

İSTANBUL BİLGİ ÜNİVERSİTY

İşletme İstatistiği

[Type the document subtitle]

Ege Yazgan ve Yüce Zerey

10/21/2003

[Type the abstract of the document here. The abstract is typically a short summary of the contents of the document. Type the abstract of the document here. The abstract is typically a short summary of the contents of the document.]

İçindekiler:

Bölüm 1: Giriş	1-8
Bölüm 2: Temel Araçlar	9-13
Bölüm 3: Tablo ve Grafiklerle Veri Sunumu (Sayısal Veriler)	14-37
Bölüm 4: Kategorik Verilerin Düzenlenmesi ve Tablolaştırılması	38-56
Bölüm 5: Betimsel İstatistikler	57-88
Bölüm 6: Olasılık ve Olasılık Dağılımları	89-129
Bölüm 7: Örneklemleme ve Örneklem Dağılımları	130-172
Bölüm 8: Tek Anakütleli Hipotez Testi	173-219
Bölüm 9: Aralık Tahmini	220-234
Bölüm 10: İki Anakütle Parametresi için Hipotez Testleri	235-270
Bölüm 11: Uyum veya Bağımlılık Testleri	271-279
Bölüm 12: İki veya Daha Fazla Anakütle için Hipotez Testleri: Varyans Analizi (ANOVA) ve F-Dağılımı	280-288
Bölüm 13: Basit Korelasyon	289-299
Bölüm 14: Basit Doğrusal Regresyon	300-333
Ek: İstatistiksel Tablolar	334-358

Bölüm 1:

Giriş

Neden İstatistik Öğrenilmeli?

Üniversitelerin farklı bölümlerine ait ders içerikleri incelendiğinde, birçok bölümde İstatistik dersinin zorunlu olarak okutulduğunu görülmektedir. İstatistik dersinin bölümlere göre değişen yönü sadece derste kullanılan matematiğin zorluk derecesi ve ders içeriğinde kullanılan örneklerden ibarettir. Bu farklılığın haricinde işlenen konularla bölümler arasında ciddi bir paralellik söz konusudur.

İşletme yönetiminde; maliyetler, gelirler, ücretler; inşaat mühendisliğinde; malzemelerin dayanıklılık katsayıları, psikolojide ise test ve deney sonuçları istatistiğin geniş uygulama sahasına giren konulardan sadece birkaçıdır.

İstatistik dersinin, çoğu bölümde zorunlu ders olarak okutulmasının sebebi, istatistiki bilgilerin yaşamın her alanında kullanılmasına gereksinim duyulmasıdır. Gazeteleri, ekonomiyi, finans haberlerini ve dergilerini, spor içeriklerini, otomobil yarışlarını, vb. dikkatli bir şekilde incelediğiniz takdirde, sayısal bilgi yağmuruna tutulduğunuzun farkına varırsınız. İşte bu sayısal bilgi yağmurunda ıslanmak istemeyenler şemsiye olarak istatistiğin gücünden faydalanmaktadırlar.

Günlük hayatımızda karşılaştığımız bu sonuçlar ve rakamlar ne kadar doğrudur? Bilgiyi elde etmek için kullanılan örneklem düzeyleri yeterince büyük müdür? Örneklemdeki veriler nasıl seçilmiştir? Bu tür sayısal bilgilerin bilinçli bir şekilde incelenebilmesi için; tablolarda, grafiklerde sunulan bilgiyi anlayabilmek, sayısal özetleri değerlendirebilmek, ve yanıltıcı bilgilerden nasıl kaçınılacağını öğrenmek gerekmektedir. Temel istatistik kavramları, bu öğrenme sürecinde bize oldukça yardımcı olacaktır.

İstatistiği kaçınılmaz kılan nedenlerden bir diğeri de günlük yaşamımızı etkileyen bazı kararların alım sürecinde istatistiksel yöntemlerin kullanımının gerekliliğidir. Örneğin:

Sigorta şirketleri, hayat, otomobil, konut sigortaları için ödenecek primleri belirlerken istatistiksel analizlere başvururlar. Otomobilinizin çalınmaya karşı sigorta primi belirlenirken, ikamet ettiğiniz semtte hırsızlık olaylarının boyutu, çalınan otomobilin polis tarafından bulunması olasılığı, ve bunlara benzer bir çok bilgi kullanılır. Bu bilgiler de istatistiki çalışmalar ile elde edilir.

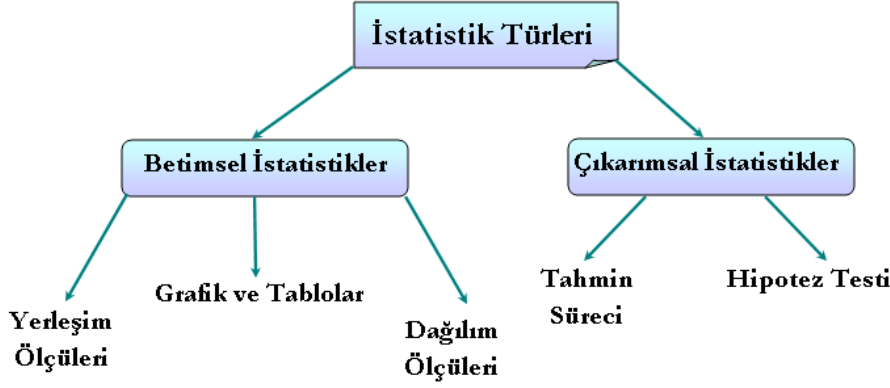
Çevre Koruma Örgütleri, Terkos Gölü'ndeki kirlilik düzeyini ölçmek ve insan sağlığına tehlike oluşturup oluşturmadığını anlamak için düzenli aralıklarla su örneklemeleri alarak istatistiksel yöntemlerle tetkik ederler.

Tıp araştırmacıları belli bir hastalığı iyileştirmek için hangi ilaçların ve yöntemlerin daha etkili olduğunu araştırırken yine istatistiksel süreçlerden yararlanmaktadırlar. Örneğin “Bel fıtığı ameliyatla mı daha kolay ve kısa zamanda iyileşir yoksa fizik tedavi metoduyla mı?”, ya da “Her gün bir aspirin almak kalp krizi riskini azaltır mı?” gibi soruların cevaplarını istatistik bilimi ile elde ederler.

Özetle, istatistik öğrenmemiz için en az 3 neden bulunmaktadır:

(1) İstatistiki bilgiler hayatımızın her döneminde her yerde karşımıza çıkacaktır.

- (2) Günlük yaşamımızı etkileyen kararları alırken istatistiki bilgilere ihtiyaç duymaktayız.
- (3) Yaptığımız iş, görevimiz ne olursa olsun, veri toplamak ve bunlardan anlamlı bilgi çıkartabilmek için istatistiki yöntemleri bilmemiz gerekmektedir.



Şekil 1.1

Betimsel İstatistikler

İstatistiksel çalışmalar genellikle betimleyici istatistikler ve tümevarımsal istatistikler olmak üzere iki ana grupta toplanmaktadır. Betimsel istatistikler, verinin toplanması, sunulması ve anlamlandırılması süreçlerinden oluşmaktadır. Veriler, anket gibi metodlar ile toplanır, tablolarla grafiklerle sunulur ve son olarak ortalama gibi istatistiklerle anlamlı bir şekilde kullanılabilir bir hale gelirler.

Örneğin, Birleşik Devletler Hükümeti'nin raporuna göre nüfus sayımında; 1960 yılında 179.323.175 kişi, 1970'te 203.302.031 kişi, 1980 de 226.542.203 kişi, 1990'da ise 248.709.783 kişi sayılmıştır. Bu bilgi betimsel istatistikleri ifade etmektedir. Aynı şekilde büyüme oranları gibi ekonomik parametrelerde betimsel istatistiklerin alanına girmektedir.

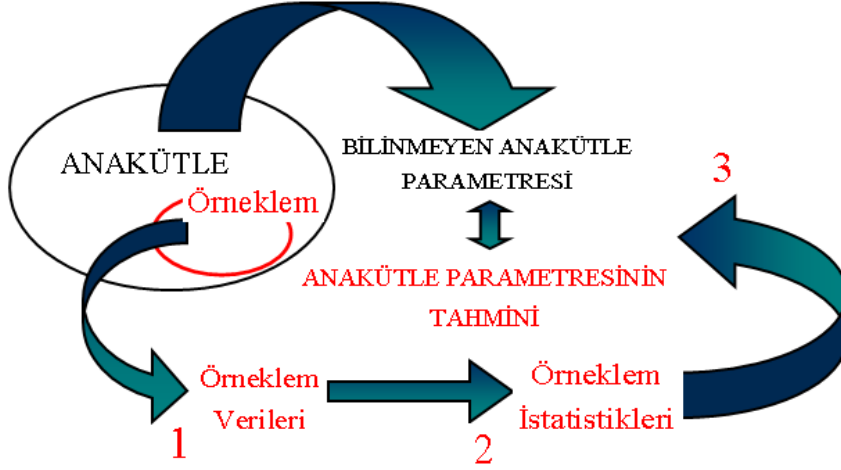
Tümevarımsal İstatistikler

Tümevarımsal istatistiklerin ana amacı elimizdeki örneklemin verilerini kullanarak anakütle parametrelerini tahmin etme prensibine dayanmaktadır. Örneğin, Marmara Bölgesinde sigara içenlerin oranını bularak Türkiye'deki sigara tüketimi hakkında bir fikir sahibi olunabilir.

Veri toplamak için öncelikle ilgilendiğimiz değişkene ait anakütle tanımlanır. Anakütle, değişkenin bütün değerlerini içeren evrensel kümedir. Örneğin, araştırmak istediğimiz konu "MBA eğitimi olan insanların ortalama işe başlama maaşları" ise, bilgi elde etmek istediğimiz anakütle, MBA yapan bütün insanların eğitimleri bitince aldıkları maaşlardan oluşur. Bu maaşların toplamı anakütledir.

MBA eğitimi alan bütün bireylerin maaşlarını birer birer öğrenmek çok zaman alacağından ve böyle bir işlemin maliyeti de yüksek olacağından, genellikle bütün anakütleyi incelemek yerine ilgililenen anakütleden örneklem seçilir. Örneklem, seçildiği anakütleyi temsil edebileceği

düşünülen bir altkümedir. Örneğin, MBA eğitimi alan bütün insanlar arasından seçilecek rastgele 500 kişinin maaşları bir örneklem grup olacaktır.



Şekil 1.2

Anakütleye ilişkin bir parametreyi (MBA eğitimi alan insanların ortalama işe başlama maaşı gibi) öğrenmek için;

1. Büyük küme anakütleden daha küçük bir altküme olan örneklem seçilir ve örneklem verileri kaydedilir.
2. Örneklem verilerden betimsel istatistikler hesaplanır.
3. Örneklem istatistiklerinden yola çıkılarak, tümevarımsal istatistik yöntemleriyle, bilinmeyen parametre tahmin edilir.

Veri Türleri

Sayısal ve Kategorik Veriler

Genel olarak verisetlerinin büyük bir kısmı nümerik olarak veya sayısal olarak ifade edilirler. Sayısal veriler, nümerik olarak bir miktar olarak tanımlanmış olan veri türleridir. Para birimleri, ölçüm birimleri, yüzdeler, maç istatistikleri, parfüm satışlarının her biri sayısal verilere birer örnek teşkil etmektedirler. Kategorik veriler, nitel olarak kategorileri ifade eden veri türleridir. Cinsiyet, gelir grupları, eğitim düzeyleri kategorik verilere örnek teşkil etmektedir.

Tablo 1.2 Sayısal ve Kategorik verilerinin örneklerle karşılaştırılması

Sayısal Veriler	Kategorik Veriler
Rakamlar	Sözel İfadeler
Sonuçlar	Yorumlar
Doğrulama	Araştırma
Kapalı Sorular	Açık Sorular

Zaman Serileri ve Yatay Kesit Verileri

Veriler ayrı bir perspektiften değerlendirildiğinde zaman serisi veya yatay kesit olup olmadıklarına göre de sınıflandırılmaktadırlar. Zaman Serileri, zaman içerisinde gözlenen, sıralanmış verisetleridir. Microsoft'un 2001 yılının ilk çeyreğinden 2002'nin son çeyreğine kadar Türkiye'de yapmış olduğu yıllık satış miktarları, tüketicilerin son bir ay içerisinde sanal marketten yapmış oldukları günlük ortalama satış miktarları, Türkiye'nin 1987'den 2003'e kadar olan tüketim miktarı, ülkelerin belli zaman aralıklarındaki büyüme oranları, enflasyon oranları, para arzları gibi verisetlerinin tamamı birer zaman serisidir. Yatay kesit verileri, sabit bir zamanda, gözlenerek elde edilmiş verisetleridir. Üniversite mezunu öğrencilerin istedikleri gibi bir iş bulması, bir firmada çalışan işçilerin memnuniyeti, e-ticarette yaşanan sorunlar, kalkınmakta olan ülkelerde okuryazar oranlarını karşılaştırılması, Kobilerde teknoloji kullanım düzeyinin, büyük ölçekli firmalarla karşılaştırılması, gibi yapılan araştırmalarda kullanılan veriler, yatay kesit verileridir.

Veri Ölçüm Düzeyleri

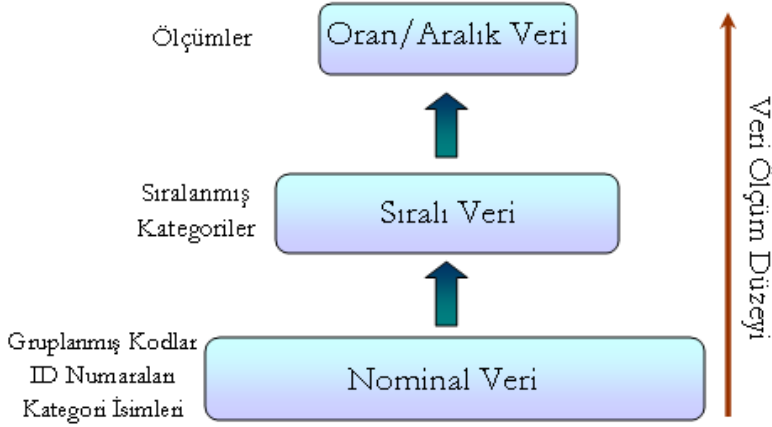
Nominal Veri : Nominal veri, herhangi bir verinin en ilkel formunu ifade etmektedir. Dolayısıyla nominal veri ile araştırmalarda, günlük yaşamda çok sık karşılaşılmaktadır. Örneğin, bir anket sorusunda bize medeni durumumuzun sorulduğunu düşünelim:

1. Bekar 2. Evli 3. Boşanmış 4. Dul

Bu soruyu cevaplayan herkesin cevabı 1, 2, 3, 4 olarak kaydedilecektir ve kaydedilen bu veriler bize nominal verileri ifade etmektedir. Rakamlar sadece sorulmuş olan kategorileri göstermektedir ve rastgele dizilmiştir. Nominal veriler sadece nümerik değerlerden oluşmazlar. Örneğin:

B = Bekar E = Evli BO=Boşanmış D= Dul

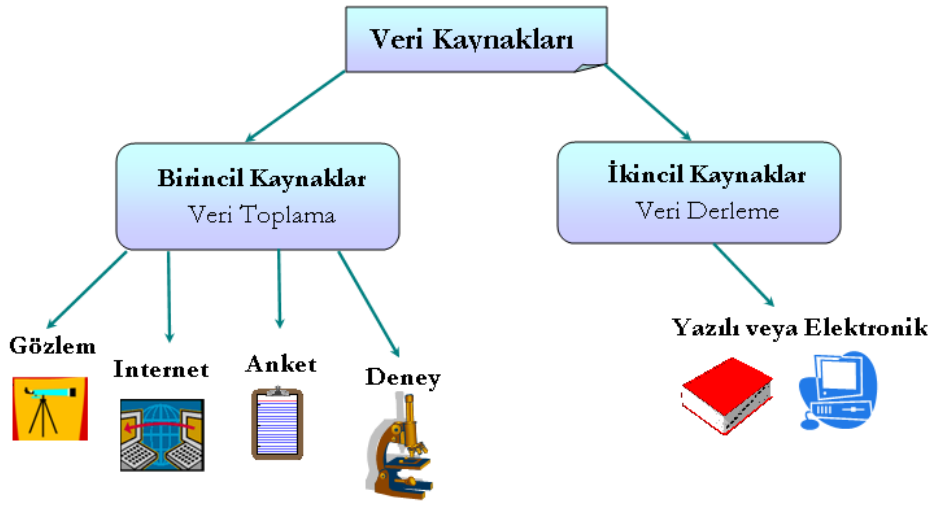
gibi kategorik verilerden oluşabilmektedir.



Şekil 1.3

Sıralı Veri: Veri elemanlarının, aralarındaki ilişkiye göre elemanlara atanan değerlerle birlikte sıralanmasıyla elde edilmektedir. Sıralı veri, veri ölçüm düzeyinde hemen nominal verinin üzerinde yer almaktadır. Örneğin Türkiye'de kredi kartı kullanımını etkileyen faktörlerini öğrenmek için bir araştırma yaptığımızı düşünelim. Kredi kartının kullanımı etkileyebilecek en önemli

Veri Kaynakları



Şekil 1.4

Veri kaynakları, birincil ve ikincil veri kaynakları olmak üzere iki ana başlıkta incelenebilir. Birincil verikaynakları araştırmanın daha çok verinin toplanması fazına odaklanmış iken, ikincil veri kaynakları ise araştırmanın verinin derlenmesi, yorumlanması fazına odaklanmıştır. Örneğin Türk Hava Yolları'nın Türkiye'deki bütün havaalanlarının teknik donanımı hakkında elinde bulundurduğu bilgi birincil bir veri iken bizim gazete Esenboğa havaalanı hakkında gazeteden okumuş olduğumuz bir bilgi de ikincil bir veriyi ifade etmektedir. Yine, BDDK murakıplarının bir banka için yapmış oldukları araştırma sonucu elde ettikleri veri birincil bir veri iken, anahaber bültenlerden o banka hakkında elde etmiş olduğumuz veri ikincil bir veriyi ifade eder. İnternet ortamından kişiler ile doğrudan iletişim kurarak (chat, mail, vs) edinilen veri, birincil veriyi ifade ederken, bu ortamdan çeşitli web sitelerini ziyaret ederek okuduğumuz, incelediğimiz ve öğrendiğimiz veri ikincil veri grubuna girmektedir.

Verinin Toplanması

Günümüzde verinin toplanmasında kullanılan çok farklı metod ve araçlar kullanılmaktadır. Genellikle kullanılan metod ve araçlar aşağıdaki gibidir:

- Doğrudan gözlem ve kişisel mülakat
- Anketler (elektronik ortamda, telefon ile veya yüzyüze)
- İnternet kanalıyla kişilere doğrudan iletişim

Deneyisel tasarım

Belirtmiş olduğumuz veri toplama metodları arasında ilk bakışta deneyisel tasarım kulağa ve göze yabancı gelmektedir. Öğrenciler genel olarak deney ortamı, deneyisel tasarım gibi kavramları görünce, haklı olarak akıllarına laboratuvarlar ve deney tüplerini vs. getirmektedirler. Bu karışıklığı netleştirmek için bir örnek ele alalım. Günümüzde hayli yaygın olan uluslararası fast-food şirketlerinin, franchising sözleşmelerinde getirdikleri çeşitli standartlar bulunmaktadır. Bu standartlar arasında patateslerin kızardığında yaklaşık olarak alması gereken renk dahi belirlenmiştir ve bu rengi tutturmak için asgari standartlar tanımlanmıştır. Bir fast food mağazasında patateslerin ortalama olarak bu standartlar doğrultusunda kızartılıp kızartılmadığına dair bir deney tasarlayalım. Bunun için öncelikle patatesler birbirlerine benzerliklerine göre

Anket tasarımı araştırma sürecinin en önemli parçalarından biridir. Araştırmacılar hazırlayacakları anketi, araştırma amaçlarının paralelinde doğru bir şekilde hazırlayabildikleri takdirde, araştırma sonuçlarında anlamlı bulgular elde edebilirler.

Internet

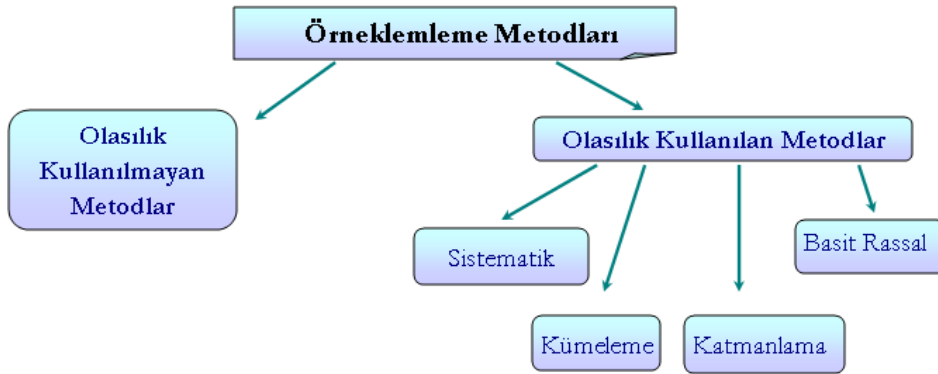
Internet sayesinde hayatımızda meydana gelen değişikliklerden biri de veri toplama alanında yaşanmaktadır. Artık tek tek insanları telefon ile aramak veya kapılarını çalarak uzun uzun anketler yapmak yerine, araştırmacılar verileri internet aracılığıyla çok daha hızlı ve kolay bir şekilde toplayabilmektedirler. Internet kullanıcıları girmiş olduklarını web sitelerini ziyaret ederken, o sitenin içeriği ile veya başka bir konu ile alakalı olan anketleri doldurabilmektedir. Internetin veri toplama literatürüne getirmiş olduğu en önemli yenilik verilerin çok daha hızlı bir şekilde veritabanları tarafından yorumlanıp sonuçların hemen incelenebilmesi ve anketlerin çok fazla sayıda insana ulaşabilmesidir. Örneğin büyük haber portallerinin açmış olduğu anketlerde bir günde 30 bin kişi oy kullanabilmektedir. Bu da çağımızın istatistiksel araştırma tekniklerine en önemli katkılarından biridir.

Bütün bu faydalarının yanı sıra bir bilgisayardan birden fazla oy atılabilmesi (artık programcılar buna da izin vermiyorlar) veya bir kişinin farklı bilgisayarlardan oy atabilmesi, programcıdan kaynaklanabilecek sistematik bir hatadan dolayı sonucu tarafında anketin yanlış yorumlanabilmesi gibi dezavantajları da bulunmaktadır.

Neden Örneklemeler?

Araştırmacılar örneklemeleri kullanarak, daha az zaman harcayarak, daha az maliyetle, daha fazla pratik yapma şansı bularak anakütle hakkında kolaylıkla fikir sahibi olabilirler.

Örneklemleme Metodları



Şekil 1.6

Örneklemleme metodları, örneklem seçilirken bilinen olasılık değerlerinin kullanılıp kullanılmamasına göre iki ana başlıkta toplanabilirler. Şekil 1.6'da da görüldüğü gibi örneklemleme metodları, olasılık kullanılan metodlar ve olasılık kullanılmayan metodlar olmak üzere iki ana başlıkta toplanmışlardır. Örneklem dağılımları konusunda örneklemleme konusu daha detaylı olarak incelenecektir.

Bölüm 2:

Temel Araçlar

Bu bölümde, istatistik dersindeki bazı temel analizleri anlamanıza yardımcı olacak temel kavramlar anlatılmaktadır. Gerekli olan temel matematiksel kavramlara ve nasıl kullanılacaklarına değinilecektir. Eğer bu kavramlara hakim olduğunuza inanıyorsanız, lütfen bu bölümü geçiniz. İspatlara ve detaylara fazla değinmeden sadece örneklerle temel kavramları inceleyeceğiz:

- Eğer tanesi 120bin Lira'dan 10 somun ekmek alırsanız 1200bin Lira ödersiniz; $(120.000.000 \times 10 = 1.200.000 \text{ TL})$
- Eğer saatte 80 kilometre hızla araba kullanırsanız, 400 kilometrelik yolu 5 saatte gidersiniz; $(400 / 80 = 5)$
- Eğer eviniz için ayda \$300, yemek için \$100, diğer şeyler için \$200 harcıyorsanız, aylık harcamalarınız ayda \$600'a $(\$300 + \$100 + \$200 = \$600)$ tekabül eder. Aynı şekilde yıllık harcamanız $600 \times 12 = \$7200$, veya haftalık harcamanız $7200 / 52 = \$138.46$ a eşit olur.
- Eğer şirketinizin geliri yılda \$1 milyon dolar ise ve yıllık maliyeti \$0.8 milyon ise, şirket yılda \$0.2 milyon $(\$1 - \$0.8 = \$0.2)$ veya \$200.000 kâr eder.

Üslü ve Köklü Sayılar

Bir sayı kendisi ile birkaç kere çarpıldığı takdirde, bu işlemi tek bir sayı olarak ifade edebiliriz. Elde etmiş olduğumuz bu sayıya üslü sayı denir. Örnek olarak, b değişkeni, kendisi ile çarpılırsa, karşımıza çıkan sonuç b değişkenin karesine eşittir ve b^2 şeklinde yazılır.

Benzer şekilde:

$$b \text{ kare} = b \times b = b^2$$

$$b \text{ küp} = b \times b \times b = b^3$$

$$b \text{ 'nin dördüncü kuvveti} = b \times b \times b \times b = b^4$$

daha genel olarak ifade edecek olursak:

$$b \text{ 'nin } n. \text{ kuvveti} = b \times b \times \dots (n \text{ kere}) = b^n$$

Eğer $b = 2$ ise

$$2 \text{ 'nin karesi} = 2 \times 2 = 2^2 = 4$$

$$2 \text{ 'nin kübü} = 2 \times 2 \times 2 = 2^3 = 8$$

$$2 \text{ 'nin dördüncü kuvveti} = 2 \times 2 \times 2 \times 2 = 2^4 = 16$$

ve genel olarak :

$$2 \text{ 'nin } n. \text{ kuvveti} = 2 \times 2 \times \dots (n \text{ kere}) = 2^n$$

Şimdi iki tane üslü sayının çarpımını inceleyim. Örneğin b^2 ile b^3 çarpılacak olursa:

$$b^2 = b \times b$$

$$b^3 = b \times b \times b$$

$$b^2 \times b^3 = (b \times b) \times (b \times b \times b) = b \times b \times b \times b \times b = b^5$$

Dolayısıyla $b^2 \times b^3 = b^5$ sonucu elde edilir. Bu işlemi genelleştirecek olursak:

$$b^m \times b^n = b^{m+n}$$

Bu yüzden $3^2 \times 3^3 = 3^5$, bu işlemi biraz daha detaylı bir şekilde inceleyelim. şöyle ki: $3^2 \times 3^3 = 9 \times 27 = 243 = 3^5$. Üslü sayı işlemlerinde yanılmamak için dikkat edilmesi gereken iki temel kural vardır:

- $b^m + b^n \neq b^{m+n}$
- $a^n + b^n \neq (a+b)^n$

Bu özellikleri $a = 1$, $b = 2$, $m = 1$ ve $n = 3$ değerlerini kullanarak kontrol edelim. Şimdi örnek olarak $b^{0.5}$ 'i alalım. Bu değeri kendisi ile çarparsak;

$$b^{0.5} \times b^{0.5} = b^{0.5+0.5} = b^1 = b$$

sonucunu elde ederiz. Kendisi ile çarpıldığında b sonucunu veren bu ifadeye ($b^{0.5}$) b 'nin karekökü denir. Dolayısıyla:

$$b^{0.5} = b^{1/2} = \sqrt{b}$$

Hatırlanacağı gibi $\sqrt{\quad}$ standart karekök işaretidir. Örneğin, $\sqrt{9} = 3$ ise $9^{0.5} \times 9^{0.5} = 3 \times 3 = 9$ eşitliği de doğrudur.

Toplama Operatörü

İstatistikte çok sık kullanılan toplama operatörü, $\sum$ sigma işareti ile gösterilir. Nasıl kullanıldığını aşağıdaki rakamları kullanarak inceleyelim :

$$5, 10, 15, 20$$

Bu rakamların bir odadaki 4 kişinin yaşlarını temsil ettiğini düşünelim. Bu değişkeni x ile ifade edelim. Bütün yaşları birbirinden ayırmamız gerektiğinden x değişkeninin altına 1 indisini yazalım (x_1). Bu indis yaş değişkenine ait serinin ilk değerini temsil eder. Bu yüzden $x_1 = 5$, $x_2 = 10$, $x_3 = 15$, $x_4 = 20$ olur. Bu değişkeni indis kullanarak daha genel bir şekilde yazabiliriz. x_i , $i = 1, 2, 3, 4$. Dolayısıyla bütün bu dört değerin toplamı aşağıdaki gibidir.

$$x_1 + x_2 + x_3 + x_4 = 5 + 10 + 15 + 20 = 50$$

Toplama Operatörü bütün bu işlemleri daha kısa ve fonksiyonel bir şekilde ifade eder.

$$\sum_{i=1}^4 x_i = x_1 + x_2 + x_3 + x_4$$

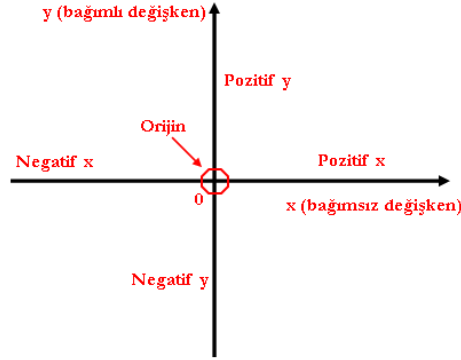
Genel olarak, eğer seride n tane rakam varsa;

$$\sum_{i=1}^n x_i = x_1 + x_2 + \dots + x_n$$

şeklinde ifade edilir.

Doğru Grafikleri

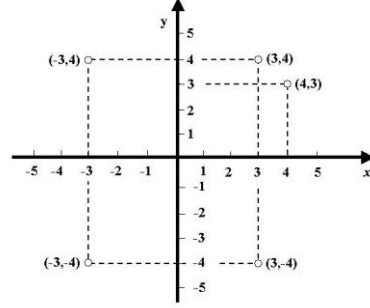
En sık karşılaşılan doğru grafikleri iki boyutlu eksenlere çizilir. Şekil 2.1’de gösterildiği gibi yatay eksen x eksenini ve dikey eksen ise y eksenini ifade eder.



Şekil 2.1

x değişkeni bizim kontrol edebildiğimiz **bağımsız bir değişken** iken, y değişkeni x değişkeni tarafından belirlenen **bağımlı bir değişkendir**. Örneğin x değişkeni reklam harcamalarını ifade ederken, y değişkeni kontrol edemediğimiz satış sonuçlarını gösterebilir; x faiz oranlarını ifade ederken, y değişkeni bankaların borçlanma oranını ifade edebilir. Eksenlerin kesiştiği ve x ile y değerlerinin sıfıra eşit olduğu noktaya **orijin** adı verilir. Orijinin üzerinde yer alan noktaların tamamında y değerleri pozitif iken altında yeralan y değerleri negatif değerlerdir. Orijinin solunda yer alan x değerleri negatif iken, sağında yer alan x değerleri pozitifdir.

Eksenlerde çizilmiş olan grafiğin üzerindeki herhangi bir noktaya koordinat adı verilir. Bu koordinatlarda yer alan ilk rakam x ekseninde yeralan rakamı ifade ederken, ikinci rakam ise y ekseninde yer alan bir rakamı ifade eder. $x = 3$, $y = 4$, noktası (3,4) şeklinde ifade edilir. Standart notasyonumuz ise (x, y) şeklindedir. Özellikle dikkat edilmesi gereken husus koordinatların yazılırken ters yazılmamasıdır. Örneğin (3,4) yerine (4,3) yazmak gibi. Bu hatayı Şekil 2.2’de inceleyebilirsiniz:

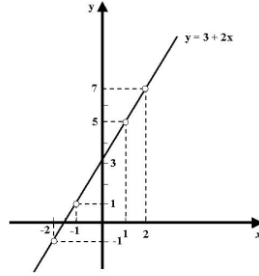


Şekil 2.2

x ekseninde yer alan noktaların koordinatları $(x; 0)$ iken y ekseninde yer alan noktaların koordinatları $(0; y)$. Orijin noktasının koordinatı ise $(0; 0)$ 'dır. Fonksiyonlar bağımlı ve bağımsız değişkenler arasındaki ilişkiyi veren matematiksel ifadelerdir. Örneğin; $y = 3+2x$ fonksiyonu x ile y arasındaki ilişkiyi gösteren matematiksel bir ifadedir. Bu fonksiyona göre x , 1 değerini aldığı anda, bağımlı değişken olan y 'nin $(y = 3+2(1) = 5)$ olduğundan) 5 değerini alması gerekir. Benzer şekilde

- $x = -2$ iken; $y = (y = 3+2(-2) = -1)$
- $x = -1$ iken; $y = (y = 3+2(-1) = 1)$
- $x = 0$ iken; $y = (y = 3+2(0) = 3)$
- $x = 2$ iken; $y = (y = 3+2(2) = 7)$

Bu yüzden, fonksiyonumuzu şekil 2.3'te olduğu gibi çizebiliriz.



Şekil 2.3

Doğru üzerinde her nokta bir y değeri ile eşlenen bir x değerini içerir.

Bölüm 3: Tablo ve Grafiklerle Veri Sunumu (Sayısal Veriler)

Birinci bölümde, verileri, veri toplama tekniklerini incelemiştik. Ancak, genellikle ham (hiçbir işlemten geçirilmemiş, toplandığı şekilde duran) verilerden bilgi edinmek zordur. Bu zorluğun nedeni verilerin genelde büyük miktarda ve karmaşık olmasından kaynaklanmaktadır. Verilerin düzenlenmesi, etkin bilgi edinme sürecini kolaylaştıracaktır. Tablo ve grafiklerden faydalanmak, verileri düzenlemenin en kolay ve en genel yoludur. Bu bölümde, tablo ve grafikleri hazırlama süreci ve nasıl yorumlanması gerektiği üzerinde duracağız.

Sayısal Veri Düzenlemesi

Kullanacağımız tablo ve grafikler, elimizdeki verilerin özelliklerine bağlı olarak değişir. Sayısal verilerin nasıl düzenlendiğini ve yorumlandığını inceledikten sonra dördüncü bölümde, nitelik belirten verilerin (kategorik veriler) nasıl düzenlendiğini ve yorumlandığını inceleyeceğiz.

Örnek olarak kendimize, bir üniversitenin kayıt işlerinde birbirini takip eden 13 günde harcanan kağıt miktarını gösteren bir verisetini ele alalım.

Tablo 3.1: Bir üniversitenin kayıt işlerinde, birbirini takip eden 13 günde harcanan kağıt miktarı

145	141	139	140	145	141	142	131	142	140	144	140	138
-----	-----	-----	-----	-----	-----	-----	-----	-----	-----	-----	-----	-----

İlk bakışta veriler, çok bilgilendirici bir veriseti gibi durmamaktadır. Dolayısıyla bu verisetini biraz daha kullanılabılır hale getirmemiz gerekmektedir.

Sıralı Dizi

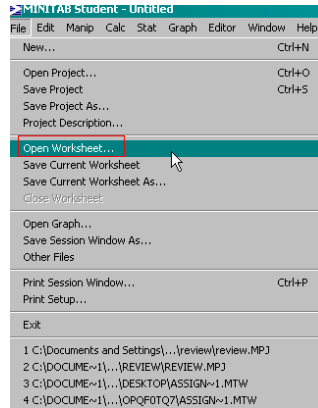
Verisetini daha bilgilendirici hale getirmek için kullandığımız en temel araçlardan biri verileri küçükten büyüğe doğru sıraladığımız sıralı dizi metodudur. Tablo 3.1’de yer alan verileri sıralayacak olursak:

Tablo 3.2 Bir üniversitenin kayıt işlerinde, birbirini takip eden 13 günde harcanan kağıt miktarının sıralı dizisi

131	138	139	140	140	140	141	141	142	142	144	145	145
-----	-----	-----	-----	-----	-----	-----	-----	-----	-----	-----	-----	-----

Sıralı Dizi Minitab Uygulamaları

Bu bölümde, sayısal verilerle oluşturulabilecek değişik tablo ve grafikleri anlatırken “*getiri1*” ismini verdiğimiz verisetini kullanacağız. “*getiri1.mtw*” dosyasını MINITAB programında açmak için, bilgisayarınızda programı çalıştırdıktan sonra aşağıda görüldüğü gibi “File”(Dosya) menüsüne girip “Open Worksheet” (Çalışma Sayfası Aç) komutunu seçiniz, açılan pencereden dosyayı yüklediğiniz klasöre girip dosyayı seçiniz. Bilgisayarınızın Windows gezgini menüsünden, yüklediğiniz dosyayı bulup üzerine çift tıklayarak da dosyayı açabilirsiniz.



Şekil 3.1

“Open Worksheet” (Çalışma Sayfası Aç) komutu sonucunda “*getiri1.mtw*” dosyasını açtığınızda ekranda şekil 3.2’de yeralan ekran görüntüsü çıkmaktadır.

	C1-T	C2	C3	C4	C5	C6	C7	C8	C9	C10
	Funds	1YearRet.1								
1	Berger SmCoGro	29.9								
2	Junk & Voytes Kaufman	26.9								
3	Lord Abbott A DevGro	39.0								
4	Freemont Growth	33.8								
5	William Blair Growth	26.6								
6	Enterprise Grow A	44.7								

Şekil 3.2

Bu dosyada ABD finansal piyasalarında işlem gören 59 kamu fonunun 1 yıllık getiri oranları yer almaktadır. Fonların isimleri C1-T sütununda, yüzdelik getirileri ise C2 sütununda yer almaktadır.

Böyle bir veriseti ile neden ilgilenmekteyiz? Potansiyel yatırımcılara fon alım-satım tavsiyeleri veren bir aracı kurumda çalıştığınızı ve görevinizin de bu fonların kârlılıklarını araştırmak, yatırım olanaklarını değerlendirmek olduğunu varsayalım. Müşterilerinize akılcı tavsiyeler verebilmek için fonların 1 yıllık getirilerini, diğer yatırım imkanları ile mukayese ederek değerlendirmeniz gerekmektedir. En temel değerlendirme, fonların getirilerinin büyükten küçüğe doğru sıralamaktır.

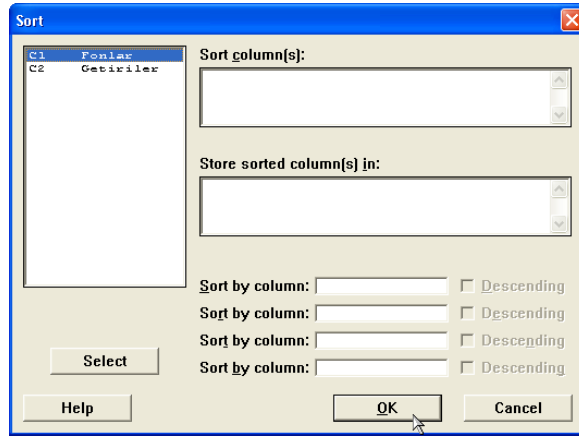
MINITAB programı bu sıralamayı sizin için hızlı ve kolay bir şekilde yapabilir. Örnek olarak “*getiri1.mtw*” dosyasındaki verileri sıralayalım.

Elinizdeki verileri sıralamadan önce, sıralı dizilerin yeracağı sütunlara birer isim vermelisiniz. Bu işlemi gerçekleştirmek için, “**Getiri**” yazılı hücrenin hemen sağındaki hücreye tıklayıp bu sütunun ismini “*SIRA - Fon*”, onun sağındaki de “*SIRA - Getiri*” olarak girin. Fonların getirileri sıralandığında, bu sütunlarda yeracaklardır.(bkz Şekil 3.3)

GETIRI.MTW ***				
	C1-T	C2	C3	C4
↓	Fonlar	Getiriler	SIRA Fon	SIRA Getiri
1	Berger SmCoGrow	29,9		
2	Jurika & Voyles Kaufmann	28,9		
3	Lord Abbett A DevGrow	39,0		
4	Freemont Growth	33,8		
5	William Blair Growth	25,6		
6	Enterprise Grow A	44,7		
7	Victory GrowStk	32,8		
8	1784 GrowInc	33,0		
9	Janus Mercury	27,6		
10	Vanguard Index Growth	32,9		
11	MFS B Research	31,6		
12	MFS B ResearchC	31,6		

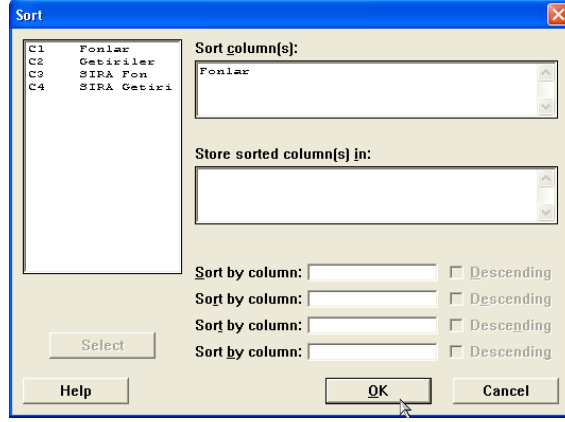
Şekil 3.3

Şimdi ise “Manip” (çevrim) menüsüne girip “Sort”(sırala) komutuna tıkladığınızda ise Şekil 3.4’teki gibi bir diyalog ekranı ile karşılaşacaksınız.



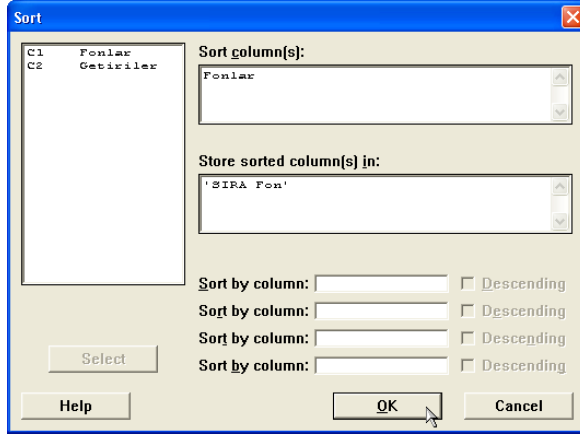
Şekil 3.4

Şekil 3.4’te yeralan diyalog ekranında, “Sort Column(s)” (sıralanacak sütunlar) kısmına tıklayınız ve sol tarafta değişkenlerin bulunduğu kutucuktan verilerin yer aldığı “C1 Fonlar” yazısının üzerine tıklayıp aşağıdaki “Select” (seç) tuşuna basınız. Aynı işlemleri “C2 Getiriler” için de gerçekleştiriniz.



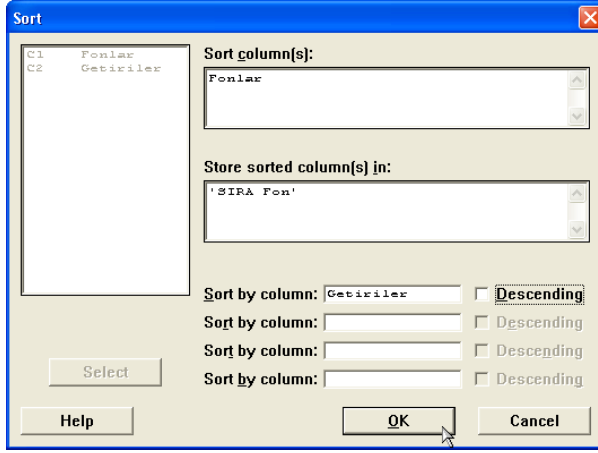
Şekil 3.5

Bu işlemten sonra, “*Store sorted column(s)*” (sırayı aşağıdaki sütun(lar)da oluştur) kısmına tıklayıp, yine solda göreceğiniz “*SIRA - Fonlar*” sütununu seçip “*Select*” komutuna tıklayınız, aynı işlemi “*SIRA - Getiriler*” için de tekrarlayınız.



Şekil 3.6

En son olarak da “*Sort by column*” (sıralama yapılacak sütun) kısmına tıklayıp, soldaki pencereden “*C2 Getiriler*”i seçip “*Select*”e basınız.



Şekil 3.7

“OK” tuşuna tıkladığınızda, veriler Tablo 3.3’teki gibi sıralanacaklardır. Buradaki sıralı dizilerin bulunduğu dosya sizin için **“getiri2.mtw”** ismiyle kaydedilmiştir.

Tablo 3.3 ABD finansal piyasalarında işlem gören 59 kamu fonunun 1 yıllık getiri oranlarının sıralı dizisi (%)

20.4	23.8	25.6	26.2	27.6	27.7	28.3	28.6	28.8	28.9	28.9	29.3
29.3	29.5	29.9	30.1	31.5	31.6	31.6	31.8	31.9	32.1	32.3	32.3
32.4	32.8	32.9	32.9	33.0	33.3	33.4	33.7	33.8	34.0	34.0	34.3
34.7	34.7	34.8	35.0	38.2	39.0	39.4	40.7	41.1	42.8	42.9	43.3
43.4	43.5	43.6	43.7	44.6	44.7	45.4	45.7	46.6	48.0	48.6	

Sayısal Verilerin Tablo Haline Getirilmesi

İstatistiksel araştırma süreçlerinde, araştırmacıların öncelikle karşılaştıkları ilk sorun ham verilerin işlenmesi ve verimli bir hale getirilmesidir. Bu süreçte, verilerin sıralanması, frekans dağılım tabloları, grafikler, ve özet tablolar kullanılan araçların başında gelirler. Verisetlerinin uygun bir şekilde özetlenmesinde en çok faydalandığımız araç frekans tablolarıdır. Frekans tablolarının nasıl oluşturulduğunu ve temel özelliklerini incelemek için öncelikle Tablo 3.1’de yeralan verisetimizi 13 günden 106 güne genişletelim. Yeni verisetimiz Tablo 3.4’teki gibidir:

Tablo 3.4 Bir üniversitenin kayıt işlerinde birbirini takip eden 106 günde harcanan kağıt miktarı

146 141 139 140 145 141 142 131 142 140 144 140 138

139 147 139 141 137 141 132 140 140 141 143 134 146
134 142 133 149 140 143 143 149 136 141 143 143 141
140 138 136 138 144 136 145 143 137 142 146 140 148
140 140 139 139 144 138 146 153 148 142 133 140 141
145 148 139 136 141 140 139 148 135 132 148 142 145
145 121 129 143 148 138 149 146 141 142 144 137 153
148 144 138 150 148 138 145 145 142 143 143 148 141
145 141

Tablo 3.4'te karşımıza çıkan 106 sayıdan oluşan liste yine bizim için çok bilgilendirici bir boyutta değildir. Bu sayıları sıraladığımızda da sorunun çok fazla giderilmediğine tanık oluyoruz. Bu farklı değerlerden hala bir eğilim, bir dağılım çıkaramamaktayız. Veriyi yorumlanabilecek hale getirebilmek için öncelikle gruplamaya çalışalım. 106 günde kullanılan kağıt miktarlarını 120 den 160'a kadar, gruplarsak, verilerin genelde 130 ile 150 kağıt grubu aralığında toplandıklarını görebiliriz. Benzer bir şekilde, tablo 3.3'te yeralan sıralanmış getiri oranları da çok fazla bilgilendirici bir konumda değildirler ve bu verilerinde gruplanıp frekans tablolarının oluşturulması gerekmektedir.

Frekans Dağılımı Tablosu

Frekans dağılımı tablolarının detaylarını incelemek için, Tablo 3.2'de yeralan bir üniversitenin kayıt işlerinde birbirini takip eden 13 günde harcanan kağıt miktarını gösteren verisetini tekrar ele alalım. gözlem sayısı küçük olan verisetlerinin, frekans dağılım tablolarını çıkarmak çok faydalı değildir. Fakat prosedürümüzün kolay anlaşılması için 13 gözlemden oluşan bu verisetini inceleyeceğiz.

Öncelikle, gözlemleri gruplara ayırmamız gerekmektedir. Verileri gruplara ayırdıktan sonra, bu gruplarda yer alan gözlem sayısının hesaplanması gerekmektedir. (ki bu gözlem sayısını frekans olarak adlandıracğız) Örneğin, Tablo 3.2'deki verisetimizi her biri 3 kağıt miktarı genişliğinde olan 5 farklı gruba ayırabiliriz. Gruplar, 131'den fazla 134'ten az, 134'ten fazla 137'den az şeklinde devam etmektedir. Birinci grupta sadece bir gözlem yeralmaktadır. Bir başka ifade ile, 131 ile 134 aralığında yer alan bir adet gözlem yeralmaktadır, bu grubun frekansı 1'dir. İkinci grupta yeralan hiçbir gözlem bulunmamasına rağmen, üçüncü grupta ise iki tane gözlem bulunmaktadır. Tablo 3.5'te oluşturduğumuz frekans dağılım tablosu detaylı bir şekilde incelenebilir.

Tablo 3.5: Özel bir üniversitenin kayıt işlerinde, birbirini takip eden 13 günde harcanan kağıt miktarının frekans tablosu

13 Günde harcanan kağıt miktarı	Frekans (Kağıt Miktarı)
131 – 134 arası (134'ten az)	1
134 – 137 arası (137'ten az)	0
137 – 140 arası (140'tan az)	2
140 – 143 arası (143'ten az)	7

143 – 146 arası (146'dan az)	3
Toplam	13

Frekans dağılım tabloları, veriyi sıralayarak birbirinden farklı gruplara ayırır ve eğilimin nerede olduğuna dair fikir verir. Peki en uygun gruplama nasıl yapılmalıdır? Veriler kaç tane farklı gruba ayrılmalıdır? Grupların (veya sınıfların) genişliği (grup aralıklarının boyutu) ne olmalıdır? Yukarıdaki örneğimizde, verileri beş farklı gruba ayırdık ve grup genişliğini ise 3 olarak seçtik. (134 -131 =3).

Bu seçimlerimizi gerçekleştirirken aşağıdaki kurallardan faydalanabiliriz:

- **Aralık, grup veya sınıfların seçimi:** Grup sayıları gözlem sayıları ile doğru orantılıdır. Gözlem sayısı artarsa, grup sayısında artması gerekmektedir. Genelde bir frekans dağılımında en az beş, en fazla yirmi grup yer alır. Dikkat edilmesi gereken bir husus, çok fazla veya çok az grup sayısına sahip olan bir frekans dağılımı, verisetinin genel şekli hakkında bir fikir sahibi olmanın zorluğudur. Grup sayısının belirlenmesi ile alakalı önemli kurallardan biride **Sturges Kuralı**'dır. Bu kural grup sayısının belirlenmesinde kullanılmaktadır.

$$\text{Grup sayısı} = 1 + \log_2(n)$$

n, verisetinde yeralan gözlem sayısını göstermektedir. Örneğin n = 13 ise bu formüle göre¹

$$1 + \log_2(13) = 1 + 3.7 = 5$$

Bu kural kesin olarak grup sayısını belirlememekte, araştırmacılara grup sayılarının belirlenmesi konusunda fikir vermektedir.

- **Grup aralıklarının belirlenmesi:** Grup sayısını belirledikten ve planladıktan sonra aşağıdaki formülü kullanarak grup genişliklerini belirleyebiliriz:

$$\text{Grup Genişliği} = \frac{\text{En büyük gözlem} - \text{En küçük gözlem}}{\text{Planlanan Grup Sayısı}}$$

Grup sayısını bir önceki adımda 5 olarak belirlediğimizden

$$\text{Grup Genişliği} = \frac{14}{5} = 2.8$$

Sonucun tam sayı çıkmaması durumunda, bulduğumuz değeri en yakın tamsayıya yuvarlayarak kullanırız. Örneğimizde de 2.8 değerinin farkını almakla uğraşmak yerine bu değeri yuvarlayarak grup genişliğini 3 olarak aldık. Artık en küçük gözleme 3 ekleyerek (131+3=134 dolayısıyla ilk grup 131-134 şeklinde olur) 5 grubu oluşturabiliriz.

Şimdi ise Tablo 3.3'teki verileri ele alalım. Grup sayısını 6 olarak belirleyelim. (Strunger kuralı bize 7 grubu önermesine rağmen: $1 + \log_2(59) = 1 + 5.88 = 7$). Dolayısıyla grup genişliği:

$$\text{Grup genişliği} = \frac{48.6 - 20.4}{6} = 4.7$$

¹ İşlemin sonucu yuvarlanmıştır.

Grup genişliğini 4.7 değerini yuvarlayarak 5 olarak ele alabiliriz. Dolayısıyla en küçük gözlem olan 20'den başlayarak ve 5 ekleyerek grupları oluşturabiliriz. Bu grupları Tablo 3.6'da görebiliriz:

Tablo 3.6: ABD finansal piyasalarında işlem gören 59 kamu fonunun 1 yıllık getiri oranlarının frekans dağılımı

1-Yıllık Toplam Getiri	Frekans (Fonların Sayısı)
%20-%25 arası (%25'ten az)	2
%25-%30 arası (%30'dan az)	13
%30-%35 arası (%35'ten az)	24
%35-%40 arası (%40'dan az)	4
%40-%45 arası (%45'ten az)	11
%45-%50 arası (%50'den az)	5
Toplam	59

Yüzdelik Frekans Dağılımı

Yüzdelik frekans dağılımı tablosu, her grupta, toplam veri sayısının yüzde kaçının bulunduğunu gösterir. Mutlak sayılardan ziyade, yüzdelik rakamlar daha kolay idrak edilebilir. Ayrıca, iki veya daha fazla veri kümesini karşılaştırma işlemi de yüzdelik dağılımlarla daha kolay gerçekleştirilebilir.

Yüzdelik frekans değerlerini hesaplamak için, her grubun frekansını toplam frekansa oranlamamız gerekmektedir. Örneğin %20 ile %25 arasında getiri sağlayan fonların yüzdelik frekansını bulmak için o gruba ait olan frekans değeri olan 2'nin toplam frekansa olan 59'a oranını bulmaz gerekmektedir. $2/59=0,034$ veya %3,4. Bu işlemi her grup için tekrarlayarak, Tablo 3.7'deki Yüzdelik Frekans Dağılımı Tablosu'nu elde edebiliriz.

Tablo 3.7: Getirilerin Yüzdelik Frekans Dağılımı Tablosu

1 Yıllık Getiri	Frekans	% Hesap	% Frekans
%20-%25 arası (%25'ten az)	2	$(2/59) \times 100$	%3,4
%25-%30 arası (%30'dan az)	13	$(13/59) \times 100$	%22,03
%30-%35 arası (%35'ten az)	24	$(24/59) \times 100$	%40,7
%35-%40 arası (%40'dan az)	4	$(4/59) \times 100$	%6,8
%40-%45 arası (%45'ten az)	11	$(11/59) \times 100$	%18,6
%45-%50 arası (%50'den az)	5	$(5/59) \times 100$	%8,5
Toplam	59	1.000(%100)	1.000(%100)

Açıkça görüldüğü gibi, verilerin büyük bir kısmı %30 ile %35 arasında yoğunlaşmıştır (%40'ı). Ayrıca verilerin %60'ından fazlası da %25 ile %35 arasında yer almaktadır. Bu tabloda yer alan değerler, 100 ile çarpılmadan da oluşturulabilir (%3,4 yerine 0,034 sayısının yeralması gibi).

Kümülatif Frekans Dağılımı

Frekans dağılımlarını göstermek için bir başka kullanışlı yöntem de kümülatif frekans dağılımıdır. Bu dağılım, frekans dağılımından elde edilir ve her bir grubun kümülatif frekansı, ondan önceki grupların frekanslarını da içermektedir. Tablo 3.8, kümülatif frekans hesaplanmasını ve bunun sonucunda ortaya çıkan kümülatif frekans dağılımını göstermektedir.

Tablo 3.8: Getirilerin Kümülatif Frekans Dağılım Tablosu:

1 Yıllık Getiri	Frekans	Kümülatif Hesap	Küm.Frekans
%20-%25 arası (%25'ten az)	2	2	2
%25-%30 arası (%30'dan az)	13	15	2+13
%30-%35 arası (%35'ten az)	24	39	15+24
%35-%40 arası (%40'dan az)	4	43	39+4
%40-%45 arası (%45'ten az)	11	54	43+11
%45-%50 arası (%50'den az)	5	59	54+5
Toplam	59	59	59

Bu tablo, yüzdelik frekans dağılımdan da yararlanılarak oluşturulabilir:

Tablo 3.9 Getirilerin Yüzdelik Kümülatif Frekans Dağılım Tablosu:

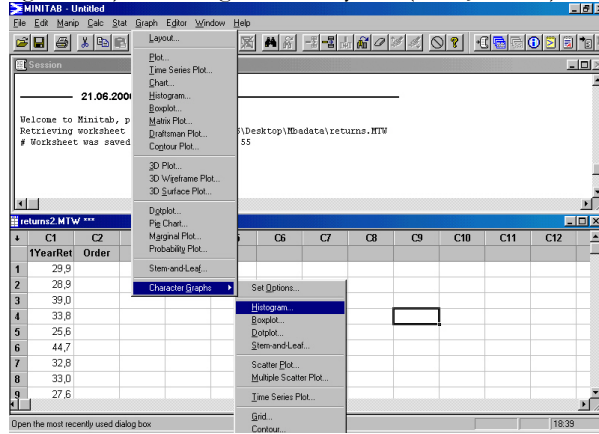
1 Yıllık Getiri	% Frekans	Kümülatif Hesap	% Küm.Frekans
%20-%25 arası (%25'ten az)	0.034 (3.4%)	0.034 (3.4%)	
%25-%30 arası (%30'dan az)	0.220 (22.0%)	0.254 (25.4%)	0.220+0.034
%30-%35 arası (%35'ten az)	0.407 (40.7%)	0.661 (66.1%)	0.254+0.407
%35-%40 arası (%40'dan az)	0.068 (6.8%)	0.729(72.9%)	0.661+0.068
%40-%45 arası (%45'ten az)	0.186 (18.6%)	0.915 (91.5%)	0.729+0.186
%45-%50 arası (%50'den az)	0.085 (8.5%)	1.000 (100%)	0.915+0.085
Toplam	1.000 (100%)		

Buradan hareket ederek, fonların %66'sının %35'in altında getiri sağladığını, %40'tan fazla getiri sağlayan fonların toplamın ancak %27'si kadar olduğunu, ayrıca fonların 4'te 1'inin de %30'un altında getiri sağladığını kolaylıkla öğrenebiliriz

Frekans Dağılımı

Minitab Uygulamaları

İncelemiş olduğumuz yüzdelik frekans ve kümülatif frekans dağılımlarını Minitab ortamında çok daha kolay bir şekilde oluşturabiliriz. Öncelikle, fon getirilerinin sıralı ve son halinin kaydedildiği “*getiri2.mtw*” dosyasını açınız. Daha sonra “*graphs*” (grafikler) menüsünden “*character graphs*” (frekans grafikleri) ve “*Histogram*”a tıklayınız. (bkz. Şekil 3.8)

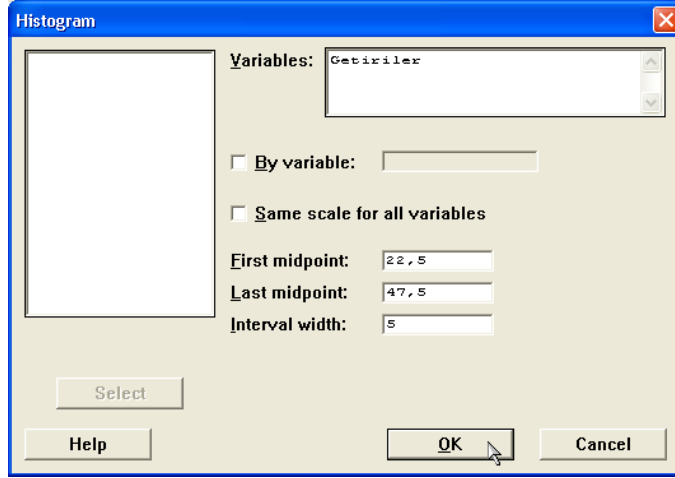


Şekil 3.8

Her zamanki verileri girmek için yine bir diyalog ekranı ile karşı karşıyayız. Bu pencerede “*Variables*” bölümünden, soldan, “*C2 Getiriler*”i tıklayıp “*Select*”e basarak frekans dağılım tablosunda yer alacak değerleri seçmiş olacağız.

MINITAB bizden bu verilerin yer alacağı grupların “*midpoint*” (orta noktalarını) isteyecektir. İlk grup aralığımız %20 - %25 olduğundan, “*First Midpoint*” (-ilk grubun orta noktası) kısmına, bu aralığın orta noktası olan 22.5'i yazmalıyız.

Bunun hemen altında yer alan “*Last Midpoint*” (son grubun orta noktası) kısmına da aynı şekilde 47.5 değerini girmeliyiz. “*Interval width*” (grup genişliği) bize grup genişliğini ifade etmektedir. Yaptığımız işlemleri Şekil 3.9'da yer alan diyalog ekranından kontrol edebilirsiniz.



Şekil 3.9

“OK” tuşuna bastığınızda Şekil 3.10’daki frekans dağılım tablosunu elde edebilirsiniz.

Histogram		
Histogram of Getiriler N = 59		
Orta Noktalar	Frekans	
22,50	2	**
27,50	13	*****
32,50	24	*****
37,50	4	****
42,50	11	*****
47,50	5	*****

Şekil 3.10

Elde ettiğimiz frekans dağılım tablosunun Tablo 3.6’da oluşturduğumuz frekans dağılım tablosu ile aynı olduğu açıkça gözükmemektedir. Burada farklı olan husus Minitab grup aralıkları yazmak yerine grupları orta noktaları ile ifade etmiştir.

Aklımıza hemen “Grup aralıkları her zaman aynı mı olmalıdır?” sorusu gelebilir. Cevabımız ise mümkün olduğunca aynı olmalıdır olacaktır. Farklı grup aralıkları okunmayı zorlaştırabileceği gibi, buradaki verileri grafiksel olarak sunmak istediğimizde de problem yaratacaktır.

Farklı grup aralıkları, ancak tamamen boş veya çok az sayıda gözlem içeren gruplar olduğunda ve bu grupların birleştirilmesinde kullanılmalıdır.

Sayısal Verilerin Grafiksel Yöntemlerle Anlatımı

Verilerin tablo haline getirilmesi her ne kadar bilgi elde etmeyi kolaylaştırırsa da, dikkat çekmek için, bu bilgileri başkalarına aktarabilmek için, ya da kısıtlı zaman içerisinde çok miktarda veriyi inceleyip yorum yapabilmek için grafiksel yöntemler ideal olacaktır.

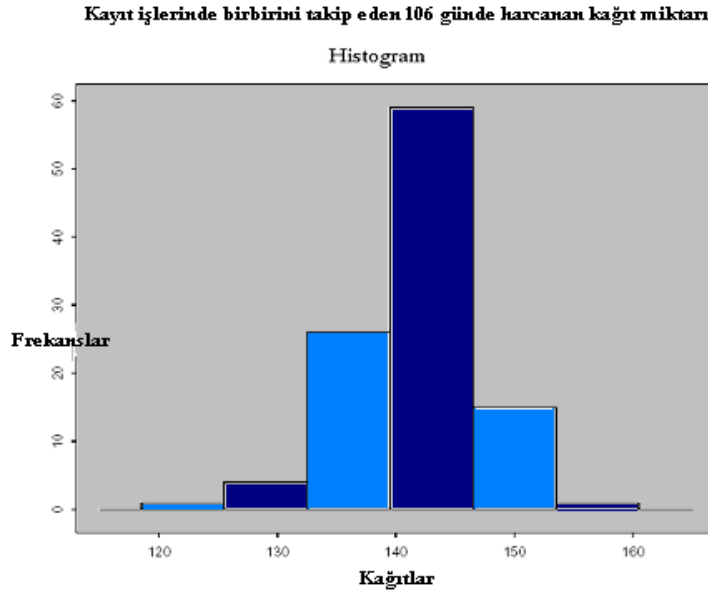
Histogram

Frekans dağılımının gösteriminde en yaygın olarak kullanılan yöntem histogramdır. Gruplara ayrılmış veriler, frekanslarının yoğunluğuna göre değişen yükseklikte, birbirlerine bitişik sütunlarla ifade edilirler.

Grafığın yatay X ekseninde, verilerin yer aldığı grupların sınırları veya orta noktaları işaretlenir. Grup genişliği, bu grubun frekansını gösteren sütunun genişliğiyle ifade edilir. Dikey Y ekseninde ise frekans değerleri yer alır ve her grubun frekans değerine göre sütun bu sayıya kadar yükseltilir.

Histogramı frekans dağılımının grafiksel gösterimi olarak ifade edebiliriz. Önemli olan hususlardan biride sütunların, grup sınırlarına ve frekanslara göre çizildiğidir. Burada esas alınan gruplar ve değerlerdir, kategoriler değildir. Bu yönü ile histogram, Bölüm 4'te işleyeceğimiz bar grafikten ayrılır.

Tablo 3.4'te yer alan ve özel bir üniversitenin kayıt işlerinde birbirini takip eden 106 günde harcanan kağıt miktarını gösteren verisetine ait histogramımız Şekil 3.11'deki gibidir.

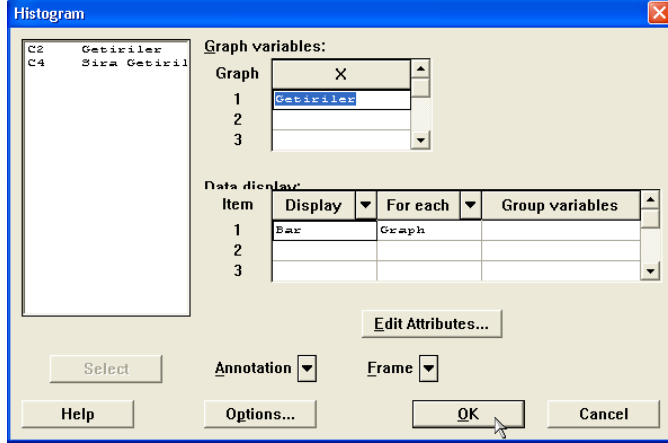


Şekil 3.11

Histogram

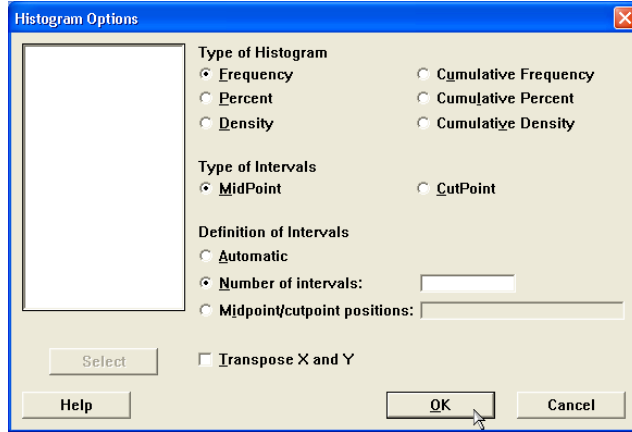
Minitab Uygulamaları

Minitab ortamında bir histogram çizerek, histogramı daha detaylı inceleyelim. Bu inceleme için veriseti olarak tekrar ABD finansal piyasalarında işlem gören 59 kamu fonunun 1 yıllık getiri oranlarını kullanacağız ve son olarak sıralı hallerini de kaydettiğimiz “*getiri2.mtw*” dosyasını açınız ve daha sonra “*Graph*” (grafikler) menüsünde yer alan “*Histogram*” seçeneğine tıklayınız.



Şekil 3.12

Karşımıza çıkan diyalog ekranında, “*Graph Variables*”(grafığı çizilecek değişken) için “*C2 Getiriler*” sütununu seçip “*Select*”e tıklayınız. Histogramın formatında daha detaylı ayarlar yapmak için de “*Options*” (seçenekler) tuşuna tıklayınız.

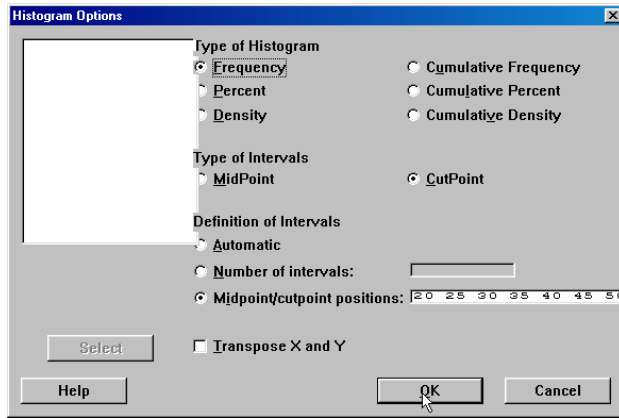


Şekil 3.13

Bu diyalog ekranında yeralan:

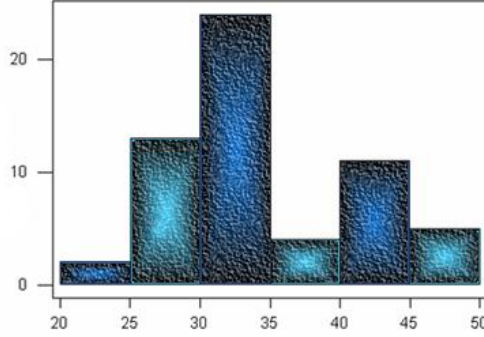
1. “Frequency” seçeneğini seçerek Frekansları,
2. “Percentage” seçeneğini seçerek yüzdelik frekansları veya
3. “Cumulative” seçeneklerini seçerek kümülatif frekansları veya kümülatif yüzdelik frekansları grafik haline getirebiliriz.

Bu seçimden sonra grupları ayırmamız gerekmektedir. Grupları ayırmak için 20, 25, 30 gibi grupların sınırlarını kullanacaksa “Type of Intervals” (aralık türü) kısmında “Cutpoint” (kesim noktası)’nı; 22.5, 27.5, 32.5 gibi grupların orta noktalarını kullanacaksa “Midpoint”(orta nokta)’i seçmemiz gerekmektedir. Son olarak ise Minitab aralıkları nasıl tanımlamak istediğimizi sormaktadır. “Automatic” seçeneğini işaretlediğimizde program kendisine göre aralık sayısını oluşturacaktır. “Number of Intervals” (aralık sayısı) seçeneğini işaretlediğimizde ise Minitab aralık sayısını bizim belirlememize izin vermektedir. “Midpoint/Cutpoint positions” (kesim noktaları veya orta noktalar) seçeneğini işaretlediğimizde ise daha önceden oluşturduğumuz frekans tablosunda yer alan orta noktaları kullanarak histogram grafiğini çizdirebiliriz. Şimdi “Midpoint/Cutpoint positions” seçeneğini işaretleyeliniz ve ilgili sayıları arada birer boşluk bırakarak giriniz. “OK” tuşuna bastığımızda, bir önceki pencereye geri döneceğiz.



Şekil 3.14

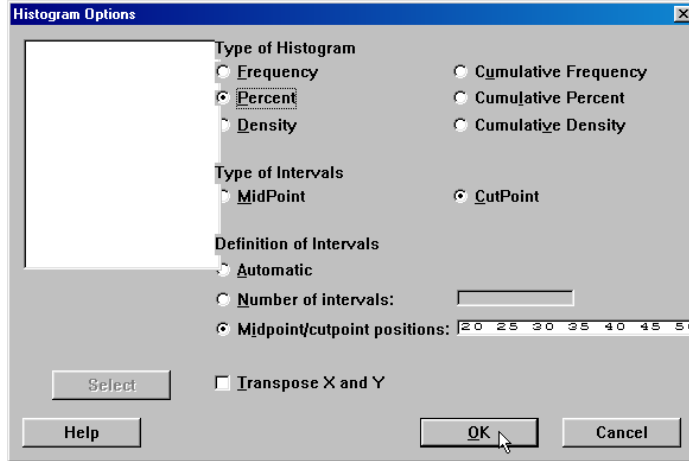
Histogram penceresinde de “OK” tuşuna tıkladığımızda, Şekil 3.15’te yer alan ve grupların frekans dağılımını gösteren histogramı elde edebiliriz.



Şekil 3.15

Histograma bakarak, %20 ile %25 arasında az sayıda getiri olduğunu (2 getiri), fonların büyük çoğunluğunun getirisinin %25 ile %35 arasında olduğunu (13+24=37), 16 fonun getirisinin de %40'tan yüksek olduğunu kısa zamanda kavrayabiliyoruz. Sütunların yükseklikleri, o gruptaki fon sayısını göstermektedir.

Siz de bu histogramı elde ettikten sonra, alıştırma olarak “Options” penceresinde yer alan diğer seçeneklerle de değişik tipte histogramlar oluşturabilirsiniz.



Şekil 3.16

Dal ve Yaprak Düzeni

Histogramlar, verisetinin özetlenmesinde, dağılımın gözlenmesinde çok faydalı olmalarına rağmen, önemli bir dezavantajı da barındırmaktadır. Histogramlarda verisetinde yer alan orjinal veri gözlenemez. Sadece verilerin gruplanmasından oluşan frekanslar gözlenebilmektedir. İstatistikçiler, bu dezavantajı ortadan kaldırmak için yeni bir metod geliştirdiler. Dal ve yaprak düzeni adı verilen bu metod, dağılımı histogram gibi gösterirken, orjinal verilere erişimi de sağlamaktadır. Dal ve yaprak düzeni, “dallar” (ana basamaklar) ve “yapraklar” (yardımcı basamaklar) olmak üzere verisetini ikiye ayırmaktadır.

Dal ve yaprak düzenini oluşturmak için:

- Verileri iki kısma böleriz. Birinci kısımda verilerin yüksek derecedeki basamakları yer alırken (ki basamaklara dallar diyoruz), ikinci kısımda ise geri kalan basamaklar yer almaktadır. (bu değerlere de yapraklar diyoruz)
- Sayı doğrusu sırasına göre yukarıdan aşağıya dalları yazıyoruz.
- Her dal değerinin yanına, o dal değeri ile başlayan bütün yaprak değerlerini yazıyoruz.

Şimdi nereden çıktı bu dallar ve yapraklar dememeniz için, anlattıklarımızı bir örnekle pekiştirelim. Aşağıdaki veriler bir mağazada çalışan satış memurlarının yaşlarına aittir.

18 21 22 19 32 33 40 41 56 57 74 28 29 29 38 39 17 75 76

Yaşları sıraladığımızda dizimiz aşağıdaki hali alır.

17 18 19 21 22 28 29 29 32 33 38 39 40 41 56 57 74 75 76

Örnek olarak, birinci gözlemin onlar basamağında yer alan 1 değerini dal değeri, birler basamağında yer alan 7 değerini ise yaprak değeri olarak düşünebiliriz. Aynı yöntemi izleyerek, ikinci gözlemin dal değeri 1 iken yaprak değeri de 8'e eşittir. Dolayısıyla onlar basamağını dal değeri, birler basamağını ise yaprak değeri olarak kullandık. Dal ve yaprak düzenimizi oluşturacak olursak:

Leaf Unit =1.0 (Yaprak Birimi)

3	1	7 8 9
8	2	1 2 8 9 9
(4)	3	2 3 8 9
7	4	0 1
5	5	6 7
3	6	
3	7	4 5 6

Dal ve yaprak düzenimizde yeralan ikinci sütun dal değerlerini temsil ederken, üçüncü sütunda yer alan veriler de yaprak değerlerini göstermektedir. Birinci satırı ele alacak olursak (birinci sütunu dikkate almayınız), bu satır verisetimizde 17, 18 ve 19 olmak üzere üç rakamın olduğunu ifade etmektedir. Aynı şekilde ikinci satır da bize, verisetinde 21, 22, 28, 29, ve 29 (tekrar) olmak üzere beş rakamın daha olduğunu ifade etmektedir.

Dal ve yaprak düzenlerinde en çok karışıklık yaratan hususların başında yaprak birimi gelmektedir. Yukarıda oluşturduğumuz yaprak düzeninde de görebileceğiniz üzere yaprak birimi 1 olarak ele alınmıştır. Yaprak biriminin 1 olması, yaprakların değerlerinin yazıldığı gibi okunması anlamına gelmektedir. Örneğin, 8 değeri 8 diye okunmaktadır. Fakat, eğer yaprak birimi 0.1 olsa

idi (bir sonraki örnekte görebileceğiniz gibi), 8 rakamı bu sefer 0.8 olarak okunacak idi. Dolayısıyla verisetimizde yer alan ilk rakam 1.7 ikinci rakam ise 1.8 olacaktı.

Birinci sütünde yer alan sayılar, o dala kadar kümülatif olarak data setimizde kaç tane sayı olduğunu göstermektedir. Örneğin 2 dalına ait olan 8 sayısı data setinde bu dala kadar sekiz tane rakamın olduğunu söylemektedir (bu rakamlar 7, 18, 19, 21, 22, 28, 29, ve 29 dur). Bu ortanca dala, (örneğimizde 3) kadar bu şekilde devam etmektedir. Ortada yer alan dal değerine gelince ise sayma işlemini bir kenara bırakıp sadece o dala ait olan kaç yaprak değeri olduğunuz tespit ederiz ve bu değeri paranteze alırız. Örneğimizde en ortadaki dalda 32, 33, 38, 39 olmak üzere 4 yaprak yer aldığından, 4 parantez içine alınmıştır. Bu dal değerinden sonra sayma işlemine tekrar devam ederiz. Fakat bir öncekinden farklı olarak aşağıdan ortaya doğru sayma işlemini gerçekleştiririz.

Peki dal ve yaprak düzeni araştırmalarımızda bize nasıl faydalı olmaktadır? Genel olarak dal ve yaprak düzeni, verilerin organize edilmesinde, nasıl dağıldığının gözlenmesinde çok faydalı bir araçtır. Örneğimizden yola çıkacak olursak, satış elemanlarının çoğunun yaşı 20 ile 30 değerleri arasındadır sonucuna varırız.

Bazen verilerin değerleri birbirine çok yakın olduğu ve frekansların bazı dal değerlerinde yoğunlaştığı durumlarda, dağılımı daha detaylı inceleyebilmek için, dal ve yaprak düzenimizde yer alan dal değerlerini ikiye böleriz. 0'dan 5'e kadar olan yaprak değerlerini bir dal değerine, 5'ten 10'a kadar olan dal yaprak değerlerini de diğer dal değerine yazarak dağılımı daha detaylı inceleyebiliriz. Biraz önce ele aldığımız örneğe bu uygulamayı gerçekleştirecek olursak:

Leaf Unit =1.0 (Yaprak Birimi)

0	1	
3	1	7 8 9
5	2	1 2
8	2	8 9 9
(2)	3	2 3
9	3	8 9
7	4	0 1
5	4	
5	5	
5	5	6 7
3	6	
3	6	
3	7	4
2	7	5 6

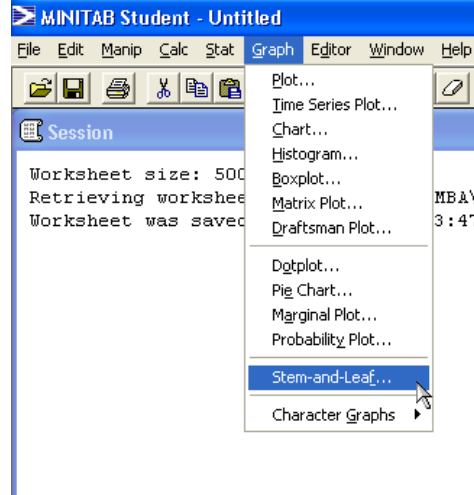
17'den daha küçük bir değer olmadığından, birinci 1 dal değerine denk gelen 0 ile 5 arasında kalan bir yaprak değeri bulunmamaktadır. Ama ikinci 1 dal değerine denk gelen ve 5 ile 10 arasında kalan üçtane yaprak değeri olduğundan 7,8,9 değerleri bu satıra girilmiştir. Tabiki dal ve yaprak düzeninin dezavantajları bulunmaktadır. Bu dezavantajların başında, bu uygulamanın büyük verisetlerinde kullanışsız olması gelir. Diğer bir dezavantajı da, dal değerlerinin sabit basamak değerlerinden oluştuğundan grup aralıklarının çok kısıtlı olmasıdır.

Dal-ve-Yaprak Düzeni

Tablo 3.3'te yer alan verileri yeniden incelediğimizde ve yaprak birimini (leaf unit) 0.1 olarak aldığımızda, dal ve yaprak düzenimiz; Tablo 3.3'te yer alan birinci gözleme ait olan dal değeri 20, yaprak değeri ise 4 olur. Benzer şekilde, ikinci gözlemimize ait olan dal değeri 23 iken yaprak değeri de 8 olur.

Minitab ortamında dal ve yaprak düzeni oluşturmak için öncelikle daha önceden kaydedilmiş olan **“getiri2.mtw”** dosyasını açınız.

Şekil 3.17’de görüldüğü gibi **“Graph”** (-grafik) menüsüne girip **“Stem-and-Leaf”** (Dal-ve-Yaprak) seçeneğine tıklayınız:



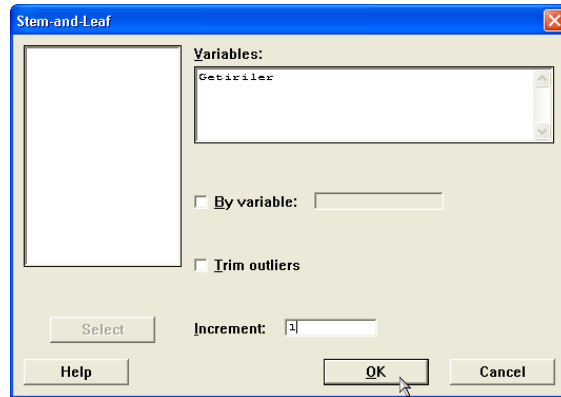
Şekil 3.17

Karşınıza çıkacak olan diyalog ekranında, **“Variables”** (değişkenler) yazan kısma tıklayıp soldaki pencereden **“C2 Getiriler”**’i seçiniz. **“Increment”** kısmına da, dallar birer birer artacağı için (28, 29, 30 gibi) 1 değerini giriniz.

DAL VE YAPRAK DÜZENİ

N = 59 gözlem sayısı
Yaprak Birimi = 0,10

```
1 20 4
1 21
1 22
2 23 8
2 24
3 25 6
4 26 2
6 27 67
11 28 36899
15 29 3359
16 30 1
21 31 56689
28 32 1334899
(5) 33 03478
26 34 003778
20 35 0
19 36
19 37
19 38 2
18 39 04
16 40 7
15 41 1
14 42 89
12 43 34567
7 44 67
5 45 47
3 46 6
2 47
2 48 06
```



Şekil 3.18

“OK” tuşuna tıkladığınızda, yanda yer alan dal ve yaprak düzenini elde edebilirsiniz.

En üstte 20 sayısı dal değerini, 4 sayısı da yaprak değerini temsil eder. Bu bize 20.4 getiriye sahip bir fonun olduğunu gösterir. Solunda yazan 1 rakamı ise o dala kadar kaç tane yaprak olduğunu ifade eder.

Dalın karşısında yaprak yoksa, (örneğin 21) bu %21 getiriye sahip hiçbir fon olmadığını belirtir.

29'a ait dalda; 3, 3, 5, ve 9 rakamları, 29.3 getiriye sahip 2 fon; 29,5 ve 29,9 getiriye sahip birer fon olduğunu gösteriyor. Soldaki 15 ise 29,9 a kadar toplam 15 getirinin sıralandığını belirtiyor.

33'e ait dalın solundaki (5) rakamı yalnızca bu daldaki yaprak sayısını gösteriyor. Bunun nedeni, verilerin yarısının sıralanmış olmasıdır.

Bundan sonraki dalların solunda yazan rakamlar sonraki dallarda bulunan toplam yaprak sayısını ifade eder, öncekileri ifade etmez.

Dal-ve-Yaprak Düzeninden ne tür bilgiler elde edebiliriz? Genellikle, dal-ve-yaprak büyük miktarda veriyi bir anda düzenlemesi ve bu verilerin nasıl dağıldığını, nerelerde yoğunlaştığını göstermesi bakımından oldukça kullanışlıdır.

Getiri dosyasındaki örnek verilerden hareket ederek şu sonuçlara varabiliriz:

1. En düşük bir yıllık getiri %20.4.
2. En yüksek bir yıllık getiri %48.6.
3. Bazı fonların getirileri aynıdır. %28.9, %31.6, %32.3, %32.9, %34 ve %34.7 getiri sağlayan birden fazla fon vardır.
4. 59 kamu fonunun dağılımı, %28 ile %34 arasında yoğunlaşmıştır.
5. %40'ın üzerinde getiri sağlayan fonların sayısı %25'in altındakilerden fazladır.

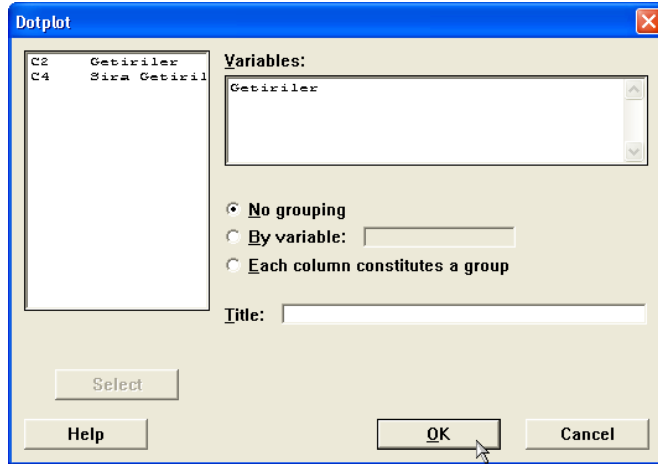
Burada yer alan yorumların ilk ikisi sıralı diziden yola çıkılarak da yapılabilir. 3. sonuç da biraz çaba ile oluşturulabilir ancak 4. ve 5. sonuçlar için Dal-ve-Yaprak düzeni kaçınılmazdır.

Nokta Diyagramı

Minitab Uygulamaları

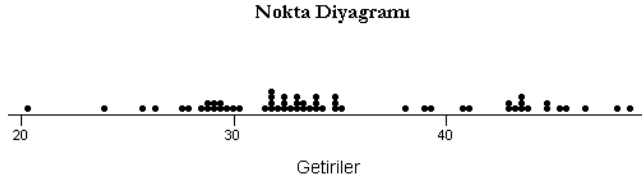
Nokta diagramı, dal yaprak diagramının grafiksel bir ifadesidir. Yatay ekseninde dallar yerilirken, dikey ekseninde ise yapraklara ait noktalar yerilmektedir. Her nokta, bir dala tekabül eden yaprakları ifade eder.

Nokta diagramı çizilirken, grup frekansları dikey ekseninde (Y ekseninde) ölçeklenmektedir. Grup limitleri yada grup orta noktaları yatay ekseninde (X ekseninde) yerilmektedir. Bütün bu işlemlerin hepsi otomatik olarak Minitab programı yardımı ile "graph" menüsünden "dot plot" (nokta diyagramı) seçeneği tıklanarak da gerçekleştirilebilirsiniz. Bu seçimler yapıldığında Şekil 3.19'daki diyalog ekranı ile karşılaşacaksınız.



Şekil 3.19

“No grouping” seçeneğini işaretledikten sonra OK tuşuna bastığınızda nokta diyagramını elde edebilirsiniz.

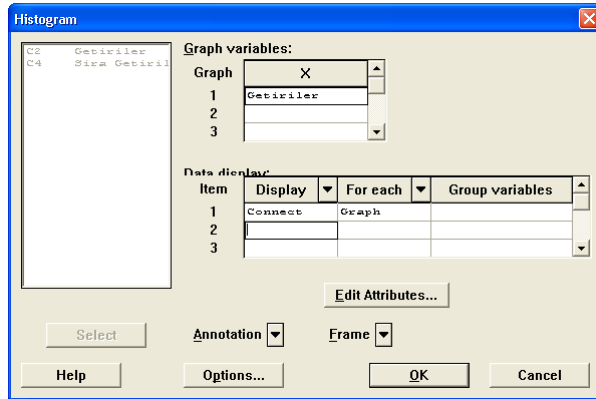


Şekil 3.20

Frekans Poligonu Minitab Uygulamaları

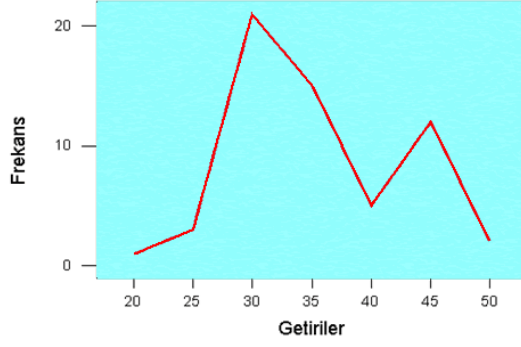
Frekans Poligonu da histogram gibi gruplara ayrılmış verilerin frekans dağılımını göstermektedir. Ancak bunu yaparken sütunlar yerine frekansların işaretlendiği bir çizgi kullanmaktadır. Aslında histogramda yer alan çubukların üst orta noktalarının birleştirilmesiyle elde edilen bir grafikten başka birşey değildir.

MINITAB programında frekans poligonu çizmek için, histogram çiziminde kullandığımız yolların aynısını kullanırız. Tek fark olarak, Şekil 3.21’de de görüldüğü üzere, “*Display*” (gösterge) kısmındaki “*Bar*” (sütun) yerine, sağ üstteki kutucuğa tıklayarak açılan menüden “*Connect*” (birleştir)’i seçmeniz gerekmektedir.



Şekil 3.21

Bu işlemi tamamladıktan sonra **“OK”** tuşuna tıkladığınızda, frekans poligonunu elde edebilirsiniz:



Şekil 3.22

Frekans Poligonu’nu da *“Options”* penceresinde görebileceğiniz değişik tiplerde elde etmek mümkündür. Alıştırma olarak Yüzdelik ve Kümülatif Frekanslardan oluşan poligonları çizmenizi tavsiye ederiz.

Histogram ve Frekans Poligonu, elimizdeki verilerin temel özelliklerini kolayca farketmemizi sağlar. Verilerin nerelerde yoğunlaştığı, nasıl dağıldığı, yükselme, alçalma noktalarını görme imkanı verir.

Bölüm 4: Kategorik Verilerin Düzenlenmesi ve Tablolaştırılması

Bir önceki bölümde, sayısal verilerin nasıl düzenlendiğini, tablo ve grafiklerle nasıl sunulduğunu incelemiştik. Bu bölümde ise sayısal olmayan kategorik verilerin nasıl düzenlendiği ve tablolarla nasıl özetlendiği üzerinde duracağız.

Kategorik veri nedir?

Kategorik veri, nesneleri veya kişileri kesin bir şekilde kategorilere indirgeyebildiğimiz bir veri türüdür. Örneğin, kategori olarak cinsiyeti (erkek ve kız olmak üzere) ele alalım, ve verimizin bir sınıfta yer alan öğrencilerden oluştuğunu düşünelim. Bu sınıfta, beş kişi bulunurken, bu beş kişiden üç tanesi erkek geri kalanı ise kızdır.

	Kategori
1. Öğrenci	Kız
2. Öğrenci	Erkek
3. Öğrenci	Kız
4. Öğrenci	Erkek
5. Öğrenci	Erkek

Tabloda da gördüğünüz gibi verilerin hiçbiri sayısal bir veri değildir, bir başka ifade ile bir önceki bölümde incelediğimiz sayısal yapının tam tersidir. Başka bir örnek ele alalım, Marmara bölgesinden rassal bir şekilde seçilmiş olan altı üniversitenin kurumsal yapısını inceleyelim.

	Kategori
1. Üniversite	Özel
2. Üniversite	Özel
3. Üniversite	Kamu
4. Üniversite	Kamu
5. Üniversite	Kamu
6. Üniversite	Yarı Özel (özel fakat devlet tarafından sübvansedeiliyor)

Şimdi ise verimiz, Özel, Kamu ve Yarı Özel olmak üzere üç kategoriden oluşmaktadır.

Aynı veri değerine ait birden fazla kategoride gözlemek mümkündür. Örneğin, sınıf örneğimizde öğrencilerinin medeni hallerini de merak ettiğimizi düşünelim.

	Kategori 1	Kategori 2
1. Öğrenci	Kız	Bekar
2. Öğrenci	Erkek	Evli
3. Öğrenci	Kız	Bekar
4. Öğrenci	Erkek	Bekar
5. Öğrenci	Erkek	Bekar

Üniversite örneğimizde, üniversitelerin kampüsünün bulunup bulunmadığını da inceleyelim. Dolayısıyla:

	Kategori 1	Kategori 2
1. Üniversite	Özel	Kampüsü var
2. Üniversite	Özel	Kampüsü yok
3. Üniversite	Kamu	Kampüsü var
4. Üniversite	Kamu	Kampüsü yok
5. Üniversite	Kamu	Kampüsü var
6. Üniversite	Yarı Özel (özel fakat devlet tarafından sübvansede ediliyor)	Şehir içinde

Özet Tablo

Kategorik veriler için özet tablo, sayısal veriler için hazırlanan frekans tablosu ile aynı formattadır. Verinin tamamını özet bir şekilde görebilmemiz imkanını verir. (1000 tane üniversite veya öğrencimiz olduğunu düşünelim) Bu tabloyu oluştururken, beş öğrencinin cinsiyetini ve medeni hallerinin dağılımını ele alalım ve bu verileri tablo 4.1 ve tablo 4.2’de olduğu gibi özetleyelim.

Tablo 4.1 Beş öğrencinin cinsiyet dağılımına ilişkin frekans ve yüzdelik özet tablosu

Cinsiyet	Öğrenci Sayısı	Öğrencilerin %’si	Oran
Kız	2	40.0	2/5
Erkek	3	60.0	3/5
Toplam	5	100.0	5/5

Tablo 4.2 Beş öğrencinin medeni halinin dağılımına ilişkin frekans ve yüzdelik özet tablosu

Medeni Hali	Öğrenci Sayısı	Öğrencilerin %'si	Oran
Evli	1	20.0	1/5
Bekar	4	80.0	4/5
Toplam	5	100.0	5/5

Bu tablolar ile frekans ve göreceli frekans dağılımı tabloları arasındaki fark dikkatinizden kaçmamıştır. Tek fark, frekans tablolarının ilk sütunda aralılar veya rakamlar yer alırken, burada kategorilerin yer almasıdır. Şimdi ikinci örneğimize ait olan özet tabloları oluşturalım.

Tablo 4.3 Altı üniversitenin kurumsal yapısının dağılımına ilişkin frekans ve yüzdelik özet tablosu

Kurumsal Yapı	Üniversite Sayısı	Üniversitelerin %'si	Oran
Özel	2	33.0	2/6
Kamu	3	50.0	3/6
Yarı Özel	1	17.0	1/6
Toplam	6	100.0	6/6

Tablo 4.4 Altı üniversitenin kampüs durumunun dağılımına ilişkin frekans ve yüzdelik özet tablosu

Kampüs Durumu	Üniversite Sayısı	Üniversitelerin %'si	Oran
Kampüsü Yok	2	33.0	2/6
Kampüsü Var	3	50.0	3/6
Şehir İçinde	1	17.0	1/6
Toplam	6	100.0	6/6

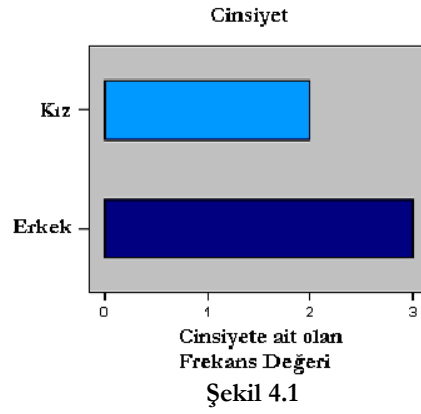
Kategorik Verilerin Grafikleri

Tablo 4.1'den Tablo 4.4'e kadar elde ettiğimiz özet verilere ait olan grafikleri oluşturarak kategorik verilerin grafiklerinin nasıl olduğunu inceleyelim.

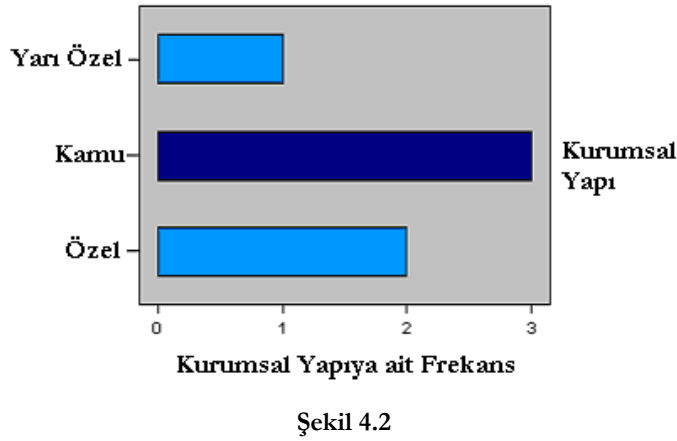
Bar Grafik

The Bar Chart

Bar grafiklerinde, her kategori bir sütundan oluşmaktadır. Bu sütunların uzunluğu, o kategoriye ait olan yüzde veya frekans miktarını temsil etmektedir. Bar grafiklerine, kategorik veriler için hazırlanmış histogramlardır da diyebiliriz. Şekil 4.1, Tablo 4.1' e ait olan bar grafiği göstermektedir.



Benzer bir şekilde, Şekil 4.2'de, Tablo 4.3'te yer alan veriye ait olan grafiği sergilemektedir.



Bar grafiğini oluştururken aşağıdaki hususlara dikkat etmemiz gerekmektedir:

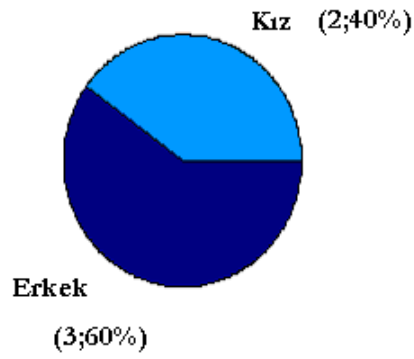
- Eğer kategorize edilmiş gözlemler, kategorik değişkenlerin birer çıktısı ise; bar grafikte yer alan sütunlar yatay olarak konulandırılır. (Şekil 4.1'de olduğu gibi) Kategorize edilmiş gözlemler, sayısal değişkenlerin birer çıktısı ise bar grafikte yer alan sütunlar dikey olarak konulandırılır. (Histogram gibi).
- Bütün sütunların genişliği birbirine eşittir. (Şekil 4.1 ve 4.2'de olduğu gibi) Sadece, uzunlukları birbirinden farklıdır.

The Pie Chart

Pasta Grafik

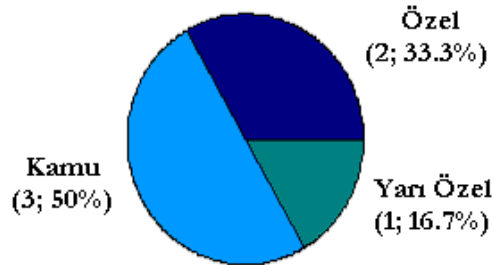
Aynı veriseti, pasta grafik kullanılarak da özetlenebilir. Şekil 4.3 ve 4.4, Tablo 4.1 ve 4.3'te yer alan verilere ait (veya Şekil4.1 ve 4.2'deki bar grafiklerindeki) olan pasta grafiklerini temsil etmektedirler.

Cinsiyetlerin Pasta Grafiği



Şekil 4.3

Kurumsal Yapının Pasta Grafiği



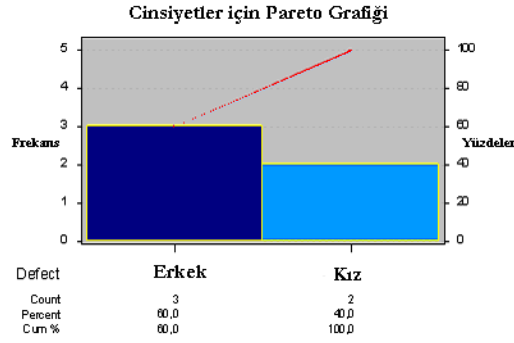
Şekil 4.4

Pasta grafik, bir daire 360^0 olduğu bilgisi temel alınarak çizilir. Dolayısıyla daire, yüzdelere göre kategorilere bölünmektedir.

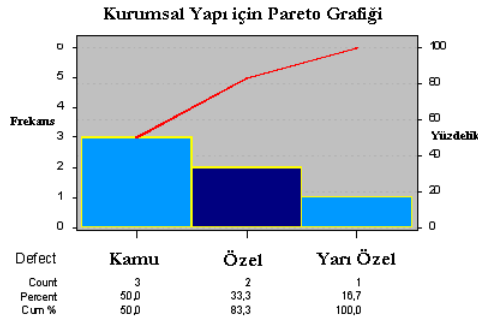
Pareto Grafiği

Kategorik verileri, yorumlama ve özetleme için kullanılan bir diğer grafik türü de Pareto grafiğidir. Pareto grafiği aslında, dikey özel bir bar grafiğidir. Frekansların büyükten küçüğe doğru dizildiği ve aynı ölçekte kümülatif poligonu da bünyesinde bulunduran bir grafik türüdür. Pareto grafiği araştırmacıya, hem bar grafiğinden hem de pasta grafiğinden daha fazla görsel bilgi iletir.

Şekil 4.5 ve 4.6, Tablo 4.1 ve 4.3'teki verilere ait olan pareto grafiğini göstermektedir.



Şekil 4.5



Şekil 4.6

Şekil 4.6'da da görüldüğü gibi Pareto grafiği, hem bar grafiğinde hem de pasta grafiğinde yer alan bilgileri araştırmacılara sunmaktadır. Ayrıca Pareto grafiği, kümülatif frekans poligonunu da göstermektedir. (sütunların üzerinde yer alan kırmızı doğru).

Kategorik Verinin Düzenlenmesi:

Kontenjan Tablosu

Genel olarak araştırmacılar, iki farklı kategorinin birbirleriyle olan etkileşimi birbirlerine göre oranlarını incelemek isterler. Örneğin, öğrencilerin cinsiyetleri ile medeni halleri arasında bir etkileşim olup olmadığını araştıralım. İşte bu tarz iki kategorinin etkileşiminin incelendiği tablolara kontenjan tabloları denilmektedir. Kontenjan tablosunu, beş öğrenciye göre çizdiğimizizi düşünelim.

Tablo 4.5 Beş öğrencinin cinsiyetini ve medeni hallerini gösteren kontenjan tablosu

Cinsiyet	Medeni Durum		
	Bekâr	Evli	Toplam

Kız	2	0	2
Erkek	2	1	3
Toplam	4	1	5

Tablo 4.5'i şöyle değerlendirebiliriz: Birinci satır, birinci sütundan başlayarak birinci satır ikinci sütuna devam etmeliyiz. Kızların ikisinde bekârdır aralarında evli yoktur. Erkeklerin ise ikisi bekâr, biri evlidir.

Benzer şekilde altı üniversite için kontenjan tablosu oluştursak:

Tablo 4.6 Altı üniversitenin kurumsal yapısını ve kampüs durumunu gösteren kontenjan tablosu

Kurumsal Yapı	Kampüs Durumu			Toplam
	Var	Yok	Şehiriçi	
Özel	1	1	0	2
Kamu	2	1	0	3
Yarı Özel	0	0	1	1
Toplam	3	2	1	6

Bu kontenjan tablosu sadece frekans değerlerine dayanmaktadır. Kategoriler arasında olası bir ilişki ortaya çıkarmak için bu sonuçları yüzdelerle dönüştürmek gerekmektedir. Ele almamız gereken üç nokta vardır:

- Genel toplamı baz alan yüzdeler
- Satır toplamlarını baz alan yüzdeler
- Sütun toplamlarını baz alan yüzdeler

Şimdi bu noktaları ayrı ayrı üç tabloda, beş öğrenci örneği ile inceleyelim.

Tablo 4.7 Beş öğrencinin cinsiyetini ve medeni hallerini gösteren kontenjan tablosu (genel toplamı baz alan yüzdeler)

Cinsiyet	Medeni Hal			Medeni Hal		
	Bekâr	Evli	Toplam	Yüzde	Yüzde	Yüzde
Kız	40	0	40	$\frac{2}{5} \times 100$	$\frac{0}{5} \times 100$	$\frac{2}{5} \times 100$
Erkek	40	20	60	$\frac{2}{5} \times 100$	$\frac{1}{5} \times 100$	$\frac{3}{5} \times 100$
Toplam	80	20	100	$\frac{4}{5} \times 100$	$\frac{1}{5} \times 100$	$\frac{5}{5} \times 100$

Tablo 4.7'yi şöyle değerlendirebiliriz: Birinci satır, birinci sütundan başlayarak birinci satır ikinci sütuna devam etmeliyiz. Sınıftaki öğrencilerin, %40'ı bekâr ve bayandır, %0'ı evli ve bayandır, %40'ı bekâr ve erkektir, %20'si ise evli ve erkektir.

Şimdi aynı verinin, satır toplamlarına göre oluşturulmuş kontenjan tablosunu inceleyelim.

Tablo 4.8: Beş öğrencinin cinsiyetini ve medeni hallerini gösteren kontenjan tablosu (satır toplamlarını baz alan yüzdelere)

Cinsiyet	Medeni Hal			Medeni Hal		
	Bekâr	Evli	Toplam	Yüzde	Yüzde	Yüzde
Kız	100	0	100	$\frac{2}{2} \times 100$	$\frac{0}{2} \times 100$	$\frac{2}{2} \times 100$
Erkek	66	34	100	$\frac{2}{3} \times 100$	$\frac{1}{3} \times 100$	$\frac{3}{3} \times 100$
Toplam	80	20	100	$\frac{4}{5} \times 100$	$\frac{1}{5} \times 100$	$\frac{5}{5} \times 100$

Tablo 4.8'i şöyle değerlendirebiliriz: Birinci satır, birinci sütundan başlayarak birinci satır ikinci sütuna devam etmeliyiz. Sınıftaki bayan öğrencilerin, %100'ü bekârdır, %0'ı evlidir; sınıftaki erkek öğrencilerin %66'sı bekârdır, %34'ü ise evlidir.

