

BETİMLEYİCİ, DURUM SAPTAYICI ARAŞTIRMALAR

MERKEZİ EĞİLİM ÖLÇÜTLERİ

Herhangi bir değişken için toplanan bireysel verilerin genellikle belirli bir merkezi nokta civarında değerler aldığı gözlemlenmektedir. Merkezi eğilim ölçütleri birimlerin hangi değer etrafında toplandığını gösteren ve dolayısıyla ana kitleyi temsil eden (niteleyen) sayısal değerlerdir.

Uygulamada en yaygın kullanılan merkezi eğilim ölçütleri ortalamalar (aritmetik ortalama, geometrik ortalama, harmonik ortalama), ortanca (medyan) ve mod (tepe değer) olarak sayılabilir.

Aritmetik Ortalama

Aritmetik ortalama bireylerin tamamının eşit ağırlığa sahip olduğunu varsayan bir merkezi eğilim ölçütüdür. Yığın için aritmetik ortalama aşağıdaki formül ile hesaplanır:

$$\mu = \frac{X_1 + X_2 + \cdots + X_N}{N} = \frac{\sum_{i=1}^N X_i}{N}$$

Burada, μ : yığın için aritmetik ortalamayı, N ise yığındaki birim sayısını göstermektedir

Yığından seçilen bir örnekten hesaplanan aritmetik ortalama ise $\bar{X}$ sembolü ile gösterilmekte olup hesaplama formülü aşağıdaki gibidir:

$$\bar{X} = \frac{X_1 + X_2 + \cdots + X_n}{n} = \frac{\sum_{i=1}^n X_i}{n}$$

Burada, n : örnekteki birim sayısıdır.

Aritmetik ortalamanın özellikleri

- ❖ Ölçme düzeyi en az eşit aralıklı olan değişkenler için hesaplanabilir. Diğer bir ifadeyle nicel değişkenler için hesaplanabilir.
- ❖ Hesaplama formülünde birimlerin tamamını kullanır. Dolayısıyla aykırı (uç) değerlerden etkilenir.

Geometrik Ortalama

Herhangi bir değişkenin zamana bağlı olarak izlediği değişkenlik oranı geometrik ortalama ile ölçülür. Hesaplama formülü aşağıdaki gibidir:

$$GO = \sqrt[N]{X_1 \cdot X_2 \cdots X_N} = (X_1 \cdot X_2 \cdots X_N)^{1/N}$$

Geometrik ortalamanın özellikleri

- ❖ Hesaplama formülünde birimlerin tamamını kullanır. Ancak aykırı (uç) değerlerden aritmetik ortalama kadar etkilenmez.
- ❖ Geometrik artış (ya da azalış) gösteren diziler için hesaplanması uygundur.
- ❖ Veriler arasında sıfır ya da negatif değerler olması durumunda geometrik ortalama hesaplanamaz.

Harmonik Ortalama

Bir aracın saatteki hızı, hane halkı başına gelir, arazi başına üretim gibi “Y başına X” olarak ifade edilen ve “X’in sabit” olduğu durumlar için kullanılan merkezi eğilim ölçütü harmonik ortalama değildir. Buna karşın “Y sabit” olduğunda tercih edilen merkezi eğilim ölçütü aritmetik ortalama değildir. Bireysel veriler için harmonik ortalamanın formülü aşağıda verilmiştir.

$$HO = \frac{N}{\sum_{i=1}^N \frac{1}{X_i}}$$

Örneğin, A şehirden B şehrine uzaklık 300 kilometre (km) olsun. A şehrinden B şehrine bir sürücü ilk 100 km’yi saatte 100 km hızla, ikinci 100 km saatte 80 km hızla ve geri kalan mesafeyi ise saatte 90 km hızla gitmiştir. Bu örnekte, Y saati ve X km’ yi göstermek üzere “Y başına X” ifadesi “saat başına km” olmaktadır. Aynı zamanda, X’ ler 100 km’ lik mesafeler ile sabit olup uygun ortalama ölçütü harmonik olmaktadır.

$$HO = \frac{3}{\frac{1}{100} + \frac{1}{80} + \frac{1}{90}} = 89.26 \text{ km/saat}$$

Harmonik ortalamanın özellikleri

- ❖ Veriler arasında sıfır ya da negatif değerler olması durumunda harmonik ortalama hesaplanamaz.
- ❖ Harmonik ortalama aritmetik ve geometrik ortalamaya göre aşağı eğilimlidir. Diğer bir deyişle $HO \leq GO \leq \mu$ olmaktadır.

Ortanca (Medyan)

Ortanca hesabından önce bireysel veriler küçükten büyüğe doğru sıralanmalıdır. Ortanca sıralı dizilerde ortaya düşen değerdir.

Gözlem sayısı N tek sayı olduğunda ortanca $X_{(N+1)/2}$ – inci değer iken,

N çift sayı olduğunda ortanca $(X_{N/2} + X_{(N+2)/2})/2$ olarak hesaplanır.

Ortancanın özellikleri

- ❖ En az sıralama düzeyinde ölçülen değişkenler için hesaplanabilir.
- ❖ Hesaplama formülü birimlerin tamamını kullanmaz. Bu nedenle uç değerlerden etkilenmez.
- ❖ Matematiksel işlemlere ortalama kadar uygun değildir.
- ❖ Birimlerin yarısı ortancadan küçük diğer yarısı ise ortancadan büyük değerler alır.

Tepe Değer (Mod)

En çok tekrar eden değere tepe değer adı verilir. 1'den fazla tepe değeri olabilir. Ya da her değer sadece 1 kez gözlemlendiğinde tepe değer hesaplanamaz.

Tepe değer özellikleri

- ❖ Sınıflama düzeyinde ölçülen değişkenler için hesaplanabilen tek merkezi eğilim ölçütüdür.
- ❖ Hesaplama formülü birimlerin tamamını kullanmaz. Bu nedenle uç değerlerden etkilenmez.
- ❖ Matematiksel işlemlere ortalama kadar uygun değildir.

MERKEZİ DAĞILIM ÖLÇÜTLERİ

Birimlerin hangi değer etrafında toplandığını gösteren eğilim ölçütleri (ortalamalar) çoğu zaman yeterli bilgiyi sağlayamamaktadır. Örneğin A ve B hisse senetleri fiyatlarına ilişkin günlük ortalamaların eşit olduğunu varsayalım. Bu durumda A ve B hisse senetleri için ortalama değerlerin kıyaslanması anlamlı olmamaktadır. Buna karşın hisse senedi fiyatları için ortalamalar aynı olsa da ortalama civarındaki dağılımları farklılaşacaktır. Hisse senedi fiyatı bakımından ortalama civarındaki dağılımı daha homojen (türdeş) olan zaman içerisinde daha istikrarlı bir seyir izlemektedir. Daha istikrarlı olan hisse senedi için risk daha az olacaktır.

Bir başka örnekte döviz kurları için verilebilir. Şöyle ki enflasyon hedeflemesi rejimi uygulayan Merkez Bankaları için döviz kurları serbest piyasada belirlenmektedir. Merkez Bankaları döviz kurunun değerine değil ancak kurlardaki oynaklığa (volatiliteye) bağlı olarak piyasaya müdahalede bulunmaktadır. Oynaklık kurların dağılımlarındaki genel gidişatın izlenmesi ile ölçülmektedir. Oynaklığın sayısal olarak ifadesi ise merkezi dağılım ölçütleri ile mümkün olmaktadır. Uygulamada en çok kullanılan merkezi dağılım ölçütleri varyans ve standart sapma ile değişim katsayısıdır.

Varyans ve Standart Sapma

Bireysel veriler ile aritmetik ortalama arasındaki farklar daima sıfırdır. Bu nedenle verilerin ortalamadan farklarına dayalı bir dağılım ölçütü hesaplanamaz. Varyans, veriler ile ortalama arasındaki fark karelerin ortalaması olarak tanımlanabilir. Bireysel veriler için yığın varyansının hesaplama formülü aşağıdaki gibidir:

$$\sigma^2 = \frac{1}{N} \sum_{i=1}^N (X_i - \mu)^2$$

Karesel bir ifadenin ortalaması daima sıfırdan büyük olacaktır. O halde varyans için tanım aralığı 0 ile $+\infty$ olmaktadır. Yorumu ise şu şekilde yapılabilir; Varyans değeri sıfıra yaklaştıkça dağılımın homojenliği artacaktır. Diğer bir ifadeyle varyans değeri sıfırdan uzaklaştıkça dağılımındaki homojenlik azalacak dolayısıyla heterojenlik artacaktır. Varyansın birimi X' in ölçü biriminin karesidir. Bu nedenle uygulamada varyans yerine onun pozitif karekökü olan standart sapma kullanılır. Böylece X' in ölçü birimi ile ifade edilen bir dağılım ölçütüne geçilmiş olmaktadır.

$$\sigma = \sqrt{\frac{1}{N} \sum_{i=1}^N (X_i - \mu)^2}$$

Yığından seçilen rastgele bir örnekten bulunan örnek varyansı için hesaplama formülü aşağıdaki gibidir:

$$s^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$$

Burada,

$\bar{X}$: Örnek ortalamasını

s^2 : Örnek varyansını

n : Örnekteki birim sayısını göstermektedir.

Değişim Katsayısı

Dağılım ölçütü olarak standart sapmanın birimi X 'in birimidir. Buna karşın iki ya da daha fazla yığın herhangi bir değişkenin dağılımları bakımından karşılaştırıldığında, ilgili değişkenin ölçme birimleri farklı olabilir. Örneğin Türkiye ile ABD'ni gelir dağılımları bakımından karşılaştırdığımızda, şayet gelirin birimi Türkiye için TL, ABD için US dolar ise, standart sapmanın kullanılması uygun olmayacaktır. Gerekli dönüşüm yapıp Türkiye için gelir US dolar cinsinden ifade edilse dahi gelirlere ilişkin ortalama değerler farklı ise yine standart sapmanın kullanılması uygun olmaz. Bu nedenle değişkenin ölçü biriminden bağımsız ve ortalamaları dikkate alan oransal bir dağılım ölçütüne ihtiyaç vardır. Bu ölçüt değişim katsayısıdır.

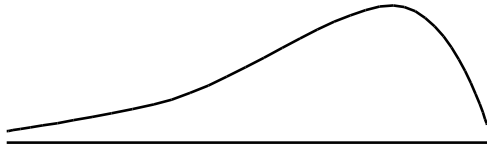
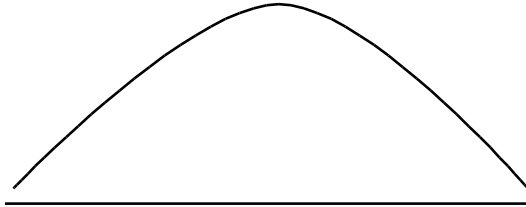
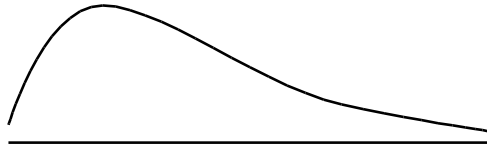
Değişim katsayısı standart sapmanın aritmetik ortalamaya oranı ile bulunan ve dolayısıyla birimden bağımsız olan bir ölçüttür. DK değişim katsayısını göstermek üzere bu katsayı aşağıdaki gibi hesaplanır:

$$DK = \frac{\sigma}{\mu}$$

Çarpıklık Katsayısı

Herhangi bir değişkenin ortalamalara göre simetrik olup olmadığı çarpıklık katsayısı ile hesaplanmaktadır. Asimetrik durum ise sola ya da sağa çarpık şeklinde iki farklı durumundan biri olarak karşımıza çıkmaktadır. Merkezi eğilim ölçütlerinden aritmetik ortalama, ortanca ve tepe değer birbirine eşitse dağılımın simetrik olduğu kararına varılır.

$$\alpha = \frac{\mu_3}{\sigma^3}$$

		
$\mu > Or > TD$ Sola Çarpık	$\mu = Or = TD$ Simetrik Dağılım	$TD < Or < \mu$ Sağ Çarpık

Çarpıklık Katsayısı

$\alpha < 0$ ise dağılım sola çarpıktır

$\alpha > 0$ ise dağılım sağa çarpıktır

$\alpha = 0$ ise dağılım simetriktir

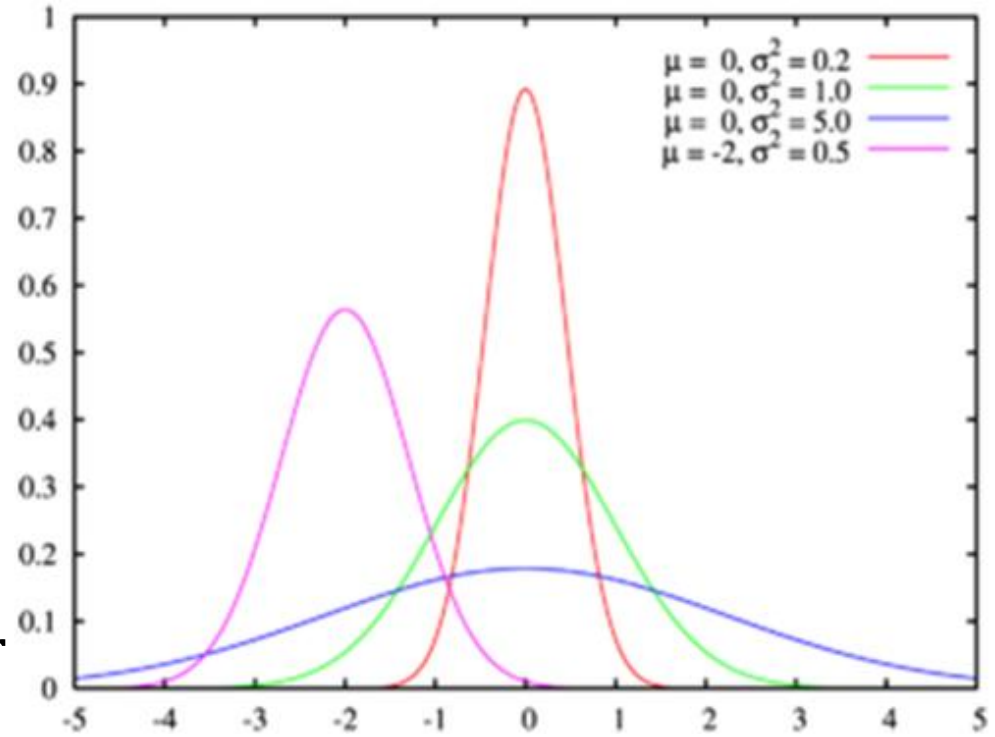
Basıklık Katsayısı

Dağılımın varyansı ile ilgili bir özelliktir. Verilerin varyansı küçükse daha sivri bir dağılıma, varyansı büyükse daha basık bir dağılıma sahiptir.

$$\gamma = \frac{\mu_4}{\sigma^4} - 3$$

- **Basıklık Katsayısı**

- $\gamma < 0$ ise dağılım basıktır
- $\gamma > 0$ ise dağılım sivridir
- $\gamma = 0$ ise dağılım normaldir



SPSS'TE VERİ GİRİŞİ

385 kişiye yapılan anket soruları aşağıdaki verilmiştir.

Soru 1. Cinsiyetiniz?

Erkek (1) Kadın (2)

Soru 2. Yaşınız?

Soru 3. Mesleğiniz?

Emekli (1) Esnaf/Zanaatkar (2) Ev Kadını (3)

İşçi (4) —İşsiz (5)— Memur (6)—————

Serbest Meslek (Doktor, Avukat, Mühendis, Eczacı vb.) (7)

Tüccar/İşadamı (8)

Soru 4. Eğitim Durumunuz?

İlköğretim (1) Lise (2) Üniversite (3)

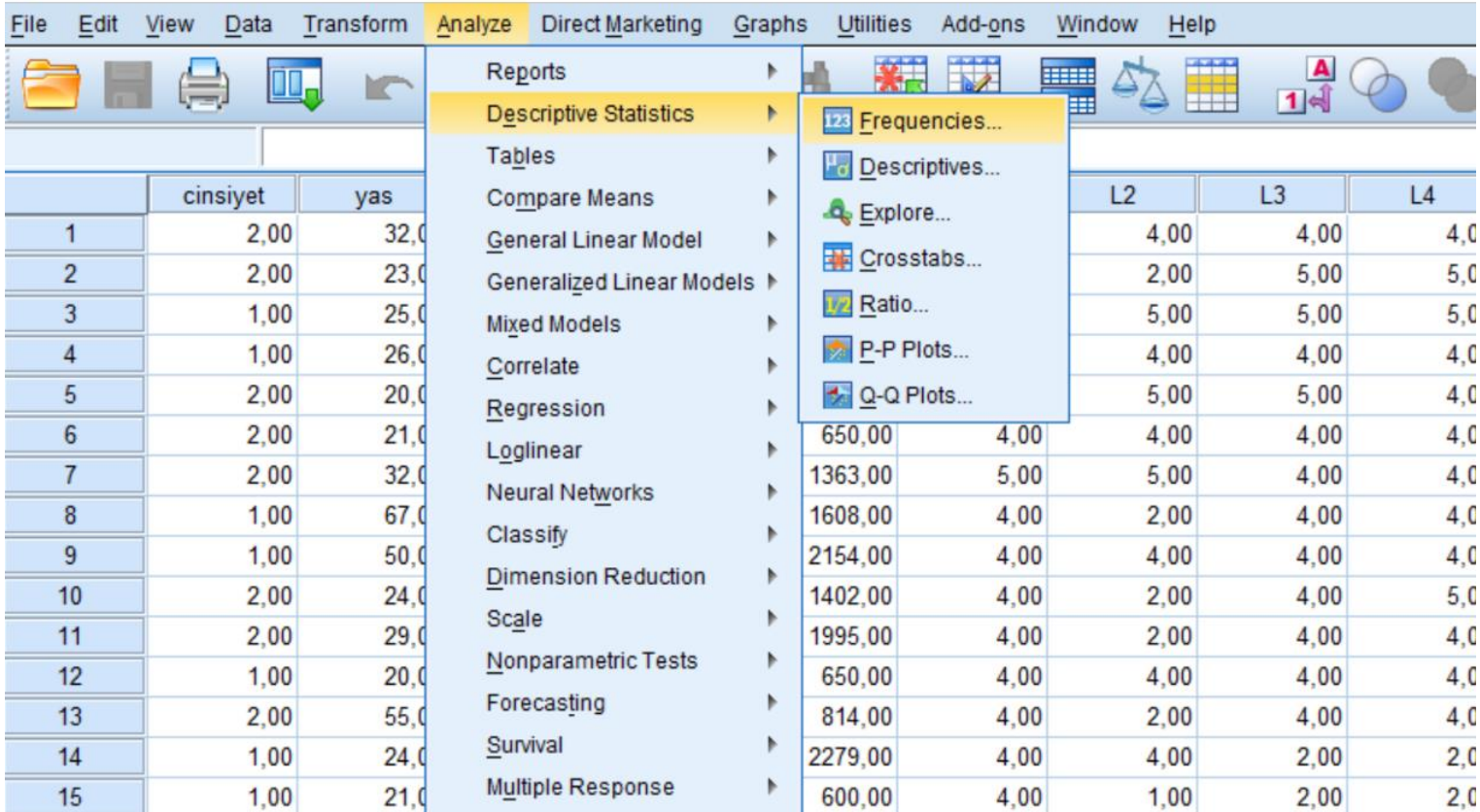
Soru 5. Aylık Toplam Geliriniz?

