Rezultatet e krahasimit të sqtrnrffjalëve sipas grupeve të nivelit të sportit

Pairwise Comparisons

Dependent Variable: sqtrnword

(I) sportlevel	(J) sportlevel	Mean Difference (I-J)	Std. Error	Sig. ^b	95% Confidence Interval for Difference ^b	
					Lower Bound	Upper Bound
0-2 here sport	3- 4 here sport	-.616*	.082	.000	-.778	-.454
	me shume se 4 here sport	-.950*	.111	.000	-1.167	-.732
3- 4 here sport	0-2 here sport	.616*	.082	.000	.454	.778
	me shume se 4 here sport	-.334*	.124	.0077	-.579	-.089
me shume se 4 here sport	0-2 here sport	.950*	.111	.000	.732	1.167
	3- 4 here sport	.334*	.124	.0077	.089	.579

Based on estimated marginal means

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

b. Adjustment for multiple comparisons: Least Significant Difference (equivalent to no adjustments).

Nuk ka arsye të mos besojme se, rritja e numrit të ushtrimeve fizike rrit numrin e fjalëve të memorizuara. Për më tepër, nese një adoleshent bën një orë më shumë ushtrime fizikese sa bëhej zakonisht në Shqipëri, e kalohet nga LG1 në LG2, atëhehe pritet të ketë statistakisht aftësi më të larta memorizuese.

7.4 Përfundime

Bazuar në rezultatet e mësipërme mund të themi:

- Ka një ndikim të ushtrimeve fizike në memorjen e adoleshentëve. Rritja e sasisë së ushtrimeve fizike, rrit shkallën e memorjes tek adoleshentët.
- Nuk ka dallim statistikor në aftësitë memorizuese mes djemve e vajzave.
- Ka një ndikim të konsumimit të duhanit në memorje.
- Por, bazuar në këtë studim, s'mund të thuhet se të ngrënit shëndetshëm dhe alkoli kanë ndikim në aftësitë memorizuese të adoleshentëve.

Referencat

- [1] Scheffé, H. (1959). *The Analysis of Variance*. Wiley, NeW York. (reprinted 1999,ISBN 0-471-34505-9), pg. 154-197
- [2]ABDI & WILLIAMS, Analysis of contrast/www.utd.edu/~herve/abdi-contrasts2010-pretty.pdf .(2010)
- [3] Afifi A.A., Azen S.P. “Statistical analysis. Computer oriented approach”, (1979): Academic press New York, San Francisco, London,pg.198-274.
- [4] An introduction to Applied Multivariate Analysis, Tenko Raikov and George Marcoulides,99-150.
- [5] Gunst R.F., Mason R.L. “Regression analysis and its application”, Marcel Dekker Inc., NewYork (1980), pg.121-189.
- [6] Afifi A.A., Azen S.P. “Statistical analysis. Computer oriented approach”, (1979): Academic press New York, San Francisco, London, pg.124-197.
- [7] Miller, I., Freund, J. E., (1978): Probability and Statistics for Engineers Prentice-Hall, Inc.,pg.261-310.
- [8] IBM SPSS Regression 22, Chapter 2, Multiple Linear Regression, 2013.
- [9] Miller, I., Freund, J. E., (1978): Probability and Statistics for Engineers Prentice-Hall, Inc.,pg.226-264.
- [10] U.S. Department of Health and Human Services. (1998): Toxicological profile for sulphur dioxide. Public Health Service, Agency for Toxic Substances and Disease Registry.
- [11] Scheffé, H. (1959). *The Analysis of Variance*. Wiley, NeW York. (reprinted 1999,ISBN 0-471-34505-9)
- [12]ftp://public.dhe.ibm.com/software/analytics/spss/documentation/statistics/22.0/en/client/Manuals/IBM_SPSS_Statistics_Command_Syntax_Reference.pdf
- [13] The Truth about Principal Components and Factor Analysis
www.stat.cmu.edu/~cshalizi/350/lectures/13/lecture-13.pdf, (2009).
(Bartholomew, David J. , 1987).
- [14] Barbara G. Tabachnick and Linda S. Fidell, Using Multivariate Statistics, Third Edition. Logistic Regression chapter.(2012) .