Şimdi aynı verinin, sütun toplamlarına göre oluşturulmuş kontenjan tablosunu inceleyelim.

Tablo 4.9 Beş öğrencinin cinsiyetini ve medeni hallerini gösteren kontenjan tablosu (sütun toplamlarını baz alan yüzdelere)

Cinsiyet	Medeni Hal			Medeni Hal		
	Bekâr	Evli	Toplam	Yüzde	Yüzde	Yüzde
Kız	50	0	40	$\frac{2}{4} \times 100$	$\frac{0}{1} \times 100$	$\frac{2}{5} \times 100$
Erkek	50	100	60	$\frac{2}{4} \times 100$	$\frac{1}{1} \times 100$	$\frac{3}{5} \times 100$
Toplam	100	100	100	$\frac{4}{4} \times 100$	$\frac{1}{1} \times 100$	$\frac{5}{5} \times 100$

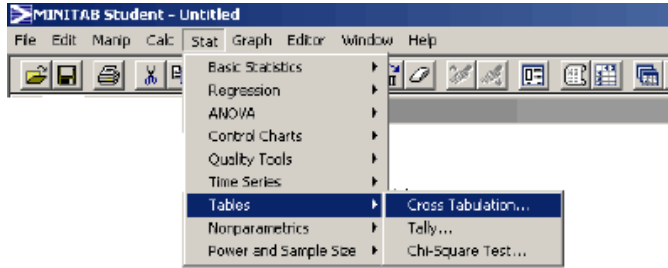
Tablo 4.9'u şöyle değerlendirebiliriz: Birinci satır, birinci sütundan başlayarak birinci satır ikinci sütuna devam etmeliyiz. Sınıftaki bekâr öğrencilerin, %50'si bayandır, sınıftaki evli öğrencilerin %0'ı bayan öğrencilerdir; bekâr öğrencilerin %50'si erkektir, evli öğrencilerin ise %100'ü erkek öğrencilerdir.

Minitab Uygulamaları

Şimdi, bu bölümde öğrendiğimiz istatistiksel araçların, Minitab yazılımı sayesinde nasıl kolaylıkla hesaplanabileceğini göreceğiz. Bu hesaplamaları gerçekleştirebilmek için öncelikle bir verisetimiz olması gerekmektedir. “*liseler.mtw*” adlı verisetimizde, Marmara Bölgesi'nden rassal olarak seçilmiş 45 tane liseye ait bilgiler yer almaktadır. Verisetinde, liseler çeşitli özelliklerine göre farklı şekilde gruplanmışlardır. Bu gruplar arasında, kurumsal yapısı (özel, kamu), lisenin yerleşkesi (köy lisesi, kasaba lisesi, şehir lisesi), akademik takvimi (dönemlik, yıllık), lisenin türü (DL: Düz Lise, AL: Anadolu Lisesi, ML: Meslek Lisesi, SL: Süper Lise, TCL: Ticaret Lisesi) yer almaktadır.

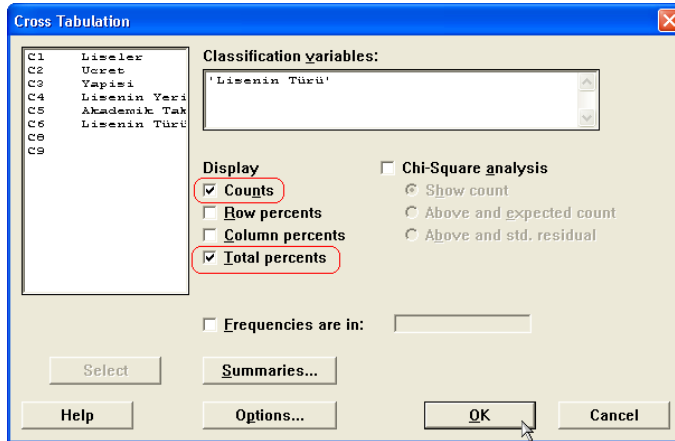
Bütün deęiřkenler verisetinde kategorik olarak yer almaktadır. Sayısal olarak yer alan tek deęiřken okul ücret veya baęıř miktarlarıdır ve dolar olarak ifade edilmiřtir.

Verisetimizdeki kategorilere ait olan özet tabloları oluřturalım. Öncelikle, “Stat” (istatistik) menüsünden “Tables”(tablolar)’ı seçiniz ve daha sonra “Cross Tabulation” (Çapraz Tablo)’a tıklayınız. (bkz Şekil 4.7)



Şekil 4.7

Karřınıza çıkacak olan diyalog ekranında, “Classification variables” (deęiřkenlerin sınıflandırılması) adlı alana tıkladıęınızda sol tarafta deęiřkenlerin çıktıđını göreceksiniz. Örneęin, lise türlerinin özet tablosunu oluřturmamak istersek. “C6 Lisenin Türü” deęiřkeninin üzerine iki kere tıklayınız veya “Select” komutu ile seçiniz. Ekranda yer alan seçeneklerden, “counts”(frekans) ve “total percents”(toplam yüzdeler) seçeneklerini iřaretleyiniz. (bkz Şekil 4.8)



Şekil 4.8

“OK” tuřuna bastıđınızda ařađıdaki çıktıyı elde edersiniz.

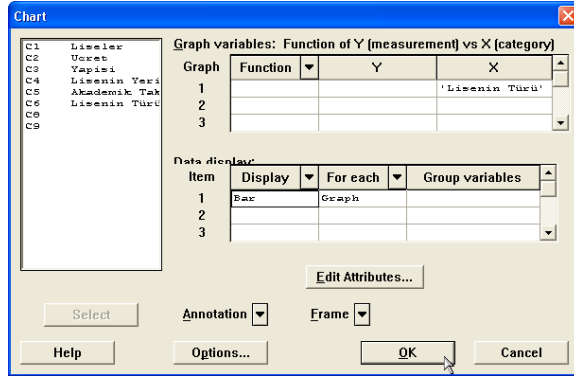
Tabulated Statistics		
Rows: Lisenin Türü		
	Frekans	%Değeri
AL	4	8,89
DL	22	48,89
ML	2	4,44
SL	16	35,56
TCL	1	2,22
Toplam	45	100,00

Şekil 4.9

Aynı yolu izleyerek diğer değişkenlerin de özet tablolarını oluşturabilirsiniz.

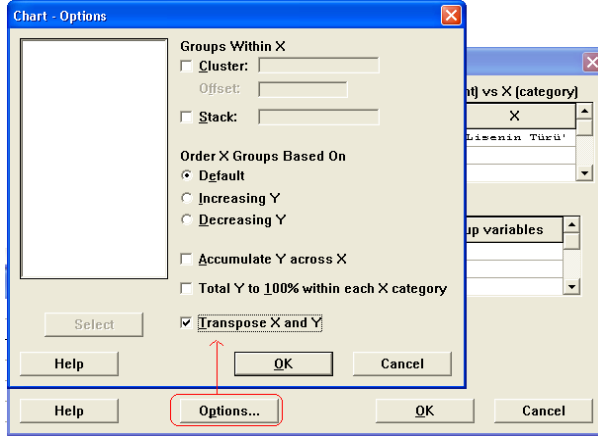
Bar Grafik

Lisenin türü değişkenine ait olan bar grafiğini çizmek için, öncelikle “Graph” menüsünden “Chart” (çizelge) seçeneğine tıklayınız. Karşınıza çıkan diyalog ekranında, “X” adlı başlığın altına tıklayınız ve sol tarafta yer alan değişkenler aktif hale gelince, “C6 Lisenin Türü” değişkenini seçiniz. (bkz Şekil 4.10).



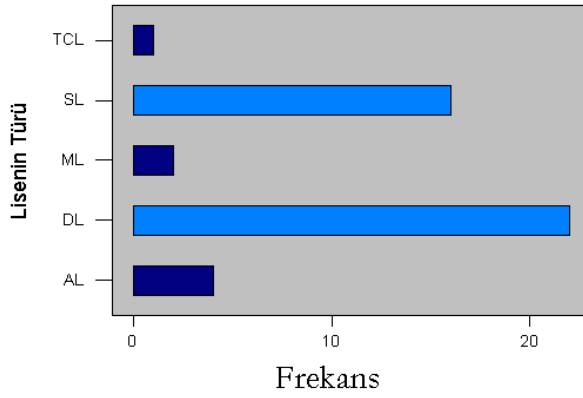
Şekil 4.10

Şimdi “Options” (seçenekler)’a tıklayınız ve “transpose X and Y” (X ve Y’nin devriği) seçeneğine tıklayınız. (bkz Şekil 4.11)



Şekil 4.11

“OK” tuşuna basıp bir önceki menüye döndükten sonra tekrar “OK” tuşuna bastığınızda Şekil 4.12’deki bar grafiği elde edebilirsiniz.

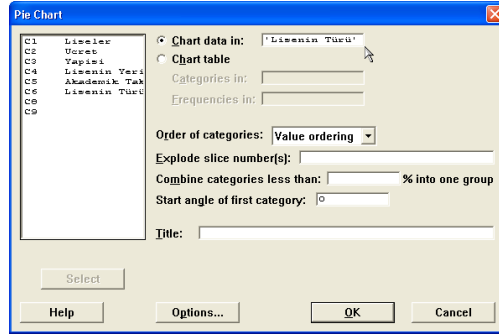


Şekil 4.12

Pasta Grafik

Lisenin türü değişkenine ait olan pasta grafiğini çizmek için, öncelikle “Graph” menüsünden “Pie Chart” (pasta grafik)

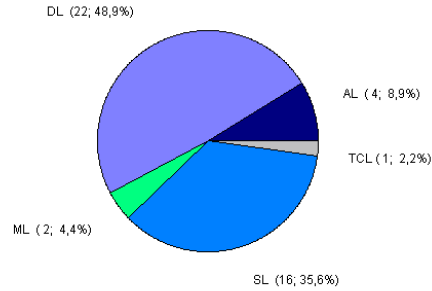
seçeneğine tıklayınız. Karşınıza çıkan diyalog ekranında, “*Chart Data in*” (çizelge verisinin alınacağı yer)’a tıklayınız ve sol tarafta yer alan değişkenler aktif hale gelince, “*C6 Lisenin Türü*” değişkenini seçiniz. (bkz Şekil 4.13).



Şekil 4.13

“OK” tuşuna bastığınızda şekil 4.14’te yer alan pasta grafiğini elde edersiniz.

Pie Chart of Lisenin Türü

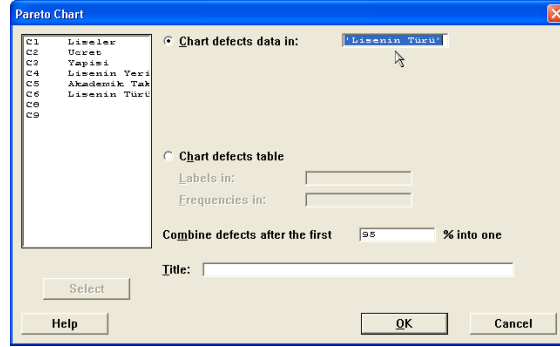


Şekil 4.14

Aynı yolu izleyerek diğer değişkenler için de pasta grafikleri oluşturabilirsiniz.

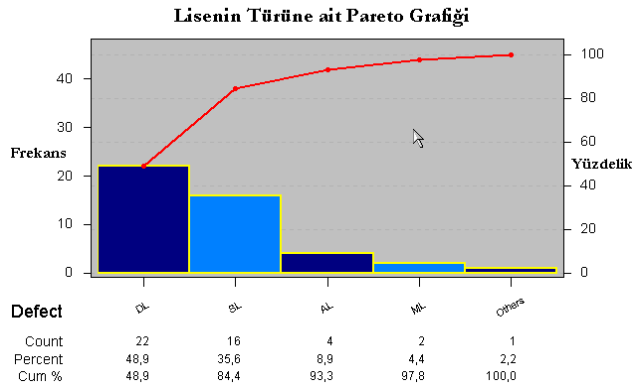
Pareto Grafiği

Lisenin türü değişkenine ait olan pareto grafiğini çizmek için, öncelikle “*Stat*”(istatistik) menüsünden “*Quality Tools*” (kalite araçları) seçeneğine ve oradan da “*Pareto Chart*” (Pareto Grafiği) seçeneğine tıklayınız. Karşınıza çıkan diyalog ekranında, “*Chart defects data in*”’e tıklayınız ve sol tarafta yer alan değişkenler aktif hale gelince, “*C6 Lisenin Türü*” değişkenini seçiniz. (bkz Şekil 4.15)



Şekil 4.15

“OK” tuşuna bastığınızda pareto grafiğini elde edebilirsiniz (bkz Şekil 4.16).

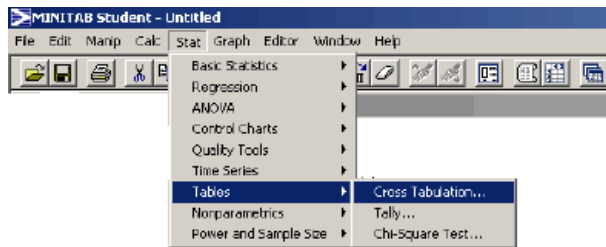


Şekil 4.16

Diğer kategoriler için de Pareto grafiğini oluşturabilirsiniz

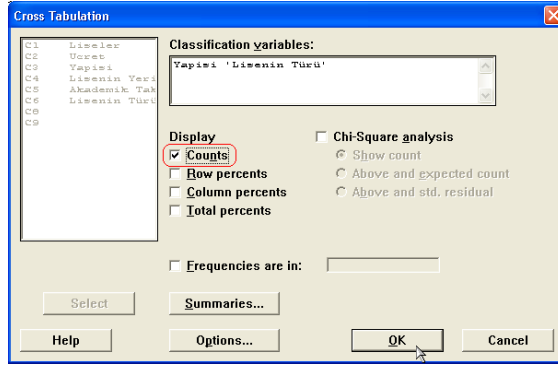
Kontenjan Tabloları

Daha önce nasıl oluşturulduğunu incelediğimiz kontenjan tablolarının Minitab ortamında nasıl oluşturulduğunu öğrenelim. Öncelikle, “Stat” (istatistik) menüsünden “Tables” (Tablolar) seçeneğini seçip daha sonra da “Cross Tabulation”(Çapraz Tablolama) seçeneğine tıklayınız. (bkz Şekil 4.17).



Şekil 4.17

Karşınıza çıkacak olan diyalog ekranında, “Classification variables”(değişkenlerin sınıflandırılması) adlı alana tıkladığınızda sol tarafta yine değişkenlerin çıktığını göreceksiniz. Şimdi ise aralarındaki ilişkiyi incelemek istediğimiz değişkenleri seçebiliriz. Örneğin, lise türleri ile lisenin yapısı arasındaki ilişkiyi incelemek için, “C3 Lisenin Yapısı” ve “C6 Lisenin Türü” değişkenlerinin üzerine iki kere tıklayınız veya “Select” komutu ile seçiniz. Ekranda yer alan seçeneklerden, “counts”(frekans) seçeneğini işaretleyiniz. (bkz Şekil 4.18).



Şekil 4.18

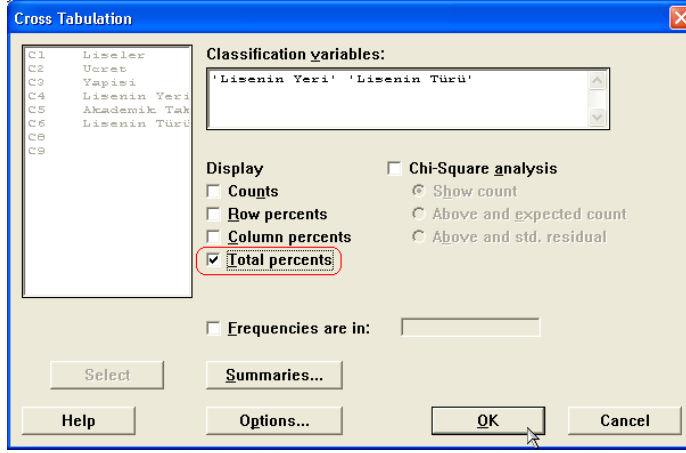
“OK” tuşuna bastığımızda çıktı aşağıdaki çıktıyı elde edebilirsiniz. (Şekil 4.19)

Tabulated Statistics						
Satırlar: Lisenin Yapısı			Sütunlar: Lisenin Türü			
	AL	DL	ML	SL	TCL	TOPLAM
Kamu	3	11	0	0	1	15
Özel	1	11	2	16	0	30
TOPLAM	4	22	2	16	1	45
Cell Contents --						
	Count					

Şekil 4.19

Elde ettiğimiz kontenjan tablosunu şöyle yorumlayabiliriz. 3’ü anadolu lisesi, 11’i düzlise, 1’i ticaret lisesi olmak üzere toplam 15 tane kamuya ait lise bulunmaktadır. 1’i anadolu lisesi, 11’i düz lise, 2’si meslek lisesi ve 16 tanesi super lise olmak üzere toplam 30 tane özel sektöre ait lise bulunmaktadır.

Şimdi diğer kontenjan tablolarını oluşturalım. (toplam, satır ve sütun yüzdeleri baz alan) Örneğin, genel toplam yüzdeleri göre bir kontenjan tablosu oluşturmak istersek, karşımıza çıkan diyalog ekranında “total percent” (toplam yüzde) kutucuğunu işaretleyiniz.



Şekil 4.20

“OK” tuşuna bastığınızda toplam yüzdelerin baz alındığı Minitab çıktısını elde edebilirsiniz.

Satırlar: Lisenin Yeri		Sütunlar: Lisenin Türü				
	AL	DL	ML	SL	TCL	TOPLAM
Kasaba Lisesi--	8,89	4,44	11,11	--	24,44	
Köy Lisesi --	13,33	--	8,89	--	22,22	
Şehir Lisesi8,89	26,67	--	15,56	2,22	53,33	
TOPLAM	8,89	48,89	4,44	35,56	2,22	100,00

Şekil 4.21

Şimdi tablomuzu yorumlayalım. Liselerin %24,44’ü kasabada bulunmaktadır. Kasabada bulunan liselerden hiçbirisi anadolu lisesi değildir. %8,89’u düz lisedir. %4,44’ü meslek lisesidir. Tablodaki verilerden hareketle yorumlar arttırılabilir. Satır yüzdelerinin ve sütun yüzdelerinin baz alındığı kontenjan tablolarını diyalog ekranında “row percent” (satır yüzdeleri) ve “column percent” (sütun yüzdeleri) kutucuklarına tıklayarak elde edebilirsiniz. Satır yüzdelerinin baz alındığı Lisenin yapısı ve türüne ait olan kontenjan tablosu aşağıdaki gibidir:

Rows:Lisenin Yapısı		Columns: Lisenin Türü				
	AL	DL	ML	SL	TCL	TOPLAM
Kamu	20,00	73,33	--	--	6,67	100,00
Özel	3,33	36,67	6,67	53,33	--	100,00
TOPLAM	8,89	48,89	4,44	35,56	2,22	100,00

Cell Contents --
% of Row

Şekil 4.22

Kamu okullarının %6,67’si ticaret lisesi iken, özel okulların ise %53,33’ü süper lisedir. Sütun yüzdelerinin baz alındığı Lisenin yapısı ve türüne ait olan kontenjan tablosu Şekil 4.23’teki gibidir:

Satırlar: Lisenin Yapısı			Sütunlar: Lisenin Türü			
	AL	DL	ML	SL	TCL	TOPLAM
Kamu	75,00	50,00	--	--	100,00	33,33
Özel	25,00	50,00	100,00	100,00	--	66,67
TOPLAM	100,00	100,00	100,00	100,00	100,00	100,00
Cell Contents --						
% of Col						

Şekil 4.23

Ticaret liselerinin %100'ü kamu okuludur. Kamuya ait bir süper lise ve meslek lisesi yoktur.

Verisetimizde yer alan diğer değişkenler için de kontenjan tablosu oluşturabilip yorumlayabilirsiniz.

Bölüm 5:

Betimsel İstatistikler

Giriş

Üçüncü ve dördüncü bölümlerde, ham verileri düzenleyip tablolarda özetleyerek, sonra da grafiklerle görsel olarak ifade ederek betimsel istatistiklere giriş yapmış olduk. Bütün bu yaptığımız işlemler, karmaşık ve kalabalık anlamsız veri topluluğundan anlamlı bilgiler çıkartabilmek amacına yönelikti.

Bu bölümde, verilerden anlamlı bilgiler çıkartabilmek için farklı yöntemler deneyeceğiz. Elimizdeki verilerin genel özelliklerini yansıtabilecek değerler bulmaya çalışacağız. Tablolarla, grafiklerle çalışmak yerine, temsili bir değer hesaplamak araştırmacıların işini daha kolaylaştırmaktadır. Bu değerler, en az tablo ve grafikler kadar, verisetinden bilgi edinmemize yardımcı olmaktadır.

Temsili değerler hesaplayarak, verilerin önemli özellikleri ile ilgili bilgilere ulaşabiliriz. Bu değerleri iki ana kategori altında inceleyeceğiz:

- Yerleşim Ölçüsü
- Dağılım Ölçüsü

Yerleşim Ölçüleri

Yerleşim ölçüsü, bir veri kümesindeki değerleri temsil edebilecek, verilerin nerede yoğunlaştığını gösterebilecek bir değerdir. Bu ölçünün, veri kümesinin özelliklerini “iyi” şekilde yansıtabilmesi için, verilerin merkezini temsil etmesi ve merkezde yeralması gerekmektedir.

Örneğin, genelde aritmetik ortalama hepimizin bildiği bir kavramdır. Aritmetik ortalama en çok bilinen yerleşim ölçülerinden biridir. “İyi” bir yerleşim ölçüsü olması verinin karakteristiği ile doğrudan ilgilidir. Bir sınıfta yer alan 5 kişinin yaşlarının 20 21 22 24 25 olduğunu düşünelim. Eğer sınıftaki öğrencilerin yaş profilini aritmetik ortalama ile temsil etmek istersek, sınıfın yaş ortalamasını 22 olarak buluruz. Bulduğumuz değer elimizdeki verilerin merkezinde (ortalarında) yeraldığından, yaş profilini gösteren iyi bir yerleşim ölçüsüdür. Verisetini 20 21 22 24 65 şeklinde değiştirirsek, aritmetik ortalama değerini 34 olarak hesaplarız. Bu değerde verisetinin merkezinden ziyade 24 ile 65 aralığında merkezin sağında kalmaktadır. Dolayısıyla, yaş profilini iyi bir şekilde temsil edemez. İstatistikçiler, değişik veri seti özelliklerine göre, bu verileri temsil edecek farklı yerleşim ölçüleri bulmuşlardır. Bu ölçüleri mümkün olduğunca basit ve kısa şekilde açıklamaya çalışacağız.

Aritmetik Ortalama

Aritmetik ortalama, merkezi eğilimi, ortalama değeri gösteren, en yaygın ve en önemli ölçüdür. Veriler örneklemden elde edilmişlerse; örneklem ortalaması, bu değerlerin toplanarak dizide yer alan eleman sayısına bölünerek elde edilebilir. Verisetinde yer alan i . değeri x_i olarak tanımlarsak, örneklem ortalamasının formülünü aşağıdaki gibi yazabiliriz:

$$\bar{x} = \frac{\sum_{i=1}^n x_i}{n}$$

Formülümüzde pay kısmını daha açık bir şekilde $\sum_{i=1}^n x_i = x_1 + x_2 + x_3 + \dots + x_n$ olarak ifade edebiliriz. Bu formülü kullanarak, bir ilkokulda yeralan sınıfların öğrenci mevcutlarının ortalamasını hesaplayalım. Sınıf mevcutları 46, 54, 42, 46, 32' ise ortalama öğrenci mevcudunu

$$\bar{x} = \frac{46 + 54 + 42 + 46 + 32}{5} = 44$$

formülünü kullanarak 44 olarak hesaplayabiliriz. Dolayısıyla bu ilkokulda bir sınıfta ortalama olarak 44 öğrenci bulunmaktadır sonucuna varırız.

Düzeltilmiş Ortalama

Verisetinde uçuk gözlem değerlerinin bulunduğu durumlarda oldukça kullanışlı olan bir ölçüdür. Çok büyük veya çok küçük değerler, aritmetik ortalamanın merkezden uzaklaşmasına, ortalamanın verileri iyi bir şekilde temsil edememesine sebebiyet verir. Bu tür durumlarda düzeltilmiş ortalama daha iyi bir merkezi eğilim ölçüsü verir. Düzeltilmiş ortalama, örneklem veri kümesindeki çok büyük ve çok küçük değerler dışarıda bırakılarak yeniden hesaplanan aritmetik ortalama değeridir.

Örneğin, %5'lik bir düzeltilmiş ortalama, verisetinde yer alan en küçük değerlerin bulunduğu %5'lik bir dilimi ve en büyük değerlerin bulunduğu %5'lik bir dilimi, verisetinden çıkararak, indirgenmiş verisetinde tekrar ortalama değerin hesaplanması ile bulunur.

Düzeltilme (indirgeme) işlemi verisetinin hem başında hem de sonunda olmak üzere simetrik olarak yapılmaktadır.

Örnek olarak, 12 yeni üniversite mezununun aylık maaşları aşağıdaki gibidir: (milyon TL)

2050, 2150, 2250, 2080, 1955, 1910, 2090, 2330, 2140, 2525, 2120, 2080

Bu verisetinin düzeltilmemiş ortalamasını hesaplayacak olursak:

$$\bar{x} = \frac{25680}{12} = 2140$$

Peki %5'lik bir düzeltilmiş ortalama değeri hesapladığımızda ne değişiklik olacak? %5'lik düzeltilmiş ortalama değeri hesaplamak için, verilerin $0.05 \times 12 = 0.6$ 'lık bir bölümü simetrik olarak çıkarılmalıdır. Bu işlemi gerçekleştirmeden önce, düzeltilmiş ortalama hesaplanırken, çıkarılacak veri miktarını (eğer tamsayı değilse) en yakın tamsayı değerine yuvarlarız. Örneğimizde de 0.6 değerini 1'e yuvarlayabiliriz.

En düşük (1910) ve en büyük (2525), birer değeri verisetimizden silersek, verisetimizde 1955, 2050, 2080, 2080, 2090, 2120, 2140, 2150, 2250, 2330 değerleri kalacaktır. Bu değerlerin ortalamasını hesaplar ise düzeltilmiş ortalama değerine aşağıdaki gibi ulaşmış oluruz.

$$\bar{x}_{düz} = \frac{21245}{10} = 2124.5$$

Ortanca Değer

Ortanca Değer, veriler büyüklük sırasına dizildikten sonra, tam ortada, yani merkezde kalan gözlemin değeridir. Eğer verisetinde bulunan eleman sayısı bir tek sayı ise, ortanca değer bu verisetinin küçükten büyüğe dizilmiş halinde en ortada yer alan terim iken, bu verisetinde bulunan eleman sayı bir çift sayı ise verisetinin küçükten büyüğe dizilmiş halinde ortada yer alan iki terimin aritmetik ortalamasıdır.

Örnek olarak yeni üniversite mezunlarının aldıkları maaşları küçükten büyüğe doğru sıralayacak olursak:

1910, 1955, 2050, 2080, 2080, 2090, 2120, 2140, 2150, 2250, 2330, 2525

Örneğimizde verisetinde yeralan veri sayısı çift bir sayı olduğundan dolayı, ortanca değeri bulabilmek için ortada yer alan iki değer olan 2090 ve 2120'nin ortalamasını hesaplamamız gerekmektedir. Bu iki değerın ortalamasını 2105 olarak hesaplayarak ortanca değeri de bulmuş oluruz.

Ortanca değer, aritmetik ortalama gibi aşırı uç değerlerden etkilenmemektedir. Bu özelliği, ortanca değerin en önemli avantajlarından. Gelir verilerinde genellikle ortanca değer kullanılmaktadır. Çünkü verisetinde bulunan yüksek bir gelir düzeyi, aritmetik ortalamayı yükseltip dağılımın merkezden ziyade sağa yatık olmasını sebebiyet verir. Beş kişiye ait olan yıllık gelir düzeylerini (milyon TL) 20000, 30000, 35000, 42000 ve 150000 olarak alırsak, ortalama gelir düzeyi 55400 olarak hesaplanır ve bu değer örneklemin ortalamasını etkin bir şekilde temsil etmemektedir. Aynı verisetinin ortanca değerini hesapladığımızda ise 35000 olarak buluruz. Açıkça görüldüğü gibi ortanca değer, beşinci kişinin yüksek gelir düzeyinden etkilenmemiştir.

Ortanca değeri bulabilmek için aşağıdaki formülden de faydalanabiliriz.

$$\text{Ortanca Değer} = \frac{(n+1)}{2} \cdot \text{gözlem değeri}$$

Formülümüzde n verisetinde yeralan gözlem sayısını ifade etmektedir. Örnek olarak, yeni üniversite mezunlarının maaşlarını gösteren verileri tekrar ele alalım. Bu örneğimizde gözlem sayımız 12 idi. Dolayısıyla ortanca değer:

$$\text{Ortanca Değer} = \frac{(12+1)}{2} = 6,5. \text{ gözlem}$$

6,5. gözlem değerini bulabilmek için, 6. gözlem ile 7. gözlemin ortalamasının alınması gerekmektedir. İşte bu yüzden burada ortanca değeri hesaplariken 2090 ile 2120'nin ortalamalarını aldık. İkinci örneğimizde n = 5 idi. Formülümüzü uyguladığımızda, ortanca değerin 3. gözleme eşit olduğunu kolaylıkla görebiliriz.

$$\text{Ortanca Değer} = \frac{(5+1)}{2} = 3. \text{ gözlem}$$

3. gözlem de verisetinde 46 değerine denk gelmektedir.

Mod

Mod, verisetinde en sık tekrar eden, frekansı en yüksek değerdir. Mod'un çoğunlukla kategorik verilerde kullanılması uygundur. Sayısal verilerde, en sık tekrar eden değer her zaman merkezi eğilimi göstermeyebilir. Örnek olarak bir spor salonuna gelen öğrencilerin yaşlarının 17, 18, 18, 15, 14 olduğunu düşünelim. Öğrencilerin yaşlarından oluşan bu verisetinin mod değeri, en

sık tekrarlanan deęer olan, 18 (iki kere tekrarlanmıřtır)'dir. Eęer verisetimizde en sık tekrarlanan iki tane gözlem var ise, bu durumda her iki gözlem de mod deęerini ifade etmektedir ve bu tarz mod deęerlerine bimod denilmektedir (17, 17, 18, 18, 15 gibi). Bir verisetinde, herhangi bir mod deęeri bulunmayadabilir (17, 12, 14, 15 gibi).

Orta Ranj

Orta Ranj, bir verisetindeki en büyük gözlem ile en küçük gözlem deęerinin ortalamasını temsil etmektedir.

Ařağıdaki gibi ifade edilir:

$$Orta\ Ranj = \frac{X_{en\ büyük\ deęer} + X_{en\ küçük\ deęer}}{2}$$

2050, 2150, 2250, 2080, 1955, 1910, 2090, 2330, 2140, 2525, 2120, 2080 verisetimizin orta ranjını hesaplayacak olursak:

$$Orta\ Ranj = \frac{1910 + 2525}{2} = 2217,5\ deęerini\ buluruz.$$

Orta Ranj, hem finansal analizcinin hemde hava durumu sunucusunun özet olarak kullandığı önemli bir ölçüm metodlarından biridir. Çünkü Orta Ranj veriseti hakkında hızlı ve yeterli bilgi vermektedir. Yine de çoęu uygulamada aritmetik ortalama gibi orta ranjı da kullanırken ihtiyatlı davranmak gerekmektedir. Sadece en küçük ve en yüksek gözlemi barındırdığından merkezi yönelim hakkında fikir sahibi olmamız zordur.

Yüzdelikler

Yüzdelik, verisetindeki bir deęerin merkezi eğilimde yeralmasına gerek kalmadan ölçülmesini sağlamaktadır. Yüzdelik, verinin en küçük deęer ile en büyük deęer arasında nasıl bir dağılım biçimi sergilediğı hakkında bize fikir verir. p.yüzdelik, veriyi yaklaşık olarak iki parçaya böler. Verinin yüzde p. kısmı, p. yüzdelikten küçükken; verinin yüzde (100 - p) kısmı, p.yüzdelikten büyüktür. Örnek olarak 50. yüzdeliğın ortanca (medyan) olduęu aşıkardır. p. yüzdeliğı hesaplamak için öncelikle veriyi küçükten büyüęe doğru sıralamamız gerekmektedir. Daha sonra i. indeks deęeri ařağıdaki gibi hesaplanabilir.

$$i = \frac{p}{100} n$$

Burada p yüzdelik iken, n ise gözlem sayısını temsil etmektedir. Eęer i deęeri bir tamsayı deęil ise bu deęeri bir sonraki tamsayı deęerine yuvarlamak gerekmektedir. Eęer i deęeri bir tamsayı ise, i'den büyük olan sonraki tamsayı deęeri ile i. gözlemin ortalaması alınır. Başka bir deyiřle i ile (i +1)'in ortalaması alınır.

Örneğın ,ařağıdaki 8 gözleme ait 30. yüzdeliğı hesaplamaya çalışalım:

5, 12, 13, 14, 17, 19, 23, 28

Ařağıdaki formülü kullanarak:

$$i = \frac{30}{100} 8 = 2.4$$

Bu yüzden 30. yüzdeler 2.4'ten bir sonraki tamsayı değeri olan 3'e eşittir. Dizimizdeki 3. gözlem 13'e tekabül etmektedir.

Daha önceden de hatırlanacağı üzere ortanca değer 14 ile 17 nin ortalaması olan 15.5'e eşittir. Bu bilginizi teyid edelim. Ortanca değer 50. yüzdeler olması gerektiğinden:

$$i = \frac{50}{100} 8 = 4$$

Bu rakam bir tamsayı olduğundan 50 . yüzdeler, 4. ile 5. gözlemlerin ortalamasına eşittir. Yani dizimizdeki 14 ile 17 nin ortalamasına eşittir.

Çeyrek Değerler

Çeyrek değerler, veri grubundaki gözlemler büyüklük sırasına dizildiğinde, bu gözlemleri 4 eşit gruba ayıran 3 değerden oluşmaktadır. Bunlara sırasıyla 1. çeyrek (25. yüzdeler), 2.çeyrek (50. yüzdeler) ve 3. çeyrek (75. yüzdeler) adı verilir.

Çeyrek değerleri hesaplamak için gerek olan formülleri bulmak için yüzdeler formüllerinden faydalanabiliriz. Çeyreklikler, genellikle Q_1, Q_2 ve Q_3 olarak ifade edilirler. Şimdi ise bu çeyreklik değerlerinin formüllerini inceleyelim. (Bu aşamada ikinci çeyreklik değerinin formülünü daha önceden bildiğimiz için incelemiyoruz. Unutulmaması gerekir ki ikinci çeyreklik zaten ortanca değere eşitti!)

$$Q_1 = \frac{25}{100}(n+1) = \frac{(n+1)}{4}$$

$$Q_3 = \frac{75}{100}(n+1) = \frac{3(n+1)}{4}$$

Yüzdeler formüllerinden yola çıkarak çeyreklik değerlerimize ait olan formülleri yukarıdaki gibi elde edebiliriz.

Çeyrek değerlerin hesaplanmasıyla, hem merkezi eğilime, hem de dağılıma ilişkin bilgi sahibi olabiliriz. Dağılım ölçüleri içerisinde çeyrekliklere tekrar değineceğiz.

Frekans Dağılım Tablolarından Yerleşim Ölçülerinin Hesaplanması:

Üçüncü bölümde, büyük verisetlerinde frekans dağılım tablolarından faydalanmanın en etkin yol olduğunu incelemiştik. Şimdi ise frekans dağılım tablolarından faydalanarak, yerleşim ölçülerini nasıl hesaplayacağımızı inceleyeceğiz. Büyük verisetlerinde verilerin gruplanıp gruplanmadığı önemlidir. Bazen gruplanmamış verilerde yerleşim ölçülerini hesaplamak daha etkin bir metod olarak karşımıza çıkabilir. Dolayısıyla, öncelikle gruplanmamış ham verilerde yerleşim ölçülerinin hesaplanmasını inceleyeceğiz. Daha sonra ise gruplanmış verilerin yerleşim ölçülerinin nasıl hesaplandığını inceleyeceğiz. Tablo 5.1'de yeralan veriler Taksim Meydanı'nda yapılan bir ankete katılan 36 kişinin yaşlarını temsil etmektedir.

Tablo 5.1

15
16, 16
17, 17, 17
18, 18, 18, 18
19, 19, 19, 19, 19
20, 20, 20, 20, 20, 20
21, 21, 21, 21, 21
22, 22, 22, 22
23, 23, 23
24, 24,
25
 $\sum_{i=1}^{36} x_i = 720$

Yapılan ankette 15 yaşında olan sadece bir kişi varken, 16 yaş grubunu temsil eden iki kişi yer almaktadır gibi tablo 5.1'i yorumlayabiliriz. Bir önceki bölümde öğrendiğimiz gibi bu verileri frekans tablosu halinde de ifade edebiliriz.

Tablo 5.2

Yaşlar (x_i)	Frekanslar (f_i)	$f_i x_i$	Kümülatif f_i
15	1	$15 \times 1 = 15$	1
16	2	$16 \times 2 = 32$	3
17	3	$17 \times 3 = 51$	6
18	4	$18 \times 4 = 72$	10
19	5	$19 \times 5 = 95$	15
20	6	$20 \times 6 = 120$	21
21	5	$21 \times 5 = 105$	26
22	4	$22 \times 4 = 88$	30
23	3	$23 \times 3 = 69$	33
24	2	$24 \times 2 = 48$	35
25	1	$25 \times 1 = 5$	36
	$\sum_{i=1}^{11} f_i = 36$	$\sum_{i=1}^{11} f_i x_i = 720$	

Tablo 5.1'i kullanarak aritmetik ortalamayı aşağıdaki gibi hesaplayabiliriz:

$$\bar{x} = \frac{\sum_{i=1}^n x_i}{n} = \frac{\sum (15 + 16 + 16 + \dots + 25)}{36} = \frac{720}{36} = 20$$

veya alternatif olarak Tablo 5.2'deki frekans değerlerini kullanarak yine aynı aritmetik ortalama değerini aşağıdaki gibi hesaplayabiliriz.

$$\bar{x} = \frac{\sum_{i=1}^{11} x_i f_i}{\sum_{i=1}^{11} f_i} = \frac{\sum (15+32+\dots+5)}{36} = \frac{720}{36} = 20$$

Burada bir önceki formülden farklı olarak toplama operatörünün sınırı 36 yerine 11 olarak alınmıştır. Bunun sebebi, toplam yapılırken artık x_i^{21} ’ların yerine $f_i x_i$ ’ların baz alınmasıdır.

Düzeltilmiş ortalama değerini hesaplamak için öncelikle, verinin ne kadarlık bir kısmını kesip çıkarmamız gerektiğine karar vermemiz gerekir. Bu kararı verebilmek için (%5’lik düzeltilmiş ortalama) $0.05 \times n$ eşitliğinden $0.05 \times 36 = 1.8$ değerini hesaplamamız gerekmektedir. Bulduğumuz 1.8 değerini en yakın tamsayıya yuvarlayacak olursak, verinin başından ve sonundan iki rakam çıkartmamız gerektiği sonucuna varırız. Verisetinin başından ve sonundan ikişer rakam çıkartığımızda örneklem boyutumuz $n_t = n - 4 = 36 - 4 = 32$ ’ye indirgenmiş olur. İndirgenmiş verisetimiz:

Tablo 5.3

16
17, 17, 17
18, 18, 18, 18
19, 19, 19, 19, 19
20, 20, 20, 20, 20, 20
21, 21, 21, 21, 21
22, 22, 22, 22
23, 23, 23
24
 $\sum_{i=1}^{n_t=32} = 640$

İndirgenmiş verisetimize ait olan frekans tablosunu oluşturacak olursak:

Tablo 5.4

Yaşlar (x_i)	Frekanslar (f_i)	$f_i x_i$
16	1	$16 \times 1 = 16$
17	3	$17 \times 3 = 51$
18	4	$18 \times 4 = 72$
19	5	$19 \times 5 = 95$
20	6	$20 \times 6 = 120$
21	5	$21 \times 5 = 105$
22	4	$22 \times 4 = 88$

² $\sum_{i=1}^{11} x_i f_i = x_1 f_1 + x_2 f_2 + \dots + x_{11} f_{11}$

23	3	$23 \times 3 = 69$
24	1	$24 \times 1 = 24$
	$\sum_{i=1}^9 f_i = 32$	$\sum_{i=1}^9 f_i x_i = 640$

Daha sonra ise düzeltilmiş ortalama değerini $\bar{x}_{düz} = 20$ ($= 640/32$) şeklinde hesaplayabiliriz.

Ortanca değeri hesaplayabilmemiz için verisetinin en ortasında yer alan gözlem değerini bulmamız gerekmektedir. Verisetimizde 36 rakam bulunduğundan, ortada yer alan 18. ve 19. gözlemlerin aritmetik ortalaması bize ortanca değeri $(20+20)/2 = 20$ olarak verir. İstedığımız gözlemlerin değerlerini bulmanın alternatif bir yolu da kümülatif frekans değerlerinden faydalanmaktır. Örneğin Tablo 5.2’de, 19 yaşına tekabül eden frekans değerinin 15 olması, 19 ve 19 yaşının altında yer alan 15 kişinin bulunduğunu gösterir. Aynı şekilde, 20 yaşına tekabül eden kümülatif frekans değerinin 21 olması, 20 ve 20 yaşının altında toplam 21 bulunduğunu gösterir. Dolayısıyla, ortanca değer de 18. ve 19. insanların yaşları 20 olduğundan, 20 değerine eşittir.

Mod değeride, en sık tekrarlanan değer olduğundan dolayı ve 20 sayısı verisetimizde 6 kere tekrarlanan bir sayı olduğundan dolayı, 20’ye eşittir. Dolayısıyla:

Ortalama	Düzeltilmiş Ortalama	Ortanca Değer	Mod
20	20	20	20

Hesaplamış olduğumuz yerleşim ölçülerinin neden birbirlerine eşit çıktıklarını merak edebilirsiniz? Bu sorunun cevabını, dağılım ölçülerini de işledikten sonra verilerin çarpıklığı hakkında yeterince fikrimiz olduktan sonra vereceğiz. Çeyreklikleri hesaplayacak olursak:

$$\frac{1}{4} (36 + 1) = 9.25 \text{ inci gözlem } Q_1 \text{’e eşittir.}$$

Bu gözlemi en yakın tam sayı değerine yuvarlarsak, 9. gözlemi 18 olarak elde ederiz. Dolayısıyla $Q_1 = 18$ olur. Benzer bir şekilde Q_3 değerini bulmak için:

$$\frac{3}{4} (36 + 1) = 27.75 \text{ inci gözlem } Q_3 \text{’e eşittir.}$$

Bu gözlemi en yakın tam sayı değerine yuvarlarsak, 28. gözlemi 22 olarak elde ederiz. Dolayısıyla $Q_3 = 22$ olur.

Frekans Dağılım Tablolarından Yerleşim Ölçülerinin Hesaplanması (Gruplanmış Veriler):

Şimdi ise aynı hesaplamaları gruplanmış veriler için yapacağız. Kilit noktalara rahatlıkla dikkat çekebilmek için, tekrar Tablo 5.1’de yer alan verisetini kullanacağız. Tablo 5.1’deki verisetini gruplamak için, öncelikle kaç adet grubumuzun olacağına karar vermeliyiz. Bu karar aşamasında da üçüncü bölümde öğrendiğimiz formülü kullanacağız.

$$\text{Grup Sayısı} = 1 + \log_2(n) = 1 + \log_2 36 = 1 + 5.16 = 6.16$$

Dolayısıyla grup sayımızı 6 olarak belirlemiş olduk. Şimdi ise grup aralıklarını belirlemeliyiz. Yine üçüncü bölümde öğrendiğimiz grup aralığı formülünü kullanarak:

$$\text{Grup Aralığı} = \frac{\text{En Yüksek Değer} - \text{En Düşük Değer}}{\text{Grup Sayısı}} = \frac{25-15}{6} = 1.66$$

Dolayısıyla bu formül yardımıyla da grup aralığını 2 olarak belirledik. Şimdi Tablo 5.2'deki verileri gruplayarak yeniden düzenleyelim.

Tablo 5.5

Yaşlar	Frekanslar(f_i)	Orta Noktalar(x_i)	$f_i x_i$	Kümülatif Frekans f_i
15-17(dahil değil)	3	16	$16 \times 3 = 48$	3
17-19(dahil değil)	7	18	$18 \times 7 = 126$	10
19-21(dahil değil)	11	20	$20 \times 11 = 220$	21
21-23(dahil değil)	9	22	$22 \times 9 = 198$	30
23-25(dahil değil)	5	24	$24 \times 5 = 120$	35
25-27(dahil değil)	1	26	$26 \times 1 = 26$	36
	$\sum_{i=1}^6 f_i = 36$		$\sum_{i=1}^6 f_i x_i = 738$	

Gruplanmış verilerde aritmetik ortalamayı nasıl hesaplayacağız? Yapacağımız tek şey orta noktaları hesaplamak ve orta noktaları, hesaplamalarda normal gözlem olarak ele almaktır. Orta noktalar, Tablo 5.5'te üçüncü sütunda hesaplanmıştır. Artık bu orta noktaları normal gözlem değerleri olarak ele alabiliriz.

$$\bar{x} = \frac{\sum_{i=1}^6 x_i f_i}{\sum_{i=1}^6 f_i} = \frac{\sum (48 + 126 + \dots + 26)}{36} = \frac{738}{36} = 20.5$$

Toplam operatörünün sınırları 11'in yerine artık 6 grup yeraldığından dolayı 6'dır. Aritmetik ortalamayı, gerçek ortalama değerinden (20) farklı bir değer 20.5 olarak hesaplarız. Bu farkın nedeni, verileri gruplarken bir miktar bilgi kaybetmemizden kaynaklanmaktadır.

Düzeltilmiş ortalamayı hesaplamak için yine öncelikle verisetinin ne kadarlık kısmını indirgememiz gerektiğini belirlememiz gerekmektedir. Bunu yapabilmek için öncelikle $0.05 \times 36 = 1.8$ değerini hesaplarız. Daha sonra verisetinin başından ve sonundan ikişer rakam çıkartığımızda örneklem boyutumuz $n_t = n - 4 = 36 - 4 = 32$ 'ye indirgenmiş olur.

Tablo 5.6

Yaşlar	Frekanslar(f_i)	Orta Noktalar(x_i)	$f_i x_i$	Kümülatif Frekans f_i
15-17(dahil değil)	1	16	$16 \times 1 = 16$	1
17-19(dahil değil)	7	18	$18 \times 7 = 126$	8
19-21(dahil değil)	11	20	$20 \times 11 = 220$	19
21-23(dahil değil)	9	22	$22 \times 9 = 198$	28
23-25(dahil değil)	4	24	$24 \times 4 = 96$	32
	$\sum_{i=1}^5 f_i = 32$		$\sum_{i=1}^5 f_i x_i = 656$	

Verisetinin başından iki gözlemin çıkarılması, birinci grubun frekans değerinin 2 birim azalmasına sebebiyet verirken, sonundan iki gözlemin çıkarılması ise beşinci grubun frekans sayısının bir birim azalmasına sebebiyet verir. Dolayısıyla düzeltilmiş ortalama değerini $\bar{x}_{düz} = 20.5 (= 656/32)$ olarak hesaplayabiliriz.

Ortanca değeri hesaplarken yine 36 gözlem bulunduğundan dolayı 18. ve 19. gözlemlerin aritmetik ortalamasını hesaplamamız gerekmektedir. Dolayısıyla, $(36+1)/2 = 18.5$ 'inci gözlem değerini hesaplamamız gerekmektedir. Bu gözlemleri yine Tablo 5.5'te bulunan kümülatif frekans değerlerinden faydalanarak bulabiliriz. Yaş grubu 17-19 iken kümülatif frekans değerinin 10 olması, örneklemde 19 yaşında veya daha küçük toplam 10 kişi vardır anlamına gelir. Bizim aradığımız gözlemler 18. ve 19. gözlemler olduğundan, kümülatif frekans değeri 21 olan 19-21 yaş grubuna bakmamız gerekir. Çünkü 17-19 yaş grubu sadece 10. gözleme kadar verisetini ele almıştır. 18. ve 19. gözlemleri içermez. Dolayısıyla rahatlıkla ortanca değerimizin 19-21 yaş grubu içerisinde yeraldığını belirtebiliriz. Ancak gruplanmış verilerde gözlemlerin kendisini göremediğimizden dolayı bu aşamada ortanca değerin tam olarak değerini bulamayız. Ancak istatistikçiler bu sorunu aşarak, gruplanmış verilerde ortanca değeri hesaplamak için orta noktadan da faydalanarak aşağıdaki formülü geliştirmişlerdir.

$$\text{Ortanca Değer} = L + \left(j - \frac{1}{2} \right) \left(\frac{U - L}{f} \right)$$

Bu formülde :

- L: ortanca değerin yeraldığı grubun alt limitini temsil etmektedir. Örneğin, örneğimizde ortanca değer 19-21 grubunda yer aldığından bu grubun alt limiti olan L değeri 19'a eşittir.
- j: ortanca değerin o grubun içinde kaçınıcı gözlem olduğunu göstermektedir. Örneğin, ortanca değerimiz bulunduğu grupta 8.5'inci gözlemdir. Normalde ortanca değerimiz 18.5'inci gözlem idi. Fakat, 18.5'ten 10 tanesi bir önceki grubun kümülatif frekansına gitti ve geriye bu grup için 8.5'inci terim kaldı. Dolayısıyla j değeri 8.5'e eşit olmaktadır.
- U: ortanca değerin yeraldığı grubun üst limitini temsil etmektedir. Örneğin, örneğimizde ortanca değer 19-21 grubunda yer aldığından bu grubun üst limiti olan U değeri 21'e eşittir.
- f: ortanca değerin yeraldığı gruba ait olan frekans değerini temsil etmektedir. Örneğimizde bu değer 11'e eşittir.

Formülümüzü kullanarak gruplanmış verilerde ortanca değeri hesaplayalım.

$$\text{Ortanca Değer} = 19 + \left(8.5 - \frac{1}{2}\right) \left(\frac{21-19}{11}\right) = 20.455$$

Görüldüğü gibi sadece orta noktalardan faydalananarak hesapladığımız ortanca değeri gayet gerçeğine yakın bir değerdir. Dolayısıyla bu da kullanmış olduğumuz formülün gücünü göstermektedir.

Aynı formülden faydalananarak gruplanmış verilerde çeyreklik değerleri de hesaplayabiliriz. Prosedürümüz aynıdır. Fakat öncelikle çeyreklik değerlerinin hangi gözlemlere tekabül ettiğinin bulunması gerekmektedir. Q_1 değeri, 9.25'inci gözleme $((36+1)/4 = 9.25)$ tekabül etmektedir. Birinci grubun (15-17) kümülatif frekans değeri 3 olduğundan birinci çeyrekliği içeremez. İkinci grubun ise (17-19) kümülatif frekans değeri 10 olduğundan birinci çeyrekliği içerir ve birinci çeyreklik bu grupta 6.25'inci gözlemdir. (çünkü ilk üç gözlem birinci gruba gitti) Kullanacağımız formülü yeniden ele alalım:

$$\text{Çeyreklik} = L + \left(j - \frac{1}{2}\right) \left(\frac{U - L}{f}\right)$$

Bu formüle :

- L: çeyrekliğin yer aldığı grubun alt limitini temsil etmektedir. Örneğin, örneğimizde çeyreklik 17-19 grubunda yer aldığından bu grubun alt limiti olan L değeri 17'ye eşittir.
- j: çeyrekliğin o grubun içinde kaçınıcı gözlem olduğunu göstermektedir. Örneğin, çeyrekliğimiz, bulunduğu grupta 6.25'inci gözlemdir. Normalde çeyrekliğimiz 9.25'inci gözlem idi. Fakat, 9.25'ten 3 tanesi bir önceki grubun kümülatif frekansına gitti ve geriye bu grup için 6.25'inci terimi kaldı. Dolayısıyla j değeri 6.25'e eşit olmaktadır.
- U: çeyrekliğin yer aldığı grubun üst limitini temsil etmektedir. Örneğimizde çeyreklik 17-19 grubunda yer aldığından bu grubun üst limiti olan U değeri 19'a eşittir.
- f: çeyrekliğin yer aldığı gruba ait olan frekans değerini temsil etmektedir. Örneğimizde bu değer 7'ye eşittir.

Dolayısıyla

$$Q_1 = 17 + \left(6.25 - \frac{1}{2}\right) \left(\frac{19-17}{7}\right) = 18.642$$

Benzer şekilde Q_3 'ü hesaplırsak:

$$Q_3 = 21 + \left(6.75 - \frac{1}{2}\right) \left(\frac{23-21}{9}\right) = 22.388$$

Mod değerini hesaplayabilmemiz için öncelikle modun yer aldığı grubu bulmamız gerekmektedir. Mod, frekans değeri en yüksek olan grupta yer almaktadır. Örneğimizde frekans sayısı 11 olan 19-21 yaş grubu modun yer aldığı yaş grubudur. Gruplanmış verilerde mod değerini hesaplamak için aşağıdaki formülü kullanarak yine orta noktalardan faydalanabiliriz.

$$\text{Mod} = L + \left(\frac{f_m - f_{m-1}}{2f_m - f_{m-1} - f_{m+1}}\right) x_i$$

- L: modun yeraldığı grubun alt limitini ifade eder. Örneğimizde, L 19-21 grubunun alt limiti olan 19'a eşittir.
- f_m : modun yeraldığı grubun frekans değeridir. Örneğimizde 11'e eşittir.
- f_{m-1} : modun yeraldığı gruptan bir önceki gruba ait olan frekans değeridir. Örneğimizde 7'ye eşittir.
- f_{m+1} : modun yeraldığı gruptan bir sonraki gruba ait olan frekans değeridir. Örneğimizde 9'a eşittir.
- x_i : modun yeraldığı grubun genişliğine eşittir. Örneğimizde 2'ye eşittir.

$$Mod = 19 + \left(\frac{11 - 7}{2(11) - 7 - 9} \right) 2 = 20.33$$

Dağılım Ölçüleri

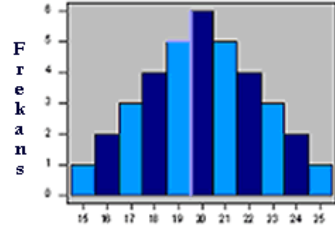
Üçüncü bölümde, ham verileri organize ederek frekans dağılım tabloları ile grafiklerle ifade ettik. Verilerin dağılımları, eğilimleri hakkında fikir sahibi olmaya çalıştık. Bu bölümde ise, merkezi eğilimlerin ölçümlerini ve dağılımlarını inceleyeceğiz verilerin değişkenliğinin hesaplanması üzerinde duracağız.

Dağılım neden önemlidir?

Aritmetik ortalama veya ortanca değer gibi değerler sadece verinin merkezinde yeralmaktadır. Ancak verinin ortasında yer alan değerler bize dağılım hakkında bir fikir vermezler. Örneğin bir doğa gezisine çıktınız, tur rehberiniz size karşınızda duran nehrin derinliğinin ortalama olarak 1.2 mt olduğunu söyledi. Nehir hakkında herhangi başka bir bilgi almadan nehrin içinden geçer misiniz? Muhtemelen geçmezsiniz. Yüzme bilmiyorsanız derinliğinin nasıl değiştiği hakkında fikir sahibi olmak istersiniz. Eğer derinlik maksimum 1.6 mt, minimum 0.4 mt olarak değişiyorsa, muhtemelen geçmeye karar verirsiniz. Fakat, eğer derinlik maksimum 3.5 mt ile minimum 0.8 mt arasında değişiyorsa, muhtemelen geçmezsiniz. Dolayısıyla bizim için verilerin ortalamasının yanı sıra nasıl farklılaştıkları, nasıl dağıldıkları da çok önemlidir.

Dağılım ölçülerinin küçük çıkması, verilerin birbirine yakın olarak, ortalamaya yakın değerler olarak dağıldığını ifade eder. Bu durumlarda ortalama değerimiz veriyi yansıtabilmesi açısından güvenilirdir. Fakat, dağılım ölçülerinin büyük çıkması durumunda veriler birbirinden uzaklaşarak dağılır ve ortalamanın etrafında toplanmazlar. Dolayısıyla bu tarz durumlarda ise ortalama veriyi temsil edebilecek kadar güvenilir değildir. Tablo 5.1'de yer alan gruplanmış yaşların ortalaması 20'dir ve bu yaşların dağılımının histogram ile gösterimi şekil 5.1'deki gibidir.³

³ Tablo 5.5'in gruplanmış verilerine ait olan histogram, şekil 5.1 deki gibi simetrik olmadığından aralarında aynı veriseti olmasına rağmen farklılık gözlenir..



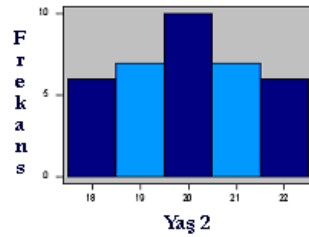
Yaş 1
Şekil 5.1

Şimdi Taksim Meydanı'nda yapılan ankete katılan 36 kişinin yaşlarının tablo 5.7'deki gibi dağıldığını düşünelim.

Tablo 5.7

18, 18, 18, 18, 18, 18
19, 19, 19, 19, 19, 19, 19
20, 20, 20, 20, 20, 20, 20, 20, 20, 20
21, 21, 21, 21, 21, 21, 21
22, 22, 22, 22, 22, 22

Yine bu verilere ait olan ortalama değeri 20'ye eşittir.



Şekil 5.2

Açıkça görüldüğü gibi, Tablo 5.7'deki verilerin dağılımı daha çok ortalama etrafında kümelenmişlerdir.

Şuana kadar, verilerin dağılımları hakkında fikir sahibi olabilmek için histogram gibi görsel araçlardan faydalandık. Fakat bu araçlar bazen bizi yanıltabilir ve dağılım hakkında kesin bir yargıya varmamızı engelleyebilir. Dağılım hakkında kesin ve güvenilir bilgiler elde edebilmek için dağılım ölçülerini incelememiz gerekmektedir.

Ranj

Ranj, verisetinde yer alan en küçük gözlem değeri ile en büyük gözlem değeri arasındaki farktır. Seride yer alan en büyük gözlem değerinden, en küçük gözlem değerinin çıkarılması ile elde edilir.

$$\text{Ranj} = X_{\text{en büyük gözlem}} - X_{\text{en küçük gözlem}}$$

Dolayısıyla hesaplanması çok kolay bir ölçüdür. Fakat, verisetinde yer alabilecek olan aşırı büyük veya aşırı küçük değerlere karşı çok hassastır, bizi yanıltabilir. Tablo 5.1’de yeralan verilerin ranjı kolaylıkla, $25-15 = 10$ olarak hesaplanabilir. Tablo 5.7’deki verilerin ranjı ise $22-18 = 4$ olarak hesaplanır. Ranj değerlerinden de tablo 5.7’deki verilerin daha az bir farklılaşma içinde olduğu verilerin birbirine daha çok yakınlaştığı gözlenmektedir.

Çeyrek Ranjı (IQR)

Çeyrek Ranjı, ranjdan farklı olarak aşırı gözlem değerlerine karşı duyarlıdır. 3. çeyrek ile 1. çeyrek arasındaki fark alınarak hesaplanır.

$$IQR = Q_3 - Q_1$$

Çeyrek ranjı veri kümesinde yer alabilecek uçdeğerlerden etkilenmedikleri için, çeyrek ranjı kararlı bir ölçüdür. Fakat, çeyrek ranjı verilerin %50’sini temsil etmektedir. Tablo 5.1’deki verilere ait olan çeyrek ranj değeri $22-18 = 4$ ’tür. Tablo 5.7’deki verilere ait olan çeyrek ranj değeri $21 - 19 = 2$ ’dir. Açıkça görüldüğü gibi ikinci tablodaki verilerin dağılımında daha az farklılaşma gözlenmektedir.

Varyans

Varyans, bütün veri değerlerini dağılım hesabına katması nedeniyle çeyrek ranjından farklı olarak en çok tercih edilen ölçüdür. Her gözlem değerinden ortalamanın farkı alınması prensibine dayanmaktadır. Bu farka sapma denilmektedir ve $x_i - \bar{x}$ olarak ifade edilir.

Varyans değeri hesaplanırken bütün sapmaların karesi alınır, sonra da bu kare sapmalar toplanır. Bir çok istatistiki uygulamada elimizde anakütle değeri değil örneklem değeri bulunmaktadır. Bu örneklem anakütle varyansının (genellikle σ^2 sembolü ile gösterilir) tahmin etmekte kullanılır. Anakütle varyansının sapmasız bir tahmin edicisi (s^2) sapmaların karesi n yerine, n-1 değerine bölünerek bulunur.

$$s^2 = \frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n-1} \quad (5.1)$$

İlkokul öğrencilerinin sınıf mevcudlarını tekrar ele alalım: 46, 54, 42, 46, 32. Örneklem ortalamasını daha önce hesaplamıştık.

$$\bar{x} = \frac{46+54+42+46+32}{5} = 44$$

Daha sonra varyans değerini aşağıdaki gibi hesaplayabiliriz:

$$s^2 = \frac{(46-44)^2 + (54-44)^2 + (42-44)^2 + (46-44)^2 + (32-44)^2}{5-1} = 64$$

Standart Sapma

Standart sapma değeri, varyansın karekökü alınarak hesaplanır. Böylece, (+) sapmalarla (-) sapmaların birbirlerini götürmelerini engelledikten sonra, sapmaların karekökünü alarak varyans değerini standardize ederiz. Dolayısıyla, varyansın ölçüm birimi dezavantajını ortadan kaldırmış oluruz.

$$s = \sqrt{s^2} \quad (5.2)$$

Standart sapma, veri kümesindeki bir değerde görülebilecek ortalama sapma miktarını belirtir. Buradan hareketle, örnek veri kümemizdeki değerlerin ortalama sapmasının $\sqrt{64} = 8$ olduğunu söyleyebiliriz.

Farklılaşma Katsayısı

Varyans ve standart sapma ölçüleri, verilerin özelliğine ve ölçüm birimlerine göre farklılık gösterebilir (insanların ağırlıklarının ortalaması ile standart sapmasının 75kg'a 15kg, ineklerin ise 240kg'a 30kg olması gibi). Aynı tiplerdeki verilerin değişkenliklerini kıyaslamak için farklılaşma katsayısı hesaplanır.

$$Farklılaşma\ Katsayısı = \frac{\text{Standart Sapma}}{\text{Ortalama}} \quad (5.3)$$

Bu katsayı, standart sapmanın ortalamaya bölümüyle elde edilir. Bu sayede, değişik ortalama ve standart sapma ölçülerine sahip veri kümelerindeki dağılım miktarları karşılaştırılabilir.

Yüzdelik Oranları

Yüzdelik oranı aşağıdaki gibi ifade edilir.

$$\text{Yüzdelik Oranı} = \frac{90. \text{ yüzdelik}}{10. \text{ yüzdelik}} \quad (5.4)$$

Bu ölçü aşırı değerleri gözardı ettiğinden dolayı kullanışlıdır.

Frekans Dağılım Tablolarından Dağılım Ölçülerinin Hesaplanması:

Frekans dağılım tablolarından faydalanarak, yerleşim ölçülerini nasıl hesaplayacağımızı inceleyeceğiz. Büyük verisetlerinde verilerin gruplanıp gruplanmadığı önemlidir. Bazen gruplanmamış verilerde yerleşim ölçülerini hesaplamak daha etkin bir metod olarak karşımıza çıkabilir. Dolayısıyla, öncelikle gruplanmamış ham verilerde yerleşim ölçülerinin hesaplanmasını inceleyeceğiz. Daha sonra ise gruplanmış verilerin yerleşim ölçülerinin nasıl hesaplandığını inceleyeceğiz. Tablo 5.1’de yeralan veriler Taksim Meydanı’nda yapılan bir ankete katılan 36 kişinin yaşlarını temsil etmektedir.

Tablo 5.1’deki verileri kullanarak yerleşim ölçüleri ile gerçekleştirdiğimiz hesaplamaların, dağılım ölçülerinde nasıl gerçekleştiğini inceleyelim. Ranj ve çeyrek ranj değerlerini daha önce 8 ve 4 olarak hesaplamıştık. Tablo 5.7’deki verilere göre hesapladığımızda 4 ve 2 olarak buluruz. Şimdi geri kalan dağılım ölçülerini her iki veriseti içinde hesaplamaya çalışalım.

Varyans değerini hesaplamak için, denklem (1)’de de görüldüğü gibi öncelikle ortalama değerinin hesaplanması gerekmektedir. Tablo 5.1’deki ortalama değerini 20 olarak hesaplamıştık, dolayısıyla varyans değerini aşağıdaki gibi hesaplayabiliriz.

$$s^2 = \frac{(15-20)^2 + (16-20)^2 + (16-20)^2 + (17-20)^2 + \dots + (24-20)^2 + (25-20)^2}{36-1} = \frac{210}{35} = 6$$

Tablo 5.2’deki verileri kullanarak frekans dağılımlarından varyans değerini nasıl hesapladığımızı inceleyelim.

Tablo 5.8

Yaşlar(x_i)	Frekanslar (f_i)	$(X_i - \bar{X})^2$	$f_i (X_i - \bar{X})^2$
15	1	$(15-20)^2=25$	$1 \times 25=25$
16	2	$(16-20)^2=16$	$2 \times 16=32$
17	3	$(17-20)^2=9$	$3 \times 9=27$
18	4	$(18-20)^2=4$	$4 \times 4=16$
19	5	$(19-20)^2=1$	$5 \times 1=5$
20	6	$(20-20)^2=0$	$6 \times 0=0$
21	5	$(21-20)^2=1$	$5 \times 1=5$
22	4	$(22-20)^2=4$	$4 \times 4=16$
23	3	$(23-20)^2=9$	$3 \times 9=27$
24	2	$(24-20)^2=16$	$2 \times 16=32$
25	1	$(25-20)^2=25$	$1 \times 25=25$
	$\sum_{i=1}^{11} f_i = 36$		$\sum_{i=1}^{11} f_i (X_i - \bar{X})^2 = 210$

Tablo 5.8’deki işlemleri yaparak varyans değerini aşağıdaki gibi hesaplayabiliriz:

$$s^2 = \frac{\sum_{i=1}^{n=11} f_i (x_i - \bar{x})^2}{\sum_{i=1}^{n=11} f_i - 1} = \frac{210}{36-1} = 6$$

Aynı prosedürü takip ederek Tablo 5.7'deki verilere ait olan varyans değerini de aşağıdaki gibi hesaplayabiliriz.

$$\begin{aligned} s^2 &= \frac{\sum_{i=1}^{n=11} f_i (x_i - \bar{x})^2}{\sum_{i=1}^{n=11} f_i - 1} \\ &= \frac{6 \times (18 - 20)^2 + 7 \times (19 - 20)^2 + 10 \times (20 - 20)^2 + 7 \times (21 - 20)^2 + 6 \times (22 - 20)^2}{36 - 1} \\ &= \frac{62}{35} = 1.771 \end{aligned}$$

Şekil 5.1 ve şekil 5.2'de yeralan histogramlardan da anlaşıldığı gibi, Tablo 5.7'deki veriler tablo 5.2'deki verilere göre daha az farklılaşmışlardır ve veriler ortalamasının etrafında toplanmışlardır. Bu tezimizi ikinci hesapladığımız varyans değerinin daha küçük olması da doğrulamaktadır. (dolayısıyla farklılaşma daha azdır) Her iki verisetine ait olan standart sapma değerlerini de, bulduğumuz varyans değerlerinin karekökünü alarak hesaplayabiliriz. Tablo 5.1 için:

$$s = \sqrt{6} = 2.449$$

Tablo 5.7 için ise:

$$s = \sqrt{1.771} = 1.331$$

Aynı ortalama değeri için açıkça Tablo 5.7'deki verilerin daha az farklılaştığı gözlenmektedir.

Her iki veriseti için farklılaşma katsayısını denklem (5.3)'ü kullanarak hesaplayabiliriz. Tablo 5.1 için:

$$\text{Farklılaşma Katsayısı} = \frac{2.499}{20} = 0.124$$

Tablo 5.7 için ise:

$$\text{Farklılaşma Katsayısı} = \frac{1.331}{20} = 0.066$$

Daha önce de belirttiğimiz gibi, farklı verisetleri farklı ortalama değerlerine sahip ise ve bu verisetlerinin dağılımını karşılaştırmak istiyorsak, farklılaşma katsayısı çok faydalı bir ölçüdür.

Yüzdelik oranların hesaplanması için 90. ve 10. yüzdelik değerlerini bulmamız gerekmektedir. Daha sonar bu yüzdelik değerlerinin birbirine olan oranını hesaplayarak yüzdelik oranını bulabiliriz. 90. ve 10. yüzdelik değerleri:

$$i = \frac{90}{100}(36+1) = 33.3' \text{üncü gözlem}$$

$$i = \frac{10}{100}(36+1) = 3.7' \text{inci gözlem}$$

33.3'üncü gözlemi 33'e, 3.7'inci gözlemi 4'e yuvarlayabiliriz. Bu gözlemlere tekabül değen gözlem değerleri 23 (tablo 5.1'den) ve 17 (tablo 5.1'den) olarak bulabiliriz. Tablo 5.1'deki verilere ait olan yüzdelik oranı:

$$\text{Yüzdelik Oranı} = \frac{23}{17} = 1.352$$

Tablo 5.1'deki verilere ait olan yüzdelik oranı:

$$\text{Yüzdelik Oranı} = \frac{22}{18} = 1.222$$

Frekans Dağılım Tablolarından Dağılım Ölçülerinin Hesaplanması (Gruplanmış Veriler):

Eğer elimizde Tablo 5.5'te olduğu gibi gruplanmış bir veriseti varsa, daha önce yerleşim ölçülerinde yaptığımız gibi, dağılım ölçülerini de hesaplarken orta noktalardan faydalanırız. Öncelikle Tablo 5.5'teki verilere ait olan varyans değerini hesaplayalım. Bu verilerin ortalamasını daha önce 20.5 ($\bar{x} = 20.5$) olarak hesaplamıştık, dolayısıyla artık varyansı bulmak için denklem (5)'i kullanabiliriz. Tek yapmamız gereken, hatırlayacağınız gibi, X_i değerlerini orta noktalarla değiştirmektir. Tablo 5.9 bu formülün süreçlerini göstermektedir.

Tablo 5.9

Yaşlar	Frekanslar (f_i)	Orta Noktalar (X_i)	$(X_i - \bar{X})^2$	$f_i (X_i - \bar{X})^2$
15-17	3	16	$(16 - 20.5)^2 = 20.25$	$3 \times 20.25 = 60.75$
17-19	7	18	$(18 - 20.5)^2 = 6.25$	$7 \times 6.25 = 43.75$
19-21	11	20	$(20 - 20.5)^2 = 0.25$	$11 \times 0.25 = 2.75$
21-23	9	22	$(22 - 20.5)^2 = 2.25$	$9 \times 2.25 = 20.25$
23-25	5	24	$(24 - 20.5)^2 = 12.25$	$5 \times 12.25 = 61.25$
25-27	1	26	$(26 - 20.5)^2 = 30.25$	$1 \times 30.25 = 30.25$
	$\sum_{i=1}^6 f_i = 36$			$\sum_{i=1}^6 f_i (X_i - \bar{X})^2 = 219$

Dolayısıyla varyans değeri:

$$s^2 = \frac{\sum_{i=1}^{n=6} f_i (x_i - \bar{x})^2}{\sum_{i=1}^{n=6} f_i - 1} = \frac{219}{36 - 1} = 6.25$$

daha önce hesapladığımızdan farklı olarak hesaplanır. Varyans

değerini 6.25 olarak hesaplarken elimizde gruplanmış veriler olduğundan, 6.25 yaklaşık bir varyans değeridir. Bu değerın hesapladığımız varyans değerlerinden farklı olmasının sebebi bu yakınsamadır. Diğer dağılım ölçüleri de aynı gruplanmış veri prosedürleri izlenerek bulunabilir.

Dağılımın Şekli ve Çarpıklığı

Daha önce merkezi eğilim hakkında fikir edinebilmek için ortalama, ortanca değer, ve mod değerlerinden faydalanmıştık. Verilerin dağılımı hakkında bize fikir veren çeşitli ölçüleri de inceledik. Şimdi ise dağılımın diğer bir karakteristiğini, çarpıklık derecesini ve nasıl ölçeceğimizi inceleyeceğiz.

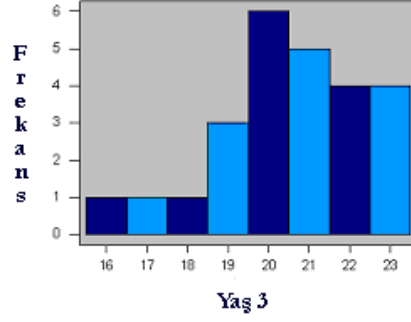
Verilerin çarpıklığına değinmeden önce simetrik dağılımlardan bahsetmek gerekmektedir. Simetrik dağılım, merkezin her iki tarafınınında eşit ölçüde aynı şekilde dağılmasıdır. Şekil 5.1 ve 5.3'te gördüğümüz dağılımlar birer simetrik dağılım örneğidir. Simetrik dağılımlarda, ortalama, ortanca değer ve mod her zaman birbirlerine eşittir. Eğer bir dağılım simetrik ise, bu dağılımın çarpıklığı yoktur veya çarpıklığı sıfırdır anlamına gelir.

Şimdi yaşların aşağıdaki gibi dağıldığını düşünelim.

Tablo 5.10

16,
17,
18,
19, 19, 19
20, 20, 20, 20, 20, 20
21, 21, 21, 21, 21
22, 22, 22, 22
23, 23, 23

Bu verilerden elde edeceğimiz histogram:



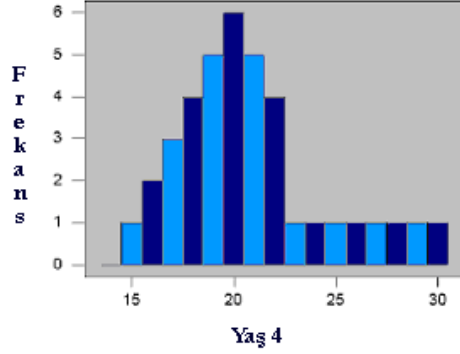
Şekil 5.3

Açıkça görüldüğü gibi dağılım sola doğru çarpık veya negatif çarpıklığa sahiptir.

Table 5.11

15
 16, 16
 17, 17, 17
 18, 18, 18, 18
 19, 19, 19, 19, 19
 20, 20, 20, 20, 20, 20
 21, 21, 21, 21, 21
 22, 22, 22, 22
 23
 24
 25
 26
 27
 28
 29
 30

Tablo 5.11'deki verilere ait olan histogram:



Şekil 5.4

Dağılım histogramdan da görüldüğü gibi sağa doğru çarpıktır veya pozitif çarpıklığa sahiptir. Herhangi bir çarpıklık durumunda, simetrik dağılım durumundan farklı olarak, ortalama, ortanca değer ve mod değerleri birbirlerine eşit değildir.

- ortalama > ortanca değer > mod: sağa doğru çarpık dağılımlar
- mod > ortanca değer > ortalama: sola doğru çarpık dağılımlar

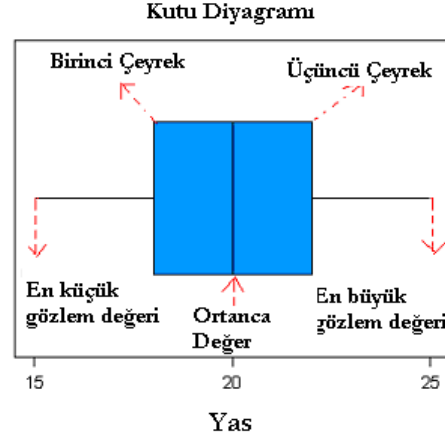
Sağa doğru çarpık dağılımlar, ortalama yüksek bir değer aldığı anda gerçekleşirken; sola doğru çarpık dağılımlar ise ortalamanın çok aşırı küçük değerler tarafından düşürüldüğü durumlarda ortaya çıkmaktadırlar. Veriler simetrik olduğunda aşırı bir değer ortaya çıkmadığından bu ölçüler birbirlerine eşit olmaktadır.

Örneğin, Tablo 5.10'daki verilere ait olan ortanca değeri = 20.5 > ortalama = 20.417, veriler negatif bir çarpıklığa sahiptir. Benzer şekilde Tablo 5.11'deki verileri ele alırsak, ortalama = 20.789 > ortanca değer = 20 olduğundan veriler pozitif bir çarpıklığa sahiptir.

Kutu Diyagramı

Kutu diyagramı, hem çarpıklığın hem de aşırı değerlerin teşhis edilmesinde çok faydalı bir araçtır. Kutu diyagramı, çeyreklikleri taban alan verisetinin genel çerçevesinin grafiksel olarak gösterimidir. Kutu diyagramını oluşturabilmemiz için, en küçük gözlem değeri, Q_1 birinci çeyreklik, ortanca değer, Q_3 üçüncü çeyreklik ve en büyük gözlem değeri olmak üzere sadece beş istatistik değerine ihtiyacımız vardır.

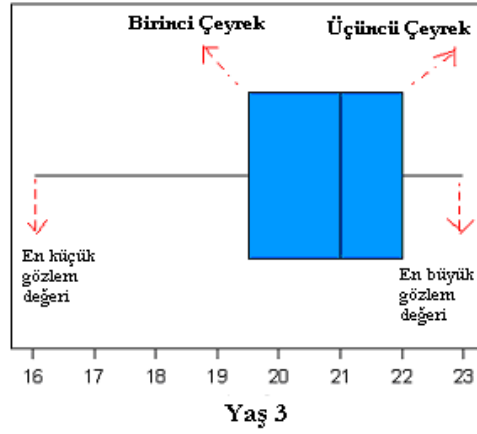
Tablo 5.1'deki verilere ait olan kutu diyagramı aşağıdaki gibidir.



Şekil 5.5

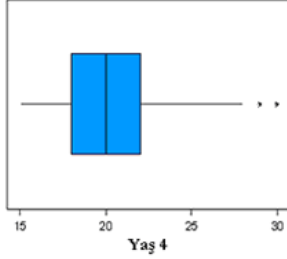
Kutu diyagramını çizmek için, öncelikle uygun ölçekte yatay eksenin oluşturulması gerekmektedir. Daha sonra, Q_1 'den Q_3 'e kadar kutuyu çizebiliriz. Bu kutunun içine ortanca değeri gösteren dikey bir çizgi çizeriz. Sonuç olarak, minimum ve maksimum değerlerini bu kutuya bağlayan yatay çizgiyi oluşturarak diyagramımızı tamamlarız.

Benzer şekilde Tablo 5.10'a ait olan kutu diyagramını çizersek:



Şekil 5.6

Bu diyagramdan dağılımın sola doğru çarpık olduğu gözükmektedir. Bu sonuca nasıl varabiliriz? Kutunun solunda yer alan birinci çeyrekte 19.5 (Q_1) minimum değer 16'ya kadar olan doğru, üçüncü çeyrekte 22 (Q_3) maksimum değer 23'e kadar olan doğrudan daha uzun olduğundan dağılım sola çarpıktır. Benzer şekilde Tablo 5.11'e ait olan kutu diyagramını çizersek:



Şekil 5.7

Şekil 5.7'deki kutu diyagramı pozitif bir dağılıma sahiptir. Bu diyagramda yer alan iki yıldız (**), verisetinde iki tane aşırı büyük değer olduğunu göstermektedir.

Standart Sapmanın Kullanımı: Ampirik Kural

Çoğu verisetlerinde, gözlemlerin büyük bir kısmı ortanca değer etrafında yer alırlar. Sağa doğru yatık verisetlerinde, verilerin kümelenmesi ortanca değer solunda gerçekleşirken sola yatık verisetlerinde verilerin kümelenmesi ortanca değer sağında yer almaktadır. Simetrik olan verisetlerinde ise ortanca değer ve ortalama aynıdır. Gözlemler eşit olarak bu ölçüm parametrelerinin etrafında dağılırlar. Gözlemlerin bu tarz kümelenmelerinde kullanılan ve standart sapmadan

da faydalanan kurala **ampirik kural** denir.

Ampirik kurala göre, çoğu verisetlerinde, her üç gözlemde ikisi (%67'si) ortalama ve standart sapma uzaklığındaki aralıkta yer alırlar ve kabaca gözlemlerin (90%-95%)'i ise ortalama ve 2 standart sapma uzaklığındaki aralıkta yer alırlar. Bu yüzden standart sapma ve ortalama olan farklılaşma bize gözlemlerin nasıl dağıldığı hakkında kabaca bilgi verebilir.

Örnek olarak, Tablo 5.1'deki verinin ortalaması 20 ve standart sapması 2.449 olduğunu düşünerek ampirik kuralı uygulayalım. Ampirik kurala göre gözlemlerin %67'si aşağıdaki aralıkta yer almaktadır.

$$(\bar{x} - 1 \times s; \bar{x} + 1 \times s) = (20 - 1 \times 2.449; 20 + 1 \times 2.449) = (17.551; 22.449)$$

Hakikaten bu gözlemlerin %67'si yukarıdaki aralıkta yer almaktadır. Bu aralıkta yer almayan aşağıdan üç gözlem (15, 16, 17), yukarıdan da üç gözlem (23, 24, 25) bulunmaktadır.

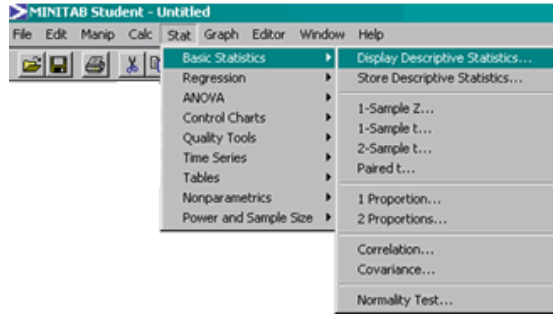
Gözlemlerin %95 ise aşağıdaki aralıkta yer alır.

$$(\bar{x} - 2 \times s; \bar{x} + 2 \times s) = (20 - 2 \times 2.449; 20 + 2 \times 2.449) = (15.102; 24.898)$$

Minitab Uygulamaları

Bu bölümde işlediğimiz yerleşim ölçüleri ve dağılım ölçülerinin çoğu MINITAB programı ile hesaplanabilir. Tablo 5.1, 5.7, 5.10, 5.11'de yer alan veriler "yaslar.mtw" dosyasına Yas1, Yas2, Yas3, Yas4, adı altında kaydedilmiştir. (Tablo 5.1'deki veriler Yas1 adı altında, tablo 5.7'deki veriler Yas2 adı altında kayıtlıdır.) N

Örnekleme istatistiklerini hesaplamak için, “Stat” (istatistik) menüsünden “Basic Statistics” (Temel istatistikler) seçeneğine ve “Display Descriptive Statistics” (Betimsel İstatistikleri Göster) komutunu seçiniz.



Şekil 5.8

Karşınıza çıkacak olan diyalog ekranından değişken olarak “yas1” değişkenini seçiniz ve “OK” tuşuna basarak aşağıdaki çıktı elde edebilirsiniz.

Descriptive Statistics					
		Örnekleme Boyutu		Düzeltilmiş Ortalama	
Variable	N	Mean	Median	Tr Mean	StDev
Age	36	20,000	20,000	20,000	2,449
Variable	Minimum	Maximum	Q1	Q3	SE Mean
Age	15,000	25,000	18,000	22,000	0,408

Ortalamanın Standart Hatası $\frac{\sigma}{\sqrt{n}}$

Şekil 5.9

N: Gözlem sayısını

Tr Mean: Düzeltilmiş Ortalamayı

StDev: Standart sapmayı

SE Mean: Ortalamanın standart hatasını

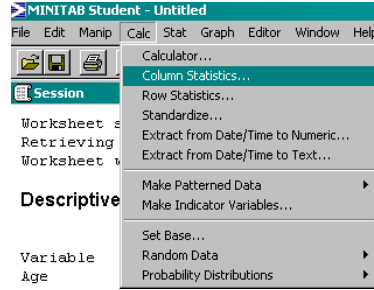
Minimum: Verisetindeki en küçük gözlem değerini

Maximum: Verisetindeki en büyük gözlem değerini

Q₁: Birinci çeyrekliği

Q₃: Üçüncü çeyrekliği

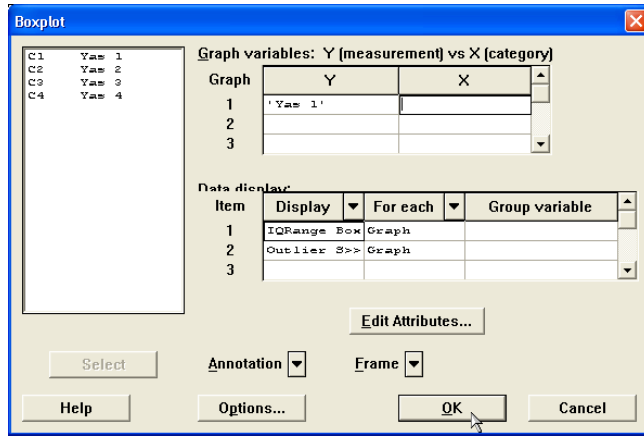
temsil etmektedir. Alternatif olarak “Row Statistics” (sıra istatistikleri) veya “Column Statistics” (sütun istatistikleri) değerlerini de kullanarak aynı ölçüleri hesaplayabiliriz. “Calc” (hesapla) menüsünden (bkz. Şekil 5.10) bu iki seçenekten birini kullanarak aynı ölçüleri hesaplayabilirsiniz.



Şekil 5.10

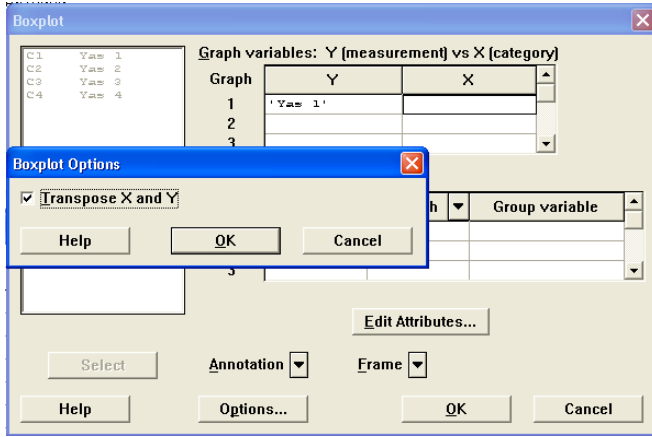
Gruplanmış verilerin olduğu durumlarda, orta noktalarını MINITAB ortamında ayrı bir sütun olarak belirtiniz.

Kutu diyagramına örnek olması amacıyla şekil 5.4’teki kutu diyagramını MINITAB yardımı ile çizebilmek için her zaman olduğu gibi öncelikle “graph” (grafik) menüsünden “Boxplot” (kutu diyagramı) seçeneğine tıklayınız. Karşınıza çıkan diyalog ekranında Yas1 değişkenini seçerek Y eksenine atayınız.



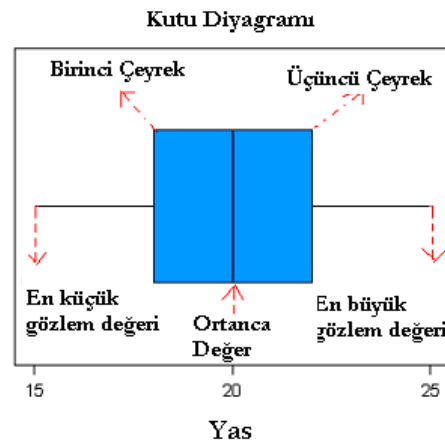
Şekil 5.11

Şekil 5.11’deki diyalog ekranında “options” (seçenekler) seçeneğinden “transpose X and Y button” (X ve Y eksenlerinin devriğini al) özelliğine tıklayınız ve “OK” tuşuna basınız.



Şekil 5.12

Tekrar “OK” tuşuna bastığınızda Şekil 5.13’teki kutu diyagramını elde edebilirsiniz.



Şekil 5.13

Bölüm 6:

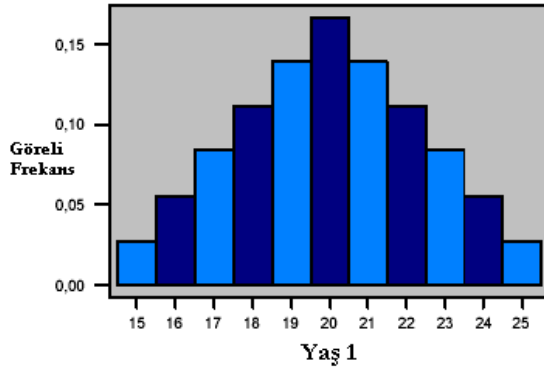
Olasılık ve Olasılık Dağılımları

Olasılık, hakkında başlı başına kitap yazılacak kadar önemli ve geniş bir konudur. Olasılık bilgisinin sadece bir kısmı, istatistiksel süreçleri anlamak için yeterlidir. Bu yüzden olasılık konusunda çok fazla detaya inmeden sadece istatistikle olan alakasını irdelleyeceğiz.

Olasılıktan bahsettiğimizde kesin olmayan bir olaydan bahsediyoruz demektir. Kesin olan olayların olasılıkları 1'e eşittir. Örneğin, bir bozuk para attığımızda (eğer dünyada iseniz) kesin bir şekilde bu para yere düşecektir. Ancak paranın yazıyı turanı geleceği kesin değildir ve paranın yazı gelme olasılığı da, tura gelme olasılığı da $1/2$ dir. Neden $1/2$? Bu sorunun cevabını birazdan öğreneceğiz.

Histogram ve Olasılık

Bir önceki bölümden de hatırlayacağınız üzere, istatistiksel analizde en çok kullanılan diagramlardan biri de histogramdır. Histogramda yaptıklarımızı tekrar hatırlayacak olursak; sınıf aralıkları, bize verinin frekans dağılımı şeklinde özetlenmiş haliydi. Her sınıfın frekansı yatay eksene herbirinin yüksekliği frekansı ifade edecek şekilde dikdörtgenler çizilerek bulunmaktaydı. Histogram, görel frekansları kullanarak da elde edilebilir. Bu durumda ise histogram Şekil 6.1'deki gibi olur.⁴



Şekil 6.1

Şimdi ise dikey ekseninde yüzde cinsinden görel frekanslar yer almaktadır.(örneğin 5% ifadesi 0.05 şeklinde yazılmıştır.). Örneğin, 36 kişilik gruptan rastgele 1 kişiyi seçelim. Bu kişinin yaşının 15 olma olasılığı $1/36 = 0.027$ dir ve hem görel frekansı hem de olasılığı temsil eder. Bu hesaplamaları neye göre yapıyoruz bunlara nasıl karar veriyoruz ? Kısaca bir sonraki bölümde değineceğiz.

Olasılık Tanımına Farklı Yaklaşımlar

⁴ Şekil 6.1'i MINITAB ortamında elde edebilmek için, "Histogram" diyalog ekranında, "options" (seçenekler) bölümüne tıkladıktan sonra "density" (yoğunluk) kutucuğuna tıklayınız ve "OK" tuşuna basınız.

İşlem ve uygulamalara geçmeden önce bazı temel tanımlamalara giriş yapalım. Olasılık problemlerini incelerken; deneylerle, olaylarla ve olası olayların koleksiyonu ile ilgileniriz. Deney, araştırmacının gözlemleri elde etmek için başvurduğu bir metodtur. Olay ise, deneyin sonuçlarından veya çıktılarından oluşan bir demettir. Basit bir olay parçalanamaz veya daha ileriye götürülemez. Örneklem uzayı, olası bütün olayları içeren bir deneydir. Örneğin, bir zar atalım bu bir deneydir. Bu işlem sonucunda 3 gelmesi ise, basit bir sonuçtur, basit bir olaydır, daha ileri götürülemez. Zar atma deneyi 1, 2, 3, 4, 5 ve 6 gibi basit olaylardan meydana gelir.

Eğer P'yi olasılık olarak ve A, B, C'yi ise kendine özgü birer olay olarak tanımlarsak; $P(A)$, A olayının olma olasılığını ifade eder. Olasılık kavramı ile alakalı iki temel yaklaşım vardır.

- **Görelî Frekans Yaklaşımı:**

Öyle bir deney tasarlayalım ki çok fazla uygulandığında, A olayının kaç kere meydana geldiğini gözlemleyebilelim. Buna bağlı olarak deneyimizde; A olayının olma olasılığı $P(A)$ aşağıdaki gibi tahmin edilir:

$$P(A) = \frac{\text{A olayının kaç kez meydana geldiği}}{\text{Rassal denemenin kaç kez tekrarlandığı}} = \frac{f}{n}$$

Yazı ve tura örneğimize geri dönelim. Parayı havaya attığımızda yazı gelme olasılığını hesaplamak için, öncelikle parayı, örneğin 1000 defa, havaya atmalıyız ve bu deney esnasında kaç kere yazı geldiğini gözlemlemeliyiz. Eğer, bu atışların 500'ünde yazı geliyorsa, parayı havaya attığımızda yazı gelme olasılığı $500/1000 = 1/2$ olur. Başka bir örnek ile açıklamak için, 36 öğrenciden yaşı 15 olanları 36000 defa seçelim. 36000 seçimimizden 1000 tanesinde yaşı 15 olan öğrenciler ile karşılaştığımızı gözlemlemişsek, 36 öğrenci arasından yaşı 15 olan öğrencileri seçmek olasılığımızı $1000/36000=1/36$ olarak hesaplayabiliriz.

- **Klasik Yaklaşım:**

Deneyimizin, n tane farklı olaya sahip olduğunu ve herbirinin gerçekleşme şansının eşit olduğunu varsayalım. Eğer A olayı, N farklı şekilde n kere gerçekleşiyorsa;

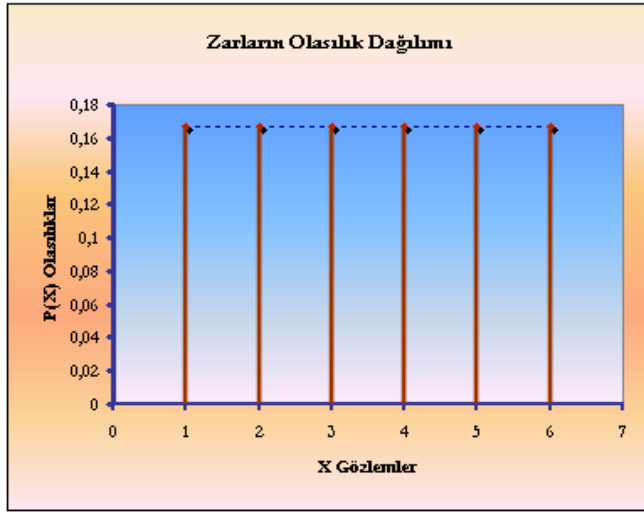
$$P(A) = \frac{\text{A Olayının Kaç Farklı Şekilde Gerçekleşebileceği}}{\text{Mümkün Olan Sonuç Sayısı}} = \frac{n}{N}$$

n, olayın kaç farklı şekilde gerçekleşebileceğini gösterirken; N ise, olası mümkün olan sonuç sayısını göstermektedir. Bu yaklaşımından yola çıkarak, para örneğimizdeki olasılığı $1/2$, öğrenci seçme örneğimizdeki olasılığı da $1/36$ olarak hesaplayabiliriz.

Klasik yaklaşımda her olayın olma olasılığı eşit olmalıdır. Eğer bir deneyde bu şart sağlanmazsa, görelî frekans yaklaşımı uygulanır. Deneyde, toplam gözlem sayısı arttıkça, yaklaşık olasılık değerleri gerçek değere doğru yaklaşma eğilimi gösterir. Bu özellik, büyük sayılar yasası olarak bilinen teoremden kullanılır. Bu yüzden, deney tekrar tekrar yapıldığında görelî frekans olasılığı, gerçek olasılık değerine yaklaşır. Eğer olasılık tahmini az sayıdaki gözleme dayandırılırsa, gerçek olasılık değerinden sapmalar görülebilir.

Olasılık Dağılımları

Olasılık dağılımı nedir? Olasılık dağılımı, örneklem uzayında yer alan her olayı tek tek gösteren bir çeşit fonksiyonun (olasılık fonksiyonu) grafiğine verilen isimdir. Bu tanımı biraz daha açıklığa kavuşturmak amacıyla zar atma örneğini tekrar ele alalım. Örneklem uzayı, örneğimizde hatırlanacağı gibi şu şekilde idi: $S = \{1, 2, 3, 4, 5, 6\}$. Olasılık dağılımı ise bize bu örneklem uzayındaki her olayın tek tek olasılığını göstermelidir. Dolayısıyla olasılık dağılımını şu şekilde gösterebiliriz: $P = \{1/6, 1/6, 1/6, 1/6, 1/6, 1/6\}$. Çünkü örneklem uzayındaki her altı rakam da $1/6$ olasılık ile gerçekleşebilir. Bu olasılık dağılımı Şekil 6.2'de olduğu gibi grafik ile gösterilebilir.



Şekil 6.2

Bu şekil, zar altıldığında her değişkenin olası alabilecekleri göstermektedir (1, 2, 3, 4, 5, 6 gibi). Bu tarz olasılık dağılımları, yeknesak (uniform) olasılık dağılımı olarak adlandırılırlar. Çünkü her bir olay eşit olasılıkla gerçekleşebilmektedir. Bir diğer olasılık dağılımına örnek ise şekil 6.1'de verilmektedir. Şekildeki histogram $S = \{15, 16, 17, 18, 19, 20, 21, 22, 23, 24, 25\}$ örnekleme ait olan olasılık değerlerini göstermektedir.

Olasılık Yoğunluk Fonksiyonu

Olasılık yoğunluk fonksiyonu, rassal bir değişkenin olasılıklarını bize gösteren fonksiyondur. Olasılık yoğunluk fonksiyonu ve olasılık dağılımı arasındaki ilişki, matematikteki bir fonksiyon ve onun grafiği arasındaki ilişki ile aynı doğrultudadır. Bir başka ifadeyle, olasılık yoğunluk fonksiyonu olasılık dağılımını bize veren matematiksel ifadedir ve matematikte karşılaşılabileceğiniz herhangi bir fonksiyondan farklı bir yapıya sahip değildir. Ancak daha ileride de göreceğimiz gibi bazı gerekli koşulları sağlaması gerekmektedir. Bu tartışmanın ardından şu soruyu sorabiliriz: Şekil 6.2'de çizilen olasılık dağılımını nasıl bir matematiksel ifade ile gösterebiliriz. Bir diğer deyişle Şekil 6.2'deki olasılık dağılımına karşılık gelen olasılık yoğunluk fonksiyonu nedir? Bu matematiksel ifade de denklem (6.1)'deki gibidir.

$$f(x) = \begin{cases} \text{eğer } x = 1, & f(x) = 1/6 \\ \text{eğer } x = 2, & f(x) = 1/6 \\ \text{eğer } x = 3, & f(x) = 1/6 \\ \text{eğer } x = 4, & f(x) = 1/6 \\ \text{eğer } x = 5, & f(x) = 1/6 \\ \text{eğer } x = 6, & f(x) = 1/6 \end{cases} \quad (6.1)$$

Veya olasılık sembolleri ile ifade edecek olursak

$$P(x) = \begin{cases} \text{eğer } x = 1, & P(x) = 1/6 \\ \text{eğer } x = 2, & P(x) = 1/6 \\ \text{eğer } x = 3, & P(x) = 1/6 \\ \text{eğer } x = 4, & P(x) = 1/6 \\ \text{eğer } x = 5, & P(x) = 1/6 \\ \text{eğer } x = 6, & P(x) = 1/6 \end{cases} \quad (6.2)$$

Bu fonksiyonun hangi özellikleri onun bir olasılık yoğunluk fonksiyonu olmasını sağlamaktadır? Temel olarak iki özellik vardır: Bunlardan birincisi aşağıdaki gibi ifade edilebilir.

$$P(x_1), P(x_2), P(x_3), P(x_4), P(x_5), P(x_6) \geq 0$$

Bir başka deyişle olasılıklar sıfır veya pozitif değerler almalıdır (yani negatif değer alamazlar). Burada $x_1 = 1, x_2 = 2, \dots$

İkinci özellik ise matematiksel olarak aşağıdaki gibi ifade edilebilir

$$\sum_{i=1}^{n=6} f(x_i) = f(x_1) + f(x_2) + f(x_3) + f(x_4) + f(x_5) + f(x_6) = 1 \text{ veya}$$

$$\sum_{i=1}^{n=6} P(x_i) = P(x_1) + P(x_2) + P(x_3) + P(x_4) + P(x_5) + P(x_6) = 1$$

Bu koşul bir olasılık fonksiyonun bütün değişkenlerine göre toplamı alındığında bu toplamın 1'e eşit olması gerekliliğini ifade etmektedir.

Bu iki koşul, bir önceki dersimizde gördüğümüz, olasılık kurallarının bir başka ifadesinden başka bir şey değildir. Açık olarak her iki koşulda, zar örneğimizde sağlanmaktadır.

$$\sum_{i=1}^{n=6} P(x_i) = \frac{1}{6} + \frac{1}{6} + \frac{1}{6} + \frac{1}{6} + \frac{1}{6} + \frac{1}{6} = 1$$

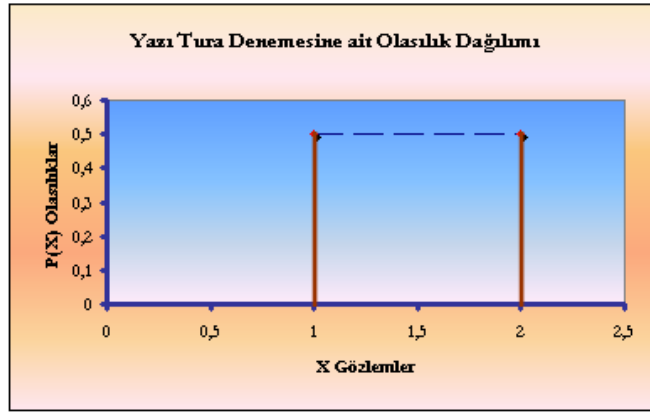
Olasılık fonksiyonlarının diğer fonksiyonlardan ayrılan bir diğer özelliği de, rassal bir değişkenin açıklayıcı değişken olarak fonksiyonda yer almasıdır. Rassal değişken nedir? Normal bir değişkenden farkı nedir? Bu konuyu bir örnek yardımıyla açıklamaya çalışalım. Açık bir şekilde gözüktüğü gibi, şekil 6.1'de ve olasılık fonksiyonlarımızda (Denklem (6.1) veya (6.2)) rassal değişkenimizi (zar atımı) göstermek için X sembolünü kullandık. Rassal değişken X ile ifade

edilirken, X 'in gerçekleşen değerleri küçük x harfi ile gösterilmektedir. Bu örnekteki rassal değişkenimiz X 'in gerçekleşen değerleri, $x = 1, x = 2, \dots, x = 6$ sayılarından (her biri $1/6$ olasılıkla olmak üzere) herhangi birisi olabilir. Örneğin zarı attığınızda deneyin gerçekleşen değeri olarak 3 gelebilir.

Bazı rassal denemelere göre, sonuçlar rakamsal bir ifade olmayabilir. Örneğin, istatistikte sıkça kullanılan bir örneği, yazı-tura atma rassal denemesini ele alalım. Bu denemenin sonucu, ya yazı gelmesi ya da tura gelmesidir. Bu durumda yapılabilecek tek şey rassal değişkenimizi X olarak tanımlayarak, yazı gelmesi durumunda X 'in 1 değerini alacağını ($x = 1$), tura gelmesi durumunda X 'in 2 değerini aldığını varsaymamızdır. Bu durumda olasılık fonksiyonumuzu aşağıdaki gibi yazabiliriz.

$$P(X) = \begin{cases} x = 1 \text{ ise } P(x) = 1/2 \\ x = 2 \text{ ise } P(x) = 1/2 \end{cases} \quad (6.3)$$

Olasılık dağılımı ise aşağıdaki şekli alır



Şekil 6.3

Bir rassal değişken, alacağı sayısal değışkene göre, sürekli veya süreksiz olarak sınıflandırılabilir. Birinci bölümden de hatırlanacağı gibi, rassal süreksiz değışkenler ancak sonlu asal sayıları gerçekleşen değerleri olarak kabul edebilirler (örneğin: 1, 2, 3, 4...). Bir malın satış miktarları, bir mağazaya giren müşteri sayısı, bir zar atıldıktan sonra gelebilecek sayılar süreksiz rassal değışkenlere birer örnek teşkil ederler. Ağırlık, uzunluk, ısı derecesi veya zaman gibi gibi rassal değışkenler sürekli rassal değışkenlerdir. Çünkü gerçekleşen değerlerini bir aralık içerisinde tanımlayabiliriz. Örneğin hava sıcaklığı bu yörede -25 derece ile +100 derece arasında değışir denilebilir.