		Kesinlikle Katılmıyorum (1)	Katılmıyorum (2)	Kararsızım (3)	Katılıyorum (4)	Kesinlikle Katılıyorum (5)
L1	Verilen Hizmetin Kalitesi Yüksek					
L2	Verilen hizmet için bekleme süresi azdır					
L3	Çalışan kişilerin hizmet sunumu yüksektir					
L4	İlgili personelin müşteri şikâyetlerini yönetme davranışı iyidir					
L5	İşyerlerine ulaşım kolaydır					
L6	Çalışan personelin müşteriye davranışı iyidir					
L7	Yürütülen işlemler başarıyla tamamlanmaktadır					
L8	Verilen hizmetlerinin çeşitliliği yeterlidir					
L9	Verilen hizmetler için uygun fiyatlar alınmaktadır					
L10	Benzer hizmet sunan birimlerden daha güvenlidir					

DESCRIPTIVE STATISTICS MENÜSÜ

Frequencies Menüsü

veri.sav [DataSet1] - IBM SPSS Statistics Data Editor



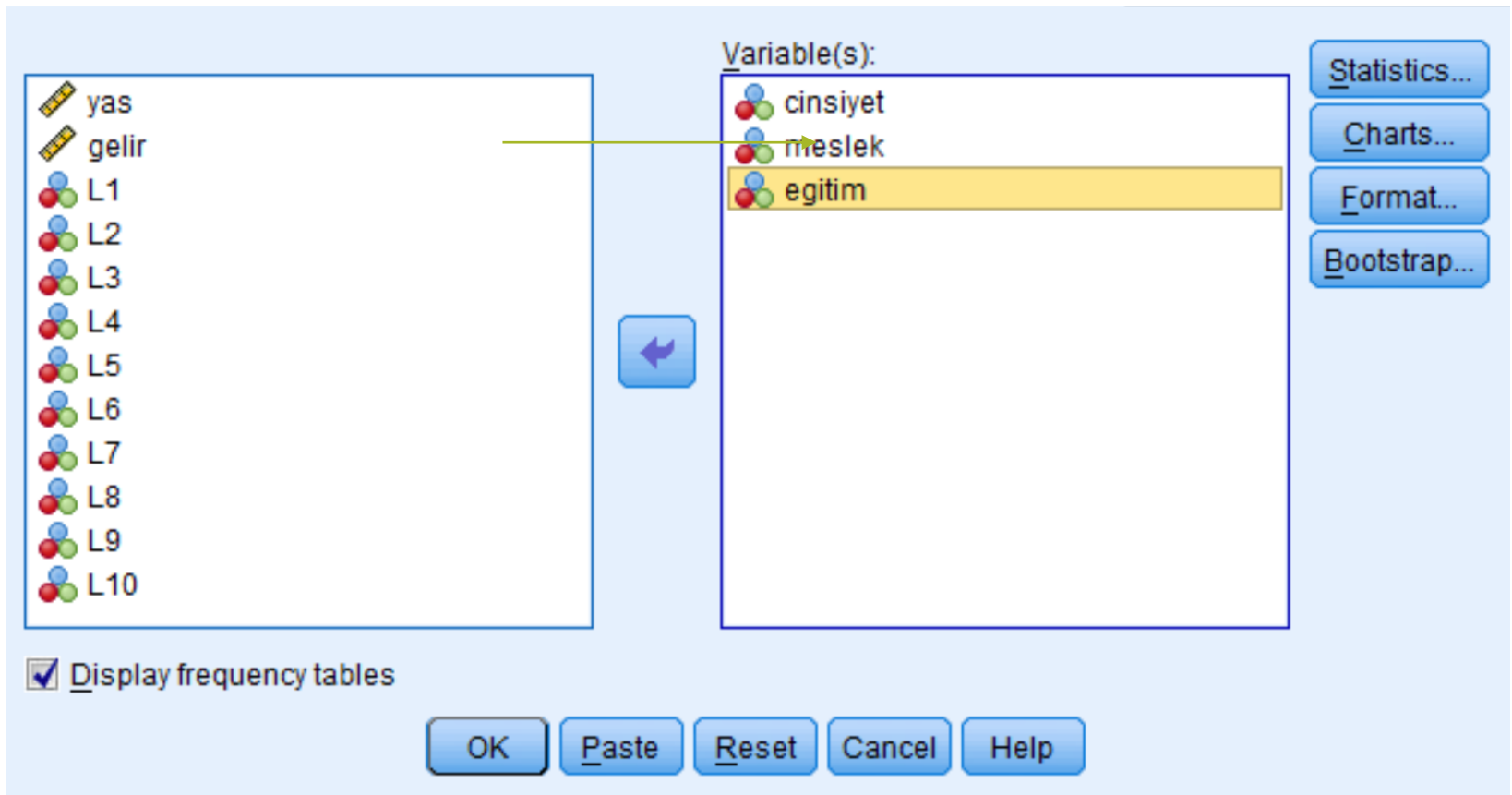
The screenshot displays the IBM SPSS Statistics Data Editor interface. The 'Analyze' menu is open, and the 'Descriptive Statistics' submenu is selected, with 'Frequencies...' highlighted. The background data table has columns 'cinsiyet' and 'yas'.

	cinsiyet	yas
1	2,00	32,0
2	2,00	23,0
3	1,00	25,0
4	1,00	26,0
5	2,00	20,0
6	2,00	21,0
7	2,00	32,0
8	1,00	67,0
9	1,00	50,0
10	2,00	24,0
11	2,00	29,0
12	1,00	20,0
13	2,00	55,0
14	1,00	24,0
15	1,00	21,0

The 'Frequencies' submenu is open, showing the following options:

- Frequencies...
- Descriptives...
- Explore...
- Crosstabs...
- Ratio...
- P-P Plots...
- Q-Q Plots...

Verilerin frekans tablosunu, merkezi eğilim ve dağılım ölçütlerini hesaplar ve grafiklerini çizer.



Yukarıdaki pencerede frekans dağılımı oluşturulmak istenilen değişkenler *Variable(s)* kısmına atılarak OK'e basılır.

Eğitim Düzeyiniz

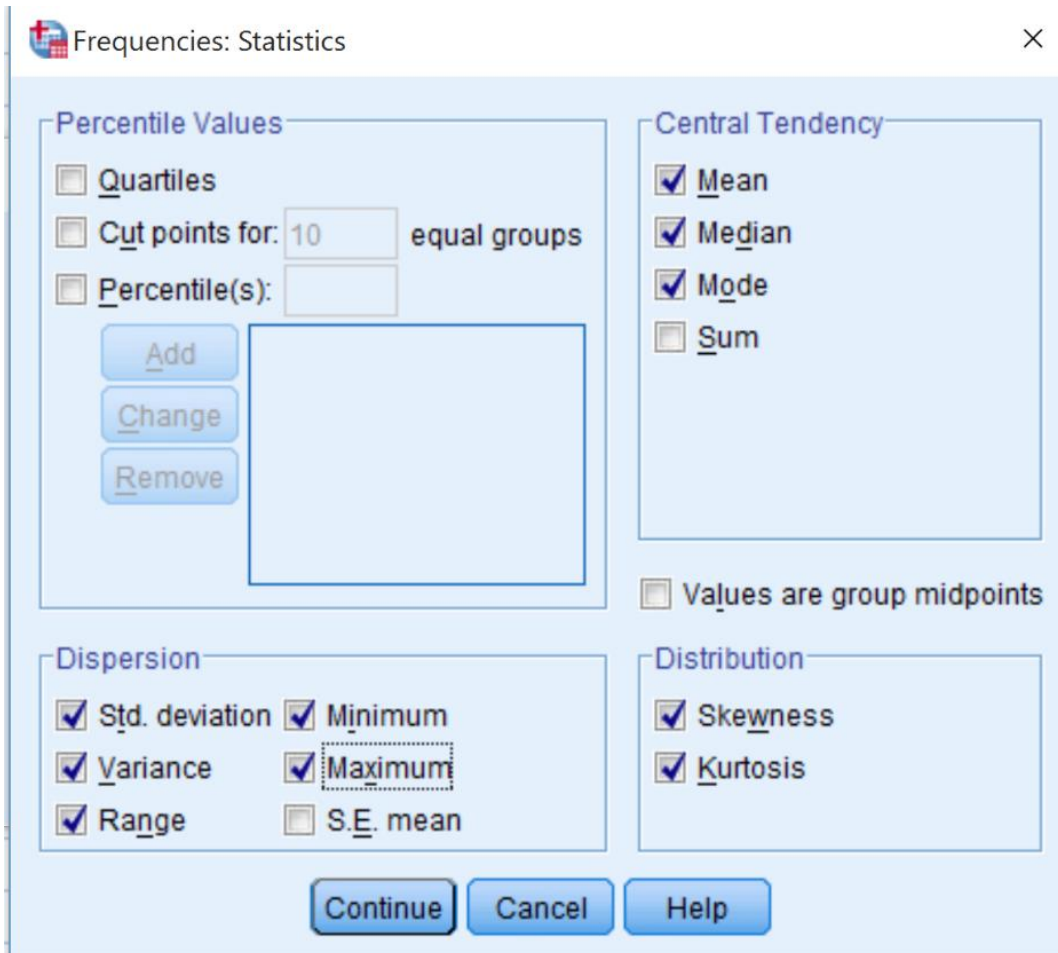
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	İlköğretim	111	28,8	28,8	28,8
	Lise	172	44,7	44,7	73,5
	Üniversite	102	26,5	26,5	100,0
	Total	385	100,0	100,0	

Yukarıdaki verilen eğitim düzeyine ait frekans dağılımı tablosunda ankete katılan 385 kişinin 111'i İlköğretim, 172'si Lise ve 102'si Üniversite mezunudur. *Percent* sütunu frekansların yüzdelik dağılımını vermektedir. Buna göre ankete katılanların %28.8'i İlköğretim, 44.7'si Lise ve 26.5'u Üniversite mezunudur.

Şayet değişkenin içinde 'kayıp gözlem' varsa *Percent* yerine *Valid Percent* sütunu yorumlanır.

Merkezi Eğilim ve Dağılım Ölçütlerinin Elde Edilmesi

SPSS'te *Analyze* → *Descriptive Statistics* → *Frequencies* → *Statistics* kısmına girilerek ilgili merkezi eğilim ve dağılım ölçütleri seçilir.



The image shows the 'Frequencies: Statistics' dialog box in SPSS. It is divided into four main sections: Percentile Values, Central Tendency, Dispersion, and Distribution. The 'Percentile Values' section has checkboxes for 'Quartiles', 'Cut points for: 10 equal groups', and 'Percentile(s):'. Below these are 'Add', 'Change', and 'Remove' buttons. The 'Central Tendency' section has checkboxes for 'Mean', 'Median', 'Mode', and 'Sum'. The 'Dispersion' section has checkboxes for 'Std. deviation', 'Variance', 'Range', 'Minimum', 'Maximum', and 'S.E. mean'. The 'Distribution' section has checkboxes for 'Skewness' and 'Kurtosis'. At the bottom are 'Continue', 'Cancel', and 'Help' buttons.

Section	Option	Selected
Percentile Values	Quartiles	☐
	Cut points for: 10 equal groups	☐
	Percentile(s):	☐
Central Tendency	Mean	☒
	Median	☒
	Mode	☒
	Sum	☐
Dispersion	Std. deviation	☒
	Variance	☒
	Range	☒
	Minimum	☒
	Maximum	☒
	S.E. mean	☐
Distribution	Skewness	☒
	Kurtosis	☒

Mean = Aritmetik Ortalama

Median = Medyan

Mode = Mod (Tepe Değer)

Std. Deviation = Standart Sapma

Variance = Varyans

Range = Açıklık

(Maksimum-Minimum)

Skewness = Çarpıklık

Kurtosis = Basıklık

Statistics

Aylık Toplam Geliriniz

N	Valid	385
	Missing	0
Mean		2009,6000
Median		1672,0000
Mode		500,00 ^a
Std. Deviation		1174,95971
Variance		1380530,319
Skewness		1,062
Std. Error of Skewness		,124
Kurtosis		,542
Std. Error of Kurtosis		,248
Range		5000,00
Minimum		500,00
Maximum		5500,00

a. Multiple modes exist. The smallest value is shown

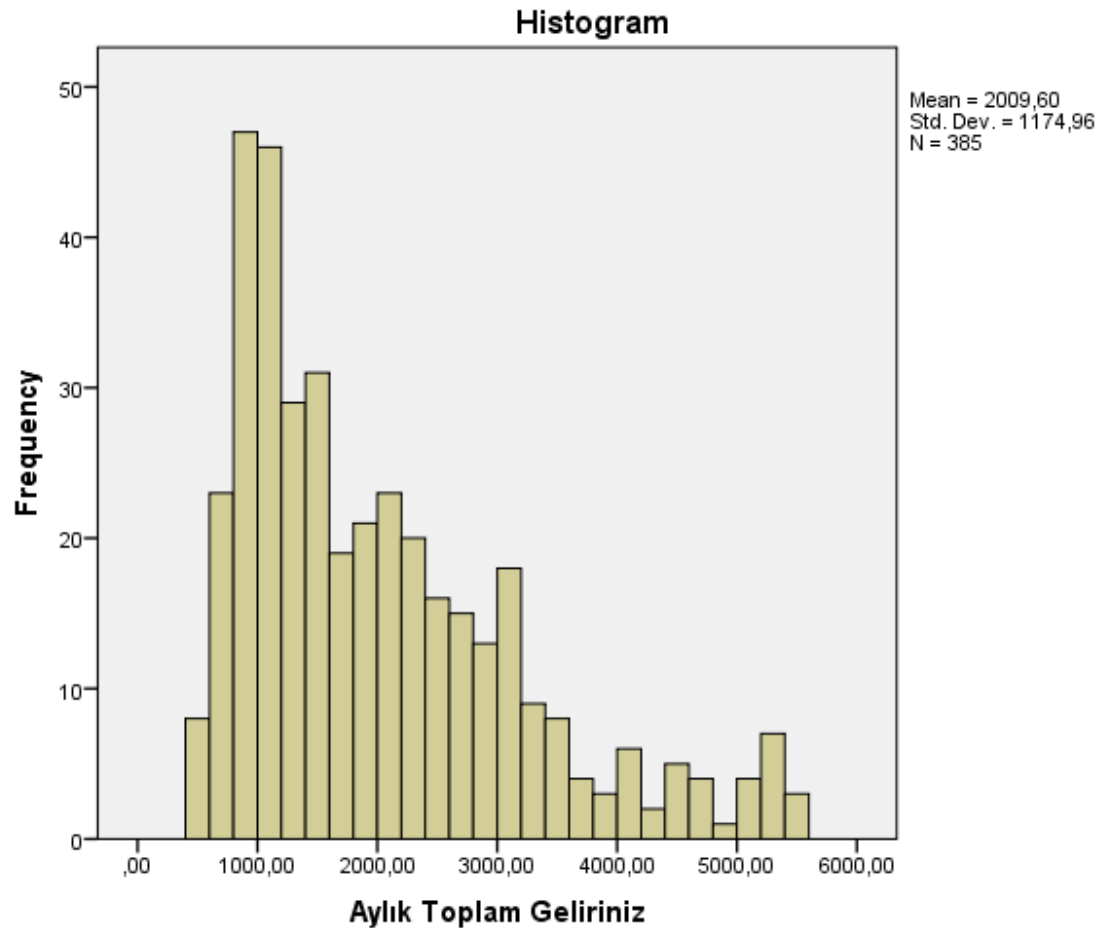
Ankete katılan 385 kişinin aylık ortalama geliri 2009 TL, Ortanca gelir 1672 TL, standart sapma 1174 TL dir. Mod değerinin üzerinde tanımlanan ‘a’ açıklaması birden fazla mod değeri olduğunu göstermektedir.

Çarpıklık ve Basıklık ölçütleri incelendiğinde aylık gelir değişkeni Çarpıklık=1.062 ile sağa çarpık bir dağılıma, Basıklık=0.542 ile de sivri bir dağılıma (varyansı küçük) sahiptir.

Hangi ölçme düzeyinde hangi merkezi eğilim ve dağılım ölçütünü kullanmak gerektiği aşağıdaki tabloda verilmiştir.

Ölçme Düzeyi	Ölçüt
Sınıflama (Nominal)	Mod
Sıralama (Ordinal)	Mod, Medyan
Eşit Aralıklı (Interval)	Tüm merkezi eğilim ve dağılım ölçütleri kullanılır
Oranlama (Ratio)	

Aylık toplam gelir değişkeninin dağılımını Histogramına bakarak görebiliriz. SPSS’de *Analyze* → *Descriptive Statistics* → *Frequencies* → *Charts* kısmına girilerek Histogram çizdirilir.