- [15] Afifi A.A., Azen S.P. "Statistical analysis. Computer oriented approach", (1979): Academic press New York, San Francisco, London,pg.301-353.
- [16] Bartholomew, David J. (1987). Latent Variable Models and Factor Analysis. New York: Oxford University Press. Stone, James V. (2004).
- [17] Independent Component Analysis: A Tutorial Introduction. Cambridge, Massachusetts: MIT Press. Thomson,(1982),sec.2.
- [18] Using Multivariate Statistics, 6/E Barbara G. Tabachnick,Linda S. Fidell, pg.215-278'
- [19] Casella Berger Statistical Inference (2002) ,pg. 600-670
- [20] <https://mathdept.iut.ac.ir/sites/mathdept.iut.ac.ir/files/AGRESTI.PDF> pg120-208
- [21]Categorical Data Analysis, Alan Agresti, ISBN:0-471-36093-7, pg.590-630
- [22] Categorical Data Analysis, Alan Agresti, ISBN:0-471-36093-7, seks3.
- [23]Altman D.G. (1991) Practical Statistics for Medical Research. Chapman & Hall, www.blackwellpublishing.com/specialarticles/jcn_10_774.pdf
- [24] London. Polit D. (1996) Data Analysis and Statistics for Nursing Research. Appleton & Lange, Stamford, pg.579-671.
- [25] Categorical Data Analysis, Alan Agresti, ISBN:0-471-36093-7, pg.169-264.
- [26]http://practicalstats.labanca.net/index.php?title=Rules_of_thumb_for_interpreting_the_size_of_a_correlation_coefficient
- [27] IBM SPSS Regression 22, Chapter 2, Logistic Regression, 2013.
- [28] Casella Berger Statistical Inference, second Edition , George Casella, Roger L. Berger,2002, pg. 600-670.
- [29] Applied Logistic Regression , David W. Hosmer, Jr., Stanley Lemeshow, Rodney X. Sturdivant, pg.13-52.
- [30] <http://www.hcdoes.org/airquality/monitoring/so2.htm>
- [31] <http://www.missouribotanicalgarden.org/gardens-gardening/your-garden/help-for-the-home-gardener/advice-tips-resources/pests-and-problems/environmental/sulfur-dioxide.aspx>
- [32] An introduction to Applied Multivariate Analysis, Tenko Raikov and George Marcoulides Chapter 9,391-427

- [33] <http://www.dokeefe.net/pub/okeefe07cmm-posthoc.pdf>
- [34] An introduction to Applied Multivariate Analysis, Tenko Raikov and George Marcoulides, 210-276.
- [35] Applied Longitudinal Analysis, G. Fitzmaurice, N. Laird and Ware, 490-513
- [36] Practical Multivariate Analysis, Fifth Edition
By Abdelmonem Afifi, Susanne May, Virginia, pg. 231-264
- [37] <http://www.brainfacts.org/across-the-lifespan/diet-and-exercise/articles/2013/physical-exercise-beefs-up-the-brain/>
- [38] <http://www.ncbi.nlm.nih.gov/pubmed/21722657>
Griffin EW, et al. Aerobic exercise improves hippocampus function and increases BDNF in the serum of young adult males. *Physiol Behav.* 2011 104(5):934-41
- [39] Devin K. Binder and Helen E. Scharfman. Brain-derived Neurotrophic Factor *Growth Factors*. 2004 September; 22(3): 123–131.
- [40] Ka Lok Chan et al., Relationship of serum brain-derived neurotrophic factor (BDNF) and health-related lifestyle in healthy human subjects (The results of this study were approved by the Research Ethnic Committee of Hong Kong Hospital Authority and the Hong Kong Polytechnic University Human Subjects Ethics Subcommittee)
- [41] <http://www.pnas.org/content/108/7/3017.short>
Kirk I. Erickson et al., Exercise training increases size of hippocampus and improves memory, PNAS February 15, 2011 vol. 108 no. 7 3017–3022
- [42] Kirk I. Erickson et al Brain-Derived Neurotrophic Factor Is Associated with Age-Related Decline in Hippocampal Volume The Journal of Neuroscience, April 14, 2010 • 30(15):5368 –5375
- [43] Applied Longitudinal Analysis, G. Fitzmaurice, N. Laird and Ware, 490-513
- [44] Practical Multivariate Analysis, A.A. Afifi, S. May, V. Clark (2011), 231-264
- [45] An introduction to applied multivariate analysis, Tenko Raikov and George Marcoulides (2008), 120-314