Süreksiz Olasılık Yoğunlukları

Süreksiz bir olasılık fonksiyonu, isminden de tahmin edilebileceği gibi, süreksiz bir rassal değışkenin, açıklayıcı değışken olduğu bir olasılık fonksiyonudur. Açıklayıcı değışkeni süreksiz bir değışken olduğundan, zar örneği süreksiz olasılık dağılımına verilebilecek bir örnektir. (alabileceği rakamlar 1, 2, 3, 4, 5, ve 6 ile sınırlandırılmıştır, Denklem 6.1 ve 6.2). Süreksiz bir olasılık dağılımına bir başka örnek ise Denklem 6.3 tarafından verilmiştir. Bu fonksiyonların olasılık dağılımları ise Şekil 6.2 ve Şekil 6.3'te gösterilmiştir.

Örnek 6.1: 300 hane halkını kapsayan bir araştırmaya göre, 54 hane halkının çocuğu bulunmamaktadır, 72 aile 2 çocuğa sahiptir, 42 ailenin 3 çocuğu vardır, 12 aile 4 çocukludur ve 3 ailenin ise 5 çocuğu bulunmaktadır. Eğer bu hane halklarından rastgele bir tanesini seçecek olursak, bu hane halkının kaç çocuğa sahip olma olasılığını hesaplayabiliriz. Tablo 6.1 bize gerekli olan bilgiyi özetlemektedir. Tabloda X değişkenini çocuk sayısını göstermek amacıyla kullandığımızı belirtirsek, $f(x)$ 'inde bu seçilecek ailenin farklı x değerleri için olasılık değerlerini gösterdiği anlaşılabilir. Örneğin $f(1)$ değeri bize rastgele seçilen ailenin 1 çocuklu olma olasılığını vermektedir. Bu örnekte bu olasılık $117/300 = 0.39$ 'a eşittir.

Bir olasılık dağılımının sahip olması gereken koşulların (yani olasılıkların negatif olamaması ve toplamlarının bire eşit olması) bu örnek için de sağlanmış olduğu Tablo 6.1'den kolaylıkla görülebilir.

Tablo 6.1: Hane Halkı Başına Düşen Çocuk Sayısının Olasılık Dağılımı

x	F	f(x)
0	54	54/300=0.18
1	117	117/300=0.39
2	72	72/300=0.24
3	42	42/300=0.14
4	12	12/300=0.04
5	3	3/300=0.01
	$\sum f = 300$	$\sum f(x) = 1$

Örnek 6.2: İkinci bir örnek olarak biri kırmızı diğeri yeşil renkte olan iki zarı beraber attığımızı düşünelim. Bu örnekte rassal değişkenimizin iki zarın üzerindeki sayıların toplamından ibaret olduğunu varsayalım. Bu durumda rassal değişkenimiz X 'in gerçekleşen değerleri, $x=2, x=3, x=4, \dots, x=12$, sayılarından ibaret olabilir. Tablo 6.2 iki zar atıldığında gerçekleşebilecek X değerlerini bir tablo halinde sunmaktadır.

Tablo 6.2: Bir çift zarın toplamının alabileceği değerler

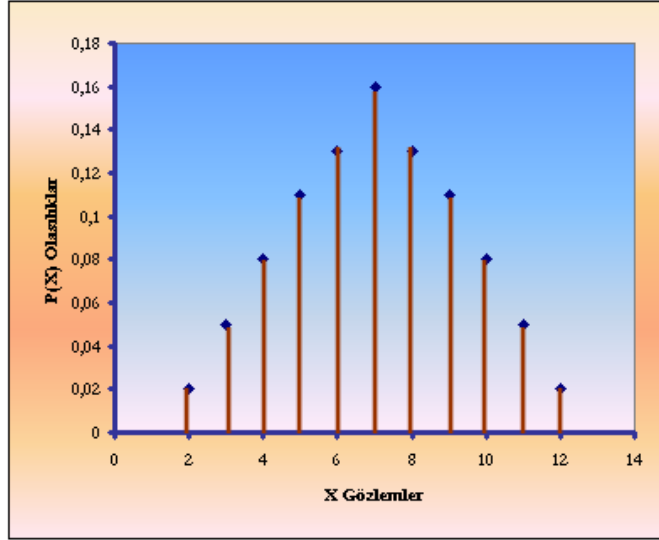
↓ Kırmızı / Yeşil →	1	2	3	4	5	6
1	2	3	4	5	6	7
2	3	4	5	6	7	8
3	4	5	6	7	8	9
4	5	6	7	8	9	10
5	6	7	8	9	10	11
6	7	8	9	10	11	12

Tablo 6.3: Bir çift zarın sıklık ve olasılık değerleri

x	f	f(x)
---	---	------

2	1	1/36=0.02
3	2	2/36=0.05
4	3	3/36=0.08
5	4	4/36=0.11
6	5	5/36=0.13
7	6	6/36=0.16
8	5	5/36=0.13
9	4	4/36=0.11
10	3	3/36=0.08
11	2	2/36=0.05
12	1	1/36=0.02
	$\sum f = 36$	$\sum f(x) = 1$

Olasılık dağılımını ise aşağıdaki gibi çizebiliriz.



Şekil 6.4

Sürekli Rassal Değişkenin Ortalaması ve Standart Sapması

Rassal bir değişkenin olasılık dağılımının ortalamasına beklenen değer denir ve $E(X)$ veya μ ile gösterilir. Eğer değişkenin olasılık dağılımını biliyorsak, anakütleyi de bildiğimiz anlamına gelir ve bu değişkenin ortalamasına da beklenen değer denir. Bu yüzden beklenen değer $E(X)$ veya μ ile gösterilerek anakütle ortalamasını ifade eder. Anakütle ortalamasını veya beklenen değerini bulmak için, göreceli frekansların p_i olasılıklar olarak kabul edildiği (6.4) formülü kullanılır.

$$E(x) = \mu = \sum P(x_i)x_i \quad (6.4)$$

Zar atma örneğimizdeki ortalamayı, (6.4)'ü kullanarak, aşağıdaki gibi hesaplayabiliriz:

$$\mu = 1 \times \frac{1}{6} + 2 \times \frac{1}{6} + 3 \times \frac{1}{6} + 4 \times \frac{1}{6} + 5 \times \frac{1}{6} + 6 \times \frac{1}{6} = 3.5$$

Dolayısıyla, zarı attığımızda 3.5 gelmesini bekleriz. Açıkçası bu sayıda zar deneyinde gözlenebilir bir sayı değildir. (hiçbir zaman zarı attığınızda 3.5 elde etmezsiniz). Bunun yanında eğer zarı çok fazla atarsanız, mesela 100.000 kere gibi, bu 100.000 atışın da ortalaması 3.5'a eşit olmalıdır. Bu nedenden dolayı beklenen değere bazen uzun-vadeli ortalama da denilmektedir. Bu ortalama daha önceki bölümlerde bahsedilen ortalamalardan, tahmin edileceği üzere farklı bir ortalama değildir. $\bar{X}$ sadece örneklem ortalamasını ifade eder. Örneklem ortalamasını hesaplamak için zarı birkaç kez havaya atarız. Öncelikle beş kere havaya attığımızı düşünelim ($n = 5$) ve bu deney sonucunda 2, 4, 5, 3, 6 sonuçları çıkmış olsun. Dolayısıyla, örneklem ortalamasını 4 olarak buluruz. $\frac{\sum_{i=1}^n x_i}{n} = 4$. Bu ortalama da daha önce hesaplanan beklenen değer (ortalama) olan 3.5'tan farklıdır.

Şimdi, Tablo 6.1'deki verilerden faydalanarak beklenen değeri hesaplayalım:

$$\mu = 0 \times 0.18 + 1 \times 0.39 + 2 \times 0.24 + 3 \times 0.14 + 4 \times 0.04 + 5 \times 0.01 = 1.5$$

Bu arada, Tablo 6.1'deki verilerin anakütleyle ait olduğunu kapalı bir şekilde varsaydık. Örneğin anakütle 300 tane hane halkı içermektedir.

Beşinci bölümde s^2 ile ifade edilen varyans hesaplamasından bahsetmiştik. Eğer bir değişkenin olasılık dağılımı hakkında fikir sahibi isek, anakütle varyansı hakkında da fikir sahibi olabiliriz. Süreksiz rassal bir değişkenin anakütle varyansını (6.5) deki gibi hesaplayabiliriz. (anakütle varyansını σ^2 ile ifade ederiz)

$$\text{var}(x) = \sigma^2 = \sum (x - \mu)^2 \times P(x) \quad (6.5)$$

Tablo 6.4: Hanehalkındaki çocuk sayısının varyansının hesaplanması

X	$x - \mu$	$(x - \mu)^2$	P(x)	$(x - \mu)^2 \times P(x)$
0	-1.5	2.25	0.18	0.4050
1	-0.5	0.25	0.39	0.0975
2	0.5	0.25	0.24	0.0600
3	1.5	2.25	0.14	0.3150
4	2.5	6.25	0.04	0.2500
5	3.5	12.25	0.01	0.1225
				$\sum (x - \mu)^2 \times P(x) = 1.25$

Buna bağlı olarak, $\sigma^2 = 1.25$ tir ve bu ifadenin karekökü her hanehalkında yeralan çocuk sayısının standart sapmasını verir:

$$\sigma = \sqrt{1.25} = 1.118$$

Bu analizin aynısını tek zar ve çift zar örneklerinde de uygulayabilirsiniz.

Binom Dağılımı

Örneklerimizi çoğaltabiliriz, ancak daha fazla ilerlemeden, şu soruyu sormalıyız: İstatistikte bu olasılık dağılımlarını hangi amaçla kullanmaktayız? Örneklerimizden de açık olarak görülebileceği gibi, eğer bir rassal değişken X 'in olasılık dağılımını biliyorsak, bu değişken hakkında sorulabilecek aşağıdaki gibi bazı sorulara rahatlıkla cevap verebiliriz.

1. X 'in belli bir sayıya eşit olma olasılığı nedir? Örneğin, zar örneğinde $x = 3$ olma olasılığı nedir? ($P(x = 3) = ?$), veya iki zarın bulunduğu bir örnekte $x = 12$ olma olasılığı nedir? ($P(x = 12) = ?$).

2. X 'in belli bir sayıdan büyük veya küçük olma olasılığı nedir? Örneğin, zar örneğinde $x > 3$ olma olasılığı nedir? ($P(x > 3) = ?$), veya iki zarın bulunduğu bir örnekte $x > 12$ olma olasılığı nedir? ($P(x > 12) = ?$).

Yukarıdaki soruları benzerleri ile çoğaltabiliriz. Ancak biz işletme istatistiğinde genellikle zar atma olayı gibi olayların olasılıklarıyla pek fazla ilgilenmemekteyiz. İşletme istatistiğinde genellikle, işletme yönetiminde karar almamıza yardımcı olabilecek başka olasılık dağılımlarıyla ilgilenmekteyiz. Bu olasılık dağılımlarının en ünlülerinden biri de Binom dağılımıdır.

Yüzyıllardan beridir kullanılan binom dağılımı, en iyi bilinen olasılık dağılımlarından biridir. İsmindeki bi-ifadesinden de anlaşılabilceği gibi (İngilizce'de bi- öntakısı ikili anlamına gelmektedir), burada ilgilendiğimiz iki sonuç vardır. Bu sonuçlardan bir tanesi başarı diğeri ise başarısızlık olarak nitelendirilir. Ancak bu adlandırmalar kelime anlamalarında genellikle kullanılmamaktadır. Başarı, sonucu mutlaka istenilen bir durumu ifade ederken, başarısızlık da istenmeyen durumu ifade etmektedir.

Binom dağılımı bize n sayıda denemede X başarılı olay yakalama olasılığını verir. Ancak binom dağılımının bunu yapabilmesi için aşağıdaki varsayımların gerçekleşmesi gerekmektedir.

1. Rassal denemenin n adet aynı denemeyi içermesi gerekmektedir.
2. Her deneme sadece iki olası sonuç barındırmaktadır (başarı veya başarısızlık)
3. Her deneme bir öncekinden bağımsızdır.
4. Eğer p terimi herhangi bir denemede ki, başarı olasılığını gösteriyor ise; ($q = 1 - p$)'da başarısızlık olasılığını göstermektedir.
5. p ve q olasılıkları denemeler boyunca aynı kalmaktadırlar.

Biraz daha ileriye gitmeden, faktöryel kavramını tanımlamamız gerekmektedir. k faktöryelini şu biçimde tanımlıyoruz:

$$k! = k \times (k-1) \times (k-2) \times (k-3) \times \dots \times 1$$

Sıfırın faktöryeli bire eşittir: $0! = 1$.

Binom katsayılarını da burada kullanmak zorundayız. n sayıdaki denemede oluşabilecek X sayıda, başarının elde edilebileceği farklı durumların sayısı aşağıdaki formül yardımıyla hesaplanabilir.

$$\frac{n!}{x!(n-x)!} \quad (6.6)$$

Bu formül genellikle kombinasyon formülü olarak bilinir ve X başarının n denemede farklı durumlarının sayısını verir. Bu noktayı bir örnek yardımıyla açıklamamız faydalı olacaktır. Varsayalım ki bir mağazadan herhangi bir müşterinin alışveriş yapma olasılığı 0.3'tür. Mağazaya girecek ilk üç müşteriden iki tanesinin alışveriş yapma olasılığı nedir? Bir başka ifade ile 3 denemede ($n = 3$, müşteri sayısı) iki başarı elde etme ($x = 2$, alışveriş yapan müşteri sayısı) olasılığı nedir? Yukarıdaki formülü kullanarak

$$\frac{3!}{2!(3-2)!} = \frac{3 \times 2 \times 1}{2 \times 1 \times 1} = \frac{6}{2} = 3$$

Mağazaya girecek olan ilk üç müşteriden ikisinin alışveriş yapma, diğerinin alışveriş yapmama olasılığını aşağıdaki gibi bulabiliriz⁵:

$$p \times p \times (1-p) = p^2 \times (1-p)$$

Eğer $p = 0.3$ ise bu olasılık 0.063 dir. Bu olasılık Tablo 6.5'de koyu harflerle belirtilen her üç durum için de geçerlidir. Dolayısıyla, n denemede x sayıda başarı veren her durum eş olasılığa sahiptir. Genel olarak n denemede x başarı elde etme olasılığı aşağıdaki formülle hesaplayabiliriz.

$$p^x \times (1-p)^{n-x} \quad (6.7)$$

Denklem (6.6) X başarıya sahip binom deneyinin sonuç sayısını, denklem (6.7)'de her X başarısına denk gelen olasılık değerlerini ifade ettiğinden dolayı her iki denklemi bir araya getirerek binom olasılık fonksiyonunu elde ederiz.

Binom Formülü:

$$P(x) = \frac{n!}{x!(n-x)!} p^x (1-p)^{n-x} \quad (6.8)$$

Burada $P(x)$, n sayıdaki deneme de x sayıda başarı elde edilme olasılığını göstermektedir. $P(x)$ fonksiyonu size n denemede x sayıda başarı elde olasılıklarını veri p ve n parametreleri için göstermektedir. Bu ifade size biraz karmaşık gözüksede, temelde, sizin lise matematiğinden de alışık olduğunuz aşağıdaki lineer fonksiyondan daha farklı değildir.

$$y = f(x) = a + bx$$

Burada a ve b iki sabit değeri ifade ettiğinden dolayı, onları sabit birer değer olarak varsayabiliriz. Örneğin, $a = 2$ ve $b = 3$ değerlerini alırsa fonksiyonumuz aşağıdaki hali alır:

$$y = f(x) = 2 + 3x$$

Fonksiyonumuz, belirli a ve b değerleri için farklı x değerleri karşısında y 'nin alabileceği değerleri temsil etmektedir. Binom fonksiyonu da aynı doğrultuda, farklı x değerleri ve belirli n ve p değerlerine denk gelen olasılık değerlerini temsil etmektedir. İki fonksiyon arasındaki tek farklılık

⁵ Temel olasılık yasasında faydalananak bulabiliriz

binom fonksiyonunun, olasılık fonksiyonlarının genel şartlarına uyması gerektirir. (olasılık toplamı 1 olacak, olasılık değerleri sadece 0 veya pozitif değerler alabileceği gibi)

Örnek 6.3 (doğrudan formülü kullanarak): Bu örneğimizde binom formülünü kullanarak olasılık değerlerini hesaplayacağız. Tablolardan ve bilgisayardan faydalananak kompleks formülleri hesaplayacağız. Aşağıdaki soruları cevaplayabilmek için binom formülünü kullanacağız.

Soru 1: Herhangi bir müşterinin alışveriş yapma olasılığı 0.4 ise, mağazaya giren yirmi müşterinin on tane ürün satın alma olasılığı nedir? $P(x=10)$ gibi.

Binom formülümüzü kullanmak için $n = 20$ ve $p = 0.4$ değerlerini kullanmamız gerekmektedir. Dolayısıyla binom formülümüz aşağıdaki gibi olur:

$$P(x=10) = \frac{20!}{10!(20-10)!} \times 0.4^{10} \times (0.6)^{10} = 0.11714$$

20 müşterinin 10 adet ürün alma olasılığını 0.11 olarak hesaplarız.

Soru 2: Mağazaya giren 20 müşterinin 3 veya daha az ürün alma olasılığı nedir? $P(x \leq 3) = ?$ gibi

$$P(x \leq 3) = P(x=0) + P(x=1) + P(x=2) + P(x=3)$$

Bu olasılığı hesaplamak için

$$P(x=0) = \frac{20!}{0!(20-0)!} \times 0.4^0 \times (0.6)^{20} = 0.000036562$$

$$P(x=1) = \frac{20!}{1!(20-1)!} \times 0.4^1 \times (0.6)^{19} = 0.00048749$$

$$P(x=2) = \frac{20!}{2!(20-2)!} \times 0.4^2 \times (0.6)^{18} = 0.0030874$$

$$P(x=3) = \frac{20!}{3!(20-3)!} \times 0.4^3 \times (0.6)^{17} = 0.01235$$

Daha sonra

$$P(x \leq 3) = 0.000036562 + 0.00048749 + 0.0030874 + 0.01235 = 0.015961$$

Dolayısıyla, mağazaya giren 20 müşterinin 3 veya daha az ürün alma olasılığı 0.01 (yüzde 1) olasılık. Bu tarz olasılık değerlerine, $P(x \leq 3)$, **kümülatif olasılık değerleri** denilmektedir. 3 değerine kadar olan olasılık değerlerini göstermektedir.

Soru 3: Mağazaya giren 20 müşterinin 3'ten daha az ürün alma olasılığı nedir?

$P(x < 3) = ?$ Gibi

Açık olarak istediğimiz olasılık değeri

$$P(x < 3) = P(x = 0) + P(x = 1) + P(x = 2)$$

ve

$$P(x < 3) = 0.000036562 + 0.00048749 + 0.0030874 = 0.0036115$$

$$P(x \leq 2) = P(x < 3).$$

Soru 4: Mağazaya giren 20 müşterinin 17'den daha fazla ürün alma olasılığı nedir? $P(x > 17) = ?$ gibi; veya mağazaya giren 20 müşterinin 17 ve daha fazla ürün alma olasılığı nedir? $P(x \geq 17) = ?$

Bu olasılık değerleri

$$P(x > 17) = P(x = 18) + P(x = 19) + P(x = 20)$$

$$P(x \geq 17) = P(x = 17) + P(x = 18) + P(x = 19) + P(x = 20)$$

ve aşağıdaki gibi hesaplanabilir

$$P(x = 17) = \frac{20!}{17!(20-17)!} \times 0.4^{17} \times (0.6)^3 = 0.00004$$

$$P(x = 18) = \frac{20!}{18!(20-18)!} \times 0.4^{18} \times (0.6)^2 = 0.000004$$

$$P(x = 19) = \frac{20!}{19!(20-19)!} \times 0.4^{19} \times (0.6)^1 = 0.0000003$$

$$P(x = 20) = \frac{20!}{20!(20-20)!} \times 0.4^{20} \times (0.6)^0 = 0.00000001$$

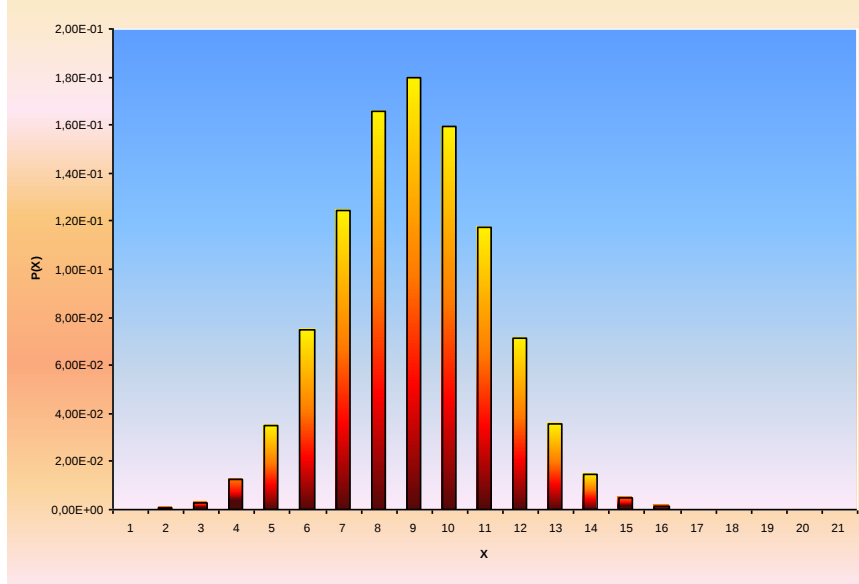
ve

$$P(x > 17) = 0.00004 + 0.000004 + 0.0000003 = 0.00004$$

$$P(x \geq 17) = 0.00004 + 0.000004 + 0.0000003 + 0.00000001 = 0.000005$$

Örneğimizin olasılık yoğunluk fonksiyonunu nasıl elde edebiliriz? Olasılık yoğunluk fonksiyonu, 0 satış miktarından 20 satış miktarına kadar olan olasılık değerlerinden oluşur. Bir başka deyişle minimum satış miktarından, maksimum satış miktarına kadar satış yapma olasılıklarını gösterir. Dolayısıyla bu yoğunluk fonksiyonunu elde edebilmek için 21 adet binom

olasılığını $P(x = 0)$; $P(x = 1)$; $P(x = 2)$; ...; $P(x = 20)$, hesaplayıp koordinat düzleminde grafiğini çizebiliriz.



Şekil 6.6

Örnek 6.3 (tablo kullanımı): Bir önceki örneğimizde hesapladığımız olasılık değerlerini hesaplamak için daha kolay yöntemler de vardır. Bu yöntemlerin en önemlilerinden biri de istatistik tablolarıdır. Belli bir değer için aradığımız olasılık değerini hiçbir hesaplama gerektirmeden doğrudan bu tablolardan elde edebiliriz. Tabloların kullanımını daha detaylı inceleyebilmek için ikinci soruyu bir daha irdелейelim. Bu soruda bizden $P(x \leq 3)$ değerini bulmamız beklenmektedir. Ekte yer almakta olan binom tablosu, 0 dan 100'e kadar olan x değerlerinin farklı n ve p parametrelerine göre kümülatif olasılık değerlerini ($P(x \leq x_0)$) gibi vermektedir. Tablodaki bloklar farklı n değerlerine göre düzenlenmiştir. Örneğimizde $n = 20$ olduğundan, öncelikle tabloda bu örneklem boyutuna denk gelen bloğu bulmamız gerekmektedir. Bu blokta her farklı sütun, p farklı olasılık değerlerini temsil etmektedir. Her farklı satır ise 0 dan 20'ye kadar olan ($n = 20$ olduğundan) kümülatif olasılık değerlerini temsil etmektedir. Aşağıdaki şekilde, $n = 20$ ve $p = 0.4$ iken $P(x \leq 3)$ değerinin nasıl bulunacağı gösterilmektedir.

$$P(x) = \frac{n!}{x!(n-x)!} p^x (1-p)^{n-x}$$

$$P(x \leq 3) = 0,0160$$

p	n	x	0,0100	0,0500	0,1000	0,1500	0,2000	0,2500	0,3000	0,3500	0,4000	0,4500	0,5000
20	0		0,8179	0,3585	0,1216	0,0388	0,0115	0,0032	0,0008	0,0002	0,0000	0,0000	0,0000
	1		0,9831	0,7358	0,3918	0,1756	0,0692	0,0243	0,0076	0,0021	0,0005	0,0001	0,0000
	2		0,9990	0,9245	0,6769	0,4049	0,2061	0,0913	0,0355	0,0121	0,0036	0,0009	0,0002
	3		1,0000	0,9841	0,8671	0,6477	0,4115	0,2252	0,1071	0,0444	0,0160	0,0049	0,0013
	4		1,0000	0,9974	0,9568	0,8299	0,6297	0,4148	0,2375	0,1182	0,0510	0,0189	0,0059
	5		1,0000	0,9997	0,9888	0,9327	0,8042	0,6172	0,4164	0,2454	0,1256	0,0553	0,0207
	6		1,0000	1,0000	0,9976	0,9781	0,9133	0,7858	0,6080	0,4166	0,2500	0,1299	0,0577
	7		1,0000	1,0000	0,9996	0,9941	0,9679	0,8982	0,7723	0,6010	0,4159	0,2520	0,1316
	8		1,0000	1,0000	0,9999	0,9987	0,9900	0,9591	0,8867	0,7624	0,5956	0,4143	0,2517
	9		1,0000	1,0000	1,0000	0,9998	0,9974	0,9861	0,9520	0,8782	0,7553	0,5914	0,4119
	10		1,0000	1,0000	1,0000	1,0000	0,9994	0,9961	0,9829	0,9468	0,8725	0,7507	0,5881
	11		1,0000	1,0000	1,0000	1,0000	0,9999	0,9991	0,9949	0,9804	0,9435	0,8692	0,7483
	12		1,0000	1,0000	1,0000	1,0000	1,0000	0,9998	0,9987	0,9940	0,9790	0,9420	0,8684
	13		1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9985	0,9935	0,9786	0,9423	0,8684
	14		1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9936	0,9793	0,9423	0,8684
	15		1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9936	0,9793	0,9423	0,8684

Şekil 6.7

Tablodan elde ettiğimiz, $P(x \leq 3) = 0.0160$ değeri daha önce ikinci sorunun cevabında elde ettiğimiz olasılık değeri ile aynıdır.

Şimdi ise aynı yöntemi izleyerek **üçüncü soruya** cevap vermeye çalışalım ve $P(x < 3)$ olasılık değerini bulalım. Bu durumda $P(x \leq 2) = P(x < 3)$ özelliğinden faydalanmamız gerekmektedir. Dolayısıyla $P(x < 3)$ olasılık değerini elde edebilmek için, tablodan 2'ye ait kümülatif olasılık değerini bulmamız yeterli olacaktır.

$$P(x) = \frac{n!}{x!(n-x)!} p^x (1-p)^{n-x}$$

$$P(x < 3) = P(x \leq 2) = 0,0036$$

p	n	x	0,0100	0,0500	0,1000	0,1500	0,2000	0,2500	0,3000	0,3500	0,4000	0,4500	0,5000
20	0		0,8179	0,3585	0,1216	0,0388	0,0115	0,0032	0,0008	0,0002	0,0000	0,0000	0,0000
	1		0,9831	0,7358	0,3918	0,1756	0,0692	0,0243	0,0076	0,0021	0,0005	0,0001	0,0000
	2		0,9990	0,9245	0,6769	0,4049	0,2061	0,0913	0,0355	0,0121	0,0036	0,0009	0,0002
	3		1,0000	0,9841	0,8671	0,6477	0,4115	0,2252	0,1071	0,0444	0,0160	0,0049	0,0013
	4		1,0000	0,9974	0,9568	0,8299	0,6297	0,4148	0,2375	0,1182	0,0510	0,0189	0,0059
	5		1,0000	0,9997	0,9888	0,9327	0,8042	0,6172	0,4164	0,2454	0,1256	0,0553	0,0207
	6		1,0000	1,0000	0,9976	0,9781	0,9133	0,7858	0,6080	0,4166	0,2500	0,1299	0,0577
	7		1,0000	1,0000	0,9996	0,9941	0,9679	0,8982	0,7723	0,6010	0,4159	0,2520	0,1316
	8		1,0000	1,0000	0,9999	0,9987	0,9900	0,9591	0,8867	0,7624	0,5956	0,4143	0,2517
	9		1,0000	1,0000	1,0000	0,9998	0,9974	0,9861	0,9520	0,8782	0,7553	0,5914	0,4119
	10		1,0000	1,0000	1,0000	1,0000	0,9994	0,9961	0,9829	0,9468	0,8725	0,7507	0,5881
	11		1,0000	1,0000	1,0000	1,0000	0,9999	0,9991	0,9949	0,9804	0,9435	0,8692	0,7483
	12		1,0000	1,0000	1,0000	1,0000	1,0000	0,9998	0,9987	0,9940	0,9790	0,9420	0,8684
	13		1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9985	0,9935	0,9786	0,9423	0,8684
	14		1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9936	0,9793	0,9423	0,8684
	15		1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9936	0,9793	0,9423	0,8684

Şekil 6.8

ve beklendiği gibi daha önce hesaplamış olduğumuz olasılık değerinin aynısını kolaylıkla elde edebiliriz.

Şimdi ise aynı şekilde **birinci soruya** cevap vermeye çalışalım. Bu durum için ise olasığın iki kümülatif olasılık değerinin farkı basit olasılık eşitliğini vermektedir. $P(x = 10) = P(x \leq 10) - P(x \leq 9)$ özelliğinden faydalanacağız. Dolayısıyla tablomuzu da kullanarak $P(x = 10) = P(x \leq 10) - P(x \leq 9) = 0.8255 - 0.7553 = 0.1172$ değerini Şekil 6.9'daki gibi hesaplayabiliriz.

$$P(x) = \frac{n!}{x!(n-x)!} p^x (1-p)^{n-x}$$

$$P(X=10) = P(X \leq 10) - P(X \leq 9) = 0,8725 - 0,7553 = 0,1172$$

p	n	x	0,0100	0,0500	0,1000	0,1500	0,2000	0,2500	0,3000	0,3500	0,4000	0,4500	0,5000
20	0	0	0,8179	0,3585	0,1216	0,0388	0,0115	0,0032	0,0008	0,0002	0,0000	0,0000	0,0000
20	1	1	0,9831	0,7358	0,3918	0,1756	0,0692	0,0243	0,0076	0,0021	0,0005	0,0001	0,0000
20	2	2	0,9990	0,9245	0,6769	0,4049	0,2061	0,0913	0,0355	0,0121	0,0036	0,0009	0,0002
20	3	3	1,0000	0,9841	0,8671	0,6477	0,4115	0,2252	0,1071	0,0444	0,0160	0,0049	0,0013
20	4	4	1,0000	0,9974	0,9568	0,8299	0,6297	0,4148	0,2375	0,1182	0,0510	0,0189	0,0059
20	5	5	1,0000	0,9997	0,9888	0,9327	0,8042	0,6172	0,4164	0,2454	0,1256	0,0553	0,0207
20	6	6	1,0000	1,0000	0,9976	0,9781	0,9133	0,7858	0,6080	0,4166	0,2500	0,1299	0,0577
20	7	7	1,0000	1,0000	0,9996	0,9941	0,9679	0,8982	0,7723	0,6010	0,4159	0,2520	0,1316
20	8	8	1,0000	1,0000	0,9999	0,9987	0,9900	0,9591	0,8867	0,7624	0,5956	0,4143	0,2517
20	9	9	1,0000	1,0000	1,0000	0,9998	0,9974	0,9861	0,9520	0,8782	0,7553	0,5914	0,4119
20	10	10	1,0000	1,0000	1,0000	1,0000	0,9994	0,9961	0,9829	0,9468	0,8725	0,7507	0,5881
20	11	11	1,0000	1,0000	1,0000	1,0000	0,9999	0,9991	0,9949	0,9804	0,9435	0,8692	0,7483
20	12	12	1,0000	1,0000	1,0000	1,0000	1,0000	0,9998	0,9987	0,9940	0,9790	0,9420	0,8684
20	13	13	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9985	0,9935	0,9786	0,9423	0,8684
20	14	14	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9935	0,9786	0,9423	0,8684
20	15	15	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9935	0,9786	0,9423	0,8684

Şekil 6.9

Son olarak **dördüncü soruyu**, incelemeye çalışalım. Öncelikle $P(x > 17)$ olasılık değerini bulmamız gerekmektedir. $P(x > 17) = 1 - P(x \leq 17)$ özelliğinden faydalanarak. Binom tablosunun yardımıyla $P(x > 17) = 1 - P(x \leq 17) = 1 - 1 = 0$ değerini hesaplayabiliriz.

$$P(x) = \frac{n!}{x!(n-x)!} p^x (1-p)^{n-x}$$

$$P(X > 17) = 1 - P(X \leq 17) = 1 - 1 \approx 0$$

p	n	x	0,0100	0,0500	0,1000	0,1500	0,2000	0,2500	0,3000	0,3500	0,4000	0,4500	0,5000
20	0	0	0,8179	0,3585	0,1216	0,0388	0,0115	0,0032	0,0008	0,0002	0,0000	0,0000	0,0000
20	1	1	0,9831	0,7358	0,3918	0,1756	0,0692	0,0243	0,0076	0,0021	0,0005	0,0001	0,0000
20	2	2	0,9990	0,9245	0,6769	0,4049	0,2061	0,0913	0,0355	0,0121	0,0036	0,0009	0,0002
20	3	3	1,0000	1,0000	0,9999	0,9987	0,9900	0,9591	0,8867	0,7624	0,5956	0,4143	0,2517
20	4	4	1,0000	1,0000	1,0000	0,9998	0,9974	0,9861	0,9520	0,8782	0,7553	0,5914	0,4119
20	5	5	1,0000	1,0000	1,0000	1,0000	0,9994	0,9961	0,9829	0,9468	0,8725	0,7507	0,5881
20	6	6	1,0000	1,0000	1,0000	1,0000	0,9999	0,9991	0,9949	0,9804	0,9435	0,8692	0,7483
20	7	7	1,0000	1,0000	1,0000	1,0000	1,0000	0,9998	0,9987	0,9940	0,9790	0,9420	0,8684
20	8	8	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9985	0,9935	0,9786	0,9423	0,8684
20	9	9	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9935	0,9786	0,9423	0,8684
20	10	10	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9935	0,9786	0,9423	0,8684
20	11	11	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9935	0,9786	0,9423	0,8684
20	12	12	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9935	0,9786	0,9423	0,8684
20	13	13	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9935	0,9786	0,9423	0,8684
20	14	14	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9935	0,9786	0,9423	0,8684
20	15	15	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9935	0,9786	0,9423	0,8684
20	16	16	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9935	0,9786	0,9423	0,8684
20	17	17	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9935	0,9786	0,9423	0,8684
20	18	18	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9935	0,9786	0,9423	0,8684
20	19	19	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9935	0,9786	0,9423	0,8684
20	20	20	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9935	0,9786	0,9423	0,8684

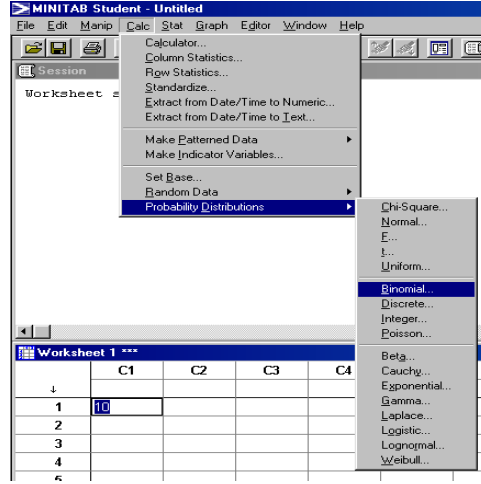
Şekil 6.10

Daha önceki hesaplamalarımızda olasılık değerini 0'dan ziyade sıfıra çok yakın bir değer (0.00004) olarak hesaplamıştık. Binom tablosu daha önce yaptığımız hesaplamalar kadar detaylı bilgi vermediğinden ve sonuç çok küçük bir değer olduğundan $P(x \geq 17) = 1 - P(x \leq 16)$ değeri sıfır olarak hesaplanmıştır.

Minitab Uygulamaları

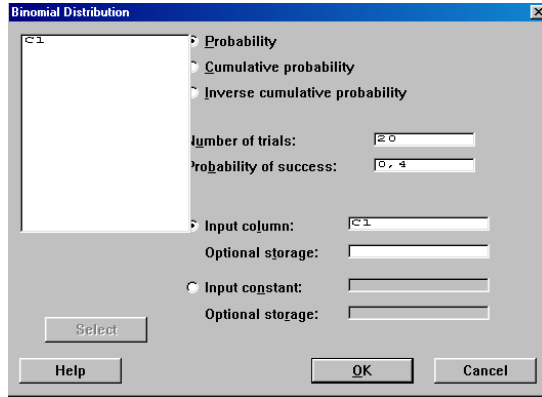
Örnek 6.3 (MINITAB): Hesaplamış olduğumuz olasılık değerlerinin bulunmasındaki en kolay yöntem istatistiksel yazılımlar kullanmaktır. Şimdi ise daha önce yaptığımız hesaplamaları Minitab yazılımı yardımıyla nasıl gerçekleştireceğimizi inceleyeceğiz. **Birinci sorudaki** ($P(x=10) = ?$) olasılık değerini MINITAB ortamında hesaplayabilmek için, öncelikle çalışma sayfasındaki ilk hücreye "10" (C1'in altına) değerini giriniz. Daha sonra ise Şekil 6.11'de de gösterildiği gibi "Calc"

(Hesapla) menüsünden “*Probability Distributions*” (Olasılık Dağılımları) seçeneğine tıklayınız ve “*binomial*” (binom)’ı seçiniz.



Şekil 6.11

Karşınıza çıkacak olan diyalog ekranında “*probability*” (olasılık) kutucuğunu işaretleyiniz ve “*Number of Trials*” (deneme miktarı) kutucuğuna “20” yazınız ($n = 20$). “*Probability of success*” (başarı olasılığı) kutucuğuna ise “0.4” yazınız ($p = 0.4$). Daha sonra “*input column*” (veri sütunu) kutucuğuna ise “C1” yazınız. Karşınızdaki diyalog ekranını Şekil 6.12’deki gibi doldurulmalıdır.



Şekil 6.12

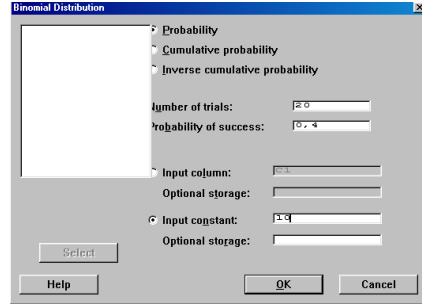
“OK” tuşuna bastığınızda Şekil 6.13’teki sonucu elde edersiniz.

Probability Density Function	
Binomial with $n = 20$ and $p = 0,400000$	
x	P(X = x)
10,00	0,1171

Şekil 6.13

Daha önce elde ettiğimiz sonucun aynısını kolaylıkla elde edebiliriz.

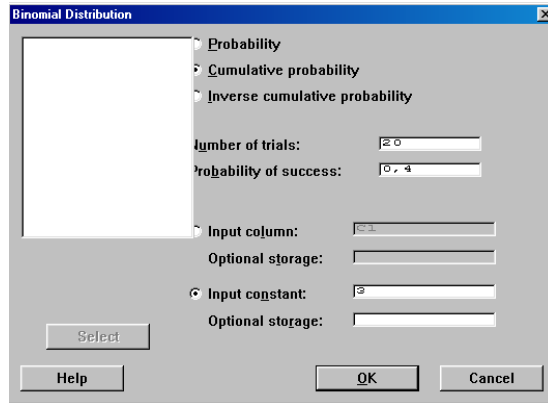
Alternatif bir metod olarak aynı sonucu Şekil 6.14'daki diyalog ekranında da gösterildiği gibi “*input column*” (veri sütunu) kutucuğuna ise “C1” yazmak yerine, “*input constant*” (veri sabiti) seçeneğine tıklayarak yazdığınızda da aynı sonucu elde edersiniz.



Şekil 6.14

Üçüncü sorudaki olasılık değerini $P(x < 3)$ elde edebilmemiz için aynı işlemleri $P(x = 0)$, $P(x = 1)$, $P(x = 2)$ değerleri için tek tek hesaplayıp bu değerleri toplamamız gerekmektedir. Veya alternatif olarak $P(x \leq 2) = P(x < 3)$ eşitliğini kullanarak $P(x \leq 2)$ kümülatif olasılık değerini Minitab ortamında hesaplayabiliriz.

İkinci sorudaki $P(x \leq 3)$ olasılık değerini hesaplayabilmek için, “*Calc*” (Hesapla) menüsünden, “*Probability Distributions*” (Olasılık Dağılımlarına) ve “*binomial*” (Binom)’a tıklayınız. Karşınıza çıkan diyalog ekranından Şekil 6.15’te olduğu gibi “*Cumulative probability*” (kümülatif olasılık) seçeneğini seçiniz ve “*Number of Trials*” (deneme miktarı) bölümüne “20” ($n = 20$) değerini giriniz. “*Probability of success*” (başarı olasılığı) kutucuğuna 0.4 ($p = 0.4$) değerini giriniz. X değerini girmek için “*input constant*” (veri sabiti) kutucuğuna geliniz ve 3 yazınız.



Şekil 6.15

“OK” tuşuna bastığınızda Şekil 6.16’daki çıktı sonucunu elde edebilirsiniz.

Cumulative Distribution Function	
Binomial with n = 20 and p = 0,400000	
x	P(X ≤ x)
3,00	0,0160

Şekil 6.16

Dolayısıyla daha önce hesapladığımız olasılık değerinin aynısını elde ederiz.

Şimdi ise **dördüncü soruda** istenilen olasılık değerlerini hesaplamaya çalışalım. Örneğin, $P(x > 17)$ olasılık değerini MINITAB ortamında nasıl hesaplayabiliriz? Bu olasılık değerini hesaplayabilmek için $P(x > 17) = 1 - P(x \leq 17)$. $P(x \leq 17)$ eşitliğini kullanmamız gerekmektedir. Minitab bize daha önce de hesapladığımız yöntemle $P(x \leq 17) = 0.99999^6$ değerini bulabiliriz. Dolayısıyla $P(x > 17) = 1 - 0.99999 = 0.00001$ değeri benzer şekilde $P(x \geq 17)$ değerini hesaplayabilmek için $P(x \geq 17) = 1 - P(x \leq 16)$ eşitliğinden faydalanabiliriz.

Sonuç olarak, olasılık yoğunluğunu elde etmek için, $n=20$ (sabit müşteri sayısı) iken x (farklı satış miktarlarını) değerlerini 0 dan 20'ye kadar seri olarak girmeniz gerekmektedir. "C1" sütununa 0'dan 20'ye kadar verileri girip diyalog ekranını da Şekil 6.8'de olduğu gibi doldurduğunuz takdirde Şekil 6.17'deki çıktıyı elde edebilirsiniz.

Probability Density Function	
Binomial with n = 20 and p = 0,400000	
x	P(X = x)
1,00	0,0005
2,00	0,0031
3,00	0,0123
4,00	0,0350
5,00	0,0746
6,00	0,1244
7,00	0,1659
8,00	0,1797
9,00	0,1597
10,00	0,1171
11,00	0,0710
12,00	0,0355
13,00	0,0146
14,00	0,0049
15,00	0,0013
16,00	0,0003
17,00	0,0000
18,00	0,0000
19,00	0,0000
20,00	0,0000

Şekil 6.17

Bu çıktı bize, mağaza giren müşteri sayısı 20 iken sıfırdan 20 satışa kadar satış yapma olasılığını göstermektedir. Çıktıda görmüş olduğunuz 17 veya daha fazla satış yapma olasılıkları dört basamaklı ve sıfırdan oluşan bir sayı ile ifade edilmektedir. Bu ifade olasılık değerlerinin

⁶ Bu rakam birden çok farklı bir rakam olmadığından, Minitab bu rakamın tam değerini doğrudan belirtmemektedir. Olasılık değerinin tam değerini Minitab ortamında göstermek ve saklamak için, "Probability Distributions-Cal" (Olasılık Dağılımları- hesapla) menüsünden "binomial" (binom) seçeneğine gelerek, karşınıza çıkan diyalog ekranında "Cumulative probability" (Kümülatif Olasılık) değerini seçmeniz ve "input column" (veri sütunu) bölümüne 1 den 17 kadar değerleri girdiğiniz sürunu seçmeniz gerekmektedir. Son olarak "optional storage" (alternatif saklama) kutucuğuna saklayacağınız yer olan "C2"yi girip OK tuşuna bastığınızda bu değerleri elde edebilirsiniz.

olmadığına delalet etmemektedir. Tam tersine bu ifade Minitab'ın dört basamaktan daha küçük değerleri raporlamadığını belirtmektedir. Bu değerlerin tam değerlerini elde etmek isterseniz biraz önce diyalog ekranında yaptığımız işlemlere ek olarak “*optional storage*” seçeneğine C2 yazınız ve OK tuşuna basınız. Bulmuş olduğumuz olasılık değerlerine ait olan grafiği de Minitab yardımıyla elde edebiliriz.⁷

Binom Dağılımının Ortalaması ve Standart Sapması

Binom dağılımının beklenen değeri (ortalaması) μ ile gösterilmekte ve aşağıdaki formül yardımı ile hesaplanmaktadır:

$$\mu = n \times p$$

Örneğin, $n = 10$ ve $p = 0.5$, ise $\mu = 5$ olarak hesaplanır. Hatırlanacağı üzere beklenen değer, belli bir p olasılık değerinde olayın beklenen ortalama değerini ifade etmektedir. Binom dağılımına ait olan dağılım aşağıdaki varyans formülü yardımıyla hesaplanabilir:

$$\sigma^2 = p \times (1 - p) \times n$$

Örneğimizdeki verilere ait varyans değerini hesaplayacak olursak:

$$\sigma^2 = 0.5 \times 0.5 \times 10 = 2.5$$

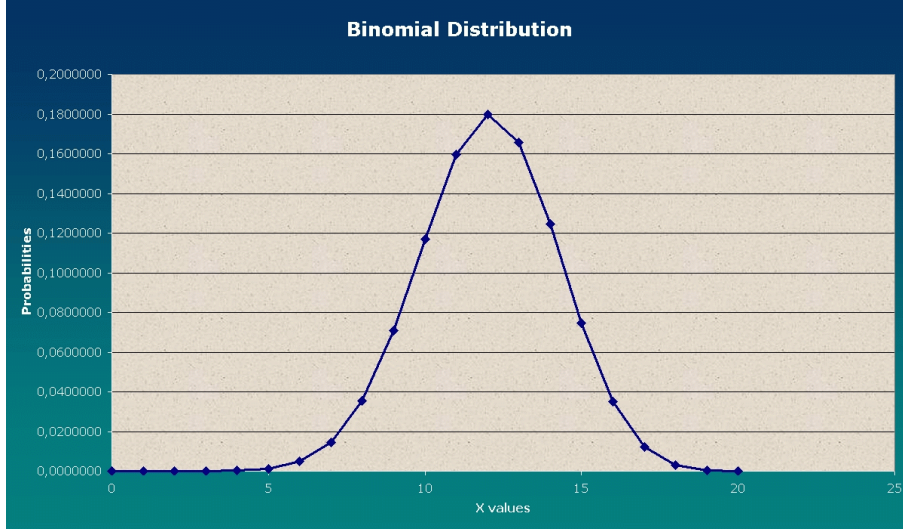
Standart sapma değeri de aşağıdaki gibi olur:

$$\sigma = \sqrt{\sigma^2} = \sqrt{2.5} = 1.58$$

Sürekli Olasılık Dağılımları

Şu ana kadar sadece süreksiz olasılık dağılımlarını inceledik. Süreksiz olasılık dağılımları ile sürekli olasılık dağılımları arasındaki tek fark; süreksiz olasılık fonksiyonları, süreksiz bir sayı kümesi üzerinden tanımlıyken, sürekli olasılık dağılımları sürekli değişkenlerden oluşan bir küme üzerinden tanımlanmasıdır. Dolayısıyla, süreksiz olasılık dağılımları gördüğümüz gibi süreksiz (kesikli) çizgilerden oluşurken, sürekli olasılık dağılımları sürekli (kesiksiz) bir çizgiden oluşur. Örneğin, binom dağılımı sürekli olarak tanımlanacak olursa (ki bu durumda 1.5 müşteri anlamsız olacaktır) dağılımın grafiği:

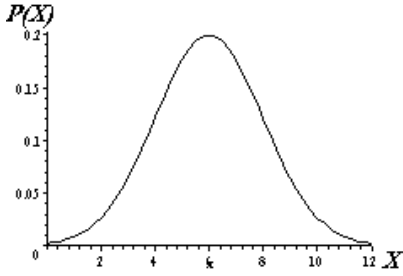
⁷ Olasılık yoğunluk grafiğini Minitab ortamında elde edebilmek için , öncelikle daha önce bahsettiğimiz yöntemle verileri “C2”de saklayınız. Daha sonra “*graph*” (grafik) menüsünden “*chart*” (çizelge) seçeneğine tıklayıp “Y” yazan yere “C2” (olasılık değerlerini) ve “X” yazan yere “C1” (0'dan 20'ye satış miktarları) giriniz. Son olarak “*Display*” (göster) seçeneğinden “*project*” (proje)’yi seçip “OK” tuşuna bastığınızda yoğunluk grafiğini elde edebilirsiniz.



Şekil 6.18

Normal Dağılım

Şuana kadar, örneklerimizde hep binom dağılımını kullandık. Ancak, gerçekte araştırmalarda en kullanılan dağılım, normal dağılımdır. Normal dağılım sürekli bir dağılımdır, başka deyişle sürekli olan bir rassal değişken üzerinden tanımlıdır. Normal dağılım, ağırlık, boy, yaşam süresi, IQ gibi pek çok insan özelliğinin dağılımına uyum sağlaması ile bilinir. Normal dağılım eğrisi Şekil 6.19'daki gibidir.

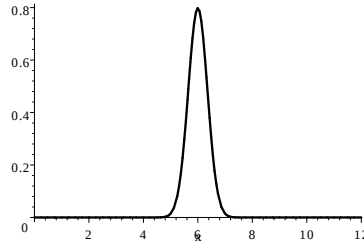


Şekil 6.19 Normal Yoğunluk $\mu = 6; \sigma = 2$

Normal dağılım eğrisine bazen çan eğrisi de denilmektedir. Değerler genellikle eğrinin orta kısımlarında yoğunlaşmıştır. X değişkeninin değerleri, eksi sonsuzdan artı sonsuza kadar değişir, fakat değerlerin büyük bir kısmı ortalamanın etrafında yoğunlaşmaktadır (μ). Örneğin Şekil 6.19'daki normal dağılımın ortalaması 6'ya eşittir. Normal bir dağılımda ortalamanın olasılığı, en yüksek değere sahiptir; başka ifade ile eğrinin tepe noktası X 'in ortalama değerine karşılık gelmelidir. Şekil 6.19'da, ortalama 6'dır ve 6'nın olasılığı, eksi sonsuz ile artı sonsuz arasında gözlemlenebilecek en yüksek değer olan 0.2'ye eşittir. Ortalamadan uzaklaştıkça, üstüne geldiğimiz noktanın (değerin) gözlemlenme olasılığı da düşecektir. Normal dağılımın bir başka özelliği de ortalamaya göre simetrik olmasıdır. Başka bir deyişle, ortalamaya sağdan ve soldan eşit

uzaklıkta iki sayının olasılıkları, eşit olmak zorundadır. Örneğin şekil 6.19'da, 5 ve 7 değerleri, ortalama olan 6'ya eşit uzaklıkta oldukları için (ikisi de 1 birim uzakta) olasılıkları aynıdır.

Binom dağılımını tanımlamak için iki parametre yeterli idi (n ve p) veya lineer bir fonksiyonu ($f(x) = a + bx$ gibi) tanımlamak için iki parametre (a ve b) yeterlidir. Dolayısıyla, normal bir dağılımı tümüyle tanımlamak için de sadece iki parametre yeterlidir: ortalaması (μ) ve standart sapması (σ). Bu parametreleri değiştirmek normal dağılımı da değiştirmek anlamına gelmektedir. Örneğin, Şekil 6.19'daki normal dağılımda ortalamayı sabit tutup standart sapmayı 2'den 0.5'e indirirsek, Şekil 6.20'yi elde ederiz.



Şekil 6.20 Normal Yoğunluk $\mu = 6$; $\sigma = 0.5$

Şekilde görüldüğü gibi, standart sapmanın azaltılması, olasılıkların ortalama etrafında daha da yoğunlaşmasına yol açmaktadır. Başka bir deyişle, ortalamanın yakınındaki değerlerin gözlemlene şansı fazlaşmıştır. Şekil 6.19'a göre, ortalama değer olan 6'yı gözlemlene olasılığı (yaklaşık 0.8) daha önceki duruma göre (yaklaşık 0.2 idi) 4 kat daha fazladır. Benzer şekilde, X 'in diğer değerleri için de olasılık daha önceki duruma göre 4 kat artmıştır. Bu durum bir tesadüf değildir. Nedeni, ortalamayı (6 sayısında) sabit tutarken standart sapmayı 4 kat (2'den 0.5'e) azaltmamızdır.

Normal bir dağılımın olasılık yoğunluk fonksiyonu aşağıdaki gibidir:

Normal Olasılık Fonksiyonu

$$P(x) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2} \frac{(x-\mu)^2}{\sigma^2}} \quad (6.9)$$

Burada (μ), X 'in ortalaması, $\pi=3.14159$, σ , X 'in standart sapmasıdır ve “e” 2.71828'e eşit sabit bir sayıdır. “e” ve π sabit sayı oldukları için, olasılık yoğunluk fonksiyonu tümüyle μ ve σ tarafından belirlenir. Bu formül size biraz karmaşık görünebilir. Ama sonuçta, verili μ ve σ değerleri için değişik X 'lerin değerlerini veren başka bir olasılık fonksiyonudur.

Örnek 6.4 (formül kullanımı) İstanbul Bilgi Üniversitesi'nde sınav notlarının normal dağıldığını düşünelim. Sınav notlarının, ortalaması (10 üzerinden) 6 ve standart sapmasının 2 olduğunu düşünelim. Rastgele seçtiğimiz bir öğrencinin notunun 5 olma olasılığını öğrenmek için, normal dağılım fonksiyonunda X yerine 5 değerini verip aşağıdaki sonucu elde ederiz:

$$P(x=5) = \frac{1}{\sqrt{2 \times \pi \times 2^2}} e^{-\frac{1(5-6)^2}{2 \times 2^2}} = 0.17603 \quad (6.10)$$

Bu öğrencinin geçer not (5) almış olma olasılığı 0.17603'tür. Ortalaması 6 ve standart sapması 2 olan normal bir dağılım yukarıdaki şekil 6.19'da gösterilmektedir. En yüksek olasılık, ortalamaya ait olan olasılık değeridir ve şöyle bulunur:

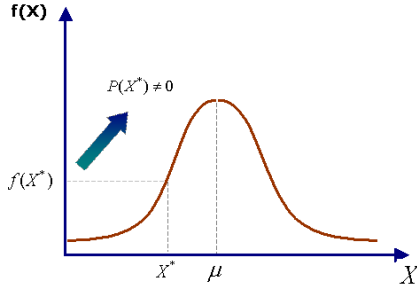
$$P(x=6) = \frac{1}{\sqrt{2 \times \pi \times 2^2}} e^{-\frac{1(6-6)^2}{2 \times 2^2}} = 0.19947 \quad (6.11)$$

Ancak burada belirtmemiz gereken çok önemli bir husus bulunmaktadır: Sürekli olasılık dağılımları tanımına göre yaptıklarımızın tamamı **yanlıştır!** Sürekli olasılık dağılımlarında fonksiyonun tek bir noktasına ait olan olasılık değerini hesaplayamamaktayız. Örnek olarak $P(x=5)$ gibi olasılık değerleri tanım gereği sıfırdır. Şimdi bu hususu bir örnek ile açıklamaya çalışalım. 200 kilometre uzunluğunda açık bir yolda kaza yapma imkanını değerlendirelim ve yolun bir kısmında veya bir yerleşkede kaza yapma olasılığımızın ne olduğunu incelemeye çalışalım. Bu deneyde ele aldığımız örneklem setimiz 0 dan 200'e kadar olan sürekli noktalardan oluşmakta olsun. Böyle bir kazanın gerçekleşme olasılığının, herhangi bir d aralığında (d kilometre olarak ölçülmektedir), $d/200$ olduğunu varsayalım. Örneğimize göre, çok kısa bir aralıkta mesela 1 santimetrelilik bir aralıkta kaza olma olasılığı, 0.00000005'tir ve bu çok küçük bir olasılık değeridir. Aralığın uzunluğu sıfıra yakınsadıkça, kaza olma olasılığı da sıfıra yakınsayacaktır. Gerçekten de sürekli dağılımlarda tek noktaların olasılık değerlerini her zaman sıfır olarak alırız. Bu durum olayların gerçekleşmeyeceği anlamına gelmemektedir. 200 kilometrelilik bir yolda elbette kaza olabilir ancak 1 cm inde kaza olma olasılığı sıfırdır. Yoksa belli aralıklarda elbette kaza olma olasılığı dikkate değerdir. Örneğin bu yolda kazalar genelde 50 ile 55. km'lerde oluyor denilebilir.

Bu yüzden (6.9)'da tanımladığımız fonksiyonu olasılık fonksiyonu olarak tanımlamamız yanlış bir tanımlamadır. Dolayısıyla olasılık fonksiyonu olarak $P(x)$ notasyonunu kullanmamız (her zaman $P(x) = 0$ olduğundan) yanlıştır. Bu notasyon yerine fonksiyonu $f(x)$ notasyonu ile gösterebiliriz.

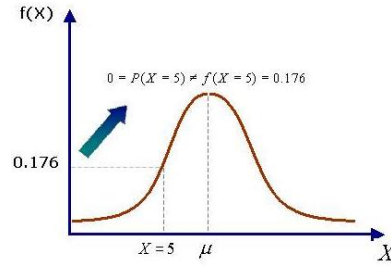
$$P(x) \neq f(x) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1(x-\mu)^2}{2\sigma^2}}$$

Şekil 6.21'de de inceleyebileceğiniz gibi $f(x)$ belli bir değer alabilirken bu değer olasılık değerine, sürekli olasılığı $P(x) = 0$ olarak tanımladığımızdan, eşit olmamaktadır.



Şekil 6.21

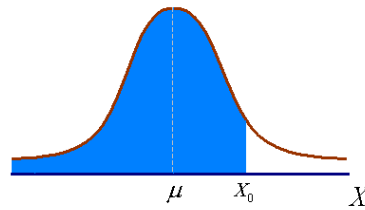
Örneğin daha önce yanlışlıkla hesapladığımız $P(x = 5) = 0.17603$ değeri ele alırsak, $P(x = 5) = 0$ tanım gereği, doğrusunu $f(x = 5) = 0.17603$ olarak gösterebiliriz.



Şekil 6.22

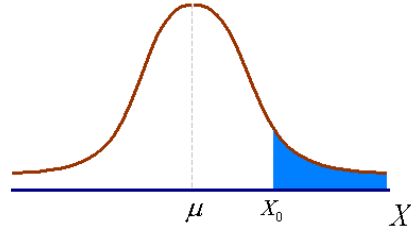
Süreklili dağılımlarda, $P(5 < x < 6)$, veya $P(x > 7)$ gibi sadece belli aralıklara ait olan olasılık değerlerini hesaplayabiliriz. Sürekli dağılımların bir diğer özelliğide belli bir noktadaki olasılık değeri her zaman sıfır olduğundan dolayı, x 'in her hangi bir noktadan büyük olma olasılığı ile x 'in her hangi bir noktaya eşit veya büyük olma olasılığının birbirine eşit olmasıdır. Örneğin $P(x = 7) = 0$ ise; $P(x > 7) = P(x \geq 7)$ ve $P(x < 7) = P(x \leq 7)$ eşitlikleri söz konusudur. Aynı özellikten dolayı $P(5 < x < 6) = P(5 \leq x \leq 6)$ eşitliği elde edilir.

Grafiksel olarak, x 'in x_0 gibi bir değerden küçük olma olasılığı, eğrinin altında kalan taralı alan ile gösterilmektedir. (bu tarz olasılıklara daha önce de öğrendiğimiz gibi kümülatif olasılık denilmektedir)



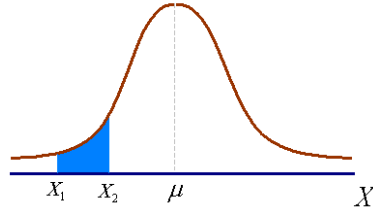
Şekil 6.23 $P(x < x_0) = P(x \leq x_0)$

Veya x 'in x_0 gibi bir değerden büyük olma olasılığı taralı alan ile ifade edilmektedir.



Şekil 6.24 $P(x > x_0) = P(x \geq x_0)$

Benzer şekilde iki değerin arasında kalan alanı da hesaplayabiliriz. Örneğin, x değerinin x_1 ve x_2 değerleri arasında kalma olasılığı aşağıdaki gibidir.



Şekil 6.25 $P(x_1 < x < x_2) = P(x_1 \leq x \leq x_2)$

Matematiksel olarak bu alanları belirli integraller yardımıyla hesaplayabiliriz. Örneğin Şekil 6.23'teki alanı aşağıdaki integral yardımı ile hesaplamaktayız.

$$P(x < x_0) = P(x \leq x_0) = \int_{-\infty}^{x_0} \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2} \frac{(x-\mu)^2}{\sigma^2}} dx \quad (6.12)$$

Şekil 6.24'teki alanı ise (6.13) numaralı integral ile hesaplayabiliriz.

$$P(x > x_0) = P(x \geq x_0) = \int_{x_0}^{\infty} \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2} \frac{(x-\mu)^2}{\sigma^2}} dx \quad (6.13)$$

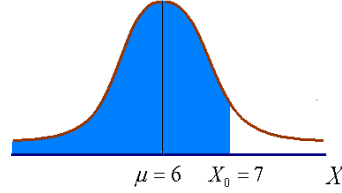
Şekil 6.25'deki alanı ise (6.14'teki) integral yardımı ile bulabiliriz.

$$P(x_1 < x < x_2) = P(x_1 \leq x \leq x_2) = \int_{x_1}^{x_2} \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2} \frac{(x-\mu)^2}{\sigma^2}} dx \quad (6.14)$$

veya alternatif olarak yukarıdaki olasılık değerlerini aşağıdaki kümülatif olasılık değerlerin farkı ile de hesaplayabiliriz.

$$P(x_1 < x < x_2) = P(x < x_2) - P(x < x_1)$$

Örnek 6.5 (formül kullanımı) Öğrencinin 7'den daha az not alma olasılığını normal dağılım formülünü kullanarak hesaplayalım $x_0 = 7$. Grafiksel olarak bu olasılığı Şekil 6.26'daki gibi inceleyebiliriz.

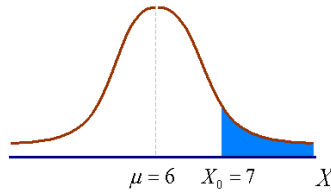


Şekil 6.26

(6.12) denkleminde faydalanarak olasılık değerini hesaplayabiliriz (normal dağılımın ortalaması, $\mu = 6$ ve standart sapması, $\sigma = 2$)⁸.

$$P(x < 7) = P(x \leq 7) = \int_{-\infty}^7 \frac{1}{\sqrt{2\pi}2^2} e^{-\frac{1(x-6)^2}{2 \cdot 2^2}} dx = 0.69146 \quad (6.12)$$

Örnek 6.6 (formül kullanımı) Öğrencinin 7'den daha fazla not alma olasılığını normal dağılım formülünü kullanarak hesaplayalım $x_0 = 7$. Grafiksel olarak bu olasılığı Şekil 6.27'deki gibi inceleyebiliriz.



Şekil 6.27

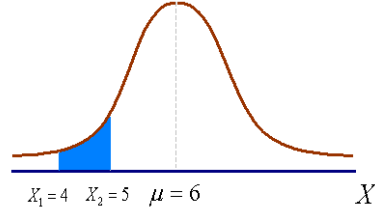
(6.13) denkleminde faydalanarak olasılık değerini hesaplayabiliriz (normal dağılımın ortalaması, $\mu = 6$ ve standart sapması, $\sigma = 2$).

$$P(x > 7) = P(x \geq 7) = \int_7^{\infty} \frac{1}{\sqrt{2\pi}2^2} e^{-\frac{1(x-6)^2}{2 \cdot 2^2}} dx = 0.30854 \quad (6.13)$$

$P(x > 7)$ olasılık değerini $P(x > 7) = 1 - P(x < 7)$ eşitliğinden, $P(x < 7)$ değerini **Örnek 6.5**'te 0.69146 olarak hesapladığımızdan elde edebiliriz.

⁸Binom dağılımı gibi süreksiz dağılımlarda, bu olasılık değerini $P(x < 7) = P(x = 6) + P(x = 5) + P(x = 4) + \dots + P(x = 10)$ olarak yazabiliriz. Ancak $P(x < 7)$ değeri $P(x = 6)$ ve $P(x = 5)$, veya $P(x = 5)$ ve $P(x = 4)$, gibi değerlerin aralarındaki olasılık değerlerini de kapsadığından dolayı sürekli fonksiyonlarda bu yazım biçimi yanlıştır.

Örnek 6.7 (formül kullanımı) İstanbul Bilgi Üniversitesi’nde yapılan İstatistik sınavının sonuçlarının 4 ile 5 arasında olma olasılığını hesaplayalım, ($x_1 = 4$ ve $x_2 = 5$ gibi). Grafikselleştirilerek bu olasılık değerini inceleyecek olursak:



Şekil 6.28

(6.14) denkleminde faydalanarak olasılık değerini hesaplayabiliriz (normal dağılımın ortalaması, $\mu = 6$ ve standart sapması, $\sigma = 2$).

$$P(4 < x < 5) = P(4 \leq x \leq 5) = \int_4^5 \frac{1}{\sqrt{2\pi}2} e^{-\frac{1(x-6)^2}{2 \cdot 2^2}} dx = 0.14988 \quad (6.14)$$

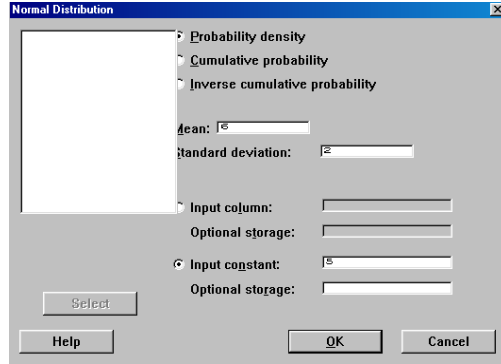
veya

$$\begin{aligned} P(4 < x < 5) &= P(x < 5) - P(x < 4) = \int_{-\infty}^5 \frac{1}{\sqrt{2\pi}2} e^{-\frac{1(x-6)^2}{2 \cdot 2^2}} dx - \int_{-\infty}^4 \frac{1}{\sqrt{2\pi}2} e^{-\frac{1(x-6)^2}{2 \cdot 2^2}} dx \\ &= 0.30954 - 0.15866 = 0.14988 \end{aligned}$$

Doğal olarak bu integral değerlerinin hesaplanmasını sizden beklememekteyiz. Binom dağılımında olduğu gibi, normal dağılımında olasılık değerlerini hesaplarken zaman zaman istatistiksel tablolardan, zaman zaman da MINITAB yazılımından faydalanacağız.

Minitab Uygulamaları

Normal dağılımın olasılık değerlerini bulmak için daha önce de müjdelediğimiz gibi karışık integrallerin sonuçlarını hesaplamıza gerek yoktur. Bu tarz hesaplamalarda MINITAB gibi istatistik programlarından faydalanabiliriz. **Örnek 6.4**’teki “yanlış” olasılık değerini (Denklem (6.10) daki $P(5)$ değeri gibi) hesaplamak için “Calc” (Hesapla) menüsünden “Probability Distributions” (Olasılık Dağılımları) ve “Normal”(Normal) dağılım seçeneğine tıklayınız. Karşınıza çıkan diyalog ekranında “Probability Density” (Olasılık Yoğunluğu) seçeneğine tıklayınız ve “mean” (ortalama) kutucuğuna “6” ($\mu = 6$), “standard deviation” (standart sapma) kutucuğuna “2” ($\sigma = 2$) değerlerini giriniz. Bir sonraki adım olarak, “Input Constant” (Veri Sabiti) kutucuğuna “5” değerini girdiğinizde Şekil 6.29’u elde edebilirsiniz.



Şekil 6.29

Diyalog ekranında “OK” tuşuna bastığınızda Şekil 6.30’daki çıktıyı elde edebilirsiniz.

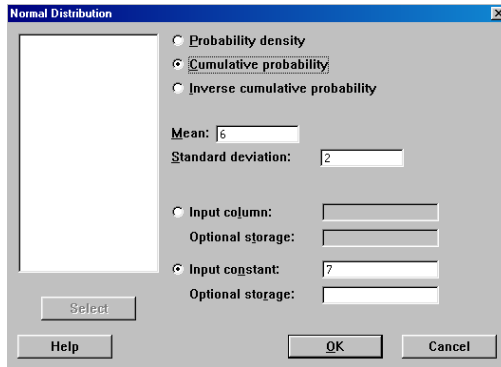
Probability Density Function	
Normal with mean = 6,00000 and standard deviation = 2,00000	
x	P(X = x)
5,0000	0,1760

Şekil 6.30

Bu çıktı değeri $P(5) = 0.176$ ’yı ifade etmektedir. Sürekli dağılımlarda, tek bir noktaya ait olan olasılık değeri daha öncede belirttiğimiz gibi sıfıra eşittir.

Örnek 6.5, 6.6, ve 6.7’deki olasılık değerleri hesaplayabilmemiz için “*Probability density*” (olasılık yoğunluğu) seçeneğinin yerine “*Cumulative probability*” (Kümülatif Olasılık) seçeneğine tıklamanız gerekmektedir. Şimdi **Örnek 6.5**’teki kümülatif olasılık değerini ($P(x < 7)$) hesaplayalım.

Bu hesaplamayı gerçekleştirebilmek için diyalog ekranını Şekil 6.31’deki gibi doldurunuz.



Şekil 6.31

Diyalog ekranında “OK” tuşuna bastığınızda Şekil 6.32’deki çıktıyı elde edebilirsiniz.

Cumulative Distribution Function	
Normal with mean = 6,00000 and standard deviation = 2,00000	
x	P(X <= x)
7,0000	0,6915

Şekil 6.32

(6.12) denkleminde hesapladığımız değere eşittir.

Örnek 6.6'daki olasılık değerini ($P(x > 7)$) hesaplamak için, $P(x > 7) = 1 - P(x \leq 7)$ eşitliğinden faydalanmamız gerekmektedir. $P(x < 7) = 0.6915$ değerini daha önce hesapladığımızdan dolayı $P(x > 7)$ değerini $P(x > 7) = 1 - 0.6915 = 0.3085$ olarak hesaplayabiliriz.

Son alıştırmamız olarak **Örnek 6.7**'deki olasılık değerini ($P(4 < x < 5)$) hesaplayalım. Bu olasılık değerini hesaplayabilmek için $P(4 < x < 5) = P(x < 5) - P(x < 4)$ eşitliğinden faydalanabiliriz. Şekil 6.31'de yer alan diyalog ekranındaki verileri değiştirerek, $P(x < 5)$ ve $P(x < 4)$ olasılık değerlerini hesaplayabiliriz.

Cumulative Distribution Function	
Normal with mean = 6,00000 and standard deviation = 2,00000	
x	P(X <= x)
5,0000	0,3085

Şekil 6.33

Cumulative Distribution Function	
Normal with mean = 6,00000 and standard deviation = 2,00000	
x	P(X <= x)
4,0000	0,1587

Şekil 6.34

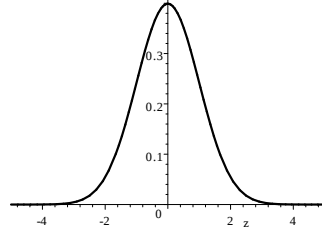
Elde ettiğimiz çıktıları kullanarak $P(4 < x < 5)$ değerini, $P(4 < x < 5) = P(x < 5) - P(x < 4) = 0.3085 - 0.1587 = 0.1498$ olarak hesaplayabiliriz.

Standart Normal Dağılım

Standart normal dağılım, bütün normal dağılıma sahip değişkenlerin dönüştürebileceği, sıfır ortalama ve bir standart sapma ile tanımlanan özel bir normal dağılım türüdür. Bütün normal dağılan değişkenlerin dönüştürebileceği standart normal dağılım değişkenini Z harfi ile gösterelim. Bu değişken, 0 ortalama ($\mu = 0$) ve 1 standart sapma ($\sigma = 1$) ile normal olarak dağılan standartlaştırılmış bir değişkendir. Normal dağılıma sahip herhangi bir X değişkeninin dönüşüm formülü aşağıdaki gibidir:

$$z = \frac{x - \mu}{\sigma} \quad (6.15)$$

Standart normal dağılımın (STND) şekli aşağıdaki gibidir:



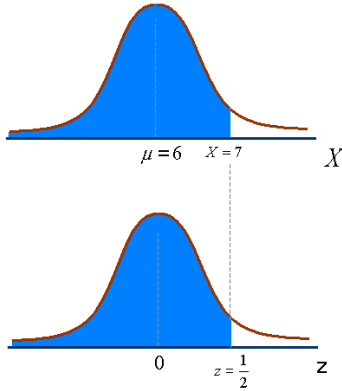
Şekil 6.35 STND ($\mu = 0, \sigma = 1$).

Z-skoru, X değerinin ortalamadan uzaklaştığı standart sapma sayısını ifade etmektedir. Eğer X değeri ortalamadan küçük ise, z skoru ise negatif bir değer; eğer X değeri ortalamadan büyük bir değer olursa, z -skoru pozitif değer alır.

Örnek olarak, X değişkenimizin $\mu = 6$ ortalaması ve $\sigma = 2$ standart sapması ile normal olarak dağıldığını düşünelim. x 'in herhangi bir değerine tekabül eden z değerini hesaplayalım. Eğer $x = 7$ ise, bu değere tekabül eden z değerini aşağıdaki gibi hesaplayabiliriz:

$$z = \frac{x - \mu}{\sigma} = \frac{7 - 6}{2} = \frac{1}{2}$$

x ve z değerlerinin dönüşümünü grafiksel olarak ifade edecek olursak



Şekil 6.36

Grafikten de kolayca algılanacağı gibi x 'in 7'den küçük olma olasılığı ile z değerinin $1/2$ 'den küçük olma olasılığı birbirine eşittir. $P(x < 7) = P(z < 1/2)$ eşitliğinden dolayı $P(x < 7)$ değerinin yerine $P(z < 1/2)$ değerini hesaplayabiliriz. $P(z < 1/2)$ değerinin hesaplanması $P(x < 7)$ değerinin hesaplanmasına çok benzemektedir. Z skoru, 0 ortalama ($\mu = 0$) ve 1 standart sapma ile ($\sigma = 1$) normal olarak dağılıyorsa, denklem (6.9)'da yer alan normal dağılım fonksiyonu aşağıdaki gibi olur:

Standart Normal Olasılık Fonksiyonu

$$f(z) = \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2} \frac{(z-0)^2}{1^2}} = \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2} z^2}$$

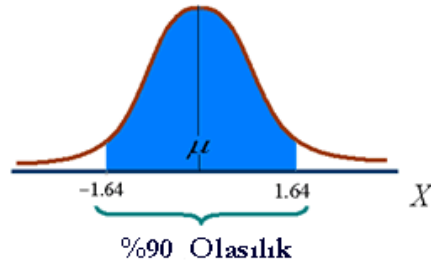
$P(z < 1/2)$ 'nin olasılık deęerini hesaplayabilmek iin ařaęıdaki belirli integralin hesaplanması gerekmektedir.

$$P(z < \frac{1}{2}) = \int_{-\infty}^{1/2} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}z^2} dz = 0.69146$$

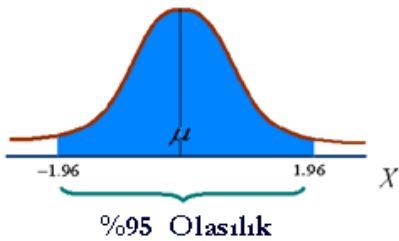
Elde etmiř olduęumuz sonu daha nce **rnek 6.5**'te elde ettięimiz $P(x < 7)$ sonucunun birebir aynısıdır.

Doęal olarak, neden bu kadar karıřık iřlemi bir de z deęeri iin yapmaktayoz sorusu aklınıza takılabilir. Eęer MINITAB gibi istatistik yazılımları ile olasılık deęerlerini hesaplıyorsanız, X deęerlerinin yerine Z deęerlerinin kullanmanın hibir avantajı yoktur. Arařtırmacıların bu tarz programlar kullanmadıęı zamanlarda, ellerinde olasılık tabloları bulunduęundan ve bu tablolardan btn kmlatif olasılık deęerleri hesaplanmış olduęundan X deęerleri bu tabloları kullanabilmek iin Z deęerlerine dnřtrlyordu. Bu yzden X yerine Z deęeri kullanmanın ciddi bir avantajı bulunmaktaydı. Ancak biz bu mekanizmanın daha iyi anlařılmasını saęlamak iin aralık tahmini ve hipotez testlerinde de X deęerlerini Z dnřm yaparak kullanmaya devam edeceęiz.

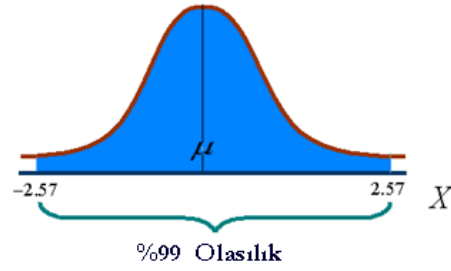
İleride kullanacaęımız bazı nemli olasılık deęerlerine tekabl eden bazı zel z deęerlerini bilmeniz gerekmektedir. Yaygın bir kullanım alanına sahip olan z deęerleri 1.645, 1.96, ve 2.575'tir. Bu deęerlere denk gelen olasılık deęerleri ise %90, %95, ve %99'dur. řimdi bu deęerleri grafik ortamında inceleyelim.



řekil 6.37



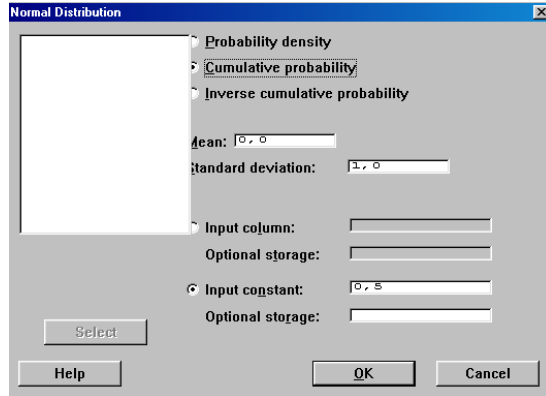
řekil 6.38



Şekil 6.39

Minitab Uygulamaları

Şimdi z değerlerini kullanarak örneklerimiz için hesapladığımız olasılık değerlerini MINITAB ortamında hesaplamaya çalışalım. $P(x < 7)$, (Örnek 6.5) olasılık değerini bir önceki MINITAB uygulamaları bölümünde gerçekleştirmiştik. (bkz. Şekil 6.20). Şimdi ise $P(z < 1/2)$ değerinin MINITAB ortamında nasıl hesaplandığını inceleyelim. $P(x < 7) = P(z < 1/2)$ eşitliğinden dolayı, $P(x < 7)$ için hesapladığımız değerin aynısını hesaplamak zorundayız. Bu olasılık değerini hesaplayabilmek için “*Calc-Probability Distributions*” seçeneğinden yine “*Normal*”a tıklayınız ve karşınıza çıkacak olan diyalog ekranını Şekil 6.40’taki gibi doldurunuz.



Şekil 6.40

Şekil 6.29 ile Şekil 6.40’ın arasındaki temel farklılık girilmiş olan ortalama ve standart sapma değerlerinden kaynaklanmaktadır. Şekil 6.40’ta standart normal dağılım ile karşı karşıya olduğumuzdan dolayı ortalama olarak 0 ve standart sapma olarak 1 değerini gireriz. $P(z < 1/2)$ olasılığını hesaplamak istediğimizden dolayı “*input constant*” kutucuğuna 0.5 değerini girdiğimizde Şekil 6.41’deki çıktı ekranını elde ederiz.

Cumulative Distribution Function	
Normal with mean = 0 and standard deviation = 1,00000	
x	P(X <= x)
0,5000	0,6915

Şekil 6.41

Elde etmiş olduğumuz çıktı, daha önce x değeri için elde ettiğimiz çıktı ekranının tamamiyle aynıdır. Dolayısıyla $x = 7$ ve $z = 1/2$ değerlerinin olasılıklarının birbirlerine eşit olduğunu göstermiş olduk.

Örnek 6.6'da $P(x > 7)$ değerini hesaplayabilmek için, $(P(x > 7) = P(z > 1/2)) P(z > 1/2) = 1 - P(z < 1/2)$ eşitliğinden faydalanabiliriz.

Örnek 6.7'deki $P(4 < x < 5) = ?$ olasılık değerini hesaplayabilmek için $P(4 < x < 5) = P(x < 5) - P(x < 4)$ özelliğinden faydalanabiliriz. Şimdi $x = 4$ ve $x = 5$ değerlerine tekabül eden z değerlerini hesaplayalım. (6.15)'teki dönüşüm formülünü kullanarak z değerlerini aşağıdaki gibi elde edebiliriz.

$$z_1 = \frac{x_1 - \mu}{\sigma} = \frac{4 - 6}{2} = -1$$

$$z_2 = \frac{x_2 - \mu}{\sigma} = \frac{5 - 6}{2} = -\frac{1}{2}$$

$P(-1 < z < -1/2) = P(z < -1/2) - P(z < -1)$ eşitliğinden faydalanarak MINITAB yardımı ile de $P(z < -1/2)$ ve $P(z < -1)$ olasılık değerlerini hesaplayabiliriz. (şekil 6.29'da 0.5 yerine, "input constant" kutucuğuna, -0.5 ve -1 değerlerini girerek hesaplayabiliriz) Bu değerleri Şekil 6.42'deki gibi elde edebiliriz.

Cumulative Distribution Function	
Normal with mean = 0 and standard deviation = 1,00000	
x	P(X <= x)
-0,5000	0,3085

Cumulative Distribution Function	
Normal with mean = 0 and standard deviation = 1,00000	
x	P(X <= x)
-1,0000	0,1587

Şekil 6.42

Dolayısıyla $P(-1 < z < -1/2)$ değeri $P(z < -1/2) - P(z < -1) = 0.3085 - 0.1587 = 0.1498$ olarak hesaplanabilir.

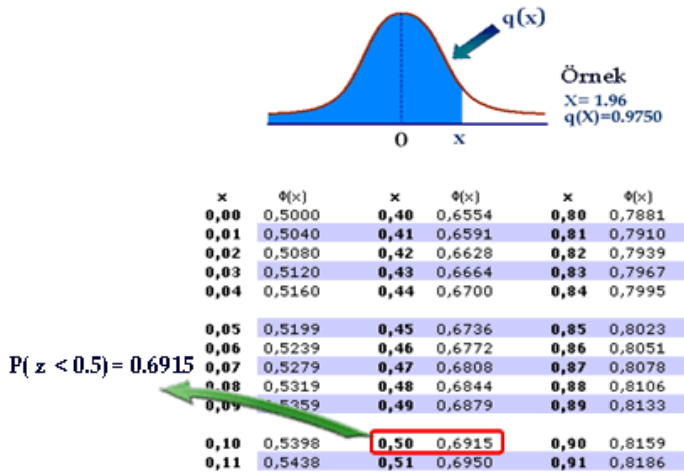
Normal Dağılımda İstatistiksel Tabloların Kullanımı

Bu bölümde **Örnek 6.5-6.7**'yi tekrar ele alarak olasılık değerlerini istatistiksel tabloları kullanarak hesaplamaya çalışacağız. **Örnek 6.5**'teki $P(x < 7)$ olasılık değerini hesaplamak için öncelikle $x = 7$ (normal dağılmış olan değişkeni) tabloda bütün olasılık değerleri listelenmiş olduğundan standardize edelim ve z değerine dönüştürelim.

$$z = \frac{x - \mu}{\sigma} = \frac{7 - 6}{2} = \frac{1}{2}$$

Dolayısıyla artık dağılım tablosunu kullanarak $P(z < 1/2)$ değerini hesaplamamız gerekecektir. Kitabınızın Ekler bölümünde standart normal tablosunu bulabilirsiniz. Bu tablo farklı z değerleri için kümülatif olasılık değerlerini vermektedir. $P(z < 1/2)$ 'nin kümülatif olasılık değerini Şekil 6.42'deki gibi bulabiliriz.

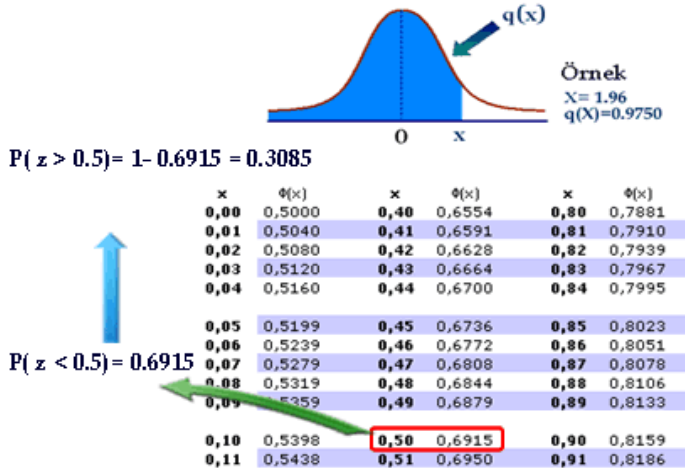
TABLO 1 NORMAL DAĞILIM FONKSİYONU



Şekil 6.42

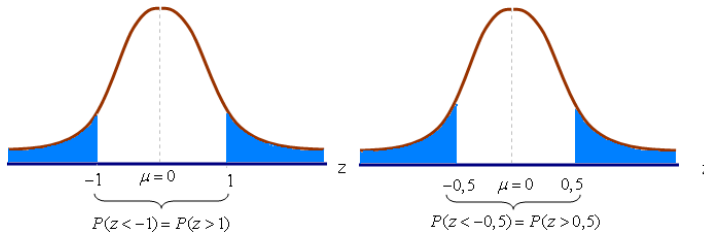
Örnek 6.6'da $P(z < 1/2)$ olasılık değerini $P(z > 1/2) = 1 - P(z < 1/2)$ özelliğinden faydalanarak hesaplayabiliriz. Dolayısıyla $P(z > 1/2) = 1 - P(z < 1/2) = 1 - 0.6915 = 0.3085$ değerini bulabiliriz.

TABLO 1 NORMAL DAĞILIM FONKSİYONU



Şekil 6.43

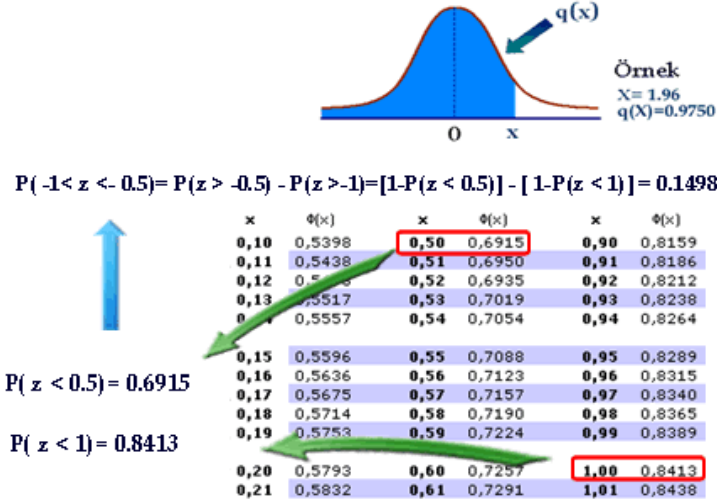
Örnek 6.7'deki $P(-1 < z < -1/2)$ olasılık değerini hesaplamak için yine $P(-1 < z < -1/2) = P(z < -1/2) - P(z < -1)$ özelliğinden faydalanabiliriz. Öncelikle $P(z < -1/2)$ ve $P(z < -1)$ olasılık değerlerini tablodan hesaplamamız gerekmektedir. Tabloda negatif değerler bulunmadığından, normal dağılımın simetri özelliğinden faydalanarak z_0 $P(z < -z_0) = P(z > z_0)$ denklğini elde edebiliriz. Dolayısıyla, $P(z < -1/2) = P(z > 1/2)$ ve $P(z < -1) = P(z > 1)$ sonucunu elde ederiz. Şimdi bu durumu grafiksel ortamda inceleyelim.



Şekil 6.44

$P(z < -1/2)$ değerini hesaplayabilmek için $P(z > 1/2)$ veya $P(z < -1/2)$ değerini hesaplamamız gerekmektedir. $P(z > 1/2)$ değerini hesaplarken $P(z > 1/2) = 1 - P(z < 1/2)$ eşitliğinden faydalanabiliriz.

TABLO 1 NORMAL DAĞILIM FONKSİYONU



Şekil 6.45

Değişkenimizin belli bir dağılıma göre dağılıp dağılmadığını nereden bilebiliriz?

Olasılık dağılımı ile uğraşan öğrencilerin en sık sordukları sorulardan biri hangi olasılık dağılımının kullanılacağına nasıl karar vereceğimiz hususundadır. Binom olasılık dağılımı bölümünde, deneyin sadece iki adet çıktısı olduğuna dikkat çekmiştik. Bu özelliğine ek olarak, n benzer deneme sonucunda bu denemeler birbirinden bağımsız olarak gerçekleştiğinden değişkenimiz binom dağılımına sahiptir fikrini öne sürebiliyorduk.

Peki değişkenimizin normal olarak dağıldığını hangi özelliklerinden anlayabiliriz? İstatistikte değişkenlerin normal dağıldığına dair kesin varsayımlar bulunmaktadır. Örneğin bir sonraki ünite de deyişeceğimizi Merkezi Limit Teoremi bu varsayımlardan birisidir. Bu teorem belli koşullar altında örneklem ortalamalarının ($\bar{x}$ lar) normal olarak dağıldığını göstermektedir.

Şimdi ise t dağılımı, ki-kare dağılımı, ve F dağılımı olmak üzere diğer önemli sürekli olasılık dağılımlarına deyişeceğiz. Bu dağılımlarda normal dağılım gibi belli istatistiksel teoremler doğrultusunda dağılmaktadır. Dolayısıyla bu teoremlerin hangi değişkenin hangi dağılım sahip olacağını anlama hususunda bize çok yardımcı dokunacaktır.

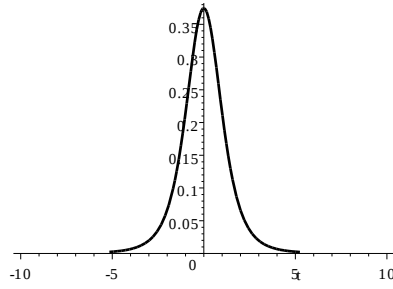
t , Ki-kare, ve F dağılımları

Şuana kadar sürekli dağılımlara örnek olarak sadece normal dağılımı irdeledik. Ancak dersimizde sürekli dağılımlar başlığı altında inceleyeceğimiz üç tane daha sürekli dağılım bulunmaktadır. Bu dağılımlar t dağılımı, ki-kare dağılımı, ve F dağılımıdır. Bütün bu sürekli dağılımlar, **serbestlik derecesi** olarak adlandırdığımız bir parametreye bağlı olarak dağılmaktadır. Belirli koşullar altında, **serbestlik derecesi sayısı** (genel olarak **df** veya **D.F.** ile ifade edilmektedir.) toplam kaç gözlem değerini serbestçe kullanabildiğimizi gösterirken, toplam gözlem sayısından aralarında bağımsız ilişki bulunan kullanabildiğimiz gözlem sayısını çıkardığımızda elde edilebilir.

t dağılımının yapısı, standart normal dağılıma çok benzemektedir. Benzer bir şekile sahiptir, sabit bir ortalaması ve varyansı vardır. Ortalama değeri, standart normal dağılımda olduğu gibi 0'a eşittir ve varyansı her zaman için $n/(n - 2)$ değerine eşittir. (hatırlanacağı üzere standart normal dağılımın varyansı da her zaman 1'e eşit idi) Standart normal dağılımdan farklı olarak, standart normal dağılım ortalama ve varyans olmak üzere iki farklı parametreye bağımlı olarak dağılırken, t dağılımı tek bir parametreye, serbestlik derecesine, bağlı olarak dağılmaktadır. t dağılımının serbestlik derecesi, n 'nin (n gözlem sayısını ifade etmektedir) farklı durumlarına göre farklı değerler almaktadır. Örneğin, serbestlik derecesi duruma ve şartlara göre $n - 1$, $n - 2$ gibi değerler olabilir.

Dersimizin amacıyla, hem t dağılımının hem de ki-kare dağılımının tek bir parametreye (serbestlik derecesine) bağlı olarak dağıldığını göstermek yeterlidir.

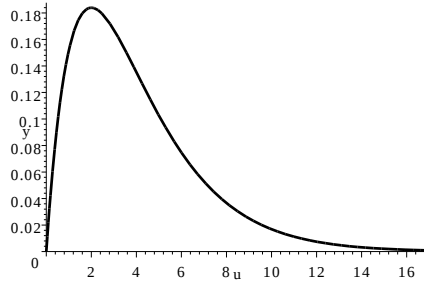
Şekil 6.46, 4 serbestlik derecesi ($n - 1 = 4$) ile dağılan t dağılımını göstermektedir:



Şekil 6.46 t dağılımı $df = 4$

Ki-Kare dağılımı da t dağılımı gibi tek bir parametreye, serbestlik derecesine bağlı olarak dağılmaktadır. **Ki-kare** dağılımının serbestlik derecesi, n 'nin (n gözlem sayısını ifade etmektedir) farklı durumlarına göre farklı değerler almaktadır. Eğer bu parametreyi biliyorsak, ki-kare ile dağılan değişkenin olasılık değerleri hakkında fikir yürütebiliriz.

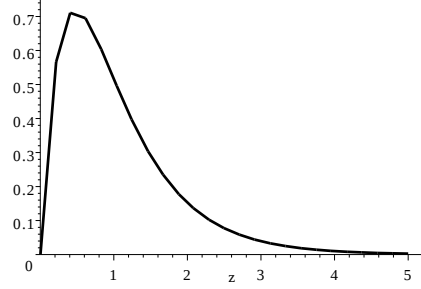
Şekil 6.47, 4 serbestlik derecesi ($n - 2 = 4$) ile dağılan ki-kare dağılımını göstermektedir:



Şekil 6.47 Ki-kare $df = 4$

Şekil 6.47'den de anlaşılabileceği gibi ki-kare dağılımı pozitif eksenlerde tanımlanmıştır. Dolayısıyla, normal ve t dağılımına sahip olan değişkenlerden farklı olarak, ki-kare ile dağılan değişkenlerin sadece pozitif değerler aldığını varsaymaktayız.

Üçüncü sürekli dağılımımız ise F dağılımıdır. F dağılımı, t dağılımı ve ki-kare dağılımından farklı olarak iki farklı serbestlik derecesine (df_1 ve df_2) bağlı olarak dağılmaktadırlar. Şekil 6.48'de da yer alan F dağılımı 4 ve 49 serbestlik derecelerine göre dağılıma göstermektedirler.



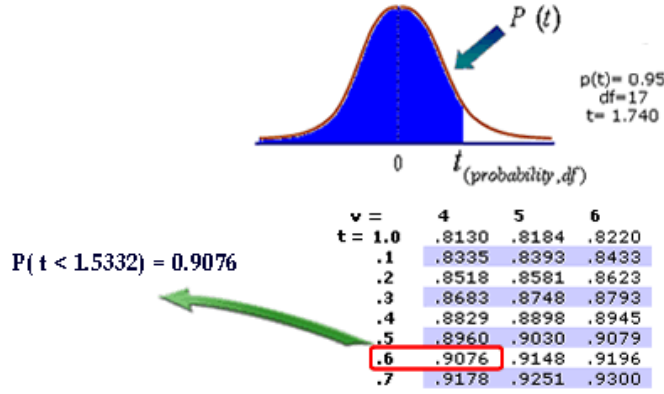
Şekil 6.48 F dağılımı $df_1 = 4$, $df_2 = 40$.

Şekilden de anlaşılacağı gibi F dağılımı da pozitif eksenlerde tanımlanmıştır. Dolayısıyla, normal ve t dağılımına sahip olan değişkenlerden farklı olarak, F ile dağılan değişkenlerin sadece pozitif değerler aldığını varsaymaktayız.

t dağılımı, Ki-kare dağılımı ve F dağılımı ile dağılan değişkenlerin istatistiksel tabloları kullanarak olasılık değerlerinin hesaplanması

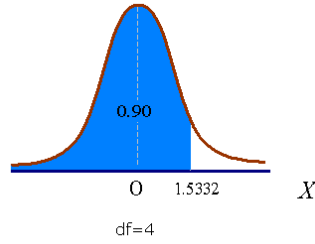
t dağılımı ile dağılan ve serbestlik derecesi $df = 4$ olan bir değişkenimiz olduğunu varsayalım. Bu değişkenin 1.5332'den küçük olma olasılığını hesaplamak istediğimizi düşünelim ($P(t < 1.5332)$). Bu hesaplamayı Ekler bölümünde yer alan t-tablosunu kullanarak gerçekleştirebiliriz. Bu tabloda farklı t değerlerine (satırlar) ve serbestlik derecelerine (sütunlar) denk gelen olasılık değerleri bulunmaktadır. $P(t < 1.5332)$ olasılık değerinin tablodan nasıl bulunacağı şekil 6.37'de gösterilmektedir.

TABLO 2 t DAĞILIMI FONKSİYONU



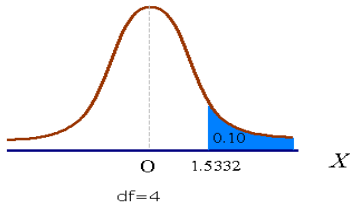
Şekil 6.49

$P(t < 1.5332) = 0.9$ değerini grafiksel olarak gösterecek olursak:



Şekil 6.50

$P(t > 1.5332)$ değerini hesaplamak istediğimizde ise bulduğumuz değeri birden çıkarabiliriz. $P(t > 1.5332) = 1 - 0.9 = 0.1$. Grafiksel olarak,

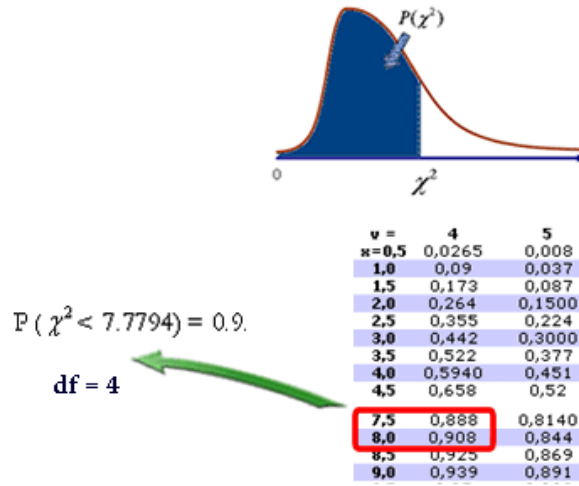


Şekil 6.51

Şimdi ise değişkenimizin ki-kare dağılımı ile ve $df = 4$ serbestlik derecesiyle dağıldığını varsayalım. Bu değişkenin 7.7794'ten küçük olma olasılığını hesaplamak istediğimizi düşünelim ($P(\chi^2 < 7.7794)$). Bu hesaplamayı Ekler bölümünde yer alan ki-kare tablosunu kullanarak gerçekleştirebiliriz. Bu tabloda farklı ki-kare değerlerine (satırlar) ve serbestlik derecelerine

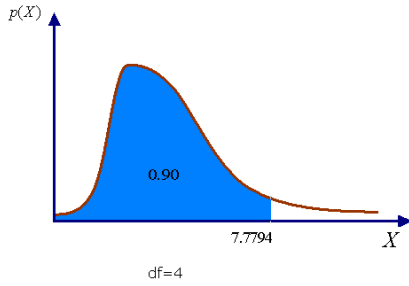
(sütunlar) denk gelen olasılık değerleri bulunmaktadır. $P(\chi^2 < 7.7794)$ olasılık değerinin tablodan nasıl bulunacağı Şekil 6.52’de gösterilmektedir.

TABLO 5 Kİ-KARE DAĞILIM FONKSİYONU



Şekil 6.52

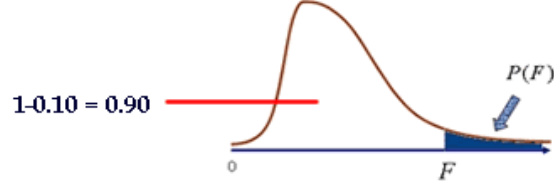
$P(\chi^2 < 7.7794) = 0.9$ değerini grafiksel olarak gösterecek olursak:



Şekil 6.53

Son olarak ise değişkenimizin F dağılımı ile ve $df_1 = 4$, $df_2 = 40$ serbestlik dereceleriyle dağıldığını varsayalım. Bu değişkenin 2.0909’dan küçük olma olasılığını hesaplamak istediğimizi düşünelim ($P(F < 2.0909)$). Malesef, F tablosunun biraz daha karışık olmasından dolayı F değerlerinin olasılık hesaplamalarında MINITAB programını kullanacağız.

TABLO 4(a) %10'na GÖRE F DAĞILIMI



v1 =	1	2	3	4	5	6	7	8	10	12	24	
v2 = 1	39.86	49.50	53.59	55.83	57.24	58.20	58.91	59.44	60.19	60.71	62.00	63.33
2	8.526	9.000	9.162	9.243	9.293	9.326	9.349	9.367	9.392	9.408	9.450	9.491
3	5.538	5.462	5.391	5.343	5.309	5.285	5.266	5.252	5.230	5.216	5.176	5.134
30	2.881	2.489	2.276	2.142	2.049	1.980	1.927	1.884	1.819	1.773	1.638	1.456
32	2.869	2.477	2.263	2.129	2.036	1.967	1.913	1.870	1.805	1.758	1.622	1.437
34	2.859	2.466	2.252	2.118	2.024	1.955	1.901	1.858	1.793	1.745	1.608	1.419
36	2.850	2.456	2.243	2.108	2.014	1.945	1.891	1.847	1.781	1.734	1.595	1.404
38	2.842	2.448	2.234	2.099	2.005	1.935	1.881	1.838	1.772	1.724	1.584	1.390
40	2.835	2.440	2.226	2.091	1.997	1.927	1.873	1.829	1.763	1.715	1.574	1.377
60	2.791	2.393	2.177	2.041	1.946	1.875	1.819	1.775	1.707	1.657	1.511	1.291
120	2.748	2.347	2.130	1.992	1.896	1.824	1.767	1.722	1.652	1.601	1.447	1.193
	2.706	2.303	2.084	1.945	1.847	1.774	1.717	1.670	1.599	1.546	1.383	1.000

Minitab Uygulamaları

Şimdi hesapladığımız olasılık değerlerini MINITAB ortamında hesaplayalım. t dağılımına sahip ve serbestlik derecesi 4 olan $P(t < 1.5332)$ olasılık değerini hesaplamak için, “Calc” menüsünden “Probability Distributions” (Olasılık Dağılımlarına) tıklayınız. “t...” dağılımını seçtikten sonra karşınıza çıkan diyalog ekranını Şekil 6.54’teki gibi doldurunuz.

t Distribution

☐ Probability density

☒ Cumulative probability

☐ Inverse cumulative probability

Degrees of freedom: 4

Input column:

Optional storage:

☒ Input constant: 1.5332

Optional storage:

Select

Help

OK

Cancel

Şekil 6.54

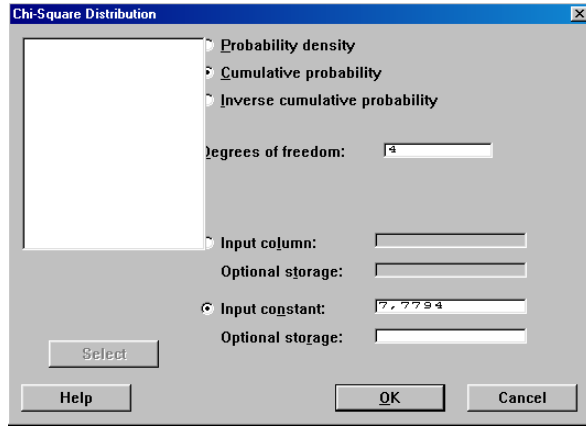
OK tuşuna bastığınızda aşağıdaki çıktıyı elde edebilirsiniz.

Cumulative Distribution Function	
Student's t distribution with 4 DF	
x	P(X <= x)
1,5332	0,9000

Şekil 6.55

$P(t < 1.5332) = 0.9$ daha önce hesapladığımız değerin aynısını elde ederiz.

Ki-kare dağılımına sahip ve serbestlik derecesi $df = 4$ olan $P(\chi^2 < 7.7794)$ değerini hesaplamak için, “Calc” menüsünden “Probability Distributions” (Olasılık Dağılımlarına) tıklayınız. “Chi-Square..” dağılımını seçtikten sonra karşınıza çıkan diyalog ekranını Şekil 6.56’daki gibi doldurunuz.