EXPLORE MENÜSÜ

Tüm birimlerin yada her bir gruptaki birimlerin betimleyici istatistiklerini hesaplar.

veri.sav [DataSet1] - IBM SPSS Statistics Data Editor

The screenshot shows the IBM SPSS Statistics Data Editor interface. The 'Analyze' menu is open, and the 'Explore...' option is selected. The data editor shows a dataset with columns 'cinsiyet' and 'yas'. The 'Explore' menu is open, showing options like Frequencies..., Descriptives..., Explore..., Crosstabs..., Ratio..., P-P Plots..., and Q-Q Plots... The 'Explore...' option is highlighted.

	cinsiyet	yas
1	2,00	32,0
2	2,00	23,0
3	1,00	25,0
4	1,00	26,0
5	2,00	20,0
6	2,00	21,0
7	2,00	32,0
8	1,00	67,0
9	1,00	50,0

4 : L2 4,00

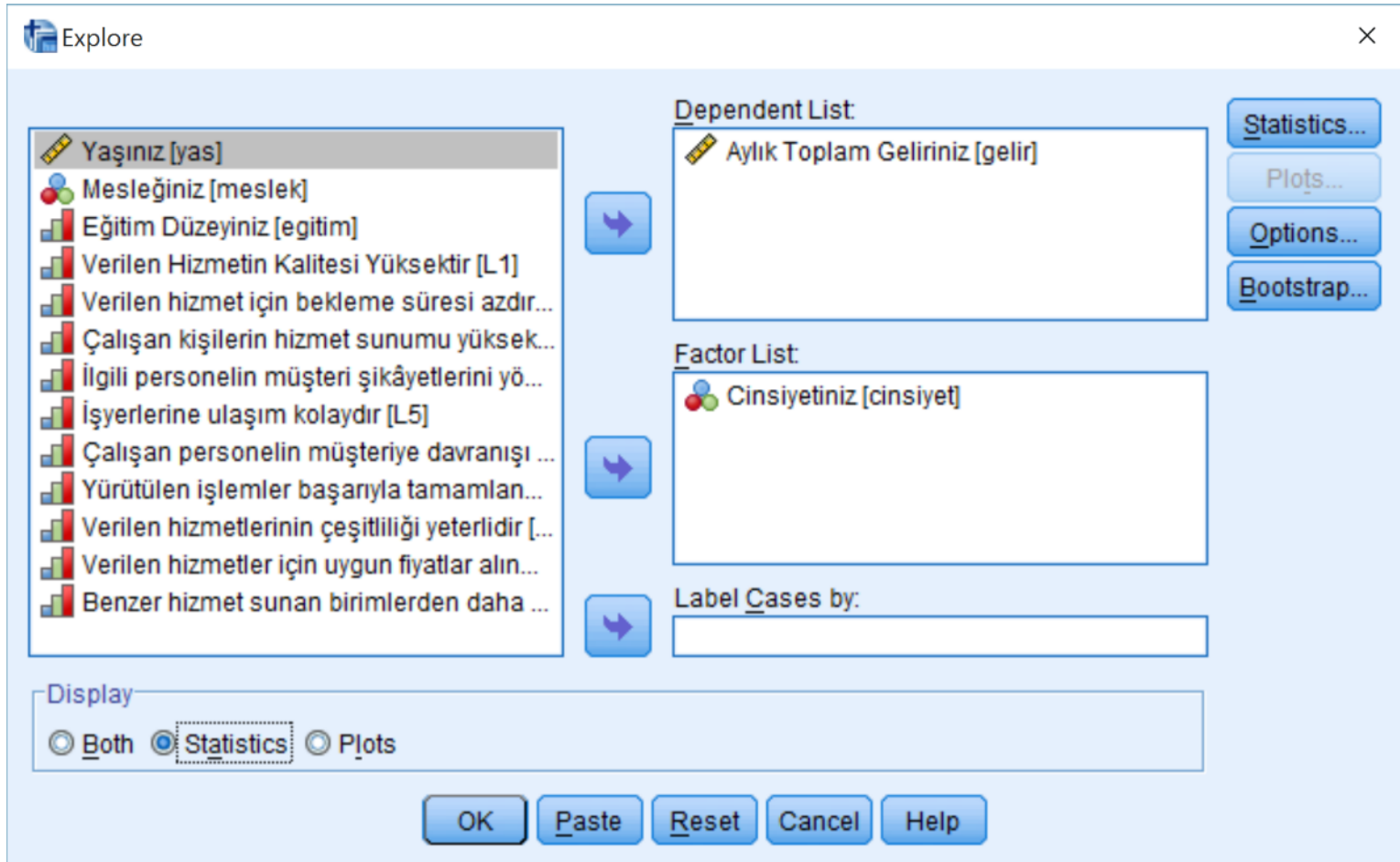
4,00

650,00 4,00 4,00

363,00 5,00 5,00

608,00 4,00 2,00

154,00 4,00 4,00



Dependent variable kısmına nicel bir değişken ve *factor list* kısmına ise sınıflama veya sıralama ölçme düzeyinde ölçülmüş nitel bir değişken atılır.

Descriptives

Cinsiyetiniz			Statistic	Std. Error
Aylık Toplam Geliriniz	Erkek	Mean	2059,0812	72,01202
		95% Confidence Interval for Mean	Lower Bound Upper Bound	
			1917,3047 2200,8576	
		5% Trimmed Mean	1964,6632	
		Median	1770,0000	
		Variance	1405332,956	
		Std. Deviation	1185,46740	
		Minimum	500,00	
		Maximum	5500,00	
		Range	5000,00	
		Interquartile Range	1622,00	
		Skewness	1,112	,148
		Kurtosis	,685	,295
	Kadın	Mean	1891,9737	107,34656
		95% Confidence Interval for Mean	Lower Bound Upper Bound	
			1679,3008 2104,6466	
		5% Trimmed Mean	1811,2934	
		Median	1392,5000	
		Variance	1313654,274	
		Std. Deviation	1146,14758	
		Minimum	500,00	
		Maximum	5500,00	
		Range	5000,00	
		Interquartile Range	1688,25	
		Skewness	,945	,226
		Kurtosis	,100	,449

Yukarıdaki tabloda hem erkek hem de kadınlar için ayrı ayrı aylık toplam gelire ait merkezi eğilim ve dağılım ölçütleri yer almaktadır.

DOĞRU-YANLIŞ SORULARI

	Doğru	Yanlış
Meslek değişkeni için hesaplanabilen tek merkezi eğilim ölçütü tepe değeridir (mod).		
Mod ve medyan aykırı (uç) değerlerden etkilenmez		
Aritmetik ortalama birimlerin tamamını kullanmaz		
Sağa çarpık bir dağılımda mod değeri en büyüktür		
Sola çarpık bir dağılımda birimlerin çoğunluğu ortalamadan daha büyük değer alır		
Standart sapmanın birimi yoktur (birimden bağımsızdır)		

TEST SORULARI

S-1) Aşağıdaki merkezi eğilim ölçütlerinden hangisi (hangileri) aykırı değerlerden etkilenmez?

A) Mod B) Medyan C) Aritmetik ortalama D) Mod ve medyan

S-2) Boy uzunluğu (cm) için hesaplanan standart sapma için birim aşağıdakilerden hangisidir?

A) Birimi yoktur B) cm C) cm^2 D) $cm^{1/2}$

S-3) Çarpıklık katsayısı -0.5 ve basıklık katsayısı 0.8 iken aşağıdaki yorumlardan hangisi doğrudur?

A) Sola çarpık ve basık B) Sağa çarpık ve sivri

C) Sola çarpık ve sivri D) Sola çarpık ve normal

S-4) Herhangi bir değişkenin zamana bağlı olarak izlediği değişkenlik oranı hangi merkezi eğilim ölçütü ile hesaplanmalıdır?

A) Aritmetik B) Geometrik C) Harmonik D) Medyan

NİCEL VERİ ANALİZİ

HİPOTEZ TESTLERİ

- ❖ Geçerliliği olasılık esaslarına göre araştırılabilen varsayımlara *istatistiksel hipotez* denir.
- ❖ Hipotezin yığından seçilen bir örneklem vasıtasıyla incelenmesine *hipotez testi* adı verilir.
- ❖ Hipotez testinde işlem, parametrenin belirli bir değere eşit olduğu şeklindeki “*sıfır hipotezi* (H_0)” ile parametrenin belirli bir değerden küçük, büyük veya farklı olduğu şeklindeki *alternatif hipotezden* (H_1) hangisinin örnekle daha iyi bağdaştığını göstermektir.
- ❖ Hipotez testinde sıfır hipotezinin doğru olduğu varsayımı ile teste başlanır.
- ❖ Örnekten elde edilen sonuç, sıfır hipotezi doğru iken büyük bir olasılıkla karşılaşılabilecek bir sonuçsa sıfır hipotezi kabul edilir. Buna karşın, örneğin verdiği sonuç sıfır hipotezi doğru iken çok zayıf bir olasılıkla karşımıza çıkabilecek bir sonuçsa, sıfır hipotezi reddedilir ve dolayısıyla alternatif hipotez kabul edilir.

❖ Bir hipotez testi sonucunda verilen karar için aşağıdaki dört durumdan birisi söz konusu olacaktır.

		Karar	
		H_0 Red Edildi	H_0 Red Edilemedi
H_0 Gerçekte	Doğru	Yanlış Karar (I. Tip Hata= α)	Doğru Karar (Güven Düzeyi= $1 - \alpha$)
	Yanlış	Doğru Karar (Testin Gücü= $1 - \beta$)	Yanlış Karar (II. Tip Hata= β)

Uygulamada α için genellikle %10, %5 veya %1 değerleri seçilir. α ile β arasında *ters yönde* bir ilişki vardır. α araştırmacı tarafından belirlendikten sonra β 'nın değerinin testi yapan tarafından belirlenmesi mümkün değildir. Burada yapılacak işlem α belli iken β 'yı minimum yapacak yollara başvurmaktır.

		Karar	
		H_0 Red Edildi (Suçlu)	H_0 Red Edilemedi (Masum)
H_0	Masum (Doğru)	Yanlış Karar (I. Tip Hata= α)	Doğru Karar (Güven Düzeyi= $1 - \alpha$)
	Suçlu (Yanlış)	Doğru Karar (Testin Gücü= $1 - \beta$)	Yanlış Karar (II. Tip Hata= β)

β 'yı Etkileyen Faktörler

- ❖ Hipotezdeki parametre değeri ile parametrenin gerçek değeri arasındaki fark arttıkça β da artar.
- ❖ Anlamlılık düzeyi α azalırken β artar.
- ❖ Yığın standart sapması σ arttıkça β artar.
- ❖ Örnek hacmi n azaldıkça β artar.

Hipotez testlerinde izlenmesi gereken aşamalar şöyledir:

- ❖ Hipotezlerin kurulması
- ❖ İlgili istatistiğin örnekleme dağılımı üzerinde ret ve kabul bölgelerinin belirlenmesi
- ❖ Test istatistiğinin hesaplanması
- ❖ Karar

Hipotezlerin Kurulması

- ❖ Hipotez testlerinde H_0 hipotezinin doğru olduğu varsayımı yapılır.
- ❖ Hipotez yığının parametresine ilişkin ortaya atılan bir iddiadır.
- ❖ H_0 hipotezi parametrenin belirli bir değere eşitliğini ifade etmektedir.

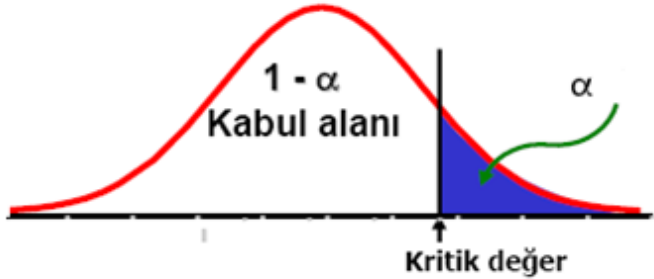
Örneğin “A-bölgesindeki hane-halklarının ortalama aylık geliri en az 2000 TL’dir” iddiasında parametre yığının ortalaması μ olup hipotezler aşağıdaki gibi kurulur:

$$H_0: \mu = 2000$$

$$H_1: \mu > 2000$$

İlgili İstatistiğin Örneklem Dağılımı Üzerinde Ret ve Kabul Bölgelerinin Belirlenmesi

İlgili istatistiğin örneklem dağılımı üzerinde ret ve kabul bölgeleri α anlamlılık düzeyi ve alternatif hipotez H_1 dikkate alınarak belirlenir.

$H_1: \theta > \theta_0 \Rightarrow$	
$H_1: \theta < \theta_0 \Rightarrow$	
$H_1: \theta \neq \theta_0 \Rightarrow$	

Test İstatistiğinin Hesaplanması

Ortalama için hipotez testinde “test istatistiği” standart normal değişken (Z) ya da serbestlik derecesi $n-1$ olan student t-değişkeni olmaktadır. Yığın oranına ilişkin hipotez testinde test istatistiği standart normal değişken (Z) iken yığın varyansı için bu istatistik serbestlik derecesi $n-1$ olan ki-kare değişkeni olmaktadır.

Karar Kuralı

İlgili istatistik için hesaplanan test istatistiğinin değeri ret bölgesinde ise H_0 hipotezi α anlamlılık düzeyinde reddedilebilir kararına varılır. Bu sonuç dolaylı olarak H_1 hipotezinin kabul edilmesi anlamına gelmektedir. Buna karşın test istatistiğinin değeri kabul bölgesine düşerse H_0 hipotezi seçilen anlamlılık düzeyinde reddedilemez.

Olasılık değeri (*p-değeri*) sıfır hipotezinin doğruluğu için hesaplanan olasılık değerine karşılık gelmektedir. Sıfır hipotezi doğru iken seçilen örnekten bulunan test istatistiğine dayalı olarak hesaplanan *p-değeri*, doğru olan sıfır hipotezini yanlışlığa düşerek reddetme olasılığıdır. O halde bu olasılık değerinin 1'e yakın olması sıfır hipotezini destekleyen bir kanıt iken o'a yaklaşması H_0 hipotezinin geçersiz olduğu anlamına gelmektedir.

$$H_1: \mu > \mu_0 \Rightarrow$$

$$p - \text{değeri} = P(t_{n-1} > t_h)$$

$$H_1: \mu < \mu_0 \Rightarrow$$

$$p - \text{değeri} = P(t_{n-1} < t_h)$$

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

$$p - \text{değeri} = 2 \times P(t_{n-1} > |t_h|)$$

Hesaplanan *p-değerine* göre karar kuralı şu şekildedir:
p-değeri < α ise H_0 hipotezi reddedilir.

COMPARE MEANS MENÜSÜ

Analyze menüsünün bir alt menüsü olan *Compare Means* ile veri setinde yer alan değişkenlerin bağımlı ve bağımsız olarak ayrılmasıyla betimleyici istatistiklerin (Means) hesaplanması, tek örneklem t testi (One-Sample T Test), bağımsız iki örneklem t testi (Independent-Samples T Test), eşleştirilmiş t testi (Paired-Samples T Test) ve tek yönlü varyans analizi (One-Way ANOVA) uygulamaları yapılmaktadır.

Means Menüsü

veri.sav [DataSet1] - IBM SPSS Statistics Data Editor

File Edit View Data Transform Analyze Direct Marketing Graphs Utilities Add-ons Window Help

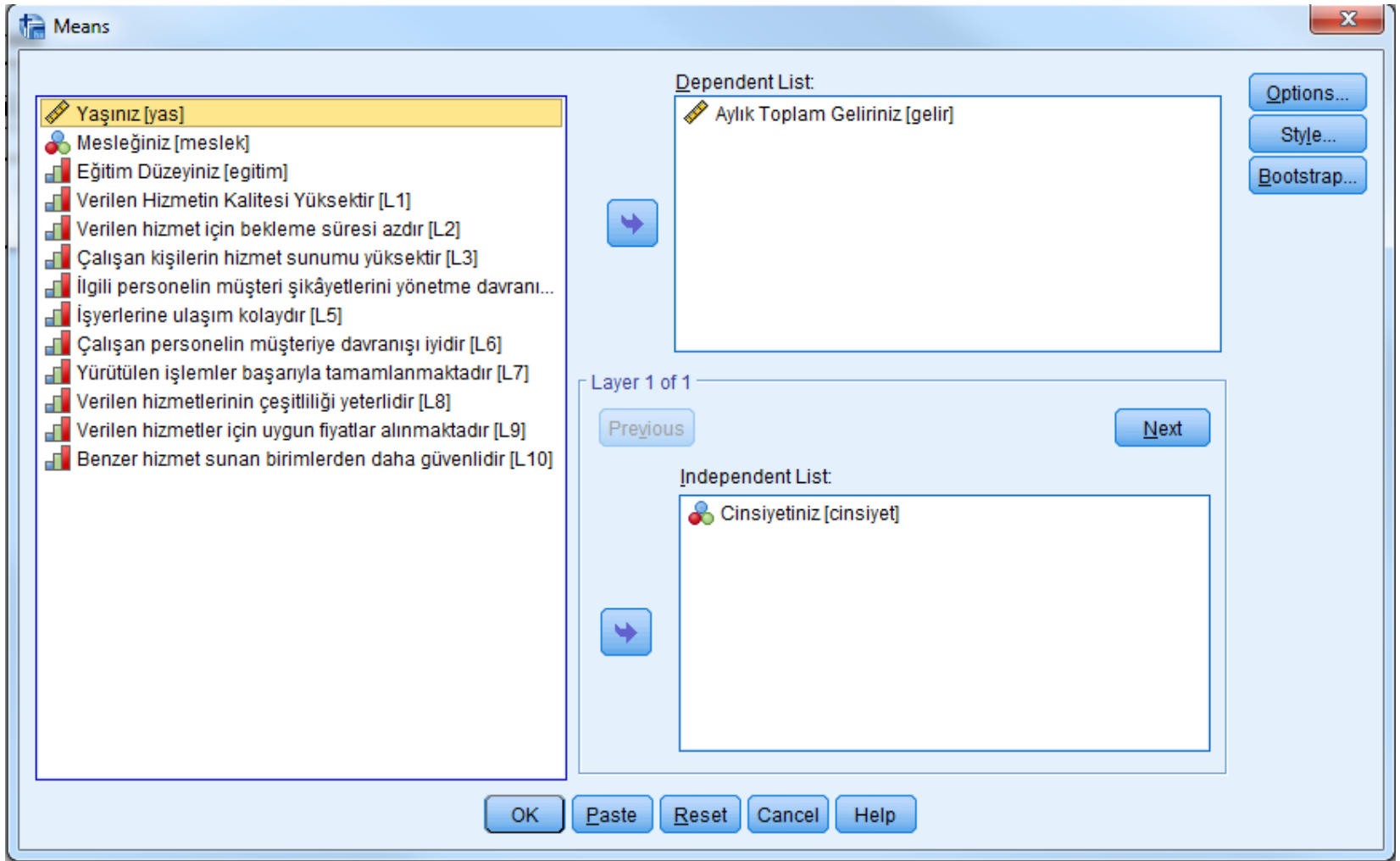
9 : L1 4,00

	cinsiyet	yas
1	2,00	32,00
2	2,00	23,00
3	1,00	25,00
4	1,00	26,00
5	2,00	20,00
6	2,00	21,00
7	2,00	32,00
8	1,00	67,00
9	1,00	50,00
10	2,00	24,00

Reports
Descriptive Statistics
Custom Tables
Compare Means
General Linear Model
Generalized Linear Models
Mixed Models
Correlate
Regression
Loglinear
Neural Networks
Classify
Dimension Reduction
Scale

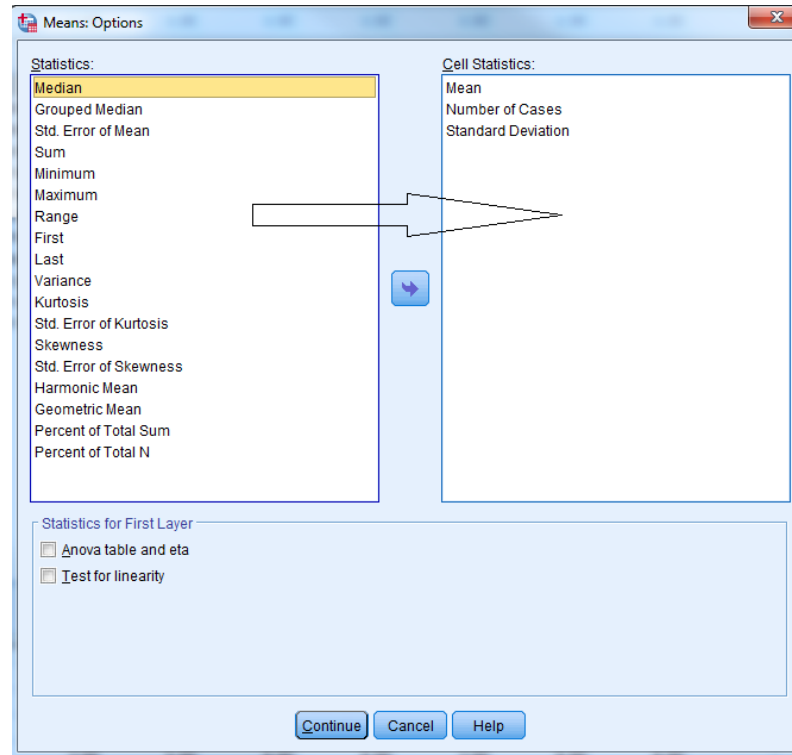
M Means...
t One-Sample T Test...
t Independent-Samples T Test...
+ Summary Independent-Samples T Test
t Paired-Samples T Test...
F One-Way ANOVA...

5,00	5,00	4,00	4,00
4,00	2,00	4,00	4,00
4,00	4,00	4,00	4,00
4,00	2,00	4,00	5,00



Bağımlı değişken aylık toplam gelir nicel değişkeni ve bağımsız değişken ise cinsiyet nitel değişkeni olmak üzere kadın ve erkekler için aylık toplam gelire göre ayrı ayrı betimleyici istatistikler elde edilir.

Means → *Options* kısmına girildiğinde istenilen betimleyici istatistikler seçilir.



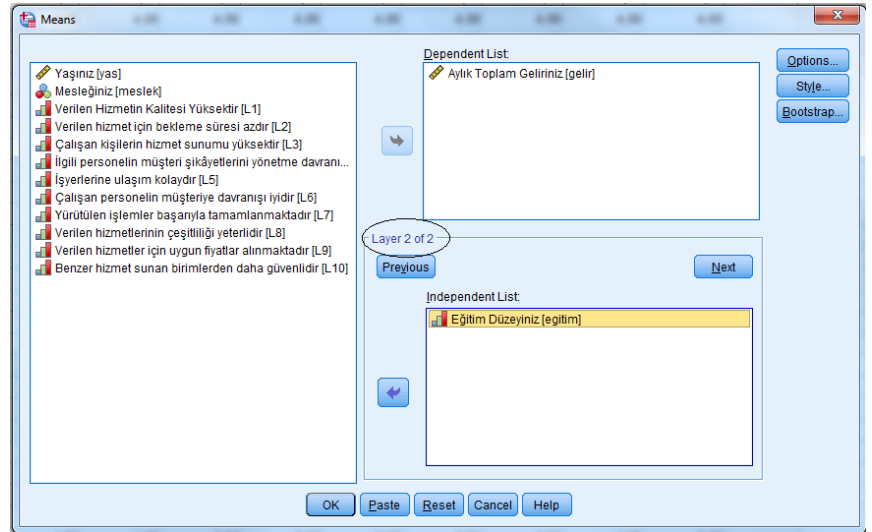
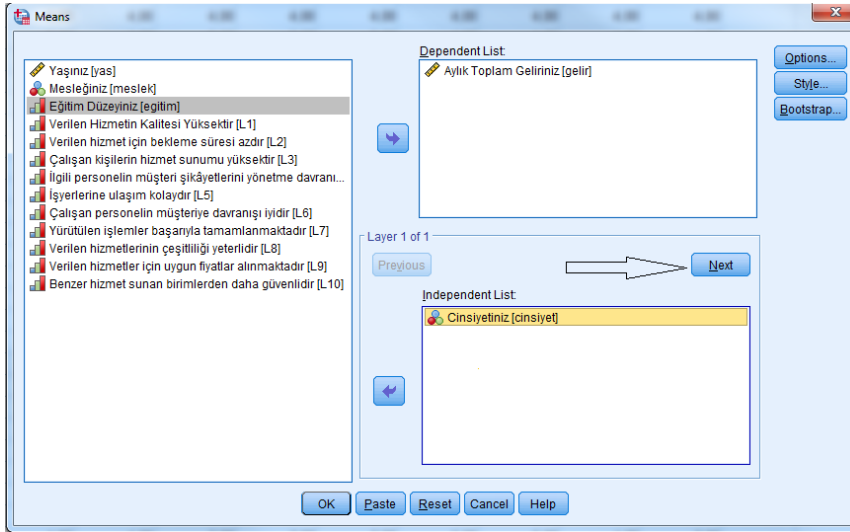
Report

Aylık Toplam Geliriniz

Cinsiyetiniz	Mean	Median	Std. Deviation	Variance	Geometric Mean	Harmonic Mean	Kurtosis	Skewness
Erkek	2059,0812	1770,0000	1185,46740	1405332,956	1761,6584	1511,3646	,685	1,112
Kadın	1891,9737	1392,5000	1146,14758	1313654,274	1585,2034	1338,3546	,100	,945
Total	2009,6000	1672,0000	1174,95971	1380530,319	1707,4550	1455,6460	,542	1,062

Bu tablo *explore* menüsündeki tablonun bir benzeridir. Farklı olarak geometrik ve harmonik ortalama değerleri tabloda yer almaktadır.

Means menüsü yardımıyla iç içe bağımsız değişkenler içinde betimsel istatistikler elde edilebilir. Bu işlem şöyle yapılır.



Kadın ve erkeklerin eğitim düzeylerine göre tabakalara ayrılarak aylık toplam gelirin betimsel istatistiklerini elde etmek için *Means* penceresinde *Independent List* kısmında *Layer* kullanılır. Soldaki şekilde Cinsiyet değişkeni *Independent List* kısmına atılır ve Next'e basılır. Sağdaki şekilde ise *Independent List* kısmına Eğitim Düzeyi değişkeni atılarak sonuçlar aşağıdaki gibi bulunur.

Report

Aylık Toplam Geliriniz

Cinsiyetiniz	Eğitim Düzeyiniz	Mean	Median	Std. Deviation	Variance	Geometric Mean	Harmonic Mean	Kurtosis	Skewness
Erkek	İlköğretim	1830,7500	1557,5000	1013,02468	1026219,000	1597,5858	1404,3119	1,583	1,293
	Lise	1931,2927	1691,0000	1001,81092	1003625,110	1702,2283	1498,1784	1,448	1,174
	Üniversite	2558,8529	2226,5000	1509,97906	2280036,754	2102,9561	1689,8166	-,912	,524
	Total	2059,0812	1770,0000	1185,46740	1405332,956	1761,6584	1511,3646	,685	1,112
Kadın	İlköğretim	1630,9355	1260,0000	950,41114	903281,329	1405,3727	1224,7591	1,434	1,247
	Lise	1847,9184	1359,0000	1049,69874	1101867,452	1570,4878	1334,6330	-,725	,683
	Üniversite	2193,4706	1795,0000	1383,41661	1913841,529	1793,0886	1468,4347	-,497	,758
	Total	1891,9737	1392,5000	1146,14758	1313654,274	1585,2034	1338,3546	,100	,945
Total	İlköğretim	1774,9459	1447,0000	995,72622	991470,706	1541,4022	1349,0766	1,492	1,270
	Lise	1907,5407	1642,5000	1013,29028	1026757,185	1663,6105	1447,6419	,754	1,008
	Üniversite	2437,0588	2190,0000	1472,33676	2167775,541	1994,1327	1608,9607	-,803	,598
	Total	2009,6000	1672,0000	1174,95971	1380530,319	1707,4550	1455,6460	,542	1,062

Sonuçlardan da görüldüğü üzere hem erkekler hem de kadınlar eğitim düzeyine göre tabakalanarak betimsel istatistikler ayrı ayrı elde edilir.

One-Sample T-Test Menüsü (Tek Örneklem İçin T-Testi)

Yığın ortalaması μ için ortaya atılan iddialar için hipotez testinde sıfır ve alternatif hipotezler aşağıdaki gibi yazılır:

$$H_0: \mu = \mu_0 \text{ iken}$$

$$H_1: \mu < \mu_0$$

$$H_1: \mu > \mu_0$$

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

Burada μ_0 herhangi bir reel sayıdır. Yığından “ n ” birimlik rastgele örnek seçildiğinde, seçilen örnekten hareketle hesaplanan ortalama

$$\bar{X} = \frac{\sum_{i=1}^n X_i}{n}$$

olacaktır.

Yığındaki aylık ortalama gelirin 2200 TL'ye eşit olup olmadığı iddiası test edilmek istenmektedir. Bu amaçla hipotezler aşağıdaki gibi kurulur.

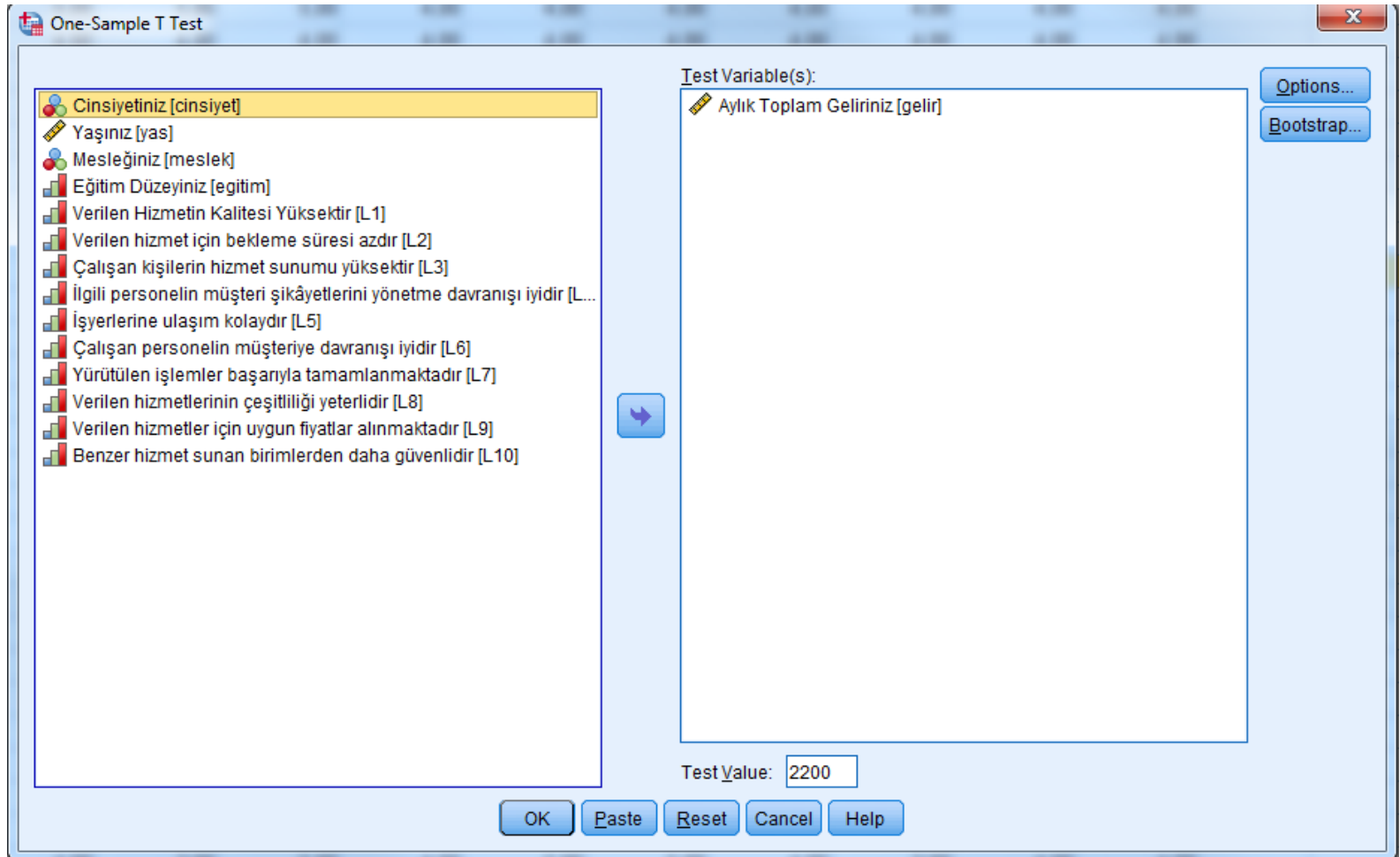
$$H_0: \mu = 2200$$

$$H_1: \mu \neq 2200$$

Test istatistiği serbestlik derecesi «n-1» olan t-değişkenidir. Aşağıdaki formül ile t-istatistiği hesaplanır:

$$t = \frac{\bar{X} - \mu}{S / \sqrt{n}}$$

SPSS'de bu hipotezleri test edelim.



One Sample T-Test penceresinde *Test Variable(s)* kısmına Aylık Toplam Gelir, *Test Value* kısmına ise μ_0 değeri yazılır ve Daha sonra OK'e basılır.

One-Sample Statistics

	N	Mean	Std. Deviation	Std. Error Mean
Aylık Toplam Geliriniz	385	2009,6000	1174,95971	59,88149



n



$\bar{X}$

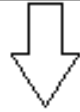


s



$s/\sqrt{n}$

	Test Value = 2200					
	t	df	Sig. (2-tailed)	Mean Difference	95% Confidence Interval of the Difference	
					Lower	Upper
Aylık Toplam Geliriniz	-3,180	384	,002	-190,40000	-308,1367	-72,6633



t_h



$n-1$



$p\text{-değeri}$



$\bar{X} - \mu_0$

Yukarıdaki tabloda verilen p-değeri alternatif hipotezin çift taraflı olduğu durumda geçerli olan olasılık değeridir.

$H_1: \mu \neq 2200$ alternatif hipotezinin testi durumunda p-değeri=0.002 < $\alpha = 0.05$ olduğundan H_0 hipotezi reddedilir.

Independent-Samples T-Test Menüsü (Bağımsız Örneklerde T-Testi)

Bağımsız iki grubun yığın ortalamalarını karşılaştırmak için kullanılır. $X_1 \sim N(\mu_1, \sigma_1^2)$ ve $X_2 \sim N(\mu_2, \sigma_2^2)$ bağımsız rastgele değişkenler olmak üzere iki yığının ortalamalarına ilişkin ortaya atılan iddialar için hipotezler

$$H_0: \mu_1 = \mu_2 \rightarrow H_0: \mu_1 - \mu_2 = 0$$

$$H_1: \mu_1 < \mu_2 \rightarrow H_1: \mu_1 - \mu_2 < 0$$

$$H_1: \mu_1 > \mu_2 \rightarrow H_1: \mu_1 - \mu_2 > 0$$

$$H_1: \mu_1 \neq \mu_2 \rightarrow H_1: \mu_1 - \mu_2 \neq 0$$

Durum	Test İstatistiği	Serbestlik Derecesi
σ_1^2 ve σ_2^2 bilinmiyor ($\sigma_1^2 = \sigma_2^2$)	$t_h = \frac{(\bar{X}_1 - \bar{X}_2)}{\sqrt{s_p^2 \left(\frac{1}{n_1} + \frac{1}{n_2} \right)}}$ $s_p^2 = \frac{s_1^2(n_1 - 1) + s_2^2(n_2 - 1)}{n_1 + n_2 - 2}$	$n_1 + n_2 - 2$
σ_1^2 ve σ_2^2 bilinmiyors a ($\sigma_1^2 \neq \sigma_2^2$)	$t_h = \frac{(\bar{X}_1 - \bar{X}_2)}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}}$	$\frac{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}{\frac{s_1^4}{n_1^2(n_1 - 1)} + \frac{s_2^4}{n_2^2(n_2 - 1)}}$

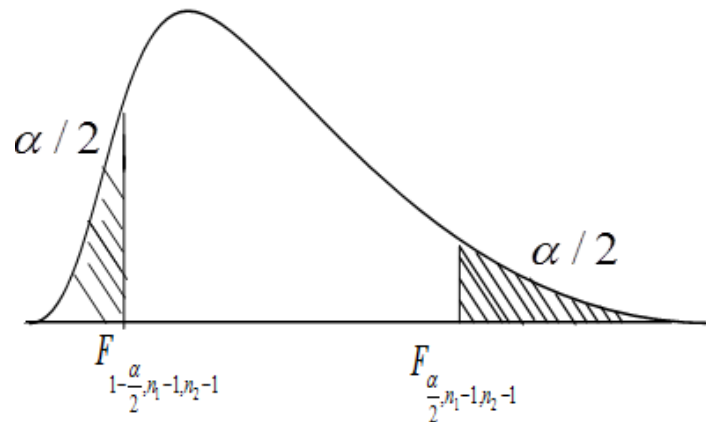
İki grup ortalamasının hipotez testinde ilk aşama yığın grup varyanslarının eşit olup olmadığının test edilmesi gereklidir. Buna ilişkin hipotezler

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

$$H_1: \sigma_1^2 \neq \sigma_2^2$$

H_0 hipotezi altında test istatistiği

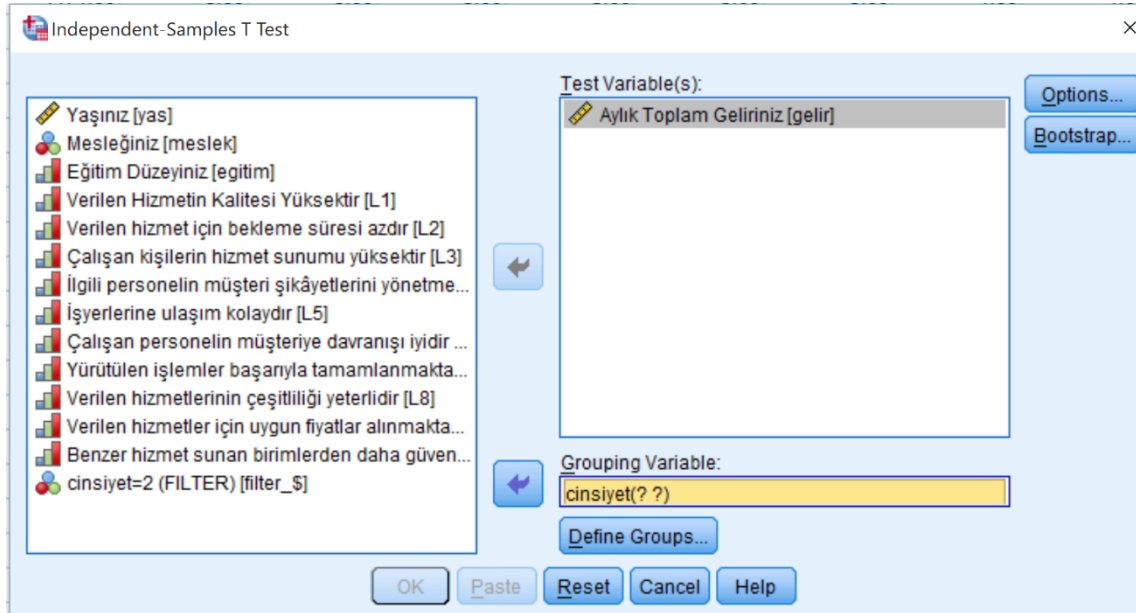
$$F_h = \frac{s_1^2}{s_2^2}$$



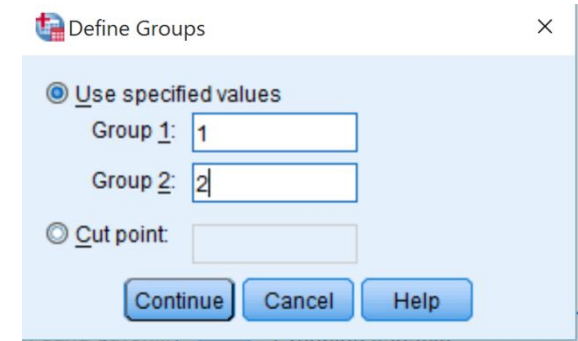
Kadın ve Erkeklerin yığındaki aylık ortalama gelirleri açısından farklılık gösterip göstermediklerini test etmek istiyoruz. Bu durumda hipotezler aşağıdaki gibidir:

$$H_0: \mu_1 = \mu_2 \rightarrow H_0: \mu_1 - \mu_2 = 0$$

$$H_1: \mu_1 \neq \mu_2 \rightarrow H_1: \mu_1 - \mu_2 \neq 0$$



Yandaki pencerede *Grouping Variable* kısmında *Define Groups* 'a girilir ve aşağıdaki gibi tanımlanır.