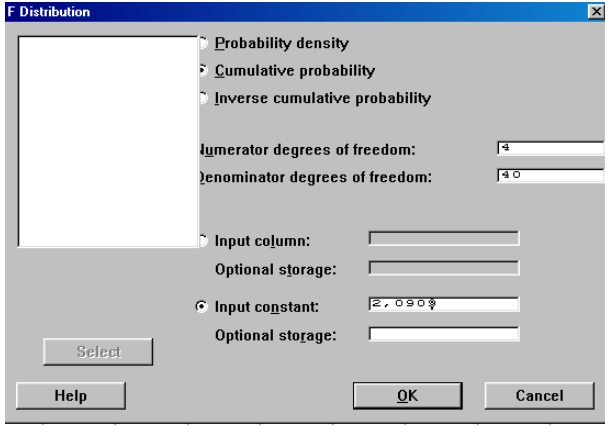
Şekil 6.56

Elde edeceğiniz çıktı aşağıdaki gibi olacaktır.

Cumulative Distribution Function	
Chi-Square with 4 DF	
x	P (X <= x)
7.7794	0.9000

Şekil 6.57

F dağılımına sahip ve serbestlik dereceleri $df_1 = 4$, $df_2 = 40$ olan t_0 olan $P(F < 2.0909)$ değerini hesaplamak için, “Calc” menüsünden “Probability Distributions” (Olasılık Dağılımlarına) tıklayınız. “F..” dağılımını seçtikten sonra karşınıza çıkan diyalog ekranını Şekil 6.58’deki gibi doldurunuz.



Şekil 6.58

Dolayısıyla elde edeceğimiz çıktı aşağıdaki gibi olacaktır.

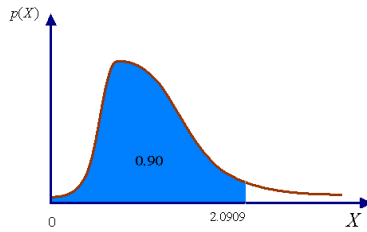
Cumulative Distribution Function

F distribution with 4 DF in numerator and 40 DF in denominator

x	P (X <= x)
2,0909	0,9000

Şekil 6.59

$P(F < 2.0909) = 0.9$ Grafiksel olarak;



Şekil 6.60

Sonuç

Bu bölümde farklı olasılık dağılımlarını inceledik. İlerleyen bölümlerde hangi dağılımı ne zaman kullanacağımızı inceleyeceğiz. Olasılık dağılımlarının nerede nasıl kullanılacağı konusu;

hipotez testlerine, aralık tahminlerine temel teşkil edecektir. Dolayısıyla bu bölümde okuduklarınızı iyice anlamadan geçmemeye çalışınız.

Bölüm 7:

Örnekleme ve Örneklem Dağılımları

Giriş

Yedinci bölümde, çıkarımsal istatistik kavramına daha detaylı bir giriş yapmaktayız. Çıkarımsal istatistik, yöntemlerinin amacı bilinmeyen bir anakütle hakkında bilgi toplama anlamına gelmektedir. Bir başka deyişle, anakütleye ait özellikleri mümkün olduğunca iyi tahmin edebilmektir. Birinci bölümden hatırlayacağınız üzere, bir anakütlenin tamamını incelemek ya tamamen imkansızdır, ya da çok maliyetlidir, bu yüzden de bilgi elde etmek için ilgili anakütleden örneklem(ler) alınır.

Örneklem verilerinden hesaplanan özet bilgiler (ortalama, standart sapma gibi istatistikler), anakütle özelliklerinin (parametrelerinin) tahmininde kullanıldığı için, bu istatistiklere aynı zamanda “*tahmin edici*” adı verilir.

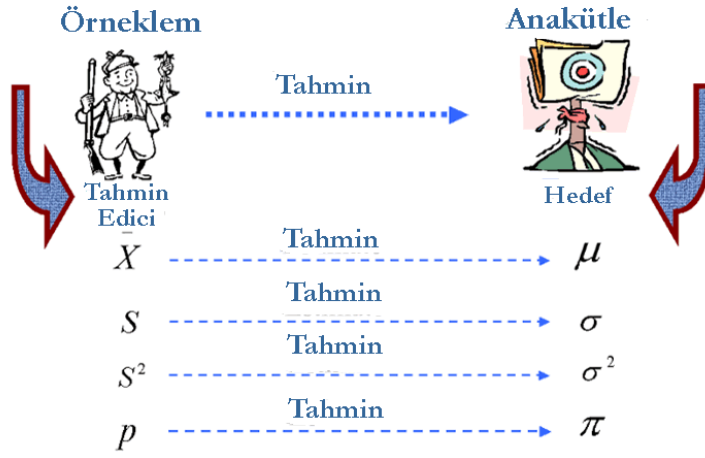
Örneğin bir ampül üretim fabrikasında, yeni tasarım ampullerin kullanım süresinin uzunluğunu araştırdığımızı düşünelim. Kullanım süresinin uzunluğu için bütün ampullerin anakütlesini incelemek hem maliyetli hem de uzun zaman alacağından dolayı test için örneklem olarak 100 adet ampul seçelim ve örneklem ortalamasını hesaplayalım. Örneklem ortalamasının 65 saat olduğunu düşünelim. 65 saat olan örneklem ortalamasını kullanarak anakütle karakteristiği hakkında fikir yürütebiliriz. Burada önemli olan tahmin edicilerimizin ne kadar iyi olduğudur. Tahmin edicilerin de ne kadar iyi olduğu uygun örneklem seçilip seçilmemesi ile alakalıdır.

Basit rassal Örnekleme, anakütleden örneklem seçilerek kullanılır. Eğer anakütlenin boyutunun sonlu olduğunu varsayarsak, N boyutlu bir anakütleden seçilen n boyutlu rassal örneklemelerin hepsinin seçilme olasılığı eşittir. Basit rassal Örnekleme geri yerleştirmeli veya geri yerleştirmesiz olarak yapılabilir.

Geri yerleştirmesiz Örnekleme metodu en sık kullanılan Örnekleme metodudur.

Nokta Tahmini

Anakütle verilerini elde edemediğimizden dolayı, örneklem verilerinden yola çıkarak anakütle karakteristiği hakkında fikir yürütebilirsiniz. Tahmin edicilerin formülleri, belli kurallar altında matematiksel olarak türetilmektedir. Tahmin süreci ise bu formüllerden faydalanarak, örneklem verileri kullanma prensibine dayanmaktadır. Bu yüzden anakütle ortalaması olan μ değerini ve anakütle standart sapması olan σ değerini tahmin edebilmek için bu değerlere denk gelen, ortalamanın örneklem istatistiği olan $\bar{X}$ değerini ve standart sapmanın örneklem istatistiği olan s değerini hesaplamamız gerekmektedir. Süreksiz değişkenlerde ise ortalama veriler oransal terimlerle ifade edilirler. π anakütle oranını ifade ederken, p ise örneklem oranını ifade etmektedir.



Şekil 7.1

Tablo 7.1 Tahmin Ediciler

	Anakütle Parametresi	Tahmin Edici
Ortalama	μ	$\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$
Varyans	σ^2	$s^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$
Standard Sapma	σ	$s = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2}$
Oran	π	$\bar{p} = \frac{x}{n}$

Tahmin ediciler konusunda özellikle vurgulanması gereken husus, tahmin edicilerinde rassal değişkenler olduğudur. Bu değişkenler rassal değişkenlerin özel bir durumunu temsil etmektedirler. Peki neden rassallar? Örneklem ortalamasının $\bar{x}$ matematiksel formülünü ele alalım:

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i = \frac{1}{n} (x_1 + x_2 + x_3 + \dots + x_n)$$

Formülümüzdeki her x değişkeni (x_1, x_2 vs.) rassal olduğundan dolayı, $\bar{X}$ örneklem ortalaması da rassal bir değişkendir.

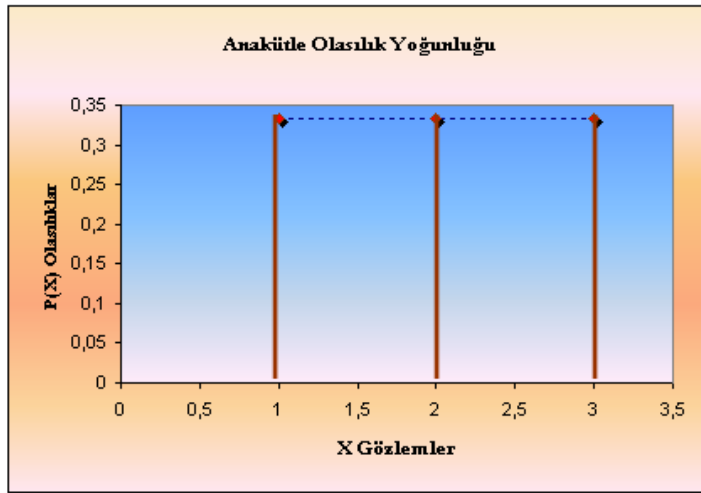
Dilerseniz şuna kadar anlattıklarımızın daha iyi anlaşılabilmesi için, anlattıklarımızı bir örnek doğrultusunda inceleyelim.

Örnek 7.1: Doğruluktan ziyade kolaylık olması açısından, yalnızca üç veriden oluşan bir anakütle düşünelim. Anakütle büyüklüğünü N ile gösterirsek, bu durumda $N = 3$ olur.

$$Anakütle = \{1, 2, 3\}$$

İlgilendiğimiz tahmin edicinin “ortalama” olduğunu düşünelim Buradan anakütle ortalamasının $\mu = 2$ olduğunu görebiliriz. Bu anakütlenin dağılımı yeknesak bir dağılım olduğundan her sayıya tekabül eden olasılık değeri $1/3$ tür.

Dağılımı aşağıdaki gibi görselleştirebiliriz:



Şekil 7.1 Anakütle olasılık Yoğunluğu

Şimdi ise anakütlemizden, geri yerleştirme ile seçilebilecek, boyutu $n = 2$ olan bütün örneklemeleri ele alalım. Tahmin edici olarak şimdilik örneklem ortalamasını ele alalım. Daha öncede belirttiğimiz gibi örneklemimizin ortalamasını aşağıdaki formül aracılığıyla hesaplayabiliriz:

$$\bar{x} = \frac{1}{n} = \sum_{i=1}^n x_i = \frac{1}{2} (x_1 + x_2)$$

Geri yerleştirme ile boyutu $n = 2$ olan 9 farklı örneklem oluşturabiliriz:

$$\begin{array}{llll} S_1 = \{1, 1\} & S_5 = \{2, 2\} & S_2 = \{1, 2\} & S_6 = \{2, 3\} \\ S_3 = \{1, 3\} & S_7 = \{3, 1\} & S_4 = \{2, 1\} & S_8 = \{3, 2\} \end{array}$$

⁹ Anakütle ortalamasını veya beklenen değeri aşağıdaki gibi hesaplayabiliriz:

$$E(X) = \mu = \frac{1}{3} \times 1 + \frac{1}{3} \times 2 + \frac{1}{3} \times 3 = 2$$

$$S_9 = \{3, 3\}$$

Bu örneklemelerin herbirini seçme olasılığımız nedir ? Geri yerleştirme metodu kullandığımızdan her birini seçme olasılığımız $1/9$ 'a eşit olur.

Bu örneklemelere bağlı olarak hesaplanabilecek örneklem ortalamaları ise aşağıdaki gibi olacaktır:

$$\begin{aligned}\bar{x}_1 &= 1, \bar{x}_2 = 1.5, \bar{x}_3 = 2 \\ \bar{x}_4 &= 1.5, \bar{x}_5 = 2, \bar{x}_6 = 2.5 \\ \bar{x}_7 &= 2, \bar{x}_8 = 2.5, \bar{x}_9 = 3\end{aligned}$$

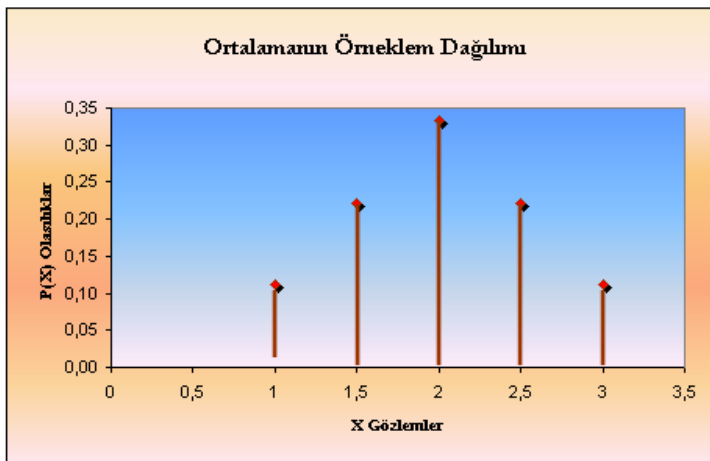
Gördüğümüz gibi, örnekleme seçilen değerlere bağlı olarak, örneklem ortalaması 1, 1.5, 2, 2.5, veya 3 olabiliyor. Dolayısıyla örneklem ortalamasının bir rassal değişken olduğunu gözlemleyebiliriz. Peki, bu değerlerin ortaya çıkma olasılıkları nedir? 9 örneklem 3 tanesi anakütle ortalamasını doğru tahmin etmiştir ($3/9$). İkişer tane örneklemimizden 1.5 ve 2.5 değerleri bulunmuştur ($2/9$), birer örneklem de 1 ve 3 değerlerini vermiştir ($1/9$).

Örneklem ortalaması, her ne kadar gerçek anakütle değerini ortalama olarak doğru tahmin etse de, ortalamaların bir kısmı örnekleme seçilen verilerin şans eseri hep küçük veya hep büyük olmaları nedeniyle (örneklem dalgalanması), bir miktar örneklem hatası barındırmaktadır (anakütle ortalamasından farklılaşmaktadır).

Bu basit örnekten yola çıkarak, $n = 2$ boyutundaki bir örneklem ortalamasına ilişkin olasılık dağılımını (örneklem dağılımı) çizebiliriz. Bu olasılık dağılımına karşılık gelen olasılık fonksiyonu aşağıdaki gibidir:

$$P(\bar{X}) = \begin{cases} \bar{X} = 1, & P(\bar{X}) = 1/9 \\ \bar{X} = 1.5, & P(\bar{X}) = 2/9 \\ \bar{X} = 2, & P(\bar{X}) = 3/9 \\ \bar{X} = 2.5, & P(\bar{X}) = 2/9 \\ \bar{X} = 3, & P(\bar{X}) = 1/9 \end{cases}$$

Olasılık dağılımını gösteren fonksiyon ise şu biçimde gösterilebilir:



Şekil 7.2 Ortalamanın Örneklem Dağılımı

Görüldüğü gibi tahmin edicinin de olasılık dağılımına örneklem dağılımı denilmektedir. Bu yüzden, şekil 7.2 örneklem ortalamasının olasılık dağılımını göstermektedir. Örneklem ortalamasının örneklem dağılımı, ileride de daha detaylı inceleyeceğimiz gibi, istatistikte çok önemli bir kavramdır.

Örnek 7.2: Şimdi ise tablo 7.1’de ikinci tahmin edici olarak gösterilen ve örneklem varyansı olarak isimlendirilen tahmin ediciyi inceleyelim. Örneklem varyansının matematiksel formülü aşağıdaki gibidir:

$$s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2 = \frac{1}{n-1} \left[(x_1 - \bar{x})^2 + (x_2 - \bar{x})^2 + \dots + (x_n - \bar{x})^2 \right]$$

Örneklem ortalamasında da bahsettiğimiz gibi her x_i değerimiz rassal birer değişken olduğundan s^2 ’de (örneklem varyansı) rassal bir değişkendir. Bir önceki örneğimizin verilerini kullanarak, örneklem varyanslarının değerlerini hesaplayalım. Örneğin, birinci örneklemin örneklem varyansını (S_1) aşağıdaki gibi hesaplayabiliriz:

$$s_1^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2 = \frac{1}{2-1} \left[(1-1)^2 + (1-1)^2 \right] = 0$$

Açıktır görüldüğü gibi anakütle varyansını $\sigma^2 = 2/3^{10}$ tahmin etmek için seçebileceğimiz en kötü örneklemlerden birini seçip varyans değerini hesapladık. Benzer hesaplamaları geri kalan örneklemde uygulayıp örneklem varyansların örneklem dağılımını elde edebiliriz.

$$\begin{aligned} s_1^2 &= 0, \quad s_2^2 = 0.5, \quad s_3^2 = 2. \\ s_4^2 &= 0.5, \quad s_5^2 = 0, \quad s_6^2 = 0.5, \\ s_7^2 &= 2, \quad s_8^2 = 0.5, \quad s_9^2 = 0, \end{aligned}$$

Tabloda da yer alan bir diğer tahmin edici, örneklem oranıdır. Adından da anlaşılacağı üzere bu tahmin edici oranlarla hesaplandığından kategorik verilerde kullanılmaktadır. Örneğin, Türkiye’deki erkek seçmenlerin oranını tahmin etmeye çalıştığımızı düşünelim. Bu oranı tahmin edebilmemiz için bir örneklem seti çekip tablo 7.1’de yeralan örneklem oranı formülünden faydalanarak tahmin de bulunabiliriz. Bu formülde yer alan X değeri örnekleme yer alan erkek seçmen sayısını temsil ederken, n değeri de örnekleme yer alan kişi sayısını göstermektedir.

Örnek 7.3: Anakütlemizin, $anakütle = \{x_p, x_y, y_1\}$ gibi üç değişkenden oluştuğunu düşünelim. x_1 , X partisine oy vermeyi düşünen birinci kişiyi temsil ederken; x_2 , X partisine oy vermeyi düşünen ikinci kişiyi, y_1 ise Y partisine oy vermeyi düşünen tek kişiyi temsil etmektedir. Bu bağlamda, X partisine oy verenlerin oranı:

$$\pi = \frac{X}{N} = \frac{2}{3}$$

¹⁰ Daha önce öğrendiğimiz gibi :

$$\sigma^2 = E(X - \mu)^2 = \frac{1}{3} \left[(1-2)^2 + (2-2)^2 + (3-2)^2 \right] = \frac{2}{3}$$

Şimdi normalde de olduğu gibi anakütleyi gözlemleyemediğimizi, sadece n boyutlu bir örneklemi gözlemleyebildiğimizi düşünelim. Seçtiğimiz birinci örneklem $S_1 = \{x_1, x_2\}$ şeklinde ise örneklem oranı formülümüzü bu örnekleme uyguladığımızda örneklem oranını aşağıdaki gibi elde edebiliriz.

$$p = \frac{X}{n} = \frac{2}{2} = 1$$

Böylelikle, 2/3 olan anakütle oranını 1 olarak tahmin etmiş olduk. Anakütlemizden boyutu 2 olan 9 farklı örneklem elde edebiliriz.

$S_1 = \{x_1, x_1\}; S_2 = \{x_1, x_2\}; S_3 = \{x_1, y_1\}; S_4 = \{x_2, x_2\}; S_5 = \{x_2, x_1\};$
 $S_6 = \{x_2, y_1\}; S_7 = \{y_1, y_1\}; S_8 = \{y_1, x_1\}; S_9 = \{y_1, x_2\}$

Örneklemlerimizin örneklem oranlarını hesaplayacak olursak;

$$p_1 = 1, p_2 = 1, p_3 = 1/2$$

$$p_4 = 1, p_5 = 1, p_6 = 1/2$$

$$p_7 = 0, p_8 = 1/2, p_9 = 1/2$$

Örnek 7.4: Daha önceki örneklerimizde anakütlelerimizin sadece üç elemmandan oluştuğunu varsaymıştık. Anakütle boyutunu bu kadar küçük almamızın nedeni seçebileceğimiz bütün olası örneklem setlerini yazabilmek idi. Bu örneklem setlerinin böyle detaylı bir şekilde yazılması bir sonraki bölümde ayrıntılı bir şekilde ele alınacak olan kavramlara bir temel teşkil etmesinden kaynaklanmaktadır. Anakütle boyutumuzu şuana kadar $N=3$ olarak aldığımızdan, bu anakütleye tekabül eden olası örneklemelerimizin sayısını 9 ($=3^2$) olarak hesaplamıştık.

Şimdi ise anakütle boyutumuzu biraz arttıralım ve anakütlemizin ($N = 10$): $P = \{1, 2, 15, 15, 17, 18, 18, 20, 21, 35\}$ gibi olduğunu düşünelim. Anakütlemizin anahtar parametrelerini hemen anakütle ortalaması ($\mu = 16.20$) ve anakütle standart sapması ($\sigma = 9.13$) olarak hesaplayabiliriz.

Yeni tanımladığımız anakütle hakkında fikrimiz olmadığını elimizde yine sadece boyutları 3 olan örneklemeler olduğunu varsayalım. $S = \{15, 17, 18\}$ gibi örneklemelerden 1000 adet elde edebiliriz. ($10^3 = 1000$ boyutu 3). Elimizde sadece örneklem olduğundan örneklem ortalamasını ve örneklem standart sapmasını $\bar{x} = 16.66$ ve $s = 1.52$ olarak hesaplayabiliriz. Örneklem ortalamasını 16,66 bularak anakütle ortalaması olan 16,20 değerine yakın bir tahminde bulunmuş olduk. Ancak, anakütle standart sapması olan 9.13 değerini tahmin ederken bulduğumuz örneklem standart sapması (s) 1.52 bu değerden hayli uzaktır.

Şimdi aynı örneklem boyutunda olan bir diğer örneklem daha alalım. ($n = 3$): $S = \{1, 2, 15\}$. Bu örneklemimize ait olan örneklem ortalaması ve standart sapma değeri $\bar{x} = 6$ ve $s = 7.81$ 'dir. Bu örnekte anakütle standart sapmasının değeri daha yakın bir şekilde tahminde bulunmuş olduk.

Bu örneğimizden de anlaşılacağı gibi seçtiğimiz örneklemelerin boyutuna ve özelliklerine göre anakütle parametrelerinin tahmininde başarı derecemiz artar veya azalır.

Örneklem Dağılımları

Anakütleden aynı boyutlarda rassal olarak örneklemeleri nasıl çektiğimizi inceledik. Nokta tahminlerimizin aynen anakütle parametrelerini tahmin etmelerini bekleyemeyiz. Örneklem ortalaması $\bar{x}$, anakütle parametrelerini tahmin deneyimizin rassal bir çıktısı olarak nitelendirilmelidir. Örneklem ortalaması rassal bir değişkendir. Bu yüzden, $\bar{X}$ 'in kendine ait bit

ortalama değeri, varyans değeri ve olasılık dağılımı bulunmaktadır. $\bar{X}$ 'ın olasılık dağılımı $\bar{X}$ 'nın örneklem dağılımı olarak da adlandırılmaktadır. Benzer bir şekilde, örneklem varyansının olasılık dağılımı, s^2 'in örneklem dağılımı olarak adlandırılırken, örneklem oranının olasılık dağılımı p 'nin örneklem dağılımı olarak adlandırılmaktadır.

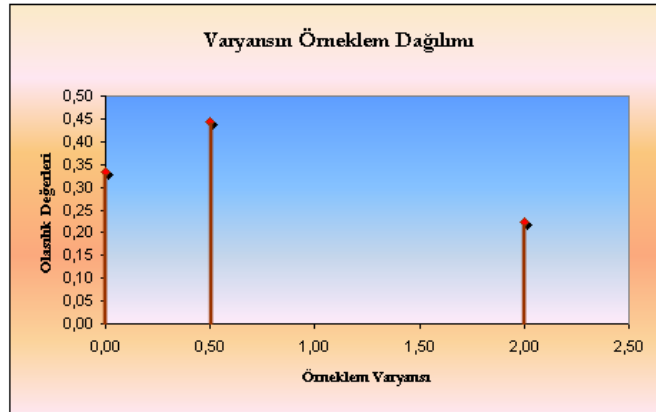
Örnek 7.5: Örnek 7.1'de yeralan $\bar{X}$ 'ın örneklem dağılımından yola çıkarak aynı örneğin örneklem varyansının örneklem dağılımını elde edelim. Olası bütün örneklem varyanslarını aşağıdaki gibi hesaplamıştık:

$$\begin{aligned} s_1^2 &= 0, s_2^2 = 0.5, s_3^2 = 2. \\ s_4^2 &= 0.5, s_5^2 = 0, s_6^2 = 0.5, \\ s_7^2 &= 2, s_8^2 = 0.5, s_9^2 = 0, \end{aligned}$$

Örneklem varyanslarına tekabül eden olasılık değerlerinden oluşan olasılık fonksiyonu aşağıdaki gibidir:

$$P(s^2) = \begin{cases} s^2 = 0 \text{ ise,} & P(s^2) = 3/9, \\ s^2 = 0.5 \text{ ise,} & P(s^2) = 4/9, \\ s^2 = 2 \text{ ise,} & P(s^2) = 2/9, \end{cases}$$

Dolayısıyla örneklem varyansının olasılık dağılımını şekil 7.3'teki gibi elde edebiliriz.

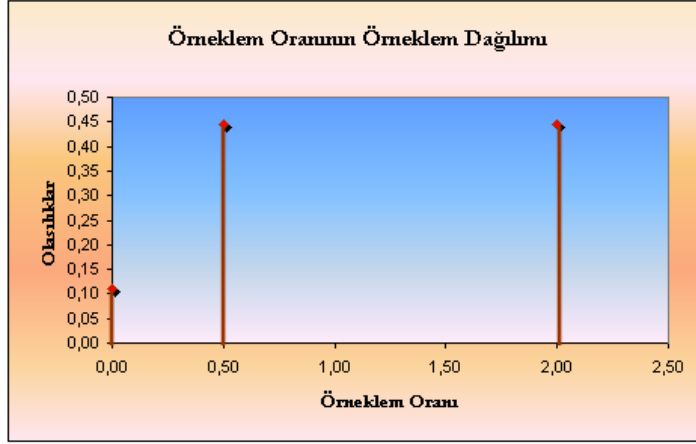


Şekil 7.3 Varyansın örneklem dağılımı

Örnek 7.6: Şimdi ise örnek 7.3'teki örneklem oranlarının örneklem dağılımını oluşturalım. Örneğimizdeki olası örneklem oranları aşağıdaki gibidir:

$$\begin{aligned} p_1 &= 1, p_2 = 1, p_3 = 1/2 \\ p_4 &= 1, p_5 = 1, p_6 = 1/2 \\ p_7 &= 0, p_8 = 1/2, p_9 = 1/2 \end{aligned}$$

Dolayısıyla örneklem oranının p olasılık dağılımını şekil 7.4'teki gibi elde edebiliriz.



Örnek 7.4 Örneklem Oranının örneklem dağılımı

$\bar{X}$ 'in Örneklem Dağılımı

Örneklem dağılımı ve özellikleri, araştırmacılara örneklem ortalamasının $\bar{X}$ anakütle ortalamasına μ ne kadar yakın olduğuna dair olasılık çıkarımları yapma imkanı verir. Her örneklemin ortalama ve varyansı olduğundan, $\bar{X}$ 'in da ortalama ve varyans değeri bulunmaktadır. Şimdi ise $\bar{X}$ 'in örneklem dağılımının ortalama ve varyansını nasıl belirlediğimizi inceleyelim.

Rassal değişkenin beklenen değerini ortalama olarak tanımlamıştık. Dolayısıyla $\bar{X}$ 'in ortalamasını $\mu_{\bar{X}}$ ile ifade edip ortalamasının (7.1)'deki gibi olduğunu söyleyebilirsiniz.

$$E(\bar{X}) = \mu_{\bar{X}} \quad (7.1)$$

Örneklem ortalamasının beklenen değerinin anakütle değerine eşit olduğu ispatlanmıştır. (bkz. EK 7.1)

$$E(\bar{X}) = \mu \quad (7.2)$$

Örneklem ortalamasının beklenen değerinin anakütle ortalamasına eşit olması özelliğine **sapmasızlık** denilmektedir. Denklem (7.1) ve (7.2)'yi bir araya getirecek olursak:

$$E(\bar{X}) = \mu_{\bar{X}} = \mu$$

Elde etmiş olduğumuz bu ifade, örneklem ortalamalarının örneklem dağılımının ortalamasının anakütle ortalamasına eşit olduğunu veya bütün olası $\bar{X}$ değerlerinin ortalamasının, anakütle ortalamasına μ eşit olduğunu göstermektedir.

Örnek 7.7 : Örnek 7.1'de yer alan örneklem ortalamasının beklenen değeri anakütle ortalamasına μ eşit idi.

$$E(\bar{X}) = 1 \times \frac{1}{9} + 1,5 \times \frac{2}{9} + 2 \times \frac{3}{9} + 2,5 \times \frac{2}{9} + 3 \times \frac{1}{9} = 2$$

Peki $\bar{X}$ 'ın örneklem dağılımının varyansı nasıl hesaplayabiliriz ?

Anakütle varyansını bildiğimiz takdirde, örneklem ortalamasının varyansını denklem (7.3)'te olduğu gibi hesaplayabiliriz. Bu formülün ispatını EK 7.2'de inceleyebilirsiniz:

$$\text{var}(\bar{X}) = \sigma_{\bar{X}}^2 = \frac{\sigma^2}{n} \quad (7.3)$$

Burada σ^2 , bilinen anakütle varyansını ifade ederken, n ise örneklem boyutunu ifade etmektedir. Örneklem ortalamasının standart sapması bazı yerlerde standart hata olarak da ifade edildiğinden, bu standart sapma diğer standart sapmalardan farklılık göstermektedir.¹¹

$$\sigma_{\bar{X}} = \sqrt{\frac{\sigma^2}{n}} = \frac{\sigma}{\sqrt{n}} \quad (7.4)$$

Örnek 7.8: Şimdi Örnek 7.1'deki örneklem ortalamalarının varyansını ele alalım. Bütün olası örneklem ortalamalarını ve bu ortalamaların ortalamasını bildiğimizden dolayı, örneklem ortalamasının varyansını aşağıdaki gibi hesaplayabiliriz.

$$\text{var}(\bar{X}) = \sigma_{\bar{X}}^2 = \frac{1}{9}(1-2)^2 + \frac{2}{9}(1.5-2)^2 + \frac{3}{9}(2-2)^2 + \frac{2}{9}(2.5-2)^2 + \frac{3}{9}(3-2)^2 = 0.33$$

Standart sapma değeri ise

$$\sigma_{\bar{X}} = \sqrt{0.33} = 0.57$$

Denklem (7.3)'ten faydalananarak $\sigma_{\bar{X}}^2$ değerini 0.33 olarak hesaplayabiliriz.

$$\sigma_{\bar{X}}^2 = \frac{\sigma^2}{n} = \frac{2}{2} = 0.33$$

Denklem (7.3) yardımıyla hesapladığımız standart sapma değerinin avantajı, elimizde kesin örneklem dağılımının olmadığı durumlarda veya bütün örneklemeleri ele alarak örneklem dağılımını elde etmenin maliyetli olduğu durumlarda ortaya çıkmaktadır.

Burada ilginç olan husus anakütle hakkında fikir sahibi olmadığımız halde anakütle varyansının kullanılmasıdır. Ancak çok nadir olarak böyle durumlarla karşılaşırız ve denklem (7.3)'ü kullanırız. Genelde ise anakütle varyansının tahmin edicisi kullanılmaktadır. Dolayısıyla formülümüz denklem (7.5)'teki gibi olur.

$$s_{\bar{X}}^2 = \frac{s^2}{n} \quad (7.5)$$

$s_{\bar{X}}^2$ değeri de $\sigma_{\bar{X}}^2$ 'in tahmin edicisini temsil etmektedir. Örneklem ortalamasının standart hatasının tahmin edicisi ise aşağıdaki gibidir:

$$s_{\bar{X}} = \frac{s}{\sqrt{n}} \quad (7.6)$$

¹¹ Bu ifadeye aynı zamanda standart hata denilmesinin nedeni, tahmin edicilerin hatalarının hesaplanmasında rol sahibi olmasıdır.

Dolayısıyla Tablo 7.1'e iki tane daha tahmin edici ekleyebiliriz.

Tablo 7.2: Tahmin ediciler

	Anakütle Karakteristiği	Tahmin Edici
Ortalamanın örneklem varyansı	$\sigma_{\bar{x}}^2 = \frac{\sigma^2}{n}$	$s_{\bar{x}}^2 = \frac{s^2}{n}; s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$
Örneklem ortalamasının standart hatası	$\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{n}}$	$s_{\bar{x}} = \frac{s}{\sqrt{n}}; s = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2}$

Örneklem 7.9: Örnek 7.1'deki örneklemlemlere geri dönelim ve $S_6 = \{2, 3\}$ örneklemine ele alalım. Anakütle varyansının tahmin edicisinin formülü tablo 7.1'den de hatırlayacağınız gibi

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

Örneklem ortalamasını $\bar{X}$, 2.5 olarak hesapladığımızdan dolayı, diğer değerleri formülde yerine koyalım.

$$s^2 = \frac{1}{2-1} [(2-2.5)^2 + (3-2.5)^2] = 0.5$$

Sonuç olarak örneklem ortalamasının örneklem dağılımının varyansının tahmin edicisini 0.25 olarak aşağıdaki gibi hesaplayabiliriz.

$$s_{\bar{x}}^2 = \frac{s^2}{n} = \frac{0.5}{2} = 0.25$$

Sonlu Anakütlerde Örneklem Süreci

Örneklem dediğimiz anakütleden örneklem seçme işlemi sonlu bir anakütleden gerçekleştiriliyorsa, $\sigma_{\bar{x}}^2$ değeri denklem (7.7)'deki gibi hesaplanmalıdır. Bu denklemdeki N, sonlu anakütlenin boyutunu ifade etmektedir.

$$\sigma_{\bar{x}}^2 = \frac{N-n}{N-1} \times \frac{\sigma^2}{n} \quad (7.7)$$

$\bar{X}$ 'ın standart sapmasının formülü de denklem (7.8)'deki gibi değişmiştir.

$$\sigma_{\bar{x}} = \sqrt{\frac{N-n}{N-1}} \times \frac{\sigma}{\sqrt{n}} \quad (7.8)$$

$\frac{N-n}{N-1}$ ifadesi sonlu anakütle düzeltme faktörü olarak adlandırılmaktadır. Pratikte, sonlu anakütlelerin boyutları çok fazla büyük olduğundan ve örneklem düzeyleri göreceli olarak bu büyüklük ile karşılaştırıldıklarında küçük kaldıklarından sonlu anakütle düzeltme faktörü $\frac{N-n}{N-1} \approx 1$ değerine yakınsamaktadır. Kitabımızda yer alan bütün örneklerde anakütlenin sonlu olduğu ve boyutunun, örneklem boyutuna göre çok fazla büyük olduğu varsayıldığından, hesaplamalarımızda bu faktörü 1 olarak ele aldık. Bu yüzden hesaplamalarımızda her zaman için denklem (7.3)'ü kullandık.

Merkezi Limit Teoremi

Sonunda çok önemli bir sorunun cevabını vermenin zamanı geldi: Tahmin edicilerin örneklem dağılımı ile neden bu kadar ilgilenmekteyiz? Bu sorunun cevabını verebilmek için spesifik bir tahmin ediciyi, örneklem ortalamasını, ele alalım. Örneklem ortalamasının dağılımı ile neden bu kadar çok ilgilenmekteyiz? Şekil 7.2'de elde edilmiş olan örneklem dağılımı ne ifade etmektedir? Bu soruların cevaplarını örneklerden yararlanarak vermeye çalışalım. Örnek 7.1'de sadece ilk örneklemi ($S_1 = \{1, 1\}$) gözlemleyebildiğimizden dolayı doğal olarak anakütle hakkında bir gözlemlerde bulunamamaktayız. Ancak anakütlenin tahmin edicisi olan örneklem ortalamasını, elimizdeki örneklemde faydalanarak 1 olarak hesaplayabiliriz. Eğer örneklem ortalamasının dağılımı hakkında bilgimiz olsaydı, 1 değerini gözlemleme olasılığımız $1/9$ olacaktı. Şimdi ise sadece üçüncü örneklemi ($S_3 = \{1, 3\}$) gözlemleyebildiğimizi, birinci örneklemi artık gözlemleyemediğimizi düşünelim. Tahmin edilmiş ortalama değeri $2/9$ olasılıkla 1.5 olarak hesaplanır. Peki bu olasılık değerleri ne ifade etmektedir? Öncelikle bu olasılık değerlerini ilerleyen bölümlerde çok fazla kullanacağımızı belirtelim. Bu olasılık değerleri çıkarımsal istatistiğin temellerini oluşturmaktadırlar. Örneğin hesaplanan $1/9$ olasılık değeri anakütle ortalamasının tahmininde **“güvenilmeyecek”** kadar küçük bir olasılık değeridir. “Güvenilmeyecek” kelimesinin ne anlama geldiğini ilerleyen bölümlerde daha detaylı bir şekilde irdedeceğiz.

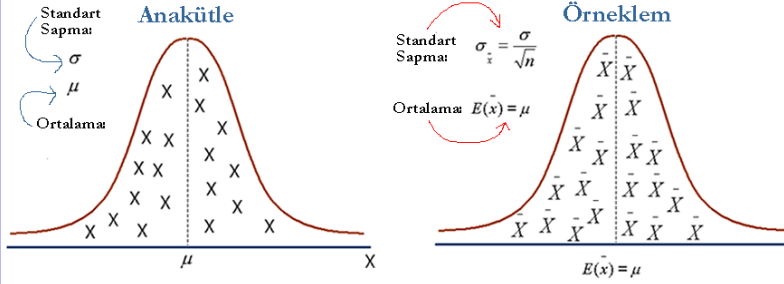
Burada sorulması gereken kritik soru, gerçek hayatta hiçbir zaman anakütleyi gözlemleyememize rağmen, anakütleden türetilmiş olduğumuz örneklem ortalamasının dağılımını nereden bilebileceğimizdir. N boyutuna sahip olan anakütleden seçilebilen olası örneklem miktarı N^n adettir. Daha da önemlisi burada halen neden örneklem dağılımı ile ilgilenmekteyiz? Eğer anakütleyi biliyorsak, herhangi bir parametreyi tahmin etmemize gerek kalmamıştır. Gerçek hayatta elimizde sadece bir örneklem olduğundan bu örneklem ile örneklem dağılımı elde etmek mümkün müdür? **EVET** mümkündür! Belli koşullar altında, **Merkezi Limit Teoremi (MLT)** örneklem ortalamasının örneklem dağılımının normal olarak dağıldığını göstermektedir. Bu koşullar nelerdir?

Merkezi Limit Teoremi'nin koşulları (MLT):

1. Bir anakütlenin dağılımı normal ise, o anakütleden seçilecek örneklemelerin dağılımı da boyutlarından bağımsız olarak kesinlikle normaldir.
2. Anakütlenin dağılımı ne olursa olsun (binom, yeknesak vs.), örneklem boyutu yeterince büyükse ($n > 30$), örneklem dağılımı yine de normal dağılıma yakınsar.

$\bar{X}$ 'in Örneklem Dağılımı

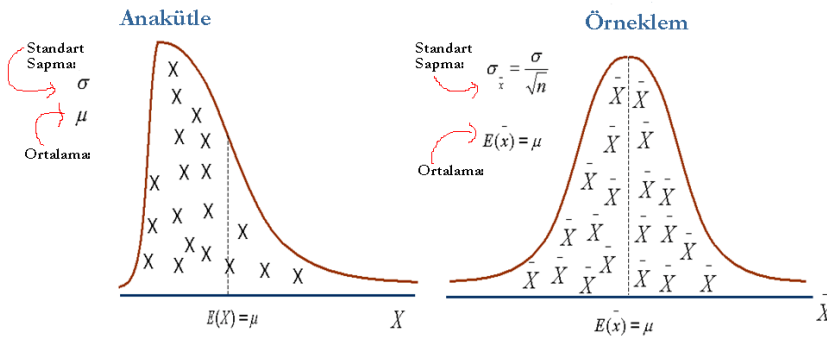
Anakütle normal olarak dağılıyorsa, örneklem ortalaması da örneklem boyutundan bağımsız olarak normal dağılır.



Comment [YZ1]: Şekil nolarını update et

$\bar{X}$ 'in Örneklem Dağılımı

Anakütle normal olarak dağılmıyorsa, örneklem ortalaması da örneklem boyutu $n > 30$ olduğu takdirde normal dağılır.



Merkezi Limit Teoreminin gücü Örnek 7.3'ü tekrar incelediğimizde daha kolay anlaşılabilir. Bu örnekte anakütlenin normal olarak dağılmadığı gözükmemektedir. Gerçekte de, Anakütle = {1, 2, 3}'ün 1/3 olasılıkla yeknesak olarak dağıldığı gözükmemektedir. Örneklem boyutu 30 dan hayli uzak bir değer olan ikidir. Bu yüzden şekil 7.2'deki örneklem dağılımı normal dağılıma benzememektedir.

Bu yüzden, MLT'ye göre eğer n sayısı 30 dan büyük ise, biz örneklem dağılımının ana dağılımdan (anakütlenin dağılımından) bağımsız olarak normal dağıldığını söyleyebiliriz. Daha önceki bölümlerde de bahsettiğimiz üzere normal dağılım, ortalama ve standart sapma gibi iki parametreye bağlıdır. Eğer örneklem dağılımının ortalamasını ve standart sapmasını biliyorsak, örneklem ortalamasını gözleme olasılığını hesaplayabiliriz. Bu olasılık değerini hesapladıktan sonra, bu tahmin edicinin anakütle ortalamasını tahmin etmek için iyi bir parametre olup olmadığını test etmiş oluruz.

Bu işlemi gerçekleştirebilmek için öncelikle, örneklem dağılımının ortalama ve standart sapmasını hesaplamamız gerekmektedir. Örneklem ortalamasının örneklem dağılımının ortalama ve standart sapmasının formüllerini daha önce türetmiştik. Hatırlanacağı gibi örneklem ortalamasının ($\bar{x}$) örneklem dağılımının ortalamasının formülü denklem (7.2)'deki gibidir.

$$E(\bar{X}) = \mu$$

Örneklem ortalamasının varyansının formülünü ise denklem (7.3)'teki gibi hesaplayabiliriz:

$$\text{var}(\bar{X}) = \sigma_{\bar{X}}^2 = \frac{\sigma^2}{n}$$

Varyans hesaplanırken kullandığımız bu formülü σ^2 anakütlenin varyansını bildiğimiz takdirde kullanabiliriz. σ^2 anakütle varyansını bilmediğimiz durumlarda ise örneklem ortalamasının dağılımının varyansını aşağıdaki gibi hesaplayabilirsiniz.

$$s_{\bar{X}}^2 = \frac{s^2}{n}; \quad s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$$

Bu durumda anakütle varyansının σ^2 yerine, tahmin edicisi olan s^2 örneklem varyansı kullanılmaktadır. Örnek olarak Türkiye'deki bebeklerin konuşma yaşını araştıracağımızı düşünelim. Ortalama yaşı tahmin edebilmek için bir örneklem ele alalım. Aldığımız örneklemde bebeklerin konuşma yaşı ortalama olarak 1.2 ($\bar{x}=1.2$) olarak hesaplanmıştır. Örneklem ortalamalarının dağılımının varyansını da hesaplayıp örneklemin güvenilirliği hakkında bir fikir edindikten sonra anakütle ortalaması hakkında fikir yürütmeye çalışabiliriz. Anakütle ortalamasını 1.5 ($\mu=1.5$) olarak ele alalım. Ortalamanın 1.5 olduğunu kabul ederek örneklem ortalamasını 1.2 olarak elde etme olasılığını hesaplayabiliriz. Böylelikle örneklemimizin güvenilirliği hakkında fikir yürütebiliriz. Dolayısıyla bu olasılık değerini hesaplamakla hipotezimizin güvenilirliği hakkında da fikir yürütebiliriz. Ele aldığımız olasılık değeri küçük bir değer ise, hipotezimizi reddederiz aksi takdirde reddedemeyiz. İzlemiş olduğumuz bu mantıksal süreç, ilerleyen bölümlerdeki hipotez testinin arkasında yatan mantığı ifade edecektir.

p'nin örneklem dağılımı

Bu bölümde, p rassal değişkenine ait olan ortalama ve varyans değerlerinin genel formülleri ele alacağız. Örneklem oranının ortalamasını ve varyansını türeteceğiz. p değişkenin beklenen değeri aşağıdaki gibidir:

$$E(p) = \pi$$

Örneklem ortalaması ile aynı doğrultuda, örneklem oranının beklenen değeri anakütle oranına π eşit olmaktadır. Bu eşitliğin nereden geldiği EK 7.3'te gösterilmiştir.

Örnek 7.10: Örnek 7.3'teki örneklem oranının (p'nin) beklenen değerinin anakütle oranına eşit olduğunu aşağıdaki gibi gösterebiliriz.

$$E(p) = 0 \times \frac{1}{9} + 0.5 \times \frac{4}{9} + 1 \times \frac{4}{9} = \frac{2}{3}$$

p değişkeninin varyansı denklem (7.9) 'daki gibi tanımlanmaktadır

$$\text{var}(p) = \sigma_p^2 = \frac{\pi(1-\pi)}{n} \quad (7.9)$$

Doğal olarak p değişkeninin standart hatası: (örneklem oranının dağılımının standart sapması)

$$\sigma_p = \sqrt{\frac{\pi(1-\pi)}{n}} \quad (7.10)$$

gibidir.

Örnek 7.11: Şimdi ise p'nin varyansını genel varyans formülünü kullanarak hesaplayalım.

$$\text{var}(p) = \sigma_p^2 = (0 - 2/3)^2 \times \frac{1}{9} + (1/2 - 2/3)^2 \times \frac{4}{9} + (1 - 2/3)^2 \times \frac{4}{9} = \frac{1}{9}$$

denklem (7.9)'daki varyans formülünden faydalanarak p'nin varyansını hesapladığımızda genel varyans formülü ile aynı sonucu elde ettiğimizi görürüz.

$$\text{var}(p) = \sigma_p^2 = \frac{\pi(1-\pi)}{n} = \frac{2/3(1-2/3)}{2} = \frac{1}{9}$$

Anakütle bilinmediği halde anakütlenin oranı hakkında yorum yapmamız aslında pratikte çok zor bir husustur. Dolayısıyla, ortalama durumunda olduğu gibi, elimizdeki örneklemini kullanarak anakütle oranının tahmin etmeye çalışarak kullanırız. Bu yüzden hesaplamalarımızda daha gerçekçi olmak için denklem (7.9) yerinde denklem (7.11)'i kullanabiliriz.

$$s_p^2 = \frac{p(1-p)}{n} \quad (7.11)$$

Burada s_p^2 , σ_p^2 'nin tahmin edicisi rolünü üstlenmektedir. Standart hatanın tahmin edicisi ise aşağıdaki gibidir.

$$s_p = \sqrt{\frac{p(1-p)}{n}} \quad (7.12)$$

Tablo 7.2'deki tahmin edicilerimizin listesine iki adet daha tahmin edici ekleyebiliriz.

Tablo 7.3: Tahmin Ediciler

	Anakütle Karakteristiği	Tahmin Edici
Örneklem oranının varyansı	$\sigma_p^2 = \frac{\pi(1-\pi)}{n}$	$s_p^2 = \frac{p(1-p)}{n}; p = \frac{X}{n}$
Örneklem oranının standart hatası	$\sigma_p = \sqrt{\frac{\pi(1-\pi)}{n}}$	$s_p = \sqrt{\frac{p(1-p)}{n}}; p = \frac{X}{n}$

Örnek 7.12: Örnek 7.3'teki örneklemelerden faydalanarak $S = \{x, y\}$ 'yi ele alalım. Örneklem oranı $p = 1/2$ ve σ_p^2 'nin tahmin edicisi aşağıdaki gibi hesaplanabilir.

$$s_p^2 = \frac{p(1-p)}{n} = \frac{1/2(1-1/2)}{2} = \frac{1}{8}$$

Bu değer sadece bir tahmin edici olduğundan aşikar bir şekilde gerçek varyans değeri olan $2/3$ değerinden farklı bir değere sahiptir.

Sonlu Anakütlelerde Örneklem Süreci

Örneklem dediğimiz işlemimiz sonlu bir anakütleden gerçekleştiriliyorsa, σ_p^2 değeri aşağıdaki gibi hesaplanmalıdır.

$$\sigma_p = \sqrt{\frac{N-n}{N-1}} \times \sqrt{\frac{\pi(1-\pi)}{n}} \quad (7.13)$$

Daha öncede bahsetmiş olduğumuz anakütle düzeltme faktörü sayesinde denklem (7.9) ve (7.13) arasındaki fark oluşmuştur. Yine genel olarak, anakütle boyutunun örneklem boyutuna göre çok daha fazla büyük olduğunu varsaydığımızdan dolayı uygulamalarımızın ve örneklerimizin büyük bir kısmında denklem (7.9)'u kullanacağız.

$p = \frac{X}{n}$ değişkeni binom dağılımına sahiptir. MLT sayesinde p değişkeninin örneklem dağılımı, örneklem boyutu yeterince büyük olduğu takdirde, normal olasılık dağılımına yakınsayabilmektedir. p değişkeninin örneklem boyutunun yeterince büyük olması için aşağıdaki iki şartın sağlanması gerekmektedir.

$$n \times \pi \geq 5 \text{ ve } n \times (1 - \pi) \geq 5$$

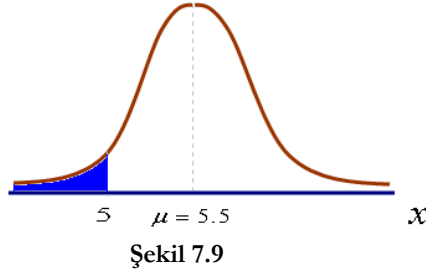
Örneklem Dağılımlarının Kullanımı:

Hipotez Testine Giriş

Örneklem dağılımı kavramı, istatistikte çok fazla önem teşkil eden bir kavramdır.

Bu bölümde, örneklem dağılımlarının istatistikteki kullanımlarını ve hipotez testine nasıl temel teşkil ettiğini örneklerle inceleyeceğiz. Örneklem dağılımlarının kullanımına geçmeden önce, normal dağılan değişkenlerin olasılık değerlerini nasıl hesapladığımızı hatırlayalım.

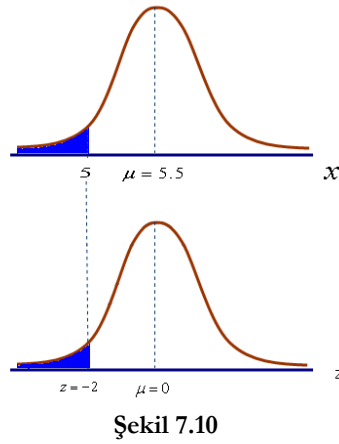
Örnek 7.13: Bir çay üreticisi firmanın poşet çay üretiminde, üretilen poşet çaylarda ortalama 5.5 gram çay bulunmaktadır. Çay miktarlarının standart sapması da 0.25'tir. (bu standartlar şirket tarafından belirlenmiştir) Poşetlerdeki çay miktarlarının normal olarak dağıldığını varsayalım. Bir bardak çay içebilmek için kullandığımız çay poşetindeki çay miktarının 5 gramdan az olma olasılığı nedir? Örneğimizde ele aldığımız X değişkeni (bir çay poşetindeki çay miktarının ağırlığını göstermektedir) $\mu = 5.5$ ortalaması ve $\sigma = 0.25$ standart sapması ile normal olarak dağılmaktadır. Anakütlemiz, çay üretim firmasında üretilen bütün çay poşetlerinden oluşmaktadır ve normal olarak dağılmaktadır. Bizden istenilen olasılık değerini grafiksel olarak göstererek olursak:



Şekil 7.9'daki taralı alan çay poşetlerinde 5 gramdan daha az çay miktarı olma olasılığını göstermektedir. Bu olasılık değerini nasıl hesaplayacağız? Daha önceden de öğrendiğimiz üzere $x = 5$ değerini standart normal dağılım değişkeni olan z değişkenine dönüştürmemiz gerekmektedir (ortalaması 0 ve standart sapması 1 olan standart normal dağılıma)

$$z = \frac{x - \mu}{\sigma} = \frac{5 - 5.5}{0.25} = -2$$

Bölüm 6'dan da hatırlanacağı üzere, $P(X < 5) = P(X < -2)$. Grafikselsel olarak ifade edecek olursak:



Normal dağılım tablosunu kullanarak (veya MINITAB kullanılarak) $P(Z < -2)$ değerini $P(Z < -2) = 0.0228$ olarak hesaplayabilirsiniz. Dolayısıyla, 5 gramdan daha az çay miktarı elde etme olasılığı yaklaşık olarak yüzde 3'tür.

Şimdi ise örneklem dağılımı ile alakalı örneğimizi inceleyelim.

Örnek 7.14: Üretilen çay poşetlerinin kalite kontrol işlemleri ile ilgilendiğimizi düşünelim. Bu bağlamda bir çay poşetindeki çay miktarı artık bizim için çok önemlidir. Çay poşetleri normal standartların altında doldurulduğu takdirde iki sorun ile karşılaşmaktadır.

Birincisi, müşteriler çayın tadını istedikleri kıvamda hissedemeyebilirler. İkinci sorun ise, etiket üzerinde belirtilen gramaja uyulmadığında etiket bildiriminde doğruluk kanunu çiğnenmiş olmaktadır. Bu örneğimizde, çay poşetlerinin etiketlerinde ortalama olarak 5.5 gram çay bulunduğu ve bu miktarın 0.25 gram standart sapmaya sahip olduğu belirtilmektedir.

Bir başka ifade ile, eğer üretilen çay miktarı etikette belirtilen ortalama standardı aşarsa, üretimden vazgeçilecektir. Çay miktarını da tam olarak belli bir standartta üretmenin ciddi zorlukları olduğu da aşıkardır. Sıcaklık, nem, farklı çay yoğunlukları bu zorlukların başını çekmektedirler.

Şimdi bu çay poşetlerinin kalite kontrolünden biz sorumlu olduğumuza göre, firma bizden üretilen çay miktarının iddia edilen ortalama gramajda olmasını temin etmemizi beklemektedir.

Anakütle ortalamasının (üretilen ortalama çay miktarının) 5.5 gram olup olmadığını anlamak için öncelikle uygun bir örneklem seçip bu örneklem ortalamasını hesaplamamız gerekmektedir.

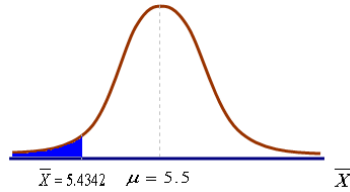
Bir saat içerisinde üretilen çay poşetlerinden 40 tanesinin örneklem olarak seçildiğinde, bu örneklemde bulunan ortalama çay miktarının 5.4342 gram olduğu gözlenmiştir. Bu bağlamda genel üretimde standartların sağlandığını, yani ortalama olarak $\mu = 5.5$ gram üretim yapıldığını söyleyebilir miyiz?

Aldığımız örneklem bilgilerini biraz daha istatistiksel sembollerle ifade edecek olursak, örneklem boyutu $n=40$ ve örneklem ortalaması $\bar{x} = 5.4342$ olarak hesaplanmıştır. Örneklem dağılımını bildiğimiz takdirde, ($\bar{x}$ 'ların olasılık dağılımı) bu örneklemi elde etme olasılığımızı hesaplayabiliriz.

$\bar{X}$ 'in örneklem dağılımı hakkında fikir yürütebilmek için MLT'den faydalanabiliriz. MLT'ye göre örneklem boyutu 30'dan büyük olduğu durumlarda, anakütlenin dağılımından bağımsız olarak, $\bar{X}$ 'in örneklem dağılımının normal olduğunu söyleyebiliriz. n örneklem boyutu 40 olduğundan bizim örneğimizde de $\bar{X}$ 'in örneklem dağılımının normal olduğunu söyleyebiliriz. MLT aynı zamanda bize her zaman için anakütle normal dağıldığı takdirde, örneklem boyutu ne olursa olsun, örneklem dağılımının da normal olacağını garanti etmektedir.

$\bar{X}$ 'in normal olarak dağıldığını belirledikten sonra daha önce öğrendiklerimizden ($\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{n}}$) faydalananarak $\bar{X}$ 'in standart sapmasını $\sigma_{\bar{x}} = \frac{0.25}{\sqrt{40}} = 0.04$ olarak hesaplayabiliriz.

$\bar{X}$ 'in ortalaması da, bilmediğimiz ama üreticinin $\mu = 5.5$ olarak iddia ettiği anakütle ortalamasına eşittir. Elimizdeki verilerden faydalananarak üreticinin bu iddiasının doğruluğunu araştırmaya çalışacağız. Dolayısıyla, $\bar{X}$ hipotetik ortalaması $\mu = 5.5$ ve $\sigma_{\bar{x}} = 0.04$ standart sapması ile normal olarak dağıldığından, $\bar{x} = 5.4342$ örneklem ortalamasını elde etme olasılığımızı hesaplayabiliriz. Bu olasılık değerini de $P(\bar{X} < 5.4342) = ?$ gibi ifade edebiliriz. Şimdi sorumuz bir önceki örneğimizdeki olasılık hesabına benzer bir hal almıştır. Tek önemli fark, bir önceki örneğimizde X (çay poşetlerinin ağırlığı) değişkenini ele almıştık, şimdi ise $\bar{X}$ (çay poşetlerinin ortalama ağırlığı) değişkeni ile çalışmaktayız. Bu olasılığı grafiksel olarak gösterecek olursak:

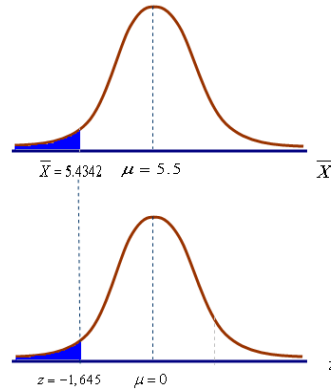


Şekil 7.11

Şekil 7.11’de taralı alan, ortalama çay miktarının 5.4342 gramdan az olma veya $P(\bar{X} < 5.4342)$ olasılığını göstermektedir. Benzer şekilde, bu olasılık değerini hesaplamak için öncelikle 5.4342 değerini z değerine dönüştürmemiz gerekmektedir¹².

$$z = \frac{\bar{x} - \mu}{\sigma_{\bar{x}}} = \frac{5.4342 - 5.5}{0.04} = -1.645$$

$P(\bar{X} < 5.4342) = P(Z < -1.645)$ olduğundan Normal Dağılım Tablosu’nu kullanabiliriz. Grafiksel olarak,

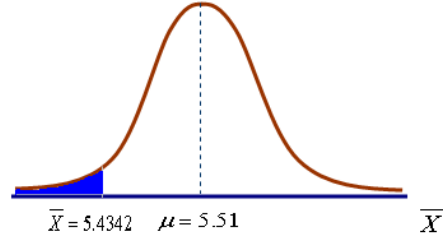


Şekil 7.12

Normal Dağılım Tablosu’nu kullandığımızda (veya MINITAB) kullandığımızda $P(Z < -1.645) = 0.05$ değerini elde edebiliriz. Bu olasılık değeri biraz düşük bir olasılık değeridir. Düşük bir olasılık değeri olması bize ne ifade etmektedir? Bu olasılık değeri, $\mu = 5.5$ olduğu takdirde $\bar{X} = 5.4342$ değere sahip bir örneklem elde etme olasılığımızın yüzde 5 olduğunu belirtmektedir. Eğer, anakütle ortalaması 5.5 değildir dediğimiz takdirde, (gerçekte anakütle ortalamasının 5.5 olduğu durumda) yüzde 5 olasılıkla hata yapmak ihtimalimiz bulunmaktadır. Bu olasılık değeri, gerçek anakütle ortalaması 5.5 olsa bile 5.5 değildir diyerek hata yapmamız için küçük bir olasılık değeridir. Dolayısıyla bu olasılık değeri, gerçek ortalama değeri 5.5 değerine eşit olamaz sonucuna varmak için düşük bir hata değeridir.

Yukarıdaki paragraftan anakütle ortalamasının 5.5 olamayacağı sonucuna varırız. Peki anakütle ortalaması 5.5’ten büyük olabilir mi? Mesela 5.51 olabilir mi? Şimdi ise hipotetik ortalaması $\mu = 5.51$ ve $\sigma_{\bar{x}} = 0.04$ standart sapması ile normal olarak dağılan, $\bar{x} = 5.4342$ örneklem ortalamasını elde etme olasılığımızı hesaplamamız gerekmektedir. Bu olasılık değerini grafiksel olarak gösterecek olursak:

¹² Bir önceki örneğimizde z’ye dönüşüm formülünü $z = (x - \mu) / \sigma$ olarak ele almıştık. Fakat burada, normal olarak dağılan değişkenimiz örneklem ortalaması ($\bar{x}$) olduğundan, dönüşüm formülü $z = (\bar{x} - \mu) / \sigma / \sqrt{n}$ halini almaktadır ve bu formülde (μ) yine anakütle ortalamasını ifade ederken $(\sigma / \sqrt{n})$ standart hatayı ifade etmektedir.

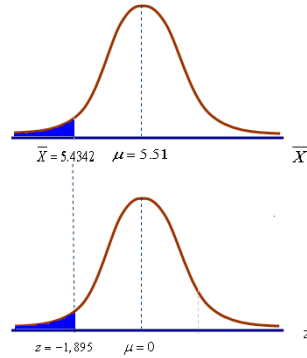


Şekil 7.13

Şekil 7.13'te taralı alan, hipotetik ortalama $\mu = 5.51$ iken ortalama çay miktarının 5.4342 gramdan az olma veya $P(\bar{X} < 5.4342)$ olasılığını göstermektedir. Bu olasılık değerini hesaplamak için öncelikle 5.4342 değerini z değerine dönüştürmemiz gerekmektedir.

$$z = \frac{\bar{x} - \mu}{\sigma_{\bar{x}}} = \frac{5.4342 - 5.51}{0.04} = -1.895$$

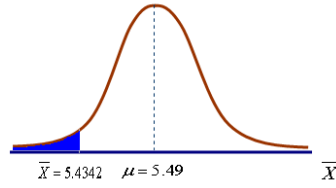
$P(\bar{X} < 5.4342) = P(Z < -1.895)$ olduğundan Normal Dağılım Tablosu'nu kullanabiliriz. Grafiksel olarak,



Şekil 7.14

Normal Dağılım Tablosu'nu kullandığımızda (veya MINITAB) kullandığımızda $P(Z < -1.895) = 0.029$ değerini elde edebiliriz. Bu olasılık değeri 0.05 değerinden küçük bir değerdir. Dolayısıyla aynı anolojiden yola çıkarak eğer 0.029 olasılıkla bir hata yapmayı kabul ediyorsak, gerçek anakütle ortalamasının 5.51 değerine eşit olamayacağı sonucuna varırız. Aynı nedenden dolayı herhangi bir hipotetik ortalama değerinin 5.5'ten büyük olamayacağı sonucuna varabiliriz. **Sonuç olarak, 0.05 veya daha az bir olasılık ile hata yapmayı kabul ediyorsak, anakütle ortalaması 5.5 değerine eşit veya bu değerden büyük olamaz sonucuna varabiliriz.**

Şimdi ise hipotetik ortalamamızın 5.5'ten küçük olma ihtimalini ele aldığımızda neler olabileceğini inceleyelim. Bu durumda ise hipotetik ortalaması $\mu = 5.49$ ve $\sigma_{\bar{x}} = 0.04$ standart sapması ile normal olarak dağılan, $\bar{x} = 5.4342$ örneklem ortalamasını elde etme olasılığımızı hesaplamamız gerekmektedir. Bu olasılık değerini grafiksel olarak gösterecek olursak:

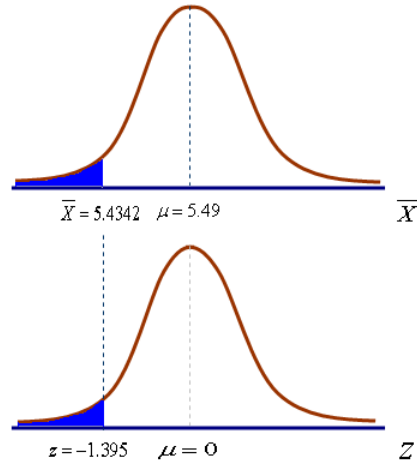


Şekil 7.15

Şekil 7.13'te taralı alan, hipotetik ortalama $\mu = 5.51$ iken ortalama çay miktarının 5.4342 gramdan az olma veya $P(\bar{X} < 5.4342)$ olasılığını göstermektedir. Bu olasılık değerini hesaplamak için öncelikle 5.4342 değerini z değerine dönüştürmemiz gerekmektedir.

$$z = \frac{\bar{x} - \mu}{\sigma_{\bar{x}}} = \frac{5.4342 - 5.49}{0.04} = -1.395$$

$P(\bar{X} < 5.4342) = P(Z < -1.395)$ olduğundan Normal Dağılım Tablosu'nu kullanabiliriz. Grafiksel olarak,



Şekil 7.16

Normal Dağılım Tablosu'nu kullandığımızda (veya MINITAB) kullandığımızda $P(Z < -1.395) = 0.0823$ değerini elde edebiliriz. Bu olasılık değeri 0.05 değerinden büyük bir değerdir. Dolayısıyla eğer 0.0823 olasılıkla bir hata yapmayı kabul ediyorsak, gerçek anakütle ortalamasının 5.49 değerine eşit olamayacağı sonucuna varırız. Eğer, yüzde 5 değerini (anakütle ortalaması doğru olduğu takdirde) yapabileceğimiz maksimum hata miktarı olarak ele alacak olursak, 5.49 değerinin doğru ortalama değeri olabileceği sonucuna varırız.

Sonuç olarak 0.05 değerini, hipotetik anakütle değerini doğru olduğu takdirde reddederek maksimum hata yapma olasılığımız olarak görüyorsak, anakütle ortalamasının 5.5 değerine eşit veya bu değerden küçük olabileceği sonucuna varabiliriz.

Şuana kadar yapmış olduğumuz analizler bir sonraki bölümde anlatılacak olan “*Hipotez Testi*” konusuna temel teşkil etmektedir. Elde ettiğimiz sonuçları özetleyerek hipotez testinin standart formülasyonunu elde etmeye çalışalım. Analizimizin başında iki farklı formda hipotez oluşturmuştuk:

- **Hipotez 1:** Anakütle ortalaması 5.5 değerine eşittir veya bu değerden büyüktür, matematiksel olarak ifade edecek olursak $\mu \geq 5.5$
- **Hipotez 2:** Anakütle ortalaması 5.5 değerinden küçüktür, matematiksel olarak ifade edecek olursak $\mu < 5.5$.

Açık bir şekilde burada iki hipotezin mücadelesi gözükmemektedir. Eğer birini reddedersek, diğerini kabul etmemiz (veya teknik ifade ile “*reddedemeyiz*”) gerekmektedir. Hipotez testi literatüründe birinci temel hipotezimiz “*Sıfır Hipotezi*” olarak adlandırılmakta ve $H_0 : \mu \geq 5.5$ şeklinde gösterilmektedir. İkinci hipotezimiz ise “*Alternatif Hipotez*” olarak adlandırılmaktadır ve $H_A : \mu < 5.5$ şeklinde gösterilmektedir.

Örneğimizde eğer sıfır hipotezini reddedersek, yani anakütle ortalamasının 5.5 değerine eşit veya daha büyük olduğunu reddedersek maksimum yüzde 5 olasılıkla hata yapmış oluruz. Araştırmacılar tarafından belirlenen maksimum hata yapma olasılığına testin “*anlamlılık düzeyi*” denilmektedir. Örneğimizde testin anlamlılık düzeyi yüzde 5 olarak alınmıştır. Anlamlılık düzeyini değiştirdiğimizde testin sonucunu da değiştirmektedir. Bu değişimin nasıl olduğuna bir sonraki bölümde değineceğiz.

Minitab Application

Örnek 7.15: Daha önce ele aldığımız çay poşetleri örneğimize geri dönüp 40 çay poşetinden oluşan yeni bir örneklem seçelim. Örneklem verileri *çaypoşetleri.mtw* adlı MINITAB dosyasında yer almaktadır. Bu dosyayı açtıktan sonra Bölüm 5’te yaptığımız gibi “descriptive statistics” (betimleyici istatistikleri) kullanarak aşağıdaki çıktı ekranını elde edebiliriz.

Results for: teabag.MTW						
Descriptive Statistics: C1						
Variable	N	Mean	Median	TrMean	StDev	SE Mean
C1	40	5,4682	5,4500	5,4625	0,2856	0,0452
Variable	Minimum	Maximum	Q1	Q3		
C1	4,9500	6,1000	5,2200	5,6225		

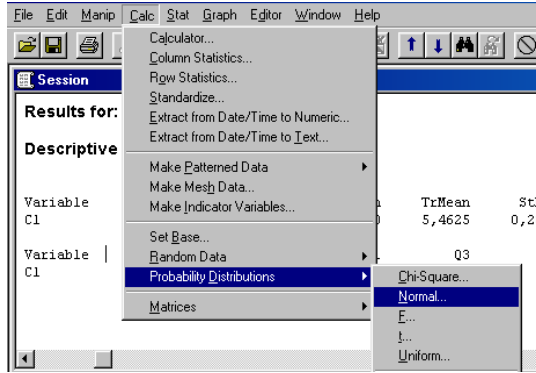
Şekil 7.17

Bu sonuçtan da takip edebileceğimiz gibi örneklem ortalaması “*Mean*” (Ortalama) 5.4682’ye eşittir. “*StDev*” (Standart sapma) ise 40 poşetin standart sapma tahminidir. Daha önceki notasyonda bu standart sapmayı s , ($=\sqrt{s^2}$) terimi ile göstermiştik (anakütle standart sapmasının $\sigma (= \sqrt{\sigma^2})$, bir tahmini). “*SE Mean*” ise the örneklem ortalamasının standart sapmasının tahminidir. Daha önceki notasyonda bunu $s_{\bar{x}} = \sqrt{\frac{s^2}{n}} = \frac{s}{\sqrt{n}}$ terimi ile göstermiştik.

Eğer isterseniz MINITAB'ın hesaplamalarını, $s_{\bar{x}} = \frac{0.2856}{\sqrt{40}}$ işlemini yaparak kontrol edebilirsiniz.

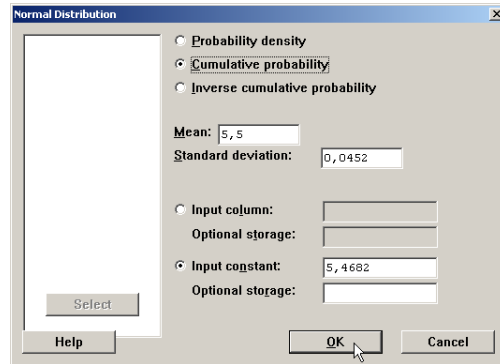
Gerçekten de bu işlem çıktıda gösterilen 0.0452 sonucunu verecektir.

Anakütle ortalaması 5.5. iken MINITAB ortamında 5.4682 değerini elde etme olasılığımızı hesaplayabilmek için “Calc” (hesapla) menüsünden “normal probability distribution” (normal olasılık dağılımına) tıklayınız.



Şekil 7.18

Açılan diyalog ekranından “cumulative probability” seçeneğini işaretleyerek, ortalama ve standart sapma olarak 5.5 ve 0.0452 değerlerini, “input constant” bölümüne ise 5.4682 değerini giriniz.



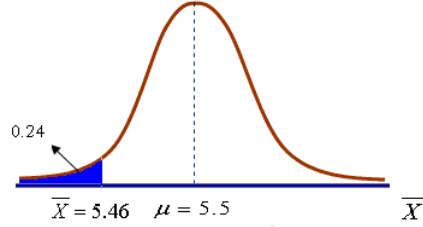
Şekil 7.19

OK tuşuna bastığınızda 5.4682 veya 5.4682'den küçük bir değeri gözlemlene olasılığını elde etmiş oluruz.

Cumulative Distribution Function		
Normal with mean = 5,50000 and standard deviation = 0,0452000		
x	P (X <= x)	
5,4682	0,2409	

Şekil 7.20

5.4682 veya 5.4682'den küçük bir değeri gözlemleme olasılığını 0.24 olarak hesapladık. Bu değeri grafiksel olarak gösterecek olursak



Şekil 7.21

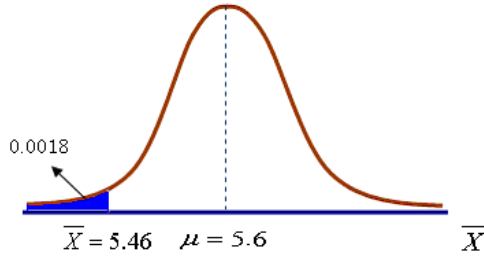
0.24 olasılık değeri, belirlemiş olduğumuz anlamlılık düzeyi olan 0.1 değerinden büyük olduğundan, $\mu \geq 5.5$ hipotezini reddedemeyiz.

Şimdi ise anakütle ortalamasının 5.6 veya daha fazla olduğu hipotezini $\mu \geq 5.6$ ele alalım ve aynı hesaplamaları tekrar ederek Şekil 7.22'deki çıktıyı elde edebiliriz.

Cumulative Distribution Function		
Normal with mean = 5,60000 and standard deviation = 0,0452000		
x	P (X <= x)	
5,4682	0,0018	

Şekil 7.22

Veya grafiksel olarak



Şekil 7.23

0.0018 olasılık değeri anlamlılık düzeyi olan 0.1 değerinden küçük olduğundan, $\mu \geq 5.6$ hipotezi reddedilmektedir.

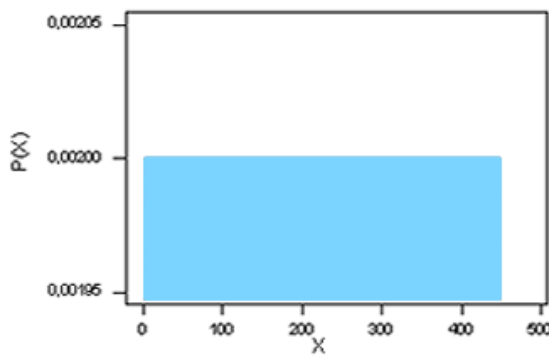
Minitab Uygulamaları

Örnek 7.3'te anakütlenin çok basit bir yapıya sahip olduğunu varsaymıştık. ($P = \{1, 2, 3\}$) Bu basit yapı sayesinde bütün olası örneklemeleri elde edebiliriz. Dahası, bütün olası örneklemeleri bildiğimizden dolayı, örneklem ortalamalarının olasılık dağılımı hakkında da fikir yürütebiliriz.

Eğer anakütle boyutunu genişletecek olursak örneklem ortalamaları ve örneklem dağılımı hakkında fikir yürütmek zordur. Alternatif bir metot olarak; simulasyon adı verilen bu metoda göre belirli bir anakütleden rassal olarak örneklemeler çekilir. Bu bölümdeki yaklaşımı MINITAB yardımıyla Merkezi Limit Teoreminine uygulayalım.

Örnek 7.3 'teki anakütle verisine benzer olarak, tekrar burada anakütlede bulunan nesnelerin numaralandırıldığını varsayıyoruz. Bunun yanında bu sefer 3 nesne yerine 450 nesne olduğunu varsayıyoruz. Bu yüzden daha önce de olduğu gibi anakütleyi 1 den 450 ye kadar olan rakamlar ile ifade edebiliriz. Bu veriler i “**anakütle.mtw**” isimli dosya içerisinde bulabilirsiniz.

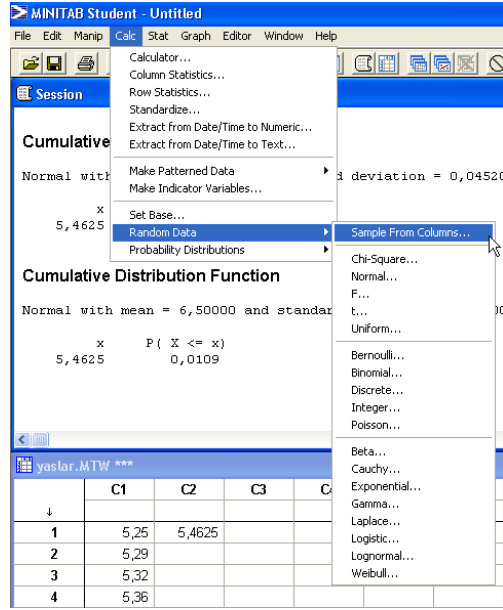
Bu anakütle Örnek 7.3'teki gibi yeknesak bir dağılıma sahiptir. Bu yüzden her sayıya ait olan olasılık değeri 1/450dir. Anakütlenin olasılık yoğunluğu ise (Şekil 7.1e benzer bir şekilde) Şekil 7.24'te gösterilmektedir:



Şekil 7.24

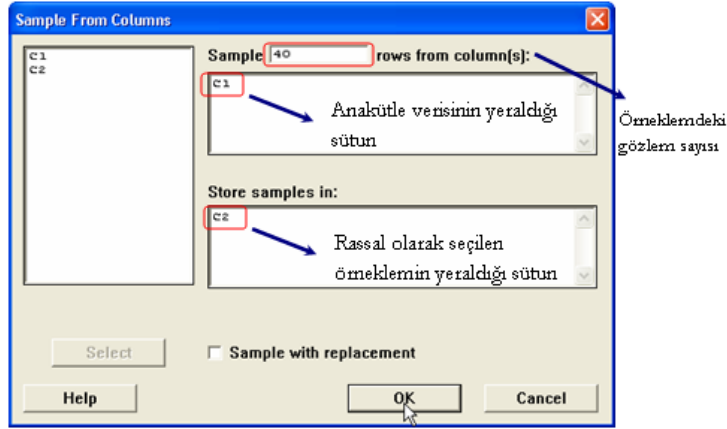
Şimdi ise örneklem boyutu olan n değerini belirlemeliyiz. Örneklem boyutumuz örnek 7.3'te $n = 2$ idi. Şimdi ise 2 den daha büyük bir örneklem seti çekebiliriz. MLT'i uygulamak için $n = 40$, olarak düşünelim ki 30 dan büyük olduğu için örneklem ortalamalarının da dağılımının normal olduğunu belirleyelim. Hatırlanacağı üzere n değeri 30 dan büyük olduğunda anakütlenin dağılımı yeknesak bile olsa örneklem ortalamalarınormal olarak dağılır. Örnek 7.3'teki metodolojiyi burada uygulamak anakütleden bütün örneklemeleri çekmenin zorluğundan dolayı imkansızdır. Çünkü burada karşımıza çıkacak olan olası örneklemelerin sayısı, anakütle boyutu 450, örneklem boyutu 40 iken $450^2 = 202500$.olarak karşımıza çıkar. Bu yüzden bütün olası örneklemelerin normal dağılıp dağılmadığını gözlemlemek olanaksızdır. Dolayısıyla MINITAB'I kullanarak simülasyona dayalı bir deney oluşturabiliriz.

Bu deneyi düzenlemek için MINITAB'tan “**anakutle.mtw**” dosyasını açınız. Şimdi “**C1**” sütununda yer alan verilerden, “**Calc**” menüsünden “**Random Data**” ya gelerek oradan da “**Sample From Columns**”ı seçerek rassal örneklem oluşturabilirsiniz.



Şekil 7.25

Şekil 7.26'daki gibi bir diyalog kutusu karşınıza gelecektir.

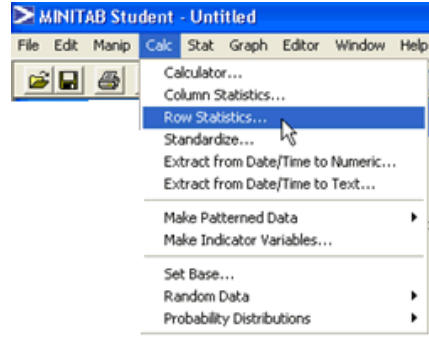


Şekil 7.26

Örneklem boyutunu ($n = 40$), olarak giriniz ve "OK" tuşuna basınız.

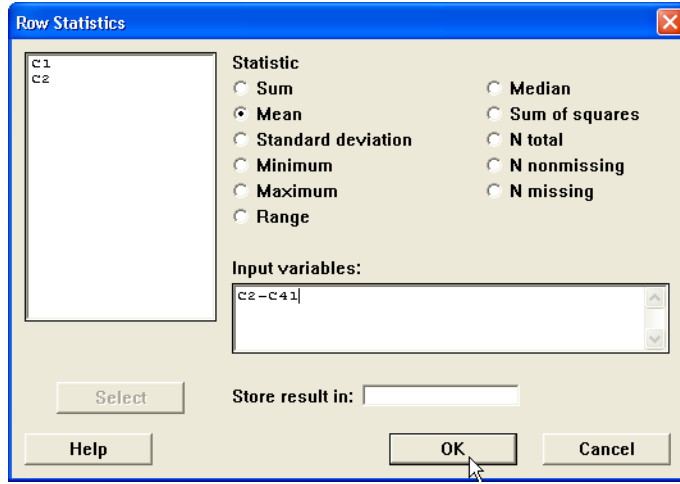
Şimdi anakütleden bir örneklem çektiniz ve bu prosedürü yaklaşık 50 ila 100 kere tekrarlamalısınız¹³, her seferinde örneklem için yeni bir sütun seçmelisiniz. Bütün bu prosedürlerden sonra, "Calc" menüsünden "Row Statistics"e gelerek bütün örneklemelere ait istatistikleri hesaplayabilirsiniz:

¹³ Gerçek simülasyonlarda bu numara 10000'e kadar ulaşır...



Şekil 7.27

Açılan diyalog ekranında “mean” opsiyonunu seçiniz. C2-C41 aralığını “*Input variables:*” kutusuna yazınız (bu sütunlar sizin örneklemelerinizi sakladığınız sütunlardır. “*Store result in*” kutucuğuna Şekil 7.28’deki gibi C42 sütununu yazınız:



Şekil 7.28

“OK” dediğiniz anda C2-C41 deki örneklemelerin ortalamaları, yani 40 örneklemenizin ortalamaları, C42 sütununda saklanmış olacaktır. Daha sonra ise histogram oluşturup örneklem dağılımı hakkında bir fikir sahibi olabilirsiniz. Dolayısıyla MLT’in doğruluğu ispatlanmış olur.

Dolayısıyla, büyük örneklem, anakütle dağılımı ne olursa olsun normal olarak dağılmaktadırlar diyebiliriz.

Tahmin edicilerin özellikleri

Örneklem ortalaması ($\bar{x}$), örneklem standart sapması (s) ve örneklem oranı (p) gibi örneklem istatistiklerinin anakütle parametrelerini μ, σ, π birer tahmin edici olarak nasıl tahmin ettiklerini inceledik. Benzer şekilde, tahmin edilmiş örneklem ortalamasının standart hatasının ($s_{\bar{x}}$) ve tahmin edilmiş örneklem oranının standart hatasının (s_p), $\sigma_{\bar{x}}$ ve σ_p ’nin nasıl birer tahmin edicisi olarak kullanıldığını da inceledik. Ancak örneğin anakütle ortalamasını tahmin ederken ortalama değeri veya düzeltilmiş ortalama da kullanabilirdik. Dolayısıyla bizim tahmin

edicileri karşılaştırabilmemiz için ve hangisi(leri)nin daha iyi tahmin gerçekleştirebileceğini anlayabilmek için, bu tahmin edicilerin özellikleri hakkında fikir sahibi olmamız gerekmektedir. Bu özelliklerden birincisi sapmasızlık özelliğidir. Sapmasızlık özelliğine sahip bir tahmin edicinin değeri tahmin edilmeye çalışılan anakütle parametre değerine eşit çıkmaktadır. Bir başka ifade ile anakütlede yeralan hedef tam olarak vurulmuştur.

Bir diğer önemli tahmin edici özelliği ise etkinliktir. Etkinlik özelliğine sahip bir tahmin edicinin örneklem dağılımı daha çok anakütle ortalamasının etrafında kümelenmiştir ve düşük bir varyansa sahiptir.

Son tahmin edici özelliğimiz ise tutarlılıktır. Bu özellik daha çok, örneklem boyutu değiştiğinde örneklem dağılımında meydana gelen değişimlerle alakalıdır. Bir örneklem boyutu sonsuza doğru yakınsarken, örneklem parametresi doğru anakütle parametresine yakınsıyorsa, tahmin edici tutarlılık özelliğine sahiptir.

Tahminlerimiz küçük örneklem boyutlarında sapmalı olsa dahi, büyük örneklem boyutlarında tahminlerimizin sapmasız olacağını garanti ettiğinden tutarlılık önemli bir tahmin edici özelliğidir.

Sapmasızlık

Tahmin edicilerin özellikleri hakkında bazı genel yorumlarda bulunabilmek için kullanacağımız notasyona odaklanalım. Anakütle parametrelerini θ ile, ve örneklem parametrelerini ise $\hat{\theta}$ ile ifade edelim. Herhangi bir örneklem tahmin edicisinin aşağıdaki özellikleri taşıması tahminin sağlıklı açısından önemlidir. Aşağıdaki eşitlik sağlandığı takdirde, örneklem istatistiğimiz $\hat{\theta}$ parametresinin sapmasız tahmin edicisidir sonucuna varabiliriz.

$$E(\hat{\theta}) = \theta$$

Bu eşitlikte $E(\hat{\theta})$ ifadesi, örneklem istatistiğinin beklenen veya ortalama değerini ifade etmektedir ve örneklem istatistiğinin beklenen bütün olası değerlerinin tahmin edilmiş olan anakütle parametresine eşit olduğunu göstermektedir.

Tahmin edici her zaman sapmasız olacak diye mutlak bir kaide olmadığı gibi tahmin edicilerdeki sapmayı ve büyüklüğünü aşağıdaki gibi ifade edebiliriz.

$$Sapma = E(\hat{\theta}) - \theta$$

Eğer bu eşitlik pozitif bir değer alırsa, örneklem istatistiği büyük bir olasılıkla anakütle parametresini olduğu değerden daha büyük olarak, negatif bir değer alırsa olduğu değerden küçük bir değer olarak tahmin etmiştir sonucuna varabiliriz.

Örneklem Ortalamasının Sapmasızlığı

$\bar{X}$ ifadesi anakütle ortalamasının sapmasız bir tahmin edicisidir.

$$E(\bar{X}) = \mu$$

Bu eşitlik açık bir şekilde EK 7.1’de ispatlanmıştır.

Örneklem Varyansının Sapmasızlığı

s^2 ifadesi anakütle varyansının sapmasız bir tahmin edicisidir.

$$E(S^2) = \sigma^2$$

Bu eşitlik açık bir şekilde EK 7.5’te ispatlanmıştır.

Örnek 7.16: Örnek 7.1’de örneklem varyansının beklenen değerini, aşağıdaki gibi hesaplamıştık.

$$E(S^2) = 0 \times \frac{3}{9} + 0.5 \times \frac{4}{9} + 2 \times \frac{2}{9} = \frac{2}{3}$$

Bütün olası örneklemelerimizin varyansları aşağıdaki gibi olduğundan

$$\begin{aligned} s_1^2 &= 0, \quad s_2^2 = 0.5, \quad s_3^2 = 2, \\ s_4^2 &= 0.5, \quad s_5^2 = 0, \quad s_6^2 = 0.5, \\ s_7^2 &= 2, \quad s_8^2 = 0.5, \quad s_9^2 = 0, \end{aligned}$$

Açık bir şekilde örneklem varyansının beklenen değeri, anakütle varyansına σ^2 , eşittir.

Örneklem Oranının Sapmasızlığı

p ifadesi anakütle oranının sapmasız bir tahmin edicisidir.

$$E(p) = \pi$$

Bu eşitlik açık bir şekilde EK 7.3’te ispatlanmıştır.

Örneklem Ortalamasının Standart Hatasının Sapmasızlığı

$s_{\bar{X}}$ ifadesi anakütle standart hatasının $\sigma_{\bar{X}}$, sapmasız bir tahmin edicisidir.

$$E(s_{\bar{X}}) = \sigma_{\bar{X}}$$

Bu eşitlik açık bir şekilde EK 7.6’da ispatlanmıştır.

Örneklem Oranının Standart Hatasının Sapmasızlığı

s_p^2 ifadesi anakütle standart hatasının σ_p^2 , sapmasız bir tahmin edicisidir.

$$E(s_p^2) = \sigma_p^2$$

Bu eşitlik açık bir şekilde EK 7.7’de ispatlanmıştır.

Etkinlik

Bir anakütle parametresini tahmin etmeye çalışan iki adet sapmasız tahmin edicinin elimizde olduğunu varsayalım. Her iki tahmin edici de sapmasız bir şekilde anakütle ortalamasını tahmin ettiğine göre hangi tahmin ediciyi seçmemiz gereklidir? İşte burada dikkat etmemiz gereken husus, tahmin edicilerden varyansı küçük olanı ele almaktır. Çünkü varyansı küçük olan tahmin edici, anakütle parametresine daha yakın değerler etrafında toplanmıştır ve daha az yayvandır. Bir tahmin edicinin daha küçük varyans değerine sahip olması, daha fazla göreceli etkinliğinin olduğunu göstermektedir. En etkin tahmin ediciler, varyansı en küçük olan tahmin edicilerdir. Aşağıdaki koşullar sağlandığından $\hat{\theta}$ ifadesi, θ parametresinin etkin tahmin edicisidir.

(a) $\hat{\theta}$ sapmasızdır

(b) $\text{Var}(\hat{\theta}) \leq \text{Var}(\tilde{\theta})$, burada $\tilde{\theta}$ tahmin edisi de θ parametresinin sapmasız olarak tahmin etmektedir.

Tutarlılık

Örneklem boyutu artarken nokta tahmin edicinin değerinin anakütle parametresinin değerine doğru yaklaşması o tahmin edicinin tutarlı olduğunu göstermektedir. Büyük örneklem boyutlarında, tahmin edicilerin tahminlerinde küçük örneklemelere göre daha iyi sonuçlar elde edilir. $\bar{x}$ veya p ’nin örneklem dağılımlarında, standart sapmaların ($\sigma_{\bar{x}}$ ve σ_p) örneklem boyutu n ile ters bir orantı içerisinde olduğunu belirtmek gerekmektedir. Sonuç olarak, örneklem boyutu arttığında tahmin edilmiş olan standart sapmalar daha küçük değerler olarak dağılımın daha çok anakütle parametresi (μ ve π) etrafında olmasını temin etmektedirler. Başka bir ifade ile bir tahmin edicinin tutarlı olabilmesi için, örneklem boyutunun sonsuza gittiği durumda tahmin edicinin değerinin gerçek anakütle parametresinin değerine yakınsaması gerekmektedir.

Tahmin Edicilerin Seçimi

Etkinlik kriterini incelerken, iki tahmin edici sapmasız ise varyansı daha küçük olanı seçeriz sonucuna varmıştık. Fakat, tahmin edicilerimizden biri sapmalı biri sapmasız ise nasıl bir seçim gerçekleştirmemiz gerekmektedir? Sapmasız olan tahmin ediciyi mi yoksa varyansı küçük olan tahmin ediciyi mi seçmeliyiz? Ortalama Kareler Toplamı (OKT) bu tarz durumlar için bir kriter belirlemiştir. OKT’yi şöyle tanımlayabiliriz.

$$OKT(\hat{\theta}) = E[\hat{\theta} - \theta]^2$$

veya bir başka şekilde ifade edecek olursak (bkz EK 7.8)

$$OKT(\hat{\theta}) = \text{var}(\hat{\theta}) + (\text{Sapma})^2$$

OKT kriteri, en düşük OKT değerine sahip olan tahmin ediciyi seçmemiz gerektiğini öğretmektedir. Eğer iki tahmin edici sapmasız ise, sapma miktarı sıfıra eşit olacak ve bu kriter minimum varyans kriterine (etkinlik kriterine) indirgenecektir.

EK

Ek 7.1 $E(\bar{X}) = \mu$ eşitliğinin ispatı

İspat: Denklem (7.2)'nin ispatını gerçekleştirebilmek için, öncelikle aritmetik ortalama formülünü denklemin içinde yerine koyalım.

$$\begin{aligned}
E(\bar{X}) &= E\left(\frac{\sum_{i=1}^n x_i}{n}\right) \\
&= E\left[\frac{1}{n}(X_1 + X_2 + \dots + X_n)\right] \\
&= \frac{1}{n}E(X_1 + X_2 + \dots + X_n) \\
&= \frac{1}{n}[E(X_1) + E(X_2) + \dots + E(X_n)]
\end{aligned}$$

n değişkene ait olan n tane elemanın $X_1; X_2; \dots; X_n$ herbiri aynı dağılıma, aynı ortalama ve varyansa sahiptirler. X_1 değeri X 'in bütün olası değerlerinden seçilmiş olan birinci gözlemi ifade ederken, X_2 değeri X 'in bütün olası değerlerinden seçilmiş olan ikinci gözlemi ifade etmektedir. $X_1; X_2; \dots; X_n$ değerlerinin hepsinin X dahil olmak üzere aynı ortalama sahip olduğunu varsaydığımızdan $E(X_1) = E(X_2) = \dots = E(X_n) = \mu$ eşitliği sağlanmaktadır. Dolayısıyla denkleminizi yeniden yazacak olursak;

$$\begin{aligned}
E(\bar{X}) &= \frac{1}{n}[\underbrace{\mu + \mu + \dots + \mu}_{n \text{ kere}}] \\
&= \frac{1}{n}n\mu \\
E(\bar{X}) &= \mu
\end{aligned}$$

sonucunu elde edebiliriz.

Ek 7.2 $\text{var}(\bar{X}) = \frac{\sigma^2}{n}$ eşitliğinin ispatı

İspat: Denklem (7.3)'ün ispatını gerçekleştirebilmek için, genel varyans formüllerinden faydalanmamız gerekmektedir.

$$\text{var}(ax) = a^2 \text{var}(x)$$

a sabit bir değeri ifade etmektedir.

$$\text{var}(x + y) = \text{var}(x) + \text{var}(y)$$

eşitliği x ve y 'nin birbirinden bağımsız olduğunu ifade etmektedir. Denklem (7.3)'ü ifade edebilmek için yukarıdaki her iki denklemden de faydalanacağız.

$$\text{var}(\bar{X}) = \text{var}\left(\frac{\sum_{i=1}^n x_i}{n}\right) = \text{var}\left(\frac{1}{n} \sum_{i=1}^n x_i\right)$$

Birinci denklemini kullanarak:

$$\begin{aligned}\text{var}(\bar{X}) &= \left(\frac{1}{n}\right)^2 \text{var}\left(\sum_{i=1}^n x_i\right) \\ &= \left(\frac{1}{n}\right)^2 \text{var}(X_1 + X_2 + \dots + X_n)\end{aligned}$$

ifadesini elde ederken, ikinci denklemi kullanarak

$$\text{var}(\bar{X}) = \left(\frac{1}{n}\right)^2 [\text{var}(X_1) + \text{var}(X_2) + \dots + \text{var}(X_n)]$$

eşitliğini elde ederiz. Aynı mantıktan yola çıkarak:

$$\text{var}(X_1) = \text{var}(X_2) = \dots = \text{var}(X_n) = \text{var}(X) = \sigma^2$$

$$\begin{aligned}\text{var}(\bar{X}) &= \left(\frac{1}{n}\right)^2 \underbrace{[\sigma^2 + \sigma^2 + \dots + \sigma^2]}_{n \text{ kere}} \\ &= \left(\frac{1}{n}\right)^2 [n\sigma^2] \\ &= \frac{\sigma^2}{n}\end{aligned}$$

Örneklem ortalamasının varyansı $\frac{\sigma^2}{n}$ olarak elde edilir. Eğer bu ifadenin karekökünü alacak olursak $\bar{x}$ 'ın standart sapmasını denklem (7.4)'te olduğu gibi elde edebiliriz.

$$\sqrt{\frac{\sigma^2}{n}} = \frac{\sigma}{\sqrt{n}}$$

Ek 7.3 $E(P) = \pi$ eşitliğinin ispatı

$$\begin{aligned}E(P) &= E\left(\frac{X}{n}\right) \\ &= \frac{1}{n} E(X)\end{aligned}$$

İspat: Öncelikle $E(X)$ değerini belirlememiz gerekmektedir. X neyi ifade etmektedir? X değeri örneklemdeki istenen kategori sayısını ifade etmektedir. Örnek olarak, örneklemdeki bayan sayısı veya oy verenlerden belli bir partiye oy verenlerin sayısı gibi. Bu yüzden X binom dağılımına

sahiptir. Bildiğiniz gibi binom dağılımına sahip olan değişkenlerin ortalamaları $E(X) = n\pi$ değerine eşit idi. Dolayısıyla;

$$\begin{aligned} E(P) &= \frac{1}{n} n\pi \\ &= \pi \end{aligned}$$

Ek 7.4 $\text{var}(p) = \frac{\pi(1-\pi)}{n}$ eşitliğinin ispatı

İspat:

$$\begin{aligned} \text{var}(p) &= \text{var}\left(\frac{X}{n}\right) \\ &= \frac{1}{n^2} \text{var}(X) \end{aligned}$$

Binom dağılımına sahip bir değişkenin varyansı aşağıdaki gibidir.

$$\text{var}(X) = n\pi(1-\pi)$$

Dolayısıyla;

$$\begin{aligned} \text{var}(p) &= \frac{1}{n^2} n\pi(1-\pi) \\ &= \frac{\pi(1-\pi)}{n} \end{aligned}$$

Ek 7.5 $E(S^2) = \sigma^2$ eşitliğinin ispatı

İspat:

Eşitliğin her iki tarafına μ değerini ekleyip çıkaralım

$$\begin{aligned} E(S^2) &= \frac{1}{n-1} E\left(\sum_{i=1}^n (x_i - \bar{x} - \mu + \mu)^2\right) \\ &= \frac{1}{n-1} E\left(\sum_{i=1}^n [(x_i - \mu) - (\bar{x} - \mu)]^2\right) \\ &= \frac{1}{n-1} E\left(\sum_{i=1}^n [(x_i - \mu)^2 - 2(x_i - \mu)(\bar{x} - \mu) + (\bar{x} - \mu)^2]\right) \\ &= \frac{1}{n-1} E\left[\sum_{i=1}^n (x_i - \mu)^2 - 2(\bar{x} - \mu)\sum_{i=1}^n (x_i - \mu) + \sum_{i=1}^n (\bar{x} - \mu)^2\right]; \quad 2(\bar{x} - \mu) \text{ değeri rassal olmadığından} \end{aligned}$$

$$\begin{aligned}
&= \frac{1}{n-1} E \left[\sum_{i=1}^n (x_i - \mu)^2 - 2(\bar{x} - \mu) \sum_{i=1}^n (x_i - \mu) + n(\bar{x} - \mu)^2 \right]; \quad 2(\bar{x} - \mu) \text{ değeri rassal olmadığından} \\
&= \frac{1}{n-1} E \left[\sum_{i=1}^n (x_i - \mu)^2 - 2(\bar{x} - \mu) \left(\sum_{i=1}^n x_i - \sum_{i=1}^n \mu \right) + n(\bar{x} - \mu)^2 \right] \\
&= \frac{1}{n-1} E \left[\sum_{i=1}^n (x_i - \mu)^2 - 2(\bar{x} - \mu) \left(n \frac{\sum_{i=1}^n x_i}{n} - n\mu \right) + n(\bar{x} - \mu)^2 \right] \\
&= \frac{1}{n-1} E \left[\sum_{i=1}^n (x_i - \mu)^2 - 2(\bar{x} - \mu)(n\bar{x} - n\mu) + n(\bar{x} - \mu)^2 \right]; \quad \bar{x} = \frac{\sum_{i=1}^n x_i}{n} \\
&= \frac{1}{n-1} E \left[\sum_{i=1}^n (x_i - \mu)^2 - 2n(\bar{x} - \mu)^2 + n(\bar{x} - \mu)^2 \right] \\
&= \frac{1}{n-1} E \left[\sum_{i=1}^n (x_i - \mu)^2 - n(\bar{x} - \mu)^2 \right] \\
&= \frac{1}{n-1} \left[\sum_{i=1}^n E(x_i - \mu)^2 - E[n(\bar{x} - \mu)^2] \right] \\
&= \frac{1}{n-1} \left[\sum_{i=1}^n \text{var}(x) - nE[(\bar{x} - \mu)^2] \right]; \quad \text{tanım gereği } \text{var}(X) = E(x_i - \mu)^2 \\
&= \frac{1}{n-1} [n \text{var}(x) - n \text{var}(\bar{x})]; \quad \text{tanım gereği } \text{var}(\bar{x}) = E(\bar{x}_i - \mu)^2 \\
&= \frac{1}{n-1} \left[n\sigma^2 - n \frac{\sigma^2}{n} \right]; \quad \text{var}(\bar{x}) = \frac{\sigma^2}{n} \text{ olduğundan} \\
&= \frac{1}{n-1} (n\sigma^2 - \sigma^2) \\
&= \frac{1}{n-1} (n-1)\sigma^2 \\
&E(S^2) = \sigma^2
\end{aligned}$$

Ek 7.6 $E(S_{\bar{x}}) = \sigma_{\bar{x}}$ eşitliğinin ispatı

İspat:

$$\begin{aligned}
E(s_{\bar{x}}) &= E\left(\frac{s}{\sqrt{n}}\right) \\
&= \frac{1}{\sqrt{n}} E(S) \\
&= \frac{1}{\sqrt{n}} E\left(\sqrt{\frac{\sum_{i=1}^n (x_i - \bar{X})^2}{n-1}}\right) \frac{\sigma}{\sqrt{n}}; \text{ daha önce ispatladığımız gibi } E\left(\sqrt{\frac{\sum_{i=1}^n (x_i - \bar{X})^2}{n-1}}\right) = \sigma
\end{aligned}$$

Ek 7.7 $E(s_p^2) = \sigma_p^2$ eşitliğinin ispatı

İspat:

$$\begin{aligned}
E(s_p^2) &= E\left[\frac{p(1-p)}{n}\right] \\
&= \frac{1}{n} E[p(1-p)] \\
&= \frac{1}{n} E(p - p^2) \\
&= \frac{1}{n} [E(p) - E(p^2)] \\
&= \frac{1}{n} [\pi - \pi^2]; \quad E(p) = \pi \text{ ve } E(p^2) = [E(p)]^2 = \pi^2 \text{ olduğundan} \\
&= \frac{1}{n} [\pi(1-\pi)] \\
&= \frac{\pi(1-\pi)}{n}
\end{aligned}$$

Ek 7.8 $OKT(\hat{\theta}) = E[\hat{\theta} - \theta]^2$ eşitliğinin ispatı

Eşitliğin her iki tarafına da θ değerini ekleyip çıkaralım

$$\begin{aligned}
OKT(\hat{\theta}) &= E\left[\left(\hat{\theta} - E(\hat{\theta})\right) - \left(\theta - E(\hat{\theta})\right)\right]^2 \\
&= E\left[\left(\hat{\theta} - E(\hat{\theta})\right)^2 - 2\left(\hat{\theta} - E(\hat{\theta})\right)\left(\theta - E(\hat{\theta})\right) + \left(\theta - E(\hat{\theta})\right)^2\right] \\
&= E\left[\left(\hat{\theta} - E(\hat{\theta})\right)^2 - 2\hat{\theta}\theta + 2\hat{\theta}E(\hat{\theta}) + 2\theta E(\hat{\theta}) - 2E(\hat{\theta})^2 + \left(\theta - E(\hat{\theta})\right)^2\right] \\
&= E\left(\hat{\theta} - E(\hat{\theta})\right)^2 - E\left[2\hat{\theta}\theta\right] + E\left[2\hat{\theta}E(\hat{\theta})\right] + E\left[2\theta E(\hat{\theta})\right] - E\left[2E(\hat{\theta})^2\right] + E\left(\theta - E(\hat{\theta})\right)^2 \\
&= E\left(\hat{\theta} - E(\hat{\theta})\right)^2 - 2\theta E(\hat{\theta}) + 2E(\hat{\theta})E(\hat{\theta}) + 2\theta E(\hat{\theta}) - 2E(\hat{\theta})^2 + \left(E(\hat{\theta}) - \theta\right)^2;
\end{aligned}$$

$E(\hat{\theta})$ ve θ rassal olmadığından

$$= E\left(\hat{\theta} - E(\hat{\theta})\right)^2 + \left(E(\hat{\theta}) - \theta\right)^2$$

$$OKT(\hat{\theta}) = \text{var}(\hat{\theta}) + (Bias)^2;$$

tanım gereği $E\left(\hat{\theta} - E(\hat{\theta})\right)^2 = \text{var}(\hat{\theta})$ ve $E(\hat{\theta}) - \theta = \text{sapma}$ olarak tanımlamıştık

Dolayısıyla OKT aşağıdaki gibi yazılabilir.

$$OKT(\hat{\theta}) = \text{var}(\hat{\theta}) + (Sapma)^2$$

Bölüm 8:

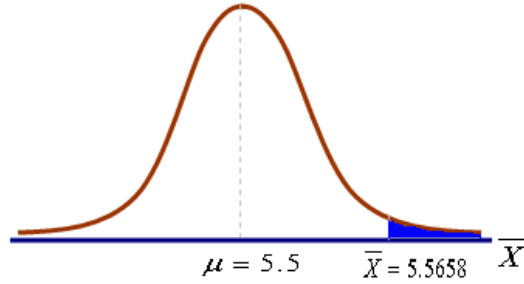
Tek anakütleli Hipotez Testi

Giriş

İstatistikte hipotez kavramı, anakütle karakteristiği hakkında bir önerme, bir çıkarım yapma imkanı sağlamaktadır. Hipotez testinde anakütle parametresi hakkında yaptığımız deneysel varsayıma sıfır hipotezi denilir ve H_0 ile gösterilir. Sıfır hipotezi aslında bize asıl test etmek istediğimiz hipotezi gösterir. Diğer hipotezi ise alternatif hipotez olarak adlandırırız ve H_A ile ifade ederiz. Alternatif hipotez genellikle sıfır hipotezinin tam tersidir. Hipotez testinde elimizdeki örneklemden yararlanarak elde ettiğimiz data ile alakalı olan iki tane önerme yazarız. Bu önermelerden biri H_0 ile ifade edilir ve diğeri H_A ile ifade edilir.