Group Statistics

	Cinsiyetiniz	N	Mean	Std. Deviation	Std. Error Mean
Aylık Toplam Geliriniz	Erkek	271	2059,0812	1185,46740	72,01202
	Kadın	114	1891,9737	1146,14758	107,34656

$$\bar{X}_1$$

$$\bar{X}_2$$

$$s_1$$

$$s_2$$

$$\frac{s_2}{\sqrt{n_2}}$$

$$\frac{s_1}{\sqrt{n_1}}$$

		Levene's Test for Equality of Variances	
		F	Sig.
Aylık Toplam Geliriniz	Equal variances assumed	,287	,592
	Equal variances not assumed		

Yandaki tabloda $H_0: \sigma_1^2 = \sigma_2^2$ hipotezinin testine ait sonuçlar verilmiştir. Sig değeri olasılık değerine karşılık gelmektedir.

$Sig = 0.592 > \alpha = 0.05$ olduğundan grup varyanslarının eşit olduğunu gösteren H_0 hipotezi reddedilememektedir.

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
Aylık Toplam Geliriniz	Equal variances assumed	,287	,592	1,275	383	,203	167,10750	131,05774	-90,57524	424,79024
	Equal variances not assumed			1,293	219,025	,197	167,10750	129,26335	-87,65171	421,86670

Levene testine göre grup varyansları eşit bulunduğundan dolayı $H_0: \mu_1 - \mu_2 = 0$ sıfır hipotezini $H_1: \mu_1 - \mu_2 \neq 0$ alternatif hipotezine karşı test etmek için yukarıdaki tablonun birinci satırı kullanılır.

t-test for Equality of Means				
t	df	Sig. (2-tailed)	Mean Difference	Std. Error Difference
1,275	383	,203	167,10750	131,05774
$\downarrow$ t_h	$\downarrow$ $n_1 + n_2 - 2$	$\downarrow$ $p - \text{değeri}$	$\downarrow$ $\bar{X}_1 - \bar{X}_2$	$\downarrow$ $s_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}$

$sig(2 - tailed) > \alpha = 0.05$ olduğundan H_0 sıfır hipotezi reddedilememektedir. Böylece erkek ve kadınların aylık ortalama harcamaları arasında fark olmadığı sonucuna varılır.

BAĞIMLI (EŞLEŞTİRİLMİŞ) ÖRNEKLERDE T-TESTİ (PAIRED SAMPLE T TEST) MENÜSÜ

Bağımlı örnekler için deney tasarımı iki farklı biçimde karşımıza çıkmaktadır. Bunlar;

1. İşlem öncesi-işlem sonrası tasarımı
2. Etkisi arındırılmak istenen değişkene göre eşleştirme tasarımı.

İşlem öncesi ve sonrası tasarımda rastgele seçilen bireyler üzerinden X değişkeni için gözlem değerleri toplandıktan sonra bu bireyler bir işleme tabi tutulmakta ve aynı bireylerin işlem sonrası gözlem değerleri de kaydedilmektedir. Bu tür araştırmalarda yapılan işlemin etkili olup olmadığı sorusuna cevap aranmaktadır.

Örneğin A işletmesinde çalışan memurlardan rastgele seçilen kişilerin iş verimliliği (üretim/hafta) şeklinde ölçülmüş olsun. Bu bireyler bir hizmet-içi eğitime tabi tutulduktan sonra tekrar iş verimliliği (üretim/hafta) ölçülmüş olsun. Hizmet içi eğitimin iş verimliliğini arttırdığı iddiası test edilecek ise bu test için uygun yöntem bağımlı (eşleştirilmiş) örneklerde t-testidir.

Ele alınan değişken üzerinde başka bir değişkenin etkisi arındırılmak istenebilir. Örneğin A ve B diyet yöntemlerinin zayıflama üzerindeki etkisini araştırır iken etkisi arındırılmak istenen değişken kişilerin kilosu olabilir. Bu durumda aynı kiloya sahip kişiler eşleştirildikten sonra bu kişiler rastgele olarak A ve B yöntemine atanır. Böylece her iki grupta da aynı kiloya sahip kişiler olduğundan zayıflama üzerinde kilonun etkisi arınmış olacaktır. Bu durumda da uygun yöntem bağımlı (eşleştirilmiş) örneklerde t-testi olacaktır.

- Yığından n gözlemlili rastgele bir örneğin seçildiği ve bu örneğe ilişkin $(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)$ eşleşmiş gözlemlerin elde edildiği varsayılınsın. Birinci ve ikinci grubun ortalamalarının eşit olduğunu iddiasının testinde, eşleşmiş (x_i, y_i) gözlem değerleri arasındaki fark $d_i = x_i - y_i$ olmak üzere sıfır hipotezi ve alternatif hipotezleri
- $H_0: \mu_d = 0$ iken $H_1: \mu_d < 0$ veya $H_1: \mu_d > 0$ veya $H_1: \mu_d \neq 0$ olacaktır.

Test istatistiği ise serbestlik derecesi $n - 1$ olan t değişkenidir.

$$t_h = \frac{\bar{d}}{S_d/\sqrt{n}} \sim t_{n-1}$$

Burada

$$\bar{d} = \frac{\sum d_i}{n} \text{ ve } S_d = \frac{1}{n-1} \sum_{i=1}^n (d_i - \bar{d})^2$$

Bu test için $p - \text{değeri}$ H_1 hipotezine bağlı olarak hesaplanır. Karar kuralı $p - \text{değeri} < \alpha$ ise H_0 hipotezi reddedilebilir şeklindedir.

Örnek 1. 2013 yılında yapılan araştırmada A kurumundan memnuniyet için hazırlanan 10 soruluk Likert ölçeğinin puanlanmasından elde edilen verilerin ortalaması alınarak her birey için genel memnuniyet skoru elde edilmiştir. Çalışma 2014 yılında aynı bireyler üzerinde uygulanmış ve genel memnuniyet skoru tekrar hesaplanmıştır. ‘*A kurumundan ortalama memnuniyet 2014 yılında bir önceki yıla göre artmıştır*’ iddiasını 0.05 anlamlılık düzeyinde test ediniz.

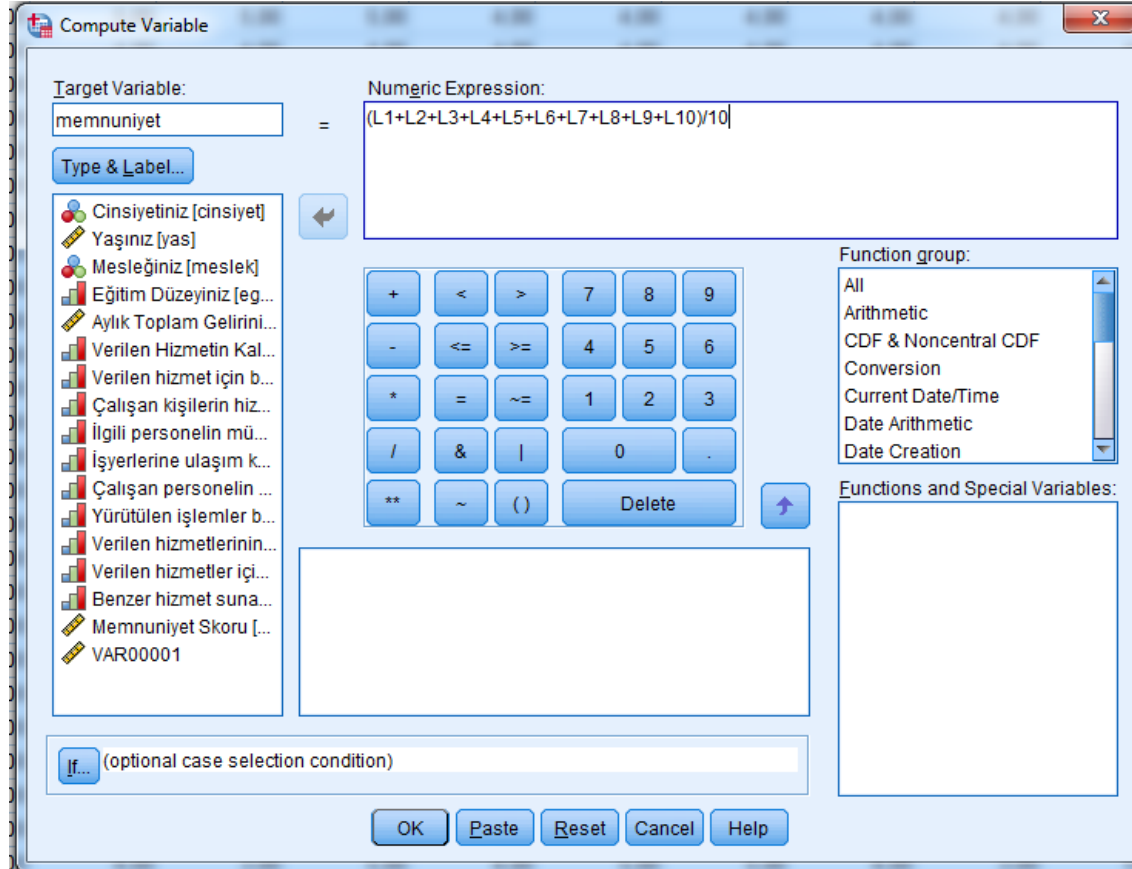
$$d_i = x_{2013i} - x_{2014i}$$

olmak üzere hipotezler:

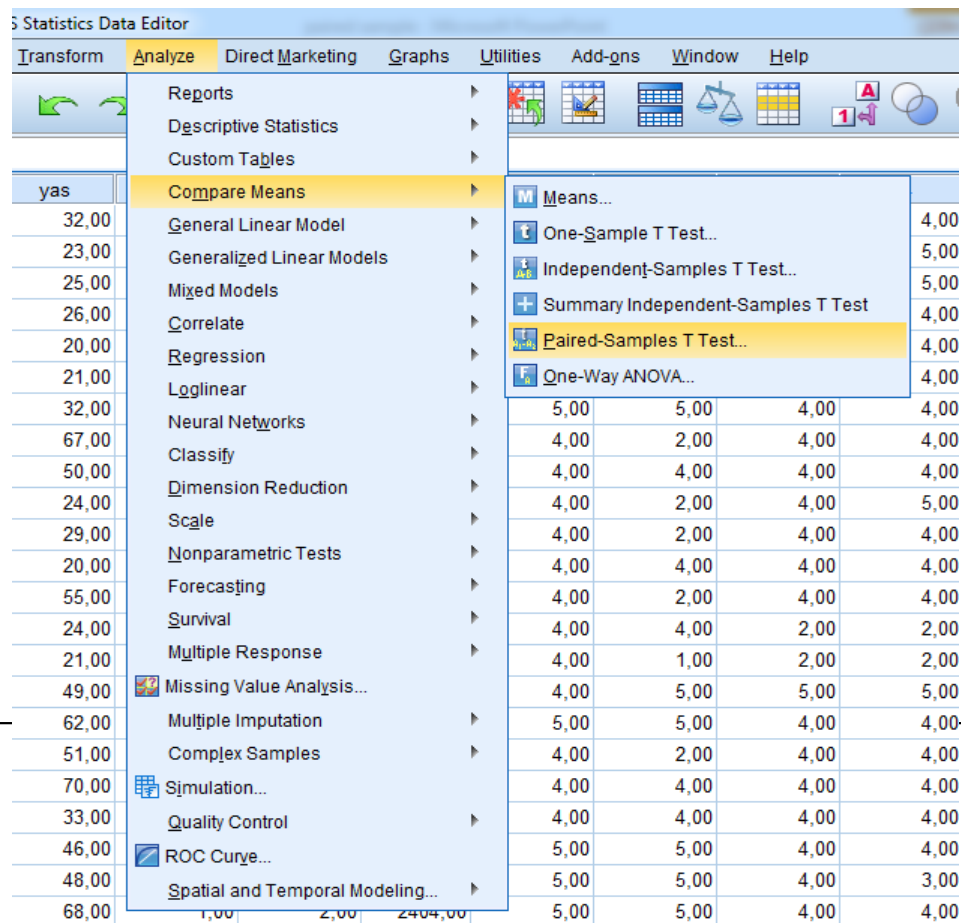
$$H_0: \mu_d = 0$$

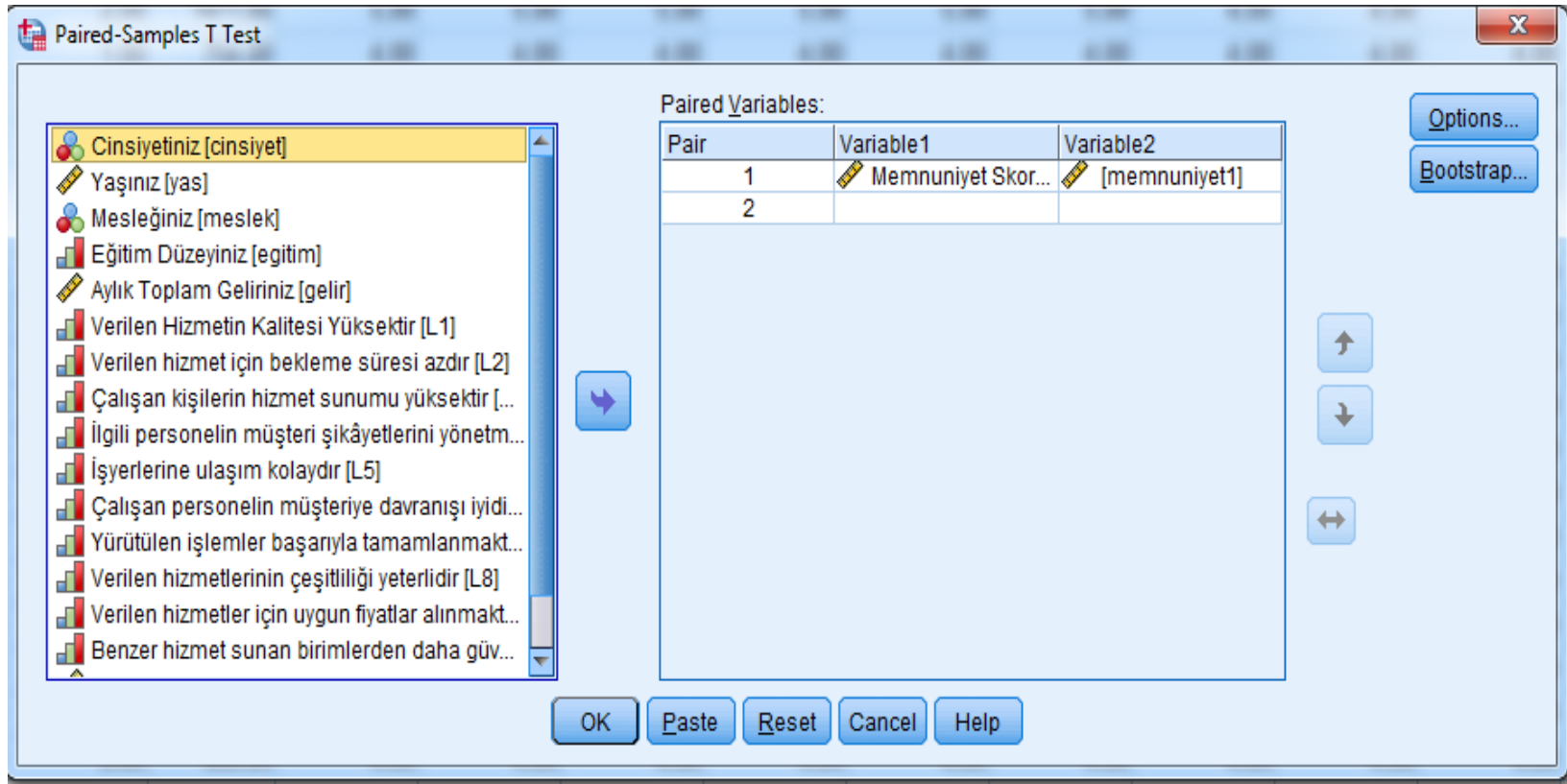
$$H_1: \mu_d < 0$$

SPSS'te *Transform* menüsünün altındaki *Compute Variable* alt menüsünden 2013 yılı için 385 kişinin ortalama memnuniyet skorları hesaplanır.



Yukarıdaki pencerede memnuniyet skoru oluşturmak için 10 Likert sorusunun ortalaması alınmıştır.





Yukarıdaki pencerede *Paired variables* kısmında *variable 1* ‘e memnuniyet skoru değişkeni, *variable 2*’e memnuniyet değişkeni atanır.

Paired Samples Test

		Paired Differences					t	df	Sig. (2-tailed)
		Mean	Std. Deviation	Std. Error Mean	95% Confidence Interval of the Difference				
					Lower	Upper			
Pair 1	Memnuniyet Skoru - memnuniyet1	-,42727	,58493	,02981	-,48589	-,36866	-14,333	384	,000

$$\begin{array}{ccccccc}
 \downarrow & & \downarrow & & \downarrow & & \downarrow \\
 \bar{d} = \bar{x}_{2013} - \bar{x}_{2014} & s_d & s_d / \sqrt{n} & & & & \frac{\bar{d}}{s_d / \sqrt{n}}
 \end{array}$$

Alternatif hipotez $H_1: \mu_d < 0$ şeklinde tek yönlü olduğundan $t_h < 0$ olmalıdır ve yukarıdaki tabloda t_h değerinin negatif olduğu görülmektedir. $\frac{p\text{-değeri}}{2} = 0.000 < 0.05$ olduğundan H_0 hipotezi 0.05 anlamlılık düzeyinde reddedilebilir. Bu sonuç A kurumunda ortalama memnuniyetin 2013 yılına göre 2014 yılında arttığını göstermektedir.

TEK FAKTÖRLÜ VARYANS ANALİZİ (ONE-WAY ANOVA) MENÜSÜ

Tek faktörlü varyans analizi, ikiden fazla grubun yığın ortalamalarının karşılaştırılmasında kullanılan bir yöntemdir ve iki bağımsız örneklem için t testinin 2’den fazla gruba genişletilmiş halidir.

Niçin grupların ortalamalarını ikişerli karşılaştırmak yerine varyans analizine ihtiyaç duyulmaktadır?

2’den fazla grubun ortalamalarının karşılaştırılması söz konusu ise çok sayıda t testinin kullanılması 1. tip hatanın artmasına yol açmaktadır.