Anakütleden rassal bir şekilde örneklem çekelim. Eğer örneklemimizin verileri sıfır hipotezi ile tutarlı ise, sıfır hipotezini reddetmeyiz. Ancak eğer, örneklemimizin verileri sıfır hipotezi ile tutarlı değilse sıfır hipotezini reddederiz ve birnevi tercihimizi alternatif hipotezden yana kullanırız. Kavramları biraz data detaylı incelemek için Bölüm 7’de yer alan Örnek 7.14’ü tekrar ele alalım.

Örnek 8.1: Örnek 7.14’teki ile aynı boyuta sahip ve ortalaması 5.5658 ($\bar{x} = 5.5658$) olan farklı bir örneklem ele alalım. Örneklem ortalamasının standart sapmasında $\sigma_{\bar{x}} = \sigma / \sqrt{n} = 0.25 / \sqrt{40} = 0.04$ ve bilinmeyen anakütle parametresinin $\mu = 5.5$ olduğu varsayımımızda hiçbir değişiklik bulunmamaktadır. Daha önce de yaptığımız gibi öncelikle, anakütle ortalaması 5.5 ($\mu = 5.5$) iken $\bar{x} = 5.5658$ olan örneklem ortalamasını elde etme olasılığını hesaplamaya çalışalım $P(\bar{X} > 5.5658) = ?$ Bu olasılık değerini grafiksel olarak gösterecek olursak:



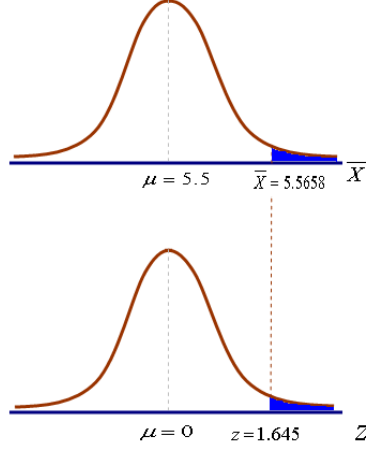
Şekil 8.1

Şekil 8.1’deki taralı alan, ortalama olarak çay poşetlerinin 5.5658 gramdan daha fazla olma olasılığını $P(\bar{X} > 5.5658)$ göstermektedir. Olasılık değerini hesaplamadan önce 5.5658 değerini kendisine denk gelen z değerine dönüştürmemiz gerekmektedir.

$$z = \frac{\bar{X} - \mu}{\sigma_{\bar{x}}} = \frac{5.5658 - 5.5}{0.04} = 1.645$$

$P(\bar{X} > 5.5658) = P(Z > 1.645)$ olduğundan Normal Dağılım Tablosu’nu kullanabiliriz.

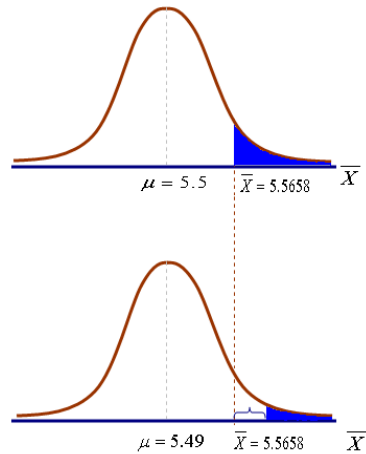
Grafiksel olarak,



Şekil 8.2

Normal Dağılım Tablosu'ndan (veya MINITAB kullanarak) $P(Z > 1.645) = 0.05$ değerini hesaplayabiliriz. Bir önceki bölümde bahsettiğimiz nedenlerden dolayı, bu olasılık değeri göreceli olarak küçük bir olasılık değeridir. Bu olasılık değeri, $\mu = 5.5$ olduğu takdirde $\bar{X} = 5.4342$ değere sahip bir örneklem elde etme olasılığımızın yüzde 5 olduğunu belirtmektedir. Eğer, anakütle ortalaması 5.5 değildir dediğimiz takdirde, (gerçekte anakütle ortalamasının 5.5 olduğu durumda) yüzde 5 olasılıkla hata yapmak ihtimalimiz bulunmaktadır. . Bu olasılık değeri, gerçek anakütle ortalaması 5.5 olsa bile 5.5 değildir diyerek hata yapmamız için küçük bir olasılık değeridir. Dolayısıyla bu olasılık değeri, gerçek ortalama değeri 5.5 değerine eşit olamaz sonucuna varmak için düşük bir hata değeridir. Dolayısıyla, anakütle ortalamasının üreticinin iddia ettiği gibi 5.5 olmadığı sonucuna varırız.

Bir önceki bölümde yer alan örnek 7.14'teki ifadelerimizi hatırlayacak olursak: Peki anakütle ortalaması 5.5'ten küçük olabilir mi? Mesela 5.49 olabilir mi? Şimdi ise hipotetik ortalaması $\mu = 5.49$ ve $\sigma_{\bar{X}} = 0.04$ standart sapması ile normal olarak dağılan, $\bar{X} = 5.5658$ örneklem ortalamasını elde etme olasılığını $P(\bar{X} > 5.5658) = 0.029$ olarak hesaplayabiliriz. Bu olasılık değerini grafiksel olarak gösterecek olursak:

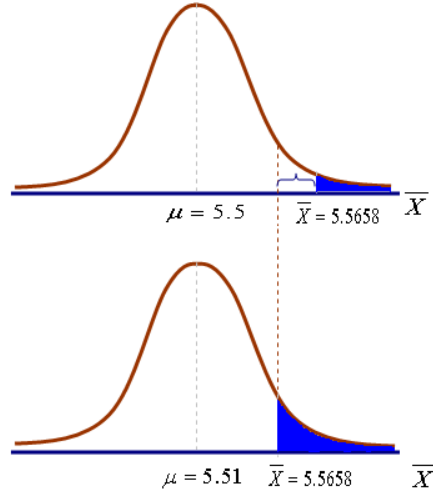


Şekil 8.3

$\mu = 5.49$ iken hesaplamış olduğumuz olasılık değeri 0.05 değerinden küçük olduğundan, gerçek anakütlesinin 5.49 değerine eşit olmadığı sonucuna varırız.

Şimdi ise anakütle ortalamasının 5.51 gibi 5.5'ten büyük olduğu durumu ele alalım. $\mu = 5.51$ durumunda örneklem ortalamasının 5.5658 olma olasılığı hesaplayalım.

$\mu = 5.51$ durumunda $P(\bar{X} > 5.5658)$ olasılık değeri 0.0815'tir ve bu değer de 0.05 değerinden büyük bir değerdir. Bu değeri grafiksel olarak ifade edecek olursak:



Şekil 8.4

Dolayısıyla gerçek anakütle ortalamasının 5.51 olduğu sonucuna varırız.

Sonuç olarak 0.05 değerini, hipotetik anakütle değerini doğru olduğu takdirde reddederek maksimum hata yapma olasılığımız olarak ele alalım.

$$\begin{aligned} H_0 : \mu &\leq 5.5 \\ H_A : \mu &> 5.5 \end{aligned} \quad (8.1)$$

Dolayısıyla hesaplamış olduğumuz olasılık değeri hata yapma olasılığımız olan 0.05 değerinden büyük bir değer ise sıfır hipotezini kabul ederiz. Bu olasılık değeri hata yapma olasılığımız olan 0.05 değerinden küçük bir değer ise sıfır hipotezini reddederiz. Eğer maksimum hata yapma olasılığımızı yüzde 10'a çıkarırsak, her iki durumda da örneğimizdeki sıfır hipotezini reddederiz. Test sonuçlarımızda anlamlılık düzeyi çok büyük bir öneme sahiptir.

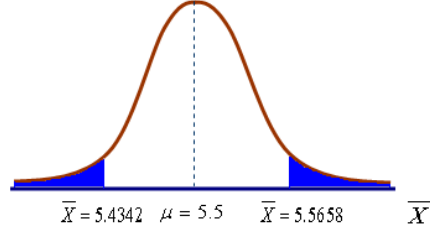
Örnek 7.14'te denklem (8.2)'de yer alan sıfır hipotezimizi, örneklem boyutu $n = 40$ ve örneklem ortalaması $\bar{x} = 5.4342$ iken reddetmiştik.

$$\begin{aligned} H_0 : \mu &\geq 5.5 \\ H_A : \mu &< 5.5 \end{aligned} \quad (8.2)$$

Denklem (8.1) ve (8.2)'deki hipotezleri tek bir hipotez formatına getirdiğimizde denklem (8.3)'ü elde ederiz.

$$\begin{aligned} H_0 : \mu &= 5.5 \\ H_A : \mu &\neq 5.5 \end{aligned} \quad (8.3)$$

Örnek 7.14'te anakütle ortalamasının μ , 5.5 değerinden daha küçük olma olasılığı 0.05 değerinden büyük iken, diğer bir yandan bu bölümde ele aldığımız örnekte ise anakütle ortalamasının μ , 5.5 değerinden daha büyük olma olasılığı 0.05 değerinden büyüktür. Dolayısıyla denklem (8.3)'te yer alan hipotezi reddettiğimizde yüzde 10'luk bir seviyede hataya düşmüş olacağımızdan, yüzde 10'luk bir anlamlılık düzeyinde $\mu = 5.5$ hipotezini reddedemeyiz.



Şekil 8.5

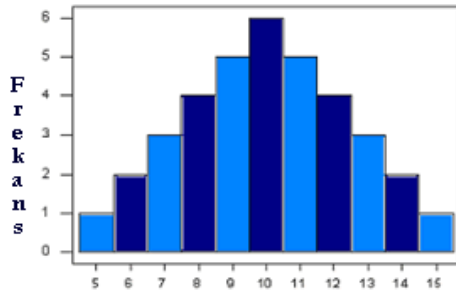
Denklem (8.1) ve (8.2)'de yeralan hipotezleri tek taraflı hipotezler iken, denklem (8.3)'te yer alan hipotezimiz ise çift taraflı bir hipotezdir. Hipotez testlerimizdeki karar mekanizmasını daha detaylı incelemeden önce bir örnek daha inceleyelim.

Örnek 8.2: Şimdi tablo 8.1'deki hipotetik anakütleyi ele alalım.

Tablo 8.1 Hipotetik anakütle

5
6 6
7 7 7
8 8 8 8
9 9 9 9 9
10 10 10 10 10 10
11 11 11 11 11
12 12 12 12
13 13 13
14 14
15

$N = 36$, N daha önceki bölümlerden de hatırlayacağınız gibi bize anakütlenin boyutunu ifade etmektedir. Anakütlenin verilerinin dağılımı tam olarak simetriktr. Anakütle dağılımını Şekil 8.6'daki gibi analiz edilebilir.



Şekil 8.6

Bu dağılımın, elde edebileceğimiz histogramın şeklinden yola çıkarak, normal olduğunu varsayalım. Burada unutulmaması gereken normal dağılımın sürekli bir dağılım olduğudur. Anakütlemizin ortalaması ve standart sapması kolayca aşağıdaki gibi hesaplanabilir: $\mu = 10$ ve $\sigma = 2.449$

Aynı anakütleden rassal bir şekilde aşağıdaki örneklemleri çektiğimizi düşünelim (tekrar yerine konulabilecek şekilde): Örneklem₁={7,9,9,10,12,15} örnekleminin boyutu $n = 6$. Örneklem ortalaması ise : $\bar{X} = 10.33$. Şimdi ise anakütle verisi hakkında bilgimiz olmadığını ve sadece elimizdeki örneklemden faydalanarak anakütle hakkında çıkarım yapmamız gerektiğini düşünelim. Ancak size bilgi olarak anakütle dağılımının normal olduğunu ve anakütle standart sapmasının 2.449 ($\sigma = 2.449$) olduğu verilmiş olsun.

Daha önceki iki bölümümüzden de hatırlayacağınız üzere eğer anakütle normal bir şekilde dağılıyorsa $\bar{X}$ 'ın örneklem dağılımı da örneklem boyutundan bağımsız olarak normal bir biçimde dağılmaktadır. Yani örneğimizde de olduğu gibi $n = 6$ da olsa $\bar{X}$ 'ın örneklem dağılımı normaldir. Anakütle ortalamasını bilmediğimizi düşünelim (normal hayatta genelde bilmiyoruz) ve anakütle ortalamasının (gerçekte 10'a eşit olduğu halde) 12 ye eşit olup olmadığını test edelim. Formal olarak hipotezimizi şöyle ifade edebiliriz:¹⁴

$$H_0 : \mu = 12 \quad (8.4)$$

Sıfır hipotezimiz, anakütle ortalamasının $E(\bar{X}) = \mu_{\bar{X}} = \mu$ 12'ye eşit olduğunu göstermektedir. Bu hipotezi reddedebiliriz (ki bu doğru bir karardır) veya reddedemeyiz (bu karar yanlış bir karardır).

Şimdi ise bu hipotezi nasıl test edeceğimizi inceleyelim. Anakütlesini tahmin edebilmek için elimizde sadece bir tahmin edici bulunduğundan, hipotezimizi test ederken bu tahmin ediciden faydalanacağız. Hipotetik anakütle ortalamamız 12 iken örneklem ortalamamızın 10.33 veya daha düşük bir değere sahip olma olasılığını hesaplamaya çalışalım.

Bu olasılık değerini hesaplayabilmek için ele aldığımız örneklem ortalamasının $\bar{x}$ 'ların olasılık dağılımı hakkında fikir sahibi olmamız gerekmektedir. Daha öncede irdelediğimiz gibi MLT'den dolayı $\bar{x}$ 'ların dağılımının normal olduğu bilinmektedir. Bir olasılık dağılımından herhangi bir olasılık hesaplamak istediğimizde, o dağılımın ortalama ve standart sapma değerlerini bilmemiz gerekmektedir. Bu yüzden olasılık değerimizi hesaplayabilmek için normal olarak dağılan örneklem ortalamasının $\bar{x}$, ortalaması ve standart sapmasını bilmemiz gerekmektedir. Örneklem ortalamasının standart sapması bulunurken aşağıdaki formülün kullanıldığını daha önce göstermiştik:

$$\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{n}}$$

Örneğimize dönecek olursak, anakütle standart sapması $\sigma = 2.449$ değerine eşit olduğundan:

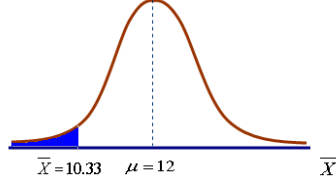
$$\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{n}} = \frac{2.449}{\sqrt{6}} = 1$$

Bir önceki bölümde örneklem ortalamasının örneklem dağılımının ortalamasının anakütle ortalamasına eşit olduğunu göstermiştik ($E(\bar{X}) = \mu_{\bar{X}} = \mu$). Fakat, anakütleyi bilmemiz mümkün olmadığından, elimizdeki örneklemden faydalanarak anakütle ortalamasını tahmin etmeye çalışırız. Bu tahmin sürecinde örneklem dağılımını nasıl kullanabiliriz? Bu sorunun cevabını bize, anakütle ortalaması hakkında hipotez testi sürecinde elimizdeki örneklemleri elde etme olasılığımızın

¹⁴ Şimdilik alternatif hipotezinin nasıl yapılandırıldığına değinmedik. İlerleyen bölümlerde alternatif hipotezin nasıl yapılandırıldığını daha detaylı irdedeceğiz.

hesaplanması verecektir. Örneğimizde anakütle ortalamasını 12 olarak düşündüğümüzden $\bar{x} = 10.33$ değerini elde etme olasılığını hesaplayabiliriz. Bu olasılık değerinin büyüklüğüne göre hipotezimizi $\mu = 12$ kabul veya red edeceğiz.

Bu olasılık değerini grafiksel olarak gösterecek olursak;



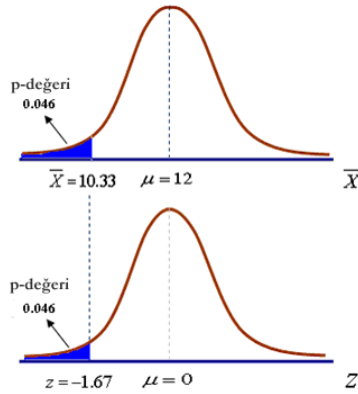
Şekil 8.7

H_0 ile anakütle ortalamasının 12 ye eşit olduğu hipotezini kurduğumuzdan dolayı şekil 8.7’de dağılımın ortasında 12 değeri yer almaktadır. Eğer bu olasılık değerini hesaplırsak ve bu olasılık değerini küçük bir değer olarak bulursak, anakütle ortalamasının 12 olduğunu belirttiğimiz sıfır hipotezimizi reddederiz. Neden reddederiz? Çünkü $\mu = 12$ iken hesaplamış olduğumuz olasılık değerinin küçük olması, elimizdeki örneklem ortalamasını gözlemleme olasılığının da düşük olduğunu göstermektedir. Bir başka deyişle, anakütle ortalamasının 12 olma olasılığı küçük bir değer olduğundan, $\mu = 12$ sıfır hipotezimizi reddederiz. Şekil 8.7’de yeralan taralı alan $\mu = 12$ hipotezini reddetmemizle birlikte yapabileceğimiz hata ihtimalini temsil etmektedir. Bu taralı alan, olasılık değeri (veya p-değeri) adlandırılmaktadır ve doğru olan hipotezi reddetme olasılığını temsil etmektedir. Aynı zamanda Tip I Hata olarak da bilinmektedir. Örneğimize geri dönecek olursak, anakütle ortalamamız 12 iken (gerçekte bu değer 10 iken) örneklem ortalamasını 10.33 olarak elde etme olasılığımız düşüktür.

Peki bu taralı alan nasıl belirlenmektedir? Daha önceki bilgilerimizden de hatırlayacağınız gibi bu olasılık değerini z dönüşüm formülü ve istatistiksel tablolardan faydalanarak hesaplayabiliriz. Dönüşüm formülümüz:

$$z = \frac{\bar{x} - \mu}{\sigma_{\bar{x}}} = \frac{10.33 - 12}{1} = -1.67$$

Veya grafiksel olarak



Şekil 8.8

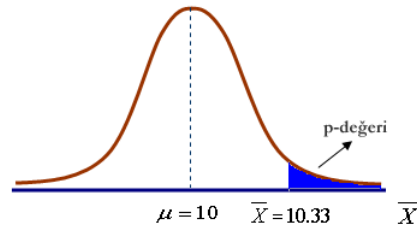
Her iki şekildeki taralı alan birbirine eşittir $P(\bar{x} < 10.33) = P(Z < -1.67)$. Dolayısıyla z-dağılım tablosunu kullanarak $P(Z < -1.67) = 0.046$ değerini hesaplayabiliriz. Anakütle ortalaması $\mu = 12$ iken örneklem ortalamasını $\bar{x} = 10.33$ olarak elde etme olasılığımız 0.047’dir ve küçük bir

değer olduğundan dolayı sıfır hipotezini $\mu = 12$ reddederiz. Olasılık değerimizin büyük veya küçük olmasını anlamlılık düzeyleri ile veya Tip I hata (doğru hipotezi reddetme) yapma olasılığımızın üst sınırları ile karşılaştırmaktayız. Genel olarak bu olasılıklar (sınırlar): %1, %5, veya %10 değerlerini almaktadırlar. Örneğimizde bu sınır değerini (anlamlılık düzeyini) %5'te alsak %10'da alsak sıfır hipotezini $\mu = 12$ reddederiz. Ancak bu sınır değerini yüzde 1 olarak alırsak sıfır hipotezimizi reddedemeyiz. Çünkü olasılık değerimiz 0.047, sınır olarak aldığımız yüzde 1 değerinden daha büyüktür.

Örnek 8.3: Şimdi ise denklem (8.5)'i test edelim

$$H_0 : \mu = 10 \quad (8.5)$$

Sıfır hipotezimiz anakütle ortalamasının 10'a eşit olduğunu iddia etmektedir (ki gerçekte de anakütle ortalamamızın 10 olduğunu belirtmiştik). Bu testte de daha önceki testlerimizde de sorguladığımız gibi; anakütle ortalaması 10 olan bir anakütleden, örneklem ortalaması 10.33 olan örneklem seçme olasılığını sorgulamalıyız. Bir başka deyişle şekil 8.9'da yeralan taralı alanın değerini sorgulamamız gerekmektedir.



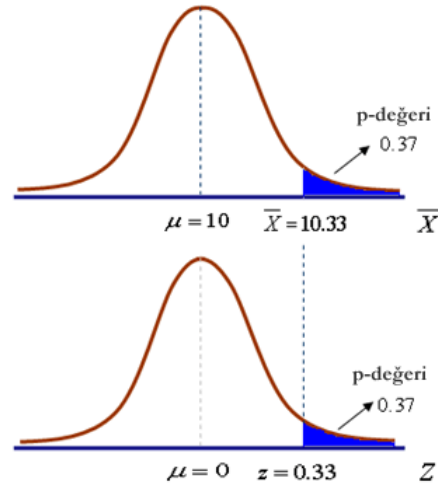
Şekil 8.9

Anakütle ortalamasının 10 olduğunu sıfır hipotezimizde iddia ettiğimizden şekil 8.9'daki örneklem dağılımının ortasına 10 değerini yerleştirdik ve bu şekildeki taralı alanda bize anakütle ortalaması 10 olan bir anakütleden seçilen örneklem ortalamasının 10.33'e eşit veya daha büyük olma olasılığını vermektedir. Bu olasılığı hesaplırsak ve bu olasılık yeterince büyük olursa, anakütle ortalamasının 10 olduğuna dair kurduğumuz hipotezimizi kabul edebiliriz.

10.33 değerine tekabül eden z değerini hesaplayacak olursak:

$$z = \frac{\bar{x} - \mu}{\sigma_{\bar{x}}} = \frac{10.33 - 10}{1} = 0.33$$

$P(\bar{X} > 10.33) = P(Z > 0.33)$, eşitliğinden p değeri Normal Dağılım Tablosu'ndan 0.37 olarak elde edebiliriz. Grafiksel olarak:



Şekil 8.10

Anakütle ortalaması $\mu = 10$ olan bir anakütleden seçilen örneklemin ortalamasının $\bar{X} = 10.33$ olma olasılığı 0.37'dir ve Tip I Hata yapmak için yüksek bir olasılıktır. Testimizin sonucu, alternatif anlamlılık düzeylerinin tamamından ($\alpha = 0.1$; $\alpha = 0.05$, veya $\alpha = 0.01$) etkilenmemektedir. Değerlerin tamamı hesaplamış olduğumuz p-değerinden küçüktür.

Hipotez Testi: p-değeri yaklaşımı

Örneğimizden yola çıkarak basit bir test kuralı geliştirebiliriz. α (alfa olarak okunur) ile temsil edilen farklı olasılık değerlerini $\alpha = 0.01$; $\alpha = 0.05$; $\alpha = 0.1$ anlamlılık düzeyi olarak ele aldığımızda aşağıdaki kuralı kullanarak testimizi gerçekleştirebiliriz.

Hipotez testinde p-değeri yaklaşımı :

Eğer p değeri $< \alpha$ ise H_0 hipotezi reddedilir
Eğer p değeri $> \alpha$ ise H_0 hipotezi reddedilemez.

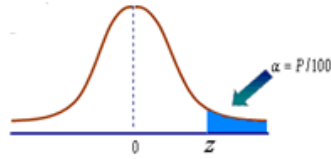
Örnek 8.4: Örnek 8.1'de, anlamlılık düzeyleri $\alpha = 0.05$ ve $\alpha = 0.10$ iken $0.043 < 0.05$; $0.043 < 0.1$ olduğundan $H_0 : \mu = 12$ hipotezini reddetmiştik. Anlamlılık düzeyini $\alpha = 0.01$ olarak ele aldığımızda ise, $0.043 > 0.01$ olduğundan $H_0 : \mu = 12$ hipotezini reddedememiz gerekirdi.

Örnek 8.5: Örnek 8.2'de, anlamlılık düzeyleri $\alpha = 0.1$, $\alpha = 0.05$, veya $\alpha = 0.10$ iken $0.37 > 0.01$; $0.37 > 0.05$; $0.37 > 0.01$ olduğundan $H_0 : \mu = 10$ hipotezini reddetmemiştik. $H_0 : \mu = 10$ hipotezini reddedebilmemiz için α değerini en az 0.4 olarak ele almalıyız.

Hipotez Testi: Kritik Değer Yaklaşımı

P-değeri yaklaşımı sayesinde hipotez testinin ne kadar basit ve kolay gerçekleştirildiğini inceledik. Hipotez testini alternatif olarak başka bir yaklaşım olan kritik değer yaklaşımı ile de gerçekleştirebiliriz. Temelde bu yaklaşım da, p-değeri yaklaşımından çok büyük farklılıklar içermez. Bu yaklaşımı daha detaylı inceleyebilmek için öncelikle kritik değerin ne olduğunu tanımlamamız gerekmektedir. Kritik değer, seçilmiş olan α anlamlılık düzeyine tekabül eden özel z değeri olarak tanımlanabilir. Örnek olarak anlamlılık düzeyini $\alpha = 0.05$ olarak seçtiğimizde, Standart Normal Dağılım Tablosu'ndan (bkz Ek Tablo 2.2) bu olasılık değerine ait olan z değerini aşağıdaki şekildeki gibi bulabiliriz: (genel olarak kritik değerleri, anlamlılık düzeylerine bağlı olarak z_{α} sembolü ile göstermekteyiz)

Yüzdelere Göre Normal Dağılım

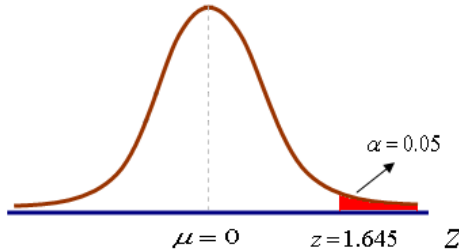


P	z	P	z	P	z	P	z	P	z	P	z
50	0,0000	5,0	1,6449	3,0	1,8808	2,0	2,0537	1,0	2,3263	0,10	3,0902
45	0,1257	4,8	1,6646	2,9	1,8957	1,9	2,0749	0,9	2,3656	0,09	3,1214
40	0,2533	4,6	1,6849	2,8	1,9110	1,8	2,0969	0,8	2,4089	0,08	3,1559
35	0,3853	4,4	1,7060	2,7	1,9268	1,7	2,1201	0,7	2,4573	0,07	3,1947
30	0,5244	4,2	1,7279	2,6	1,9431	1,6	2,1444	0,6	2,5121	0,06	3,2389
25	0,6745	4,0	1,7507	2,5	1,9600	1,5	2,1701	0,5	2,5758	0,05	3,2905
20	0,8416	3,8	1,7744	2,4	1,9774	1,4	2,1973	0,4	2,6521	0,04	3,7190
15	1,0364	3,6	1,7991	2,3	1,9954	1,3	2,2262	0,3	2,7478	0,005	3,8906
10	1,2816	3,4	1,8250	2,2	2,0141	1,2	2,2571	0,2	2,8782	0,001	4,2649
5	1,6449	3,2	1,8522	2,1	2,0335	1,1	2,2904	0,1	3,0902	0,0005	4,4172

$$z_{0.05} = 1.6449$$

Şekil 8.11

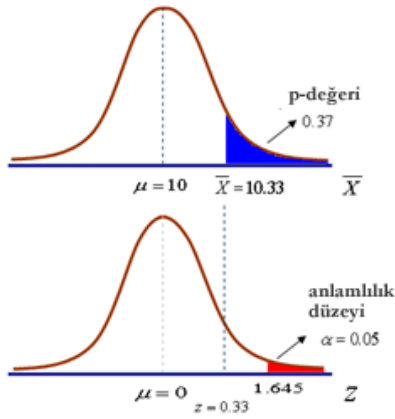
Tabloda da gösterildiği gibi yüzde 5 anlamlılık düzeyine tekabül eden kritik z değeri 1.645 ($P(z > 1.645) = 0.05$) olarak hesaplanmaktadır. Kritik değeri grafiksel olarak gösterecek olursak:



Şekil 8.12

Şekil 8.12'deki taralı alan, Tip I Hata yapabilmemize izin veren azami alanı, anlamlılık düzeyini α ifade etmektedir. **Örnekleme ortalamasına $\bar{x}$ tekabül eden z değerinin taralı alanın içerisinde yeralması; p değerinin, α anlamlılık düzeyinden küçük olduğu anlamına gelmektedir. Dolayısıyla, bu sonuç doğrultusunda sıfır hipotezini reddederiz. Reddetme kararımızda etkin bir rol oynadığından dolayı Şekil 8.12'deki taralı alana REDDETME alanı da denilmektedir.**

Kritik değer yaklaşımını bir örnek 8.3 ile ele alalım. Şekil 8.10'dan, $\bar{x} = 10.33$ değerine tekabül eden z değerinin 0.33 olduğunu bilmekteyiz. Bu değer, reddetme alanının sınırlarının dışında kalmasından dolayı, sıfır hipotezini reddedemeyiz (kabul ederiz). Bu çıkarımı Şekil 8.13'te görsel olarak inceleyim:



Şekil 8.13

Anlamlılık düzeyini $\alpha = 0.1$ olarak belirlesek bile bu çıkarımımız değişmeyecektir, çünkü bu anlamlılık düzeyine denk gelen kritik z değeri 2.326'dır ve bu değer halen 0.33 değerinden daha büyük bir değerdir.

Bu kuralı genelleştirebilmek için, hesaplanan örneklem ortalamasına $\bar{X}$ denk gelen z değerini (örneğimizde, z-skoru = 0.33 idi), **z-skoru** olarak ve α anlamlılık düzeyine denk gelen z değerini (örneğimizde $z_\alpha = 1.645$), z_α olarak adlandıralım. Genel olarak aşağıdaki hipotezi ele alalım:

$$H_0 : \mu = \mu_0$$

μ_0 ifadesi, anakütle ortalamasının μ eşit olduğu varsayılan sayısal değeri temsil etmektedir (Örnek 8.2'de bu değer 10 iken, Örnek 8.3'te 12'dir). α anlamlılık düzeyinde hipotezimiz için reddetme (veya reddetmeme) kuralı aşağıdaki gibidir.

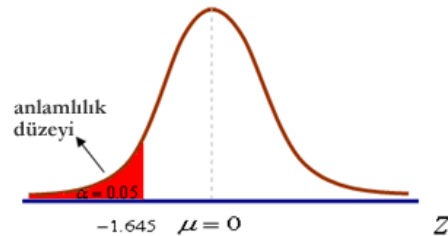
Hipotez Testinde Kritik Değer Yaklaşımı:

Eğer z-istatistik skoru $> z_\alpha$; H_0 Reddedilmektedir

Eğer z-istatistik skoru $< z_\alpha$; H_0 Reddedilmemektedir

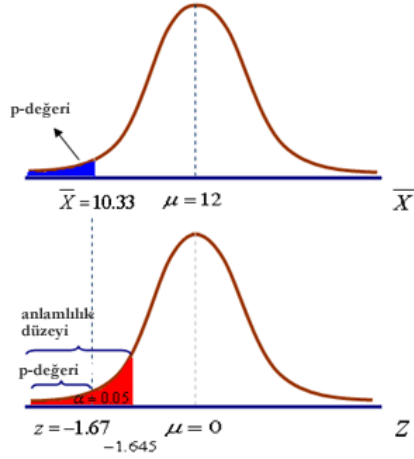
Tekrar Örnek 8.2'yi ele alalım. Tablodan 0.05 olasılık değerine tekabül eden iki adet z değeri bulunmaktadır. Çünkü bu olasılık değeri olasılık dağılımında sağ kuyrukta veya sol kuyrukta da yer alabilmektedir. Bu olasılık değerleri de normal dağılımın simetri özelliğinden dolayı birbirlerine eşittir. Şekil 8.12'de de olduğu gibi, $\alpha = 0.05$ olasılık değerine tekabül eden sağ kuyruktaki z-değeri 1.645 ve sol kuyruktaki z değeri -1.645'tir.

($P(z < -1.645) = 0.05$) Grafiksel olarak gösterecek olursak:



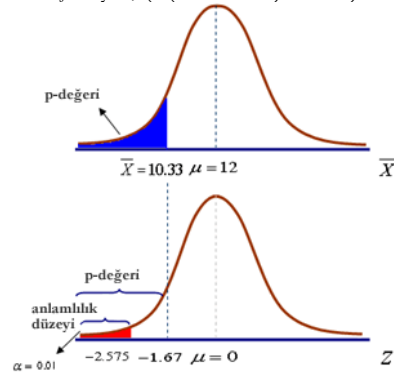
Şekil 8.14

Örnek 8.2'deki $\bar{X} = 10.33$ değerine tekabül eden z değeri Şekil 8.15'te görüldüğü gibi, -1.67 (z-skoru = -1.67) olduğundan, ve bu değer reddedilme alanının içerisinde yeraldığından sıfır hipotezimizi reddederiz.



Şekil 8.15

Eğer anlamlılık $\alpha = 0.01$ olarak seçilseydi, ($P(z < -2.575) = 0.01$)



Şekil 8.16

-1.67 değeri reddetme alanında yeralmadığından sıfır hipotezini kabul ederiz (reddetmeyiz).

α anlamlılık düzeyindeki, örneğimizde ele aldığımız reddetme prosedürümüzü genelleştirecek olursak:

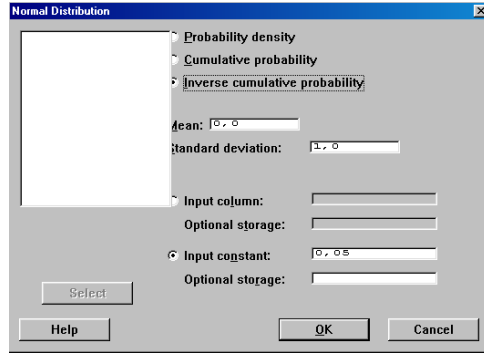
Hipotez Testinde Kritik Değer Yaklaşımı:

Eğer z-istatistik skoru $> -z_{\alpha}$; H_0 reddedilmektedir.

Eğer z-istatistik skoru $< -z_{\alpha}$; H_0 reddedilmemektedir.

Bu prosedüründe z istatistik skorunun işareti belirtilmemiştir (bu değer pozitif de olabilir negatif de olabilir, nitekim örnek 8.2'de negatif işarete sahiptir. - 1.67).

Kritik değerler MINITAB programı ile kolaylıkla hesaplanabilmektedir. Bu değerleri hesaplayabilmek için, “Calc-Probability distributions” (Hesapla-Olasılık Dağılımları) menüsünden “Normal” seçeneğine tıklayınız. Karşınıza çıkacak olan diyalog ekranında “inverse cumulative probability” (ters kümülatif olasılık) seçeneğine tıklayınız ve “input constant” (veri sabiti) bölümüne 0.05 değerini giriniz. (anlamlılık düzeyi $\alpha = 0.05$).



Şekil 8.17

“OK” tuşuna bastığınızda Şekil 8.18’deki çıktıyı elde edebilirsiniz.

Inverse Cumulative Distribution Function		
Normal with mean = 0 and standard deviation = 1,00000		
P (X <= x)	x	
0,0500	-1,6449	

Şekil 8.18

Elde etmiş olduğumuz kritik değer normal dağılımda sol kuyruğa ait bir olasılık değeridir. Sağ kuyruğa ait olan olasılık değerini hesaplamak için “input constant” bölümüne 0,05 yerine 0.95 (1-0,05) değerini girmeniz gerekmektedir.

Dolayısıyla

Inverse Cumulative Distribution Function		
Normal with mean = 0 and standard deviation = 1,00000		
P (X <= x)	x	
0,9500	1,6449	

Şekil 8.19

çıktısı elde edilebilir.

Alternatif Hipotezin Seçimi: H_A

Şuana kadar ele aldığımız örneklerde sadece sıfır hipotezlerine deyindik. Herhangi bir sıfır hipotezini, kabul veya reddettiğimizde alternatif hipotez hakkında da bir yorumda bulunmuş olmaktadır. Bu bağlamda örnek 8.1'deki sıfır hipotezini tekrar ele alalım.

$$H_0 : \mu = 12$$

Sıfır hipotezimize ait 3 farklı alternatif vardır.

$$H_A : \mu \neq 12$$

$$H_A : \mu > 12$$

$$H_A : \mu < 12$$

$\mu = 12$ önermesinin karşısında bulunan üç farklı durum alternatif olarak belirtilmiştir. Dolayısıyla, araştırmacılar bu üç alternatiften birini seçmek zorundadırlar.

Hipotez Testinin Kurulumu

Hipotez testi kurulurken önce sıfır hipotezi sonra alternatif hipotezi belirlenmektedir. Sıfır hipotezi her zaman tek bir değere eşit olmaktadır. Genel olarak sıfır hipotezini ifade edecek olursak:

$$H_0 : \mu = \mu_0$$

μ_0 burada anakütle ortalamasının μ eşit olduğu varsayılan değeri temsil etmektedir (bir önceki örneğimizde bu değer 12 idi).

Sıfır hipotezini belirlediğimize göre, amacımız doğrultusunda alternatif hipotezimizi belirleyebiliriz. Bu hipotezin belirlenmesinde üç farklı seçenek vardır.

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

Yukarıdaki alternatif hipotezini, anakütle ortalamasının belli bir değerden μ_0 farklı olup olmadığını test etmek için kullanmaktayız. Anakütle parametresinin, bu belirli değerden büyük veya küçük olup olmaması ile ilgilenmeyiz. Bizim için önemli olan anakütle parametresinden farklı olup olmamasıdır. Bu tarz hipotez testlerine çift taraflı test denilmektedir.

$$H_A : \mu < \mu_0$$

Bu durumda ise, anakütle ortalamasının belli bir değerden μ_0 küçük olup olmadığını test etmek için kullanmaktayız. Bu tarz hipotez testlerine tek taraflı test denilmektedir.

$$H_A : \mu > \mu_0$$

Bu durumda ise, anakütle ortalamasının belli bir değerden μ_0 büyük olup olmadığını test etmek için kullanmaktayız. Bu tarz hipotez testlerine de tek taraflı test denilmektedir.

Dolayısıyla, hipotez testini yapılandırabilmek için öncelikle sıfır ve alternatif hipotezleri belirlemeli ve bu doğrultuda testimizin tek taraflı mı çift taraflı mı olup olmadığına karar vermeliyiz.

Alternatif Hipotezin p-değeri ve kritik değer Yaklaşımlarına Etkisi

Hipotez testimize ait olan belirlenmiş alternatif hipotez, hem p-değeri hem de kritik değer yaklaşımlarını etkilemektedir. Şimdi bu etkinin nasıl gerçekleştiğini incelemek için tekrar Örnek 8.2 ve 8.3 ele alalım.

Örnek 8.6: Örnek 8.3'te hipotez testimize ait olan p-değerini 0.37 olarak hesaplamıştık. Bu olasılık değeri hesaplanırken, şekil 8.10'da olduğu gibi *sadece dağılımın sağ kuyruğundaki olasılık değerini ele almıştık*. Bir başka deyişle, anakütle ortalaması 10 iken örneklem ortalamasının 10.33 veya daha büyük bir değer alma olasılığını hesaplamış olduk. *Yalnız burada dikkat edileceği üzere dağılımın sol kuyruğunda yer alan olasılık değerini hesaplamadık*. Denklem (8.5)'te belirtilen sıfır hipotezine karşı olarak, alternatif hipotezimizi aşağıdaki gibi belirleyecek olursak:

$$H_A : \mu > 10 \quad (8.6)$$

Bu alternatif hipotez bize, örnek 8.3'te yeralan hipotez testinin *sağ taraflı* bir test olduğunu belirtmektedir. Anlamlılık düzeylerini, $\alpha = 0.01$, 0.05 , veya 0.1 olarak ele aldığımızda, sıfır hipotezimi reddedemeyiz veya alternatif hipotezimizi $\mu > 10$ reddedebiliriz. Fakat elde etmiş olduğumuz olasılık değeri (p-değeri) dağılımın sağ kuyruğunu içermemektedir. Bir başka ifade ile, sıfır hipotezini kabul etme kararını verdiğimizde, bu karar aynı zamanda alternatif hipotezimizi de reddetme anlamına gelmektedir. Alternatif hipotezimizi denklem (8.6)'daki gibi belirlediğimizde, anakütle ortalamasının 10'dan küçük olma ihtimalini $\mu < 10$ gözardı etmiş oluruz. Şimdi bu durumu da içerecek şekilde sıfır hipotezimizi tekrar belirleyelim.

$$H_0 : \mu \leq 10 \quad (8.7)$$

alternatif olarak

$$H_A : \mu > 10 \quad (8.8)$$

Hipotez testimiz tamamiyle *sağ taraflı bir testtir*. Dolayısıyla, sıfır hipotezini reddettiğimiz takdirde, doğru hipotezi reddetme riskinde artık $\mu < 10$ durumu yeralmamaktadır. Kabul ettiğimizde ise anakütle ortalamasının 10 değerine eşit veya büyük olduğu sonucuna varırız.

Örnek 8.7: Örnek 8.2'de hipotez testimize ait olan p-değerini 0.046 olarak hesaplamıştık. Bu olasılık değeri hesaplanırken, şekil 8.8'de olduğu gibi *sadece dağılımın sol kuyruğundaki olasılık değerini ele almıştık*. Bir başka deyişle, anakütle ortalaması 10 iken örneklem ortalamasının 10.33 veya daha küçük bir değer alma olasılığını hesaplamış olduk. *Yalnız burada dikkat edileceği üzere dağılımın sağ kuyruğunda yer alan olasılık değerini hesaplamadık*. Denklem (8.4)'te belirtilen sıfır hipotezine karşı olarak, alternatif hipotezimizi aşağıdaki gibi belirleyecek olursak:

$$H_A : \mu < 12 \quad (8.9)$$

Bu alternatif hipotez bize, örnek 8.2'de yeralan hipotez testinin *sol taraflı* bir test olduğunu belirtmektedir. Anlamlılık düzeylerini, $\alpha = 0.05$, veya 0.1 olarak ele aldığımızda, sıfır hipotezimi reddederiz. Fakat elde etmiş olduğumuz olasılık değeri (p-değeri) dağılımın sol kuyruğunu

içermemektedir. Bir başka ifade ile, sıfır hipotezini reddetme kararını verdiğimizde, bu karar aynı zamanda alternatif hipotezimizi de kabul etme (reddetmeme)¹⁵ anlamına gelmektedir. Alternatif hipotezimizi denklem (8.9)'daki gibi belirlediğimizde, anakütle ortalamasının 12'den büyük olma ihtimalini $\mu > 12$ gözardı etmiş oluruz. Şimdi bu durumu da içerecek şekilde sıfır hipotezimizi tekrar belirleyelim.

$$H_0 : \mu \geq 12 \quad (8.10)$$

alternatif olarak

$$H_A : \mu < 12 \quad (8.11)$$

Hipotez testimiz tamamiyle *sol taraflı bir testtir*. Dolayısıyla, sıfır hipotezini reddettiğimiz takdirde, doğru hipotezi reddetme riskinde artık $\mu > 12$ durumu yeralmamaktadır. Kabul ettiğimizde ise anakütle ortalamasının 12 değerine eşit veya büyük olduğu sonucuna varırız.

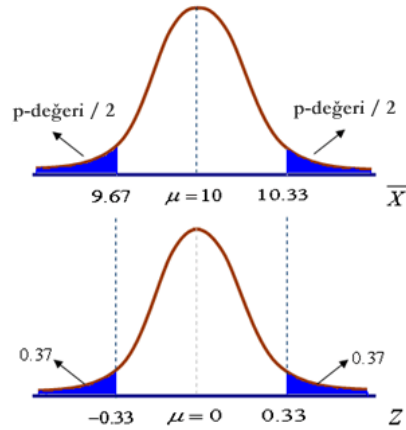
Örnek 8.8: Şimdiye kadar örneklerimizde hep tek taraflı hipotezleri ele aldık. Şimdiki örneğimiz ile birlikte *çift taraflı test* kuralını nasıl oluşturduğumuzu inceleyeceğiz. Örnek 8.2'deki sıfır hipotezimizi ele alacak olursak:

$$H_0 : \mu = 10$$

alternatif olarak

$$H_A : \mu \neq 10$$

Hipotezlerini kurduğumuzda, alternatif hipotezimizin açık bir şekilde $\mu < 10$ iki durumu ihtiva ettiğini görebiliriz. Dolayısıyla artık çift taraflı bir test ile karşı karşıya bulunmaktayız. Dağılımın her iki kuyruğunda da yer alan olasılık değerlerini dikkate almamız gerekmektedir. Bu iki olasılık değerini de dikkate alabilmek için 10.33 tahmin edicisi gibi dağılımın diğer kuyruğuna tekabül eden simetrik bir tahmin edici daha hesaplanmalıdır.



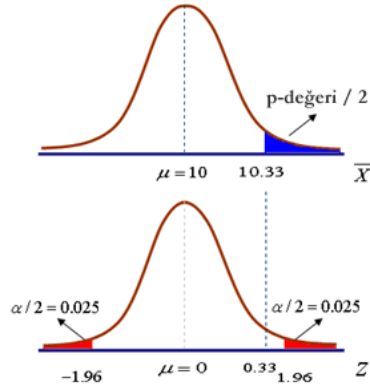
Şekil 8.20

¹⁵ Her ne kadar “kabul etme” yerine terminolojik olarak “reddetmeme” deyimini kullanmamız gerekse de bazı yerlerde öğrencinin anlamasını kolaylaştırmak için bu terminolojik detayın aşıldığını görebilirsiniz.

Şekil 8.20, 10.33 değerinin simetrik tahmin edicisinin de gösterilmesi eklenerek şekil 8.12'den elde edilmiştir. Bu simetrik tahmin edicinin değeri 9.67'dir. Bu durumda p-değeri daha önceki değerin tam iki katı bir değer almaktadır. Dolayısıyla p-değeri = $0.37 + 0.37 = 0.74$ değerine eşit olur. Bu bağlamda, p-değerini kullanarak sıfır hipotezini $\mu = 10$ kabul ederiz (reddedemeyiz).

Aynı durumu kritik değer yaklaşımında ele alalım. Analizimize anlamlılık düzeyini $\alpha = 0.05$ olarak başlayalım. α değerinin dağılımın her iki tarafına eşit olarak paylaşılması gerektiğinden, her kuyruktaki maksimum yapabileceğimiz Tip I Hata miktarı $\alpha / 2 = 0.025$ olmaktadır. Sonuç olarak toplamda maksimum yapabileceğimiz Tip I Hata miktarı yine α değerine eşit olacaktır.

Daha önceki bölümlerden de hatırlayabileceğiniz gibi, Normal Dağılım Tablosu'ndan her iki kuyruk içinde 0.025 olasılık değerine denk gelen z değerleri ± 1.96 'dır ($P(z > 1.96) = P(z < -1.96) = 0.025$). Grafiksel olarak gösterecek olursak:



Şekil 8.21

Örneklem ortalamasına $\bar{x}$ tekabül eden z değeri, $z_{\alpha/2}$ kritik değerinden küçük bir değer alıyorsa ve böylece reddedilme alanının içerisinde yer almıyorsa, sıfır hipotezi kabul edilir (reddedilemez).

Bu kuralı genelleştirmek için, örneklem ortalamasına $\bar{x}$ tekabül eden z değerini z skoru olarak (örneğimizde z-skoru = 0.33), ve α anlamlılık düzeyine tekabül eden z değerlerini de $z_{\alpha/2}$ olarak adlandıralım. (örneğimizde $z_{\alpha/2} = \pm 1.96$). Genel olarak sıfır hipotezimizi aşağıdaki gibi belirleyelim:

$$H_0 : \mu = \mu_0$$

alternatif olarak

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

μ_0 ifadesi, anakütle ortalamasının μ eşit olduğu varsayılan sayısal değeri temsil etmektedir. α anlamlılık düzeyinde hipotezimiz için reddetme (veya reddetmeme) kuralı aşağıdaki gibidir:

Hipotez Testinde Kritik Değer Yaklaşımı:

Eğer z-istatistik skoru $> z_{\alpha/2}$ veya z-istatistik skoru $< -z_{\alpha/2}$; H_0 reddedilir

Eğer z-istatistik skoru $< z_{\alpha/2}$ veya z-istatistik skoru $< -z_{\alpha/2}$; H_0 **reddedilemez**

Test Prosedürünün Özeti:

Öncelikle anakütle ortalaması hakkındaki hipotezimizi belirlememiz gerekmektedir. Bu hipotez belirlenirken önümüzde üç olasılık yer almaktadır: μ değeri belirlenmiş μ_0 değerinden büyük bir değer olabilir, μ değeri belirlenmiş μ_0 değerinden küçük bir değer olabilir, μ değeri belirlenmiş μ_0 değerinden farklı bir değer olabilir. Bu üç durumdan birine karar verdikten sonra hipotezimizi sembolik olarak ifade edebiliriz. Eğer μ değerinin belirlenmiş μ_0 değerinden büyük bir değer olduğunu iddia ediyorsak, alternatif tezimizi de belirleyerek hipotezimizi aşağıdaki gibi yazabiliriz.

$$H_0 : \mu \leq \mu_0$$

$$H_A : \mu > \mu_0$$

Eğer μ değerinin belirlenmiş μ_0 değerinden küçük bir değer olduğunu iddia ediyorsak, alternatif tezimizi de belirleyerek hipotezimizi aşağıdaki gibi yazabiliriz.

$$H_0 : \mu \leq \mu_0$$

$$H_A : \mu < \mu_0$$

Eğer μ değerinin belirlenmiş μ_0 değerinden farklı bir değer olduğunu iddia ediyorsak, alternatif tezimizi de belirleyerek hipotezimizi aşağıdaki gibi yazabiliriz.

$$H_0 : \mu = \mu_0$$

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

Bir sonraki adımda, aşağıdaki kurallardan uygun olanı teste adapte edilmelidir.

Sağ Taraflı Testler:

$$H_0 : \mu \leq \mu_0$$

$$H_A : \mu > \mu_0$$

Öncelikle z skoru hesaplanmalıdır

$$z = \frac{\bar{x} - \mu}{\sigma_{\bar{x}}}$$

Burada $\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{n}}$ değeri örneklem ortalamasının standart sapmasını temsil etmektedir.

Şimdi ise bu z-skoru kullanılarak, z dağılım tablosundan (veya MINITAB programından) p-değeri hesaplanmalı ve p-değeri kuralı uygulanmalıdır.

P-değeri kuralı (α anlamlılık düzeyinde):

Eğer p-değeri $< \alpha$; H_0 **reddedilir**

Eğer p-değeri $> \alpha$; H_0 **reddedilemez**

Alternatif olarak z-skoru kullanılarak kritik deęer kuralı uygulanabilir (z kritik deęerini z_{α} bulmak için yine z daęılım tablosundan veya MINITAB programından faydalanabilirsiniz)

Kritik deęer kuralı (α anlamlılık düzeyinde):

z- istatistik skoru $> z_{\alpha}$; H_0 *reddedilir*

z- istatistik skoru $< z_{\alpha}$; H_0 *reddedilemez*

Sol taraflı test:

$$H_0 : \mu \geq \mu_0$$

$$H_A : \mu < \mu_0$$

Öncelikle z skoru hesaplanmalıdır

$$z = \frac{\bar{x} - \mu}{\sigma_{\bar{x}}}$$

Burada $\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{n}}$ deęeri örneklem ortalamasının standart sapmasını temsil etmektedir.

Şimdi ise bu z-skoru kullanılarak, z daęılım tablosundan (veya MINITAB programından) p-deęeri hesaplanmalı ve p-deęeri kuralı uygulanmalıdır.

P-deęeri kuralı (α anlamlılık düzeyinde):

Eęer p-deęeri $< \alpha$; H_0 *reddedilir*

Eęer p-deęeri $> \alpha$; H_0 *reddedilemez*

Alternatif olarak z-skoru kullanılarak kritik deęer kuralı uygulanabilir (z kritik deęerini $-z_{\alpha}$ bulmak için yine z daęılım tablosundan veya MINITAB programından faydalanabilirsiniz)

Kritik deęer kuralı (α anlamlılık düzeyinde):

z- istatistik skoru $< -z_{\alpha}$; H_0 *reddedilir*

z- istatistik skoru $> -z_{\alpha}$; H_0 *reddedilemez*

Çift Taraflı Testler:

$$H_0 : \mu = \mu_0$$

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

Öncelikle z skoru hesaplanmalıdır

$$z = \frac{\bar{x} - \mu}{\sigma_{\bar{x}}}$$

Burada $\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{n}}$ deęeri örneklem ortalamasının standart sapmasını temsil etmektedir.

Şimdi ise bu z-skoru kullanılarak, z daęılım tablosundan (veya MINITAB programından) p-deęeri hesaplanmalı ve p-deęeri kuralı uygulanmalıdır.

P-değeri kuralı (α anlamlılık düzeyinde):Eğer p-değeri $< \alpha$; H_0 **reddedilir**Eğer p-değeri $> \alpha$; H_0 **reddedilemez**

Alternatif olarak z-skoru kullanılarak kritik değer kuralı uygulanabilir (z kritik değerlerini $\pm z_{\alpha/2}$ bulmak için yine z dağılım tablosundan veya MINITAB programından faydalanabilirsiniz)

Kritik değer kuralı (α anlamlılık düzeyinde):z- istatistik skoru $> z_{\alpha/2}$ veya z- istatistik skoru $< -z_{\alpha/2}$; H_0 **reddedilir**z- istatistik skoru $< z_{\alpha/2}$ veya z- istatistik skoru $> -z_{\alpha/2}$; H_0 **reddedilemez**

Şuana kadar yapmış olduğumuz prosedürlerde üzerinde durmadığımız husus uygun anlamlılık düzeyinin α seçimidir. Bu konuya Tip I ve Tip II Hata'nın anlatıldığı bölümde daha detaylı olarak değineceğiz.

Tablo 8.2'de farklı anlamlılık düzeyleri için $z_{\alpha/2}$ ve z_{α} (kritik değerleri) ifade edilmektedir.

Tablo 8.2 Standart Normal Dağılımın Kritik Değerleri

Anlamlılık Düzeyi	α	z_{α} (Kritik Değer)	$\alpha / 2$	$z_{\alpha/2}$ (Kritik Değer)
10%	0.10	1.2816	0.05	1.6449
5%	0.05	1.6449	0.025	1.96
2.5%	0.025	1.96	0.0125	2.24165 ¹⁶
1%	0.01	2.3263	0.005	2.5758

Örnek 8.9: Örnek 8.3'te anakütle ortalamasının 12'den küçük olduğunu iddia etmiştik. Dolayısıyla *sol taraflı* bir test gerçekleştirmiştik. Şimdi ise anakütle ortalamasının 12'den büyük olduğunu iddia edelim. Bu iddia doğrultusunda hipotezlerimiz:

$$H_0 : \mu \leq 12$$

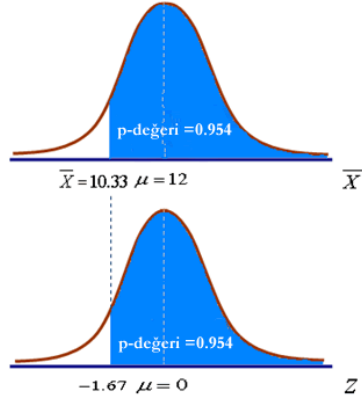
$$H_A : \mu > 12$$

Standart prosedürümüzü kullanarak z-skorunu hesaplayacak olursak:

$$z = \frac{10.33 - 12}{1} = -1.67$$

Testimiz artık *sağ taraflı* bir test halini almıştır. P-değeri anakütle ortalamasının sağında kalan alana eşittir. $P(\bar{X} > 10.33) = P(Z > -1.67) = 0.954$. Grafikselsel olarak gösterecek olursak:

¹⁶ Doğrusal enterpolasyon ile bulunmuştur

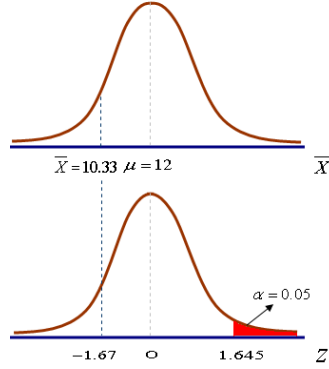


Şekil 8.22

Örnek 8.3'te aynı anakütle parametresi ile ($\mu = 12$) sol taraflı test gerçekleştirmiştik ve p-değerini örneklem ortalamasının solunda kalan alan olarak hesaplamıştık. $P(\bar{X} < 10.33) = P(Z < -1.67) = 0.046$ (bkz Şekil 8.8).

Bu örneğimizde ise, p-değeri (=0.954) seçilen herhangi bir anlamlılık düzeyinden büyük olduğundan sıfır hipotezini kabul ederiz (reddedemeyiz).

Aynı sonuca kritik değer yaklaşımı ile de ulaşabiliriz. Testimiz sağ taraflı olduğundan dolayı anlamlılık düzeyini α olarak seçebiliriz ve $\alpha = 0.05$ olarak belirleyebiliriz



Şekil 8.23

Dolayısıyla -1.67 reddedilme alanının dışında kaldığından yine sıfır hipotezimizi kabul ederiz (reddedemeyiz).

Örnek 8.10: Örnek 8.2'de anakütle ortalamasının 10'dan büyük olduğunu iddia etmiştik. Dolayısıyla *sağ taraflı* bir test gerçekleştirmiştik. Şimdi ise anakütle ortalamasının 10'dan küçük olduğunu iddia edelim. Bu iddia doğrultusunda hipotezlerimiz:

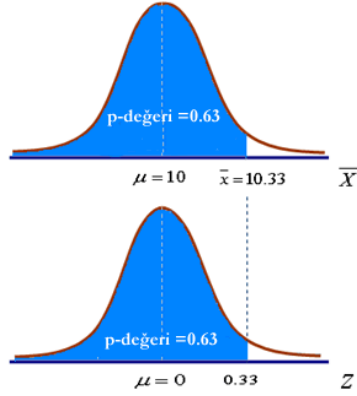
$$H_0 : \mu \geq 10$$

$$H_A : \mu < 10$$

Standart prosedürümüzü kullanarak z-skorunu hesaplayacak olursak:

$$z = \frac{10.33 - 10}{1} = 0.33$$

Testimiz artık *sol taraflı* bir test halini almıştır. P-değeri anakütle ortalamasının solunda kalan alana eşittir. $P(\bar{X} < 10.33) = P(Z < 0.33) = 0.63$. Grafiksel olarak gösterecek olursak:

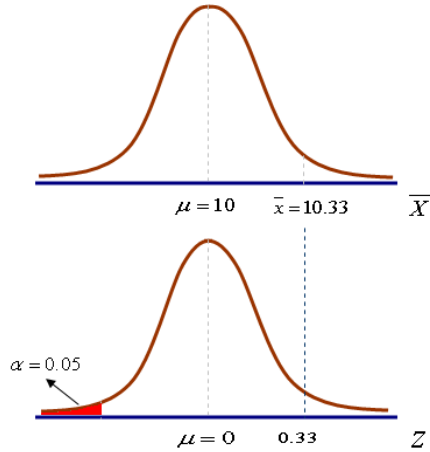


Şekil 8.24

Örnek 8.2’de aynı anakütle parametresi ile ($\mu = 10$) sağ taraflı test gerçekleştirmiştik ve p-değerini örneklem ortalamasının sağında kalan alan olarak hesaplamıştık. $P(\bar{X} > 10.33) = P(Z > 0.33) = 0.37$ (bkz Şekil 8.8).

Bu örneğimizde ise, p-değeri (=0.63) seçilen herhangi bir anlamlılık düzeyinden büyük olduğundan sıfır hipotezini kabul ederiz (reddedemeyiz).

Aynı sonuca kritik değer yaklaşımı ile de ulaşabiliriz. Testimiz sol taraflı olduğundan dolayı anlamlılık düzeyini α olarak seçebiliriz ve $\alpha = 0.05$ olarak belirleyebiliriz

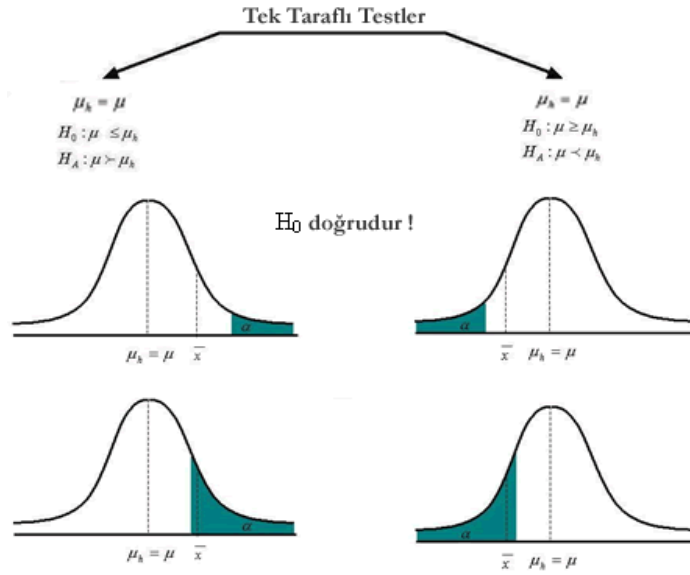


Şekil 8.25

Dolayısıyla 0.33 reddedilme alanının dışında kaldığından yine sıfır hipotezimizi kabul ederiz (reddedemeyiz).

Tip I ve Tip II Hataları

Anlamlılık düzeyleri belirtilirken neden sadece 1%, 5%, 10% gibi değerlerden bahsettik? Başka rakamlar anlamlılık düzeyi olarak ele alınamaz mı? Daha öncede kısaca deyindiğimiz gibi bu olasılık değerleri, **doğru hipotezi reddederek** maksimum yapabileceğimiz hata miktarını **Tip I Hata**'yı ifade etmektedir. **Yanlış olan bir hipotezi kabul ettiğimizde** ise **Tip II Hata** işlemiş olmaktadır. Tip II Hata'yı β (beta olarak okunur) ile ifade etmekteyiz. α ve β değerleri arasında bir ikilem yaşanmaktadır. α olasılık değerini arttırdığımızda, β olasılık değerini düşürürüz. Veya, α olasılık değerini azalttığımızda, β olasılık değerini arttırırız. Bu ikilemi grafiksel olarak ifade edecek olursak:



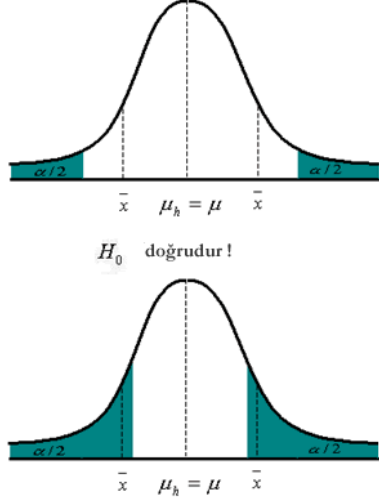
Şekil 8. 26 A

Çift Taraflı Testler

$$\mu_h = \mu$$

$$H_0 : \mu = \mu_h$$

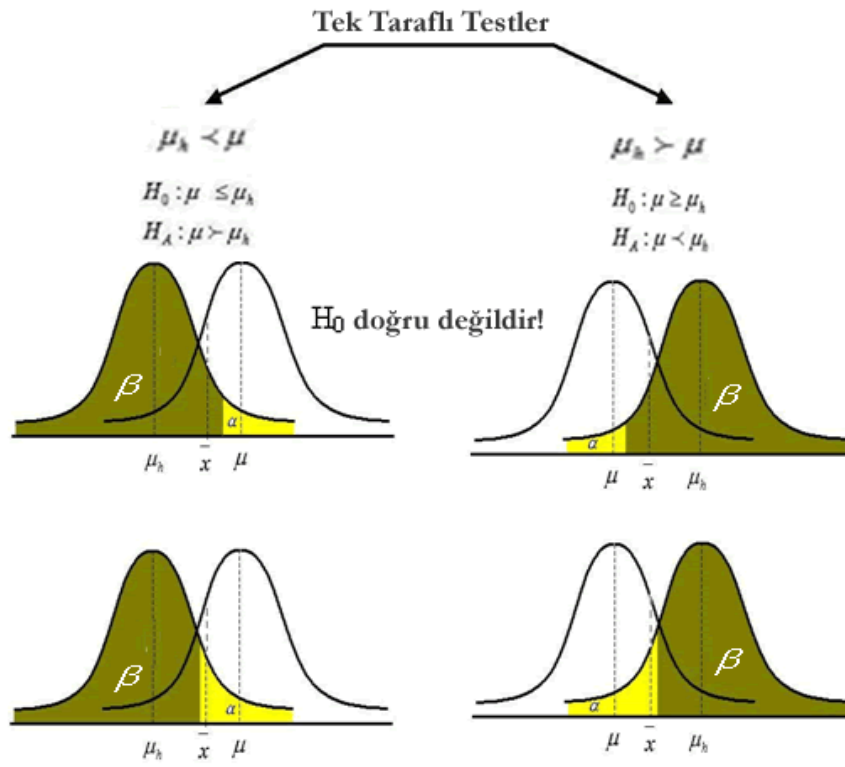
$$H_A : \mu \neq \mu_h$$



Şekil 8.26 B

Şekil 8.26'daki A ve B panellerinde, hipotetik anakütle ortalamasının (μ_h), gerçek anakütle ortalamasına (μ) eşit olduğu gözlenmektedir. Dolayısıyla, $H_0 : \mu = \mu_h$ hipotezini reddettiğimizde gerçekten $\mu = \mu_h$ eşitliği söz konusu ise yanlışlıkla Tip I Hata yapmış oluruz. Aynı şekilde, eğer H_0 hipotezini kabul edersek (reddetmezsek) doğru sonuca varmış oluruz. Şekilde de gösterildiği gibi bu durumda α değerini yüksek bir değer olarak alırsak, Tip I Hata olasılığını da arttırmış oluruz. Bu durumda sıfır hipotezi doğru bir şekilde ifade edildiğinden Tip II Hata yapma imkanımız bulunmamaktadır.

Bu yüzden şekil 8.26'da β gösterilmemiştir.



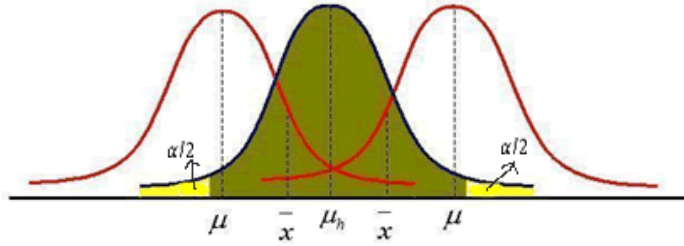
Şekil 8.27 A

Çift Taraflı Testler

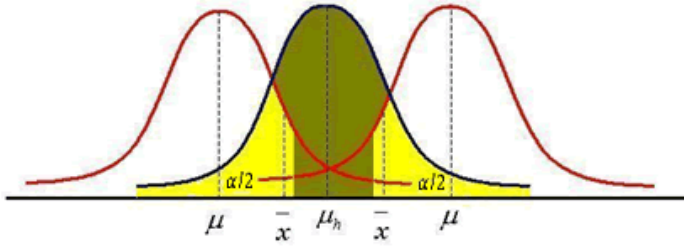
$$\mu_h \neq \mu$$

$$H_0 : \mu = \mu_h$$

$$H_A : \mu \neq \mu_h$$



H_0 doğru değildir!



Şekil 8.27 B

Şekil 8.27'deki A ve B panellerinde, hipotetik anakütle ortalamasının (μ_h), gerçek anakütle ortalamasından (μ) farklı olduğu gözlenmektedir. Dolayısıyla, μ_h değeri anakütle ortalamasından μ büyük bir değer de alabilir, küçük bir değer de alabilir. $H_0 : \mu = \mu_h$ hipotezini reddettiğimizde doğru sonuca varmış oluruz. Aynı şekilde, eğer H_0 hipotezini kabul edersek (reddetmezsek), gerçekte hipotetik anakütle ortalaması ile gerçek anakütle ortalaması birbirinden farklı değerlere sahip olduğundan Tip II Hata yapmış oluruz. Şekilde de gösterildiği gibi bu durumda α değerini yüksek bir değer olarak alırsak bu durumda tehlikeli bir şekilde Tip II Hata yapma olasılığı azalmaktadır.

Sıfır hipotezi doğru bir şekilde ifade edildiğinden ve doğru hipotez reddedilmediğinden, bu durumda Tip I Hata yapma olasılığımız bulunmamaktadır. Şekil 8.27'te α değerleri gösterilmesine rağmen, aslen Tip I hata yapma olasılığımız bulunmamaktadır, buradaki α değerleri anlamlılık düzeyini ifade etmektedir. Dolayısıyla, α anlamlılık düzeyine karar verirken iki tip olasılık değerini dikkate almamız gerekmektedir. Dikkat etmemiz gereken hususlardan birincisi, yanlış bir hipotezi kabul etme olasılığını β maksimize ederek, doğru bir hipotezi reddetme olasılığımızı α minimize etme durumu; ikincisi ise yanlış bir hipotezi kabul etme olasılığımızı β minimize ederek, doğru bir hipotezi reddetme olasılığımızı α maksimize etme hususudur. Dolayısıyla, eğer sıfır hipotezinin daha çok doğru olabileceğine inanıyorsanız, uygulamanız gereken ilk strateji (α değerini olabildiğince küçük tutmaktır) gerçek tehlikeyi doğru hipotezi reddetmeyi minimize

etmektedir. Bir diğer taraftan, eğer sıfır hipotezinin doğruluğundan tam olarak emin değilseniz, ikinci strateji olarak yanlış hipotezi kabul (reddetmeyerek) etmeyi minimize etmeniz (olası miktarda yüksek α değeri, dolayısıyla β değerini düşürmek) uygun olacaktır. Genel olarak istatistikçiler, Tip I Hata'nın Tip II Hata'dan daha önemli olduğuna inanmaktadırlar. Dolayısıyla olabildiğince α değerini küçük tutmaya çalışmaktadırlar. Sıfır hipotezini masum bir suçlu olarak ele alırsak; Tip I Hata, masum olan birinin suçlanması anlamına gelirken; Tip II Hata ise suçlu bir kişinin serbest bırakılması anlamına gelmektedir. Genel olarak bakıldığında masum birinin suçlanması daha önemli olduğundan Tip I Hata, Tip II Hata'ya oranla daha önemli olarak gözükmemektedir.

Bir işadınının piyasaya girip girmemesi ile alakalı olarak bir karar vermesi gerekmektedir. Alternatif olarak başka bir piyasaya girdiği takdirde bu işadınının kar oranı 5% olacaktır. Kararını aşağıdaki hipotez testine göre belirleyecektir:

$$H_0 : \mu \leq 0.05$$

$$H_A : \mu > 0.05$$

Sıfır hipotezimizi reddettiğimiz takdirde, diğer piyasaya girmeyi düşünmektedir. H_0 doğru ise ve reddedilirse, Tip I Hata işlenmektedir. Bu hata sonucunda, diğer piyasaya girer ve zarar eder. H_0 doğru olmamasına rağmen kabul edilirse, Tip II Hata işlenir ve yüzde 5'ten fazla kar etme olasılığı kaçırılmış olur.

Tablo 8.3 hipotez testlerinde Tip I Hata ve Tip II Hata'yı özet bir şekilde inceleyebilirsiniz.

Table 8.3: Hipotez Testlerinde Tip I ve Tip II Hataları

Sonuç		Durum	
		H_0 doğru ise Doğru Sonuç	H_0 yanlış ise Tip II Hata
	H_0 Kabul edilir		
	H_0 Reddedilir	Tip I Hata	Doğru Sonuç

Hipotez Testinin Gücü

H_0 hipotezini, doğru olduğu halde reddetme olasılığımızı minimize edebilmemiz için dikkate alacağımız son kriter, testin gücüdür. H_0 hipotezi yanlış olduğu takdirde, bu hipotezi kabul etme olasılığımız ne kadar küçük olursa testimiz o kadar güçlü olmaktadır. Dolayısıyla testimizin gücü, H_0 hipotezi yanlış olduğu halde bu hipotezi reddetme olasılığımız ile ölçülmektedir. Tip II Hata olasılığı, yanlış olan bir hipotezin kabul edilme olasılığını temsil ettiğinden, testin gücü de $1 - P(\text{Tip II Hata})$ olarak hesaplanabilir.

Ortalama Hakkındaki İddiamızın Testi: Örneklem Boyutunun Büyük Olduğu Durum

Bir önceki bölümlerde de bahsettiğimiz gibi hipotez testinin amacı örneklem nokta tahmininin (örneklem ortalaması gibi) iddia edilen değerden istatistiksel olarak anlamlı derecede farklı olup olmadığına karar vermektir. Geçerli olan örneklem istatistiği ($\bar{X}$), test istatistiğine dönüştürülür ve kritik değerle karşılaştırılır.

Örneklerimizde daha önce hep anakütle standart sapmasının bilindiğini (σ biliniyor), ve anakütlenin normal dağıldığını varsaydık. Normal olarak test prosedürümüzde standart sapmanın bilinip bilinmediğine göre veya örneklem boyutlarına göre değişiklikler meydana gelecektir. Öncelikle büyük örneklemelerin olduğu durumu ele alalım ve bu durumda standart sapmanın bilinip bilinmediği durumları inceleyelim.

σ biliniyorsa

Örneklem boyutunun büyük olduğu durumlarda yani $n \geq 30$ olduğu durumlarda, ana dağılımın şekline bağımsız olarak yine merkezi limit teoremi kullanılır. Bu yüzden doğrudan $\bar{X}$ ların örneklem dağılımının normal olarak dağıldığını varsayabiliriz. Daha önceki durumlardan farklı olarak, $n \geq 30$ olduğundan ve bu durumda MLT'te göre örneklem dağılımının normal olduğunu garantilediğinden, ayrıca bir normallik varsayımı yapmamıza gerek yoktur.

Örnek 8.11: Bir kahve üreticisi ürettiği kahve kavanozlarının etiketlerinde iddia ettiğine göre her kahve kavanozu kahvenin 3 Lirasını içermektedir. Araştırmacı ise bu iddianın geçerliliğini test etmek istemektedir. Öncelikle sıfır ve alternatif hipotezler kurulacaktır. Sıfır ve alternatif hipotezler aşağıdaki gibi ifade edilmiştir:

$$H_0 : \mu \geq 3$$

$$H_A : \mu < 3$$

Açıkça görüldüğü üzere tek taraflı bir hipotez testidir. Daha önceki çalışmalar anakütle standart sapmasının 0.20 lira olduğunu açıklamışlardır. Araştırmacı 49 kahve kavanozunu örneklem olarak almıştır ve bu örneklemelerin ortalamasını da kavanoz başına 2.95 lira ($\bar{X} = 2.95$) olarak hesaplamıştır. Araştırmacı, Tip I hata yapma olasılığını azami $\alpha = 0.01$ olarak düşünmüştür. Reddetme alanının sınırı 1% lik anlamlılık düzeyindedir ve z tablosu kullanılarak, -2.33 olarak belirlenmiştir. Bu değer testimiz için kritik değerdir ve reddetme alanı ile reddetmeme alanını birbirinden ayırmaktadır. Eğer hesaplanan z skoru bu değerden küçükse, reddedilme alanında yer alır ve sıfır hipotezini alternatif hipoteze göre reddederiz.

Örneklem değerlerimizi (denklikte yerine koyarak z-skorunun değerine ulaşabiliriz.

$$z = \frac{\bar{X} - \mu}{\sigma_{\bar{X}}} = \frac{2.95 - 3}{0.2/\sqrt{49}} = -1.75$$

ve bunu z-skor istatistik diye düşünebiliriz. z-skoru $> -z_{0.01}$. Burada sıfır hipotezini reddedemeyiz. Bu yüzden üreticinin kavanoz etiketlerinde verdiği bilgiyi reddedecek yeterli istatistiksel kanıtı sahip değiliz.

Testimizin p değerini hesaplayacak $P(z > -1.75) = 0.0401$ ve anlamlılık düzeyi $\alpha = 0.01$ ile karşılaştıracak olursak, yine sıfır hipotezimizi reddetmememiz gerektiği sonucuna varırız.

σ bilinmiyorsa

Eğer σ^2 bilinmiyorsa, , örneklem dağılımının standart hatasını aşağıdaki formülü kullanarak hesaplayamayız.

$$\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{n}}$$

Dolayısıyla z-istatistik skorunu hesaplayamayız ve testimizi sonuçlandıramayız. Bu sorunun üstesinden gelebilmek için σ^2 değerinin bir tahmin edicisi olan örneklem varyansını s^2 kullanmamız gerekmektedir. Örneklem varyansını daha önceden de hatırlayacağınız üzere aşağıdaki gibi hesaplayabilirsiniz.

$$s^2 = \frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n-1} \quad (8.17)$$

Bu yüzden , örneklem dağılımının standart sapması

$$s_{\bar{x}} = \frac{s}{\sqrt{n}} \quad (8.18)$$

Analizimize daha önce olduğu gibi z skorunu, paydada $\sigma_{\bar{x}}$ yerine $s_{\bar{x}}$ ifadesini koyarak hesaplayarak devam edebilir miyiz?

$$z = \frac{\bar{x} - \mu}{s_{\bar{x}}} \quad (8.19)$$

Bu sorunun cevabı malesef olumsuzdur. Çünkü, denlem (8.19)'un paydasında gerçek değerinin $\sigma_{\bar{x}} = \sigma / \sqrt{n}$ yerine örneklem ortalamasının tahmin edicisi olarak $s_{\bar{x}} = s / \sqrt{n}$ ¹⁷ kullandığımızdan dolayı artık z nin standart normal dağılımını kullanamamaktayız.

$$\frac{\bar{x} - \mu}{s / \sqrt{n}} \quad (8.20)$$

Bu konunun nedenini biraz daha açacak olursak; $n \geq 30$ olduğundan $\bar{X}$ ın normal dağıldığını bilmemize rağmen formülde s gibi başka bir tahmin edici yer almaktadır. Peki s nasıl bir dağılıma sahiptir? s değişkeni Gamma adı verilen bir dağılıma sahiptir. Bizim burada sorunumuz denklem (8.20)'de normal olarak dağılan bir değişken ile Gamma dağılımı¹⁸ ile dağılan

¹⁷ Örneklem ortalamasının standart sapmasının $\sigma / \sqrt{n}$, tahmin edicisi olduğundan, s değişkeni de σ standart sapmasının tahmin edicisi olduğundan bu şekilde ifade edilmektedir.

¹⁸Denklem (8.20)'nin pay ve paydasını σ değerine bölerek ifadeyi tekrar düzenlediğimizde $\frac{(\bar{x} - \mu)\sqrt{n} / \sigma}{\sqrt{(n-1)s^2 / (n-1)\sigma^2}}$ ifadesini elde edebiliriz. Bu ifadenin payı, standart normal olarak dağılan z değişkenini ifade ederken, paydası da

bir değişkenin oranının nasıl dağılacığı hususudur. Denklem (8.20)'de bahsedilen bu oran t dağılımı ile dağılmaktadır. Dolayısıyla standart sapmanın σ bilinmediği durumlarda dönüşüm formülümüz t dağılımı ile dağılmaktadır ve dönüşüm formülü aşağıdaki gibidir.

$$t_{n-1} = \frac{\bar{x} - \mu}{s / \sqrt{n}}$$

t dağılımı, normal dağılım gibi özel bir dağılımdır. t dağılımında Bölüm 6'da da deyindiğimiz gibi serbestlik derecesi, $(n - 1)$ değerine eşittir.

Bu teste ait olan kritik değerler, serbestlik derecesi tarafından belirlenmektedir. Bu değerler z skoruna benzer niteliktedirler ancak z skorunun kritik değerlerinden temel olarak farklılaştığı husus bu değerlerin serbestlik derecesine göre belirlenmesidir. Bunun yanısıra sıfır ve alternatif hipotezler aynı şekilde belirlenmektedir.

Serbestlik derecesi arttığında (n'in artmasıyla) t dağılımı ile standart normal dağılım arasındaki fark küçülmektedir. Dolayısıyla büyük örneklem boyutlarında, t ve z değerleri arasında çok fazla fark görülmemektedir.

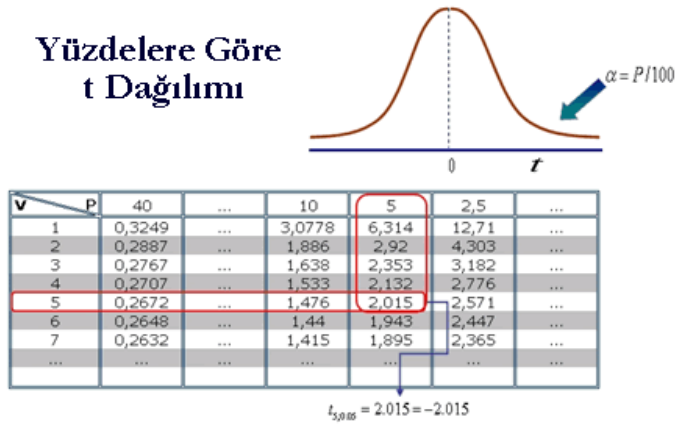
Tahmin edebileceğiniz gibi t değerlerinin de test kuralları z testininkilere benzemektedir. Farklılaşma, kritik değerlerin ve p-değerlerinin hesaplanmasında görülmektedir. Artık kritik değerler ve p-değerleri hesaplanırken t-dağılımı tablolarından faydalanmak gerekmektedir.

t Dağılımı Tabloları ve MINITAB Kullanılarak Kritik Değerlerin Bulunması

Minitab Uygulamaları

t dağılımına ait olan kritik değerleri $t_{df,\alpha}$ sembolü ile ifade etmekteyiz. Kritik değerlerin bulunabileceği t-dağılım tablosu Ek (Tablo 3.2)'te yer almaktadır. Bu tablo farklı olasılık değerlerini (α ları) ve bu değerlere her farklı serbestlik derecesi için tekabül eden t değerlerini göstermektedir.

Yüzdelere Göre t Dağılımı



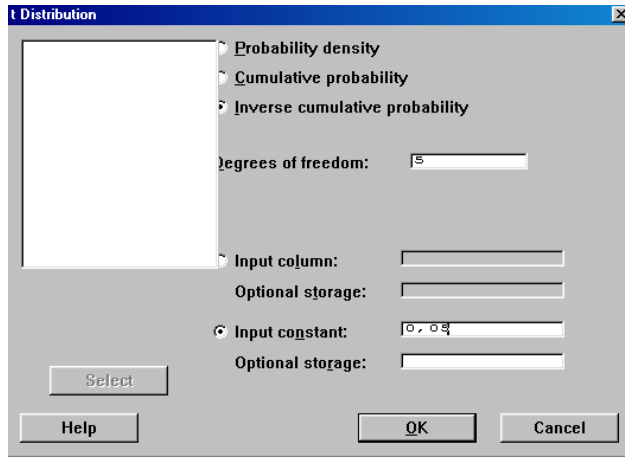
Şekil 8.28

Gamma dağılımı ile dağılan $(n-1)s^2 / \sigma^2$ ki-kare dağılımını göstermektedir. $\frac{N(0,1)}{\sqrt{\chi_{n-1}^2 / (n-1)}}$, elde edilen bu değişken

ise t dağılımı ile dağılmaktadır. $\frac{(\bar{x} - \mu)}{s / \sqrt{n}} \sim t_{n-1}$.

Tabloda da gösterildiği gibi $t_{5,0.05} = 2.015$ kritik değerini kolaylıkla elde edebiliriz. Bulduğumuz bu değer sağ taraflı bir test içindir, eğer sol taraflı bir test için kritik değere ihtiyacımız olsa idi, normal dağılımın simetrik özelliğinden faydanalarak bu kritik değeri -2.015 olarak bulacaktık.

t kritik değerleri MINITAB yardımıyla da kolaylıkla hesaplanabilmektedir. BU hesaplamayı gerçekleştirebilmek için “Calc-Probability distributions” bölümünden “t..” yi seçiniz. Karşınıza çıkacak olan diyalog ekranında “inverse cumulative probability” seçeneğine tıklayıp, “degrees of freedom” (serbestlik derecesi) bölümüne 5 değerini ve “input constant” bölümüne de 0.05 (anlamlılık düzeyi $\alpha = 0.05$) değerini giriniz.



Şekil 8.29

“OK” tuşuna tıkladığınızda aşağıdaki çıktıyı elde edebilirsiniz.

Inverse Cumulative Distribution Function	
Student's t distribution with 5 DF	
P (X <= x)	x
0,0500	-2,0150

Şekil 8.30

Elde ettiğimiz kritik değer sol taraflı bir test olduğundan, sağ taraflı bir teste ait bir kritik değeri elde etmek istersek “input constant” kutucuğuna 0.05 değerinin yerine 0.95 değerini girmemiz gerekmektedir.

Inverse Cumulative Distribution Function	
Student's t distribution with 5 DF	
P (X <= x)	x
0,9500	2,0150

Şekil 8.31

t Testi Prosedürünün Özeti

Sağ Taraflı Testler:

$$H_0 : \mu \leq \mu_0$$
$$H_A : \mu > \mu_0$$

Öncelikle t skoru hesaplanmalıdır

$$t = \frac{\bar{x} - \mu}{s_{\bar{x}}} \quad (8.21)$$

Burada $s_{\bar{x}} = \frac{s}{\sqrt{n}}$ ¹⁹ değeri örneklem ortalamasının standart sapmasını temsil etmektedir. Şimdi ise bu t-skoru kullanılarak, t dağılım tablosundan (veya MINITAB programından) p-değeri hesaplanmalı ve p-değeri kuralı uygulanmalıdır.

P-değeri kuralı (α anlamlılık düzeyinde):

Eğer p-değeri $< \alpha$; H_0 **reddedilir**

Eğer p-değeri $> \alpha$; H_0 **reddedilemez**

Alternatif olarak t-skoru kullanılarak kritik değer kuralı uygulanabilir (t kritik değerini $t_{n-1;\alpha}$ bulmak için yine t dağılım tablosundan veya MINITAB programından faydalanabilirsiniz)

Kritik değer kuralı (α anlamlılık düzeyinde):

t- istatistik skoru $> t_{n-1;\alpha}$; H_0 **reddedilir**

t- istatistik skoru $< t_{n-1;\alpha}$; H_0 **reddedilemez**

Sol taraflı test:

$$H_0 : \mu \geq \mu_0$$
$$H_A : \mu < \mu_0$$

Öncelikle t skoru hesaplanmalıdır

$$t = \frac{\bar{x} - \mu}{s_{\bar{x}}}$$

¹⁹ Ayrıca örneklem boyutunun, anakütle boyutuna oranla oluştuğunu ($n/N \leq 0.05$) ve sonlu bir anakütle olduğunu varsaymaktayız.

Burada $s_{\bar{x}} = \frac{s}{\sqrt{n}}$ değeri örneklem ortalamasının standart sapmasını temsil etmektedir. Şimdi ise bu t-skoru kullanılarak, t dağılım tablosundan (veya MINITAB programından) p-değeri hesaplanmalı ve p-değeri kuralı uygulanmalıdır.

P-değeri kuralı (α anlamlılık düzeyinde):

Eğer p-değeri $< \alpha$; H_0 *reddedilir*

Eğer p-değeri $> \alpha$; H_0 *reddedilemez*

Alternatif olarak t-skoru kullanılarak kritik değer kuralı uygulanabilir (t kritik değerini $-t_{n-1;\alpha}$ bulmak için yine t dağılım tablosundan veya MINITAB programından faydalanabilirsiniz)

Kritik değer kuralı (α anlamlılık düzeyinde):

t- istatistik skoru $< -t_{n-1;\alpha}$; H_0 *reddedilir*

t- istatistik skoru $> -t_{n-1;\alpha}$; H_0 *reddedilemez*

Çift Taraflı Testler:

$$H_0 : \mu = \mu_0$$

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

Öncelikle t skoru hesaplanmalıdır

$$t = \frac{\bar{x} - \mu}{s_{\bar{x}}}$$

Burada $s_{\bar{x}} = \frac{s}{\sqrt{n}}$ değeri örneklem ortalamasının standart sapmasını temsil etmektedir. Şimdi ise bu t-skoru kullanılarak, t dağılım tablosundan (veya MINITAB programından) p-değeri hesaplanmalı ve p-değeri kuralı uygulanmalıdır.

P-değeri kuralı (α anlamlılık düzeyinde):

Eğer p-değeri $< \alpha$; H_0 *reddedilir*

Eğer p-değeri $> \alpha$; H_0 *reddedilemez*

Alternatif olarak t-skoru kullanılarak kritik değer kuralı uygulanabilir (t kritik değerini $\pm t_{n-1;\alpha/2}$ bulmak için yine t dağılım tablosundan veya MINITAB programından faydalanabilirsiniz)

Kritik değer kuralı (α anlamlılık düzeyinde):

t- istatistik skoru $> t_{n-1;\alpha/2}$ veya t- istatistik skoru $< -t_{n-1;\alpha/2}$ H_0 *reddedilir*

t- istatistik skoru $< t_{n-1;\alpha/2}$ veya t- istatistik skoru $> -t_{n-1;\alpha/2}$ H_0 *reddedilemez*

Tablo 8.4, α anlamlılık düzeyinde $t_{n-1;\alpha/2}$ ve $t_{n-1;\alpha}$ değerleri için en çok kullanılan kritik değerleri göstermektedir. Tabloda genel olarak kullanılan anlamlılık düzeyleri yüzde 1, 2.5, 5 ve 10 değerleri alınmıştır. $t_{n-1;\alpha/2}$ ve $t_{n-1;\alpha}$ değerleri aynı zamanda t dağılımı için kritik değerlerdir. Bu değerler (n-1) serbestlik derecesine göre değişmektedir.

Tablo 8.4 Standart t-Dağılımının Kritik Değerleri

Serbestlik derecesi (ν) = $n-1 = 6-1=5$

Anlamlılık Düzeyi	α	$t_{n-1;\alpha}$ (Kritik Değer)	$\alpha/2$	$t_{n-1;\alpha/2}$ (Kritik Değer)
10%	0.10	1.476	0.05	2.015
5%	0.05	2.015	0.025	2.571
2.5%	0.025	2.571	0.0125	-
1%	0.01	3.365	0.005	4.032

Serbestlik derecesi (ν) = $n-1 = 39-1 = 38$

Anlamlılık Düzeyi	α	$t_{n-1;\alpha}$ (Kritik Değer)	$\alpha/2$	$t_{n-1;\alpha/2}$ (Kritik Değer)
10%	0.10	1.304	0.05	1.686
5%	0.05	1.686	0.025	2.024
2.5%	0.025	2.024	0.0125	2.429
1%	0.01	2.429	0.005	2.712

Serbestlik derecesi (ν) = $n-1 = 121-1 = 120$

Anlamlılık Düzeyi	α	$t_{n-1;\alpha}$ (Kritik Değer)	$\alpha/2$	$t_{n-1;\alpha/2}$ (Kritik Değer)
10%	0.10	1.289	0.05	1.686
5%	0.05	1.658	0.025	1.98
2.5%	0.025	1.98	0.0125	2.358
1%	0.01	2.358	0.005	2.617

Tablo 8.4'ten de görüldüğü gibi serbestlik derecesi arttıkça n ile birlikte t dağılımı ile standart normal dağılım değerleri arasındaki fark azalmaktadır.

Örnek 8.12: Örnek 8.3'ü ele alarak t-testini uygulamaya çalışalım. Örneğimizde anakütle standart sapmasının bilindiğini varsaymıştık. Şimdi ise anakütle standart sapmasını bilmediğimizi varsayalım. Hipotez testimiz yine:

$$H_0 : \mu \geq 12$$

$$H_A : \mu < 12$$

Denklem (8.21)'deki t skorunu hesaplayabilmek için örneklem standart sapmasını s hesaplamamız gerekmektedir. $S_1 = \{7, 9, 9, 10, 12, 15\}$. Denklem (8.17)'den faydalanırsak:

$$s^2 = \frac{1}{6-1} \left[(7-10.33)^2 + (9-10.33)^2 + (9-10.33)^2 + (10-10.33)^2 + (10-10.33)^2 + (15-10.33)^2 \right] = 7.3307$$

$$\text{Dolayısıyla, } s = \sqrt{7.3307} = 2.707$$

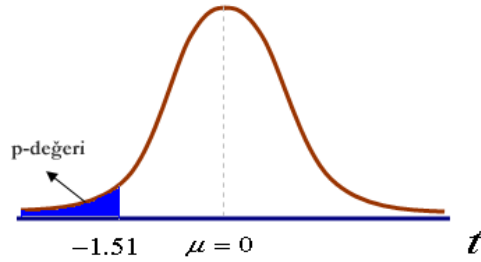
Tahmin edicisinin standart hatası da aşağıdaki gibi hesaplanabilir:

$$s_{\bar{x}} = \frac{2.707}{\sqrt{6}} = 1.105$$

Sonuç olarak t-skoru:

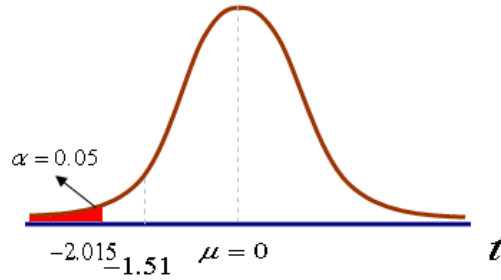
$$t = \frac{\bar{x} - \mu}{s_{\bar{x}}} = \frac{10.33 - 12}{1.105} = -1.51$$

Şimdi ise t-tablolarını (veya MINITAB'ı) kullanarak sol tarafta kalan p-değerini hesaplayabiliriz. T-tablosundan bu değeri 5 (=n-1) serbestlik derecesi ile $P(t < -1.51) = 0.097$ olarak bulabiliriz.



Şekil 8.32

p-değeri, standart sapmanın σ bilindiği duruma göre neredeyse iki kat bir değer almıştır. Dolayısıyla yüzde 1 ve 5 anlamlılık düzeylerinde sıfır hipotezi reddedilmemektedir. Kritik değer kuralı ile değerlendirmemizi yapacak olursak:



Şekil 8.33

Tablo 8.4'ten (Ek'te yer alan Tablo 3.2'den veya MINITAB'tan) $\alpha = 0.05$ anlamlılık düzeyinde t kritik değeri $t_{n-1;\alpha} = t_{5;0.05} = -2.015$ olarak bulunabilir.

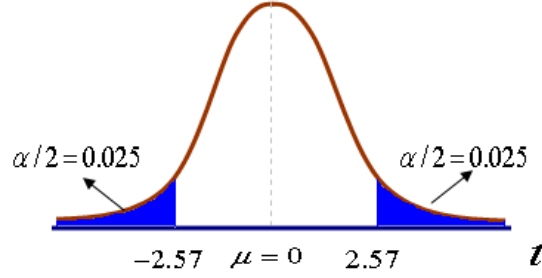
Örnek 8.13: Şimdi ise Örnek 8.12'yi kullanarak aşağıdaki testi gerçekleştirelim

$$H_0 : \mu = 12$$

$$H_A : \mu \neq 12$$

Testimizin p-değeri bir önceki örnekte yer alan p değerinin iki katıdır. Dolayısıyla, p-değeri = $2 \times 0.097 = 0.194$ ve bu p-değerinden hareketle herhangi bir uygun anlamlılık düzeyinde sıfır hipotezi reddedilmemektedir.

Kritik değer kuralına göre testimizi gerçekleştirsek



Şekil 8.34

Tablo 8.4'ten (Ek'te yeralan Tablo 3.2'den veya MINITAB'tan) $\alpha/2 = 0.025$ anlamlılık düzeyinde t kritik değeri $t_{n-1;\alpha/2} = t_{5;0.05} = \pm 2.571$ olarak bulunabilir.

Örnek 8.14: 1990 yılında Türkiye'deki hanehalkının ortalama yıllık benzin tüketimi \$600 idi. Üst gelir düzeyinde yer alan kişilerden rassal olarak 36 kişilik bir örneklem seçilmiştir. Bu grubun ortalama benzin tüketimi \$825 olarak, tahmin edilmiş standart sapması da \$150 olarak hesaplanmıştır. Acaba üst gelir düzeyinde yer alan hanehalklarının benzin tüketimi daha fazla mı olmaktadır? Testimiz açıkça görüldüğü gibi sağ taraflı bir testtir. Araştırmamızda tahmin edilmiş standart sapma yer aldığından, t istatistik değeri kullanmamız gerekmektedir. Bu durumdaki sıfır hipotezimiz:

$H_0: \mu \leq 600$ (ortalama harcama miktarı, genel harcama miktarından daha fazla değildir)

$H_0: \mu > 600$ (ortalama harcama miktarı, genel harcama miktarından daha fazladır)

Anlamlılık düzeyini de $\alpha = 0.05$ olarak seçtiğimiz takdirde, bu anlamlılık düzeyine denk gelen ve serbestlik derecesi 35 olan t-kritik değeri $t_{35;0.05} = 1.6896$ 'dır.

t-istatistik skorunu dönüşüm formülünü kullanarak aşağıdaki gibi hesaplayabiliriz:

$$t = \frac{825 - 600}{150 / \sqrt{36}} = 9$$

t istatistik değeri, t kritik değerinden daha büyük olduğundan sıfır hipotezini H_0 reddederiz. Dolayısıyla, yüksek gelir düzeyindeki ailelerin yıllık benzin tüketimi genel ortalamadan daha yüksektir sonucuna varabiliriz.

Anakütle ortalamasının Hipotez Testi: Küçük Örneklem Boyutu Durumunda

Merkezi Limit Teorem (MLT) tahmin aralığı yaparken çok önemli bir rol oynamıştı, ancak bu sadece $n \geq 30$ durumu için sözkonusu idi. Eğer $n < 30$ ise, $\bar{X}$ ların örneklem dağılımı

anakütlenin dağılımına bağlıdır. Daha önceki bilgilerimize dayanarak eğer x değerlerinin anakütle dağılımı normal ise $\bar{X}$ larında örneklem dağılımı da örneklem boyutundan bağımsız olarak normal dağılır. Bu yüzden örneklem boyutu küçük olduğu durumlarda $\bar{X}$ ların normal dağıldığına dair varsayım yapmamız gerekir. Bu varsayımın ötesinde analiz yine aynen devam ettirilir.

Eğer anakütle standart sapması (σ) biliniyorsa z dağılımını ve denklem (8.13) kullanırız. denklik kullanılır. Eğer anakütle sapması σ bilinmiyorsa, t dağılımı ve dolayısıyla denklem (8.21) kullanılmalıdır.

Küçük örneklem için en önemli varsayımımız anakütle dağılımının normal olduğu varsayımdır. Unutulmamalıdır ki örneklem boyutu arttığında, anakütle ortalamamız normal dağılmasa bile test için iyi sonuçlar elde edilebilir.

Anakütle Ortalaması μ için Hipotez Testi (Özet)

Bu bölümde, z ve t testlerinde detaylı olarak işlediğimiz test prosedürlerini özet olarak tekrar hatırlamaya çalışacağız. Heriki testte de özet olarak aşağıdaki adımları takip etmeniz gerekmektedir.

1. Öncelikle hipotezimizi, anakütle hakkındaki iddiamıza göre kurmamız gerekmektedir. Bu işlem için önümüzde üç olasılık vardır: anakütle ortalaması μ , belli bir μ_0 değerinden büyük olabilir, anakütle ortalaması μ , belli bir μ_0 değerinden küçük olabilir veya anakütle ortalaması μ , belli bir μ_0 değerinden farklı olabilir. Bu üç durumdan birine karar verdikten sonra, hipotezimizi sembolik olarak ifade ederiz. Eğer anakütle ortalamasının μ , belli bir μ_0 değerinden büyük olduğu iddiasını savunuyorsak, bu iddiamızı aşağıda da olduğu gibi alternatif hipotezimizde belirtebiliriz.

$$H_0 : \mu \leq \mu_0$$

$$H_A : \mu > \mu_0$$

Eğer anakütle ortalamasının μ , belli bir μ_0 değerinden küçük olduğu iddiasını savunuyorsak, bu iddiamızı yine alternatif hipotezimizde belirtebiliriz.

$$H_0 : \mu \geq \mu_0$$

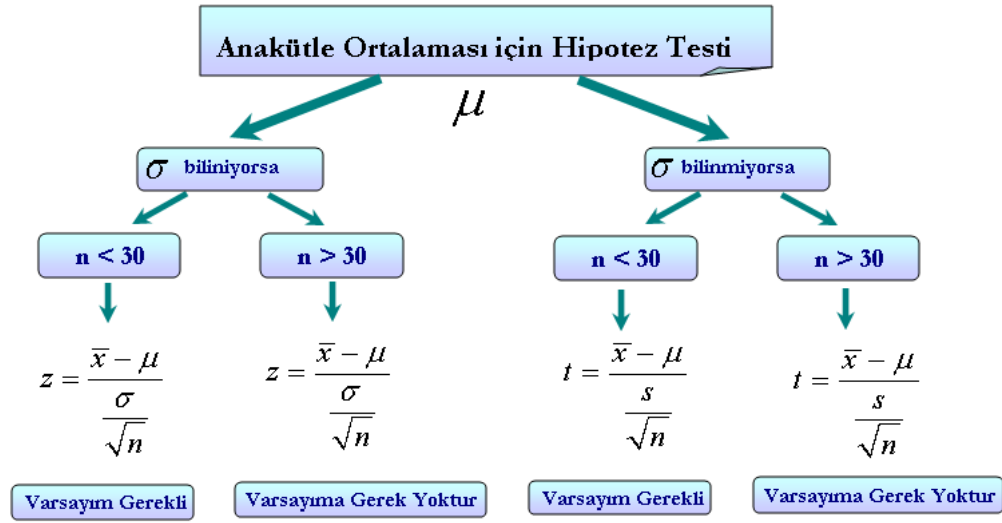
$$H_A : \mu < \mu_0$$

Eğer anakütle ortalamasının μ , belli bir μ_0 değerinden farklı olduğu iddiasını savunuyorsak, bu iddiamızı alternatif hipotezimizde belirttiğimizde:

$$H_0 : \mu = \mu_0$$

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

2. Testimiz için uygun olan test istatistiği (z -skoru veya t -skoru) ve dolayısıyla örneklem dağılımı belirlenmelidir. Şekil 8.27 hangi durumlarda hangi testin kullanılabileceği özetlenmektedir.



Şekil 8.35

3. Seçimimize göre, bahsettiğimiz gibi z veya t skorunu hesaplanmalıdır.
4. Anlamlılık düzeyi (α), dolayısıyla Tip I Hata olasılığı hesaplanmalıdır.
5. Seçilmiş olan test istatistiğine denk gelen kritik değerleri hesaplanmalıdır.
6. Test istatistik değerleri kritik değerlerinden daha büyük ise sıfır hipotezi H_0 reddedilmektedir. Aksi takdirde reddedilmemektedir.
7. İstenildiği takdirde p-değeri hesaplanabilir ve p-değeri, anlamlılık düzeyinden küçük olduğu takdirde sıfır hipotezi H_0 reddedilmektedir. Aksi takdirde reddedilmemektedir.

Anakütle Ortalaması için Hipotez Testi

Minitab Uygulamaları

Örnek 8.2'deki verileri ele alalım. Öncelikle örneklem değerlerini $S_1 = \{7, 9, 9, 10, 12, 15\}$ aşağıda gösterildiği gibi sütun "C1" e giriniz.

Worksheet 1 ***		
	C1	C2
↓		
1	7	
2	9	
3	9	
4	10	
5	12	
6	15	
7		

Şekil 8.36

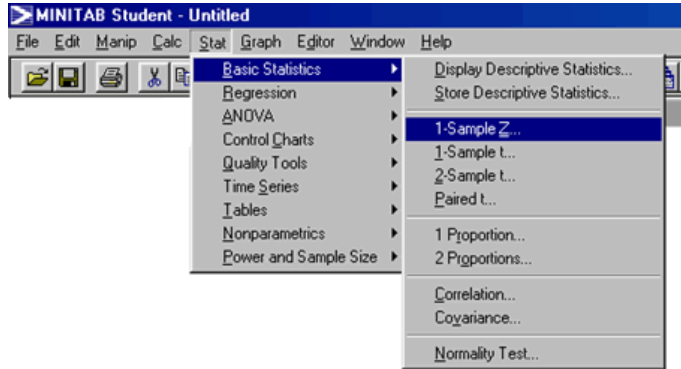
Sıfır hipotezimiz:

$$H_0 : \mu = 12$$

ve buna karşılık gelen alternatif hipotezimiz:

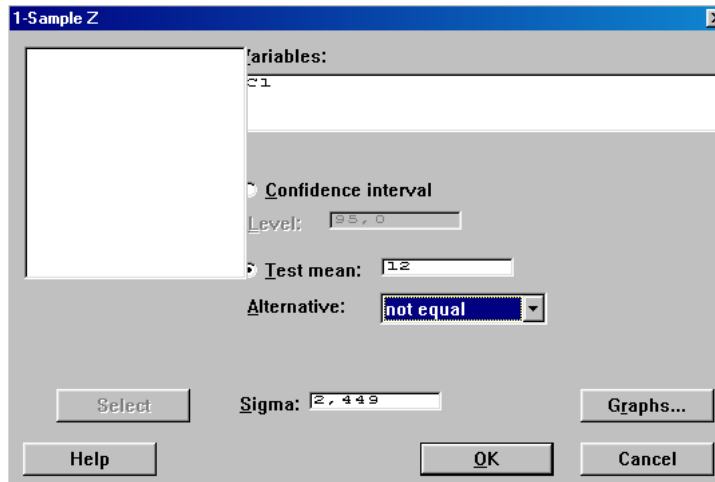
$$H_A : \mu \neq 12$$

Şekil 8.37'daki gibi “Stat” menüsünde “basic statistics” ten “1-SampleZ” ye tıklayınız.