Örneğin, 4 tane grubumuz olsun ve bu grupların yığın ortalamalarını ikişerli karşılaştırmak isteyelim. Bu durumda şayet t testi kullanacak olsaydık 1-2, 1-3, 1-4, 2-3, 2-4 ve 3-4 grupları için ayrı ayrı 6 tane t testi yapmamız gerekecekti. Her bir testte 1. tip hata yapma olasılığı %5 için güven düzeyi 0.95 (1. tip hata yapmama olasılığı) olmak üzere 6 test için 1. tip hata yapmama olasılığı

$$0.95 \times 0.95 \times 0.95 \times 0.95 \times 0.95 \times 0.95 = 0.735$$

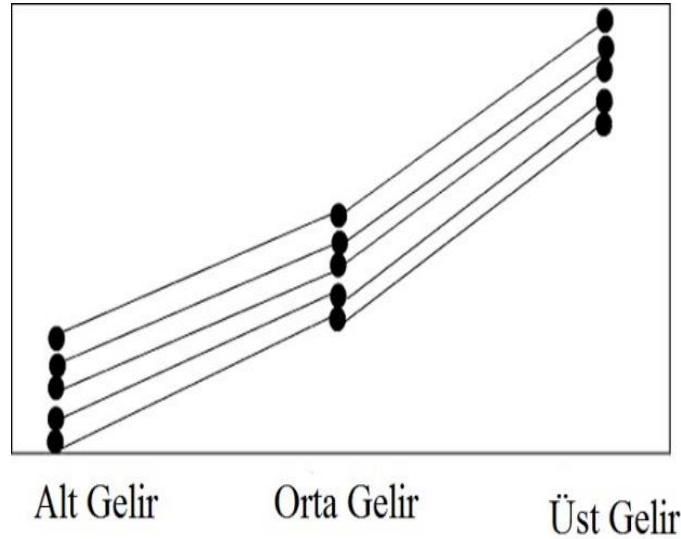
olur. 1. tip hata yapma olasılığı ise $1 - 0.735 = 0.265$ olarak bulunur. Böylece α anlamlılık düzeyi %5’den %26.5’e yükselmiş olur.

Çok sayıda t testi uygulamanın 1. Tip hata (gerçekte doğru olan bir sıfır hipotezini reddetme olasılığı) yapma olasılığını arttırdığına göre, bir çok gruba ait ortalamaları tek bir adımda eş zamanlı olarak karşılaştırabilecek bir yönteme ihtiyaç duyulmaktadır. Bu yöntem varyans analizi yöntemi olarak adlandırılır. Varyans analizi yönteminde bağımlı değişkenin ölçme düzeyi eşit aralıklı veya orantılı olmalıdır.

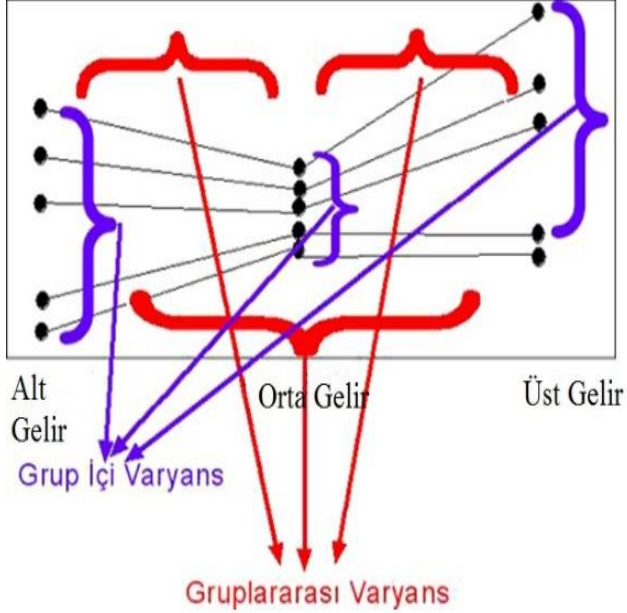
Tek faktörlü varyans analizinin varsayımları şunlardır:

- Her bir grubun seçildiği yığın normal dağılıma sahip olmalıdır.
- Her bir grubun seçildiği yığınların varyansları birbirine eşittir.
- Her bir gruptaki örnekler birbirinden bağımsızdır.

Örneğin üç farklı gelir grubunda kişilerin iş memnuniyetini ölçmeye çalışalım. Şayet gruplar arasında fark varsa, her bir gelir grubu kendi içinde küçük bir varyansa sahip olmalı yani memnuniyet puanları birbirine yakın olmalıdır. Ayrıca gruplardaki her bir bireyin memnuniyet puanı diğer gruptaki herhangi bir bireyden önemli ölçüde farklı olmalıdır.



Bazı durumlarda farklı gruplarda yer alan bireyler arasındaki farklılıktan çok, aynı grupta yer alan bireyler arasındaki farklılık daha büyük olabilir.



Burada grupların farklılığından çok, aynı grupta yer almasına rağmen birbirinden çok farklı memnuniyet puanına sahip bireyler söz konusudur.

Varyans analizinin yapmaya çalıştığı şey şudur: Gruplar arasındaki varyansı ve grupların kendi içlerindeki varyansı hesaplayarak birbirine oranlamak ve bu varyansların büyüklüklerine göre bir karar vermektir.

$\underline{G_1}$	$\underline{G_2}$	$\dots$	$\underline{G_k}$
x_{11}	x_{12}	$\dots$	x_{1k}
x_{21}	x_{22}	$\dots$	x_{2k}
x_{31}	x_{32}		x_{3k}
$\vdots$	$\vdots$	$\dots$	$\vdots$
$x_{n_1 1}$	$x_{n_2 2}$	$\dots$	$x_{n_k k}$

$$GAKT = \sum_{j=1}^k n_j (\bar{x}_{.j} - \bar{x})^2$$

$$TKT = \sum_{j=1}^k \sum_{i=1}^{n_j} (x_{ij} - \bar{x})^2$$

$$GIKT = \sum_{j=1}^k \sum_{i=1}^{n_j} (x_{ij} - \bar{x}_{.j})^2$$

x_{ij} : j . grupta i . birimin aldığı değer

$\bar{x}_{.j}$: j . gruptaki birimlerin ortalaması

n_j : j . gruptaki birimlerin sayısı

$\bar{x}$: Genel ortalama

ANOVA Tablosu

Değişimin Kaynağı	Serbestlik Derecesi	Kareler Toplamı	Kareler Ortalaması	F
Gruplar Arası	$k-1$	GAKT	$GAKO = GAKT / (k-1)$	$GAKO / GIKO$
Gruplar İçi	$n-k$	GIKT	$GIKO = GIKT / (n-k)$	
Toplam	$n-1$	TKT		

Tek faktörlü varyans analizinde ortalamalarının karşılaştırılmasında kullanılan hipotezler şöyledir:

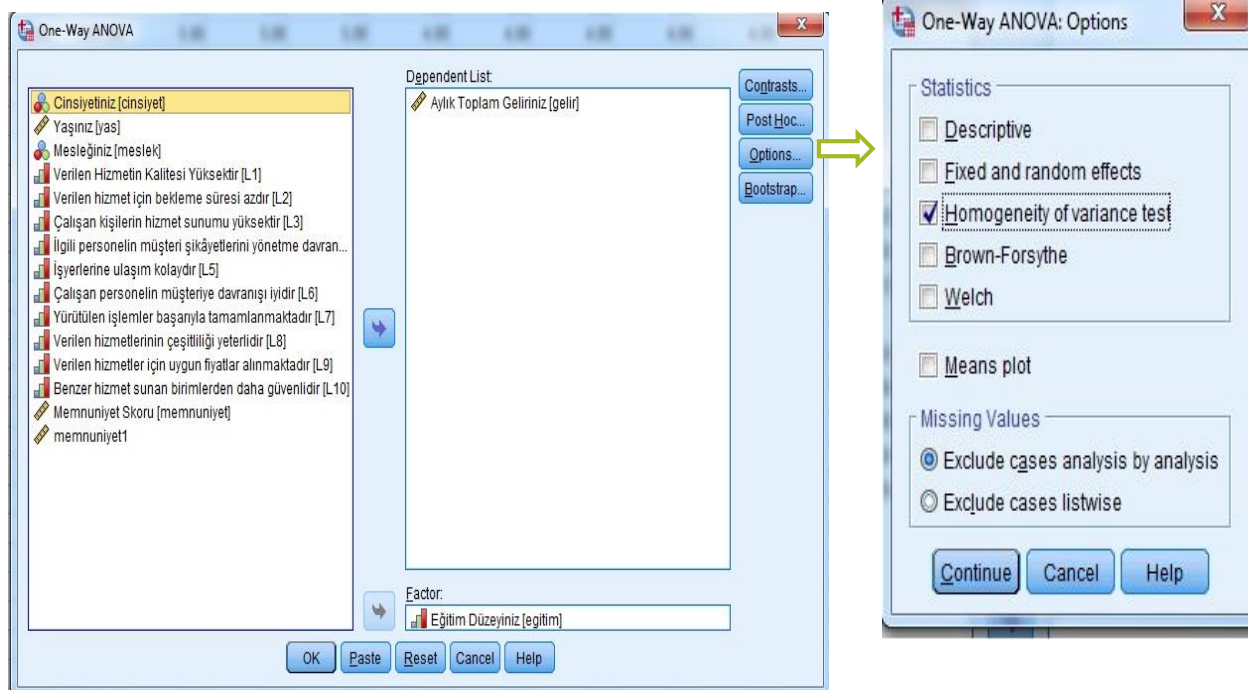
$$H_0: \mu_1 = \mu_2 = \dots = \mu_k$$
$$H_1: \mu'_j \text{lerin en az biri diğerlerinden farklıdır } j = 1, 2, \dots, k$$

ANOVA tablosuna göre

Şayet $F > F_{(k-1), (n-k), \alpha}$ ise H_0 hipotezi reddedilir. Yani gruplar yığın ortalamalarına göre farklılık göstermektedir.

Varyans analizinin önemli varsayımlarından birisi grupların geldikleri yığın varyanslarının eşit olduğudur. Gruplara ait yığın varyanslarının eşit olmadığı durumda F testi yerine Welch veya Brown-Forsythe tarafından önerilen ve asimptotik olarak F dağılıma sahip testler kullanılır. Welch testi Brown-Forsythe testinden daha güçlüdür.

Örnek 2. Eğitim düzeyine göre yığındaki aylık ortalama gelirler farklılaşmakta mıdır?



İlk olarak her bir eğitim düzeyinin yığın varyanslarının eşitliği test edilir. *Options* penceresinin altından *Homogeneity of variance test* kısmı seçilir.

Grupların yığın varyanslarının eşitliğinin testi için hipotezler şöyledir:

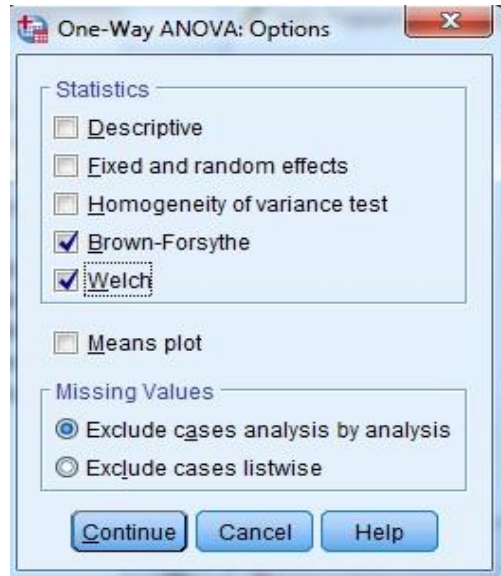
$$H_0: \sigma_1^2 = \sigma_2^2 = \sigma_3^2$$
$$H_1: \sigma_i^2 \neq \sigma_j^2 \text{ bazı } i \text{ ve } j' \text{ ler için}$$

Test of Homogeneity of Variances

Aylık Toplam Geliriniz

Levene Statistic	df1	df2	Sig.
16,916	2	382	,000

$p - \text{değeri}(\text{sig}) < \alpha = 0.05$ olduğundan yığın varyanslarının eşit olduğu H_0 hipotezi reddedilir.



Yığın varyansları eşit olmadığından yığın ortalamalarının karşılaştırılmasında F testi kullanılamaz. F testi yerine Welch veya Brown-Forsythe testlerinden birisi tercih edilir. Bu işlem yandaki *Options* penceresinden yapılır.

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

$$H_1: \mu'_j \text{lerin en az biri diğerlerinden farklıdır } j = 1, 2, 3$$

Robust Tests of Equality of Means

Aylık Toplam Geliriniz

	Statistic ^a	df1	df2	Sig.
Welch	7,402	2	212,221	,001
Brown-Forsythe	9,257	2	260,527	,000

a. Asymptotically F distributed.

Welch testi sonucuna göre $p - \text{değeri} = 0.001 < \alpha = 0.05$ olduğundan yığın ortalamalarının eşit olduğu H_0 hipotezi reddedilebilir. Yani eğitim düzeyine göre yığında aylık ortalama gelirler %5 anlamlılık düzeyinde farklılaşmaktadır sonucuna varılır. Aynı sonuç Brown-Forsythe testinde de bulunmuştur.

DOĞRU-YANLIŞ SORULARI

	Doğru	Yanlış
Birinci ve ikinci tip hata arasında ilişki yoktur		
Geçerliliği olasılık esaslarına göre araştırılabilen varsayımlara istatistiksel hipotez denir		
Bağımsız örneklerde t-testi, ikiden fazla grubun yığın ortalamalarının karşılaştırılmasında kullanılan bir yöntemdir		
Anlamlılık düzeyi (I. Tip hata) α azalırken ikinci tip hata β artar		
İşlem öncesi-işlem sonrası tasarımında iki yığın ortalamasının eşit olduğu iddiası eşleştirilmiş örneklerde t-testi ile araştırılır		
Sıfır hipotezi (H_0) parametrenin belirli bir değerden farklı olduğunu ifade eder		

TEST SORULARI

S-1) İki den fazla grubun ortalamaları karşılaştırıldığında kullanılan yöntem aşağıdakilerden hangisidir?

- A) Bağımsız örneklerde t-testi B) Eşleşmiş örneklerde t-testi
C) Tek örneklem t-testi D) Tek faktör varyans analizi

S-2) I. tip hata aşağıdakilerden hangisine karşılık gelir?

- A) Gerçekte doğru olan sıfır hipotezini reddetme olasılığı
B) Gerçekte doğru olan sıfır hipotezini kabul etme olasılığı
C) Gerçekte yanlış olan sıfır hipotezini reddetme olasılığı
D) Gerçekte yanlış olan sıfır hipotezini kabul etme olasılığı

S-3) Yığın standart sapması (σ) ile ikinci tip hata (β) arasındaki ilişki için aşağıdakilerden hangisi doğrudur?

- A) İlişki yoktur B) İlişki aynı yöndedir
C) İlişki ters yöndedir D) Çoğunlukla aynı nadiren ters yöndedir

S-4) A kurumundan memnuniyet skorunun cinsiyete göre farklılaşmadığı iddiası test edilecektir. Aşağıdaki testlerden hangisi uygulanmalıdır?

- A) Bağımsız örneklerde t-testi B) Eşleşmiş örneklerde t-testi
C) Tek örneklem t-testi D) Tek faktör varyans analizi

KORELASYON ANALİZİ

Pearson Korelasyon Katsayısı

Doğrusal ilişkili iki değişken arasındaki ilişkinin gücünü ve yönünü ölçen bir istatistiktir ve -1 ile +1 arasında değerler almaktadır.

$$\rho_{XY} = \frac{Cov(X, Y)}{\sigma_Y \sigma_X}$$

Burada

ρ_{XY} : X ve Y değişkenleri arasındaki Pearson korelasyon katsayısını

$Cov(X, Y)$: X ve Y değişkenleri arasındaki kovaryansı

σ_Y : Y değişkeninin standart sapmasını

σ_X : X değişkeninin standart sapmasını

göstermektedir.

Kovaryans katsayısı da iki değişken arasındaki ilişkinin yönünü gösteren bir ölçüttür. Bu katsayı negatif olduğunda ilişki ters yönlü iken pozitif değerler ilişkinin aynı yönlü olduğunu gösterir. Ancak kovaryans katsayısının tanım aralığı $(-\infty, +\infty)$ olduğundan ilişkinin derecesi hakkında bilgi veremez. Ayrıca kovaryansın birimi $X'in birimi \times Y'nin birimidir$. Oysaki korelasyon katsayısının birimi yoktur. Diğer bir ifadeyle kovaryans X ve Y'nin birimlerine bağlı bir değer alırken korelasyon birimden bağımsız olmaktadır. Uygulamalı çalışmalarda ilişki katsayısının birimden bağımsız olması bir avantajdır. Bu iki durum dikkate alındığında, iki değişken arasındaki ilişkinin kovaryans yerine korelasyon katsayısı ile araştırılması daha uygun olmaktadır.

Varsayımlar:

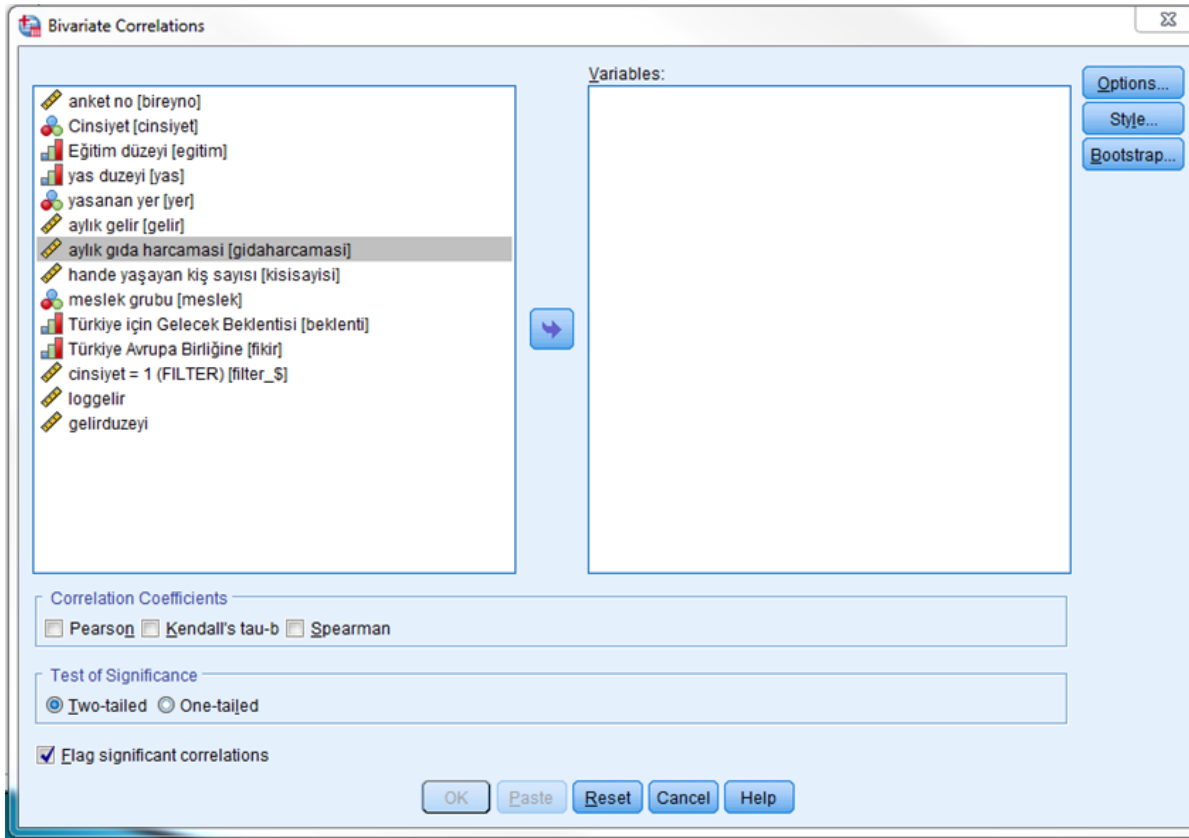
- Her iki değişken de normal dağılıma sahip olmalı
- Her iki değişken en az eşit aralıklı ölçme düzeyinde ölçülmüş olmalıdır.
- İki değişken arasındaki ilişki doğrusal olmalı

Özellikler:

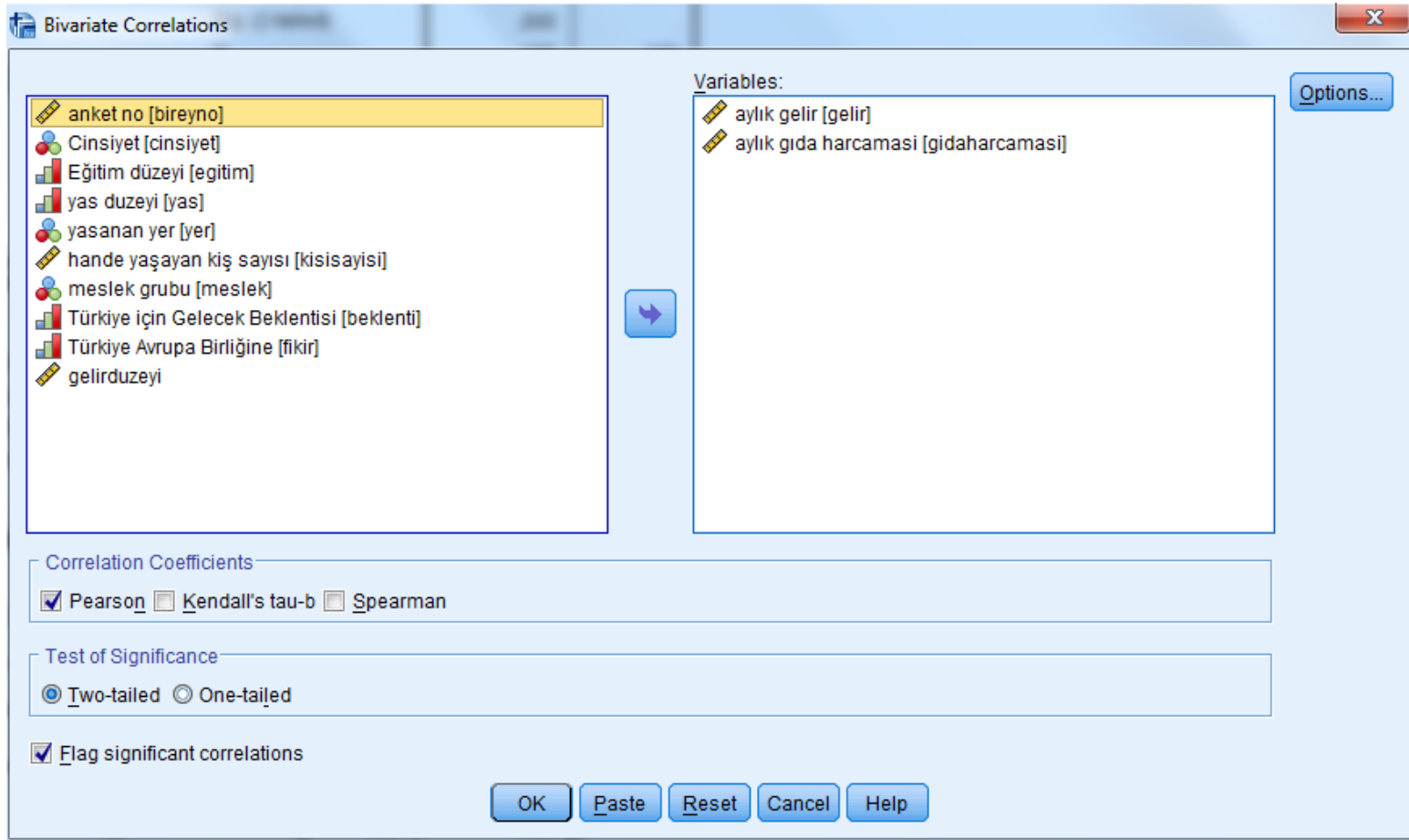
- ❖ Cohen'in standardı dikkate alındığında Pearson korelasyon katsayısı 0.10-0.29 arasında ise zayıf bir ilişkiyi, 0.30-0.49 arasında ise orta düzeyde bir ilişkiyi, 0.50 den yukarıda ise güçlü bir ilişkiyi göstermektedir.
- ❖ $\rho_{XY} < 0$ ise X ve Y değişkenleri arasında ters yönlü,
 $\rho_{XY} > 0$ ise X ve Y değişkenleri arasında aynı yönlü bir ilişki vardır.
- ❖ Korelasyon katsayısı simetriktir ($\rho_{XY} = \rho_{YX}$). Diğer bir ifade ile bağımlı bağımsız değişken ayrımı yoktur.

SPSS'te Korelasyon Analizi

SPSS'te iki deęiřken arasındaki korelasyon analizi *Analyze* → *Correlate* → *Bivariate* kısmından yapılmaktadır.



Yandaki pencerede *Variables* kısmına korelasyonu bulunacak iki deęiřken atılır ve daha sonra hangi korelasyon  l  t  kullanılacaksa se ilir.



Örnek: Aylık gelir ve aylık gıda harcaması değişkeni arasındaki uygun korelasyon katsayısını hesaplayınız.

Yukarıdaki pencerede *Variables* kısmına aylık gelir ve aylık gıda harcaması değişkenleri atılır. Daha sonra değişkenler normal dağılıma sahip ve nicel olduğundan Pearson korelasyon katsayısı işaretlenir.

Correlations

		aylık gelir	aylık gıda harcamasi
aylık gelir	Pearson Correlation	1	,950**
	Sig. (2-tailed)		,000
	N	100	100
aylık gıda harcamasi	Pearson Correlation	,950**	1
	Sig. (2-tailed)	,000	
	N	100	100

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

H_0 :Aylık gıda harcaması ile aylık gelir arasında ilişki yoktur. ($\rho_{XY}= 0$)

H_1 :Aylık gıda harcaması ile aylık gelir arasında ilişki vardır. ($\rho_{XY}\neq 0$)

Sig. (2-tailed) < 0.05 olduğundan H_0 reddedilebilir.

DOĞRUSAL REGRESYON ANALİZİ

Regresyon analizi bağımlı değişken (Y) ile bağımsız değişkenler ($X_1, X_2, \dots, X_k$) arasındaki ilişkiyi doğrusal formda ifade eden bir modeldir. Doğrusal regresyon analizinde bağımlı değişken nicel iken bağımsız değişkenler hem nicel hem de nitel olabilir. Yığın için doğrusal regresyon eşitliği aşağıdaki gibi tanımlanır;

$$Y_i = \alpha_0 + \alpha_1 X_{1i} + \alpha_2 X_{2i} + \dots + \alpha_k X_{ki} + \varepsilon_i$$

Burada α 'lar modelin bilinmeyen parametreleridir. ε ise kesin ilişkiyi bozan hata terimidir. Bu ilişkide Y bağımlı değişkeni $X_1, X_2, \dots, X_k$ bağımsız değişkeninin bir fonksiyonudur.

Diğer bir ifadeyle $Y = f(X_1, X_2, \dots, X_k)$ şeklinde yazılabilir. Bu durum Y ile X arasında bir nedensel ilişki olduğuna işaret etmektedir. Bu nedensel ilişkinin yönü ise X'lerden Y'ye doğrudur. Diğer bir ifadeyle X değişkenleri Y değişkeninin nedenidir. Nedenselliğin yönü hakkında ise önsel bilgi ile karar verilmektedir. Önsel bilgi, ilgili bilim dalından gelen teorilere dayalı olup verilerle ortaya çıkartılamaz. Örneğin tüketim ile gelir arasındaki ilişkide tüketim bağımlı gelir ise bağımsız değişken olacaktır. Çünkü iktisat bilimi tüketimi gelirin bir fonksiyonu olduğunu varsaymaktadır (Tüketim=f(Gelir))

Regresyon analizinin amacı bağımlı değişkendeki varyansın kaynağını araştırmaktır. Bu modeller karar vericilerin kullanacağı bazı bilgileri sağlamaktadır. Regresyon modelleri genellikle yapısal analiz ve öngörü amaçlı olarak kullanılmaktadır. Model kurucu bir parametreyi modeldeki diğer parametrelerden fonksiyonel olarak bağımsız varsayarak modeli kuruyorsa, o parametre yapısal olarak adlandırılır. Bir parametre yapısal ise, diğer her şey sabitken, X_j 'nin Y üzerindeki marjinal etkisi α_j kadar olacaktır. Diğer bir ifadeyle; X_j ve Y sürekli değişkenler olmak üzere Y 'in X_j 'e göre kısmi türevi X_j 'nin Y üzerindeki marjinal etkisidir;

$$X_j\text{'nin } Y \text{ üzerindeki marjinal etkisi} = \frac{\partial Y}{\partial X_j} = \alpha_j$$

O halde, diğer her şey sabitken, X_j değişkeni 1 birim değiştiğinde Y değişkeni α_j birim kadar değişecektir. Modeldeki α_0 katsayısı sabit terim (veya kesim katsayısı) olarak adlandırılmaktadır. $X_1 = X_2 = \dots = X_k = 0$ olduğunda $Y = \alpha_0$ değerini alacaktır. Sabit terim çoğu zaman anlamlı bir yoruma sahip değildir. Örneğin, konut fiyatının bağımlı, konut yüzölçümünün (metrekare) bağımsız değişken olduğu bir regresyon modelinde, α_0 , konutun yüzölçümü sıfır olduğunda konutun fiyatına karşılık gelecektir. Bu da uygulamada karşılaşılmayacak bir durum olduğundan, α_0 'ın yorumlanması anlamlı olmayacaktır. Buna karşın tüketim-gelir ilişkisinde, gelir sıfır olduğunda tüketim α_0 kadar olacaktır. Bu da zorunlu tüketime karşılık gelmektedir. Bu örnekte α_0 'ın yorumlanması anlamlı görülmektedir. Aynı zamanda α_0 katsayısı, $X_1 = X_2 = \dots = X_k = 0$ olduğunda Y 'nin ortalama değerine karşılık gelmektedir. Bu nedenle Y 'nin ortalaması sıfır olmadıkça sabit terimin modelden dışlanması uygun değildir.

Regresyon eşitliğinde yer alan hata terimi, ε , kesin ilişkiyi bozan rastgele değişkendir. Aynı zamanda hata terimi, bağımlı değişkendeki değişimin bağımsız değişkenler tarafından açıklanmayan kısmını da tanımlamaktadır. Dolayısıyla, hata teriminin regresyon eşitliğine katılmasının nedenleri şu şekilde sıralanabilir;

- İhmal edilen (modele alınmayan) önemli bağımsız değişkenler
- Önemli olmayan bağımsız değişkenlerin modele katılması
- Değişkenlerin ölçülmesinde yapılan hatalar (Ölçme hataları)
- Modelin matematiksel kalıbında kesinlik olmaması

Doğrusal regresyon modelinin sağlanması gereken varsayımlar aşağıda verilmiştir.

- ❖ Y ile X'ler arasındaki ilişki $Y_i = \alpha_0 + \alpha_1 X_{1i} + \alpha_2 X_{2i} + \dots + \alpha_k X_{ki} + \varepsilon_i$ şeklinde doğrusaldır.
- ❖ Hata terimi rastgele değişkendir.
- ❖ Hata teriminin ortalaması sıfırdır.
- ❖ Hata teriminin varyansı sabittir.
- ❖ Hata terimleri arasında ardışık bağımlılık (otokorelasyon) yoktur. $i \neq j$ için $Cov(\varepsilon_i, \varepsilon_j) = 0$
- ❖ Bağımsız değişken X'ler tekrar eden örneklerde sabittir. Bunun bir sonucu olarak hata terimi ile bağımsız değişkenler arasında ilişki yoktur.
- ❖ Hata teriminin dağılımı normal dağılımdır.
- ❖ Gözlem sayısı bağımsız değişken sayısından en az bir fazla olmalıdır. ($n > k$)

Modelin Açıklama Gücü: Belirleme Katsayısı

Belirleme katsayısı (R^2) bağımlı değişkendeki toplam değişimin yüzde kaçının bağımsız değişkenler tarafından açıklandığını gösteren istatistiki bir ölçüttür. Bu katsayı modelin açıklama gücü olarak da yorumlanmaktadır. Diğer bir ifadeyle, belirleme katsayısı 1'e yaklaştıkça modelin açıklama gücü artmakta ve seçilen bağımsız değişkenler yardımıyla bağımlı değişkendeki toplam değişimlerin büyük bir bölümü açıklanabilmektedir.

Bağımlı değişken için toplam kareler toplamı (TKT) aşağıdaki gibi yazılabilir:

$$TKT = \sum_{i=1}^n (Y_i - \bar{Y})^2$$
$$TKT = RKT + AKT$$

Burada toplamı kareler toplamı (TKT), regresyonla açıklanan kareler toplamı (RKT) ve regresyonla açıklanamayan kareler toplamı veya artık kareler toplamı (AKT) şeklinde iki kısma ayrılır. Belirleme katsayısı (R^2)

$$R^2 = \frac{RKT}{TKT} = 1 - \frac{AKT}{TKT}$$

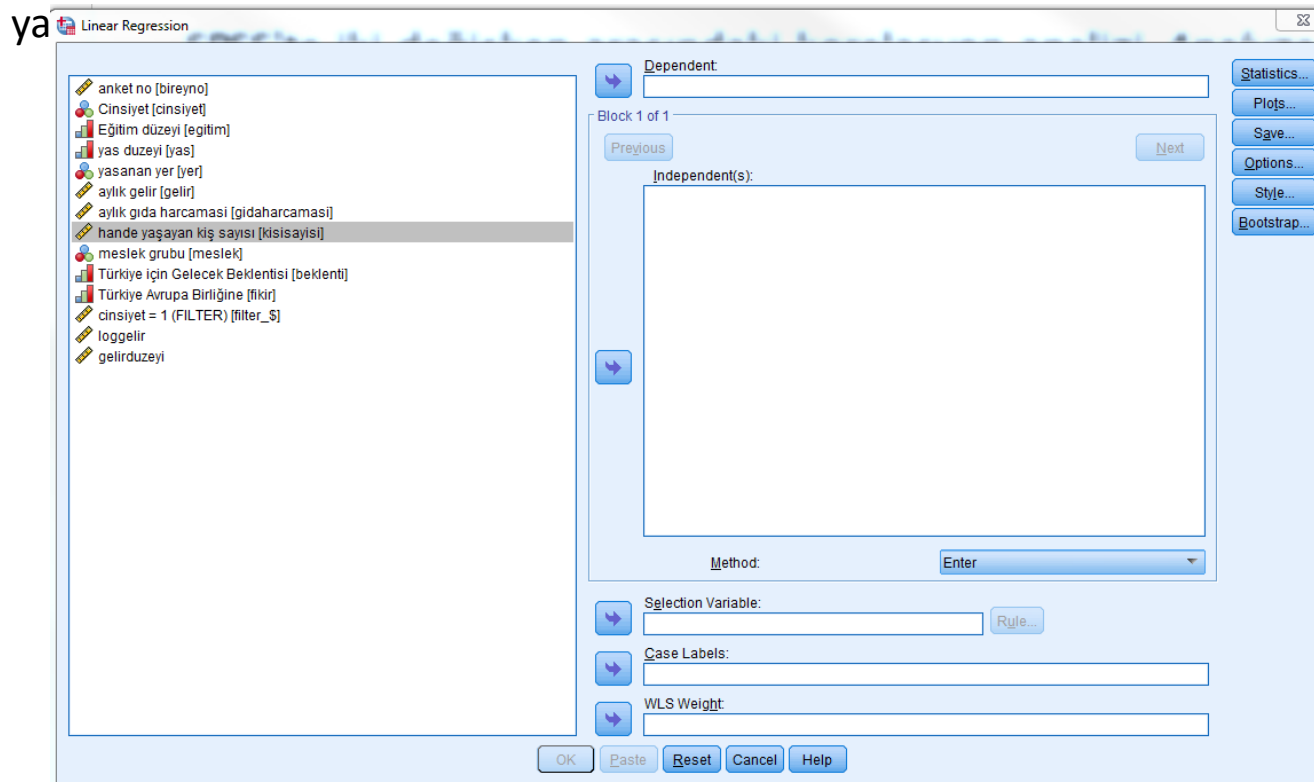
şeklinde hesaplanır. Burada

$$RKT = \sum_{i=1}^n (\hat{Y}_i - \bar{Y})^2$$

$$AKT = \sum_{i=1}^n (Y_i - \hat{Y}_i)^2$$

SPSS'te Doğrusal Regresyon Analizi

SPSS'te iki değişken arasındaki korelasyon analizi *Analyze* → *Regression* → *Linear* kısmından



Yandaki pencerede *Dependent* kısmına bağımlı değişken, *Independent(s)* kısmına ise bağımsız değişkenler atılır.

Örnek 3. Aylık gelir ve hanede yaşayan kişi sayısının aylık gıda harcaması üzerindeki etkisini doğrusal regresyon analizi ile inceleyiniz.

Coefficients^a

Model	Unstandardized Coefficients		Standardized Coefficients	t	Sig.
	B	Std. Error	Beta		
1 (Constant)	-15,614	27,509		-,568	,572
aylık gelir	,385	,014	,892	27,274	,000
hanede yaşayan kişi sayısı	25,458	6,320	,132	4,028	,000

a. Dependent Variable: aylık gıda harcaması

$$y_i = \alpha_0 + \alpha_1 X_{i1} + \alpha_2 X_{i2} + \varepsilon_i$$

$$\hat{y}_i = -15.614 + 0.385 X_{i1} + 25.458 X_{i2}$$

$\hat{\alpha}_0 = -15.614 \rightarrow X_1 = X_2 = 0$ iken ortalama aylık gıda harcaması -15.61 TL dir.

$\hat{\alpha}_1 = 0.385 \rightarrow$ Diğer her şey sabitken aylık gelir 1 TL arttığında aylık gıda harcaması 0.385 TL artar.

$\hat{\alpha}_2 = 25.458 \rightarrow$ Diğer her şey sabitken hanede yaşayan kişi sayısı 1 artarsa aylık gıda harcaması 25.458 TL artar.

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Durbin-Watson
1	,957 ^a	,916	,914	61,938	2,078

a. Predictors: (Constant), hanede yařayan kiři sayısı, aylık gelir

b. Dependent Variable: aylık gıda harcaması

Aylık gıda harcamasını %91.6 oranında aylık gelir ve hanede yařayan kiři sayısı deęiřkenleri açıklamaktadır.

$$H_0: \alpha_1 = 0$$

$$H_1: \alpha_1 \neq 0$$

$t_{\hat{\alpha}_1} = 27.274$ ve p-deęeri=0 < 0.05 olduęundan H_0 reddedilebilir. Yani aylık gelir deęiřkeninin katsayısı istatistiksel olarak anlamlıdır.

$$H_0: \alpha_2 = 0$$

$$H_1: \alpha_2 \neq 0$$

$t_{\hat{\alpha}_2} = 4.028$ ve p-deęeri=0 < 0.05 olduęundan H_0 reddedilebilir. Yani hanede yařayan kiři sayısı deęiřkeninin katsayısı istatistiksel olarak anlamlıdır.