Şekil 8.37

Ekrandaki menüde gözüken “C1” i seçiniz ve daha sonra “Test mean” tuş una basıp “12” yazınız. “Alternative” olduğu gibi bırakıp “sigma” hücresine “2.449” u yazın. Bu sigma bize bilinen anakütle standart sapmasını göstermektedir ve σ şeklinde ifade edilir:



Şekil 8.38

“OK” tuşuna bastığınızda aşağıdaki çıktı ekranı ile karşılaşmanız gerekmektedir.

Z-Test						
Test of mu = 12,00 vs mu not = 12,00						
The assumed sigma = 2,45						
Variable	N	Mean	StDev	SE Mean	Z	P
C1	6	10,33	2,80	1,00	-1,67	0,096

Şekil 8.39

Son iki satır yukarıda hesapladığımız z skorunu ve yukarıda hesaplanan p-değerini gösterir. Çift taraflı test yaptığımızdan dolayı yukarıdaki olasılık değerini 2 ile çarpalım ve değerimizdeki farklılıklar programın hataları yuvarlamasından kaynaklanmaktadır. Diyalog menüsündeki “Alternative” i değiştirerek MINITAB programında tek taraflı testte yapabilirsiniz. Bunu yaparsanız bulacağınız p-değerinin yaklaşık olarak yukarıda hesaplanan değerle aynı olduğunu göreceksiniz. Eğer p-değeri anlamlılık derecesinin (α) altında ise sıfır hipotezini reddedersiniz. Dolayısıyla eğer $\alpha = 0.05$ ise, sıfır hipotezini kabul edersiniz.

Örneğimizde anakütle standart sapmasını (σ) bilmediğimizi düşünelim. Böyle bir durumda t dağılımını kullanırsınız. Bu da MINITAB da “Stat” menüsünden “basic statistics” seçilerek oradan da “1-Sample t” opsiyonu seçilerek gerçekleştirilebilir. Bir önceki durumdan farklı olarak, standart sapma örneklemden hesaplandığından dolayı diyalog ekranında “sigma” standart sapmasının girilebileceği bir kutucuk bulunmamaktadır.

T-Test of the Mean						
Test of mu = 12,00 vs mu not = 12,00						
Variable	N	Mean	StDev	SE Mean	T	P
C1	6	10,33	2,80	1,15	-1,46	0,21

Şekil 8.40

Örneklem boyutunun küçük olmasından dolayı z ve t testlerinin p-değerleri arasında büyük farklılıklar görülebilir. Örneklem boyutu arttığında bu farklılıklarda azalmaktadır. Alıştırma olarak, Bölüm 7’deki MINITAB uygulamasının verilerinin bulunduğu “teabag.mtw” dosyasını ele alalım.

Aynı verileri kullanarak aşağıdaki çift taraflı testleri gerçekleştirelim. (hatırlayacağınız gibi anakütle standart sapmasını $\sigma = 0.25$ olarak varsaymışık) Bu hipotezi test etmek için:

$$H_0 : \mu \geq 12$$

alternatif olarak

$$H_A : \mu < 12$$

Bir önceki bölümde elde ettiğimiz p-değeri ile aynı değeri elde ettiğimiz aşağıdaki çıktıdan da gözükmektedir.

Z-Test						
Test of mu = 5,5000 vs mu < 5,5000						
The assumed sigma = 0,250						
Variable	N	Mean	StDev	SE Mean	Z	P
C1	40	5,4682	0,2856	0,0395	-0,80	0,21

Şekil 8.41

Aynı testi t dağılımını kullanarak gerçekleştirdiğimizde

T-Test of the Mean						
Test of mu = 5,5000 vs mu < 5,5000						
Variable	N	Mean	StDev	SE Mean	T	P
C1	40	5,4682	0,2856	0,0452	-0,70	0,24

Şekil 8.42

Açıkça görüldüğü gibi heriki testin de p-değerleri birbirlerine çok yakındır.

Anakütle Oranının Test Edilmesi : π

Anakütle oranı (π) hakkındaki test prosedürü daha incelediğimiz test prosedürüne çok benzemektedir.

$$H_0 : \pi = \pi_0$$

$$H_A : \pi \neq \pi_0$$

π_0 değeri, anakütle oranının π hipotetik oranıdır.

Bir önceki bölümde de belirttiğimiz gibi p'nin örneklem dağılımı binom ile dağılmaktadır.

Örnek 8.15: Yeni bir tedavi metodundan 20 hastanın 4 tanesi fayda görememiştir. 0.05 anlamlılık düzeyinde sıfır hipotezini $\pi = 0.5$, alternatif olarak $\pi \neq 0.5$ hipotezi ile birlikte test edelim.

$$H_0 : \pi = 0.5$$

$$H_1 : \pi \neq 0.5$$

Örneğimizden yola çıkarak, örneklem oranı $p = 4/20 = 0.2$ değerine eşittir ve binom ile dağılmaktadır. (x değişkeni $p = (1/n)x$ ifadesi ile binom dağılımı ile dağılmaktadır)

Binom dağılım formülü:

$$b(x; n, \theta) = \binom{n}{x} \pi^x (1 - \pi)^{n-x}, \text{ for } x = 0, 1, 2, \dots, n$$

olduğundan

$$P(X \leq 4) = \binom{20}{4} 0.5^4 (1-0.5)^{20-4} + \binom{20}{3} 0.5^3 (1-0.5)^{20-3} + \dots + \binom{20}{0} 0.5^0 (1-0.5)^{20-0}$$

$$P(X \leq 4) = 0.0059$$

Bu olasılık değeri 4'ten daha az hasta miktarının bu tedavi metodundan yararlanamama olasılığını temsil etmektedir. Birbaşka deyişle $\pi = 0.5$ iken örneklem ortalamasını $p = 0.2$ olarak elde etme olasılığımızı temsil etmektedir. Bu bağlamda p -değeri $= 2(0.0059) = 0.0118 < 0.05$ (test çift taraflı olduğundan hatırlayacağınız gibi p -değerini iki ile çarpmaktayız) olduğundan sıfır hipotezini H_0 reddederiz.

Büyük x değerleri söz konusu olduğunda binom formülünün kullanımı bilgisayar yardımı olmadan çok zordur. Dolayısıyla bu tarz durumlarda binom dağılımının normal dağılıma yakınsaması kullanılmaktadır. Bu yakınsama aşağıdaki teorem ile gerçekleştirilmektedir.

Teorem: Eğer X değişkeni, n ve π (başarı olasılığı) parametreleri ile rassal bir şekilde binom dağılımına sahip ise,

$$Z = \frac{X - n\pi}{\sqrt{n\pi(1-\pi)}}$$

$n \rightarrow \infty$ olduğunda teoremdaki z formülü de normal dağılıma yakınsamaktadır. (hatırlanacağı üzere binom dağılımının ortalaması ve varyansı $\mu = n\theta$ and $\sigma^2 = n\pi(1-\pi)$)

Dolayısıyla n büyük değerler aldığında binom dağılımının normal dağılıma yakınsama formülünü kullanabiliriz.

$$Z = \frac{x - n\pi}{\sqrt{n\pi(1-\pi)}}$$

veya bu ifadeyi n 'e bölerek aşağıdaki gibi kullanabiliriz.

$$Z = \frac{p - \pi}{\sqrt{\frac{\pi(1-\pi)}{n}}}$$

Örnek 8.17: Bir önceki örneğimizi binom dağılımının normal dağılıma yakınsama formülüne uygularsak:

$$Z = \frac{p - \pi}{\sqrt{\frac{\pi(1-\pi)}{n}}} = \frac{0.2 - 0.5}{\sqrt{\frac{0.5(1-0.5)}{20}}} = -2.683$$

Standart normal dağılım tablosunu kullanarak $P(Z < -2.683) = 0.0036$ değerini bulabiliriz. Bu bağlamda p -değeri $2 \times 0.0036 = 0.0072$ 'dir. p değeri, binom dağılımında hesapladığımızdan biraz daha küçük bir değere sahiptir. n 'in boyutu küçük olduğundan yakınsama iyi çalışmamaktadır. Ama tüm bunun yanında hala sıfır hipotezini reddetmekteyiz.

Oranların hipotez testi için test istatistiği olarak denklem (8.22)'yi kullanmaktayız:

$$Z = \frac{p - \pi}{\sigma_p} \quad (8.22)$$

Burada $\sigma_p = \sqrt{\frac{\pi(1-\pi)}{n}}$ ifadesinin anakütle oranının standart hatası olarak bir önceki bölümden hatırlayabilirsiniz.

Anakütle oranının π hipotez testi prosedürü

Sağ taraflı test:

$$H_0 : \pi \leq \pi_0$$

$$H_A : \pi > \pi_0$$

Öncelikle z skoru hesaplanmalıdır

$$Z = \frac{p - \pi}{\sigma_p}$$
$$\sigma_p = \sqrt{\frac{\pi(1-\pi)}{n}}$$

Şimdi ise bu z-skoru kullanılarak, z dağılım tablosundan (veya MINITAB programından) p-değeri hesaplanmalı ve p-değeri kuralı uygulanmalıdır.

P-değeri kuralı (α anlamlılık düzeyinde):

Eğer p-değeri $< \alpha$; H_0 **reddedilir**

Eğer p-değeri $> \alpha$; H_0 **reddedilemez**

Alternatif olarak z-skoru kullanılarak kritik değer kuralı uygulanabilir (z kritik değerini z_α bulmak için yine z dağılım tablosundan veya MINITAB programından faydalanabilirsiniz)

Kritik değer kuralı (α anlamlılık düzeyinde):

z- istatistik skoru $> z_\alpha$ H_0 **reddedilir**

z- istatistik skoru $< z_\alpha$ H_0 **reddedilemez**

Sol taraflı test:

$$H_0 : \pi \geq \pi_0$$

$$H_A : \pi < \pi_0$$

Öncelikle z skoru hesaplanmalıdır

$$Z = \frac{p - \pi}{\sigma_p}$$

$$\sigma_p = \sqrt{\frac{\pi(1-\pi)}{n}}$$

Şimdi ise bu z-skoru kullanılarak, z dağılım tablosundan (veya MINITAB programından) p-değeri hesaplanmalı ve p-değeri kuralı uygulanmalıdır.

P-değeri kuralı (α anlamlılık düzeyinde):

Eğer p-değeri $< \alpha$; H_0 *reddedilir*

Eğer p-değeri $> \alpha$; H_0 *reddedilemez*

Alternatif olarak z-skoru kullanılarak kritik değer kuralı uygulanabilir (z kritik değerini $-z_\alpha$ bulmak için yine z dağılım tablosundan veya MINITAB programından faydalanabilirsiniz)

Kritik değer kuralı (α anlamlılık düzeyinde):

z- istatistik skoru $< -z_\alpha$ H_0 *reddedilir*

z- istatistik skoru $> -z_\alpha$ H_0 *reddedilemez*

Çift Taraflı Testler:

$$H_0 : \pi = \pi_0$$

$$H_A : \pi \neq \pi_0$$

Öncelikle z skoru hesaplanmalıdır

$$Z = \frac{p - \pi}{\sigma_p}$$

$$\sigma_p = \sqrt{\frac{\pi(1-\pi)}{n}}$$

Şimdi ise bu z-skoru kullanılarak, z dağılım tablosundan (veya MINITAB programından) p-değeri hesaplanmalı ve p-değeri kuralı uygulanmalıdır.

P-değeri kuralı (α anlamlılık düzeyinde):Eğer p-değeri $< \alpha$; H_0 **reddedilir**Eğer p-değeri $> \alpha$; H_0 **reddedilemez**

Alternatif olarak z-skoru kullanılarak kritik değer kuralı uygulanabilir (z kritik değerini $\pm z_{\alpha/2}$ bulmak için yine z dağılım tablosundan veya MINITAB programından faydalanabilirsiniz)

Kritik değer kuralı (α anlamlılık düzeyinde):z- istatistik skoru $> z_{\alpha/2}$ veya z- istatistik skoru $> -z_{\alpha/2}$; H_0 **reddedilir**z- istatistik skoru $< z_{\alpha/2}$ veya z-istatistik skoru $< -z_{\alpha/2}$; H_0 **reddedilemez**

Burada dikkat edilmesi gereken husus örneklem oranının hesaplanması ile alakalıdır. Örneklem oranının standart sapması oranların birer fonksiyonudur.

$$\sigma_p = \sqrt{\frac{\pi(1-\pi)}{n}}$$

Ancak gerçek anakütle oranının bilinmediği durumlarda, z testi uygulanırken anakütle oranının yerine onun hipotetik değeri kullanılabilir.

Örnek 8.18: Bu yıl içerisinde İstatistik dersinden öğrencilerin başarısızlık oranı 10%'dur. Geçen sene İstatistik dersinden 75 öğrenciden 12'si başarısız olmuştur. Bu senenin daha önceki senelere oranla daha yüksek bir başarısızlık oranına sahip olup olmadığını test ediniz.

Bu amaç doğrultusunda sıfır ve alternatif hipotezlerimizi aşağıdaki gibi belirtebiliriz.

$$H_0 : \pi = 0.10$$

$$H_A : \pi > 0.10$$

Testimiz açıkça görüldüğü gibi sağ taraflı bir hipotezdir. Anlamlılık düzeyimizi $\alpha = 0.05$ olarak belirlersek, bu anlamlılık düzeyine tekabül eden geçerli kritik z değerimizi de 1.645 olarak buluruz. z-skoru, 1.645 değerinden büyük olduğu takdirde öğrencilerin başarısızlık oranının önceki yıllara göre, daha yüksek olduğu sonucuna varabiliriz.

75 öğrenciden 12 tanesinin başarısız olma oranı 16% veya $p = 0.16$ 'dır. $\pi = 0.10$ değerini bildiğimizden, $\pi = 7.5$ ve $n(1 - \pi) = 67.5$ olarak hesaplanabilir. Her iki durumda normal dağılım kullanımını sağladığından, bu değerleri denklem (8.22)'de yerine koyalım:

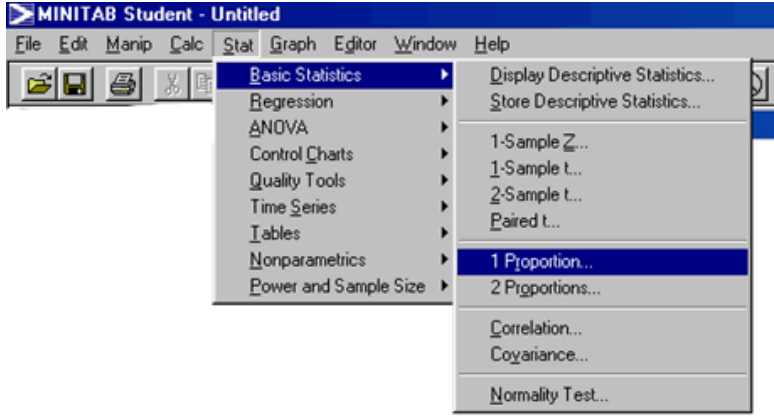
$$z = \frac{0.16 - 0.10}{\sqrt{\frac{0.10(1-0.10)}{75}}} = 1.73$$

z-skoru 1.73, 1.645 kritik değerinden büyük olduğundan, sıfır hipotezi reddedilmektedir. Dolayısıyla, İstatistik dersinin başarısızlık oranı daha önceki senelere oranla artmıştır sonucuna varabiliriz. Aynı sonuca p değerini 0.0418 hesaplayarak da varabiliriz

Anakütle Oranı π için Hipotez Testi

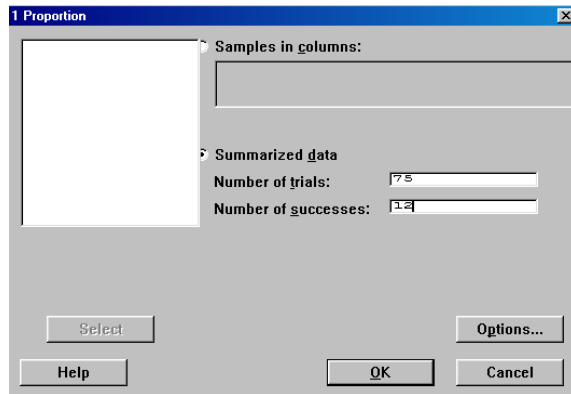
Minitab Uygulamaları

MINITAB ortamında Örnek 8.18'deki oran testini gerçekleştirebilmek için “Stat- basic statistics (temel istatistikler)” menüsünden “1-Proportion” seçeneğine tıklayınız.



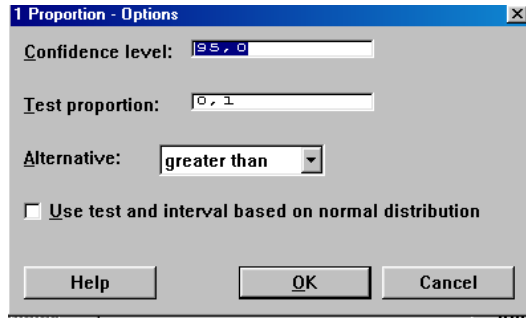
Şekil 8.43

Daha sonra “Summarized data” (özet veriye) geliniz ve “Number of trials” (deneme sayısı) bölümüne 75 yazınız ve “Number of successes” (başarı sayısı) bölümüne de 12 değerini örnek 18'deki verilere dayanarak giriniz.



Şekil 8.44

Daha sonra “Options” (Seçenekler) bölümünden hipotezi belirleyebiliriz. Şimdilik “Confidence Level” (Anlamlılık düzeyi) bölümünü boş bırakabilirsiniz, bu buraya Bölüm 9’da değineceğiz. “Test proportion” (oran testi) bölümüne ise 0.1 değerini sıfır hipotezi olarak girebiliriz. Alternatif hipotezi “greater than” (büyüktür), seçeneğinden faydalananarak oluşturabilir.



Şekil 8.45

Şimdilik şekil 8.37’deki gibi “Use test and interval based on normal distribution” (normal dağılımı baz alarak testi ve aralık tahminini gerçekleştiriniz) seçmemekteyiz. Çünkü şimdilik binom dağılımının normal dağılıma yakınsamasını kullanmamaktayız. “OK” tuşuna bastığımızda aşağıdaki çıktıyı elde edebiliriz.

Test and Confidence Interval for One Proportion						
Test of p = 0,1 vs p > 0,1						
Sample	X	N	Sample p	95,0 % CI	Exact P-Value	
1	12	75	0,160000	(0,085502; 0,262813)	0,069	

Şekil 8.46

Çıktıdan da görüldüğü gibi p-değeri, seçilen anlamlılık düzeyinden büyük olduğundan sıfır hipotezini reddedemeyiz. Aynı örneğin normal dağılım yakınsamasını elde etmek için MINITAB’taki bütün adımları aynen tekrar ediniz ve “Use test and interval based on normal distribution” kutucuğuna tıklayınız.

Test and Confidence Interval for One Proportion						
Test of p = 0,1 vs p > 0,1						
Sample	X	N	Sample p	95,0 % CI	Z-Value	P-Value
1	12	75	0,160000	(0,077031; 0,242969)	1,73	0,042

Şekil 8.47

$n\pi = 7.5$ ve $n(1 - \pi) = 67.5$ olduğundan, verileri Şekil 8.48'deki gibi giriniz.



	C1-T	C2	C3
↓			
1	F		
2	F		
3	F		
4	F		
5	F		
6	F		
7	F		
8	F		
9	F		
10	F		
11	F		
12	F		
13	S		
14	S		
15	S		
16	S		
17	S		
18	S		
19	S		
20	S		
21	S		
22	S		
23	S		
24	S		
25	S		

Şekil 8.48

Tek yapmamız gereken “*Summarized data*” yerine “*Samples in column*” kutucuğunu seçmek ve Şekil 8.44'teki gibi ve bir önceki adımları aynen tekrar etmektir.

Bölüm 9: Aralık Tahmini

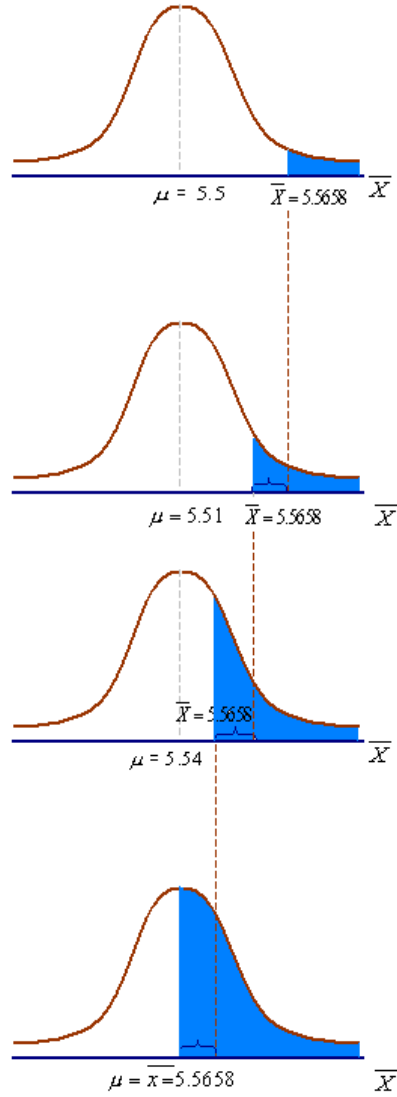
Giriş

7. bölümde nokta tahmini konusunu işlerken anakütle ortalaması (μ) gibi parametrelerin, nokta tahminlerinin nasıl yapılabileceğini incelemiştik. Yine 7.bölümde belirttiğimiz gibi, anakütle parametrelerinin nokta tahminlerinin ($\bar{x}$ gibi), tahmin ettikleri parametrelere (μ gibi) tam olarak eşit olmaları beklenilmemektedir.

Örnek 7.1'de de gördüğümüz gibi sadece üç elemanlı bir anakütleden alınan örneklemlerden hesaplanan örnek ortalamalarının değerleri 1 ve 3 arasında değişmekteydi. Bu örnekte, alınan örneklemeye bağlı olarak, örneklem ortalamasını 1, 1.5, 2, 2.5 veya 3 değerlerinden birisi olarak hesaplayabilir ve an kütle ortalamasını bunlardan biri olarak tahmin etmiş oluruz. Ancak bu tahminlerden sadece birisi hariç, hepsi anakütle ortalamasının (tahmin edilen parametre) değerinden farklıdır ($\mu = 2$). Bu son derece basit örnek dahi, göstermektedir ki nokta tahmin edicileri her zaman güvenilir bir tahmin oluşturmaktan uzaktır. Bazı örneklemelerde tahmin edilen parameter değerinden çok farklı olabilirler.

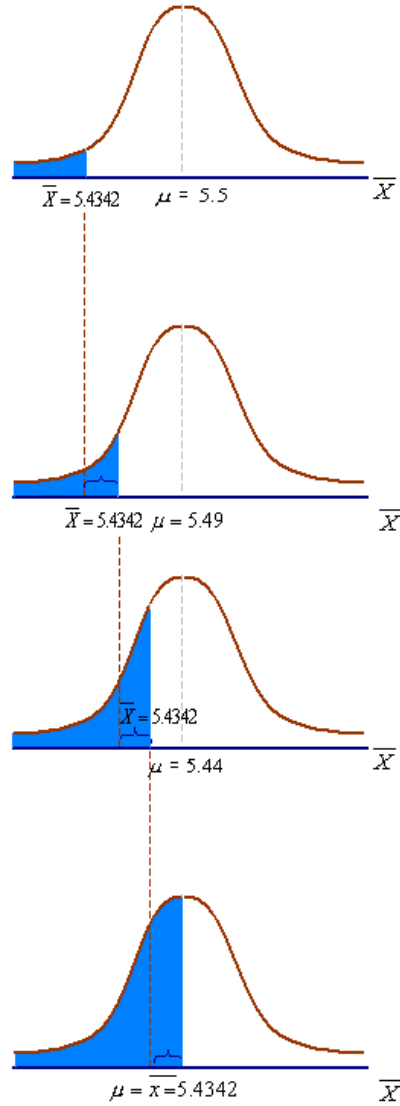
Aralık tahmini, nokta tahmin edicilerinin bu dezavantajını gidermeye çalışmaktadır. Aralık tahmin edicileri, adından da anlaşılabilirliği gibi, tek bir tahmin sunmak (örneğin, $\bar{x} = 1.5$) yerine, bir aralık içindeki sayılardan oluşan bir tahmin sunmaktadır (örneğin, μ , 1 ile 2.5 arasındadır gibi). Aralık tahmininin bir diğer özelliği de, anakütle parameteresinin içinde yer alacağı aralığı belirtirken, bu aralığa bağlı bir olasılık değeri de söyleyebilmemizdir. Örnek olarak, anakütle ortalamasının % 90 olasılıkla 1 ile 2.5 aralığında olduğunu söyleyebiliriz. Bu bölümde bu tür bir aralık tahmininin nasıl oluşturulabileceği üzerinde duracağız. Hemen belirtelim ki aralık tahmininin bir diğer ismi de güven aralığıdır. İleride bu iki terimi birbirlerini ikame eder biçimde kullanacağız.

Bir önceki bölümde incelediğimiz hipotez testi ve aralık tahmini arasında yakın bir ilişki bulunmaktadır. Örnek 8.1'i hatırlayacak olursanız, 5.5'ten büyük olarak öngördüğümüz bir μ değeri bize 0.05'ten büyük bir p değeri verecektir. Benzer bir şekilde μ değerini 5.5658 olarak öngörürsek, bu öngörümüz örneklem ortalamasının tamamıyla aynıdır (örneklem ortalaması ile öngördüğümüz anakütle ortalaması birbirine eşittir), dolayısıyla p değeri de 0.5'e eşit olur.



Şekil 9.1

Yine Bölüm 7’de yeralan Örnek 7.14’ü hatırlayacak olursanız, herhangi bir öngörölmüş μ değeri 5.5’ten küçük ise p değeri yine 0.05’ten büyük bir değeri olur. Eğer μ değerini 5.4342 olarak öngörürsek, ki bu değeri yine ortalamasının tamamiyle aynısıdır (örneklem ortalaması ile öngördüğümüz anakütle ortalaması birbirine eşittir), ve dolayısıyla p değeri aynı şekilde 0.5’e eşit olur.



Şekil 9.2

Bu yüzden eğer %10 olasılıkla Tip I hata yapmayı kabul edersek, öngördüğümüz anakütle ortalamasının 5.4342 ile 5.5658 arasında olduğunu söylersek, bu öngörümüzü reddedemeyiz. Bir başka ifade ile, %90 ihtimalle anakütle ortalamasının 5.4342 ile 5.5658 değerleri arasında olduğundan eminiz. Anakütlenin belli bir olasılık değeri ile verdiğimiz aralıkta olduğuna dair güvenimiz var. İşte bu nedenden dolayı bu aralıklara aynı zamanda güven aralığı da denilmektedir.

Anakütle Ortalamasının Aralık Tahmini: Örneklem Sayısının Yeterince Yüksek Olduğu Durum ($n \geq 30$)

Aralık tahmini, örneklem sayısının küçük ($n < 30$) veya büyük olduğu ($n \geq 30$) her iki durumda da yapılabilir. Ancak bu iki durumda uygulanacak formüller birbirlerinden, küçük de olsa, farklılıklar gösterecektir. Bunun dışında uygulayacağımız formül, yine anakütle standart sapmasının bilinip bilinmemesine (σ) göre farklılaşacaktır.

Büyük örneklemelerin en önemli karakteristiği, X değişkeninin anakütle dağılımı hakkında herhangi bir varsayım yapmamıza gerek olmamasıdır.

σ biliniyor ($n \geq 30$) ise

Yukarıda da belirtildiği gibi örneklem sayısı (n) 30'un üzerinde olduğu durumu, $n \geq 30$, büyük örneklem olarak kabul etmekteyiz ve bir önceki bölümde de incelediğimiz gibi, Merkezi Limit Teoremine (MLT) dayanarak $\bar{X}$ örneklem dağılımının normal dağılıma yakınsadığını kabul etmekteyiz. Bu bölümde ayrıca anakütle standart sapmasının (σ) bilindiği şeklinde bir diğer önemli varsayımda daha bulunuyoruz. Şimdi bu iki varsayım altında ($n \geq 30$ ve σ biliniyor), $\bar{X}$ 'nin güven aralığını oluşturmaya çalışalım. Bir örnek konuyu ele alalım.

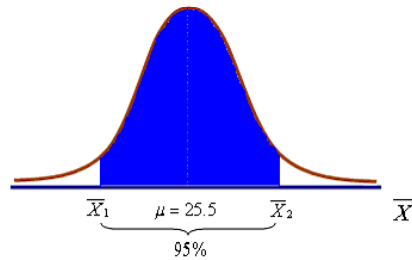
İstanbul Bilgi Üniversitesi'nin 36 öğrencisinden rassal bir şekilde oluşan bir örneklem seçelim. ($n = 36$). Amacımız Bilgi Üniversitesi öğrencilerinin yaş ortalaması için bir aralık tahmini gerçekleştirmek olsun. Örneklemden hesapladığımız ortalama yaş 25.5 yıldır ($\bar{X} = 25.5$). Bu örneklem ait veriler tablo 9.1'de özetlenmektedir

Tablo 9.1 İstanbul Bilgi Üniversitesinin 36 öğrencisinin yaşları:

18	25	19	18	33	45	20	22	19
23	24	26	28	29	40	19	31	32
26	25	24	27	28	29	22	21	26
23	25	20	21	19	33	35	21	11

Daha önceki çalışmalardan bilinmektedir ki öğrenci yaşlarının anakütle standart sapması 6'ya eşittir ($\sigma = 6$). Bir önceki bölümden de bilmekteyiz ki $n \geq 30$ olduğu durumlarda; $\bar{X}$ 'in örneklem dağılımı, μ ortalama, ve $\sigma_{\bar{X}} = \frac{\sigma}{\sqrt{n}}$ standart sapma (bu örnekte $\frac{6}{\sqrt{36}} = 1$) ile yaklaşık olarak normal dağılmaktadır.

$\bar{X}$ normal dağılımına göre, anakütle ortalamasını örnekten hesaplanan ortalamaya eşit varsayarsak ($\mu = 25.5$), $\bar{X}$ 'in dağılımını aşağıdaki şekildeki gibi gösterebiliriz.

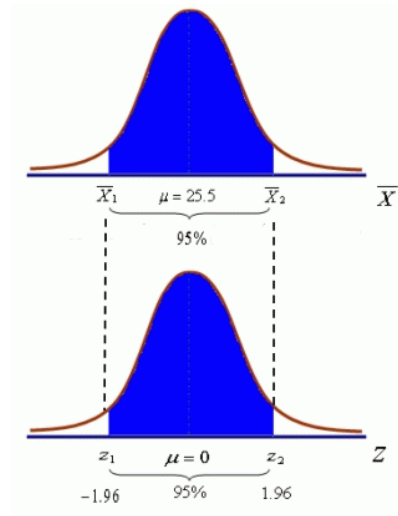


Şekil 9.3

Burada dikkat edilmesi gereken nokta: $\bar{X}$ 'in örneklem dağılımının ortalamasını (ki bildiğimiz gibi anakütle ortalamasına eşittir) 25.5 olarak kabul ederek, gerçek anakütle

ortalamasının 25.5'e, yani örneklem ortalamasına eşit olduğunu iddia etmiyoruz. Böyle bir iddia, bütün çabamızı anlamsız kılar, çünkü böyle bir durumda zaten bildiğimiz birşeyi tahmin etmeye çalışıyor olurduk. Yapmak istediğimiz, anakütle ortalaması 25.5 olsaydı, yani örneklem ortalamanıza (nokta tahminimize) bire bir eşit olsaydı, bu ortalamanın % 42.5 olasılığı sağında ve solunda (toplam % 95) bırakan değerler ne olurdu sorusuna cevap aramaktır. Bir başka deyişle, ortalaması 25.5 ve standart sapması 1 olan, bir normal dağılımda %95 olasılığa tekabül eden güven aralığı değerleri nedir sorusuna cevap arıyoruz. Şekil 9.3'ten de görülebileceği gibi eğer $\bar{X}_1$ ve $\bar{X}_2$ değerlerini bulabilirsek bu soruya cevap vermiş oluruz. Bu değerler bize aradığımız aralık tahminini (veya bir başka deyişle güven aralığını vereceklerdir), çünkü bu değerler bize eğer anakütle ortalaması 25.5 olsaydı, örneklem ortalamalarının % 95 olasılıkla yer alacağı aralığı vermektedir. Bu ifade zaman zaman şu şekilde de ifade edilmektedir: % 95 olasılıkla anakütle ortalaması $\bar{X}_1$ ve $\bar{X}_2$ aralığında bulunmaktadır. Her ne kadar bu ikinci ifade tam olarak doğru kabul edilemez ise de, sık sık kullanıldığı için birinci ifade ile aynı anlam da kullanıldığı düşünülebilir. Bu konu ileride daha fazla açıklığa kavuşturulacaktır.

Şimdi bu değerleri nasıl bulabileceğimizi inceleyelim. Bu noktada Standart Normal Dağılım bize yardımcı olacaktır. 6.bölümde Standart Normal Dağılım konusunu irdlemiştik. Standart Normal Dağılım, normal dağılımın sadece özel bir biçimidir. Standart Normal Dağılım'da, normal dağılımı tanımlayan iki parametre değerinden; ortalama sıfıra, standart sapma ise bire eşittir. Yani, Standart Normal Dağılım ortalamanın sıfıra, standart sapmanın ise bire eşit olduğu normal dağılımın özel bir biçimidir. Normal dağıldığını bildiğimiz herhangi bir X değişkeni kolaylıkla Standart Normal Dağılım'a dönüştürebilir. Bölüm 6'da da gösterildiği gibi % 95 olasılığa tekabül eden z değerleri 1.96 ve -1.96'dır. Şekil 9.3'ü şekil 6.24 ile birleştirirsek aşağıdaki şekli elde ederiz.



Şekil 9.4

Daha önce gördüğümüz z dönüştürme formülünü kullanırsak

$$z = \frac{\bar{x} - \mu}{\sigma_{\bar{x}}} = \frac{\bar{x} - \mu}{\sigma / \sqrt{n}} \quad (9.1)$$

Burada dikkat edilmesi gereken husus: daha önce incelediğimiz genel z dönüştürme formülünde normal dağılan değişken X olarak yazılmış idi ve ortalaması μ , standard sapması ise σ idi. Bu özel örnekte ise normal dağılan değişkenimiz $\bar{X}$ ’dır. Ortalaması μ , ve standart sapması ise σ ’dır. Dolayısıyla denklem (9.1)’i örneğimizde kullanacak olursak:

$$-1.96 = \frac{\bar{x}_1 - 25.5}{1}; \quad 1.96 = \frac{\bar{x}_2 - 25.5}{1}$$

$\bar{X}_1$ ve $\bar{X}_2$ değerlerini aşağıdaki gibi belirleyebiliriz

$$\bar{x}_1 = 25.5 - 1.96 \times 1; \quad \bar{x}_2 = 25.5 + 1.96 \times 1$$

veya

$$\bar{x}_{1,2} = 25.5 \pm 1.96 \times 1 = 25.5 \pm 1.96$$

Sonuç olarak, % 95 olasılık ile anakütle ortalaması [23.54; 27.46] aralığında olduğunu söyleyebiliriz. Ancak, yukarıda da görüldüğü gibi bu sonuç, anakütle ortalaması örneklem ortalamasına eşit olarak kabul edildiğinde bulundu. Bir başka örneklem alınsaydı, büyük bir olasılıkla 25.5’den farklı bir örneklem ortalaması bulunacağı için bu aralık tahmini [23.54; 27.46] de farklı bir değer alacaktı. Ancak, 25.5 toplanıp çıkarıldığı 1.96×1 sayısı bütün örneklem için sabit kalacaktır. Bu sayı bize bütün örneklem ortalamalarının aşağı ve yukarı (artı ve eksi olarak) % 95 ihtimalle sapacağı aralığı vermektedir. Sonuç olarak, % 95 olasılıkla örneklem ortalamaları 1.96×1 yıl fazla veya eksik olacaktır. Sadece % 5 olasılıkla örneklem hatası 1.96×1 yıldan daha fazla olabilir diyebiliriz. Bu sonuç bize elde ettiğimiz tahmin edicinin hassasiyeti üzerine bir bilgi verebilir. Eğer araştırmacı olarak bu aralığın çok geniş olduğunu ve bize bir bilgi vermekten uzak olduğunu düşünüyorsak, yapabileceğimiz bir tek şey vardır. O da, n sayısını arttırılarak örneklem ortalamalarının standart sapma değerinin düşürülmesidir. Örneğimizde standart sapma görüldüğü gibi 1’e eşittir. Eğer [23.54; 27.46] aralığının çok geniş olduğunu düşünüyorsak, n sayısını arttırarak standart sapma değerini 1’den aşağıya çekebilir ve aralığı küçültebiliriz.

İstatistikçiler α sembolünü örneklem hatası olasılığını kullanırlar. Dolayısıyla, $\alpha = 0.05$ ifadesi örneklem hatasının % 5’e eşit olduğunu göstermektedir. Normal dağılım simetrik olduğu için, normal dağılımın her iki kuyruğundaki alan da $\alpha / 2$ ’ye eşit olacaktır. Bu durumda $1 - \alpha$ de bizim güven aralığının olasılık değerini verecektir.

Daha önce de belirttiğimiz gibi, z harfini standart normal dağılıma sahip değişkeni temsil etmesi için kullanıyorsak, z harfine ilişilecek bir $\alpha / 2$ indisi ($z_{\alpha/2}$), standard normal dağılımın pozitif kuyruğunda $\alpha / 2$ olasılığına denk gelen z değerini göstermek amacıyla kullanabiliriz. Bu durumda eğer $(1 - \alpha) = 0.95$ olarak kabul edilirse, $z_{\alpha/2} = z_{0.025} = 1.96$ olacaktır. Dolayısıyla, $(1 - \alpha)$ olasılıkla örneklem ortalaması (artı, eksi) $z_{\alpha/2} \sigma_{\bar{X}}$ kadar bir örneklem hatası taşıyacaktır

Bu notasyonları kullanarak, aralık tahmini genel bir ifade ile aşağıdaki biçimde ifade edilebilir. $(1 - \alpha)$ olasılıkla gerçek ana kütle ortalaması μ , aşağıdaki aralıkta bulunacaktır.

$$\bar{x} \pm z_{\alpha/2} \sigma_{\bar{X}}$$

Bu aralık tahminini bir diğer biçimde ifade etmenin yolu ise, α olasılıkla aşağıdaki aralığın gerçek anakütle ortalamasını içermediğini söylemektir.

Bu ifade $\sigma_{\bar{X}}$ ’nin bir önceki bölümde öğrendiğimiz formülü kullanılarak, aşağıdaki gibi yazılabilir:

$$\bar{x} \pm z_{\alpha/2} \frac{\sigma}{\sqrt{n}} \quad (9.2)$$

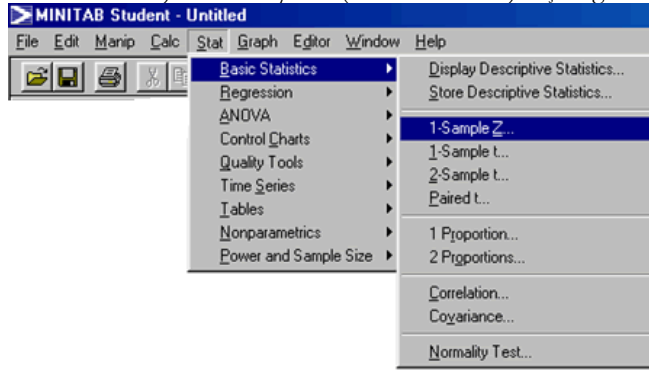
Aşağıdaki tablo, $(1-\alpha)$ 'nın genellikle kullanılan 0.90, 0.95 ve 0.99 değerleri için $z_{\alpha/2}$ değerlerini göstermektedir.

Tablo 9.2 Standart Normal Dağılımı için Güven Aralığı Değerleri

Güven Aralığı	α	$\alpha / 2$	$z_{\alpha/2}$ (Kritik Değer)
90%	0.10	0.05	1.645
95%	0.05	0.025	1.96
99%	0.01	0.005	2.575

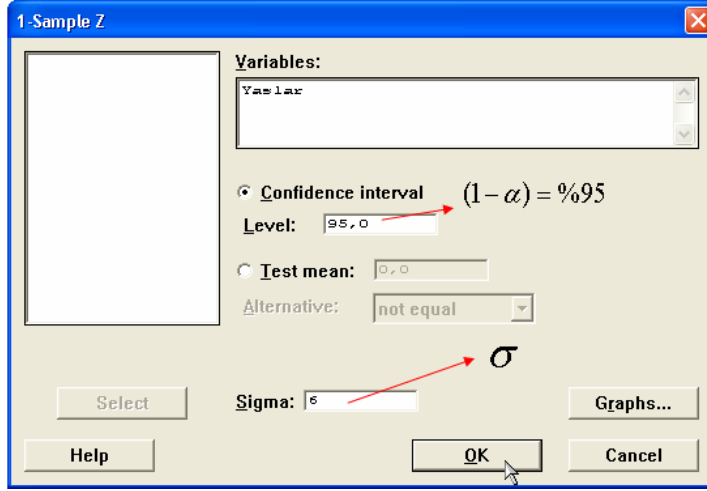
Minitab Uygulamaları

MINITAB ortamında güven aralıklarını kolaylıkla hesaplayabilirsiniz. Bu işlemin nasıl yapıldığını incelemiden önce, lütfen Tablo 9.1'de yer alan verilerin bulunduğu "*yaslar.mtn*" adlı MINITAB çalışma sayfasını açınız. Güven aralığını oluşturmak için "*Stat*" (istatistik) menüsünden "*basic statistics*" (temel istatistikler) ve "*1-Sample Z*" (1 Örneklemli Z) seçeneğine tıklayınız.



Şekil 9.5

Karşınıza çıkacak olan diyalog ekranını şekil 9.6'daki gibi doldurunuz



Şekil 9.6

“Yaslar” öğrenci yaşlarının gösterildiği veri sütununu temsil etmektedir. “Confidence Level” (Güven Aralığı) kutusunu tıklamanız gerekmekte ve sizin için uygun olan olasılık düzeyini yazmanız beklenmektedir. Biz burada % 95 olasılığını uygun bulduk, çünkü örneğimizde güven aralığı % 95 olasılık değeri için bulunmuş idi. Sigma kutusuna 6 yazıldı. Bu kutu bilindiği varsayılan anakütle standard sapmasının değeri için ayrılmıştır. Örneğimizde öğrenci yaşlarının standard sapmasının 6’ya eşit olduğunu varsaymıştık.

“OK” tuşuna tıkladıktan sonra aşağıdaki güven aralığını elde etmeniz gerekmektedir.

Z Güven Aralıkları					
Varsayılan sigma(anakütle standart sapması) = 6,00					
Variable	N	Mean	StDev	SE Mean	95,0 % CI
Yaslar	36	25,50	6,24	1,00	(23,54; 27,46)
	(örneklem boyutu)	(örneklem ortalaması)	(örneklem standart sapması)		

$$\frac{\sigma}{\sqrt{n}} = \frac{6}{\sqrt{36}} = 1$$

Şekil 9.7

σ bilinmiyor ($n \geq 30$) ise

Bir önceki alt başlıkta, anakütle varyansının (σ^2) bilindiği varsayımı altında çalışmıştık. Ancak, daha gerçekçi bir yaklaşımla, σ^2 ’nin bilinmediği durum üzerinde odaklanmamız gerekmektedir. Eğer σ^2 bilinmiyor ise, tek çare, bir önceki bölümde de yaptığımız gibi anakütle varyansının bir tahminin örnekleminden hesaplanmasıdır. σ^2 tahminin formülü, hatırlayacağınız üzere, aşağıdaki gibi idi:

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

Bu durumda, $\bar{x}$ 'ın standart sapmasını aşağıdaki formülle bulabiliriz.

$$s_{\bar{x}} = \frac{s}{\sqrt{n}}$$

Tekrar örneğimize geri dönecek olursak, anakütle varyans tahmini

$$s^2 = \frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n-1} = 38.93 \quad \text{olarak, } \bar{x} \text{ 'ın standart sapması ise}$$

$$s_{\bar{x}} = \frac{s}{\sqrt{n}} = \frac{6.24}{6} = 1.04 \quad \text{olarak hesaplanacaktır.}$$

Şimdi $s/\sqrt{n}$ ifadesini, (9.1) nolu denklemin paydasında $\sigma/\sqrt{n}$ - yerine- kullanarak bütün aşamaların tekrar edilmesi gerektiğini düşünebilirsiniz. Ancak burada bir farklılık meydana gelecektir. $s/\sqrt{n}$ ifadesini (9.1) nolu denklemin paydasında kullandığımızda, yani $\bar{x}$ gerçek standard sapması ($\sigma/\sqrt{n}$) yerine onun tahminini kullandığımızda ($\frac{s}{\sqrt{n}}$), maalesef artık standard normal dağılımı kullanmamız mümkün olmamaktadır. Neden ? Bunu açıklamak için (9.1) nolu denklemin sağ tarafını buraya tekrar yazalım

$$t_{n-1} = \frac{\bar{x} - \mu}{s/\sqrt{n}}$$

$t_{\alpha/2}$ kritik değeri ile birlikte, t dağılımı kullanarak güven aralığı formülü denklem (9.3)'teki gibidir:

$$\bar{x} \pm t_{n-1;\alpha/2} \frac{s}{\sqrt{n}} \quad (9.3)$$

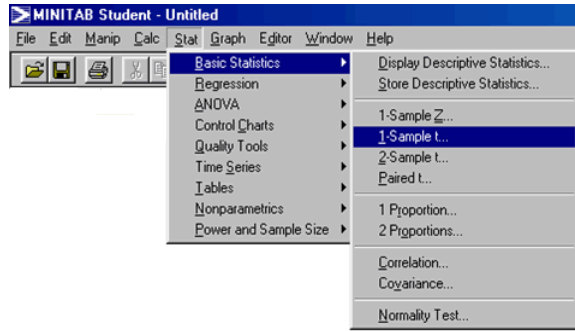
(9.3)'nolu denklemi uygulayarak % 95 olasılığa tekabül eden güven aralığını bulabiliriz. Ancak bu değeri bulabilmemiz için öncelikle, nasıl denklem (9.1)'de $z_{\alpha/2}$ değerini bilmemiz gerekemekteyse burada da $t_{n-1;\alpha/2}$ değerini bilmemiz gerekmektedir. En sık kullanılan $z_{\alpha/2}$ Tablo 9.2'de gösterilmiş idi. Ancak $t_{n-1;\alpha/2}$ değerleri daha öncede belirttiğimiz gibi serbestlik derecesine bağlı olarak değişmektedir (serbestlik derecesi de veri sayısına bağlı olarak değişmektedir). Oysa $z_{\alpha/2}$ değerleri, Tablo 9.2'den açıkça görülebileceği gibi, sadece seçilen olasılık değerine (α) bağlı olmaktadır. Dolayısıyla burada t değerleri için benzer bir tablo size göstermiyoruz (Kitabın arkasında yer alan ekte bu tabloyu bulabilirsiniz). Ancak yukarıdaki işlemi bitirmemiz için gerekli $t_{n-1;\alpha/2}$ değerini vermekle yetineceğiz. Bu değer: $t_{36-1;0.05/2} \approx \pm 2.030$. Dolayısıyla güven aralığını aşağıdaki gibi hesaplayabiliriz.

$$x_{1,2} = 25.5 \pm 2.030 \times 1.04 = 25.5 \pm 2.11$$

Anakütle ortalaması 95% ihtimalle [23.39; 27.61] aralığında yer almaktadır.

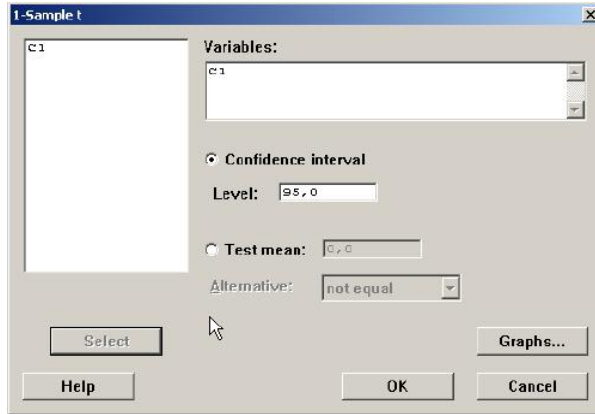
Minitab Uygulamaları

Standart sapma değerinin bilinmediği durumuna ait güven aralığını MINITAB içerisinde elde etmek için bütün aşamaları tekrar etmememiz gerekmektedir. Güven aralığını oluşturmak için “Stat” menüsünden “basic statistics” ve “1-Sample t” seçeneğine tıklayınız.



Şekil 9.8

Daha sonra karşılaştığınız ekranı Şekil 9.9’da gösterildiği gibi doldurunuz.



Şekil 9.9

Burada dikkat edileceği gibi anakütle standard sapması bilinmediği için böyle bir kutucuk da bir önceki durumda olduğu gibi aktif olmamıştır. “OK” tuşuna bastığınızda aradığımız güven aralığını elde edebilirsiniz.

T Guven Araliklari					
Variable	N	Mean	StDev	SE Mean	95,0 % CI
Yaslar	36	25,50	6,24	1,04	(23,39; 27,61)
	(örneklem boyutu)	(örneklem ortalamasi)	(örneklem standart sapmasi)		

$$\frac{s}{\sqrt{n}} = \frac{6,24}{\sqrt{36}} = 1,24$$

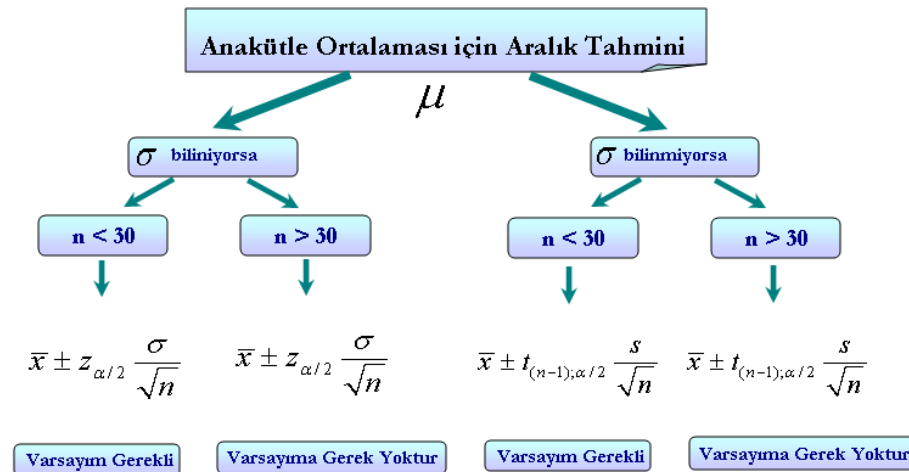
Şekil 9.10

Aralık Tahmini Küçük Örneklem Durumu ($n < 30$)

MLT güven aralığının oluşturulması tartışmalarında son derece önemli bir rol oynamaktadır. Ancak, sadece $n \geq 30$ durumunu ele aldık. Bu durumda bilindiği gibi anakütlede X 'lerin dağılımı hakkında bir varsayım yapmamıza gerek kalmadan $\bar{X}$ ' in normal dağıldığını (σ biliniyorsa) veya t dağılımı ile dağıldığını (σ bilinmiyorsa) varsayabiliyorduk. Ancak n sayısı 30 dan küçük olursa bu durumda anakütlede X 'lerin normal dağıldığını varsaymak zorundayız ki $\bar{X}$ 'in normal dağıldığını (σ biliniyorsa) veya t dağılımı ile dağıldığını (σ bilinmiyorsa) varsayabilelim.

Aralık Tahminleri için Uygun Dağılımın Seçimi

Aralık tahmini için uygun dağılımı seçerken, hipotez testinde kullandığımız prosedürün aynısını kullanmaktayız. Aşağıdaki tablo, aralık tahmini için uygun dağılımın seçiminde bize yardımcı olmaktadır. $n < 30$ durumunda anakütle ortalamasının normal olarak dağıldığını kapalı olarak varsaymaktayız.



Şekil 9.11

Anakütle Oranı için Aralık Tahmini

Şuana kadar, anakütle ortalaması için aralık tahminlerini heriki dağılımı da kullanarak da inceledik. Aralık tahminlerinin genel yapısını formülüze edecek olursak:

$$\left(\begin{array}{c} \text{örneklem} \\ \text{istatistiği} \end{array} \right) \pm \left(\begin{array}{c} \text{Uygun olan dağılımın} \\ \text{kritik değeri} \end{array} \right) \times \left(\begin{array}{c} \text{Örneklem parametresinin} \\ \text{standart sapması} \end{array} \right)$$

Anakütle oranının π tahmin edicisi, binom dağılımı ile dağılmaktadır (bir önceki bölümü inceleyebilirsiniz) ve standart sapması $\sigma_p = \sqrt{\frac{\pi(1-\pi)}{n}}$ formülü ile hesaplanmaktadır. Bir önceki bölümde de tartıştığımız gibi binom dağılımının yerine z dağılımı kullanılmaktadır. Dolayısıyla anakütle oranı için güven aralığı formülünü gösterecek olursak:

$$p \pm z_{\alpha/2} \sqrt{\frac{\pi(1-\pi)}{n}}$$

π değerini bilmediğimizden dolayı, bu parametrenin yerine tahmin edicisi olan p kullanılmalıdır. Dolayısıyla formülümüz aşağıdaki gibi değişmektedir.

$$p \pm z_{\alpha/2} \sqrt{\frac{p(1-p)}{n}}$$

Örnek 9.1: 8. bölümdeki örnek 8.18'i tekrar ele alalım. İstatistik sınavındaki başarısızlık oranı %10 idi. Geçen sene ise 75 öğrenciden 12'si başarısız olmuştur. Anakütle oranı için %95'lik bir güven aralığı oluşturalım.

75 öğrenciden 12'si başarısız ise başarısızlık oranı %16 veya $p = 0.16$ 'dır. Bu doğrultuda güven aralığını aşağıdaki gibi elde edebiliriz.

$$0.16 \pm 1.96 \sqrt{\frac{0.16(1-0.16)}{75}} = [0.077; 0.242]$$

8. bölümdeki MINITAB uygulamasını, güven aralığı hesaplanmasına kolaylıkla adapte edebilirsiniz.

Bölüm 10:

İki Anakütle Parametresi için Hipotez Testleri

Giriş

9. bölümde güven aralığı kavramına giriş yapmış ve aralık tahminin anakütle parametrelerini tahmin etmek için nasıl yapılacağını göstermiştik. Bu metod çıkarımsal istatistiğin önemli bir özelliğidir. 8.bölümde ise anakütle parametrelerinin tahmini için çıkarımsal istatistiğin diğer bir özelliği olan hipotez metodunu kullandık, fakat burada kullandığımız hipotez testi tek bir anakütle parametresi hakkında çıkarım yapmaya yönelikti. Bu konuyu irdelerken zımni olarak sadece bir anakütle olduğunu varsaymıştık. Gerçekte iki veri setinin karşılaştırılması çok daha fazla kullanılan ve karşılaşılan bir durumdur. Örneğin, sınavda kalan öğrencilerin farklı cinsiyetlere göre oranları, özel sektörde çalışanların aldığı ücretlerle kamu sektöründe çalışanların aldıkları ücretlerin karşılaştırılması gibi. Bu tip durumlarda tek anakütle parametresi için kullandığımız hipotez testi yerine çift anakütle parametresi için hazırlanmış hipotez testini kullanmalıyız.

İki örneklem dayalı hipotez testi metodlarını göz önünde bulundururken, örneklemelerin birbirinden bağımsız olup olmadıklarını bilmemiz önemlidir. Eğer bir anakütleden seçilen örneklem diğer anakütle ile alakalı değilse, iki örneklem bağımsızdır. Eğer bu örneklem diğer örneklem ile alakalı ise bu örneklem birbiriyle bağımlıdır denilir. Bu tip örneklemelere, ikili veya eşlenik örneklem denilir. Hipotez testinin mantığı aslında seçilen örneklem bağımlı ya da bağımsız olup olmadığına bağlıdır.

İki Anakütle Parametresi için Çıkarımlar: Varyans biliniyorsa (Bağımsız örneklem)

Büyük Örneklerde

İki anakütle ortalaması ile alakalı hipotez testlerinde ve güven aralıkları yapılandırılırken, genellikle iki anakütle ortalaması arasında istatistiksel olarak anlamlı derecede bir fark olup olmadığı ile ilgileniriz. Eğer iki tane bağımsız (bir ve iki diye tanımlanacak) ve her biri büyük örneklem grubuna girebilecek büyüklükte ($n_1 \geq 30$ ve $n_2 \geq 30$ gibi) örneklemim varsa, uygulanacak olan hipotez testi daha önceki derste öğrendiğimiz tarza benzer bir şekilde uygulanacaktır.

Genel yapısı aşağıdaki gibidir:

$$\frac{(\text{örneklem istatistiği}) - (\text{iddia edilen anakütle parametresi})}{\text{örneklem parametresinin standart sapması}}$$

Daha öncede olduğu gibi istatistiksel testi uygularken öncelik sıfır ve alternatif hipotezlerin formülse bir şekilde ifade edilmesinden yanadır. Örneğin, eğer iki ayrı anakütleden alınmış iki ayrı örneklemi (1 ve 2) kullanarak anakütle ortalamalarının farklı olup olmadığını test etmek istersek sıfır hipotezimizi ve alternatif hipotezimizi aşağıdaki gibi kurarız:

$$H_0 : \mu_1 - \mu_2 = \mu_0$$

$$H_A : \mu_1 - \mu_2 \neq \mu_0$$

alternatif olarak ise :

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

Sıfır hipotezini ve alternatif hipotezi daha farklı bir şekilde ifade edebiliriz:

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

$$H_A : \mu_1 - \mu_2 \neq 0$$

Veya alternatif olarak anakütle ortalamaları arasındaki farkın belli bir değere eşit olup olmadığını test etmek istersek:

$$H_0 : \mu_1 - \mu_2 = \mu_0$$

$$H_A : \mu_1 - \mu_2 \neq \mu_0$$

Bunun akabinde bu testte kullanılacak test istatistiği $(\bar{X}_1 - \bar{X}_2)$ şeklinde yazılır. Burada da yine Merkezi Limit Teoremini kullanmak gerekmektedir (MLT). Eğer n yeterince büyükse, örneklem ortalamaları ($\bar{X}$ lar) normal dağılıma eğilimi gösterirler. Aynı teoremi kullanarak örneklem ortalamaları arasındaki farkı ifade edebiliriz, $(\bar{X}_1 - \bar{X}_2)$ (burada $\bar{X}_1$, birinci örneklemin ortalaması ve $\bar{X}_2$ de ikinci örneklemin ortalamasıdır), ve bu fark $(\mu_1 - \mu_2)$ da anakütle ortalamalarıyla birlikte

$$\sigma_{\bar{X}_1}^2 = \frac{\sigma_1^2}{n_1}, \quad \sigma_{\bar{X}_2}^2 = \frac{\sigma_2^2}{n_2} \quad (\sigma_1^2 \text{ ve } \sigma_2^2 \text{ birinci ve ikinci anakütlenin varyanslarını temsil etmektedir})$$

varyanslarıyla normal olarak dağılmaktadır. Anakütlenin dağılımından bağımsız olarak, eğer $n_1, n_2 > 30$ ise $(\bar{X}_1 - \bar{X}_2)$ 'in örneklem dağılımını bulmak için aşağıdaki teoremden faydalanmaktayız.

Teorem: Eğer $X_1, X_2, \dots, X_n$ değişkenleri birbirinden bağımsız ve normal olarak dağılıyorsa, $X_i \rightarrow (\mu_i, \sigma^2)$ $X_i \rightarrow (\mu_i, \sigma^2)$, k_i sıfırdan farklı sabit bir sayı olmak üzere $X = k_1 X_1 + k_2 X_2 + \dots + k_n X_n$ ifadesi de $E(X) = \mu = k_1 \mu_1 + k_2 \mu_2 + \dots + k_n \mu_n$ ortalaması ve $\sigma^2 = k_1^2 \sigma_1^2 + k_2^2 \sigma_2^2 + \dots + k_n^2 \sigma_n^2$ varyansı ile $Z \rightarrow (\mu, \sigma^2)$ normal olarak dağılmaktadır.

Teoreminizi, iki örneklem farkının dağılımı durumuna uyguladığımız takdirde, $(\bar{X}_1 - \bar{X}_2)$ değişkeninin

$$E(\bar{X}_1 - \bar{X}_2) = \mu_1 - \mu_2$$

ortalaması

$$\text{var}(\bar{X}_1 - \bar{X}_2) = \sigma_{\bar{X}_1 - \bar{X}_2}^2 = \sigma_{\bar{X}_1}^2 + \sigma_{\bar{X}_2}^2 = \frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}$$

ve varyansı ile normal olarak dağıldığı açıkça gözükmektedir. (örneklem ortalamalarının farkı teoremimizin $k_1 = 1, k_2 = -1$ ve $X_1 = \bar{X}_1, X_2 = \bar{X}_2$ olduğu özel bir durumdur.) Dolayısıyla

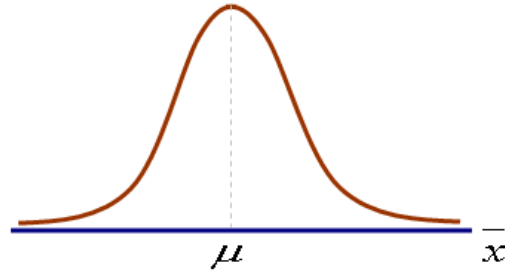
$$(\bar{X}_1 - \bar{X}_2) \rightarrow N\left(\mu_1 - \mu_2, \frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}\right)$$

Hipotez testinin geriye kalan prosedürü daha önce işlediğimiz test prosedürlerine çok benzemektedir. $\bar{X}$ ve $(\bar{X}_1 - \bar{X}_2)$ 'in örneklem dağılımlarının karşılaştırılmasında fayda bulunmaktadır.

Tek anakütleli durumda:

$$\bar{X} \rightarrow N\left(\mu, \frac{\sigma^2}{n}\right)$$

ifadesi, $\bar{X}$ değişkeninin μ ortalaması ve $\frac{\sigma^2}{n}$ varyansı ile normal olarak dağıldığı anlamına gelmekteydi. Grafiksel olarak gösterecek olursak:



Şekil 10.1

Bu durum için z skorumuz ise

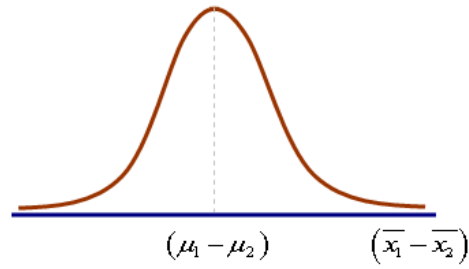
$$z = \frac{\bar{x} - \mu}{\sqrt{\sigma^2/n}} \quad (10.1)$$

İki anakütlenin bulunduğu durumlarsa ise:

$$(\bar{X}_1 - \bar{X}_2) \rightarrow N\left(\mu_1 - \mu_2, \frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}\right)$$

ifadesi $(\bar{X}_1 - \bar{X}_2)$ değişkeninin, $\mu_1 - \mu_2$ ortalaması ve $\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}$ varyansı ile normal olarak dağıldığını göstermektedir.

Grafiksel olarak gösterecek olursak:



Şekil 10.2

Dolayısıyla denklem (10.1)'deki genel yapıdan hareketle, iki anakütleli durumda z skorunu denklem (10.2)'de ifade edildiği gibi kullanabiliriz.

$$z = \frac{(\bar{x}_1 - \bar{x}_2) - (\mu_1 - \mu_2)}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}} \quad (10.2)$$

Testi gerçekleştirmek için sıfır hipotezindeki anakütle parametrelerinin (μ lerin) değerlerini formüle eklemeliyiz. Sıfır hipotezinde, elimizde $\mu_1 - \mu_2 = 0$ olduğundan dolayı bu denkliği (10.1) de yerine koyarız ve istenen z-skorunu elde etmiş oluruz. Formülümüzde aşağıdaki gibi olur:

$$z = \frac{(\bar{x}_1 - \bar{x}_2)}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}} \quad (10.3)$$

Bir sonraki adımda, hipotez testinin doğruluğunu araştırmak için z-tablolarından gerekli kritik z değerlerinin elde edilmesidir. Daha önceki bilgilerimize dayanarak, öncelikle anlamlılık derecesi olan α ' yı belirleriz. Çift taraflı bir test uygulayacağımızdan testimizi aşağıdaki gibi formülüne edebiliriz.

$$H_0 : \mu_1 - \mu_2 = \mu_0$$

$$H_A : \mu_1 - \mu_2 \neq \mu_0$$

Denklem (10.2)'de hesapladığımız z skorunu, z dağılım tablosundan (veya MINITAB'tan) faydalanarak, p değerini hesaplamak için kullanabiliriz.

Hipotez testinde p-değeri yaklaşımı :

Eğer p değeri $< \alpha$ ise H_0 hipotezi reddedilir
Eğer p değeri $> \alpha$ ise H_0 hipotezi reddedilemez.

Alternatif olarak z-skoru kullanılarak kritik değer kuralı uygulanabilir (z kritik değerlerini $\pm z_{\alpha/2}$ bulmak için yine z dağılım tablosundan veya MINITAB programından faydalanabilirsiniz)

Hipotez Testinde Kritik Değer Yaklaşımı:

Eğer z-istatistik skoru $> z_{\alpha/2}$ veya z-istatistik skoru $> -z_{\alpha/2}$; H_0 **reddedilir**
Eğer z-istatistik skoru $< z_{\alpha/2}$ veya z-istatistik skoru $< -z_{\alpha/2}$; H_0 **reddedilemez**

Ortalama farkları için Güven Aralıkları

9.bölümden de hatırlayacağınız gibi, büyük örneklem boyutlarında ve anakütle standart sapmasının bilindiği durumlarda , anakütle ortalamasını μ için güven aralığını aşağıdaki gibi formülize etmiştik.

$$\bar{x} \pm z_{\alpha/2} \frac{\sigma}{\sqrt{n}}$$

Güven aralıklarının genel yapısını yazacak olursak:

$$\left(\begin{array}{c} \text{örneklem} \\ \text{istatistiği} \end{array} \right) \pm \left(\begin{array}{c} \text{Uygun olan dağılımın} \\ \text{kritik değeri} \end{array} \right) \times \left(\begin{array}{c} \text{örneklem parametresinin} \\ \text{standart sapması} \end{array} \right)$$

Bu bağlamda iki anakütlenin farkına ($\mu_1 - \mu_2$) ait olan güven aralığını denklem (10.4)'teki gibi elde edebiliriz.

$$(\bar{x}_1 - \bar{x}_2) \pm z_{\alpha/2} \left(\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}} \right) \quad (10.4)$$

$\alpha = 0,05$ anlamlılık düzeyinde güven aralığımızı; gerçek anakütle ortalamalarının arasındaki farkın ($\mu_1 - \mu_2$) % 95 ihtimalle bu aralıkta yer aldığı şeklinde yorumlayabiliriz.

Test Prosedürünün Özeti:

Test prosedürlerimizi sağ, ve sol taraflı testleri de ele alarak özetlemeye çalışalım.

Sağ Taraflı Testler:

$$H_0 : \mu_1 - \mu_2 \leq \mu_0$$
$$H_A : \mu_1 - \mu_2 > \mu_0$$

Öncelikle z skoru hesaplanmalıdır

$$z = \frac{(\bar{x}_1 - \bar{x}_2) - \mu_0}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}}$$

Şimdi ise bu z-skoru kullanılarak, z dağılım tablosundan (veya MINITAB programından) p-değeri hesaplanmalı ve p-değeri kuralı uygulanmalıdır.

P-değeri kuralı (α anlamlılık düzeyinde):

Eğer p-değeri $< \alpha$; H_0 **reddedilir**

Eğer p-değeri $> \alpha$; H_0 **reddedilemez**

Alternatif olarak z-skoru kullanılarak kritik değer kuralı uygulanabilir (z kritik değerini z_α bulmak için yine z dağılım tablosundan veya MINITAB programından faydalanabilirsiniz)

Kritik deęer kuralı (α anlamlılık düzeyinde):

z- istatistik skoru $> z_{\alpha}$; H_0 *reddedilir*

z- istatistik skoru $< z_{\alpha}$; H_0 *reddedilemez*

Sol taraflı test:

$$H_0 : \mu_1 - \mu_2 \geq \mu_0$$

$$H_A : \mu_1 - \mu_2 < \mu_0$$

Öncelikle z skoru hesaplanmalıdır

$$z = \frac{(\bar{x}_1 - \bar{x}_2) - \mu_0}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}}$$

Şimdi ise bu z-skoru kullanılarak, z dağılım tablosundan (veya MINITAB programından) p-deęeri hesaplanmalı ve p-deęeri kuralı uygulanmalıdır.

P-deęeri kuralı (α anlamlılık düzeyinde):

Eęer p-deęeri $< \alpha$; H_0 *reddedilir*

Eęer p-deęeri $> \alpha$; H_0 *reddedilemez*

Alternatif olarak z-skoru kullanılarak kritik deęer kuralı uygulanabilir (z kritik deęerini $-z_{\alpha}$ bulmak için yine z dağılım tablosundan veya MINITAB programından faydalanabilirsiniz)

Kritik deęer kuralı (α anlamlılık düzeyinde):

z- istatistik skoru $< -z_{\alpha}$; H_0 *reddedilir*

z- istatistik skoru $> -z_{\alpha}$; H_0 *reddedilemez*

Çift Taraflı Testler:

$$H_0 : \mu_1 - \mu_2 = \mu_0$$

$$H_A : \mu_1 - \mu_2 \neq \mu_0$$

Öncelikle z skoru hesaplanmalıdır

$$z = \frac{(\bar{x}_1 - \bar{x}_2) - \mu_0}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}}$$

Şimdi ise bu z-skoru kullanılarak, z dağılım tablosundan (veya MINITAB programından) p-değeri hesaplanmalı ve p-değeri kuralı uygulanmalıdır.

P-değeri kuralı (α anlamlılık düzeyinde):

Eğer p-değeri $< \alpha$; H_0 **reddedilir**

Eğer p-değeri $> \alpha$; H_0 **reddedilemez**

Alternatif olarak z-skoru kullanılarak kritik değer kuralı uygulanabilir (z kritik değerlerini $\pm z_{\alpha/2}$ bulmak için yine z dağılım tablosundan veya MINITAB programından faydalanabilirsiniz)

Kritik değer kuralı (α anlamlılık düzeyinde):

z- istatistik skoru $> z_{\alpha/2}$ veya z- istatistik skoru $> -z_{\alpha/2}$; H_0 **reddedilir**

z- istatistik skoru $< z_{\alpha/2}$ veya z- istatistik skoru $< -z_{\alpha/2}$; H_0 **reddedilemez**

Güven Aralığı:

$$(\bar{x}_1 - \bar{x}_2) \pm z_{\alpha/2} \left(\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}} \right)$$

Örnek 10.1: İki tane benzer makinenin ürettiği parçalar mm cinsinden yapılmaktadır. Anakütlelerin normal dağıldığını ve birinci makinenin varyansının 9, ikinci makinenin varyansının 16 olduğunu varsayalım. Her iki makinadan da örneklem olarak 35 tane parça alınmıştır. $n_1 = n_2 = 35$. Gerekli veriler aşağıdaki tabloda verilmiştir:

Makina 1

20,39	22,03	24,28	20,49	29,27	22,16	26,04	23,60	29,18	27,66
23,81	23,55	24,76	26,53	27,58	26,89	26,02	28,27	24,16	23,51
27,43	24,53	25,43	21,95	26,49	23,24	28,99	29,31	26,68	25,03
22,92	29,88	26,16	21,72						

Makina 2

20,24	25,78	21,56	24,40	29,72	27,09	25,17	21,38	26,92	27,54	27,81	20,44	28,36
	24,31	24,57	24,50	21,09	20,10	27,12	23,09	26,36	29,38	28,79	23,03	27,03
	25,35	22,96	26,44	28,79	20,81	20,25	24,33	24,34	28,54	29,49		

İki tane makinanın ürettiği parçaların ortalama uzunlukları arasındaki fark için yüzde 90lık bir güven aralığını bulunuz. Aşağıdaki testleri gerçekleştiriniz.

- (a) her iki makinanında ürettiği parçaların ortalama uzunluklarının eşit olduğunu
- (b) her iki makinanında ürettiği parçaların ortalama uzunlukları arasındaki farkın yüzde 5 anlamlılık düzeyinde 1mm'den az olduğunu test ediniz.

Ortalamaları alacak olursak $\bar{x}_1 = 25.174$ ve $\bar{x}_2 = 25.059$. Bu değerleri denklem (10.4)'te yerine koyduğumuzda:

$$(25.174 - 25.059) \pm 1.645 \left(\sqrt{\frac{9}{35} + \frac{16}{35}} \right)$$

$z_{0.1/2} = 1.645$ olduğundan dolayı güven aralığını aşağıdaki şekilde hesaplayabiliriz:

$$0.115 \pm 0.845 \times 1.645 = [-1.275; 1.505]$$

Buna göre, yüzde 90 olasılıkla ortalama farkları $[-1.275; 1.505]$ aralığında yer alır.

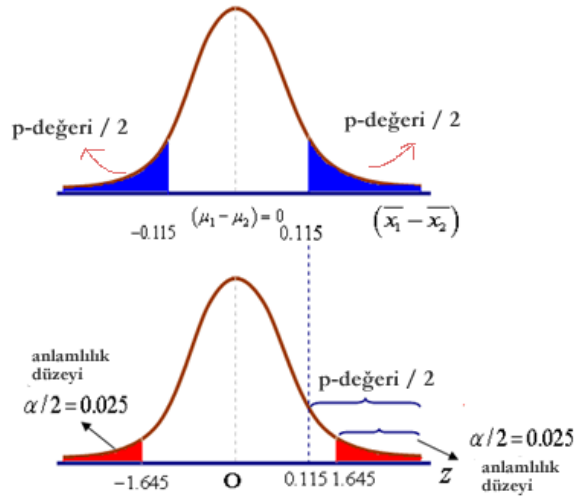
Hipotez aşağıdaki gibi yapılandırılır

$$\begin{aligned} H_0 : \mu_1 - \mu_2 &= 0 \\ H_A : \mu_1 - \mu_2 &\neq 0 \end{aligned}$$

(10.3) formülden z-skorunu hesaplayabiliriz

$$z = \frac{(25.174 - 25.059)}{\sqrt{\frac{9}{35} + \frac{16}{35}}} = \frac{0.115}{0.845} = 0.136$$

Standart normal dağılım tablosundan p değerini kolaylıkla $2 \times P(Z > 0.138) = 2 \times 0.446 = 0.89$ olarak hesaplayabiliriz ve bu değer $0.1 < 0.89$ olduğundan sıfır hipotezini H_0 reddederiz. Veya, alternatif olarak kritik değer kuralından faydalanarak, z istatistik skoru $< z_{\alpha/2}$ ($0.136 < 1.645$) olduğundan yine aynı sonuca varabiliriz. Test prosedürünü grafiksel olarak özetleyecek olursak:



Şekil 10.3

İkinci hipotez ise aşağıdaki gibi yapılandırılabilir:

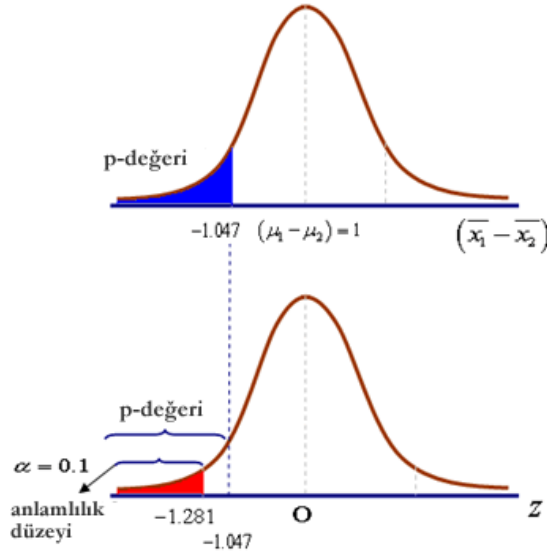
$$H_0 : \mu_1 - \mu_2 \geq 1$$

$$H_A : \mu_1 - \mu_2 < 1$$

Denklem (10.3)'teki z skorunu uyguladığımızda:

$$z = \frac{(25.174 - 25.059) - 1}{\sqrt{\frac{9}{35} + \frac{16}{35}}} = \frac{0.115 - 1}{0.845} = -\frac{0.885}{0.845} = -1.047$$

Standart normal dağılım tablosundan p değerini kolaylıkla $P(Z < -1.047) = 0.1475$ olarak hesaplayabiliriz ve bu değer $0.1 < 0.1475$ olduğundan sıfır hipotezini H_0 reddedemeyiz. Veya, alternatif olarak kritik değer kuralından faydalanarak, z istatistik skoru $> z_{\alpha/2}$ ($1.047 > 1.281$) olduğundan yine aynı sonuca varabiliriz. Test prosedürünü grafiksel olarak özetleyecek olursak:



Şekil 10.4

%10'luk bir anlamlılık düzeyinde hipotezlerimizin hiçbirini reddedilememektedir. Bu iki durumdan, neden “kabul etme” yerine “reddedememe” terimini kullandığımız anlaşılmaktadır. Çünkü, mantık olarak hem $\mu_1 - \mu_2 = 0$ hipotezini hem de $\mu_1 - \mu_2 = 1$ hipotezini aynı anda kabul edemeyiz.

Küçük Örneklerde

Örneklem boyutunun her örneklem için 30 dan küçük olduğu ($n_1 \leq 30$ ve $n_2 \leq 30$), **eğer anakütle varyansları biliniyorsa**, yukarıdaki formülü yeterli bir şekilde uygulayabilmek için anakütlelerin normal dağıldığını varsayımını extra olarak yapmamız gerekmektedir.

Örnek 10.2: Örnek 10.1'deki makinelerin ürettiği parçalar arasından mm cinsinden, birinci makinadan 8 ($n_1 = 8$) ve ikinci makinadan 7 ($n_2 = 7$) tane örneklem alınmıştır. Anakütlelerin 9 ve 16 varyansları ile normal olarak dağıldığını varsaymaktayız. Makina parçalarına ait olan veriler aşağıdaki gibidir:

1. Makina: 23.7 23.0 22.2 24.0 21.2 23.1 27.1 24.0

2. Makina: 23.4 15.3 30.9 18.8 25.3 25.2 32.1

İki tane makinanın ürettiği parçaların ortalama uzunlukları arasındaki fark için yüzde 90lık bir güven aralığını bulunuz ve aşağıdaki testleri gerçekleştiriniz.

- (a) her iki makinanında ürettiği parçaların ortalama uzunluklarının eşit olduğunu
- (b) her iki makinanında ürettiği parçaların ortalama uzunlukları arasındaki farkın yüzde 10 anlamlılık düzeyinde 1mm'e eşit olduğunu test ediniz.

Ortalamaları alacak olursak $\bar{x}_1 = 23.54$ ve $\bar{x}_2 = 24.43$. Bu değerleri denklem (10.4)'te yerine koyduğumuzda:

$$(23.54 - 24.43) \pm \left(\sqrt{\frac{9}{8} + \frac{16}{7}} \right) 1.645$$

$z_{0.1/2} = 1.645$ olduğundan dolayı aralığı aşağıdaki şekilde hesaplayabiliriz:

$$-0.89 \pm 1.847 \times 1.645 = [-3.93; 2.15]$$

Buna göre, yüzde 90 olasılıkla ortalama farkları $[-3.93; 2.15]$. aralığında yer alır. Hipotez aşağıdaki gibi yapılandırılır

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

$$H_A : \mu_1 - \mu_2 \neq 0$$

(10.3) formülden z-skorunu hesaplayabiliriz

$$z = \frac{(23.54 - 24.43)}{\sqrt{\frac{9}{8} + \frac{16}{7}}} = \frac{-0.89}{1.847} = -0.48$$

P değeri $2 \times P(Z < -0.48) = 2 \times 0.3156 = 0.6312$ olarak hesaplanabilmektedir. Böylelikle sıfır hipotezini reddedemeyiz. Veya, alternatif olarak kritik değer kuralından faydalanarak, z istatistik skoru $< z_{\alpha/2}$ ($-0.48 > 1.645$) olduğundan yine aynı sonuca varabiliriz. Dolayısıyla, yüzde 5'lik anlamlılık düzeyinde ortalamaların eşitliğini kabul etmekteyiz.

Standart normal dağılım tablosundan p değerini kolaylıkla $P(Z < 1.047) = 0.1475$. olarak hesaplayabiliriz ve bu değer $0.1 < 0.1475$ olduğundan sıfır hipotezini H_0 reddedemeyiz. Veya, alternatif olarak kritik değer kuralından faydalanarak, z istatistik skoru $> z_{\alpha/2}$ ($1.047 > 1.281$) olduğundan yine aynı sonuca varabiliriz. Test prosedürünü grafiksel olarak özetleyecek olursak:

İkinci hipotez ise aşağıdaki gibi yapılandırılabilir:

$$H_0 : \mu_1 - \mu_2 = 1$$

$$H_A : \mu_1 - \mu_2 \neq 1$$

Denklem (10.3)'teki z skorunu uyguladığımızda:

$$z = \frac{(23.54 - 24.43) - 1}{\sqrt{\frac{9}{8} + \frac{16}{7}}} = \frac{-1.89}{1.847} = -1.02$$

Standart normal dağılım tablosundan p değerini kolaylıkla $P(Z < 1.02) = 2 \times 0.1539 = 0.3078$ olarak hesaplayabiliriz ve bu değer $0.1 < 0.3078$ olduğundan sıfır hipotezini H_0 reddedemeyiz. Veya, alternatif olarak kritik değer kuralından faydalanarak, z istatistik skoru $< -z_{\alpha/2}$ ($-1.02 > -1.645$) olduğundan yine aynı sonuca varabiliriz.

Dolayısıyla, her iki makinasında ürettiği parçaların ortalama uzunlukları arasındaki farkın yüzde 10 anlamlılık düzeyinde 1mm'e eşit olduğunu hipotezi kabul edilmektedir.

İki Ortalama Hakkında Çıkarımlar: Varyanslar Bilinmiyorsa (bağımsız örneklerde)

Varyansların bilinmediği durumlarda (ki pratikte dahat gerçekçi bir durumdur), basit bir sonuç elde edebilmek için varyansların eşit olduğunu varsaymaktayız. Bu varsayım kısıtlayıcı bir varsayım olmasına rağmen, varyansların farklı olması tam güven aralığı değerlerini bulmayı zorlaştırmaktadır.

Büyük Örneklerde : z dağılımı

Eğer varyanslar bilinmiyorsa, ki bu daha gerçekçi bir durum, ancak anakütle varyanslarının eşit olduğunu varsayarsak bir sonuca varabiliriz. Bu varsayım gerçekte doğru olmayabilen kısıtlayıcı bir varsayımdır. Anakütle varyanslarının eşit olmadığına dair bir kanıt varsa, güven aralığını bulacak ya da hipotez testini gerçekleştirecek kesin bir metot kullanamayız. Büyük örneklerde ise örneklem varyansları anakütle varyanslarının bilinen değerlerini yansıttığından dolayı anakütle varyanslarının eşit olduğu varsayımına gerek yoktur. Bu varsayım altında tahmin edilen s^2 örneklem varyansı ve dolayısıyla (10.2) 10.4'teki gibi gerçekleşir:

$$z = \frac{(\bar{x}_1 - \bar{x}_2) - (\mu_1 - \mu_2)}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}} \quad (10.5)$$

s_1^2 örneklem varyansı aşağıdaki formül ile hesaplanmaktadır

$$s_1^2 = \frac{\sum (x_{1i} - \bar{X}_1)^2}{n_1 - 1},$$

s_2^2 örneklem varyansı aşağıdaki formül ile hesaplanmaktadır

$$s_2^2 = \frac{\sum (x_{2i} - \bar{X}_2)^2}{n_2 - 1}.$$

Aşağıdaki hipotezi test etmek istediğimizde

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

$$H_A : \mu_1 - \mu_2 \neq 0$$

z-skoru

$$z = \frac{(\bar{x}_1 - \bar{x}_2)}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}}$$

şeklini alır ve aynı test prosedürü uygulanmaktadır.

Örnek 10.4: Öncelikle Örnek 10.2'yi tekrar ele alalım ve anakütle varyanslarının bilinmediğini varsayalım. Bu yüzden anakütle varyanslarının yerine tahmin edilmiş olan örneklem varyanslarını kullanmamız gerekmektedir. Tahmin edilmiş varyans değerleri $s_1^2 = 3.026^2 = 9.156$ ve $s_2^2 = 2.7^2 = 7.29$ olarak hesaplanmaktadır. Örneklem boyutlarımız çok küçük olduğundan ($n_1 = 8$; $n_2 = 7$), tahmin edilmiş olan varyansları bize bilinen gerçek varyanslar gibi yorumlamamıza izin vermezler. Örneklem boyutlarımız en azından 30 olsalardı, varyanslarımız örneklemeleri yeteri derecede ifade edebilecek güçte olacaktırlardı. Denklem (10.4)'te bulduğumuz değerleri yerine koyacak olursak:

$$(25.174 - 25.059) \pm \left(\sqrt{\frac{9.156}{35} + \frac{7.29}{35}} \right) 1.645$$

$z_{0.1/2} = 1.645$ olduğundan ve daha önceki ünitelerde öğrendiğimiz gibi aralığı aşağıdaki gibi hesaplarız

$$0.115 \pm 0.685 \times 1.645 = [-1.011; 1.241].$$

Örneklem ortalamalarının farkı %90 olasılık ile $[-1.011; 1.241]$ aralığında yer alır.

Hipotezimiz aşağıdaki gibi oluşturulmalıdır.

$$\begin{aligned} H_0 : \mu_1 - \mu_2 &= 0 \\ H_A : \mu_1 - \mu_2 &\neq 0 \end{aligned}$$

(10.3) 'ten z skorunu hesaplayalım

$$z = \frac{(25.174 - 25.059)}{\sqrt{\frac{9.156}{35} + \frac{7.29}{35}}} = \frac{0.115}{0.685} = 0.167$$

P değeri $2 \times P(Z > 0.167) = 2 \times 0.4337 = 0.8674 > 0.1$ olarak hesaplayabiliriz ve bu değer $0.1 < 0.1539$ olduğundan sıfır hipotezini H_0 reddedemeyiz. Veya, alternatif olarak kritik değer kuralından faydalananarak, z istatistik skoru $< z_{\alpha/2}$ ($0.167 > 1.645$) olduğundan yine aynı sonuca varabiliriz.

İkinci hipotez aşağıdaki gibi yapılandırılabilir:

$$\begin{aligned} H_0 : \mu_1 - \mu_2 &\geq 1 \\ H_A : \mu_1 - \mu_2 &< 1 \end{aligned}$$

Denklem (10.3)'ü kullanarak z skorunu hesaplayacak olursak:

$$z = \frac{(25.174 - 25.059) - 1}{\sqrt{\frac{9.156}{35} + \frac{7.29}{35}}} = \frac{0.115 - 1}{0.685} = -1.291$$

p-değeri = $P(Z < -1.291) = 0.098 > 0.1$ olduğundan, sıfır hipotezini reddedebiliriz. Veya, alternatif olarak kritik değer kuralından faydalananarak, z istatistik skoru $< -z_{\alpha}$ ($-1.291 < -1.281$) aynı sonuca varabiliriz ve ortalamaların yüzde 10 anlamlılık düzeyinde eşitliğini reddederiz.

Büyük Örneklerde: t dağılımı

Bazı istatistikçiler z dağılımının yerine t dağılımını kullanmayı tercih etmektedir. Dolayısıyla t skorunu hesaplamak için denklem (10.6)'dan faydalanmaktadırlar.

$$t = \frac{(\bar{x}_1 - \bar{x}_2) - (\mu_1 - \mu_2)}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}} \quad (10.6)$$

dağılımın serbestlik derecesini hesaplamak için ise

$$df = \frac{(s_1^2 / n_1 + s_2^2 / n_2)^2}{\left[\frac{(s_1^2 / n_1)^2}{n_1 - 1} \right] + \left[\frac{(s_2^2 / n_2)^2}{n_2 - 1} \right]}$$

denklemini kullanılmaktadır. Bu bağlamda daha önce deyinmiş olduğumuz test prosedüründe hiçbir değişiklik bulunmamaktadır. Diğer bir taraftan güven aralığı da aşağıdaki gibi hesaplanabilir:

$$(\bar{x}_1 - \bar{x}_2) \pm \left(\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}} \right) t_{\alpha/2} \quad (10.7)$$

Örnek 10.5: Örnek 10.4'teki güven aralığını t dağılımını kullanarak tekrar hesaplayalım. Bu hesaplama da yapmamız gereken tek şey $z_{\alpha/2}$ değerlerinin yerine $t_{\alpha/2}$ değerlerini denklemde yerine koymaktır. Dolayısıyla denklem (10.7)'yi kullanarak güven aralığını aşağıdaki gibi elde edebiliriz:

$$(25.174 - 25.059) \pm \left(\sqrt{\frac{7.29}{35} + \frac{9.156}{35}} \right) 1.6679 = [-1.0283; 1.2583]$$

$t_{0.1/2;67} = \pm 1.6679$ serbestlik derecesini ise

$$df = \frac{\left(\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2} \right)^2}{\left[\frac{\left(\frac{s_1^2}{n_1} \right)^2}{n_1 - 1} \right] + \left[\frac{\left(\frac{s_2^2}{n_2} \right)^2}{n_2 - 1} \right]} = \frac{\left(\frac{9.156}{35} + \frac{7.29}{35} \right)^2}{\left[\frac{\left(\frac{9.156}{35} \right)^2}{(35-1)} \right] + \left[\frac{\left(\frac{7.29}{35} \right)^2}{(35-1)} \right]} = 67.13$$

şekilde hesaplayarak 67 olarak buluruz.

Anakütle ortalamaları arasındaki farkın $[-1.0283; 1.2583]$ arasında yerlacağından %90 ihtimalle eminiz sonucuna varırız.

Hipotezlerin oluşturalım

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

$$H_A : \mu_1 - \mu_2 \neq 0$$

ve bu bağlamda denklem (10.6)'yı kullanarak t skorunu hesaplayalım.

$$t = \frac{(25.174 - 25.059)}{\sqrt{\frac{9.156}{35} + \frac{7.29}{35}}} = \frac{0.115}{0.685} = 0.167$$

P-değeri = $2 \times P(t > 0.167) = 2 \times 0.4339 = 0.867 > 0.1$, olduğundan, sıfır hipotezini reddedemeyiz. Veya, alternatif olarak kritik değer kuralından faydalanarak, t istatistik skoru $< t_{0.1/2; 67}$ ($0.167 < 1.6679$) aynı sonuca varabiliriz. İkinci hipotezimizlerimiz ise:

$$H_0 : \mu_1 - \mu_2 \geq 1$$

$$H_A : \mu_1 - \mu_2 < 1$$

(10.6)'dan t skorunu hesaplayacak olursak

$$t = \frac{(25.174 - 25.059) - 1}{\sqrt{\frac{9.156}{35} + \frac{7.29}{35}}} = \frac{0.115 - 1}{0.685} = -1.291$$

P-değeri = $P(t < -1.291) = 0.1006 > 0.1$, olduğundan, sıfır hipotezini reddedemeyiz. Veya, alternatif olarak kritik değer kuralından faydalanarak, t istatistik skoru $< -t_{0.1/2; 67}$ ($-1.291 > -1.6679$) aynı sonuca varabiliriz ve yine sıfır hipotezini reddedemeyiz.

Küçük Örneklerde: t dağılımı

Eğer örneklem boyutları küçükse, örneklem varyanslarını anakütle varyanslarının bilinen değerleriymiş gibi düşünemeyiz (10.4) veya (10.6) i artık kullanamayız. Bu yüzden küçük örneklerde **eşit anakütle varsayımını** ($\sigma_1^2 = \sigma_2^2$) kullanmak gereklidir. Anaküteler normal dağılmaktadır. Burada ayrıca dikkat edilmesi gereken bir başka husus, z testi yerine geçerli dağılım olarak t dağılımı kullanılmasıdır ve t testi yaparken ortak varyans kullanılmasıdır:

$$t = \frac{(\bar{x}_1 - \bar{x}_2) - (\mu_1 - \mu_2)}{\sqrt{\frac{s_p^2}{n_1} + \frac{s_p^2}{n_2}}} \quad (10.8)$$

$$df = n_1 + n_2 - 2$$

$$s_p^2 = \frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{(n_1 - 1) + (n_2 - 1)} \quad (10.9)$$

Diğer bir taraftan güven aralığı denklem (10.10)'daki gibidir.

$$(\bar{x}_1 - \bar{x}_2) \pm \left(\sqrt{\frac{s_p^2}{n_1} + \frac{s_p^2}{n_2}} \right) t_{\alpha/2} \quad (10.10)$$

Örnek 10.6: Örnek 10.2'yi ele alalım ve anakütle varyanslarını bilmediğimizi varsayalım. Küçük örnekler durumunda olduğumuzdan, anakütle varyanslarının eşit olarak dağıldığını varsaymak zorundayız. Bu bağlamda denklem (10.9)'u uygulayarak ortak varyans değerini hesaplayabiliriz.

$$s_p^2 = \frac{(8-1)2.98 + (7-1)36.36}{(8-1) + (7-1)} = 18.38$$

Denklem (10.10) yardımıyla güven aralığını hesaplayabiliriz:

$$(23.54 - 24.43) \pm \left(\sqrt{\frac{18.38}{8} + \frac{18.38}{7}} \right) 1.77$$

$t_{0.1/2; 13} = 1.77$ ve $df = 7 + 8 - 2 = 13$ olduğundan güven aralığı:

$$-0.89 \pm 2.218 \times 1.943 = [-4.82; 3]$$

Anakütle ortalamaları arasındaki farkın $[-4.82; 3]$ arasında yeracağından %90 ihtimalle eminiz sonucuna varırız.

İlk hipotezi oluşturacak olursak:

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

$$H_A : \mu_1 - \mu_2 \neq 0$$

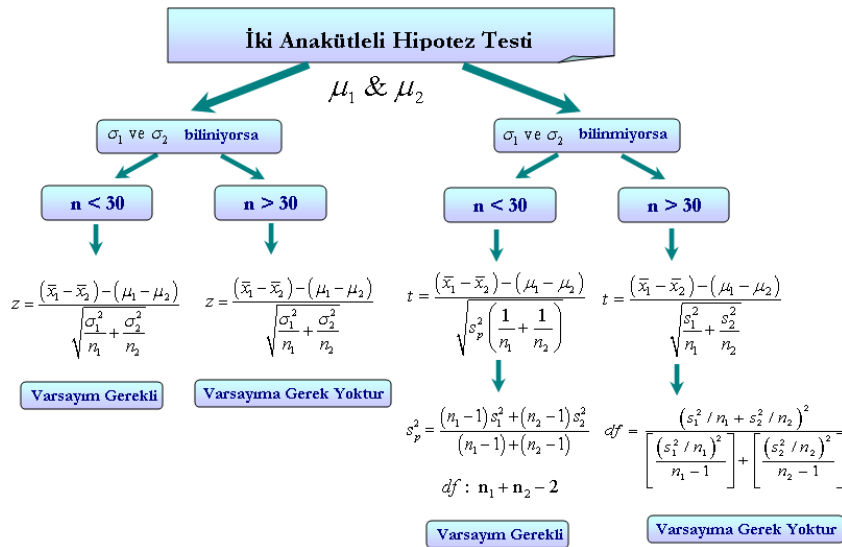
Daha sonra denklem (10.8)'den faydalanarak t skorunu hesaplayabiliriz

$$t = \frac{(23.54 - 24.43)}{\sqrt{\frac{18.38}{8} + \frac{18.38}{7}}} = \frac{-0.89}{2.218} = -0.4$$

t istatistik skoru $< -t_{n-1;\alpha/2}$ ($-0.4 > -1.77$) olduğundan, sıfır hipotezini reddedemeyiz. İkinci hipotezi de benzer şekilde oluşturup test edebilirsiniz.

Anakütle Ortalamaları için Test İstatistiklerinin Özeti

Şekil 10.5'te, iki anakütle ortalamalarının hipotez testlerinde, ne zaman hangi testin kullanılacağını, hangi durumlarda varsayımların gerektiğini özetlemektedir.

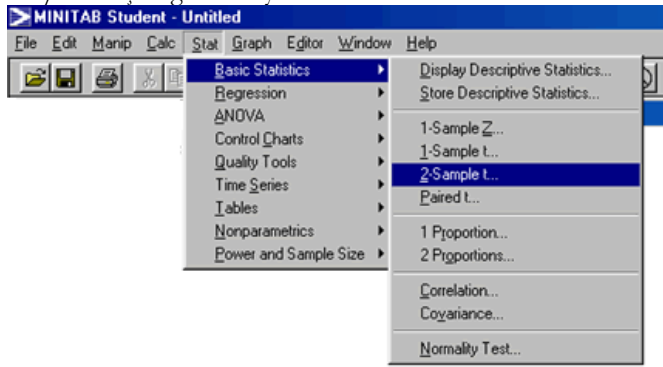


Şekil 10.5

İki Anakütlenin Ortalama Farkları

Minitab Uygulamaları

Örnek 10.2 ve 10.5'teki t testlerin MINITAB ortamında gerçekleştirilebilmesi için öncelikle kullanacağımız verisetini "*two-machines1.mtw*" dosyasında bulabilirsiniz. İki anakütle ortalamasının farkının hipotez testi ve güven aralığını hesaplamak için "*Stat*" "*Basic Statistics*" menüsünden "*2-Sample t*" seçeneğine tıklayınız.



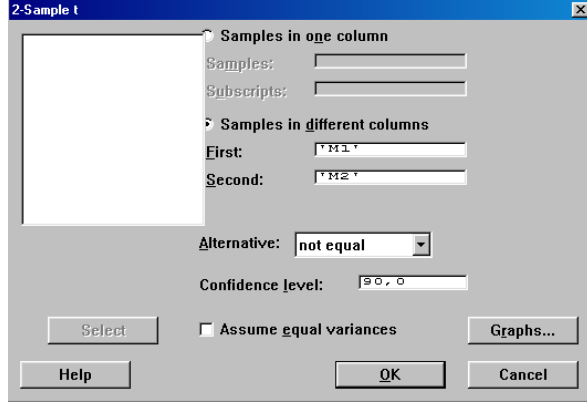
Şekil 10.6

Test edeceğimiz hipotez:

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

$$H_A : \mu_1 - \mu_2 \neq 0$$

karşınıza gelen diyalog ekranında "*samples in different columns*" seçeneğine tıklayınız ve birinci sütun olarak "*First*" M1, ikinci sütun olarak "*Second*" M2 değerini giriniz ve "*Alternative*" bölümünü varsayılan değer gibi bırakınız. ("*not equal*" olarak bırakınız). Güven aralığı düzeyini ise "90" olarak değiştirdikten sonra "*Assume equal variances*" seçeneğine tıklamayınız. Bu işlemler sonucunda karşınızdaki ekran 10.7'deki gibi olmalıdır.



Şekil 10.7

“OK” tuşuna bastığınızda aşağıdaki çıktı ekranını elde edebilirsiniz.

Two Sample T-Test and Confidence Interval				
Two sample T for M1 vs M2				
	N	Mean	StDev	SE Mean
M1	35	25,17	2,70	0,46
M2	35	25,06	3,03	0,51
90% CI for mu M1 - mu M2: (-1,03; 1,26)				
T-Test mu M1 = mu M2 (vs not =): T = 0,17 P = 0,87 DF = 67				

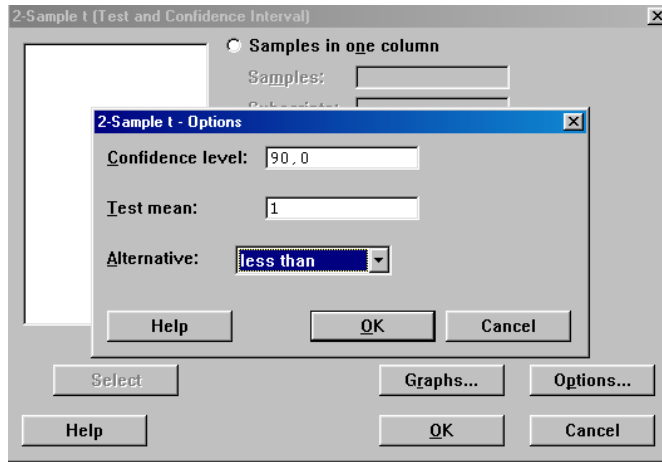
Şekil 10.8

Çıktı ekranındaki elde ettiğimiz güven aralığı Örnek 10.5’tekinin [-1.0283; 1.2583] aynısıdır. Örnek 10.5’teki t değerini elde ettiğimiz çıktıda da 0.17 olarak, p-değerini de yine 0.87 olarak hesaplayabiliriz. Dolayısıyla, p değeri anlamlılık düzeyinden büyük olduğundan sıfır testini reddedemeyiz. Serbestlik derecemiz de Örnek 10.5’te olduğu gibi serbestlik derecesi (DF) 67 olarak belirtilmiştir. Örnek 10.5’teki ikinci testi gerçekleştirmek için :

$$H_0 : \mu_1 - \mu_2 \geq 1$$

$$H_A : \mu_1 - \mu_2 < 1$$

MINITAB programının full sürümü olan 13. versiyonunu kullanmaya ihtiyacınız vardır. Şekil 10.9’daki diyalog ekranını elde ettikten sonra “Options” bölümüne tıklayıp, “Confidence level” (Güven aralığı) bölümüne 90, “Test mean” bölümüne 1 değerini girerek “Alternative” bölümünden de “less than” (daha az) seçeceğini seçiniz.



Şekil 10.9

“OK” tuşuna basarak öncelikle anakütle diyalog ekranına geri dönünüz ve daha sonra aşağıdaki çıktıyı elde etmek için tekrar “OK” tuşuna basınız.

Two-Sample T-Test and CI: M1; M2

Two-sample T for M1 vs M2

	N	Mean	StDev	SE Mean
M1	35	25,17	2,70	0,46
M2	35	25,06	3,03	0,51

Difference = μ M1 - μ M2

Estimate for difference: 0,115

90% upper bound for difference: 1,002

T-Test of difference = 1 (vs <): T-Value = -1,29 P-Value = 0,100 DF = 67

Şekil 10.10

Sonuçlar incelendiğinde çıktıda belli bir güven aralığının olmadığı yalnızca %90’lık aralığın üst sınırının verildiği açıkça gözlenebilmektedir. Bunun nedeni testimizin tek taraflı bir test olmasındandır. Testimiz çift taraflı bir test olsaydı hem üst sınırın hem de alt sınırın belirtildiği güven aralığı sonuçlar arasında belirtilecekti. Şimdi örnek 10.6’ya geri dönelim ve 10.7’deki diyalog ekranını aynı şekilde doldurarak, ancak bu kez örneklem boyutu 30’dan küçük olduğundan “Assume equal variances” seçeneğine tıklayalım. Birinci hipotezimizin çıktısı bu durumda:

Two Sample T-Test and Confidence Interval				
Two sample T for M1 vs M2				
	N	Mean	StDev	SE Mean
M1	8	23,54	1,73	0,61
M2	7	24,43	6,03	2,3
90% CI for μ M1 - μ M2: (-4,82; 3,0)				
T-Test μ M1 = μ M2 (vs not =): T = -0,40 P = 0,69 DF = 13				
Both use Pooled StDev = 4,29				

Şekil 10.11

Diğer bir taraftan ikinci hipotezin çıktısı aşağıdaki gibidir.

Two-Sample T-Test and CI: M1; M2				
Two-sample T for M1 vs M2				
	N	Mean	StDev	SE Mean
M1	8	23,54	1,73	0,61
M2	7	24,43	6,03	2,3
Difference = mu M1 - mu M2				
Estimate for difference: -0,89				
90% CI for difference: (-4,82; 3,04)				
T-Test of difference = 1 (vs not =): T-Value = -0,85 P-Value = 0,410 DF = 13				
Both use Pooled StDev = 4,29				

Şekil 10.12

Anakütle Ortalamaları Hakkında Çıkarımlar: Bağımlı ve Küçük veya Büyük Örneklemeler

Bağımlı örneklemeler, her konudan iki değer çıkardığımız durumlarda ortaya çıkmaktadır. Örnek olarak, öğrencilerin mikroekonomi ve makroekonomi derslerindeki performansları ölçmek istediğimizde her öğrencinin iki dersten aldıkları notları birden inceleriz. Dolayısıyla elde ettiğimiz örneklemeler birbirlerine bağımlıdır. Öğrencilerin iki farklı derste aldıkları notlar arasında bir farklılaşma var mıydı? Bu testi gerçekleştirebilmek için bu farkın anakütle ortalamasını temsil etmek amacıyla μ_d değişkenini tanımlayalım. Artık test etmek istediğimiz hipotezi sembollerle ifade edebiliriz:

$$H_0 : \mu_d = 0$$

$$H_A : \mu_d \neq 0$$

Örneklem olarak seçilmiş olan öğrencilerin mikroekonomi ve makroekonomi derslerinden aldıkları notlar ikililer halinde (42,44), (48,54), (55,51), (64,64), (71,70), (39,29), (57,48), (58,53), (68,61), (60,57), (50,58) olarak ifade edilebilir.