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	4065375,233	2	2032687,616	529,851	,000 ^b
	Residual	372124,767	97	3836,338		
	Total	4437500,000	99			

a. Dependent Variable: aylık gıda harcaması

b. Predictors: (Constant), hanede yaşayan kişi sayısı, aylık gelir

$H_0: \alpha_1 = \alpha_2 = 0$ (Kurulan regresyon modeli anlamsızdır)

$H_1: \alpha_j$ 'lerin en az biri sıfırdan farklıdır. ($j = 1,2$)

F=529.851

$P=0.000 < \alpha = 0.05$ olduğu için H_0 reddedilebilir. Yani kurulan regresyon modeli istatistiksel olarak anlamlıdır.

DOĞRU-YANLIŞ SORULARI

	Doğru	Yanlı ş
Pearson korelasyon katsayısı simetriktir. Yani X ile Y yerine Y ile X arasındaki korelasyonu hesaplamak sonucu deęiřtirmez.		
Korelasyon X ile Y deęiřkenleri arasındaki nedensel iliřkinin yönü hakkında bilgi vericidir.		
Pearson korelasyon katsayısı sadece doğrusal iliřkiler hakkında bilgi verebilir.		
Belirleme katsayısı R^2 sıfıra yaklařtıkça modelin açıklama gücü azalır		
Regresyon analizinde ihmal edilen önemli bağımsız deęiřkenler hata teriminin önemli bir kaynağıdır.		
Regresyon analizinde hata terimi deterministik bir bileřendir		

TEST SORULARI

S-1) Y bağımlı ve X bağımsız değişkeni arasındaki doğrusal regresyon eşitliğinin parametreleri tahmin edilmiş ve sabit terim 3, eğim katsayısı ise 0.6 olarak bulunmuştur. X'in Y üzerindeki marjinal etkisi için aşağıdaki yorumlardan hangisi doğrudur??

- A) Diğer her şey sabitken, X %1 arttığında Y %0.6 artar
- B) Diğer her şey sabitken, X 1 birim arttığında Y %0.6 artar
- C) Diğer her şey sabitken, X 1 birim arttığında Y 0.6 birim artar
- D) Diğer her şey sabitken, X %1 arttığında Y 0.6 birim artar

S-2) Aşağıdakilerden hangisi X ve Y değişkenleri arasındaki ilişkiyi ölçmek için hesaplanan Pearson korelasyon katsayısının birimidir?

- A) Birimden bağımsızdır
- B) X'in birimi (Y'nin birimi)
- C) X'in birimi
- D) Y'nin birimi

S-3) Regresyon analizi ařağıdaki amalardan hangisi (hangileri) iin kullanılır?

- A) ngr ve yapısal analiz B) Sadece ngr
C) Sadece yapısal analiz D) Bağımlı deęiřkeni aıklamak

4) X ve Y deęiřkenleri arasındaki Pearson korelasyon katsayısı 0.98 olarak hesaplanmıřtır. Ařağıdaki yorumlardan hangisi tam olarak doęru yapılmıřtır?

- A) X ve Y arasında aynı ynde iliřki vardır B) X ve Y arasında aynı ynde kuvvetli bir iliřki vardır
C) X ve Y arasında aynı ynde orta derecede iliřki vardır D) X ve Y arasında aynı ynde olduka kuvvetli doęrusal bir iliřki vardır

SPSS'TE ÇAPRAZ TABLO

Çapraz tablo temel olarak, iki kategorik değişken arasındaki ilişkiyi analiz etmek için kullanılır. Örneğin cinsiyet ve oy verilen parti arasında ilişki olup olmaması gibi.

Çapraz tablolar izlenen amaca göre üç türlü yapılmaktadır. (Darcy ve Rohrs, 1995)

- Bir değişkenin bir başka değişken üzerindeki etkisini göstermek (yüzdelemenin yönü eğer satır -yatay yönündeki- değişkeni bağımsız değişken ise bu yönde, yok eğer bağımsız değişken sütun – dikey yönündeki- değişkeni ise bu yönde yapılır).
- Bir grubun kompozisyonunu (dağılımını) belirlemek için.
- Çaprazlanan değişkenler sonucu ortaya çıkan olası alt grupların bütün içindeki kompozisyonunu belirlemek için.

Örneğin cinsiyet ile gelir düzeyi arasındaki ilişkiyi gösteren çapraz tablonun oluşturulmasında SPSS'te aşağıdaki adımlar uygulanır.

Analyze → Descriptive Statistics → Crosstabs...

SS Statistics Data Editor

Transform Analyze Direct Marketing Graphs Utilities Add-ons Window Help

cinsiyet		na		kisisayis
1				00
1				00
1				00
1				00
1				50
1				50
1				550
1				600
1				650
1				650
1				850

Reports

Descriptive Statistics

Custom Tables

Compare Means

General Linear Model

Generalized Linear Models

Mixed Models

Correlate

Regression

Loglinear

Neural Networks

Classify

Dimension Reduction

Scale

Nonparametric Tests

123 Frequencies...

Descriptives...

Explore...

Crosstabs...

TURF Analysis

1/2 Ratio...

P-P Plots...

Q-Q Plots...

Crosstabs

anket no [bireyno]

Eğitim düzeyi [egitim]

yas duzeyi [yas]

yasanan yer [yer]

aylık gelir [gelir]

aylık gıda harcaması [gidaharcaması]

hande yaşayan kişi sayısı [kisisayisi]

meslek grubu [meslek]

Türkiye için Gelecek Beklentisi [beklenti]

Türkiye Avrupa Birliğine [fikir]

Row(s):

Cinsiyet [cinsiyet]

Column(s):

gelirduzeyi

Layer 1 of 1

Previous

Next

☐ Display clustered bar charts

☐ Suppress tables

Display layer variables in table layers

OK

Paste

Reset

Cancel

Help

Exact...

Statistics...

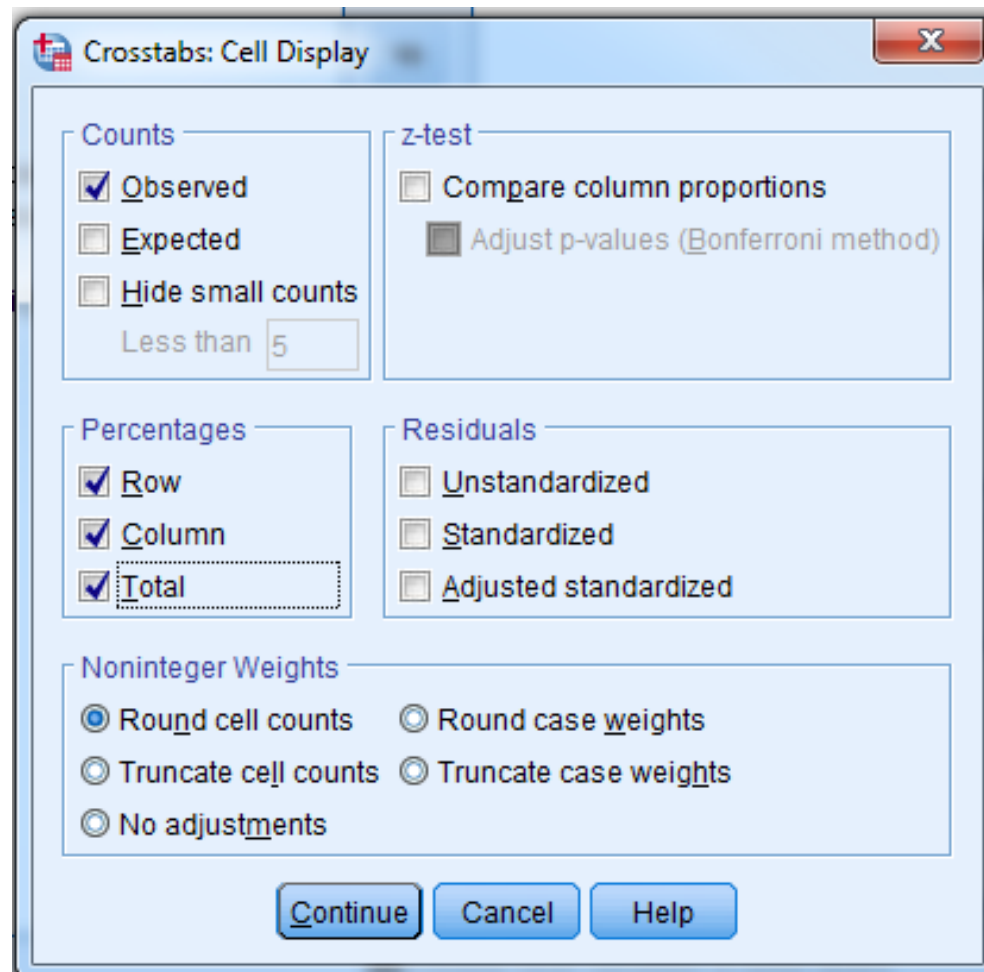
Cells...

Format...

Style...

Bootstrap...

Cells... menüsünde aşağıdaki kısımlar işaretlenir.



Cinsiyet * gelirduzeyi Crosstabulation

			gelirduzeyi			Total
			1100-1800	1801-2500	2501-3200	
Cinsiyet	erkek	Count	22	20	6	48
		% within Cinsiyet	45,8%	41,7%	12,5%	100,0%
		% within gelirduzeyi	47,8%	51,3%	40,0%	48,0%
		% of Total	22,0%	20,0%	6,0%	48,0%
	kadın	Count	24	19	9	52
		% within Cinsiyet	46,2%	36,5%	17,3%	100,0%
		% within gelirduzeyi	52,2%	48,7%	60,0%	52,0%
		% of Total	24,0%	19,0%	9,0%	52,0%
Total	Count		46	39	15	100
	% within Cinsiyet		46,0%	39,0%	15,0%	100,0%
	% within gelirduzeyi		100,0%	100,0%	100,0%	100,0%
	% of Total		46,0%	39,0%	15,0%	100,0%

Erkeklerin %45.8'inin geliri
1100-1800 aralığındadır

Geliri 1801 - 2500 TL
aralığında olanların %48.7'i
kadındır

Ankete katılanların %24'ü kadın ve
geliri 1100 - 1800 TL aralığındadır

Kİ-KARE BAĞIMSIZLIK TESTİ

Uygulamalı çalışmalarda değişkenlerin büyük bir bölümü sınıflama ya da sıralama ölçme düzeyinde nitel değişkenlerdir. Eğer iki değişken arasında ilişki yoksa bu iki değişkenin bağımsız olduğu söylenebilir. İki değişken bağımsız ise değişkenlerden birinin değerini bilmek, diğer değişkenin alacağı değeri tahmin etmemize yardımcı olmaz. Sınıflama ya da sıralama ölçme düzeyinde gözlemlenmiş iki değişken arasındaki ilişki Ki-kare (χ^2)bağımsızlık testi ile araştırılabilir.

χ^2 bağımsızlık testinde çapraz tablonun oluşturulması önemli bir yer tutar. Bu tablonun oluşturulmasında öncelikle değişkenlerin kaç farklı değer alacağı saptanır. Birinci değişken (X_1) düzeyleri c ve ikinci değişken (X_2) r düzeyli olsun. Bu durumda

n_{ij} : X_1 değişkeninin i . ve X_2 değişkeninin j . düzeyindeki örnek birimlerinin sayısı (Gözlenen frekanslar)

$n_{.i}$: X_1 değişkeninin i . düzeyindeki örnek birimlerinin sayısı

$n_{.j}$: X_2 değişkeninin j . düzeyindeki örnek birimlerinin sayısı

Gözlenen frekanslar $G_{ij} = n_{ij}$ ile gösterilirsin.

H_0 : Değişkenler bağımsızdır.

H_1 : Değişkenler bağımsız değildir.

χ^2 bağımsızlık testi gözlenen ve beklenen frekanslar arasındaki farklara dayalı olan bir testtir. İki değişken bağımsız ise i . satır ve j . sütunda yer alan hücre için beklenen frekans B_{ij}

$$B_{ij} = \frac{n_{.i} \times n_{.j}}{n}$$

şeklinde hesaplanır.

H_0 hipotezi doğru iken test istatistiği

$$\chi_h^2 = \sum_{i=1}^c \sum_{j=1}^r \frac{(G_{ij} - B_{ij})^2}{B_{ij}}$$

Bu test istatistiği $(c - 1)(r - 1)$ serbestlik dereceli χ^2 dağılımına sahiptir. Yukarıda tanımlanan χ_h^2 değeri, araştırmacı tarafından önceden belirlenen 1. tip hata düzeyine (α) karşılık gelen χ^2 tablo değeri ile karşılaştırılarak H_0 hipotezi test edilir.

$\chi_h^2 > \chi_{Tablo}^2$ ise H_0 hipotezi reddedilebilir.

$\chi_h^2 \leq \chi_{Tablo}^2$ ise H_0 hipotezi reddedilemez.

Ki-kare bağımsızlık testinde dikkat edilmesi gereken hususlar

(i) Beklenen frekanslar 1'den küçük olmamalıdır.

(ii) Beklenen frekansların en fazla %20'si 5'den daha küçük olabilir.

Bu koşullar sağlanmadığında

➤ Örnek çapı arttırılabilir

➤ Satırlar veya sütunlar birleştirilebilir

(iii) Her iki değişken için düzeylerin 2 olması durumunda (2×2) Yates düzeltmesi yapılmalıdır.

$$\chi_h^2 = \sum_{i=1}^c \sum_{j=1}^r \frac{(|G_{ij} - B_{ij}| - 0.5)^2}{B_{ij}}$$

Örnek 1. Eğitim düzeyi ile gelir düzeyi arasında %5 anlamlılık düzeyinde ilişki var mıdır?

Eğitim düzeyi * gelirdüzeyi Crosstabulation

			gelirdüzeyi			Total
			1100-1800	1801-2500	2501-3200	
Eğitim düzeyi	ilkogretim	Count (Gözlenen Frekans)	29	4	1	34
		Expected Count (Beklenen Frekans)	15,6	13,3	5,1	34,0
	lise	Count (Gözlenen Frekans)	15	18	1	34
		Expected Count (Beklenen Frekans)	15,6	13,3	5,1	34,0
	yuksekokul	Count (Gözlenen Frekans)	2	17	13	32
		Expected Count (Beklenen Frekans)	14,7	12,5	4,8	32,0
Total	Count (Gözlenen Frekans)		46	39	15	100
	Expected Count (Beklenen Frekans)		46,0	39,0	15,0	100,0

H_0 : Eğitim düzeyi ile gelir düzeyi arasında ilişki yoktur.

H_1 : Eğitim düzeyi ile gelir düzeyi arasında ilişki vardır.

Chi-Square Tests

	Value	df	Asymptotic Significance (2-sided)
Pearson Chi-Square	52,829 ^a	4	,000
Likelihood Ratio	57,887	4	,000
Linear-by-Linear Association	43,022	1	,000
N of Valid Cases	100		

a. 1 cells (11,1%) have expected count less than 5. The minimum expected count is 4,80.

Yukarıdaki tabloda H_0 hipotezinin testinde Pearson Chi-Square değerine bakılır. $\chi_h^2 = 52.829$ ve $p\text{-değeri} = 0.000 < 0.05$ olduğundan 'Eğitim düzeyi ile gelir düzeyi arasında ilişki yoktur' sıfır hipotezi reddedilebilir.

Örnek 2. Cinsiyet ile yaşanan yer arasında %5 anlamlılık düzeyinde ilişki var mıdır?

Cinsiyet * yaşanan yer Crosstabulation

			yasanan yer		Total
			kır	kent	
Cinsiyet	erkek	Count	26	22	48
		Expected Count	23,5	24,5	48,0
	kadın	Count	23	29	52
		Expected Count	25,5	26,5	52,0
Total		Count	49	51	100
		Expected Count	49,0	51,0	100,0

H_0 : Cinsiyet ile yaşanan yer arasında ilişki yoktur.

H_1 : Cinsiyet ile yaşanan yer arasında ilişki vardır.

Chi-Square Tests

	Value	df	Asymptotic Significance (2-sided)	Exact Sig. (2- sided)	Exact Sig. (1- sided)
Pearson Chi-Square	,986 ^a	1	,321	,423	,214
Continuity Correction ^b	,629	1	,428		
Likelihood Ratio	,988	1	,320		
Fisher's Exact Test					
Linear-by-Linear Association	,976	1	,323		
N of Valid Cases	100				

a. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 23,52.

b. Computed only for a 2x2 table

Çapraz tablo 2×2 olduğundan Pearson Chi-Square değeri yerine Continuity Correction (Yates Düzeltmeli χ_h^2) katsayısına bakılır. $\chi_h^2 = 0.629$ ve p-değeri $= 0.426 > 0.05$ olduğundan 'Cinsiyet ile yaşanan yer arasında ilişki yoktur' sıfır hipotezi reddedilememektedir.

Çapraz tablo 2×2 olduğunda ve 5'den küçük beklenen frekansların sayısı %20'den büyük olduğunda birleştirme yapılamadığından dolayı yukarıdaki tabloda Fisher Exact testi kullanılmalıdır.

DOĞRU-YANLIŞ SORULARI

	Doğru	Yanlış
Sınıflama ya da sıralama ölçme düzeyinde gözlemlenmiş iki değişken arasındaki ilişki Ki-kare (χ^2)bağımsızlık testi ile araştırılabilir.		
Ki-kare bağımsızlık testi gözlenen ve beklenen frekanslar arasındaki farklara dayalı olan bir testtir.		
Ki-kare bağımsızlık testinde beklenen frekansların tamamı 5'den küçük olamaz		
Çapraz tablo 3x3 boyutunda ise Ki-kare bağımsızlık testinde Yates düzeltmesi yapılmalıdır		
Cinsiyet ile oy verilen parti arasında ilişki yoktur iddiası Ki-kare bağımsızlık testi ile araştırılır.		
Döviz kuru ile enflasyon arasındaki ilişki Ki-kare bağımsızlık testi ile araştırılır.		

TEST SORULARI

S-1) Eğitim düzeyi ile Avrupa Birliğine üyelik (katılsın, fark etmez, katılmasın) hakkındaki tutumlar arasındaki ilişki olup olmadığı Ki-kare bağımsızlık testi ile araştırılacaktır. Aşağıdakilerden hangisi sıfır (H_0) hipotezidir?

- A) Eğitim düzeyi arttıkça AB'ye üye olma yönündeki eğilim artar
- B) Eğitim düzeyi ile AB'ye üye olma hakkındaki tutumlar arasında ilişki vardır
- C) Eğitim düzeyi ile AB'ye üye olma hakkındaki tutumlar arasında ilişki yoktur
- D) Eğitim düzeyi ile AB'ye üye olma hakkındaki tutumlar birbirinden bağımsız değildir

S-2) Kikare bağımsızlık testinde beklenen frekansların en az yüzde kaç 5'den küçük olamaz?

- A) %5 B) %10 C) %15 D) %20