Testi gerçekleştirebilmek için öncelikle bu iki dersin notları arasındaki farkların hesaplanması gerekmektedir. Bu farkları alacak olursak:

$$-2, -6, 4, 0, 1, 10, 9, 5, 7, 3, \text{ ve } -8.$$

Sıfır hipotezi reddedildiği takdirde, iki dersin notları arasında bir farklılık olduğu sonucuna varırız. Fark değerleri için örneklem ortalamasını ve örneklem standart sapmasını aşağıdaki gibi hesaplayabiliriz:

örneklem ortalaması için,

$$\bar{d} = \frac{\sum_{i=1}^n d_i}{n}$$

ve örneklem standart sapması için

$$s_d = \sqrt{\frac{\sum_{i=1}^n (d_i - \bar{d})^2}{n-1}}$$

formülleri kullanılabilir.

Verilerimizi ele aldığımızda, farkların ortalama ve standart sapma değerlerini aşağıdaki gibi hesaplayabiliriz:

$$\bar{d} = \frac{23}{11} = 2.091, s_d = \sqrt{\frac{336.909}{10}} = 5.804$$

Anakütlenin normal olarak dağıldığı varsayıldığı takdirde, anakütle ortalaması hakkındaki test gerçekleştirilirken n-1 serbestlik derecesinde t testi kullanılmalıdır. Hipotez testinde kullanılacak olan test istatistiği denklem (10.11)'deki gibi olacaktır:

$$t = \frac{\bar{d} - \mu_d}{\frac{s_d}{\sqrt{n}}} \quad (10.11)$$

μ_d değerinin yerine sıfır hipotezimizi dikkate alarak $\mu_d = 0$ değerini koyduğumuz takdirde test istatistiğimiz:

$$t = \frac{\bar{d}}{\frac{s_d}{\sqrt{n}}} \quad (10.12)$$

Verilerimizi denklem (10.12)'de yerine koyduğumuzda istatistik değerimiz:

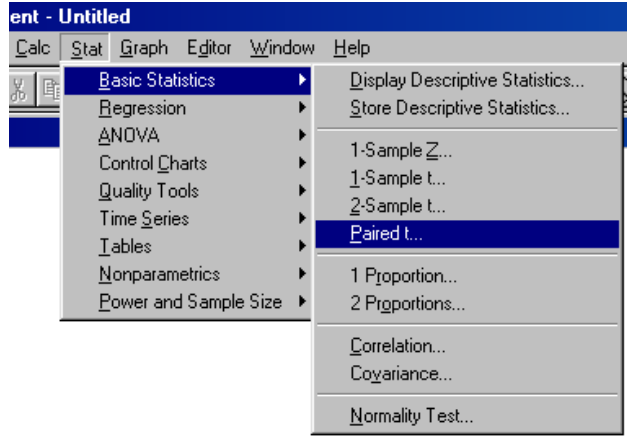
$$t = \frac{2.091}{\frac{5.804}{\sqrt{11}}} = 1.1949$$

Hipotez testimizin yapılanmasından da anlaşılacağı gibi testimiz çift taraflı bir testtir. 10 (n-1) serbestlik derecesinde ve $\alpha = 0.05$ anlamlılık düzeyinde t kritik değerimiz $t_{0.05/2;10} = 2.228$ 'dir. $1.1949 < 2.228$ olduğundan anakütle ortalamalarının arasında bir fark olmadığı sıfır hipotezi reddedilememektedir. Alternatif olarak p-değerini de hesapladığımızda da p-değeri $= 2 \times P(t < 1.1949) = 2 \times 0.1298 = 0.2596$ olarak bulunur ve bu değer de 0.05 anlamlılık düzeyinden büyük olduğundan yine aynı sonuca ulaşmış oluruz. Dolayısıyla, öğrencilerin mikroekonomi ve makroekonomi notlarından aldıkları notlar arasında bir farklılık yoktur.

Anakütle Ortalamaları Arasındaki Farklar (Bağımlı Örneklemeler)

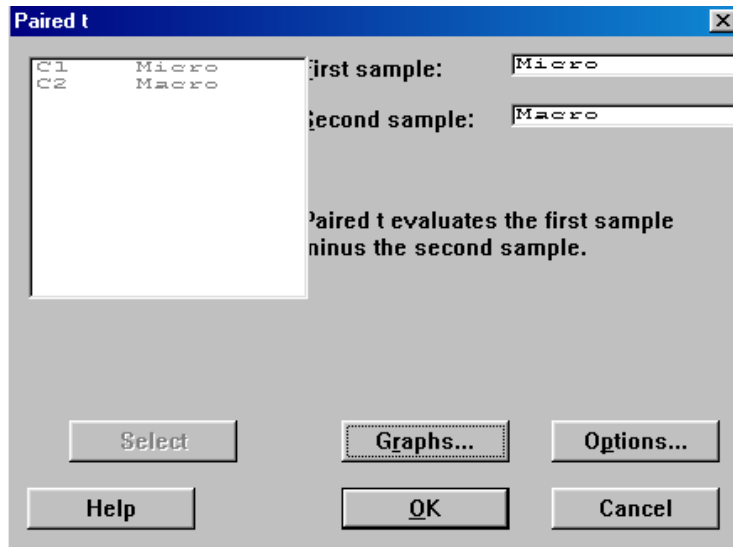
Minitab Uygulamaları 

Mikroekonomi ve makroekonomi derslerinde öğrencilerin aldıkları notların farklılığını MINITAB ortamında test etmek istersek öncelikle “*microandmacro.mtw*” verisetini açmamız gerekmektedir. Hemen akabinde “Stat” - “Basic Statistics” menüsünden “Paired t” (eşli t) seçeneğine tıklayınız.



Şekil 10.13

“*First sample*” (Birinci Örneklem) kutucuğuna “*Micro*” ; “*Second sample*” (İkinci Örneklem) kutucuğuna “*Macro*” değişkenini Şekil 10.14’te görüldüğü gibi atayınız.



Şekil 10.14

“OK” tuşuna bastığınızda aşağıdaki çıktı ekranını elde edebilirsiniz.

Paired T-Test and Confidence Interval				
Paired T for Micro - Macro				
	N	Mean	StDev	SE Mean
Micro	11	55,64	10,19	3,07
Macro	11	53,55	10,95	3,30
Difference	11	2,09	5,80	1,75
95% CI for mean difference: (-1,81; 5,99)				
T-Test of mean difference = 0 (vs not = 0): T-Value = 1,19 P-Value = 0,260				

Şekil 10.15

İki Oran Hakkında Çıkarımlar

İki farklı anakütlenin oranları arasındaki ilişki hakkında çıkarımlarda bulunabilmek için, $n\pi \geq 5$ ve $n(1-\pi) \geq 5$ koşullarının sağlandığı rassal olarak seçilen birbirinden bağımsız iki örneklem setinin olduğu varsayılmalıdır.

Birinci anakütle için:

π_1 = birinci anakütlenin oranı

n_1 = birinci örneklemin boyutu

x_1 = birinci örneklemdaki başarı sayısı

$p_1 = \frac{x_1}{n_1}$ = birinci örneklemin oranı

İkinci anakütle için:

π_2 = ikinci anakütlenin oranı

n_2 = ikinci örneklemin boyutu

x_2 = ikinci örneklemdaki başarı sayısı

$p_2 = \frac{x_2}{n_2}$ = ikinci örneklemin oranı

p_1 ve p_2 'nin ortak oranını tanımlarsak:

$$p = \frac{x_1 + x_2}{n_1 + n_2}$$

Ortak oran, oranlar için hipotez testi gerçekleştirildiğinde eşit oranlar varsayımı sözkonusu olduğunda iki örneklem için en iyi tahmin edici olarak kullanılmaktadır.

Eşit anakütle oranlarının eşitliğini test etmek istediğimizde hipotez testini aşağıdaki gibi oluşturabiliriz:

$$H_0 : \pi_1 - \pi_2 = 0$$

$$H_A : \pi_1 - \pi_2 \neq 0$$

Veya alternatif olarak anakütle oranlarının belli bir değere eşit olduğunu test etmek istediğimizde ise hipotez testini aşağıdaki gibi oluşturabiliriz:

$$H_0 : \pi_1 - \pi_2 = \pi_0$$

$$H_A : \pi_1 - \pi_2 \neq \pi_0$$

Oranlar arasındaki farkı test etmek için, örneklem istatistiğini $(p_1 - p_2)$ şeklinde ifade edebilmemiz gerekmektedir. Bu bağlamda öncelikle $(p_1 - p_2)$ 'nin örneklem dağılımını bulmamız gerekmektedir. Bir önceki bölümden hatırlayacağınız üzere oranlardan her biri p_1 ve p_2 binom dağılımı ile dağılmaktadır. Büyük örneklem boyutlarında bu oranlar normal olarak dağılmaktadır. n 'in büyük değerleri için $(n\pi \geq 5$ ve $n(1-\pi) \geq 5)$ p_1 ve p_2 'nin herbiri; π_1 , π_2 ortalamaları ve $\frac{\pi_1(1-\pi_1)}{n_1}, \frac{\pi_2(1-\pi_2)}{n_2}$ varyansları ile normal olarak dağılmaktadırlar.

$$p_1 \rightarrow N\left(\pi_1, \frac{\pi_1(1-\pi_1)}{n_1}\right)$$

$$p_2 \rightarrow N\left(\pi_2, \frac{\pi_2(1-\pi_2)}{n_2}\right)$$

Örneklem oranları arasındaki fark $(p_1 - p_2)$,

$$E(p_1 - p_2) = \pi_1 - \pi_2$$

ortalaması

$$\text{var}(p_1 - p_2) = \sigma_{p_1 - p_2}^2 = \sigma_{p_1}^2 + \sigma_{p_2}^2 = \frac{\pi_1(1-\pi_1)}{n_1} + \frac{\pi_2(1-\pi_2)}{n_2}$$

ve varyansı ile normal olarak dağılmaktadır.

Dolayısıyla,

$$(p_1 - p_2) \rightarrow N\left(\pi_1 - \pi_2, \frac{\pi_1(1-\pi_1)}{n_1} + \frac{\pi_2(1-\pi_2)}{n_2}\right)$$

Denklem (10.1)'deki genel yapıyı ele alacak olursak, testte kullanılacak olan z-skorunun formülü denklem (10.13)'teki gibidir:

$$z = \frac{(p_1 - p_2) - \pi_0}{\sqrt{\frac{\pi_1(1-\pi_1)}{n_1} + \frac{\pi_2(1-\pi_2)}{n_2}}} \quad (10.13)$$

π_1 ve π_2 anakütle oranları bilinmediğinden, bu parametrelerin yerine örneklem tahminleri olan p_1 ve p_2 kullanılmaktadır. Dolayısıyla denklem (10.13) aşağıdaki gibi olmaktadır.

$$z = \frac{(p_1 - p_2) - \pi_0}{\sqrt{\frac{p_1(1-p_1)}{n_1} + \frac{p_2(1-p_2)}{n_2}}} \quad (10.14)$$

Güven aralığı ise aşağıdaki gibi hesaplanmaktadır.

$$(p_1 - p_2) \pm \left(\sqrt{\frac{p_1(1-p_1)}{n_1} + \frac{p_2(1-p_2)}{n_2}} \right) z_{\alpha/2} \quad (10.15)$$

Aşağıdaki gibi hipotezleri test edebilmek için:

$$H_0 : \pi_1 - \pi_2 = 0 \text{ veya } \pi_1 = \pi_2$$

$$H_A : \pi_1 - \pi_2 \neq 0 \text{ veya } \pi_1 \neq \pi_2$$

$$H_0 : \pi_1 - \pi_2 \geq 0 \text{ veya } \pi_1 = \pi_2$$

$$H_A : \pi_1 - \pi_2 < 0 \text{ veya } \pi_1 \neq \pi_2$$

$$H_0 : \pi_1 - \pi_2 \leq 0 \text{ veya } \pi_1 = \pi_2$$

$$H_A : \pi_1 - \pi_2 > 0 \text{ veya } \pi_1 \neq \pi_2$$

π_0 değerinin 0'a eşit olduğu varsayıldığında, her iki anakütlenin oranları π_1 , π_2 eşit olmaktadır. Dolayısıyla denklem (10.13) $\pi_1 = \pi_2 = \pi$ eşitliği kullanılarak aşağıdaki skoru elde edebiliriz.

$$z = \frac{(p_1 - p_2) - \pi_0}{\sqrt{\frac{\pi_1(1-\pi_1)}{n_1} + \frac{\pi_2(1-\pi_2)}{n_2}}} = \frac{(p_1 - p_2) - \pi_0}{\sqrt{\pi(1-\pi)\left(\frac{1}{n_1} + \frac{1}{n_2}\right)}}$$

Fakat π değerini bilmediğimizden bu değer yerine bu değer ortak tahmin edicisini p 'yi kullanırız:

$$z = \frac{(p_1 - p_2) - \pi_0}{\sqrt{\frac{p(1-p)}{n_1} + \frac{p(1-p)}{n_2}}} = \frac{(p_1 - p_2) - (\pi_1 - \pi_2)}{\sqrt{p(1-p)\left(\frac{1}{n_1} + \frac{1}{n_2}\right)}} \quad (10.16)$$

$$p = \frac{x_1 + x_2}{n_1 + n_2}$$

Örnek 10.7: Firmamız üretim parkı için alacağı makinanın seçiminde iki farklı markaya sahip makina arasında kararsız kalmıştır. Her iki makinanın üretim güvenilirliği için her makinanın ürettiği mallardan 200 adet örneklem seçilmiştir. Birinci makinanın ürettiği mallardan 9'u, ikinci makinanın ürettiği mallardan 14 tanesi defoludur. Her iki makinanın defoluluk oranları arasında bir farklılık olup olmadığını test ediniz?

Fark olup olmadığı ile alakalı olan sıfır ve alternatif testlerimizi oluşturacak olursak:

$$H_0 : \pi_1 - \pi_2 = 0$$

$$H_A : \pi_1 - \pi_2 \neq 0$$

anlamlılık düzeyini de $\alpha = 0.05$ olarak alıp, ortak oranı hesaplayalım p :

$$p = \frac{x_1 + x_2}{n_1 + n_2} = \frac{9 + 14}{200 + 200} = 0.0575$$

Hesapladığımız ortak oran değerini denklem (10.16)'da yerine koyacak olursak, z istatistik skorunu aşağıdaki gibi hesaplayabiliriz:

$$z = \frac{(0.045 - 0.07)}{\sqrt{0.0575(1 - 0.0575)\left(\frac{1}{200} + \frac{1}{200}\right)}} = -1.0739$$

Şimdi ise elde ettiğimiz z istatistik skorunu, z istatistik tablosundan elde ettiğimiz z kritik değeri ($z_{0.05/2} = -1.96$) ile karşılaştırmamız gerekmektedir. z -istatistik skoru $> -z_{\alpha/2}$ olduğundan her iki makinanın defoluluk oranlarının arasında bir fark olmadığı sıfır hipotezini reddederiz.

Güven aralığını hesaplamak için, $p_1 = 9/200 = 0.045$ ve $p_2 = 14/200 = 0.07$ değerlerini denklem (10.15)'te kullanmamız gerekmektedir.

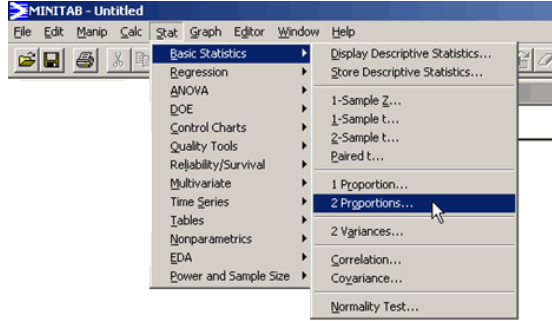
$$(0.045 - 0.07) \pm \left(\sqrt{\frac{0.045(1 - 0.045)}{200} + \frac{0.07(1 - 0.07)}{200}} \right) 1.96$$

Dolayısıyla güven aralığı aşağıdaki gibi bulunabilir.

$$-0.025 \pm 0.023 \times 1.96 = [-0.07; 0.02]$$

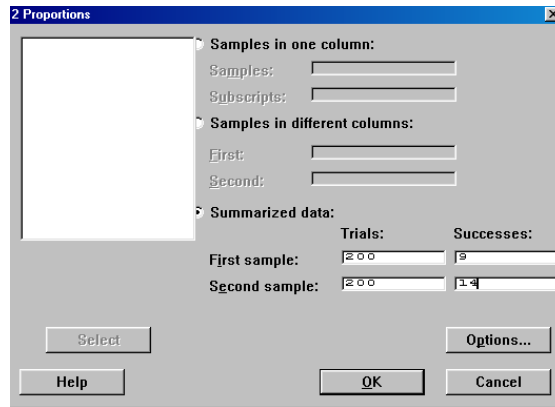
İki Oran Hakkında

MINITAB ortamında örnek 10.7'yi uygulamak için öncelikle “*Stat-Basic statistics*” menüsünden “*2 proportions*” seçeneğine tıklayınız.



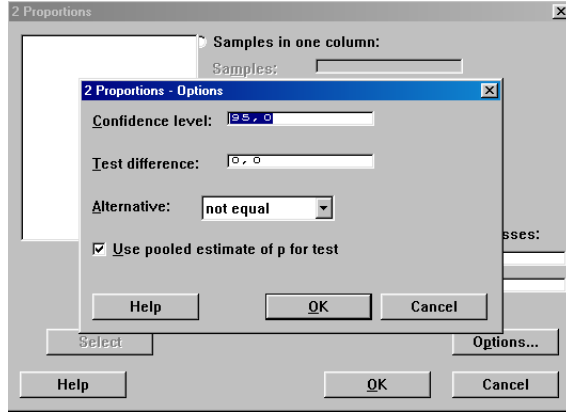
Şekil 10.16

Karşınıza çıkacak olan diyalog ekranından “*Summarized data*” (Özetlenmiş veri) seçeneğine tıkladıktan sonra, “*Trials*” (Deneme Sayısı) kutucuğuna 200, “*Successes*” (Başarı) kutucuğuna birinci örneklem için 9, ikinci örneklem için 14 değerini giriniz.



Şekil 10.17

Daha sonra “*Options*” tuşunu kullanarak “*Used pooled estimate of p for test*” (*p’nin ortak oranını kullan*) seçeneğini aktif hale getiriniz.



Şekil 10.18

OK tuşuna basarak anakütle diyalog ekranına dönüp, ana ekranda da OK tuşuna bastığınızda istediğimiz çıktı ekranını elde ederiz.

Test and Confidence Interval for Two Proportions				
Sample	X	N	Sample p	
1	9	200	0,045000	
2	14	200	0,070000	
Estimate for $p(1) - p(2)$: -0,025				
95% CI for $p(1) - p(2)$: (-0,0705613; 0,0205613)				
Test for $p(1) - p(2) = 0$ (vs not = 0): Z = -1,08 P-Value = 0,282				

Şekil 10.19

Bölüm 11: Uyum veya Bağımlılık Testleri

Giriş

Bu bölüme kadar hep sayısal verilerle alakalı olan testleri ele aldık. Bu bölümde ise kategorik verilerin kullanıldığı testlerin nasıl yapıldığını inceleyeceğiz. Kategorik veriler, sayısal olmayan farklı hücrelere göre gruplandırıldığından bazen nitel veriler olarak da adlandırılabilirler. Kategorik verilerden oluşan tablolara kontenjan tablosu adı altında Bölüm 4'te biraz değinmiştik.

Ki-Kare Dağılımı ve Uyum Bağımlılığı Testleri

Ki-kare dağılımı, belli kategorilerde yer alan gözlenen veya gerçek gözlem değerlerini, tahmin edilmiş veya beklenen değerlerle karşılaştırılmasının testlerinde kullanılır. Uyum bağımlılığı testi ise gözlenen frekans değerleri ile beklenen frekans değerleri arasındaki farklılaşmalara odaklanmıştır. Beklenen ve gözlenen frekans değerleri arasındaki büyük farklar, kategorilerin birbirinden bağımsız olduğu hipotezini kuvvetlendirmektedir. Ki-kare testi, gözlenen frekans değerlerinde meydana gelen farklılıkların anlamlı olup olmadığını test etme çabasıdır. Gözlenen frekansların, bilinen dağılımlara uyumunu incelediğinden dolayı aynı zamanda uyum bağımlılığı testi denilmektedir. Uyum bağımlılığı testinin test istatistiği aşağıdaki gibidir:

$$\chi^2 = \frac{\sum_{i=1}^c (f_i - e_i)^2}{e_i} \quad (11.1)$$

f_i : i kategorisindeki gözlenen frekans değeri,

e_i : i kategorisindeki beklenen frekans değeri, (sıfır hipotezinin doğru olduğu baz alınmıştır)

c : kategori sayısı

Test istatistiğinin serbestlik derecesi $(c-1)$ 'dir ve ki-kare dağılımına sahiptir. Ki-kare testi gerçekleştirilirken aşağıdaki adımlar izlenir:

1. Sıfır hipotezini oluşturunuz.
2. n tane gözlemden oluşan rassal örneklemini kullanarak, her c kategorisinde yer alan gözlenen frekansları kaydediniz.
3. Sıfır hipotezinin doğru olduğu varsayımı altında, her c kategorisine denk gelen olasılığı veya oranı hesaplayınız.
4. Üçüncü adımda hesapladığımız kategori oranlarını kullanarak beklenen frekans değerlerini hesaplayınız.
5. Ki-kare test değerini hesaplamak için daha önceki adımlarda hesapladığımız gözlenen ve beklenen frekans değerlerini hesaplayınız..
6. Reddetme kuralı için aşağıdaki kritik kuralı kullanınız:

Eğer $\chi^2 > \chi^2_{\alpha, c-1}$ ise, H_0 hipotezi reddedilir. (α anlamlılık düzeyini göstermektedir)

Örnek

Geçtiğimiz yılda birbirine rakip olan üç perakende firmasının (Migros, Tesco, Gima) piyasa payları aşağıdaki gibidir: Tesco'nun piyasadaki payı (PT olarak ifade edelim) 0.4, Migros'un piyasadaki payı (PM olarak ifade edelim) 0.4, ve Gima'nın piyasadaki payı (PG olarak ifade edelim) ise 0.2'dir. Piyasa paylarının açıklanmasından sonra Gima yöneticileri, piyasa paylarını arttırmak için indirim yapan ve taksit yapma imkanı tanıyan bir kredi kartı piyasaya sürmeye karar verir. Bu uygulamanın, Gima'nın piyasa payını artırıp arttırmayacağını öğrenmek için Gima yönetimi 200 tüketiciyi kapsayacak bir anket hazırlamıştır. Acaba yeni kart uygulaması Gima'nın piyasa payına bir katkıda bulunacak mı?

Sıfır ve alternatif hipotezlerini oluşturacak olursak (adım (1)):

H_0 : PT = 0.4; PM = 0.4; PG = 0.2:

H_A : Anakütle oranları sıfır hipotezindeki değerlere eşit olmadığı durum

Anket sonuçlarına göre, 75 kişi Tesco'da alışveriş yaparken, 71 kişi Migros'ta, 54 kişi ise Gima'da alışveriş yapmaktadır (adım (2)). Eğer sıfır hipotezi doğru ise (piyasa payları değişmemişse), 80 kişinin Tesco'dan ($e_1 = 0.4 \times 200 = 80$), 80 kişinin Migros'tan ($e_2 = 0.4 \times 200 = 80$), ve 40 kişinin Gima'dan ($e_3 = 0.2 \times 200 = 40$) alışveriş etmesini bekleriz. (adım (3) ve (4)). Birinci adımı kullanarak test istatistikimizi yazacak olursak (adım (5)):

$$\begin{aligned}\chi^2 &= \frac{(75-80)^2}{80} + \frac{(71-80)^2}{80} + \frac{(54-40)^2}{40} \\ &= 0.3125 + 1.0125 + 4.9 = 6.225\end{aligned}$$

Üç farklı kategori bulunduğundan ($c = 3$), ki-kare dağılımı $c - 1$ veya 2 serbestlik derecesi ile dağılmaktadır. Şimdi ise ki-kare tablosunu kullanarak, $\alpha = 0.05$ anlamlılık düzeyi ve 2 serbestlik derecesine tekabül eden kritik değeri bulmamız gerekmektedir.

Bulduğumuz ki-kare kritik değeri $\chi^2_{0.05, 2} = 5.99$ olduğundan ve ki-kare istatistik değeri olan 6.225'ten küçük olduğundan sıfır hipotezimizi reddetmemiz gerekmektedir. Böylelikle, Gima'nın piyasaya sürmeyi düşündüğü kart sayesinde piyasa payında bir değişiklik meydana gelecektir, sonucuna varılır (adım (6)).

Ki-kare Dağılımı ve Bağımsızlık Testleri

Bu durumda ise iki kategorik değişken arasında bir ilişki olup olmadığını (birbirinden bağımsız olup olmadığını) test ederiz. Bağımsızlık testinin test istatistiği denklem (1)'e çok benzemektedir. Fakat bu durumda iki farklı değişken olduğundan dolayı, r kategori bir değişken için, c kategoride bir değişken için yer almaktadır. Bu durum Bölüm 4'te incelediğimiz kontenjan tabloları konusuna çok benzemektedir. Gözlenen ve beklenen frekanslara dayalı χ^2 değerini hesaplamak için:

$$\chi^2 = \frac{\sum_{i=1}^r \sum_{j=1}^c (f_{ij} - e_{ij})^2}{e_{ij}} \quad (11.2)$$

Test istatistiğimiz, $(r-1) \times (c-1)$ serbestlik derecesi ile ki-kare dağılımına sahiptir. Hipotez testinin geri kalan adımlarında bir değişiklik yoktur.

Örnek

Şimdi bağımsızlık testini bir örnek ile incelemeye çalışalım. Öncelikle, politik bir anket düzenleyelim ve bu anketle üç ana siyasi partinin oylarının dağılımında insanların tercihlerinin cinsiyetlerinden bağımsız olup olmadığını araştıralım. Bu amacımızı istatistiksel olarak sıfır hipotezi ve alternatif hipotez kavramlarını göz önünde bulundurarak aşağıdaki gibi oluşturabiliriz:

H_0 : Seçim tercihleri insanların cinsiyetlerinden bağımsızdır

H_A : Seçim tercihleri insanların cinsiyetlerinden bağımsız değildir:

150 kişilik rassal bir şekilde potansiyel oy verecek seçmenlerden oluşan bir anket yapıldığını düşünelim. Bu anket hakkındaki bilgi Tablo 11.1’de özetlenmiştir. Tabloyu inceleyecek olursak; 3 sütun insanların oy vermeyi düşündükleri parti türü hakkında, iki satır ise ankete katılan insanların cinsiyeti hakkında bize bilgi vermektedir. Buna bağlı olarak birinci satır ikinci sütundaki 25 sayısı bize, sağ eğilimden olan bir partiye kaç tane erkeğin oy vereceğini göstermektedir.

Tablo 11.1: Cinsiyete göre seçmenin gözlenen oy frekansları

	Sosyal Demokrat	Sağ Eğilim	Merkez	Toplam
Erkek	20	25	30	75
Bayan	30	35	10	75
Toplam	50	60	40	150

Tablo 11.1, kategorilere ait altı (2×3) sınıfın gözlenen frekanslarını içermektedir. Aynı zamanda bu tablo bütün potansiyel oy tercihleri ve cinsiyet kombinasyonlarını veya bütün kontenjanları listeleyen bir tablodur. Bu yüzden bu tarz tablolar Bölüm 4’te de değindiğimiz gibi kontenjan tabloları da denilir.

Analizimizin temel noktası oy tercihleri ile cinsiyetin bağımsız olduğu varsayımı altında beklenen frekansları belirlemektir. Beklenen frekanslardan kastımız, sıfır hipotezimizin bağımsızlığı (oy tercihleri ile cinsiyet birbirinden bağımsız ise) gerçekten doğru olduğunda karşımıza çıkacak olan frekans değerleridir. Uyum veya bağımlılık testi, gözlenen frekanslarla beklenen frekansların arasındaki farka odaklanır. Gözlenen ve beklenen değerler arasındaki büyük farklılıklar elimizdeki çıktıların doğru olup olmadığına dair şüphe uyandırır. Bu farklılıkların istatistiksel olarak anlamlı olup olmadığını araştırın teste “ki-kare” testi denir (t veya Z gibi bir olasılık dağılımıdır). Bu test genel olarak uyumluluk testi olarak da ifade edilir.

Bu test gözlenen ve beklenen değerler arasında anlamlı bir fark olup olmadığını belirlemeye yarayan bir testtir. Bu testi gerçekleştirmek için gerekli olan adımlar:

1. Sıfır hipotezinde yeralan oy tercihleri ile cinsiyetin birbirinden bağımsız olduğu doğru olarak varsayılır.

2. Örneklemedeki toplamlar ve her partiye toplam oy verilme olasılıkları (cinsiyeti dikkate almadan) hesaplanır. Mesela, Sosyal Demokratların oy alma olasılığı $50/150 = 1/3$ tür. Bağımsızlık varsayımı altında bulduğumuzu beklenen frekans olarak not ederiz. Örneğimizde olduğu gibi Sosyal Demokratlara oy veren 50 insan olduğundan, eğer cinsiyet önemli değilse erkek veya bayan insanların bu partiye oy verme olasılıkları $1/3$ tür. Sağ eğilimdeki bir parti için ise $60/150 = 6/15$, ve merkez kesime ait bir partinin olasılığı ise $40/150 = 4/15$ tir. [Olasılık toplamlarına dikkat ediniz $1/3 + 6/15 + 4/15 = 1$].

3. Eğer bağımsızlık varsayımı geçerli ise (2.) maddede hesaplanan oranlar da her iki cinsiyet grubu için de geçerli olmalıdır. Eğer oy tercihleri cinsiyetten bağımsız ise, bu oranları kullanarak, her kategori için beklenen frekansları hesaplayabiliriz. . Örneğin, oy tercihleri ile cinsiyetin bağımsız olduğu varsayımı altında, Sosyal Demokratlara oy veren erkeklerin sayısı örnekleminizde 75 erkek olduğundan, $(1/3) \times 75 = 25$ tir. Benzer şekilde Sosyal Demokratlara oy veren bayanların sayısı örnekleminizde 75 bayan olduğundan $(1/3) \times 75 = 25$ tir.

4. (3.) maddede yapılan hesaplamalara dayanılarak beklenen frekans değerleri için kontenjan tablosu oluşturulur.

5. Bağımsızlığı test etmek için ki-kare dağılımı kullanılır.

Örneğimizde bir sonraki adıma geçmeden belli notasyonları tanımlayalım.

- f_{ij} : i satırında ve j sütununda gözlenen frekans değerini ifade etsin. Örneğimizden yola çıkarak i = 1; 2 ve j = 1; 2; 3. Tablodaki f_{12} nin eşit olduğu değer 25 tir.
- e_{ij} : i satırında ve j sütununda beklenen frekans değerini ifade etsin. Örneğimizden yola çıkarak i = 1; 2 ve j = 1; 2; 3. Bu frekansların rakamları yukarıda ikinci maddede hesaplanmıştır. Tablodaki hesaplanan $e_{11} = 25$ ve $e_{21} = 25$ tir.

Şimdi birinci satır ve ikinci sütuna ait olan beklenen frekans değerini hesaplayalım. Bağımsızlık varsayımı altında bu değer daha önceden bulduğumuz $(60/150)$ oranını sayısı (bu da o satırın toplamına tekabül eder) ile çarparak bulunur:

$$e_{11} = \left(\frac{75}{150} \right) \times 50 = 25, \quad e_{21} = \left(\frac{75}{150} \right) \times 50 = 25$$

$$e_{12} = \left(\frac{75}{150} \right) \times 60 = 30, \quad e_{22} = \left(\frac{75}{150} \right) \times 60 = 30$$

$$e_{13} = \left(\frac{75}{150} \right) \times 40 = 20, \quad e_{23} = \left(\frac{75}{150} \right) \times 40 = 20$$

Genel terimlerle:

$$e_{ij} = \left(\frac{\text{Toplam Satır } i}{\text{Örneklem Boyutu}} \right) \times \text{Toplam Sütun } j$$

e_{12} için yaptığımız hesaplamayı biraz daha geliştirelim ve üstü ile alttaki oranları örneklem boyutu ile çarparsak ifade de bir değişiklik olmaz.

$$e_{12} = \left(\frac{60}{150} \right) \times \left(\frac{75}{150} \right) \times 150 = 30$$

Aslında bu ifade marjinal olasılıkların örneklem boyutu ile çarpımından başka bir şey değildir. 75/150 rassal olarak seçilen örneklemimizde erkek gelme olasılığı iken 60/150 ise rassal olarak seçilen örneklemimizde sağ eğilim partisine oy verme olasılığıdır. Eğer iki olay birbirinden bağımsız ise (ki biz öyle varsaydık) , erkek bir sağ eğilim partisi mensubunun birleşik olasılığı kendi marjinal olasılıklarının çarpımına eşittir. Tabiki biz bu hesabı bağımsızlık varsayımı altında yapmış bulunuyoruz. Eğer birleşik olasılığı örneklem boyutu ile çarparsak (örneğimize göre 150) her hücreye ait beklenen frekans değerlerini hesaplamış oluruz. Beklenen frekanslar sıfır hipotezinin bağımsızlığı varsayımı altında hesaplandığından dolayı ki-kare bu testi gerçekleştirmek için kullanmamız gereken metodur. Bu da bize sıfır hipotezinin geçerliliği hakkında fikir verecektir.

Genel terimlerle ifade edilen denklemleri kullanarak beklenen frekanslar için kontenjan tablosunu oluşturalım:

Tablo 11.2: Cinsiyete göre seçmenin beklenen oy frekanslar

	Sosyal Demokrat	Sağ Eğilim	Merkez	Toplam
Erkek	25	30	20	75
Bayan	25	30	20	75
Toplam	50	60	40	150

Artık ki-kare testini gerçekleştirmek için elimizde yeterince bilgi bulunmaktadır. Bu testin mantığına göre; eğer 11.2’de hesaplanan beklenen frekans değerleri 11.1’de hesaplanan gözlenen frekans değerlerine yeterince yakınsa bu farklılık örneklem dalgalanmasından kaynaklanmaktadır. Dolayısıyla sıfır hipotezimiz (bağımsızlık) doğrudur. Aksi takdirde ise sıfır hipotezimizi reddederiz.

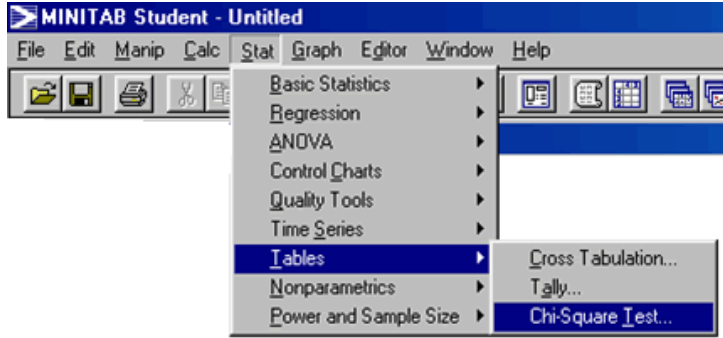
χ^2 test istatistiği beklenen ve gözlenen frekanslar arasındaki farkların karelerini alıp toplayıp herbirini kendi beklenen frekanslarına bölünerek hesaplanır.

$$\chi^2 = \frac{(20-25)^2}{25} + \frac{(25-30)^2}{30} + \frac{(30-20)^2}{20} + \frac{(30-25)^2}{25} + \frac{(35-30)^2}{30} + \frac{(10-20)^2}{20} = 13.67$$

Bu hesaplamaları yapmak zaman aldığından dolayı biz sadece p-değerine bakarak sıfır hipotezinin doğruluğunu test edebiliriz.

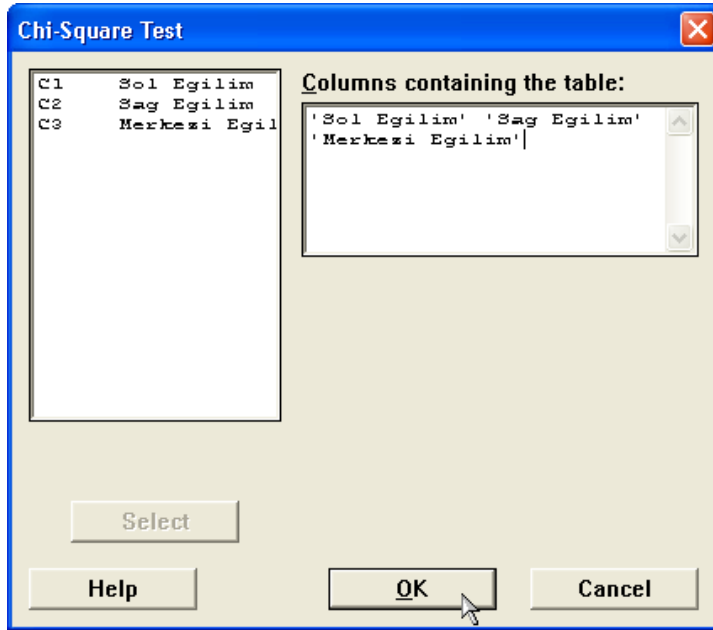
Minitab Uygulamaları

Tablo 11.1’deki verilerin yer aldığı “*ki_kare.mtw*” Minitab çalışma sayfasını açınız. Daha sonra “Stat” (istatistik) menüsünden “Tables”(tablolar)’ı seçip “Chi-square test”(ki-kare testi) seçeneğine tıklayınız.



Şekil 11.1

Karşınıza çıkan diyalog ekranında, Sol Eğilim, Sağ Eğilim ve Merkezi Eğilime ait olan değişkenleri seçiniz.



Şekil 11.2

“OK” tuşuna bastığınızda Şekil 11.3’teki çıktıyı elde edebilirsiniz.

Ki-Kare Testi				
Beklenen frekanslar Gözlenen Frekansların Altında Yeralmaktadır				
	Sol Egilim	Sag Egilim	Merkezi	Toplam
1	20	25	30	75
	25,00	30,00	20,00	
2	30	35	10	75
	25,00	30,00	20,00	
Toplam	50	60	40	150
Ki-Kare = 1,000 + 0,833 + 5,000 +				
1,000 + 0,833 + 5,000 = 13,667				
DF = 2, P-degeri = 0,001				

Şekil 11.3

P-değeri, 0.05 olan anlamlılık düzeyinden küçük olduğundan dolayı sıfır hipotezimizi reddederiz. Dolayısıyla cinsiyet ile oy tercihlerinin birbirinden bağımsız olmadığı sonucuna varırız. Bir başka deyişle cinsiyet oy tercihlerinin bir belirleyicisidir.

Bölüm 12

İki veya daha fazla Anakütle için Hipotez Testleri: Varyans Analizi (ANOVA) ve F-Dağılımı

Daha önceki bölümlerde, hipotez testinin mantığını, tek bir anakütle parametresini nasıl test edeceğimizi, iki anakütle parametresini nasıl test edeceğimizi öğrenmiştik. Peki karşımıza ikiden daha fazla anakütle parametresinin farkı ile alakalı bir test çıkarsa ne yapacağız? Bu sorunun cevabını verebilmemiz için öncelikle bu bölümü bitirmemiz gerekmektedir.

Şimdi gerekli olan hipotez testini nasıl uygulayacağımızı öğrenelim. Daha önceki bölümlerde iki anakütle ortalaması arasındaki eşitliği aşağıdaki gibi test ediyorduk:

$$H_0 : \mu_1 = \mu_2$$

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

Ancak buradaki test prosedürümüze ikiden daha fazla anakütle ortalamasının eklenmesi gerekmektedir. Buna bağlı olarak sıfır ve alternatif hipotezleri aşağıdaki gibi değişir:

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

$$H_A : \text{en az bir } \mu \text{ diğerlerinden farklıdır}$$

Sıfır hipotezimiz anakütle ortalamalarının hepsinin birbirine eşit olduğu, alternatif hipotez ise bu anakütle ortalamalarından en az birinin farklı olduğunu göstermektedir. Konuyu biraz daha açıklamak amacıyla bir örnek verelim. Birinci dönemde, Bilgi Üniversitesi Ekonomi Bölümü'ne 600 tane öğrenci, 3 farklı zorunlu derse kayıt olmuştur. Öğrencilerin bu derslerdeki başarıları bu derslerden alacakları final notuna göre değerlendirilecektir. Bütün öğrencilerin bir derse ait final notları rassal olarak değişerek dağılmaktadır. Her dersin ortalama notu anakütle ortalaması olsun ve buna bağlı olarak da 3 farklı dersin ortalama notları aynı öğrencilere ait olduklarından birbiri ile alakalıdır. Biz burada bu ortalama notlar arasında bir fark olup olmadığı ile ilgileceğiz.

Eğer bu birbiri ile alakalı anakütle parametreleri arasında bir fark varsa bu farkın nedeni ne olabilir? Anakütle parametreleri arasında fark oluşmasına sebep te kil eden bir faktör vardır. ANOVA bu faktörün anakütle ortalamaları üzerinde ne kadar etkili olup olmadığını test eder. Az önceki örneğimizde de derslerin zorluğunun (faktör) öğrencilerin başarıları veya aldıkları final notları üzerinde bir etkisi olup olmadığını irdeledik. Faktör etkisini test etmek için öncelikle örneklem verisini elde etmeliyiz. Çünkü bizim elimizde anakütleri verileri yoktur (bütün dünyadaki ekonomi öğrencilerinin notları).

Şimdi gerekli veriyi elde ettiğimizi ve aşağıdaki şekilde organize ettiğimizi düşünelim. Bu tablo şu şekilde değerlendirilmelidir: gözlemlerimiz üç dersde aynı zamanda alan bir öğrenciye tekabül etmekte ve derslere ait her hücredeki rakamlar ise öğrencinin bu derslerden aldıkları final notlarını göstermektedir. Bazı öğrencilerin bazı derslere ait final notları olmadığından o hücreler boş bırakılmıştır

Gözlem/Ders	Ekonomi	İşletme	Matematik
1	57	60	42

2	65	58	48
3	54	68	60
4	48	66	
5		58	
Toplam	224	310	150
Örneklem Ortalaması ($\bar{x}$)	$\bar{x}_1 = 224/4 = 56$	$\bar{x}_2 = 310/5 = 62$	$\bar{x}_3 = 150/3 = 50$

Örneklem ortalamaları yukarıdaki tablodan da anlaşılacağı üzere birbirinden farklıdır. Peki bu farkı yaratan neden ne olabilir? Bu sorunun iki olası cevabı olabilir:

- **Faktör Etkisi:** Farklı zorluk düzeyleri, bu grupların ortalamaları arasındaki farkın temelinde yatan etmendir. Matematik en zor derstir, bu yüzden bu dersin örneklem ortalaması diğer derslerden daha düşüktür. Aynı yaklaşım derslerin anakütle ortalamaları için de geçerlidir.
- **Örneklem Dalgalanması :** Öğrenci notları değişkendir, bu yüzden rassal olarak örneklem seçtiğimizde İşletme'nin gözlemlerinin ortalaması, Matematik'in gözlemlerinin ortalamasına oranla daha yüksek olur. Bu rassal seçim örneklem ortalamaları arasındaki farklılaşmanın kaynağıdır ama buradan anakütle ortalamalarının da birbirinden farklı olduğu anlamı çıkarılmamalıdır.

Örneklem dağılımı örneklem ortalamalarının arasında fark oluşmasının önemli bir nedeni olmasına rağmen, faktör etkisi de örneklem ortalamaları arasındaki farklılaşmanın artmasına neden olur. Faktör etkisinin olup olmadığını anlamak için öncelikle denklik şartı ile başlarız. Yani, bütün grupların aynı anakütle ortalamasına sahip olduğunu ve faktör etkisinin olmadığını varsayarsınız. Daha sonra örneklemimizi seçer ve bu varsayımımızın ne kadar doğru olup olmadığını test ederiz.

Bu hipotezleri örneğimize uygulayacak olursak:

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

$$H_A : \text{en az bir } \mu \text{ diğerlerinden farklıdır}$$

Şimdi hipotezleri nasıl test edebileceğimizi öğrenelim. Hipotezlerimizi test ederken öncelikle çok önemli bir varsayımda bulunmamız gerekmektedir. Bu varsayım, her üç grubun varyanslarında birbirine eşit olduğu, ortak bir varyansları olduğu varsayımdır ($\sigma_1^2 = \sigma_2^2 = \sigma_3^2 = \sigma^2$). Ortak varyans σ^2 'nin tahmin edicisini aşağıdaki formülü kullanarak hesaplayabiliriz.

$$\frac{\sum_{i=1}^c \sum_{j=1}^n (x_{ij} - \bar{X})^2}{T-1}$$

Formülümüzde yeralan **T**, toplam gözlem sayısını, **n** ise her sütunda (her grupta) yer alan toplam gözlem sayısını ifade ederken **c** ise grup sayısını (veya sütun sayısını) ifade etmektedir. $\bar{X}$ ise genel ortalamayı göstermektedir. Tablomuzda yer alan 12 verinin genel ortalamasını hesaplayacak olursak:

$$\bar{X} = \frac{57 + 65 + \dots + 60}{12} = \frac{684}{12} = 57$$

σ^2 'nin tahmin edicisinin payında yer alan ifade veride meydana gelen toplam farklılaşmayı göstermektedir. Bu farklılaşmayı aşağıdaki gibi iki parça halinde yazalım:

$$\underbrace{\sum_{i=1}^c \sum_{j=1}^n (x_{ij} - \bar{X})^2}_{TKT} = \underbrace{\sum_{i=1}^c n_i (\bar{x}_i - \bar{X})^2}_{GAKT} + \underbrace{\sum_{i=1}^c \sum_{j=1}^n (x_{ij} - \bar{x}_i)^2}_{GİKT}$$

TKT : Toplam Kareler Toplamı, verilerde meydana gelen toplam farklılaşmayı göstermektedir.

GAKT : Gruplar Arası Kareler Toplamı, gruplar arasında oluşan farklılaşmayı göstermektedir.

GİKT: Gruplar İçi Kareler Toplamı, grupların içerisinde meydana gelen farklılaşmayı göstermektedir.

Gruplar Arası Farklılaşma (GAKT)

Örneklem dalgalanmasından kaynaklanan farklılıklar, verinin farklılaşması ile doğrudan alakalıdır. Eğer gözlemler geniş bir dağılıma sahipse, seçilme şansına göre bazı gruplarda büyük örneklem gözlemleri yer alırken, bazı gruplarda da küçük örneklem gözlemleri yer alır. Eğer gözlemlerin dağılımı daha dar ise, örneklem değerleri birbirlerine yakındır dolayısıyla grup ortalamaları arasında önemli farklara rastlamayız.

Öncelikle gruplar arasındaki farklılaşmayı ölçelim. Bu ölçümü gerçekleştirmek için,

1. Genel ortalamayı ($\bar{X}$) hesaplayınız (hangi grupta yer aldığına bakılmaksızın bütün verilerin ortalamasıdır).
2. Her Grup ortalamasının genel ortalamadan sapmasını hesaplayınız ($\bar{x}_i - \bar{X}$).
3. Bu sapmaların karelerini alınız ($\bar{x}_i - \bar{X}$)².
4. Kareleri alınmış bu sapma değerlerini her gruba ait olan örneklem boyutu (n_j) ile çarpınız.
5. Gruplar arası kareler toplamını (GAKT) hesaplamak için bir önceki maddede hesapladığınız değerleri toplayınız $\sum n_i (\bar{x}_i - \bar{X})^2$.
6. Hesapladığımız GAKT değerini, $(c - 1)^{20}$ serbestlik derecesine bölersek, Gruplar Arası Ortalama Kare (GAOK) değerini elde ederiz:

$$GAOK = \frac{GAKT}{c - 1}$$

Örneğimizdeki GAOK değerini hesaplamak için yukarıda belirttiğimiz prosedürü izlersek:

$$GAOK = \frac{4 \times (56 - 57)^2 + 5 \times (62 - 57)^2 + 3 \times (50 - 57)^2}{3 - 1} = 138$$

Eğer sıfır hipotezi doğru ise, üç grup için sadece ortak bir ortalama olmalıdır ($\mu_1 = \mu_2 = \mu_3 = \mu$), dolayısıyla GAKT ve tabi GAOK değerlerinin çok küçük miktarlarda olması beklenir. Çünkü her

²⁰ (n - 1) serbestlik derecesi mantığından yola çıkarak, n gibi c tane grup ortalaması olduğundan serbestlik derecesi (c - 1) olmaktadır.

grup ortalamasının tahmin edicisi ($\bar{x}_1, \bar{x}_2, \bar{x}_3$) ile genel ortalamanın ($\bar{X}$) tahmin edicisi arasındaki farklılaşma sadece örneklem dalgalanmasından meydana gelecektir.

Gruplar İçi Farklılaşma (GİKT)

Gruplar arası farklılaşmaya benzer bir şekilde, gruplar içi farklılaşmayı da aşağıdaki adımları izleyerek hesaplayabiliriz:

1. Gözlemlerin kendi grup ortalamalarından sapmasını hesaplayınız ($x_{ij} - \bar{x}_i$).
2. Bu sapmaların karesini alınız ($x_{ij} - \bar{x}_i$)².
3. Karelerini aldığımız sapma değerlerini toplayınız $\sum_j (x_{ij} - \bar{x}_i)^2$. Elde etmiş olduğumuz değer bir grubun içindeki farklılaşma değeridir. Gruplar içindeki toplam farklılaşmayı elde etmek için:
4. Her grup için $\sum_j (x_{ij} - \bar{x}_i)^2$ değerlerini toplayınız. (GİKT = $\sum_{i=1}^c \sum_{j=1}^n (x_{ij} - \bar{x}_i)^2$).
5. Hesapladığımız GİKT değerini, (N -c) serbestlik derecesine²¹ bölersek, Gruplar İçi Ortalama Kare (GİOK) değerini elde ederiz:

$$GİOK = \frac{GİKT}{N - c}$$

Örneğimizdeki verileri kullanarak GİOK değerini hesaplayalım:

$$GİOK = \frac{[(57-56)^2 + \dots + (48-56)^2] + [(60-62)^2 + \dots + (58-62)^2] + [(42-50)^2 + \dots + (60-50)^2]}{12-3}$$
$$= \frac{406}{9} = 45.11$$

Daha önce de bahsettiğimiz gibi eğer sıfır hipotezi doğru ise, GAKT değerinin küçük bir farklılaşma göstereceğini beklediğimizi söylemiştik. Bu şartlar altında veride meydana gelen farklılaşmanın tamamının GİKT değerinden geldiği gözükmemektedir. Dolayısıyla sıfır hipotezinin ölçümü bu değerlerin karşılaştırılması ile gerçekleştirilebilir.

Bu iki temel terimi GAKT ile GİOK oranlayarak ve bu orana da F-istatistiği adını vererek karşılaştırabiliriz.

$$F = \frac{GAKT}{GİOK}$$

Eğer sıfır hipotezi doğru ise, F istatistik değerinin küçük çıkmasını bekleriz (küçük GAKT ve büyük GİOK değerinden dolayı). Eğer sıfır hipotezi doğru değilse, tam tersini bekleriz²². Örneğimizdeki verilere ait F istatistik değeri:

Aslında bu serbestlik derecesi her grubun kendi serbestlik derecelerinin toplamıdır: c tane grubumuz olduğundan, bu grupların serbestlik dereceleri toplamı [($n_1 - 1$) + ($n_2 - 1$) + ($n_3 - 1$) + ... + ($n_4 - 1$) = (N - c)] değerine eşit olur. N burada bütün n değerlerinin toplamını ifade etmektedir.

²² İki örneklem varyansının oranı F-dağılımı olarak dağılır. Bu yüzden bu orana F oranı denilmektedir ve bu dağılımı kullanarak iki veya daha fazla anakütle arasındaki farklılıkları test edebiliriz.

$$F = \frac{138}{45.11} = 3.05$$

F istatistik değeri büyük bir değer olduğundan, anakütle parametreleri üzerinde faktör etkisinden bahsedebiliriz. Peki faktör etkisinden bahsedebilmek için F-istatistik değerinin ne kadar büyük olması gerekmektedir? Bu sorunun cevabını bulabilmek için F dağılımı tablosundan elde edilecek olan F kritik değerini bulmamız gerekmektedir. GAOK için serbestlik derecesi (c - 1) ve GİOK için serbestlik derecesi (T - c) iken, bu serbestlik derecelerinden birincisi paya ikincisi paydaya aittir. Bu serbestlik derecelerinden yararlanarak F dağılımı tablosundan F kritik değerini bulabiliriz ve bulduğumuz bu F kritik değerini, F istatistik değeri ile karşılaştırarak faktör etkisinden bahsedebiliriz veya bahsedemeyiz. Kuralımızı tekrar hatırlayacak olursak:

Eğer $F_{\text{istatistik}} > F_{\alpha, c-1, T-c}$ ise H_0 hipotezimizi redederiz α yine burada anlamlılık düzeyini göstermektedir. Alternatif olarak p-değeri kuralını uygulayabiliriz.

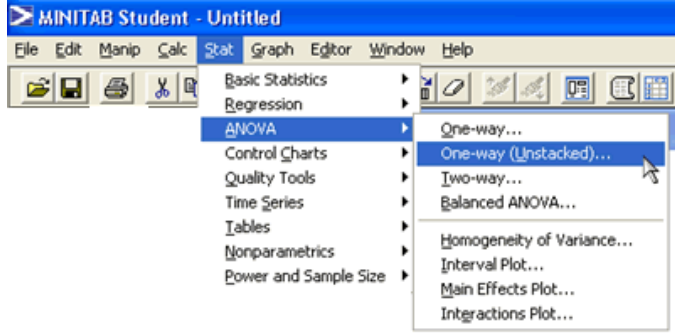
Minitab Uygulamaları

MINITAB programını kullanarak, faktör etkisini ortaya çıkarmak için yaptığımız işlemleri çok daha kısa bir sürede gerçekleştirebiliriz. Öncelikle aşağıdaki gibi örneklem verilerimizi giriniz:

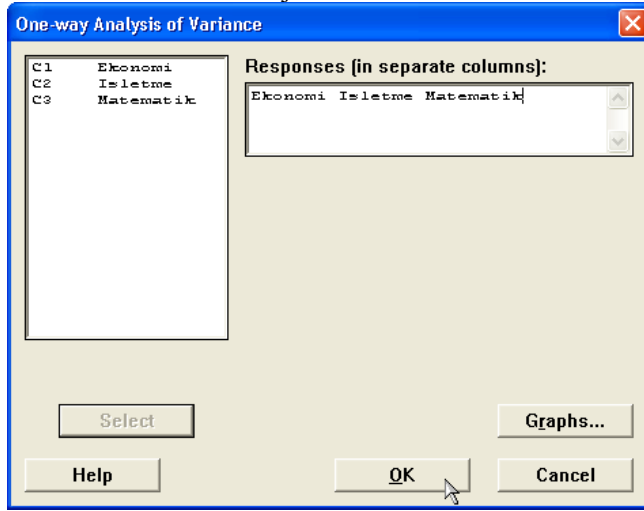
Worksheet 1 ***			
	C1	C2	C3
↓	Ekonomi	İşletme	Matematik
1	57	60	42
2	65	58	48
3	54	68	60
4	48	66	
5		58	
6			

Şekil 12.1

“Stat” (istatistik) menüsünden “ANOVA”ya tıklayıp, “One Way (Unstacked)” (tek yönlü ANOVA) seçeneğine basınız.



Şekil 12.2



Şekil 12.3

Karşınıza çıkacak olan diyalog ekranında verileri girdiğiniz sütunları seçip “OK” tuşuna basınız.

Bütün bu işlemler sonucunda MINITAB bize şekil 12.4’teki çıktıyı verecektir.

One-way Analysis of Variance (Tek Yönlü Varyans Analizi)					
Source	DF	SS	MS	F	P
Factor	2	276,0	138,0	3,06	0,097
Error	9	406,0	45,1		
Total	11	682,0			

Individual 95% CIs For Mean Based on Pooled StDev			
Level	N	Mean	StDev
Ekonomi	4	56,000	7,071
İsletme	5	62,000	4,690
Matematik	3	50,000	9,165

Pooled StDev =	6,716	48,0	56,0	64,0
----------------	-------	------	------	------

Şekil 12.4

Bu çıktıdan faydalanarak hipotez testimizin sonucu bulabiliriz. Hipotezimizi test edebilmek için daha önceki bölümlerde de kullandığımız p-değerini kullanalım.

Çıktımızdaki p-değeri 0.097'dir. Eğer anlamlılık düzeyini 0.05 olarak alırsak, sıfır hipotezimizi (H_0) reddetmeyiz. Bu da bize örneklem verisinin faktör etkisini tam olarak yansıtmak için yeterli olmadığını gösterir.

Alternatif olarak anlamlılık düzeyini yüzde 10 olarak alırsak sıfır hipotezini reddederiz ve faktör etkisi olduğu sonucuna varırız.

Varyans analizi, ikiden çok ortalamanın karşılaştırılması gereken alanlarda uygulanabilen önemli bir metottur. Örneğin tarım sektöründe, aynı tarım alanlarına uygulanan farklı tarım yöntemlerinin ortalama olarak alınan ürünler üzerindeki etkinliğini araştırabiliriz. Veya bir araba üzerinde denenen farklı motor yağlarının, arabanın hızı üzerine olan ortalama etkisini araştırabiliriz.

Bölüm 13: Basit Korelasyon

Giriş

Bu bölümde, değişkenler arasındaki ilişkinin miktarını inceleyerek, bu ilişkinin kuvveti üzerinde duracağız. Tüm bunları gerçekleştirebilmemiz için öncelikle örneklem kovaryansı ve basit doğrusal korelasyon katsayısı kavramlarını tanımlamamız gerekmektedir.

Kovaryans ve Doğrusal Korelasyon Katsayısı

Bazı durumlarda araştırmacılar iki değişkenin arasındaki ilişkinin şiddetini bilmeye ihtiyaç duyarlar. Korelasyon analizi bu ihtiyacı karşılayacak en önemli araçlardan biridir. Korelasyon analizi, iki değişkenin arasındaki nedensellikten ziyade, bu ilişkinin kuvvetini incelemeye yöneliktir. Doğal olarak bu analiz doğrusal ilişkiler baz alındığında çok faydalı olmaktadır. Fakat, doğrusal olmayan ilişkilerin incelenmesinde çok fazla başarı oranı beklenemez.

Anakütke kovaryansını (anakütke boyutu N olmak üzere) aşağıdaki gibi hesaplayabiliriz:

$$\sigma_{XY} = \frac{\sum_{i=1}^N (X_i - \mu_X)(Y_i - \mu_Y)}{N} \quad (13.1)$$

Örneklem kovaryansını ise aşağıdaki hesaplayabiliriz

$$s_{XY} = \frac{\sum_{i=1}^n (X_i - \bar{X})(Y_i - \bar{Y})}{n-1} \quad (13.2)$$

Bu formülde her X gözlemi kendisine tekabül eden Y gözlemi ile eşleşmiştir.

Kovaryansın büyüklüğü, gözlemlerin ölçüm birimlerinden etkilenmektedir. Analizimizi daha etkin bir şekilde yapabilmek için bu etkiyi ortadan kaldırmamız gerekmektedir. İşte Pearson korelasyon katsayısı bu problemimize bir çözüm getirmektedir. Pearson korelasyon katsayısı örneklem verileri için aşağıdaki gibi hesaplanmaktadır:

$$r_{XY} = \frac{s_{XY}}{s_X s_Y} \quad (13.3)$$

r_{XY} = örneklem korelasyon katsayısı

s_{XY} = örneklem kovaryansı

s_X = X değişkeninin örneklem standart sapması

s_Y = Y değişkeninin örneklem standart sapması

Değişkenlerin standart sapmalarını hatırlayalım

$$s_X = \sqrt{\frac{\sum_{i=1}^n (X_i - \bar{X})^2}{n-1}} \text{ ve } s_Y = \sqrt{\frac{\sum_{i=1}^n (Y_i - \bar{Y})^2}{n-1}}$$

Dolayısıyla Pearson korelasyon katsayısını tekrar ifade edecek olursak

$$r_{XY} = \frac{\sum_{i=1}^n (X_i - \bar{X})(Y_i - \bar{Y}) / n - 1}{\sqrt{\frac{\sum_{i=1}^n (X_i - \bar{X})^2}{n - 1}} \sqrt{\frac{\sum_{i=1}^n (Y_i - \bar{Y})^2}{n - 1}}} \quad (13.4)$$

Denklem 13.4'te ölçüm birimi standardize edilmiştir ve kovaryansta karşılaştığımız bu sorun ortadan kalkmıştır. Boyutu N olan bir anakütleyle ait Pearson korelasyon katsayısı aşağıdaki gibi hesaplanmaktadır:

$$\rho_{XY} = \frac{\sigma_{XY}}{\sigma_X \sigma_Y} \quad (13.5)$$

ρ_{XY} = anakütle korelasyon katsayısı

σ_{XY} = anakütle kovaryansı

σ_X = X değişkeninin anakütle standart sapması

σ_Y = Y değişkeninin anakütle standart sapması

Şimdi bu korelasyon analizini bir örnek ile inceleyelim. Tablo 13.1'deki veriler belli bir mala olan talep ile fiyatları göstermektedir.

Tablo 13.1

Fiyat	Talep
1	300.0
2	298.0
3	290.0
4	297.5
5	288.0
6	291.0
7	285.5
8	280.5
9	286.0
10	275.0

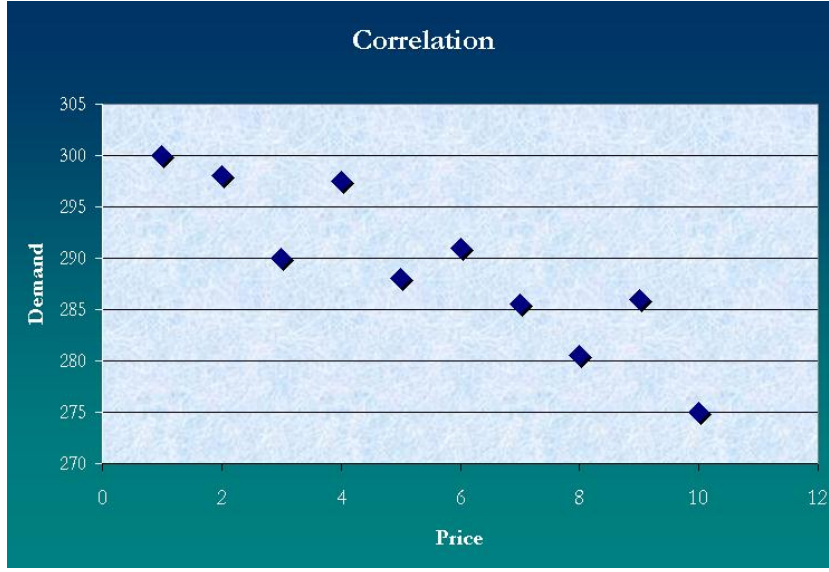
X değişkeni fiyatları, Y değişkeni ise talebi temsil etmektedir. X_i ise $i = 1$ den 10'a kadar olan fiyatlar arasında i. fiyatı temsil ederken; Y_i , $i = 1$ den 10'a kadar olan talep miktarları arasında i. talep miktarını göstermektedir. $\bar{X}$, fiyatların örneklem ortalamasını temsil ederken; $\bar{Y}$ talep edilen miktarın örneklem ortalamasını göstermektedir. Özet istatistik değerleri aşağıdaki gibidir:

$$\bar{X} = 5.5, s_x = 3.028$$

$$\bar{Y} = 289.15, s_y = 7.95$$

Ortalama fiyat düzeyi 5.5 iken, ortalama talep miktarı 289.15'tir.

İlk etapta bu iki değişken arasındaki korelasyon incelenirken en faydalı metodlardan biri scatter diyagramı çizmektedir. Şekil 13.1'de, X ekseninde fiyat yer alırken, Y ekseninde talep miktarı yer almaktadır. Şekilde yer alan on tane nokta, fiyat ve talep kombinasyonlarını göstermektedir. Açık bir şekilde bu noktaların bir paternde dağıldığı, rastgele dağılmadıkları gözükmemektedir. Fiyatlar artınca talep miktarının düştüğü gözlenmektedir.



Şekil 13.1 Fiyat ve Talep Miktarının Scatter Diyagramı

Şimdi denklem (13.2)'nin pay kısmına geri dönelim ve paya odaklanalım:

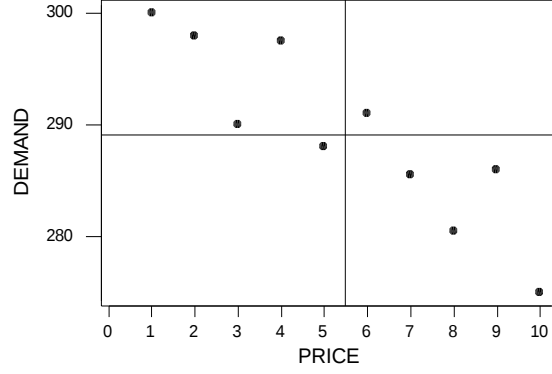
$$(X_i - \bar{X})(Y_i - \bar{Y}):$$

Eğer $X_i > \bar{X}$ ve $Y_i > \bar{Y}$ ise $(X_i - \bar{X})(Y_i - \bar{Y}) > 0$ (I.Durum)

Eğer $X_i < \bar{X}$ ve $Y_i > \bar{Y}$ ise $(X_i - \bar{X})(Y_i - \bar{Y}) < 0$ (II.Durum)

Eğer $X_i < \bar{X}$ ve $Y_i < \bar{Y}$ ise $(X_i - \bar{X})(Y_i - \bar{Y}) > 0$ (III.Durum)

Eğer $X_i > \bar{X}$ ve $Y_i < \bar{Y}$ ise $(X_i - \bar{X})(Y_i - \bar{Y}) < 0$ (IV.Durum)



Şekil 13.2 Fiyat ve Talep Miktarının Scatter Diyagramı (Bölgelere Göre)

Şekil 13.2’de yer alan dikey doğru, X_i değerlerinin ortalaması (5.5) tarafından belirlenirken, yatay doğru ise Y_i değerlerinin ortalaması (289.15) tarafından belirlenir. Şimdi biraz önce bahsettiğimiz dört farklı durumu Şekil 13.2’yi de kullanarak inceleyelim.

- Kuzeydoğu bölgesinde bulunan noktalar birinci durumu temsil ederler.
- Güneydoğu bölgesinde bulunan noktalar dördüncü durumu temsil ederler.
- Güneybatı bölgesinde bulunan noktalar üçüncü durumu temsil ederler.
- Kuzeybatı bölgesinde bulunan noktalar ikinci durumu temsil ederler.

Şekil 13.1’den de açıkça görüldüğü gibi noktaların çoğu kuzeybatı ve güneydoğu bölgelerinde toplanmıştır. Bu bölgelerde ikinci ve dördüncü durum söz konusu olduğundan, denklem (13.2)’nin pay kısmı $(X_i - \bar{X})(Y_i - \bar{Y})$ negatif bir değer alacaktır ve buda bize X ve Y değişkenleri arasında negatif bir ilişki olduğunu gösterecektir. Bulgumuzu, aşağıdaki gibi örneklem kovaryansını hesaplayarak teyid edelim:

$$s_{XY} = \frac{\sum_{i=1}^{10} (X_i - \bar{X})(Y_i - \bar{Y})}{10 - 1} = -21.63889$$

Örneklem kovaryansının sayısal değeri ölçüm birimlerinden etkilendiğinden böyle büyük çıkabilir. Ama işaretinin negatif olması bize X ve Y değişkenlerinin arasında negatif bir ilişki olduğunu göstermektedir. Peki bu ilişkinin şiddeti nedir? Bu sorunun cevabını alabilmek için öncelikle örneklem kovaryans değerinin değişkenlerin standart sapmalarına bölünerek standardize edilmesi gerekmektedir:

$$r_{XY} = \frac{s_{XY}}{s_X s_Y} = \frac{-21.63889}{3.028 \times 7.95} = -0.898$$

Hesaplamış olduğumuz örneklem korelasyon katsayısı, fiyat ve talep miktarı arasında kuvvetli bir şekilde negatif bir ilişki olduğunu göstermektedir. Korelasyon katsayısı -1 ile 1 arasında değerler alabilmektedir. Örneklem korelasyon katsayısının +1 olması değişkenler arasında tam pozitif ilişkinin olduğunu, -1 olması ise tam negatif ilişkinin olduğunu göstermektedir. Korelasyon katsayısının istatistiksel anlamlılık testine daha ilerleyen bölümlerde değinilecektir.

Basit Doğrusal Korelasyon Katsayısına Alternatif bir Yaklaşım

X değişkeninin, bir ürünün Cuma geceleri televizyonda yayımlanan reklam sayısını, ve Y değişkeninin de bu ürünün Cumartesi günleri olan satış miktarını (\$ 000) temsil ettiğini varsayalım. Bu örnekte, her iki değişken için de sadece yedi gözlem yer almaktadır.

Tablo 13.2: Örneklem Kovaryansının Hesaplanması

X_i	Y_i	$X_i - \bar{X}$	X_i^2	$Y_i - \bar{Y}$	$X_i Y_i$	$(X_i - \bar{X})(Y_i - \bar{Y})$
2	24	-1	4	-1	48	1
5	28	2	25	3	140	6
1	22	-2	1	-3	22	6
3	26	0	9	1	78	0
4	25	1	16	0	100	0
1	24	-2	1	-1	24	2
5	26	2	25	1	130	2
$\sum = 21$	$\sum = 75$	$\sum = 0$	$\sum = 81$	$\sum = 0$	$\sum = 542$	$\sum = 17$

Elimizdeki verilerden, $\bar{X} = 21/7 = 3$ ve $\bar{Y} = 75/7 = 10.71$, ayrıca toplam $(X_i - \bar{X})^2$ değerini 18, toplam $(Y_i - \bar{Y})^2$ değerini 22 olarak hesaplayabiliriz. Bulduğumuz sonuçları denklem (13.2)'de kullanacak olursak:

$$s_{XY} = \frac{\sum_{i=1}^n (X_i - \bar{X})(Y_i - \bar{Y})}{n-1} = \frac{17}{6} = 2.8333 \quad (13.6)$$

Kovaryans terimi reklamlar ile satışlar arasında kuvvetli bir pozitif ilişki olduğunu göstermektedir. X değişkeni reklam sayısı cinsinden ölçülürken, Y değişkeni ödenilen dolar cinsinden ölçüldüğü için yine kovaryans terimimizin boyu ölçü birimlerinden etkilenmektedir.

Bu sorunun aşılması için daha öncede bahsettiğimiz gibi bu terimin standardize olması gerekmektedir.

Hatırlanacağı gibi bu standardizasyonu sağlamamızda yardımcı olan en önemli araç Pearson korelasyon katsayısı idi:

$$r_{XY} = \frac{s_{XY}}{s_X s_Y}$$

r_{XY} = örneklem korelasyon katsayısı

s_{XY} = örneklem kovaryansı

s_X = X'in örneklem standart sapması

$s_Y = Y$ 'nin örneklem standart sapması

X ve Y değişkenlerinin standart sapmaları:

$$s_X = \sqrt{\frac{\sum (X_i - \bar{X})^2}{n-1}} = \sqrt{\frac{18}{6}} = 1.7321$$

$$s_Y = \sqrt{\frac{\sum (Y_i - \bar{Y})^2}{n-1}} = \sqrt{\frac{22}{6}} = 1.9149$$

Dolayısıyla Pearson örneklem korelasyon katsayısını aşağıdaki gibi hesaplayabiliriz:

$$r_{XY} = \frac{s_{XY}}{s_X s_Y} = \frac{2.8333}{1.7321 \times 1.9149} = 0.854$$

Tahmin edilmiş olan örneklem korelasyon katsayısı 0.854 bize Cuma günü olan reklam sayısı ile Cumartesi günü olan satışlar arasında pozitif bir korelasyon öngörmektedir. Analizimizi bir adım ilerleterek, örneklem korelasyon katsayısının istatistiksel olarak sıfırdan farklı olup olmadığını inceleyelim. Örneklem korelasyon katsayısı (r_{XY}), anakütle korelasyon katsayısının (ρ_{XY}) bir tahmin edicisidir. X ve Y arasındaki doğrusal ilişkinin anlamlılığına yönelik bir test aşağıdaki gibi gerçekleştirilebilir:

$$H_0 : \rho_{XY} = 0$$

$$H_1 : \rho_{XY} \neq 0$$

Burada kullanacağımız istatistik değerinin serbestlik derecesi, X 'in standart sapması hem de Y 'nin standart sapması elimizdeki verilerden hesaplandığından dolayı $(n-2)$ 'dir. Örneklem korelasyon katsayımızı da kullanarak, uygun olan test aşağıdaki gibi ifade edilmiştir.

$$r_{XY} \sqrt{\frac{n-2}{1-r_{XY}^2}} \quad (13.7)$$

Değerleri yerine koyduğumuzda ise:

$$r_{XY} \sqrt{\frac{n-2}{1-r_{XY}^2}} = 0.854 \sqrt{\frac{7-2}{1-(0.854)^2}} = 3.67$$

$\alpha = 0.05$ ve serbestlik derecesi $n-2 = 5$ iken çift taraflı bu testin kritik değeri 0.878'tir. Dolayısıyla, iki değişken arasında doğrusal bir ilişki olmadığını belirten sıfır hipotezi reddedilir. İstatistikçiler, r_{XY} 'nin örneklem dağılımının sadece örneklem boyutuna (n ye) bağlı olduğunu

göstermişlerdir. İstatistik tabloları da r_{xy} 'nin, doğrusal ilişki yoktur sıfır hipotezini reddedecek kadar büyük bir örneklem düzeyinden başlaması prensibine göre oluşturulmuştur. Örneğin, $\alpha = 0.05$ ve $n = 7$ iken, tahmin edilmiş örneklem korelasyon katsayısı sıfır hipotezini reddedebilmek için 0.878 değerinden büyük olmalıdır. Genelde karşılaşılan sorun ise $\alpha = 0.05$ ve $n, 100$ 'e doğru artarken, tahmin edilmiş örneklem korelasyon katsayısı sıfır hipotezini reddedebilmek için sadece 0.196 değerinden büyük olmalıdır.

Tablo 13.3: Pearson Korelasyon Katsayısı için Seçilmiş Kritik Değerler

N	$\alpha = 0.05$	$\alpha = 0.01$
5	0.878	0.959
6	0.811	0.917
7	0.754	0.875
8	0.707	0.834
9	0.666	0.798
10	0.632	0.765
20	0.444	0.561
30	0.361	0.463
40	0.312	0.402
50	0.279	0.361
60	0.254	0.330
100	0.196	0.256

$H_0 : \rho_{xy} = 0$ hipotezini, $H_0 : \rho_{xy} \neq 0$ hipotezine karşılık olarak reddedebilmemiz için, r_{xy} 'nin mutlak değeri tablo 13.3'deki ikinci veya üçüncü sütunda yeralan kritik değerlerden büyük olmalıdır.

Şimdi bu hipotezlerin istatistiksel anlamlılığını ölçmek için her iki yaklaşımı da ele alalım. X (fiyat) ve Y (talep) arasında istatistiksel olarak anlamlı bir doğrusal ilişki olup olmadığını test edelim:

$$H_0 : \rho_{xy} = 0$$

$$H_1 : \rho_{xy} \neq 0$$

test istatistik değerimizi bulmak için

$$r_{xy} \sqrt{\frac{n-2}{1-r_{xy}^2}}$$

elimizdeki değerleri formüle yerine koyarsak:

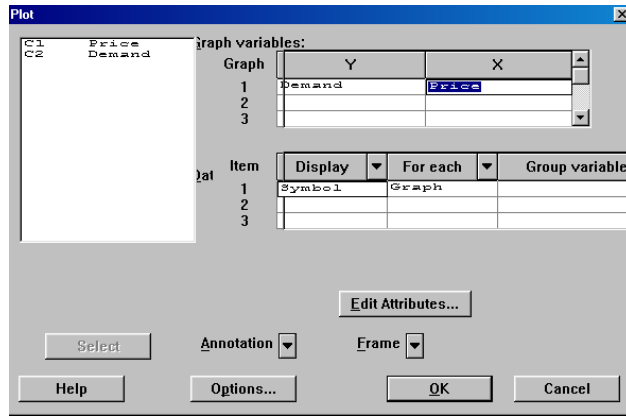
$$-0.854 \sqrt{\frac{10-2}{1-(-0.854)^2}} = -1.8368$$

sonucunu elde ederiz. Serbestlik derecesi 8 ve $\alpha = 0.05$ iken çift taraflı t-testinin kritik değerleri ± 0.811 'dir. Dolayısıyla sıfır hipotezini reddederiz.

Aynı sonuca varmak için tablo 13.3'teki Pearson korelasyon katsayısının kritik değerlerinden de faydalanabilirsiniz. $n = 10$ iken korelasyon katsayısının kritik değeri 0.707 ve tahmin değerimiz ise 0.854 olduğundan dolayı sıfır hipotezimizi bu yaklaşımla da reddederiz.

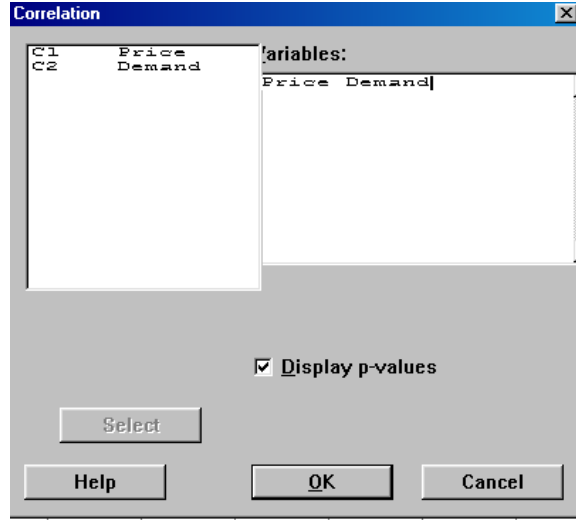
Minitab Uygulamaları

Minitab programı bize bu hesaplamalarda çok yardımcı olmaktadır. Tablo 13.1'de yer alan verilerin yer aldığı "*talep.mtw*" adlı dosyayı açınız. Şekil 13.1'deki scatter grafiğiniz çizmek için "*Graph*" (Grafik) menüsünden "*Plot*"(Çizime)'a tıklayınız ve "*X*" başlıklı olan bölüme "*Fiyat*" değişkenini seçiniz. "*Y*" başlıklı olan bölüme de "*Talep*" değişkenini seçiniz. Daha sonra "*OK*" tuşuna tıklayınız.



Şekil 13.3

Minitab ile korelasyon katsayısını hesaplamak için "*Stat*"(İstatistik) menüsünden "*Basic statistics*" (Temel İstatistikler)'i seçiniz ve "*Correlation*"(Korelasyon)'a tıklayınız. Karşınıza çıkan diyalog ekranında da "*Talep*" ve "*Fiyat*" değişkenlerini seçtikten sonra "*OK*" tuşuna basınız.



Şekil 13.4

Aşağıdaki çıktı ekranını elde etmeniz gerekmektedir:

Correlations (Pearson)

Correlation of Price and Demand = -0,898; P-Value = 0,000

Şekil 13.5

Minitab çıktısında yeralan “*P-Value*” sıfır hipotezine ait olan p-değerini (olasılık değerini) temsil etmektedir. Bu p-değeri, 0.05 anlamlılık düzeyinden küçük olduğu için $H_0 : \rho_{XY} = 0$ sıfır hipotezini reddederiz ve fiyat ile talep arasında anlamlı bir şekilde negatif ilişki olduğu sonucuna varırız.

Bölüm 14: Basit Doğrusal Regresyon

Giriş

Bir önceki bölümde, basit doğrusal korelasyon katsayılarını kullanarak iki değişken arasındaki doğrusal ilişkinin nasıl ölçüldüğünü inceledik. Bu bölümün amacı ise, iki değişken arasındaki doğrusal ilişkiyi bir adım ileri götürerek, bu ilişkiyi daha detaylı açıklamak ve doğrusal ilişkinin nedenselliğini incelemektir. Korelasyon analizinde ilişkinin nedenselliği incelenmemektedir. Bu doğrultuda, basit doğrusal regresyon analizinin altında yatan prensipler, basit doğrusal ilişkiyi tahmin etmek için kullanılacak olan regresyon tahmin edicilerinin hesaplanması incelenecektir.

Korelasyon analizi, iki değişkenin arasındaki derece hakkında bir çerçeve çizmektedir. Regresyon analizi ise, iki veya daha fazla değişken arasındaki fonksiyonel ilişkiyi belirlemeye, açıklamaya çalışmaktadır. En çok kullanılan ve en temel regresyon modeli basit doğrusal regresyon modelidir.

Bu modelde, bir değişken diğer bir değişken tarafından tahmin edilmeye çalışılmaktadır. Dolayısıyla, korelasyon analizinde yaptığımız ve nedenselliğin yönünü belirten analiz bu durumda yetersiz kalmaktadır.

Tahmin edilen değişken “*bağımlı değişken*” (veya regresand) olarak tanımlanır ve genellikle Y harfi ile ifade edilir. Tahmin edici değişken ise “*bağımsız değişken*”, “*açıklayıcı değişken*” (veya regresor), olarak tanımlanır ve X harfi ile ifade edilir.

Örnek 14.1: Tüketim Harcamaları

Regresyon analizi, bilinen ya da sabit olan açıklayıcı değişkenleri (X leri) temel alarak, anakütle ortalamasını ve/veya bağımlı değişkenlerin (Y lerin) ortalama değerlerini tahmin etme mücadelesidir. Bu mücadele doğrultusunda, istatistiğin her alanında olduğu gibi, genellikle erişemediğimiz ve açıklamak istediğimiz anakütlelerden seçilen X ve Y örneklem değerlerini kullanırız. Bu prosedürün daha iyi anlaşılabilmesi için aşağıdaki örneği ele alalım.

Toplam nüfusun 60 aileden oluştuğu bir ülke düşünelim. Bu ülkedeki ailelerin tüketim harcamaları (Y) ile harcanabilir gelirleri (vergiden arındırılmış) (X) arasındaki ilişkiyi inceleyelim. Bu ilişkiyi inceleyerek, elimizdeki haftalık gelir verilerini kullanarak, haftalık ortalama (anakütle) tüketim düzeyini tahmin etmeye çalışacağız.

Bu çalışmayı gerçekleştirirken, 60 aileyi yaklaşık olarak aynı gelir düzeyine sahip 10 ayrı gruba böldüğümüzü düşünelim ve her gelir grubunda yer alan ailelerin tüketim harcamalarını inceleyelim. Elde ettiğimiz veriler, Tablo 14.1’deki gibidir.

Tablo 14.1: Ailelerin Haftalık Gelir Düzeyi X ,

$\downarrow Y / X \rightarrow$	80	100	120	140	160	180	200	220	240	260
Ailelerin	55	65	79	80	102	110	120	135	137	150
Haftalık	60	70	84	93	107	115	136	137	145	152
Tüketim	65	74	90	95	110	120	140	140	155	175
Miktarları	70	80	94	103	116	130	144	152	165	178
$Y, \$$	75	85	98	108	118	135	145	157	175	180
	-	88	-	113	125	140	-	160	189	185
	-	-	-	115	-	-	-	162	-	191

Toplam	325	462	445	707	678	750	685	1043	966	1211
--------	-----	-----	-----	-----	-----	-----	-----	------	-----	------

Yukarıdaki tabloyu şöyle yorumlamalıyız:

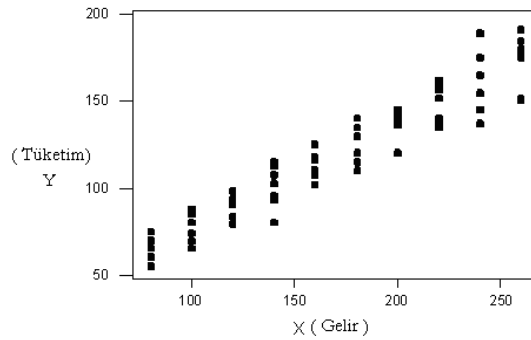
Örnek olarak, 80\$'lık haftalık gelire sahip beş ailenin, tüketim harcamaları 55\$ ile 75\$ aralığında değişmektedir. Benzer bir şekilde, $X = 220\$$ haftalık gelir düzeyinde, haftalık tüketim miktarı 137\$ ile 162\$ arasında değişen altı aile yer almaktadır. Sonuç olarak, Tablo 14.1'de yer alan her sütun sabit bir gelir düzeyine (X değerine) tekabül eden, tüketim harcamalarının (Y lerin) dağılımını göstermektedir. Bu gösterim de bize, verili X değerlerine tekabül eden Y'lerin **koşullu dağılımını** göstermektedir.

Yukarıdaki verilerimizi koşullu olasılık dağılımı prensiplerine göre yeniden düzenleğimizde aşağıdaki tabloyu elde ederiz.

Tablo 14.2

X	Y
80	55
80	60
80	65
80	70
80	75
100	65
100	70
100	74
100	80
100	85
100	88
...	...
260	191

Tablo 14.2'ye ait olan scatter diyagramı aşağıdaki gibi çizilmiştir:



Şekil 14.1

Anakütle scatter grafiği

Scatter grafikte açık bir şekilde pozitif doğrusal ilişki gözükmemektedir. Bu doğrultuda bir önceki bölümde açıkladığımız formülü kullanarak, anakütle korelasyon katsayısını 0.951 olarak hesaplayabiliriz.

Anakütle verilerinin yer aldığı Tablo 14.1'i kullanarak, Y'ye ait olan koşullu olasılık değerlerini hesaplayabiliriz. $X = 80$ değerine tekabül eden beş tane Y değeri vardır: 55 60 65 70 ve 75. Bu yüzden, verilmiş $X = 80$ iken, o sütuna ait olan herhangi bir tüketim değerini elde etme olasılığı $1/5$ 'tir. Sembolik olarak, $P(Y = 55 / X = 80) = 1/5$. Aynı şekilde diğer olasılık değerlerini de hesaplayabiliriz. $P(Y = 150 / X = 260) = 1/7$, gibi. Tablo 14.1'deki verilerin koşullu olasılıkları Tablo 14.3'te yer almaktadır.

Tablo 14.3: Tablo 14.1'deki değerlerin koşullu olasılık değerleri $P(Y / X_i)$

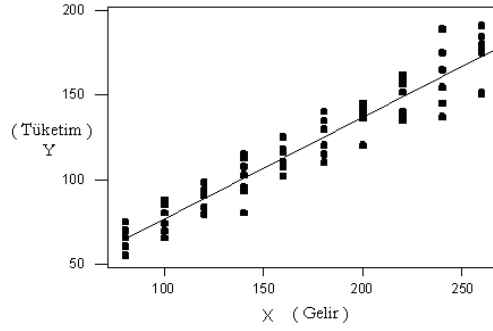
$\downarrow P(Y / X_i) / X \rightarrow$	80	100	120	140	160	180	200	220	240	260
Koşullu Olasılıklar	$\frac{1}{5}$	$\frac{1}{6}$	$\frac{1}{5}$	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{5}$	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{7}$
$P(Y / X_i)$	$\frac{1}{5}$	$\frac{1}{6}$	$\frac{1}{5}$	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{5}$	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{7}$
	$\frac{1}{5}$	$\frac{1}{6}$	$\frac{1}{5}$	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{5}$	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{7}$
	$\frac{1}{5}$	$\frac{1}{6}$	$\frac{1}{5}$	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{5}$	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{7}$
	$\frac{1}{5}$	$\frac{1}{6}$	$\frac{1}{5}$	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{5}$	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{7}$
	$\frac{1}{5}$	$\frac{1}{6}$	$\frac{1}{5}$	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{5}$	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{7}$
	$\frac{1}{5}$	$\frac{1}{6}$	$\frac{1}{5}$	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{5}$	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{7}$
	$\frac{1}{5}$	$\frac{1}{6}$	$\frac{1}{5}$	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{5}$	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{7}$
	-	$\frac{1}{6}$	-	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{6}$	-	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{7}$
	-	-	-	$\frac{1}{7}$	-	-	-	$\frac{1}{7}$	-	$\frac{1}{7}$
Y için koşullu ortalama değerleri	65	77	89	101	113	125	137	149	161	173

Şimdi, her Y değerinin koşullu dağılımları için ortalama değeri hesaplayabiliriz. Bu ortalama değer, koşullu ortalama veya koşullu beklenti olarak tanımlanır ve $E(Y / X = X_i)$ olarak ifade edilir. Bu ifade “verili bir özel X_i değeri için Y'nin beklenen değeri” olarak okunur ve $E(Y / X_i)$ şeklinde de daha kısa olarak yazılabilir. (Not: Beklenen değer anakütle ortalaması olduğu unutulmamalıdır.) Elimizdeki verileri kullanarak, Tablo 14.1'deki Y değerleri ile bu değerlere ait Tablo 14.2'deki koşullu olasılık değerlerini çarpıp, topladıktan sonra koşullu ortalama veya koşullu beklentilerin değerlerini hesaplayabiliriz.

$$E(Y | X = 80) = 55\left(\frac{1}{5}\right) + 60\left(\frac{1}{5}\right) + 65\left(\frac{1}{5}\right) + 70\left(\frac{1}{5}\right) + 75\left(\frac{1}{5}\right) = 65$$

Örnek olarak, verili $X = 80$ değeri için Y 'nin koşullu ortalaması veya koşullu beklentisi yukarıdaki gibi hesaplanabilir.

Dolayısıyla koşullu ortalamalar, Tablo 14.2'nin son satırında hesaplanmıştır. Tablo 14.1'deki veriler ile koşullu ortalamalarını bir arada gösteren scatter grafiğini çizildiğinde bu grafiğin şekil 14.1'den çok farklı olmadığı görülmektedir. Ailelerin tüketim miktarlarında farklılıklar olmasına rağmen, gelir arttıkça ortalama harcama miktarının arttığı açıkça şekil 14.2'de gözükmemektedir. Başka bir şekilde ifade edecek olursak, X değerleri arttıkça, Y 'nin (koşullu) ortalama değerleri de artmaktadır. Scatter grafiği, pozitif eğimli bir doğrunun üzerinde yeralan koşullu ortalama değerlerini göstermektedir. İşte bu doğru **anakütle regresyon doğrusu**, veya biraz daha genel bir ifade ile, **anakütle regresyon eğrisidir**. Daha basit bir ifade ile **Y nin X üzerine regresyonudur**.



Şekil 14.2
Anakütle Regresyon Doğrusu

Geometrik olarak, anakütle regresyon doğrusu, bağımlı değişkenin koşullu ortalamaları ile açıklayıcı değişkenin sabit değerlerinin birlikteliğinden oluşur.

Anakütle Regresyon Fonksiyonu Kavramı

Anakütle regresyon doğrusu, şekil 14.2'de doğrusal bir fonksiyon olarak çizildiğinden ve aşağıdaki gibi doğrusal bir fonksiyon olarak ifade edilebildiğinden, bir fonksiyondur.

$$E(Y_i / X_i) = \alpha + \beta X_i \quad (14.1)$$

Denklem (14.1) bize **anakütle regresyon fonksiyonu**'nu (ARF) göstermektedir. Anakütle regresyon doğrusu üzerinde 10 tane verimiz yereldiğinden bu veriler aşağıdaki gibidir.

X_i	$E(Y / X_i)$
-------	--------------

80	65
100	77
120	89
140	101
160	113
180	125
200	137
220	149
240	161
260	173

α ve β değerlerini hesaplayabiliriz. Bu değerleri bulabilmek için regresyon doğrusu üzerindeki sadece iki noktayı bilmemiz yeterlidir. Örneğin, α ve β aşağıdaki denklem sistemi çözülerek bulunabilir.

$$\begin{aligned} 65 &= \alpha + \beta 80 \\ 77 &= \alpha + \beta 100 \end{aligned}$$

Dolayısıyla $\alpha = 17$ ve $\beta = 0.6$ olarak α ve β değerlerini hesaplayabiliriz. Buna bağlı olarak da Denklem (14.1)'i aşağıdaki gibi yazabiliriz.

$$E(Y_i / X_i) = 17 + 0.6X_i$$

ARF'nin Stokastik Olarak Düzenlenmesi

Şekil 14.2'den de açıkça anlaşıldığı gibi, ailelerin gelir düzeyi arttıkça, ortalama tüketim düzeylerinin de arttığı gözlenmektedir. Ancak, bir ailenin tüketim harcaması ile gelir düzeyinin arasındaki ilişki nedir? Tablo 14.1'de de görüldüğü gibi bir ailenin tüketim harcamalarının artması için ile de gelir düzeyinin artmasına gerek yoktur. Örneğin, Tablo 14.1'de 100\$'lık gelir düzeyine sahip ve tüketim harcamaları 65\$ olan bir ailenin harcama miktarı, gelir düzeyi 80\$ olan iki ailenin tüketim harcamalarına kıyasla daha azdır. Fakat, gelir düzeyi 100\$ olan ailelerin ortalama harcama miktarlarının, gelir düzeyi 80\$ olan ailelere nazaran daha fazla olduğu da gözden kaçırılmamalıdır.

Peki, tek bir ailenin tüketim düzeyi ile verili gelir düzeyi arasındaki ilişki hakkında nasıl bir yorum getirmeliyiz? Şekil 14.2'yi incelediğimizde, verili bir X_i gelir düzeyi için, bir ailenin tüketim harcamaları, bütün ailelerin ortalama harcama miktarı olan X 'in etrafında veya koşullu beklentisinin etrafında kümelenmiştir. Dolayısıyla, hata terimini u_i olarak tanımlayabiliriz.

$$u_i \rightarrow N(0, \sigma_u^2)$$

Bu terim, bir Y_i değerinden beklenen değerinin sapma miktarını ifade etmektedir.

$$u_i = Y_i - E(Y_i / X_i)$$

bir başka ifade ile

$$Y_i = E(Y_i / X_i) + u_i \quad (14.2)$$

olarak gösterilebilir. u_i , gözlenemeyen (örnekleme) ve rassal olarak pozitif ve negatif değerler alan bir değişkendir. Dolayısıyla, u_i 'ya **stokastik (rassal) bozucu terim** veya **stokastik hata terimi** de denilmektedir.

Peki Denklem (14.2)'yi nasıl yorumlayabiliriz? Bir ailenin tüketim harcamaları, verili bir gelirleri varken, aşağıdaki iki kısmın toplamı olarak ifade edilebilir:

1. $E(Y / X_i)$, aynı gelir düzeyine sahip olan ailelerin ortalama tüketim harcamalarını ifade etmektedir. Normal şartlar altında insanlar bir gelire sahip olacağından bu parçayı sistematik kısım olarak ifade edebiliriz.

2. u_i 'yu ise rassal ya da sistematik olmayan kısım olarak tanımlayabiliriz. Bu kısım gelir değişkeninden başka, tüketim harcamalarını etkileyen değişkenleri temsil etmektedir. (modelde açıklayıcı olarak bulunmayan bütün faktörleri temsil etmektedir.) Sistematik olmayan faktörlerin yer almadığı bir durumda, aile ortalama olarak $E(Y / X_i)$ kadar tüketim yapmayı bekler.

Örneğin, bir ailenin gelir düzeyi 80\$ ise bu ailenin beklenen tüketim harcamaları miktarı, $E(Y / X_i) = 17 + 0.6(80) = 65$ değerine eşittir.

$$u_1 = Y_1 - E(Y / X = 80) = 55 - 65 = -10$$

$$u_2 = Y_2 - E(Y / X = 80) = 60 - 65 = -5$$

$$u_3 = Y_3 - E(Y / X = 80) = 65 - 65 = 0$$

$$u_4 = Y_4 - E(Y / X = 80) = 70 - 65 = 5$$

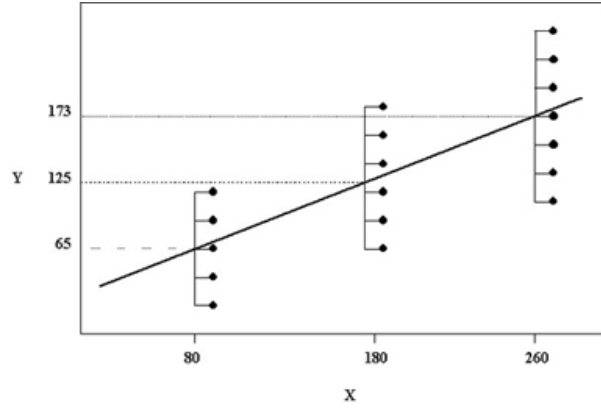
$$u_5 = Y_5 - E(Y / X = 80) = 75 - 65 = 10$$

Bu rassal hataların en önemli özelliği, ortalamalarının sıfıra eşit olmasıdır. $[-10 + (-5) + 0 + 5 + 10] / 5 = 0$. Bu özellikten dolayı regresyon modellerinde aşağıdaki varsayımı yaparız.

$$E(u_i) = 0$$

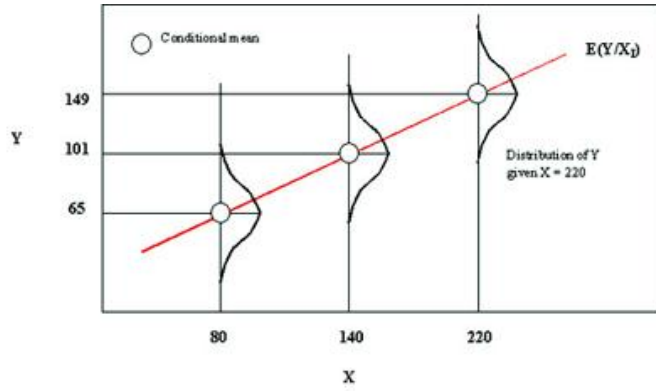
Örneğimizdeki verileri kullanarak bir önceki sayfada gerçekleştirdiğimiz alıştırmayı tekrarlayabilirsiniz ve aynı özelliğin hala geçerli olup olmadığını deneyebilirsiniz.

Ek olarak, hata terimlerinin (u_i) dağılımları hakkında da bir varsayımda bulunmamız gerekmektedir. Her gelir düzeyinden örneklem olarak bir aile seçelim. Anakütledeki 60 aileden 10 tanesini örneklem olarak alabiliriz. (10 farklı gelir seviyesi bulunduğundan). 80\$'lık gelir düzeyinden ve tüketim miktarı 70\$ olan; 100\$'lık gelir düzeyinden ve tüketim miktarı \$65 olan vb. aileleri örneklem olarak seçebiliriz. Her gelir düzeyine ait olan ailelerin tüketim harcamalarının rassal olarak seçilme olasılıkları nelerdir? 80\$'lık bir gelir düzeyinden 70\$'lık tüketim harcamasına sahip olan bir ailenin seçilme olasılığı, bu gelir düzeyinde 5 tane aile bulunduğundan, $1/5$ 'tir. Benzer şekilde, 100\$'lık bir gelir düzeyinden 70\$'lık tüketim harcamasına sahip olan bir ailenin seçilme olasılığı, bu gelir düzeyinde 6 tane aile bulunduğundan, $1/6$ 'dır. Bu prosedür aynı şekilde bütün örneklemde devam etmektedir. Dolayısıyla u_i 'ların olasılık dağılımı her gelir düzeyi için yeknesak (tekdüze) bir dağılıma sahiptir. Bu dağılım aşağıdaki şekilde de incelenebilir.



Şekil 14.3
 u_i 'ların dağılımı

u_i 'lar burada yeknesak bir dağılıma sahip olmasına rağmen, regresyon modelimizde, hata terimlerinin normal olarak dağıldığını varsaymıştık. u_i terimlerinin normal olarak dağıldığı durum Şekil 14.4'teki gibidir.



Şekil 14.4
Normal Dağılan u_i Terimleri

u_i terimlerinin normal olarak dağıldığını varsaymamızın birçok nedeni vardır. Bu nedenlerin arasında başı Merkezi Limit Teoremi çekmektedir. Bu teoreme göre, eğer yüksek sayıda bağımsız ve birbirine denk olarak dağılan rassal değişken varsa, birkaç istisna dışında, bu değişkenlerin sayısı sonsuza yakınsadıkça toplamalarının normal olarak dağılma eğilimi gösterdiği gözlenir. Daha öncede belirttiğimiz gibi u_i 'lar, bağımlı değişkeni etkileyen, ihmal edilen veya reddelilen rassal değişkenlerin toplam etkisini temsil etmektedir.

Gerçek hayatta anakütle değerlerini göremediğimizden, u_i değerlerini de gözlemleyememekteyiz. Dolayısıyla bu da varsayımlarımızın geçerliliğini test etmemizi imkansız kılmaktadır.

Normal dağılım varsayımımızı bir önceki varsayımı $E(u_i) = 0$ da kullanarak matematiksel olarak yeniden ifade edersek, aşağıdaki şekilde olur.

$$u_i \rightarrow N(0, \sigma_u^2)$$

Bu ifade bize u_i terimlerinin 0 ortalama ve σ_u^2 varyansı ile normal olarak dağıldığını göstermektedir.

α ve β Tahmin Edicilerini Bulmak için Normal Denklemlerin Kullanımı

Örneğimize tekrar dönelim ve biraz daha gerçekçi bir hale getirelim. Normalde, anakütle değerlerini elde etmemiz çok zor olduğundan, anakütleden örneklem çekerek işlemlerimizi gerçekleştirmekteyiz. Dolayısıyla aşağıdaki tablo da, tablo 14.1’de yeralan anakütle değerlerinden çekilmiş olan bir örneklem verilerini göstermektedir.

Tablo 14.4: Tablo 14.1’deki anakütleden çekilmiş rassal bir örneklem

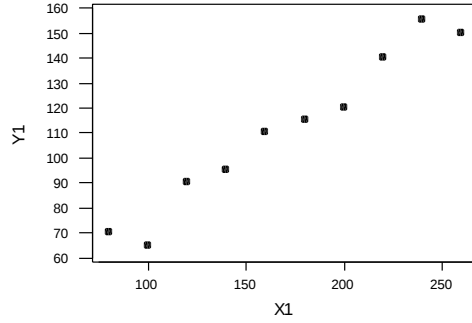
Y	X
70	80
65	100
90	120
95	140
110	160
115	180
120	200
140	220
155	240
150	260

Şimdi ise çıkarımsal istatistiğin çok karşılaştığı sorulardan biriyle karşı karşıyayız. Acaba, elimizdeki örnekleme kullanarak ortalama haftalık tüketim miktarını $E(Y / X_i)$ tahmin edebilir miyiz?. Başka bir şekilde ifade edecek olursak, elimizdeki örneklemin verileri ile ARF’yi tahmin edebilir miyiz? Veya α ve β değerlerini tahmin edebilir miyiz? Örneklem dalgalanmalarından dolayı bu tahmini kusursuz bir şekilde gerçekleştiremeyebiliriz. Fakat, en iyi, en doğruya yakın tahminini bulabiliriz.

Her zaman için Y ve X arasında doğrusal bir ilişki olduğunu varsaydığımızdan dolayı, Örneklem Regresyon Fonksiyon’u (ÖRF), ARF’ye benzemektedir ve bu fonksiyon aşağıdaki gibidir.

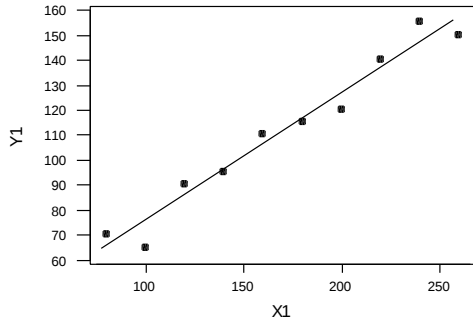
$$Y_i = a + b X_i + e_i \quad (14.3)$$

Bu denklemde; a örnekleme tahmin edilmiş α katsayısını ifade ederken, aynı şekilde b’de tahmin edilmiş olan β katsayısını ifade etmektedir. e_i ise, u_i hata terimlerinin örnekleme yer alan kısmını oluşturmaktadır. Şimdi ise a ve b değerlerini hesaplamak için gerekli olan formülleri türetilim. Öncelikle, örneklem verilerimizi gösteren scatter diagramını çizelim. Bu scatter diagramı, şekil 14.5’te olduğu gibidir.



Şekil 14.5
Birinci Rassal örneklemin grafiği

ÖRF'nin doğrusal bir yapısı olduğundan dolayı, bu verisetine en uygun olan doğruyu çizmek istemekteyiz. Peki en uygun doğru ne demektir? En uygun doğru, bütün veri noktalarına en yakından geçen doğrudur ve scatter diagramda aşağıdaki gösterilmiştir.



Şekil 14.6
En uygun doğru

Eğrinin üzerindeki noktalar (örneklem regresyon doğrusu) aşağıdaki denklem ile ifade edilebilir.

$$\hat{Y}_i = a + b X_i$$

Burada $\hat{Y}_i$, verili X_i değeri için Y_i 'nin tahmin edilmiş değerini ifade etmektedir. Bir başka ifade ile $\hat{Y}_i$, $E(Y_i / X_i)$ 'nin örnekleimde yer alan kısmıdır. $\hat{Y}_i$ ile Y_i 'nin arasındaki fark e_i 'yi verir, bu terimde örneklemin hata terimidir. Peki, şekil 14.6'da olduğu gibi, doğrumuzun bütün verilere en yakın noktalardan geçeceğini nasıl garanti edebiliriz? Bu garantiyi aslında rahatlıkla verebiliriz. Öncelikle, e_i 'lerin toplamını bulalım, ve bu toplamı ($\sum e_i$) minimize etmeye çalışalım. Ancak, hatırlayacağınız gibi bu toplam bize her zaman için sıfır değerini vermektedir. Dolayısıyla öncelikle bu sorunu halletmemiz gerekmektedir. Hata terimleri, Y_i gerçek değer ile $\hat{Y}_i$ tahmin

edilmiş değer arasındaki fark olduğundan bu farkın (yani hata terimlerinin) karelerini alırsak, terimler pozitif de olsa negative de olsa, birbirlerini sadeleştirmelerini ve toplamın sıfır çıkmasını engelleyebiliriz. Hata terimlerinin kareleri alınmış bu toplamına, “hataların kareler toplamı” denilir ve kısaca HKT olarak gösterilir. HKT’yi aşağıdaki gibi tanımlayabiliriz.

$$HKT = \sum_{i=1}^n (Y_i - \hat{Y}_i)^2 = \sum_{i=1}^n e_i^2 \quad (14.4)$$

$\hat{Y}_i$ değerini denklemde yerine koyarsak, $\hat{Y}_i = a + b X_i$ olduğundan

$$HKT = \sum_{i=1}^n (Y_i - a - bX_i)^2 \quad (14.5)$$

HKT, basit ve çoklu regresyon analizinde, regresyon doğrusunun standart hatasının hesaplanmasında önemli bir role sahiptir.

Artık, denklem (14.5)’i bir başka ifade ile hataların kareler toplamını, minimize eden a ve b değerlerini bulma zamanı geldi. Bu işlemi HKT’nin a ve b değerlerine göre kısmi türevlerini alarak ve sıfıra eşitleyerek gerçekleştirebiliriz. b katsayısının formülünü aşağıdaki gibi gösterebiliriz (Ek 14.1’e bakınız).

$$b = \frac{\sum XY - \frac{\sum Y \sum X}{n}}{\sum X^2 - \frac{\sum X \sum X}{n}} \quad (14.6)$$

b’nin formülü aynı zamanda daha toplu bir ifade ile aşağıdaki gibi de gösterilebilir:

$$b = \frac{\sum (X - \bar{X})(Y - \bar{Y})}{\sum (X - \bar{X})^2} \quad (14.7)$$

Denklem (14.7), b için kısaltılmış bir formül olmasına rağmen denklem (14.6)’i kullanarak hesaplama yapmak daha kolaydır.

a katsayısının formülü ise aşağıdaki gibi gösterilebilir (Ek 14.1’e bakınız).

$$a = \bar{Y} - b\bar{X} \quad (14.8)$$

Şimdi, tablo 14.3’teki örneklemeimize geri dönelim ve bu değerler için a ve b katsayılarını hesaplayalım. Bu hesaplamalar tablo 14.5’te gerçekleştirilmiştir.

Tablo 14.5

Y	X	$(X - \bar{X})$	$(X - \bar{X})^2$	$(Y - \bar{Y})$	$(X - \bar{X})(Y - \bar{Y})$
70	80	-90	8100	-41	36900
65	100	-70	4900	-46	3220
90	120	-50	2500	-21	1050
95	140	-30	900	-16	480
110	160	-10	100	-1	10

115	180	10	100	4	40
120	200	30	900	9	270
140	220	50	2500	29	1450
155	240	70	4900	44	3080
150	260	90	8100	39	3510
$\sum Y = 1110$	$\sum X = 1700$	$\sum (X - \bar{X}) = 0$	$\sum (X - \bar{X})^2 = 33000$	$\sum (Y - \bar{Y}) = 0$	$\sum (X - \bar{X})(Y - \bar{Y}) = 16800$

Bu deęerleri denklem (14.7)'de yerine koyarsak:

$$b = \frac{\sum (X - \bar{X})(Y - \bar{Y})}{\sum (X - \bar{X})^2} = \frac{16800}{33000} = 0.51$$

řimdi de, bulduęumuz b deęerini denklem (14.8)'da yerine koyalım:

$$\begin{aligned}\bar{Y} &= a + b\bar{X} \\ 111 &= a + (0.51) \times (170) \\ a &= 24.45\end{aligned}$$

Dolayısıyla ARF'yi ařaęıdaki gibi yazabiliriz:

$$Y_i = 24.45 + 0.51X + e_i$$

Grafiksel olarak ifade edecek olursak,



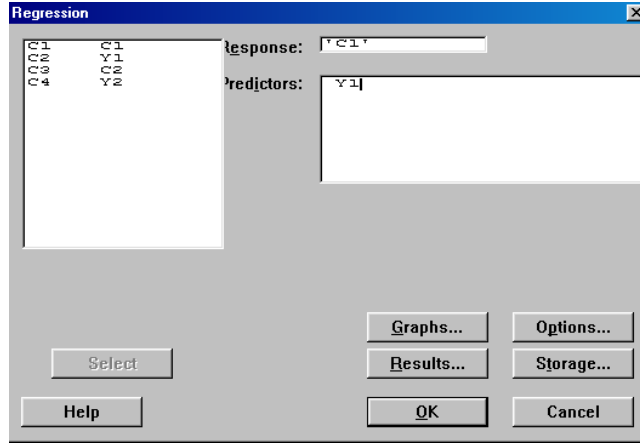
Şekil 14.7
Regresyon Doğrusu

α ve β 'nın tahmin edicilerini (a ve b) bulduğumuza göre, artık kolaylıkla u_i 'lerin tahmin edicilerini, e_i leri, hesaplayabiliriz. Bu hesaplamalar bir sonraki tabloda olduğu gibidir.

$Y_i - a - bX_i =$	e_i
$70 - 24.45 - 0.51(80) =$	4.8182
$65 - 24.45 - 0.51(100) =$	-10.3636
$90 - 24.45 - 0.51(120) =$	4.4545
$95 - 24.45 - 0.51(140) =$	-0.7273
$110 - 24.45 - 0.51(160) =$	4.0909
$115 - 24.45 - 0.51(180) =$	-1.0909
$120 - 24.45 - 0.51(200) =$	-6.2727
$140 - 24.45 - 0.51(220) =$	3.5455
$155 - 24.45 - 0.51(240) =$	8.3636
$150 - 24.45 - 0.51(260) =$	-6.8182

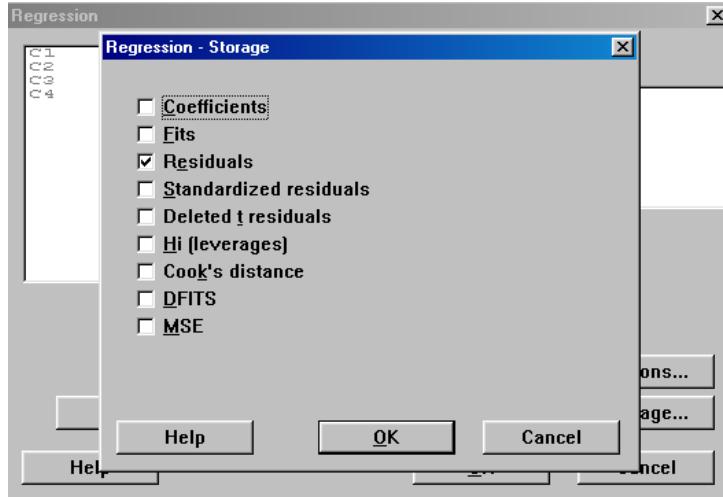
Minitab Uygulamaları

Regresyon katsayılarının hesaplaması MINITAB programında çok daha kısa ve kolay bir şekilde gerçekleştirilir. Aynı örneklem verilerini kullanmak için lütfen “*tuketim1.mtn*” adlı dosyayı açınız. Öncelikle, “Stat” menüsünden “*regression*”ı seçenerek “*regression*” seçeneğinin üzerine tıklayınız. Karşınıza çıkacak olan diyalog ekranında “*response*” yazan yere (bağımlı değişkeni temsil eder) “*C1*”(bizim örnekleminizdeki *Y1* değişkeni, *tüketim harcamaları*), “*predictor*” yazan yere de (bağımsız değişkeni ifade eder) “*Y1*”(bizim örnekleminizdeki *X1* değişkeni, *gelir düzeyi*) yazınız.



Şekil 14.8

Diyalog ekranındaki menüden “storage” seçerek aşağıda görüldüğü gibi “residuals”ı işaretleyiniz. (residuals bizim örnekleminizdeki hata terimlerinizi ifade etmektedir)



Şekil 14.9

İki kere “OK” tuşuna bastığınızda aşağıdaki çıktıyı elde etmeniz gerekmektedir.

Regression Analysis

The regression equation is
 $C1 = 24,5 + 0,509 Y1$

ÖRF
 Regresyon
 Katsayıları

Predictor	Coef	StDev	T	P
Constant	24,455	6,414	3,81	0,005
Y1	0,50909	0,03574	14,24	0,000

S = 6,493 R-Sq = 96,2% R-Sq(adj) = 95,7%

Analysis of Variance

Source	DF	SS	MS	F	P
Regression	1	8552,7	8552,7	202,87	0,000
Residual Error	8	337,3	42,2		
Total	9	8890,0			

Şekil 14.10

Karşılaştığımız regresyon çıktısında, regresyon doğrusu hata terimlerinin haricinde tamamen hesaplanmıştır. Hata terimlerinin bu program yardımı ile nasıl hesaplandığını daha detaylı bir şekilde bu bölümde irdelenecektir. Regresyon analizini gerçekleştirirken “storage” bölümünden “residuals” seçeneğini tıkladığımızdan dolayı Minitab hata terimlerini bulacak ve aşağıda olduğu gibi çalışma sayfasında “RESID1” adı altında kaydedecektir.

C4	C5
Y2	RESID1
80	4,8182
100	-10,3636
120	4,4545
140	-0,7273
160	4,0909
180	-1,0909
200	-6,2727
220	3,5455
240	8,3636
260	-6,8182

Şekil 14.11

Bu değerler bulmak istediğimiz (e_i ’ler) hata terimleridir.

Daha farklı bir yol takip edip aynı regresyon denklemini, yine Minitab’ı kullanarak, regresyon doğrusunu scatter diyagram ile birlikte çizdirerek de hesaplayabilirsiniz. Bu işlemi gerçekleştirmek için, “Stat” menüsünden “regression”ı seçiniz ve “fitted line plot” (en uygun doğru) seçeneğine tıklayınız. Karşınıza çıkan diyalog ekranında “response” yazan yere (bağımlı değişkeni temsil eder) “C1”(bizim örnekleminizdeki Y1 değişkeni, tüketim harcamaları), “predictor” yazan yere de (bağımsız değişkeni ifade eder) “Y1”(bizim örnekleminizdeki X1 değişkeni, gelir düzeyi) yazıp “OK” tuşuna bastığınızda şekil 14.7’yi elde edersiniz. Bu şekil üzerinde regresyon denklemini de görebilirsiniz.

Şimdi ise tablo 14.1’deki anakütlemizden başka bir örneklem seçelim. Bu örnekleme ait olan veriler tablo 14.6’da yer almaktadır.

Tablo 14.6: Tablo 14.1'deki anakütleden elde edilmiş diğer bir rassal örneklem

Y	X
55	80
88	100
90	120
80	140
118	160
120	180
145	200
135	220
145	240
175	260

Bu örnekleme kullanılarak, aşağıdaki ÖRF elde edilebilir.

$$Y = 17.2 + 0.576X + e_i$$

Açıkça gözükmektedir ki ÖRF'lerin her ikisinde örneklem dalgalanmalarından dolayı ARF'den sayısal olarak farklıdır. Ayrıca, ikinci örneklemden elde ettiğimiz ÖRF, ARF'ye daha yakın olduğundan dolayı ikinci örneklemin daha iyi olduğu sonucuna varabiliriz. Ancak, gerçek hayatta bu kadar şanslı olmayabiliriz, dolayısıyla istatistiksel test prosedürlerini kullanarak anakütle parametrelerini tahmin etmeye çalışacağız. Şimdi bu testlerin nasıl gerçekleştirildiğini inceleyelim.

b Katsayısının İstatistiksel Anlamlılığının Test Edilmesi

Bu bölümde, b katsayısının istatistiksel anlamlılığını nasıl test edeceğimizi inceleyeceğiz.

Ashnda test ettiğimiz anakütle parametresi β 'nin değerinin sıfıra eşit olup olmadığıdır. b katsayısının selimizdeki örneklem verilerinden b katsayısını hesaplayıp, bu b katsayısını kullanarak testimizi gerçekleştirebiliriz. Peki β anakütle parametresinin sıfıra eşit olması bize ne ifade etmektedir? Eğer β değeri, sıfıra eşit olması, açıklayıcı değişkenimizin bağımlı değişken üzerinde hiçbir etkisi olmadığını gösterir. Dolayısıyla, biz bu testi gerçekleştirirken açıklayıcı değişkenin bağımlı değişkeni gerçekten açıklayıp açıklamadığını test etmekteyiz. Bu testi aşağıdaki gibi oluşturabiliriz:

$$H_0 : \beta = 0$$

$$H_A : \beta \neq 0$$

Peki neden bu tarz testleri gerçekleştirirken bu tarz testlere istatistiksel anlamlılık testi denilmektedir? Örneğin, b katsayısını 0.2 olarak hesapladığımızı düşünelim. Bu katsayı bize X değişkeninin Y değişkeni üzerinde pozitif bir etkisi olduğunu göstermektedir. X değişkenini bir birim arttırdığımızda, Y değişkeni de 0.2 birim artmaktadır. İlk örneğimizi kullanacak olursak, b = 0.2 olması, gelir düzeyimizin bir birim arttığında, tüketim harcamalarımızın 0.2 birim artması

anlamına gelmektedir. Tüm bu yorumların yanısıra, b katsayısı her ne kadar 0.2'ye eşit olursa olsun anakütle parametresi gerçekte 0'a eşit olabilir. Dolayısıyla, açıklayıcı değişkenimiz olan X (gelir düzeyi) değişkeninin, Y değişkeni (tüketim harcamaları) üzerinde bir etkisinin olmadığını sonucuna varırız. Böyle bir durumda X değişkeni, Y değişkenini açıklama gücünü yitirmektedir. Bir başka ifade ile “anlamlılığını” kaybetmiştir. Bu nedenden dolayı bu tarz testlere istatistiksel anlamlılık testleri denir.

Aynı yorumlar, a katsayısının hesaplanmasında ve testinde de geçerlidir. Fakat, açıklayıcı değişkenin katsayısı olduğundan, testlerde b katsayısına daha fazla odaklanmaktayız.

Bu anlamlılık testlerini gerçekleştirebilmek için, modelin standart sapmasını bilmemiz gerekmektedir. (t testini nasıl yaptığımızı hatırlayınız. Anakütle ortalaması (μ) değeri için test

yapmak istediğimizde, standart sapmasını da bilmemiz gerekmektedir $s_{\bar{x}} = \frac{s}{\sqrt{n}}$). Bu yüzden, b katsayısının standart sapması olan s_b 'yi hesaplamamız gerekmektedir. s_b değerini hesaplamak için öncelikle regresyonun standart sapması olan s'i bulmamız gerekmektedir.

Hataların Kareleri Toplamını (HKT), uygun olan serbestlik derecesine böldüğümüzde Ortalama Hataların Karelerini (OHK) elde ederiz. Bu ifadenin karekökü ise bize regresyon modelimizin standart hatasını vermektedir.

$$s = \sqrt{\frac{\sum (Y - a - bX)^2}{n - 2}} = \sqrt{\frac{HKT}{n - 2}} = \sqrt{OHK} \quad (14.9)$$

Regresyonun standart hatası, anakütle bozucu (hata) terimlerinin, u_i 'lerin tahmin edicisinin standart hatası anlamına gelmektedir ve her gözlem için sabittir.²³ u_i 'in tahmin edicisinin standart hatası ile regresyonun standart hatasının aynı isme sahip olmasının nedeni; Y_i 'nin (koşullu) varyansı, u_i 'lerin varyansına denk olmasındandır. Bu standart sapma (veya standart hata), b katsayısının da standart sapmasını aşağıdaki gibi hesaplamamızda bize yardımcı olur:

$$s_b = \frac{s}{\sqrt{\sum (X - \bar{X})^2}} = \frac{s}{\sqrt{\sum X^2 - \frac{\sum X \sum X}{n}}} \quad (14.10)$$

Tablo 14.2'deki verileri kullanarak, standart hatayı hesaplamaya çalışalım.

$$s = \sqrt{\frac{HKT}{10 - 2}}$$

Hesaplamayı gerçekleştirebilmemiz için önce HKT'yi, $\sum (Y_i - a - bX_i)^2 = \sum e_i^2$ denkliliğini kullanarak bulmamız gerekmektedir. Bu hesaplamada aşağıdaki gib olmalıdır. (Unutmayınız ki Minitab “residuals” (hata terimleri) seçeneğini seçtiğimizde bize “RESID1” adı altında ayrı bir sütun açarak bu değerleri otomatik olarak hesaplamıştı.)

$$Y_i - a - bX_i = e_i \quad e_i^2$$

²³ u_i 'lerin varyansının sabitliği Şekil 14.4'te incelenebilir. Bu şekilde, normal olarak dağılan u_i terimleri, sabit varyansa sahiptirler. (hepsi aynı şekle ve sıfır ortalamaya sahiptir).

70-24.45-0.51(80)=4.8182	23.215
65-24.45-0.51(100)=-10.3636	107.405
90-24.45-0.51(120)=4.4545	19.843
95-24.45-0.51(140)=-0.7273	0.529
110-24.45-0.51(160)=4.0909	16.736
115-24.45-0.51(180)=-1.0909	1.190
120-24.45-0.51(200)=-6.2727	39.347
140-24.45-0.51(220)=3.5455	12.570
155-24.45-0.51(240)=8.3636	69.950
150-24.45-0.51(260)=-6.8182	46.488
	$HKT = \sum e_i^2 = 337.27$

$$s = \sqrt{\frac{337.27}{10-2}} = 6.493$$

Daha önce bu değer ile Minitab çıktısında karşılaşmıştık. Minitab çıktısında da bu değeri S ile gösterilmekteydi. Şimdi bulduğumuz bu değeri, denklem (14.10)'te yerine koyarak b katsayısının standart sapmasını bulalım.

$$s_b = \frac{s}{\sqrt{\sum (X - \bar{X})^2}} = \frac{6.493}{\sqrt{33000}} = 0.03574$$

Yine hesapladığımız bu değerle daha önce Minitab çıktısı sayesinde tanışmıştık. Bizi tanıştıran Minitab çıktısını hatırlayacak olursak:

Regression Analysis

The regression equation is
C1 = 24,5 + 0,509 Y1

Predictor	Coef	StDev	T	P
Constant	24,455	6,414	3,81	0,005
Y1	0,50909	0,03574	14,24	0,000

S = 6,493 R-Sq = 96,2% R-Sq(adj) = 95,7%

Analysis of Variance

Source	DF	SS	MS	F	P
Regression	1	8552,7	8552,7	202,87	0,000
Residual Error	8	337,3	42,2		
Total	9	8890,0			

Şekil 14.12

Eğer β değerinin sıfırdan farklı olup olmadığını test etmek istiyorsak, bu testi gerçekleştirmemiz için β değerinin tahmin edicisini ve bu tahmin edicinin standart sapmasını s_b kullanmamız gerekmektedir. Formal olarak bu hipotezi şöyle ifade ederiz:

$$H_0 : \beta = 0$$

$$H_A : \beta \neq 0$$

ve bu hipotezi test etmek için t-testini kullanırız. T-testi, $(n-2)$ serbestlik derecesiyle aşağıdaki gibi ifade edilebilir:

$$t = \frac{b - \beta}{s_b}$$

Sıfır hipotezimiz $\beta = 0$ olduğundan:

$$t = \frac{b}{s_b} = \frac{0.51 - 0}{0.03574} = \frac{0.51}{0.03574} = 14.24$$

Bu t değeri, olasılık değeri ile birlikte Minitab çıktımızda yer alan değer ile birebir aynıdır.

Anlamlılık düzeyi $\alpha = 0.05$ ve serbestlik derecesi 8 olan, t kritik değeri 1.86'ya eşittir. Dolayısıyla, sonuç olarak açık bir şekilde sıfır hipotezimizi reddederiz. (Minitab çıktısında yer alan p-değeri de kullanılarak aynı sonuca varılabilmektedir). Yapılan t-testi, korelasyon katsayısı kullanılarak yapılan teste paralel sonuçlar vermektedir. Ancak burada kullandığımız t-testinin avantajı alternatif hipotez belirtilebilmesidir. Örnek olarak, aşağıdaki hipotezi test etmeye çalışalım.

$$H_0 : \beta = 1$$

$$H_A : \beta \neq 1$$

Bu durumda yine t-testi aşağıdaki gibi olur:

$$t = \frac{b - \beta}{s_b}$$

Sıfır hipotezimiz $\beta = 1$ olduğundan t-istatistik değeri aşağıdaki gibi hesaplanır:

$$t = \frac{b - 1}{s_b} = \frac{0.51 - 1}{0.03574} = -13.71$$

Yukarıda belirlediğimiz t-kritik değerini yine kullanacağımızdan, tekrar sıfır hipotezimizi reddederiz.

Belirleyicilik Katsayısı

En küçük kareler yöntemi, X ve Y değişkenleri arasındaki ilişkiye doğrusal bir yakınsama sağlamaktadır. Burada akla gelen önemli soru, regresyon doğrusunun bu yakınsamayı ne kadar iyi bir şekilde gerçekleştirdiği, bu doğrunun ne kadar iyi çizildiğidir? Bu sorunun cevabını verebilmek için bazı kavramlara değinmemiz gerekmektedir. Hataların Kareleri Toplamı'nı daha önce aşağıdaki gibi tanımlamıştık (Denklem 14.5):

$$HKT = \sum_{i=1}^n e_i^2 = \sum_{i=1}^n (Y_i - a - bX_i)^2$$

Bu miktar, regresyonun açıklanamayan kısmını oluşturmuştur. (bu kısım regresyon doğrusu tarafından açıklanamamaktadır). Toplam kareler toplamı (TKT) ise aşağıdaki gibi tanımlanabilir:

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

Bu ifade bize, bağımlı değişkendeki (Y) toplam farklılaşmayı vermektedir. Toplam Kareler Toplamı'nı (n-1)'e (ait olduğu serbestlik derecesine) böldüğümüzde, Y değişkeninin varyansını elde ederiz.

Regresyon doğrusu tarafından açıklanan kareler toplamı, Regresyon Kareler Toplamı olarak tanımlanır:

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

Bu ifade, regresyonun açıklanan kısmını temsil etmektedir. Tanımlamış olduğumuz bu üç miktar arasındaki ilişki aşağıdaki gibidir:

$$TKT = RKT + HKT \quad (14.11)$$

Bir başka ifade ile, toplam karelerin toplamı, regresyon kareleri toplamı (açıklanan kısım) ve hataların kareleri toplamı (açıklanamayan kısım) olmak üzere iki parçadan oluşmaktadır. Dolayısıyla, regresyon doğrumuzun ne kadar iyi olduğunu ölçmek için denklem (14.11)'i kullanabiliriz. Eğer eşitliğin her tarafını toplam kareler toplamına (TKT) bölersek:

$$1 = \frac{RKT}{TKT} + \frac{HKT}{TKT}$$

$\frac{RKT}{TKT}$ miktarı, Y değişkeninde meydana gelen farklılaşmanın, regresyon doğrusu tarafından açıklanan kısmını ifade etmektedir. Bu bize doğrunun uygunluğuna dair iyi bir ölçektir. Bu ölççe, farklılaşma katsayısı denir ve R^2 ile gösterilir. $\frac{RKT}{TKT} = R^2$ Yukarıdaki denkliği yeniden düzenleyecek olursak:

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

R^2 , 0 ile 1 arasında değerler almaktadır. Hataların kareleri toplamı azaldığında, R^2 artar ve daha iyi bir regresyon doğrusu çizilebilir.

Tekrar, tablo 14.1'deki Tüketim/Gelir örneklemimizdeki verileri kullanarak R^2 değerini hesaplayabiliriz. Daha önceden $HKT = \sum_{i=1}^n e_i^2 = 337.27$ değerini hesaplamıştık. Şimdi ise

hesaplamamız gereken TKT dır. $TKT = \sum_{i=1}^n (Y_i - \bar{Y})^2 = 8890$

Belirleyicilik katsayısını ise aşağıdaki gibi hesaplayabiliriz:

$$R^2 = 1 - \frac{HKT}{TKT} = 1 - \frac{337.27}{8890} = 0.962$$

Bu değeri (R-Sq olarak) Minitab çıktısından da görebiliriz.

Regression Analysis

The regression equation is
C1 = 24,5 + 0,509 Y1

Predictor	Coef	StDev	T	P
Constant	24,455	6,414	3,81	0,005
Y1	0,50909	0,03574	14,24	0,000

S = 6,493 R-Sq = 96,2% R-Sq(adj) = 95,7%

Analysis of Variance

Source	DF	SS	MS	F	P
Regression	1	8552,7	8552,7	202,87	0,000
Residual Error	8	337,3	42,2		
Total	9	8890,0			

Şekil 14.13

Belirleyicilik Katsayısı'nın yaklaşık olarak %96 olarak çıkması, tüketim harcamalarındaki farklılaşmanın, gelir düzeyindeki farklılaşmadan kaynaklandığının yüzde 96'sının açıklandığı anlamına gelmektedir.

Düzeltilmiş- R^2 (genellikle $\bar{R}^2$ sembolü ile ifade edilir), toplam kareler toplamını ile serbestlik derecesini de ele alarak inceler. Basit doğrusal durumu incelediğimizden dolayı $k=2$ olarak, aşağıdaki gibi hesaplanabilir:

$$\bar{R} = 1 - \frac{\frac{HKT}{(n-k)}}{\frac{TKT}{(n-1)}} = 1 - \frac{HKT}{TKT} \times \frac{(n-1)}{(n-k)} = 1 - \frac{337.27}{8890} \times \frac{(10-1)}{(10-2)} = 0.9578$$

Bu ölçüm yöntemi, birden fazla açıklayıcı değişken durumu söz konusu olduğunda, çoklu regresyon modellerinde belirleyicilik katsayısına göre daha faydalıdır.

Kestirim

Regresyon analizinde amacımız, değişkenler arasındaki ilişkiyi ortaya çıkararak, verili bir bağımsız değişken değeri ile bağımlı değişkeni tahmin etmek idi. ÖRF, Y'nin tahmin edilmiş değerini hesaplayan aşağıdaki gibi matematiksel bir denklemdir:

$$\hat{Y}_i = a + b X_i$$

Herhangi bir verili bağımsız değişken değeri için (X_0), regresyon fonksiyonu, (Y_0) bağımlı değişken değerini aşağıdaki denklik ile tahmin eder:

$$\hat{Y}_0 = a + b X_0$$

Örneğin, daha önceden hesapladığımız ÖRF'yi ele alalım $C = 24.5 + 0.509Y$. Gelir düzeyi $X_0 = 350$ iken, tüketim miktarını tahmin etmeye çalışalım, $C = ?$. ÖRF fonksiyonunda X_0 'ı yerine koyarsak:

$$\begin{aligned} C &= 24.5 + 0.509 \times 350 \\ &= 202.65 \end{aligned}$$

ÖRF fonksiyonu, tüketim miktarını \$202.65 olarak tahmin etmektedir. Bu aile gerçekte aşağı yukarı bu tüketim düzeyinde bir tüketim gerçekleşmektedir.

$E(\hat{Y}_0) = Y_0$ olduğundan $\hat{Y}_0$, gerçek anakütle bağımlı değişken değerinin Y_0 sapmasız bir tahmin edicisidir (Ek 14.3'e bakınız).

Bu yüzden, ÖRF aşağı yukarı Y_0 'ı aşağıdaki varyans ile birlikte doğru bir şekilde tahmin etmektedir. Varyans formülü ise aşağıdaki gibi verilmektedir (Ek 14.4'e bakınız).

$$\text{var}(\hat{Y}_0) = \sigma_u^2 \left[\frac{1}{n} + \frac{(X_0 - \bar{X})^2}{\sum x_i^2} \right]$$

Hata terimlerinin varyansı σ_u^2 , bilinmiyorsa:

$$\text{est. var}(\hat{Y}_0) = s_u^2 \left[\frac{1}{n} + \frac{(X_0 - \bar{X})^2}{\sum x_i^2} \right]$$

kullanılır. Hatırlanacağı (mutlaka hatırlamanız gerektiği) üzere, standart hatayı veya tahmin edilmiş standart hatayı bulabilmemiz için, varyansın veya tahmin edilmiş varyansın karekökünü almak gereklidir.

Şimdi ise Y_0 değeri için güven aralığı oluşturabiliriz:

$$\begin{aligned} Y_0 &= \hat{Y}_0 \pm z_{\alpha/2} SE(\hat{Y}_0) \\ Y_0 &= \hat{Y}_0 \pm t_{\alpha/2; n-2} \text{est. SE}(\hat{Y}_0) \end{aligned}$$

Bu güven aralığı, ortalama olarak Y_0 değerinin $(1-\alpha)$ olasılık ile yer aldığı aralığı ifade etmektedir. Bir başka ifade ile, gelir düzeyi \$1000 olan aileleri ele aldığımızda, ortalama tüketim düzeyi % $(1 - \alpha)$ kere bu aralığın içerisinde yer alacaktır. Bu nedenden dolayı, bu aralığa **Ortalama Kestirim Aralığı** denilmektedir.

Eğer Y değişkeninin ortalama değerinden ziyade, tek bir değeri için bir güven aralığı oluşturmak istiyorsak, **Tekil Kestirim Aralığı** aralığını oluşturmamız gerekmektedir. Bu aralık verili bir X_0 değerine tekabül edecek olan Y değerini içermesini beklediğimiz aralıktır.

Örneğimizden yola çıkarak bir başka ifade ile, sadece bir ailenin tüketim düzeyinin $(1 - \alpha)$ olasılıkla bu aralıkta olacağını göstermektedir.

Tekil kestirim gerçekleştirilirken, gerçek bağımlı değişken değeri ile tahmin ettiğimiz bağımlı değişken değeri arasındaki bir fark oluşmaktadır. Bu fark *kestirim hatası* olup aşağıdaki gibi ifade edilebilir:

$$e_0 = Y_0 - \hat{Y}_0$$

Tahminlerimizin doğru olması için ortalama olarak kestirim hatalarının beklenen değerinin 0 olması gerekmektedir. Dolayısıyla, $\hat{Y}_0$ değişkeni Y_0 'ın sapmasız bir tahmin edicisidir. Bu sapmasızlığın gösterimini ortalama kestirim aralığının gösterimi esnasında gerçekleştirmiştik. Bu yüzden, şimdi ise sadece güven aralığında kullanmak üzere kestirim hatasının varyansını inceleyelim. Bu varyans aşağıda verilmektedir.

$$\text{var}(e_0) = \sigma_u^2 \left[1 + \frac{1}{n} + \frac{(X_0 - \bar{X})^2}{\sum x_i^2} \right]$$

Dolayısıyla, tekil kestirim aralığı aşağıdaki gibi elde edilir:

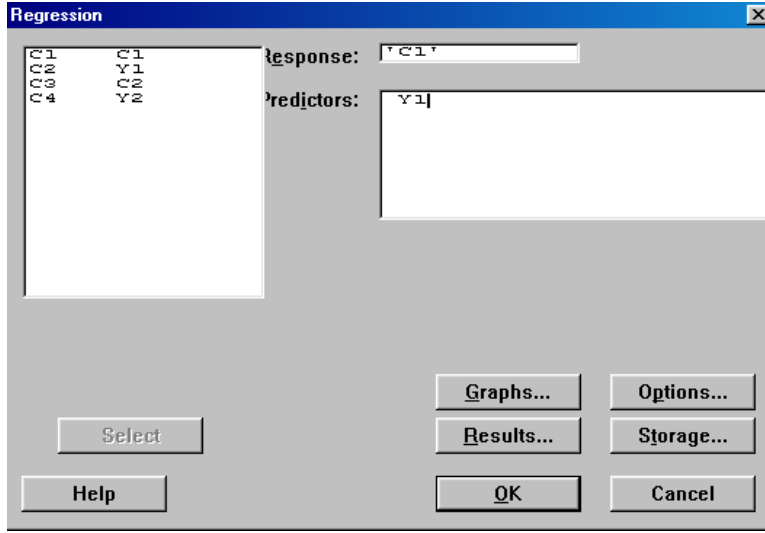
$$Y_0 = \hat{Y}_0 \pm z_{\alpha/2} SE(Y_0 - \hat{Y}_0)$$

Hata terimlerinin varyansının bilinmediği durumda ise, varyansın tahmin edicisi s_u^2 ve dağılım olarak da $(n - 2)$ serbestlik derecesi ile t-dağılımı kullanılmaktadır.

$$Y_0 = \hat{Y}_0 \pm t_{\alpha/2; n-2} \text{est.} SE(Y_0 - \hat{Y}_0)$$

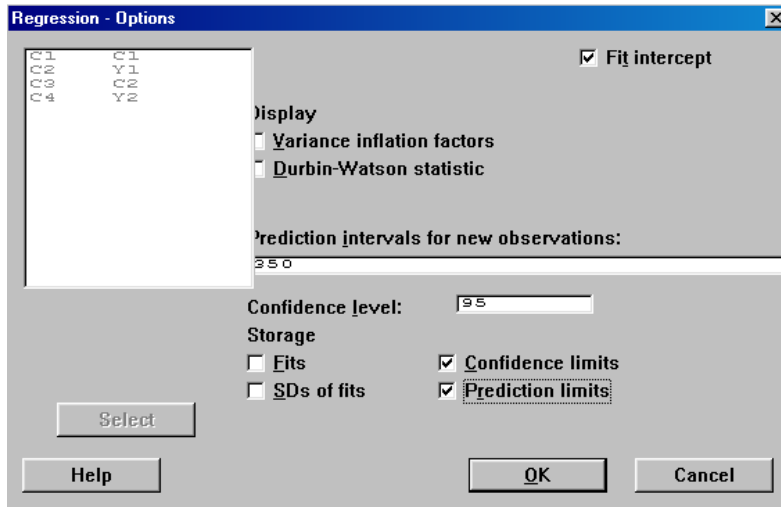
Minitab Uygulamaları

Gerekli veriler için öncelikle “*tuketim1.mtw*” dosyasını açınız. “Stat” menüsünden “*regression*”ı seçenerek “*regression*” seçeneğinin üzerine tıklayınız. Karşınıza çıkacak olan diyalog ekranında “*response*” yazan yere (bağımlı değişkeni temsil eder) “Y1” değişkenini ve “*predictor*” yazan yere de (bağımsız değişkeni ifade eder) “X1” yazınız.



Şekil 14.14

Diyalog ekranında yer alan “Options” menüsüne tıklayarak özellikleri aşağıdaki gibi belirleyiniz.



Şekil 14.15

OK tuşuna basıp anakütle diyalog ekranına geri döndükten sonra tekrar OK tuşuna basarak aşağıdaki çıktı elde edersiniz.

Regression Analysis

The regression equation is
C1 = 24,5 + 0,509 Y1

Predictor	Coef	StDev	T	P
Constant	24,455	6,414	3,81	0,005
Y1	0,50909	0,03574	14,24	0,000

S = 6,493 R-Sq = 96,2% R-Sq(adj) = 95,7%

Analysis of Variance

Source	DF	SS	MS	F	P
Regression	1	8552,7	8552,7	202,87	0,000
Residual Error	8	337,3	42,2		
Total	9	8890,0			

Predicted Values

Fit	StDev Fit	95,0% CI	95,0% PI
202,64	6,75	(187,06; 218,21)	(181,03; 224,24)

XX denotes a row with X values away from the center
XX denotes a row with very extreme X values

Şekil 14.15

Çıktının çerçeve içine alınmış kısmına odaklanırsak, tahmin edilmiş değeri 202,64 (*“Fit”* adı altında) ve bu değere ait olan standart sapma değerini, ortalama kestirim aralığını (*“CI”* adı altında), kestirim aralığını (*“PI”* adı altında) görebiliriz. Bu çıktıdaki ortalama kestirim aralığı, gelir düzeyi 350\$ iken, *ortalama olarak* tüketim düzeyinin %95 olasılık ile (187,06; 218,21) aralığında yer aldığını göstermektedir. Kestirim aralığı ise, gelir düzeyi 350\$ iken, tüketim düzeyinin %95 olasılık ile (181,03; 224,24) aralığında yer aldığını göstermektedir.

Ekler

Ek 14.1

Kolaylık sağlama açısından, i indislerini şimdilik görmezlikten gelelim. Her minimizasyon veya maksimizasyon işleminde olduğu gibi ifadenin türevini alıp sıfıra eşitleyelim.

$$\frac{\partial HKT}{\partial a} = 2 \sum (Y - a - bX) \times (-1) = 0 \quad (E.1)$$

$$\frac{\partial HKT}{\partial b} = 2 \sum (Y - a - bX) \times (-X) = 0 \quad (E.2)$$

Denklem (E.1) ve (E.2), normal denklemlerimizi oluşturmaktadırlar. İki bilinmeyen (a ve b) ve iki denklem bulunduğundan bu denklem sistemi birkaç küçük işlemden sonra kolaylıkla çözülebilir. Öncelikle, denklem (E.1)'i $-(1/2)$ ile çarpalım:

$$\begin{aligned} \sum (Y - a - bX) &= 0 \\ \sum Y - \sum a - b \sum X &= 0 \end{aligned}$$

$\sum a = na$ olduğundan:

$$\sum Y - na - b \sum X = 0 \quad (E.3)$$

$\frac{\sum Y}{n} = \bar{Y}$ (Y değişkeninin ortalaması) ve $\frac{\sum X}{n} = \bar{X}$ (X değişkeninin ortalaması) olduğundan denklem (E.3)'ü de kullanarak a için gerekli çözümü aşağıdaki gibi bulabiliriz.

$$a = \frac{\sum Y}{n} - b \frac{\sum X}{n} = \bar{Y} - b\bar{X} \quad (E.4)$$

Bu şekilde Denklem (14.8)'i de türetmiş olduk. Şimdi tekrar denklem (E.2)'ye dönelim ve her tarafı $-1/2$ ile çarpalım:

$$\sum (Y - a - bX) \times X = 0$$

Gerekli işlemleri yaptıktan sonra aşağıdaki denklemi elde ederiz:

$$\sum (YX - aX - bX^2) = 0$$

toplam operatörünü kullanarak parantezi açacak olursak:

$$\sum YX - a \sum X - b \sum X^2 = 0 \quad (E.5)$$

Denklem (E.4)'de elde ettiklerimizi denklem (E.5)'de yerine koyacak olursak:

$$\sum XY - \left(\frac{\sum Y}{n} - b \frac{\sum X}{n} \right) \sum X - b \sum X^2 = 0$$

Parantezi açalım:

$$\sum XY - \frac{\sum Y \sum X}{n} - b \frac{\sum X \sum X}{n} - b \sum X^2 = 0$$

Eşitliğin her iki tarafında -1 ile çarpıp ifadeni tekrar düzenlersek:

$$\frac{\sum Y \sum X}{n} - \sum XY + b \sum X^2 - b \frac{\sum X \sum X}{n} = 0$$

b'yi sol tarafta yalnız bırakarak denklemi çözdüğümüzde aşağıdaki denkleği elde ederiz:

$$b \sum X^2 - b \frac{\sum X \sum X}{n} = \sum XY - \frac{\sum Y \sum X}{n}$$
$$b = \frac{\sum XY - \frac{\sum Y \sum X}{n}}{\sum X^2 - \frac{\sum X \sum X}{n}}$$

Bu şekilde de Denklem (14.6)'yı göstermiş olduk.

Ek 14.2

Denklem (14.6)'ın pay kısmını aşağıdaki gibi yeniden düzenleyebiliriz:

$$\sum XY - \frac{\sum Y \sum X}{n} = \sum (X - \bar{X})(Y - \bar{Y})$$

ve payda kısmını da aşağıdaki gibi yeniden düzenleyebiliriz:

$$\sum X^2 - \frac{\sum X \sum X}{n} = \sum (X - \bar{X})^2$$

Bu kısaltılmış ifadeler kullanılarak Denklem (14.7) elde edilebilir.

Ek 14.3 $E(\hat{Y}_0) = Y_0$ ifadesinin kanıtlanması

$E(\hat{Y}_0) = Y_0$ koşulu aşağıdaki biçimde de yazılabilir

$$E(Y_0 - \hat{Y}_0) = 0$$

Dolayısıyla

$$\begin{aligned} E(Y_0 - \hat{Y}_0) &= E(\alpha + \beta X_0 + u_i - a - bX_0) \\ &= E(\alpha) + E(\beta X_0) + E(u_i) - E(a) - E(bX_0) \\ &= \alpha + \beta X_0 + 0 - \alpha - \beta X_0 \\ &= 0 \end{aligned}$$

Ek 14.4 $\text{var}(\hat{Y}_0) = \sigma_u^2 \left[\frac{1}{n} + \frac{(X_0 - \bar{X})^2}{\sum x_i^2} \right]$ ifadesinin kanıtlanması

$$\begin{aligned}
 \text{var}(\hat{Y}_0) &= \text{var}(a + bX_0) \\
 &= \text{var}(a) + X_0^2 \text{var}(b) + 2X_0 \text{cov}(a, b) \\
 &= \left[\frac{\sum X_i^2}{n \sum x_i^2} \right] \sigma_u^2 + X_0^2 \left[\frac{\sigma_u^2}{\sum x_i^2} \right] + 2X_0(-\bar{X}) \left(\frac{\sigma_u^2}{\sum x_i^2} \right) \\
 &= \sigma_u^2 \left[\frac{\sum X_i^2}{n \sum x_i^2} + \frac{X_0^2}{\sum x_i^2} - 2 \frac{X_0 \bar{X}}{\sum x_i^2} \right] \\
 &= \sigma_u^2 \left[\frac{\sum X_i^2 + nX_0^2 - 2nX_0 \bar{X}}{n \sum x_i^2} \right]
 \end{aligned}$$

Payda, $2n\bar{X}^2$ ekleyip $2n\bar{X}^2$ çıkarsak ve $\bar{X} = \frac{\sum X_i}{n}$ dönüşümünü yapsak:

$$\begin{aligned}
 \text{var}(\hat{Y}_0) &= \sigma_u^2 \left[\frac{\sum X_i^2 + nX_0^2 - 2nX_0 \bar{X} + 2n\bar{X}^2 - 2n\bar{X}^2}{n \sum x_i^2} \right] \\
 &= \sigma_u^2 \left[\frac{\sum X_i^2 + nX_0^2 - 2nX_0 \bar{X} + 2n\bar{X}^2 - 2\bar{X} \sum X_i}{n \sum x_i^2} \right]
 \end{aligned}$$

$2n\bar{X}^2$ 'i ayırarak denklemi tekrar düzenlediğimizde:

$$\text{var}(\hat{Y}_0) = \sigma_u^2 \left[\frac{\sum X_i^2 - 2\bar{X} \sum X_i + n\bar{X}^2 + nX_0^2 - 2nX_0\bar{X} + n\bar{X}^2}{n \sum \tilde{x}_i^2} \right]$$

$\left[\sum X_i^2 - 2\bar{X} \sum X_i + n\bar{X}^2 = \sum (X_i - \bar{X})^2 = \sum x_i^2 \right]$ dönüşümünü yapalım ve
 $\left[nX_0^2 - 2nX_0\bar{X} + n\bar{X}^2 = n(X_0 - \bar{X})^2 \right]$ ifadesini kullanarak aşağıdaki formül elde edilebilir.

$$\text{var}(\hat{Y}_0) = \sigma_u^2 \left[\frac{\sum \tilde{x}_i^2 + n(X_0 - \bar{X})^2}{n \sum x_i^2} \right]$$

Ek 14.4 $\text{var}(e_0) = \sigma_u^2 \left[1 + \frac{1}{n} + \frac{(X_0 - \bar{X})^2}{\sum x_i^2} \right]$ ifadesinin kanıtlanması

$$\begin{aligned} \text{var}(e_0) &= \text{var}(Y_0 - \hat{Y}_0) \\ &= \text{var}(Y_0) + \text{var}(\hat{Y}_0) \end{aligned}$$

$\hat{Y}_0$ ve Y_0 birbirinden bağımsız olduğundan dolayı varyans:

$$\text{var}(e_0) = \sigma_u^2 + \sigma_u^2 \left[\frac{1}{n} + \frac{(X_0 - \bar{X})^2}{\sum x_i^2} \right]$$

Dolayısıyla

$$\text{var}(e_0) = \sigma_u^2 \left[1 + \frac{1}{n} + \frac{(X_0 - \bar{X})^2}{\sum x_i^2} \right]$$

Bölüm 15: Çoklu Doğrusal Regresyon

Giriş

Bir önceki bölümde tek bir bağımlı değişken ile tek bir açıklayıcı değişken arasındaki ilişkiyi inceledik. Çoklu regresyon modelinde ise, açıklayıcı değişkenlerden oluşan bir setin tek bir bağımlı değişken ile ilişkilerini inceleyeceğiz. Örneğin, talep miktarını bağımlı değişken olarak incelerken, fiyat seviyesi ile gelir düzeyini birer açıklayıcı değişken olarak ele alabiliriz.

Talep miktarı ile fiyat seviyesi ve gelir düzeyi arasındaki ilişkiyi incelerken, açıklayıcı değişkenlerin birlikte etkileri ile kısmi etkilerini birbirinden ayırmak gereklidir. Örnek olarak, gelir düzeyini sabit tutarak sadece fiyatların talep üzerindeki etkisini veya fiyatları sabit tutarak sadece gelir düzeyinin talep miktarı üzerindeki etkisini incelemek isteyebiliriz. Çoklu regresyon modeli bize bu isteklerimiz çerçevesinde çalışma imkanı verir.

Çoklu Regresyon Modelinin Varsayımları

Modelimizde iki tane açıklayıcı değişken olduğunu düşünelim ve modelimizi aşağıdaki gibi oluşturalım:

$$Y_i = \alpha + \beta X_i + \gamma Z_i + u_i \quad (15.1)$$

Çoklu regresyonda, α , β ve γ 'nın tahmin edici değerlerini (a, b, ve c gibi) seçerek, hataların kareleri toplamını minimize etmeye çalışırız. Bölüm 14'te de incelediğimiz gibi tahmin edicileri bulmak için normal doğrusal denklem setini, a, b, ve c değerleri çözmemiz gerekmektedir. Bu bölümde yapacağımızın bölüm 14'tekinden tek farkı burada çözmemiz gereken iki değil üç denklemin bulunmasıdır.

Tahmin edicileri nasıl bulacağımızı incelemeden önce, yukarıda yer alan denklem hakkında birkaç yorumda bulunalım. Hata terimleri, Y'de meydana gelen ölçüm hataları, Y, X ve Z arasındaki ilişkinin belirlenmesinde meydana gelen hataların hepsini kapsamaktadır. Çoklu regresyon analizinde, aynı basit regresyon analizinde olduğu gibi, hata terimleri ile alakalı olarak birkaç temel varsayımda bulunmamız gerekmektedir. Bu varsayımlar:

- $E(u_i) = 0$

(hataların beklenen değerinin 0 olduğu, hata ortalamalarının sıfır olduğu). Bu varsayımın işlemediği durumlarda katsayı tahmin edicileri sapmalı değerlere sahip olurlar.

- $E(u_i^2) = \sigma^2$

(hataların karelerinin beklenen değeri sabit bir değerdir, bu varsayıma hata terimlerini homoskedastik (eş dağılımı) olması varsayımında denilebilir). Bu varsayımın işlememesi halinde heteroskedastisite (farklı dağılım) durumu söz konusu olur ve tahmin edilmiş standart hata değerleri sapmalı olurlar.

- $E(u_i u_j) = 0 \quad i \neq j$ iken

(hata terimleri arasındaki kovaryans sıfırdır. Bu varsayım özellikle zaman serileri ile alakalı uygulamalarda çok önem teşkil eder. Bu varsayıma göre hatalar birbirlerinden bağımsızdır). Bu varsayımın işlememesi halinde seri korelasyon durumu görülür ve tahmin edilmiş standart hatalar sapmalı çıkar.

- $E(u_i, X_i) = E(u_i, Z_i) = 0$

(hata terimleri ile açıklayıcı değişkenler arasındaki kovaryans değeri sıfıra eşittir, başka ifade ile u_i terimi, X ve Z değişkenlerinden bağımsızdır). Bu varsayımın işlememesi halinde tahmin edilmiş katsayılar sapmalı değerler alırlar.

- $u_i \sim N(0, \sigma^2)$

(hata terimleri, sıfır ortalama ve sabit varyans ise normal olarak dağılırlar). Bu varsayımın işlememesi halinde, t ve F test sonuçları tutarlılığını ve geçerliliğini yitirir.

Kitabımızın hedefleri arasında bu varsayımları detaylı bir şekilde incelemek yoktur. Fakat, hata terimleri hakkındaki bu varsayımların herhangi birinin işlemediği durumlarda sonuçlarında neler olabileceği bilinmelidir.

Çoklu Regresyon Analizini Kullanarak Tahmin Edicilerin Bulunması

İki değişkenli regresyon analizinde, anakütle terimleri olan α ve β 'nin tahmin edicilerinin formüllerini türettik ve onları a ve b diye adlandırdık. Şimdi ise aynı işlemleri çoklu regresyon analizi için gerçekleştireceğiz.

$$Y_i = \alpha + \beta X_i + \gamma Z_i + u_i \quad (15.2)$$

Çoklu regresyon modelimizde ise a ve b tahmin edicileri, α ve β anakütle parametrelerini tahmin etmeye çalışırken; c tahmin edicisi ise γ anakütle parametresini tahmin etmeye çalışmaktadır.

$$Y_i = a + bX_i + cZ_i + e_i \quad (15.3)$$

Tahmin edicilerimizin formülleri bir önceki bölümde bulduklarımıza benzerdir. Sadece şimdi artık elimizde bir tane daha açıklayıcı değişken daha vardır. Bir önceki bölüme benzer bir şekilde a tahmin edicisinin formülünü aşağıdaki gibi elde ederiz:

$$a = \bar{Y} - b\bar{X} - c\bar{Z}$$

Bu formülün değerini hesaplayabilmemiz için öncelikle b ve c değerlerini bulmamız gerekmektedir. Şimdi ise b ve c değerlerini nasıl bulacağımızı inceleyelim. b tahmin edicisini hesaplamak için aşağıdaki formül kullanılırken:

$$b = \frac{\sum (X_i - \bar{X})(Y - \bar{Y}) \times \sum (Z_i - \bar{Z})^2 - \sum (Z_i - \bar{Z})(Y - \bar{Y}) \times \sum (X_i - \bar{X})(Z_i - \bar{Z})}{\sum (X_i - \bar{X})^2 \times \sum (Z_i - \bar{Z})^2 - (\sum (X_i - \bar{X})(Z_i - \bar{Z}))^2}$$

c tahmin edicisini hesaplamak için aşağıdaki formül kullanılır:

$$c = \frac{\sum (Z_i - \bar{Z})(Y - \bar{Y}) \times \sum (X_i - \bar{X})^2 - \sum (X_i - \bar{X})(Y - \bar{Y}) \times \sum (X_i - \bar{X})(Z_i - \bar{Z})}{\sum (X_i - \bar{X})^2 \times \sum (Z_i - \bar{Z})^2 - \left(\sum (X_i - \bar{X})(Z_i - \bar{Z}) \right)^2}$$

Notasyonlarda biraz değişiklik yaparak bu sevimli formülleri biraz daha sevecen bir hale getirelim. Sapma notasyonlarını kullanarak $x_i = X_i - \bar{X}$, $y_i = Y_i - \bar{Y}$, ve $z_i = Z_i - \bar{Z}$, (küçük harfler ortalama sapmaları göstermektedir) formüllerde yerine koyalım.

$$b = \frac{\sum x_i y_i \times \sum z_i^2 - \sum z_i y_i \times \sum x_i z_i}{\sum x_i^2 \times \sum z_i^2 - \left(\sum x_i z_i \right)^2}$$

$$c = \frac{\sum z_i y_i \times \sum x_i^2 - \sum x_i y_i \times \sum x_i z_i}{\sum x_i^2 \times \sum z_i^2 - \left(\sum x_i z_i \right)^2}$$

Ancak hala bu formüller son derece karışıktır, neyse ki hatırlanmaları gerekmemektedir²⁴. Formüllerin türetimindeki prensip tamamen iki değişkenli regresyon analizinde olduğu gibidir. HKT'nin minimize edilmesi prensibine dayanır.

Çoklu regresyon analizinde regresyonun standart hata formülü aşağıdaki gibidir.

$$s = \sqrt{\frac{\sum (Y_i - a - bX_i - cZ_i)^2}{n-3}} = \sqrt{\frac{\sum e_i^2}{n-3}} = \sqrt{\frac{HKT}{n-3}}$$

Yine bu standart hata değerini hem b'nin hem de c'nin standart sapmalarını bulurken kullanacağız. b tahmin edicisinin standart hatası s_b :

$$s_b = s \sqrt{\frac{\sum z_i^2}{\sum x_i^2 \times \sum z_i^2 - \left(\sum x_i z_i \right)^2}} \quad (15.4)$$

c tahmin edicisinin standart hatası s_c :

$$s_c = s \sqrt{\frac{\sum x_i^2}{\sum x_i^2 \times \sum z_i^2 - \left(\sum x_i z_i \right)^2}} \quad (15.5)$$

Formülü biraz daha işimize yarayacak bir hale getirmeye çalışalım. Denklem 4'ün payını ve paydasını karekökün içerisinde, $1/\sum z_i^2$ ifadesine bölelim.

²⁴ Genel olarak bu tarz operasyonlarda matrisler kullanıldığından bu formüllere ihtiyaç duyulmamaktadır. Ancak, matris cebiri dersimizin alanı içerisine girmediğinden dolayı bu hususa değinilmeyecektir.

$$s_b = s \sqrt{\frac{\sum z_i^2 / \sum z_i^2}{\frac{\sum x_i^2 \times \sum z_i^2 - (\sum x_i z_i)^2}{\sum z_i^2}}} = s \sqrt{\frac{1}{\sum x_i^2 - \frac{(\sum x_i z_i)^2}{\sum z_i^2}}}$$

Daha sonra payı ve paydayı $1/\sum x_i^2$ ile çarpalım

$$s_b = s \sqrt{\frac{1/\sum x_i^2}{\frac{\sum x_i^2 - (\sum x_i z_i)^2}{\sum x_i^2 \sum z_i^2}}} = s \sqrt{\frac{1/\sum x_i^2}{1 - \frac{(\sum x_i z_i)^2}{\sum x_i^2 \sum z_i^2}}} \quad (15.6)$$

$\frac{(\sum x_i z_i)^2}{\sum x_i^2 \sum z_i^2}$ ifadesi, Bölüm 13'te belirttiğimiz, X ve Y değişkenleri arasındaki korelasyon katsayısının karesidir.

$$r_{xz} = \frac{\sum x_i z_i}{\sqrt{\sum x_i^2 \sum z_i^2}}$$

Dolayısıyla denklem 15.6'yı tekrar ifade edecek olursak:

$$s_b = s \sqrt{\frac{1/\sum x_i^2}{1 - r_{xz}^2}} = s \sqrt{\frac{1}{(1 - r_{xz}^2) \sum x_i^2}} \quad (15.7)$$

ve denklem 15.5'i yeniden ifade edecek olursak

$$s_c = s \sqrt{\frac{1/\sum z_i^2}{1 - r_{xz}^2}} = s \sqrt{\frac{1}{(1 - r_{xz}^2) \sum z_i^2}} \quad (15.8)$$

15.7. ve 15.8. denklemler, b ve c tahmin edicilerinin standart hatalarına alternatif bir gösterim tarzı getirmesinin yanısıra, önemli bir gerçeği de ortaya çıkarmışlardır. Eğer açıklayıcı değişkenler arasında doğrusal bir ilişki yoksa, bir başka ifade ile aralarındaki korelasyon katsayısı sıfır ise ($r_{xz} = 0$); denklemler, iki değişkenli regresyon modelindeki katsayıların standart hata formülleri ile birebir aynı olmaktadır. Dolayısıyla eğer iki açıklayıcı değişken arasında, doğrusal bir ilişki yok ise katsayıların standart hataları iki değişkenli basit regresyon analizinde olduğu gibidir.

$$s_b = s \sqrt{\frac{1}{\sum x_i^2}}$$

$$s_c = s \sqrt{\frac{1}{\sum z_i^2}}$$

Eğer, iki açıklayıcı değişken arasında pozitif veya negatif yönde tam bir ilişki varsa ($r_{xz} = 1$ veya $r_{xz} = -1$), denklem 15.7 ve 15.8 sonsuz değerini alırlar ($s \sqrt{\frac{1}{0}} = \infty$).

Sonuç olarak, X ve Z açıklayıcı değişkenlerinin arasındaki ilişki arttıkça, r_{xz} arttıkça, standart hataları azalmaktadır.

Çoklu Regresyon Modelinde

Çoklu Belirleyicilik Katsayısı

Çoklu regresyon modelinde belirleyicilik katsayısı, bir önceki bölümde olduğu gibi, RKT regresyon kareler toplamı ile HKT hatalar kareler toplamının ayrı ayrı TKT'ye toplam kareler toplamına bölünmesi ile elde edilebilir:

$$TKT = RKT + HKT$$

Eşitliğin her iki tarafını da toplam kareler toplamına bölelim (TKT):

$$1 = \frac{RKT}{TKT} + \frac{HKT}{TKT}$$

Daha öncede belirttiğimiz gibi, Y'deki farklılaşmanın açıklanan kısmı $\frac{RKT}{TKT}$, R^2 veya düzeltilmemiş R^2 'dir. Yukarıdaki eşitliği tekrar düzenleyecek olursak, $\frac{SSR}{SST} = R^2$, aşağıdaki ifadeyi elde ederiz:

$$R^2 = 1 - \frac{SSE}{SST}$$

Şimdiye kadar bir önceki bölümden farklı birşey yapmadık. Bu konunun çoklu regresyon analizindeki farkı, sadece HKT, RKT ve TKT'nin tanımlamalarında ortaya çıkmaktadır. Bu tanımlamalar aşağıdaki gibidir:

$$HKT = \sum_{i=1}^n e_i^2 = \sum_{i=1}^n (Y_i - a - bX_i - \gamma Z_i)^2$$

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

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

Aynı şekilde katsayımız yine, 0 ile 1 aralığındadır. HKT'nın küçük olması, R²'nin değerini yükseltecek ve daha uygun bir tahmin sağlanacaktır.

R²'nin diğer bir önemli özelliği, açıklayıcı değişkenlerin sayılarının azalmayan bir fonksiyonu olmasıdır. Açıklayıcı değişkenlerin sayısı arttıkça, genellikle R²'de artmaktadır ve hiçbir zaman azalmamaktadır. Bir başka ifade ile yeni eklenen bir açıklayıcı değişken hiçbir zaman R²'nin değerini düşürmemektedir. Bu yüzden çoklu regresyonda, iki tane R²'yi karşılaştırıyorsak, ve iki farklı modelinde açıklayıcı değişkenleri birbirinden farklı ise bir önceki bölümde de gördüğümüz düzeltilmiş -R²'yi kullanmamız gerekmektedir. Düzeltilmiş R²'nin formülünü hatırlayacak olursak:

$$\bar{R} = 1 - \frac{\frac{HKT}{(n-k)}}{\frac{TKT}{(n-1)}} = 1 - \frac{HKT}{TKT} \times \frac{(n-1)}{(n-k)}$$

Çoklu Regresyon Modellerinde İstatistiksel Çıkarım

Çoklu regresyon modellerinde de, bir önceki bölümde deyindiğimiz gibi, regresyon katsayılarının istatistiksel anlamlılığına dair testler gerçekleştiririz. Bu testler b ve c katsayılarının istatistiksel anlamlılığına dair olan testlerdir. (daha öncede olduğu gibi sabit katsayıyı dikkate almıyoruz) Bu anlamlılık testlerini ayrı ayrı gerçekleştirebiliriz.

$$H_0 : \beta = 0$$

$$H_A : \beta \neq 0$$

ve

$$H_0 : \gamma = 0$$

$$H_A : \gamma \neq 0$$

olarak hipotezlerimizi belirtebilir ve bu hipotezleri test edecek olan uygun t-istatistik değerlerini de aşağıdaki gibi belirtebiliriz:

$$t = \frac{b}{s_b}$$

ve

$$t = \frac{c}{s_c}$$

bölümümüzün başında standart hata formüllerine (s_b ve s_c) deyinmiştik. Her zaman olduğu gibi testimizin sonucunu bulmak için hesaplanmış olan t-istatistik değerleri ile tablodan baktığımız t-kritik değerlerini karşılaştırırız. (veya p değeri yaklaşımını da kullanabiliriz).

Bu testlerin sonuçları bize, X ve Z açıklayıcı değişkenlerinin, Y değişkeni üzerinde anlamlı bir etkisinin olup olmadığını göstermektedir.

İlişkinin Anlamlılığının Bütünsel Olarak Test Edilmesi

Denklem (15.3)'te yer alan modeli kullanarak, tahmin edilmiş regresyondaki ilişkinin bütünsel olarak anlamlılığını test etmek istediğimizi düşünelim. Denklemi yeniden yazacak olursak:

$$Y_i = a + bX_i + cZ_i + e_i$$

Her iki açıklayıcı değişkenin, bağımlı değişken üzerindeki birlikte olan ilişkinin anlamlılığını test etmek istemekteyiz. Bu testi gerçekleştirirken, t-testinin yerine Bölüm 12'den de hatırlayacağınız F-testini kullanacağız. Hatırlanacağı üzere F-testi iki varyansın birbirine olan oranı ile hesaplanmaktaydı. Bu test kullanılırken hipotezler ile alakalı olarak bir takım kısıtlamalar yapılacaktır.

Yukarıdaki denklemden, hataların kareleri toplamını elde ettiğimizi ve HKT olarak ifade ettiğimizi varsayalım. HKT, kısıtlanmamış modelden elde edilmiş olan, hataların kareleri toplamını temsil eder. Bu modelin kısıtlanmamış olmasının nedeni, denklem 2'ye herhangi bir kısıtlayıcı parametrenin eklenmemiş olmasındandır. Açıklayıcı değişkenlerin bütünsel anlamlılığını test edebilmek için gerekli olan sıfır ve alternatif hipotezini oluşturalım:

$$H_0 : \beta = \gamma = 0$$

$$H_A : H_0 \text{ 'ın yanlış olması}$$

Dolayısıyla tahmin edilmiş model aşağıdaki gibidir:

$$Y_i = a + e_i \quad (15.9)$$

Bu modele göre tahmin edilmesi gereken sadece sabit terimdir. Bu bağlamda,

$$a = \bar{Y} - b\bar{X} - c\bar{Z}$$

Sıfır hipotezi $\beta = \gamma = 0$ olduğundan:

$$a = \bar{Y}$$

Şimdi ise bu modele ait olan hataların kareleri toplamını, KHKT kısıtlanmış hataların kareleri toplamı olarak ifade edebiliriz ve bu değer:

$$KHKT = \sum_{i=1}^n (Y_i - a)^2 = \sum_{i=1}^n (Y_i - \bar{Y})^2 = TKT$$

Hipotezi test edebilmek için, gerekli olan F testinin istatistik değerini hesaplamaya yönelik formül aşağıdaki gibidir:

$$F = \frac{(KHKT - HKT) \div (df_{kısıtlı} - df_{kısıtsız})}{HKT \div df_{kısıtsız}}$$

Burada, $df_{kısıtsız}$, kısıtlanmamış modelin serbestlik derecesini gösterir ve bu model için $n - 3$ 'e eşittir. $df_{kısıtlı}$, kısıtlanmış modelin serbestlik derecesini gösterir ve bu model için $n - 1$ 'e eşittir. F testinin dağılımında, pay ve paydada olmak üzere iki tane serbestlik derecesine ihtiyacımız olduğundan; $(df_{kısıtlı} - df_{kısıtsız})$ paya ait olan serbestlik derecesini, $df_{kısıtsız}$ 'da paydaya ait olan serbestlik derecesini ifade etmektedir.

Modelimizde, paya ait olan serbestlik derecesi $(n - 1) - (n - 3) = 2$ 'dir. Testimize ait olan kısıtlama sayısı 2'dir ($\beta = 0$ ve $\gamma = 0$). Paydaya ait olan serbestlik derecesi ise $n - 3$ 'e eşittir. Genel ifade ile F-testini aşağıdaki gibi belirtebiliriz:

$$F \sim F(g, n - k)$$

g testimizde yer alan kısıtlama sayısını göstermektedir (kısıtlanmış ile kısıtlanmamış modellerinin serbestlik dereceleri arasındaki fark). k ise kısıtlanmamış modelimizdeki parameter sayısını ifade etmektedir. Şimdi F testimizin formülünü yeniden yazalım:

$$F = \frac{(KHKT - HKT) \div g}{TKT \div (n - k)}$$

Denklemler 15.3'teki modeli tahmin etmek istersek, öncelikle $KHKT = TKT$ eşitliğini kullanarak F testini tekrar ifade edelim:

$$F = \frac{(TKT - HKT) \div g}{TKT \div (n - k)}$$

Eğer payı ve paydayı TKT değerine bölersek:

$$F = \frac{(1 - HKT / TKT) \div g}{1 \div (n - k)}$$

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

$$\frac{HKT}{TKT} = R^2 + 1$$

Eşitliklerini kullanarak, F testi başka bir şekilde de ifade edebiliriz:

$$F = \frac{R^2 \div g}{(1 - R^2) \div (n - k)}$$

Dolayısıyla, açıklayıcı değişkenlerin bütünsel anlamlılığına dair testi R^2 'de kullanarak yapabiliriz. Eğer, F test değeri, F kritik değerinden büyük olursa sıfır hipotezini reddederiz ve açıklayıcı değişkenlerin bütünsel olarak bağımlı değişkeni etkilediği sonucuna varırız. F değeri tahmin edilen modelin anlamlılığı ile alakalı olarak çok faydalı bilgiler edinmemize yardımcı olur. Eğer R^2 düşük çıkarsa, F testinin değeri de küçülecektir dolayısıyla, testin anlamlılığında da bir düşüş gözlenecektir.

Çoklu Regresyon Modelinde Regresyon Katsayılarının Yorumu

Denklem 10'da ifade ettiğimiz regresyon modelini kullanarak, tahmin edilmiş katsayıların nasıl yorumlanması gerektiğini inceleyelim. Q, belli bir mala olan talep miktarını; P fiyat düzeyini gösterirken Y'nin de gelir düzeyini gösterdiğini düşünelim:

$$Q_t = a + bP_t + cY_t + e_t \quad (15.10)$$

Modelimizde yine, e_t hata terimini göstermektedir. Şimdi aklımıza takılan soru, b ve c katsayılarını nasıl yorumlamamız gerektiğidir. b katsayısı, gelir düzeyi sabit iken fiyatların talep miktarı üzerindeki saf etkisini göstermektedir. c katsayısı ise, fiyatlar sabit iken, gelir düzeyinin talep miktarı üzerindeki saf etkisini göstermektedir.

Dikkat edilmesi gereken hususlardan biri de b katsayısının elastikiyet olarak yorumlanmaması gerektiğidir. b'nin yorumu yapılırken, b'nin hem fiyat hem de talep miktarında kullanılan ölçü birimine bağlı olduğunun unutulmamasıdır. b katsayısı, talep denkleminin fiyata göre türevini temsil etmektedir. Bir başka ifade ile:

$$\frac{\partial Q_t}{\partial P_t} = b$$

c katsayısı ise talep denkleminin gelir düzeyine göre türevini temsil etmektedir. Bir başka ifade ile:

$$\frac{\partial Q_t}{\partial Y_t} = c$$

Bu ifadelerden hiçbirisi elastikiyeti temsil etmemektedir. Bu tahmin edicileri, elastikiyet değerini hesaplamak için de kullanabiliriz. Örneğin, talebin fiyat esnekliğini hesaplamak için, aşağıdaki yöntemi kullanırız:

$$\eta = \frac{\partial Q_t}{\partial P_t} \frac{\bar{P}}{\bar{Q}} = b \frac{\bar{P}}{\bar{Q}}$$

Burada $\bar{P}$ ve $\bar{Q}$, ortalama fiyat ve ortalama talep miktarını temsil etmektedir.

Alternatif bir yaklaşımı da, denklem 3'ün doğal logaritmasını alarak (ln'nini alarak) gerçekleştirebiliriz:

$$\ln(Q_t) = a + b \ln(P_t) + c \ln(Y_t) + e_t$$

Doğal logaritmasını aldığımız denklemin, fiyata göre türevini alacak olursak:

$$\frac{\partial \ln(Q_t)}{\partial \ln(P_t)} = b = \frac{\partial Q_t \div Q}{\partial P_t \div P} = \frac{\partial Q_t}{\partial P_t} \frac{P_t}{Q_t}$$

Dolayısıyla tahmin edilmiş olan katsayı burada, aynı zamanda fiyata göre talep esnekliğine eşittir. Doğal logaritmasını aldığımız denklemin, gelir düzeyine göre türevini alacak olursak:

$$\frac{\partial \ln(Q_t)}{\partial \ln(Y_t)} = c = \frac{\partial Q_t \div Q}{\partial Y_t \div Y_t} = \frac{\partial Q_t}{\partial Y_t} \frac{Y_t}{Q_t}$$

Burada dikkat etmemiz gereken husus elastikiyetin ölçüm birimlerinden bağımsız olduğudur.

Minitab Uygulamaları

Şimdi ise çoklu regresyon modelimizdeki tahminleri Minitab ile gerçekleştirelim. Bir pasta fabrikasının genel müdürü olduğumuzu ve Marmara bölgesinde ürünümüzü etkileyen faktörleri araştırdığımızı düşünelim. Modelimiz aşağıdaki gibi olsun:

$$Q_t = \alpha + \beta P_t + \gamma P_t^c + \delta Y_t + \theta A_t + u_t$$

Q_t : *Talep miktarı*. Bu bölgede satılan pasta sayısı (Minitab dosyasında Talep olarak ifade edilmiştir)

P_t : *Bir pastanın fiyatı* (Minitab dosyasında fiyat olarak ifade edilmiştir)

P_t^c : *Rakip firmanın ürettiği pastanın fiyatı* (Minitab dosyasında Rfiyat olarak ifade edilmiştir)

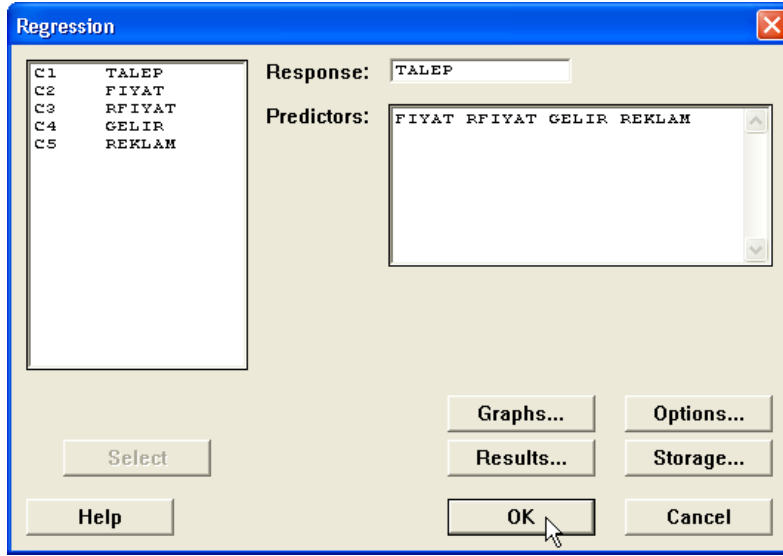
Y_t : *Bölgedeki kişi başına düşen ortalama gelir düzeyi* (Minitab dosyasında Gelir olarak ifade edilmiştir)

A_t : *Reklam Harcamaları* (Minitab dosyasında Reklam olarak ifade edilmiştir)

Örneklem verimiz, “talep2.xls”, “talep2.mtw” ve “talep2.sav” dosyalarında yer almaktadır. Anakütle denkleminizi tahmin etmek için kullanacağımız örneklem modeli aşağıdaki gibidir.

$$Q_t = a + bP_t + cP_t^c + dY_t + fA_t + e_t$$

Tahmin için öncelikle kullanacağımız verilerin bulunduğu çalışma sayfasını açalım. Daha sonra Minitab programında “stat” (istatistik) menüsünden “regression” (regresyon)’u seçelim ve tekrar “regression” (regresyon)’a tıklayalım. Karşımıza çıkacak olan diyalog ekranında, “response” yazan kutucuğa (bu kutucuk bağımlı değişkeni ifade eder “Talep” değişkenini, “predictor” yazan kutucuğa da (bu kutucuk da bağımsız açıklayıcı değişkenlerin yazıldığı yeri ifade eder) “Fiyat, Rfiyat, Gelir, Reklam” açıklayıcı değişkenlerini girelim.



Şekil 15.1

“OK” tuşuna tıkladığımızda aşağıdaki çıktı ekranını elde etmemiz gerekmektedir.

Regression Analysis

The regression equation is

$$\text{TALEP} = -108262 - 0,0296 \text{ FIYAT} + 0,0412 \text{ RFIYAT} + 0,00863 \text{ GELIR} + 0,225 \text{ REKLAM}$$

Predictor	Coef	StDev	T	P
Constant	-108262	73227	-1,48	0,148
FIYAT	-0,02962	0,01010	-2,93	0,006
RFIYAT	0,041217	0,009821	4,20	0,000
GELIR	0,0086253	0,0008138	10,60	0,000
REKLAM	0,2247	0,1118	2,01	0,052

S = 17955 R-Sq = 89,8% R-Sq(adj) = 88,6%

Analysis of Variance

Source	DF	SS	MS	F	P
Regression	4	98915469291	24728867323	76,70	0,000
Residual Error	35	11283905709	322397306		
Total	39	1,10199E+11			

R denotes an observation with a large standardized residual

Şekil 15.2

Regresyon katsayılarını aşağıdaki gibi yorumlayabiliriz:

- $a = -108262$. Regresyon fonksiyonunun sabit terimidir her durum için yorumlama yapılamaz. Bu örnekte ise, istatistiksel olarak anlamlı olmadığından dolayı herhangi bir yorum getirilemez.
- $b \approx -0,03$. Fiyat'ın katsayısı olarak, pasta fiyatlarında bir birimlik bir artış olduğunda, talep miktarının yaklaşık olarak 0,03 birim düşeceğini göstermektedir.

- $c \approx 0,04$. Eğer rakip firma, fiyat düzeyini bir birim arttırırsa, Marmara Bölgesi'nde pastalarımıza olan talep miktarı 0,04 birim artacaktır.
- $d \approx 0,0086$. Bölgedeki kişi başına düşen ortalama gelir düzeyi bir birim artarsa, insanların 0,0086 birimlik bir ekstra pasta talebi olacaktır.
- $f \approx 0,225$. Reklam için yapılan her ekstra birimlik harcama, satış miktarında 0,225 birimlik bir artışa yol açacaktır.

Bu katsayıların istatistiksel anlamlılığını test etmek için ister t-istatistiklerini isterseniz de p-değerlerini kullanabilirsiniz. Örneğin, fiyat değişkeninin modelimiz için anlamlı olup olmadığını test edelim:

$$H_0 : \beta = 0$$

$$H_A : \beta \neq 0$$

Her zaman olduğu gibi yine t-istatistik değerlerini veya p-değerlerini kullanacağız. Minitab çıktısından, fiyat değişkeninin anlamlı bir etkisi olmadığı sıfır hipotezini açık bir şekilde reddetmemiz gerektiğini anlamaktayız.

Regression Analysis

The regression equation is
TALEP = - 108262 - 0,0296 FİYAT + 0,0412 RFIYAT + 0,00863 GELİR + 0,225 REKLAM

Predictor	Coef	StDev	T	P
Constant	-108262	73227	-1,48	0,148
FİYAT	-0,02962	0,01010	-2,93	0,006
RFİYAT	0,041217	0,009821	4,20	0,000
GELİR	0,0086253	0,0008138	10,60	0,000
REKLAM	0,2247	0,1118	2,01	0,052

S = 17955

R-Sq = 89,8%

R-Sq(adj) = 88,6%

Analysis of Variance

Source	DF	SS	MS	F	P
Regression	4	98915469291	24728867323	76,70	0,000
Residual Error	35	11283905709	322397306		
Total	39	1,10199E+11			

R denotes an observation with a large standardized residual

Şekil 15.3

Benzer bir şekilde, sabit terimin haricinde bütün katsayıların yüzde 10'luk anlamlılık düzeyinde, anlamlı olduğunu söyleyebiliriz. Yüzde beşlik anlamlılık düzeyinde ise sabit terim ve reklam katsayısı haricinde diğer katsayıların tamamı anlamlıdır. R^2 ve $\bar{R}^2$ 'nin her ikisinde modelimizin açıklayıcılığının yüksek olduğuna işaret etmektedir.

Son olarak, katsayıların bütünsel anlamlılığını F test kullanarak incelemeye çalışalım. Hipotezlerimiz aşağıdaki olmalıdır:

$$H_0 : \beta = \gamma = \delta = \theta = 0$$

$$H_A : H_0 \text{ 'in yanlış olması}$$

Minitab çıktısında “*Analysis of Variance*” (Varyans Analizi) başlığının altında F istatistik değerini ve p-değerini bulabiliriz.

Regression Analysis

The regression equation is

TALEP = - 108262 - 0,0296 FIYAT + 0,0412 RFIYAT + 0,00863 GELIR + 0,225 REKLAM

Predictor	Coef	StDev	T	P
Constant	-108262	73227	-1,48	0,148
FIYAT	-0,02962	0,01010	-2,93	0,006
RFIYAT	0,041217	0,009821	4,20	0,000
GELIR	0,0086253	0,0008138	10,60	0,000
REKLAM	0,2247	0,1118	2,01	0,052

S = 17955 R-Sq = 89,8% R-Sq(adj) = 88,6%

Analysis of Variance

Source	DF	SS	MS	F	P
Regression	4	98915469291	24728867323	76,70	0,000
Residual Error	35	11283905709	322397306		
Total	39	1,10199E+11			

R denotes an observation with a large standardized residual

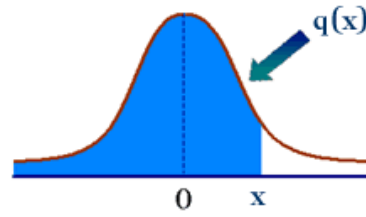
Şekil 15.4

Çıktıdan da görüldüğü gibi, p değeri çok küçük olduğundan veya F-istatistik değeri çok büyük olduğundan dolayı sıfır hipotezimizi reddederiz ve regresyon katsayılarının bütünsel olarak anlamlı olduğu sonucuna varırız.

EK

İSTATİSTİK TABLOLARI

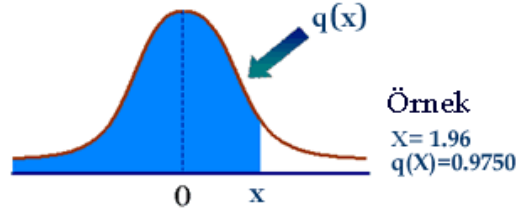
TABLO 1 NORMAL DAĞILIM FONKSİYONU



Örnek
 $X = 1.96$
 $q(X) = 0.9750$

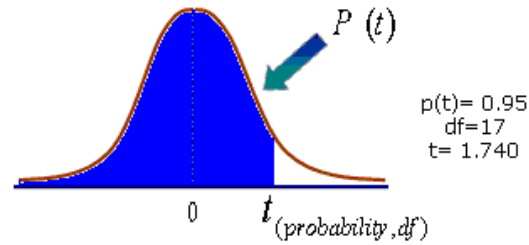
x	$\Phi(x)$	x	$\Phi(x)$	x	$\Phi(x)$	x	$\Phi(x)$	x	$\Phi(x)$	x	$\Phi(x)$
0,00	0,5000	0,40	0,6554	0,80	0,7881	1,20	0,8849	1,60	0,9452	2,00	0,97725
0,01	0,5040	0,41	0,6591	0,81	0,7910	1,21	0,8869	1,61	0,9463	2,01	0,97778
0,02	0,5080	0,42	0,6628	0,82	0,7939	1,22	0,8888	1,62	0,9474	2,02	0,97831
0,03	0,5120	0,43	0,6664	0,83	0,7967	1,23	0,8907	1,63	0,9484	2,03	0,97882
0,04	0,5160	0,44	0,6700	0,84	0,7995	1,24	0,8925	1,64	0,9495	2,04	0,97932
0,05	0,5199	0,45	0,6736	0,85	0,8023	1,25	0,8944	1,65	0,9505	2,05	0,97982
0,06	0,5239	0,46	0,6772	0,86	0,8051	1,26	0,8962	1,66	0,9515	2,06	0,98030
0,07	0,5279	0,47	0,6808	0,87	0,8078	1,27	0,8980	1,67	0,9525	2,07	0,98077
0,08	0,5319	0,48	0,6844	0,88	0,8106	1,28	0,8997	1,68	0,9535	2,08	0,98124
0,09	0,5359	0,49	0,6879	0,89	0,8133	1,29	0,9015	1,69	0,9545	2,09	0,98169
0,10	0,5398	0,50	0,6915	0,90	0,8159	1,30	0,9032	1,70	0,9554	2,10	0,98214
0,11	0,5438	0,51	0,6950	0,91	0,8186	1,31	0,9049	1,71	0,9564	2,11	0,98257
0,12	0,5478	0,52	0,6985	0,92	0,8212	1,32	0,9066	1,72	0,9573	2,12	0,98300
0,13	0,5517	0,53	0,7019	0,93	0,8238	1,33	0,9082	1,73	0,9582	2,13	0,98341
0,14	0,5557	0,54	0,7054	0,94	0,8264	1,34	0,9099	1,74	0,9591	2,14	0,98382
0,15	0,5596	0,55	0,7088	0,95	0,8289	1,35	0,9115	1,75	0,9599	2,15	0,98422
0,16	0,5636	0,56	0,7123	0,96	0,8315	1,36	0,9131	1,76	0,9608	2,16	0,98461
0,17	0,5675	0,57	0,7157	0,97	0,8340	1,37	0,9147	1,77	0,9616	2,17	0,98500
0,18	0,5714	0,58	0,7190	0,98	0,8365	1,38	0,9162	1,78	0,9625	2,18	0,98537
0,19	0,5753	0,59	0,7224	0,99	0,8389	1,39	0,9177	1,79	0,9633	2,19	0,98574
0,20	0,5793	0,60	0,7257	1,00	0,8413	1,40	0,9192	1,80	0,9641	2,20	0,98610
0,21	0,5832	0,61	0,7291	1,01	0,8438	1,41	0,9207	1,81	0,9649	2,21	0,98645
0,22	0,5871	0,62	0,7324	1,02	0,8461	1,42	0,9222	1,82	0,9656	2,22	0,98679
0,23	0,5910	0,63	0,7357	1,03	0,8485	1,43	0,9236	1,83	0,9664	2,23	0,98713
0,24	0,5948	0,64	0,7389	1,04	0,8508	1,44	0,9251	1,84	0,9671	2,24	0,98745
0,25	0,5987	0,65	0,7422	1,05	0,8531	1,45	0,9265	1,85	0,9678	2,25	0,98778
0,26	0,6026	0,66	0,7454	1,06	0,8554	1,46	0,9279	1,86	0,9686	2,26	0,98809
0,27	0,6064	0,67	0,7486	1,07	0,8577	1,47	0,9292	1,87	0,9693	2,27	0,98840
0,28	0,6103	0,68	0,7517	1,08	0,8599	1,48	0,9306	1,88	0,9699	2,28	0,98870
0,29	0,6141	0,69	0,7549	1,09	0,8621	1,49	0,9319	1,89	0,9706	2,29	0,98899

TABLO 1 NORMAL DAĞILIM FONKSİYONU (Devam)



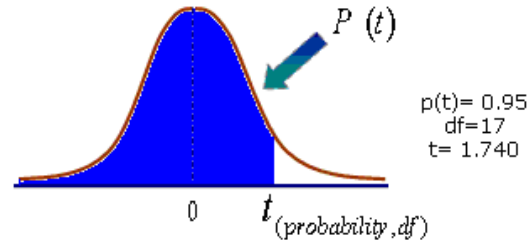
0,30	0,6179	0,70	0,7580	1,10	0,8643	1,50	0,9332	1,90	0,9713	2,30	0,98928
0,31	0,6217	0,71	0,7611	1,11	0,8665	1,51	0,9345	1,91	0,9719	2,31	0,98956
0,32	0,6255	0,72	0,7642	1,12	0,8686	1,52	0,9357	1,92	0,9726	2,32	0,98983
0,33	0,6293	0,73	0,7673	1,13	0,8708	1,53	0,9370	1,93	0,9732	2,33	0,99010
0,34	0,6331	0,74	0,7704	1,14	0,8729	1,54	0,9382	1,94	0,9738	2,34	0,99036
0,35	0,6368	0,75	0,7734	1,15	0,8749	1,55	0,9394	1,95	0,9744	2,35	0,99061
0,36	0,6406	0,76	0,7764	1,16	0,8770	1,56	0,9406	1,96	0,9750	2,36	0,99086
0,37	0,6443	0,77	0,7794	1,17	0,8790	1,57	0,9418	1,97	0,9756	2,37	0,99111
0,38	0,6480	0,78	0,7823	1,18	0,8810	1,58	0,9429	1,98	0,9761	2,38	0,99134
0,39	0,6517	0,79	0,7852	1,19	0,8830	1,59	0,9441	1,99	0,9767	2,39	0,99158
0,40	0,6554	0,80	0,7881	1,20	0,8849	1,60	0,9452	2,00	0,9772	2,40	0,99180
x	Φ(x)	x	Φ(x)	x	Φ(x)	x	Φ(x)	x	Φ(x)	x	Φ(x)
2,40	0,99180	2,55	0,99461	2,70	0,99653	2,85	0,99781	3,00	0,99865	3,15	0,99918
2,41	0,99202	2,56	0,99477	2,71	0,99664	2,86	0,99788	3,01	0,99869	3,16	0,99921
2,42	0,99224	2,57	0,99492	2,72	0,99674	2,87	0,99795	3,02	0,99874	3,17	0,99924
2,43	0,99245	2,58	0,99506	2,73	0,99683	2,88	0,99801	3,03	0,99878	3,18	0,99926
2,44	0,99266	2,59	0,99520	2,74	0,99693	2,89	0,99807	3,04	0,99882	3,19	0,99929
2,45	0,99286	2,6	0,99534	2,75	0,99702	2,90	0,99813	3,05	0,99886	3,20	0,99931
2,46	0,99305	2,61	0,99547	2,76	0,99711	2,91	0,99819	3,06	0,99889	3,21	0,99934
2,47	0,99324	2,62	0,99560	2,77	0,99720	2,92	0,99825	3,07	0,99893	3,22	0,99936
2,48	0,99343	2,63	0,99573	2,78	0,99728	2,93	0,99831	3,08	0,99896	3,23	0,99938
2,49	0,99361	2,64	0,99585	2,79	0,99736	2,94	0,99836	3,09	0,99900	3,24	0,99940
2,50	0,99379	2,65	0,99598	2,80	0,99744	2,95	0,99841	3,10	0,99903	3,25	0,99942
2,51	0,99396	2,66	0,99609	2,81	0,99752	2,96	0,99846	3,11	0,99906	3,26	0,99944
2,52	0,99413	2,67	0,99621	2,82	0,99760	2,97	0,99851	3,12	0,99910	3,27	0,99946
2,53	0,99430	2,68	0,99632	2,83	0,99767	2,98	0,99856	3,13	0,99913	3,28	0,99948
2,54	0,99446	2,69	0,99643	2,84	0,99774	2,99	0,99861	3,14	0,99916	3,29	0,99950
2,55	0,99461	2,70	0,99653	2,85	0,99781	3,00	0,99865	3,15	0,99918	3,30	0,99952

TABLO 2 t DAĞILIMI FONKSİYONU



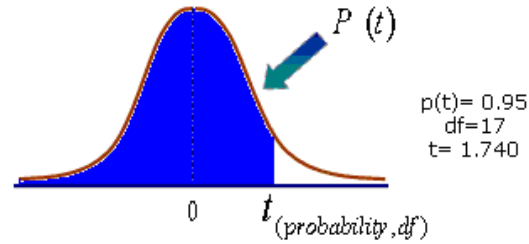
$\nu =$	1	$\nu =$	1	$\nu =$	2	$\nu =$	2	$\nu =$	3	$\nu =$	3
$t = 0.0$	0.5000	4.0	0.9220	$t = 0.0$	0.5000	4.0	0.9714	$t = 0.0$	0.5000	4.0	0.9860
.1	.5317	4.2	.9256	.1	.5353	4.1	.9727	.1	.5367	4.1	.9869
.2	.5628	4.4	.9289	.2	.5700	4.2	.9739	.2	.5729	4.2	.9877
.3	.5928	4.6	.9319	.3	.6038	4.3	.9750	.3	.6081	4.3	.9884
.4	.6211	4.8	.9346	.4	.6361	4.4	.9760	.4	.6420	4.4	.9891
.5	.6476	5.0	.9372	.5	.6667	4.5	.9770	.5	.6743	4.5	.9898
.6	.6720	5.5	.9428	.6	.6953	4.6	.9779	.6	.7046	4.6	.9903
.7	.6944	6.0	.9474	.7	.7218	4.7	.9788	.7	.7328	4.7	.9909
.8	.7148	6.5	.9514	.8	.7462	4.8	.9796	.8	.7589	4.8	.9914
.9	.7333	7.0	.9548	.9	.7684	4.9	.9804	.9	.7828	4.9	.9919
1.0	.7500	7.5	.9578	1.0	.7887	5.0	.9811	1.0	.8045	5.0	.9923
.1	.7651	8.0	.9604	.1	.8070	5.1	.9818	.1	.8242	5.1	.9927
.2	.7789	8.5	.9627	.2	.8235	5.2	.9825	.2	.8419	5.2	.9931
.3	.7913	9.0	.9648	.3	.8384	5.3	.9831	.3	.8578	5.3	.9934
.4	.8026	9.5	.9666	.4	.8518	5.4	.9837	.4	.8720	5.4	.9938
.5	.8128	10.0	.9683	.5	.8638	5.5	.9842	.5	.8847	5.5	.9941
.6	.8222	10.5	.9698	.6	.8746	5.6	.9848	.6	.8960	5.6	.9944
.7	.8307	11.0	.9711	.7	.8844	5.7	.9853	.7	.9062	5.7	.9946
.8	.8386	11.5	.9724	.8	.8932	5.8	.9858	.8	.9152	5.8	.9949
.9	.8458	12.0	.9735	.9	.9011	5.9	.9862	.9	.9232	5.9	.9951
2.0	.8524	12.5	.9746	2.0	.9082	6.0	.9867	2.0	.9303	6.0	.9954
.1	.8585	13.0	.9756	.1	.9147	6.1	.9871	.1	.9367	6.1	.9956
.2	.8642	13.5	.9765	.2	.9206	6.2	.9875	.2	.9424	6.2	.9958
.3	.8695	14.0	.9773	.3	.9259	6.3	.9879	.3	.9475	6.3	.9960
.4	.8743	14.5	.9781	.4	.9308	6.4	.9882	.4	.9521	6.4	.9961
.5	.8789	15	.9788	.5	.9352	6.5	.9886	.5	.9561	6.5	.9963
.6	.8831	16	.9801	.6	.9392	6.6	.9889	.6	.9598	6.6	.9965
.7	.8871	17	.9813	.7	.9429	6.7	.9892	.7	.9631	6.7	.9966
.8	.8908	18	.9823	.8	.9463	6.8	.9895	.8	.9661	6.8	.9967
.9	.8943	19	.9833	.9	.9494	6.9	.9898	.9	.9687	6.9	.9969
3.0	.8976	20	.9841	3.0	.9523	7.0	.9901	3.0	.9712	7.0	.9970
.1	.9007	21	.9849	.1	.9549	7.1	.9904	.1	.9734	7.1	.9971
.2	.9036	22	.9855	.2	.9573	7.2	.9906	.2	.9753	7.2	.9972
.3	.9063	23	.9862	.3	.9596	7.3	.9909	.3	.9771	7.3	.9973
.4	.9089	24	.9867	.4	.9617	7.4	.9911	.4	.9788	7.4	.9974
.5	.9114	25	.9873	.5	.9636	7.5	.9913	.5	.9803	7.5	.9975
.6	.9138	30	.9894	.6	.9654	7.6	.9916	.6	.9816	7.6	.9976
.7	.9160	35	.9909	.7	.9670	7.7	.9918	.7	.9829	7.7	.9977
.8	.9181	40	.9920	.8	.9686	7.8	.9920	.8	.9840	7.8	.9978
.9	.9201	45	.9929	.9	.9701	7.9	.9922	.9	.9850	7.9	.9979
4.0	.9220	50	.9936	4.0	.9714	8.0	.9924	4.0	.9860	8.0	.9980

TABLO 2 t DAĞILIM FONKSİYONU (Devam)



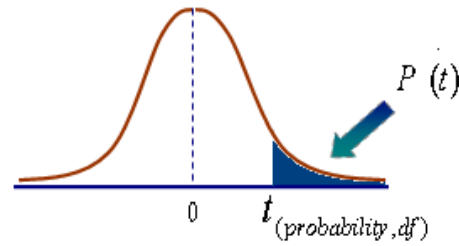
$\nu =$	4	5	6	7	8	9	10	11	12	13	14
$t = 0.0$	0.5000	0.5000	0.5000	0.5000	0.5000	0.5000	0.5000	0.5000	0.5000	0.5000	0.5000
.1	.5374	.5379	.5382	.5384	.5386	.5387	.5388	.5389	.5390	.5391	.5391
.2	.5744	.5753	.5760	.5764	.5768	.5770	.5773	.5774	.5776	.5777	.5778
.3	.6104	.6119	.6129	.6136	.6141	.6145	.6148	.6151	.6153	.6155	.6157
.4	.6452	.6472	.6485	.6495	.6502	.6508	.6512	.6516	.6519	.6522	.6524
.5	.6783	.6809	.6826	.6838	.6847	.6855	.6861	.6865	.6869	.6873	.6876
.6	.7096	.7127	.7148	.7163	.7174	.7183	.7191	.7197	.7202	.7206	.7210
.7	.7387	.7424	.7449	.7467	.7481	.7492	.7501	.7508	.7514	.7519	.7523
.8	.7657	.7700	.7729	.7750	.7766	.7778	.7788	.7797	.7804	.7810	.7815
.9	.7905	.7953	.7986	.8010	.8028	.8042	.8054	.8063	.8071	.8078	.8083
1.0	.8130	.8184	.8220	.8247	.8267	.8283	.8296	.8306	.8315	.8322	.8329
.1	.8335	.8393	.8433	.8461	.8483	.8501	.8514	.8526	.8535	.8544	.8551
.2	.8518	.8581	.8623	.8654	.8678	.8696	.8711	.8723	.8734	.8742	.8750
.3	.8683	.8748	.8793	.8826	.8851	.8870	.8886	.8899	.8910	.8919	.8927
.4	.8829	.8898	.8945	.8979	.9005	.9025	.9041	.9055	.9066	.9075	.9084
.5	.8960	.9030	.9079	.9114	.9140	.9161	.9177	.9191	.9203	.9212	.9221
.6	.9076	.9148	.9196	.9232	.9259	.9280	.9297	.9310	.9322	.9332	.9340
.7	.9178	.9251	.9300	.9335	.9362	.9383	.9400	.9414	.9426	.9435	.9444
.8	.9269	.9341	.9390	.9426	.9452	.9473	.9490	.9503	.9515	.9525	.9533
.9	.9349	.9421	.9469	.9504	.9530	.9551	.9567	.9580	.9591	.9601	.9609
2.0	.9419	.9490	.9538	.9572	.9597	.9617	.9633	.9646	.9657	.9666	.9674
.1	.9482	.9551	.9598	.9631	.9655	.9674	.9690	.9702	.9712	.9721	.9728
.2	.9537	.9605	.9649	.9681	.9705	.9723	.9738	.9750	.9759	.9768	.9774
.3	.9585	.9651	.9694	.9725	.9748	.9765	.9779	.9790	.9799	.9807	.9813
.4	.9628	.9692	.9734	.9763	.9784	.9801	.9813	.9824	.9832	.9840	.9846
.5	.9666	.9728	.9767	.9795	.9815	.9831	.9843	.9852	.9860	.9867	.9873
.6	.9700	.9759	.9797	.9823	.9842	.9856	.9868	.9877	.9884	.9890	.9895
.7	.9730	.9786	.9822	.9847	.9865	.9878	.9888	.9897	.9903	.9909	.9914
.8	.9756	.9810	.9844	.9867	.9884	.9896	.9906	.9914	.9920	.9925	.9929
.9	.9779	.9831	.9863	.9885	.9901	.9912	.9921	.9928	.9933	.9938	.9942
3.0	.9800	.9850	.9880	.9900	.9915	.9925	.9933	.9940	.9945	.9949	.9952
.1	.9819	.9866	.9894	.9913	.9927	.9936	.9944	.9949	.9954	.9958	.9961
.2	.9835	.9880	.9907	.9925	.9937	.9946	.9953	.9958	.9962	.9965	.9968
.3	.9850	.9893	.9918	.9934	.9946	.9954	.9960	.9965	.9968	.9971	.9974
.4	.9864	.9904	.9928	.9943	.9953	.9961	.9966	.9970	.9974	.9976	.9978
.5	.9876	.9914	.9936	.9950	.9960	.9966	.9971	.9975	.9978	.9980	.9982
.6	.9886	.9922	.9943	.9956	.9965	.9971	.9976	.9979	.9982	.9984	.9986
.7	.9896	.9930	.9950	.9962	.9970	.9975	.9979	.9982	.9985	.9987	.9988
.8	.9904	.9937	.9955	.9966	.9974	.9979	.9983	.9985	.9987	.9989	.9990
.9	.9912	.9943	.9960	.9971	.9977	.9982	.9985	.9988	.9989	.9991	.9992
4.0	.9919	.9948	.9964	.9974	.9980	.9984	.9987	.9990	.9991	.9992	.9993
.1	.9926	.9953	.9968	.9977	.9983	.9987	.9989	.9991	.9993	.9994	.9995
.2	.9932	.9958	.9972	.9980	.9985	.9988	.9991	.9993	.9994	.9995	.9996
.3	.9937	.9961	.9975	.9982	.9987	.9990	.9992	.9994	.9995	.9996	.9996
.4	.9942	.9965	.9977	.9984	.9989	.9991	.9993	.9995	.9996	.9996	.9997
.5	.9946	.9968	.9979	.9986	.9990	.9993	.9994	.9995	.9996	.9997	.9998
.6	.9950	.9971	.9982	.9988	.9991	.9994	.9995	.9996	.9997	.9998	.9998
.7	.9953	.9973	.9983	.9989	.9992	.9994	.9996	.9997	.9997	.9998	.9998
.8	.9957	.9976	.9985	.9990	.9993	.9995	.9996	.9997	.9998	.9998	.9999
.9	.9960	.9978	.9986	.9991	.9994	.9996	.9997	.9998	.9998	.9999	.9999
5.0	.9963	.9979	.9988	.9992	.9995	.9996	.9997	.9998	.9998	.9999	.9999

TABLO 2 t DAĞILIM FONKSİYONU (Devam)



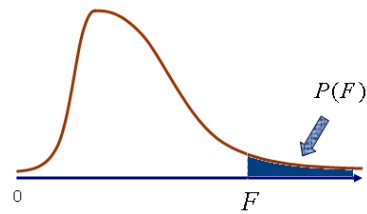
$\nu =$	15	16	17	18	19	20	24	30	40	60	
$t = 0.0$	0.5000	0.5000	0.5000	0.5000	0.5000	0.5000	0.5000	0.5000	0.5000	0.5000	0.5000
.1	.5392	.5392	.5392	.5393	.5393	.5393	.5394	.5395	.5396	.5397	.5398
.2	.5779	.5780	.5781	.5781	.5782	.5782	.5784	.5786	.5788	.5789	.5793
.3	.6159	.6160	.6161	.6162	.6163	.6164	.6166	.6169	.6171	.6174	.6179
.4	.6526	.6528	.6529	.6531	.6532	.6533	.6537	.6540	.6544	.6547	.6554
.5	.6878	.6881	.6883	.6884	.6886	.6887	.6892	.6896	.6901	.6905	.6915
.6	.7213	.7215	.7218	.7220	.7222	.7224	.7229	.7235	.7241	.7246	.7257
.7	.7527	.7530	.7533	.7536	.7538	.7540	.7547	.7553	.7560	.7567	.7580
.8	.7819	.7823	.7826	.7829	.7832	.7834	.7842	.7850	.7858	.7866	.7881
.9	.8088	.8093	.8097	.8100	.8103	.8106	.8115	.8124	.8132	.8141	.8159
1.0	.8334	.8339	.8343	.8347	.8351	.8354	.8364	.8373	.8383	.8393	.8413
.1	.8557	.8562	.8567	.8571	.8575	.8578	.8589	.8600	.8610	.8621	.8643
.2	.9756	.8762	.8767	.8772	.8776	.8779	.8791	.8802	.8814	.8826	.8849
.3	.8934	.8940	.8945	.8950	.8954	.8958	.8970	.8982	.8995	.9007	.9032
.4	.9091	.9097	.9103	.9107	.9112	.9116	.9128	.9141	.9154	.9167	.9192
.5	.9228	.9235	.9240	.9245	.9250	.9254	.9267	.9280	.9293	.9306	.9332
.6	.9348	.9354	.9360	.9365	.9370	.9374	.9387	.9400	.9413	.9426	.9452
.7	.9451	.9458	.9463	.9468	.9473	.9477	.9490	.9503	.9516	.9528	.9554
.8	.9540	.9546	.9552	.9557	.9561	.9565	.9578	.9590	.9603	.9616	.9641
.9	.9616	.9622	.9627	.9632	.9636	.9640	.9652	.9665	.9677	.9689	.9713
2.0	.9680	.9686	.9691	.9696	.9700	.9704	.9715	.9727	.9738	.9750	.9772
.1	.9735	.9740	.9745	.9750	.9753	.9757	.9768	.9779	.9790	.9800	.9821
.2	.9781	.9786	.9790	.9794	.9798	.9801	.9812	.9822	.9832	.9842	.9861
.3	.9819	.9824	.9828	.9832	.9835	.9838	.9848	.9857	.9866	.9875	.9893
.4	.9851	.9855	.9859	.9863	.9866	.9869	.9877	.9886	.9894	.9902	.9918
.5	.9877	.9882	.9885	.9888	.9891	.9894	.9902	.9909	.9917	.9924	.9938
.6	.9900	.9903	.9907	.9910	.9912	.9914	.9921	.9928	.9935	.9941	.9953
.7	.9918	.9921	.9924	.9927	.9929	.9931	.9937	.9944	.9949	.9955	.9965
.8	.9933	.9936	.9938	.9941	.9943	.9945	.9950	.9956	.9961	.9966	.9974
.9	.9945	.9948	.9950	.9952	.9954	.9956	.9961	.9965	.9970	.9974	.9981
3.0	.9955	.9958	.9960	.9962	.9963	.9965	.9969	.9973	.9977	.9980	.9987
.1	.9963	.9966	.9967	.9969	.9971	.9972	.9976	.9979	.9982	.9985	.9990
.2	.9970	.9972	.9974	.9975	.9976	.9978	.9981	.9984	.9987	.9989	.9993
.3	.9976	.9977	.9979	.9980	.9981	.9982	.9985	.9988	.9990	.9992	.9995
.4	.9980	.9982	.9983	.9984	.9985	.9986	.9988	.9990	.9992	.9994	.9997
.5	.9984	.9985	.9986	.9987	.9988	.9989	.9991	.9993	.9994	.9996	.9998
.6	.9987	.9988	.9989	.9990	.9990	.9991	.9993	.9994	.9996	.9997	.9998
.7	.9989	.9990	.9991	.9992	.9992	.9993	.9994	.9996	.9997	.9998	.9999
.8	.9991	.9992	.9993	.9993	.9994	.9994	.9996	.9997	.9998	.9998	.9999
.9	.9993	.9994	.9994	.9995	.9995	.9996	.9997	.9997	.9998	.9999	
4.0	.9994	.9995	.9995	.9996	.9996	.9996	.9997	.9998	.9999	.9999	
.1	.9995	.9996	.9996	.9997	.9997	.9997	.9998	.9999	.9999	.9999	
.2	.9996	.9997	.9997	.9997	.9998	.9998	.9998	.9999	.9999	.9999	
.3	.9997	.9997	.9998	.9998	.9998	.9998	.9999	.9999	.9999	.9999	
.4	.9997	.9998	.9998	.9998	.9998	.9999	.9999	.9999	.9999	.9999	
.5	.9998	.9998	.9998	.9999	.9999	.9999	.9999	.9999	.9999	.9999	

TABLO 3 t DAĞILIM FONKSİYONUN YÜZDELİK NOKTALARA GÖRE DAĞILIMI



p = (%)	40	30	25	20	15	10	5	2.5	1	0.5	0.1	0.05
v = 1	0.3249	0.7265	1.0000	1.3764	1.963	3.078	6.314	12.71	31.82	63.66	318.3	636.6
2	.2887	.6172	.8165	1.0607	1.386	1.886	2.920	4.303	6.965	9.925	22.33	31.60
3	.2767	.5844	.7649	.9785	1.250	1.638	2.353	3.182	4.541	5.841	10.21	12.92
4	.2707	.5686	.7407	.9410	1.190	1.533	2.132	2.776	3.747	4.604	7.173	8.610
5	.2672	.5594	.7267	.9195	1.156	1.476	2.015	2.571	3.365	4.032	5.893	6.869
6	.2648	.5534	.7176	.9057	1.134	1.440	1.943	2.447	3.143	3.707	5.208	5.959
7	.2632	.5491	.7111	.8960	1.119	1.415	1.895	2.365	2.998	3.499	4.785	5.408
8	.2619	.5459	.7064	.8889	1.108	1.397	1.860	2.306	2.896	3.355	4.501	5.041
9	.2610	.5435	.7027	.8834	1.100	1.383	1.833	2.262	2.821	3.250	4.297	4.781
10	.2602	.5415	.6998	.8791	1.093	1.372	1.812	2.228	2.764	3.169	4.144	4.587
11	.2596	.5399	.6974	.8755	1.088	1.363	1.796	2.201	2.718	3.106	4.025	4.437
12	.2590	.5386	.6955	.8726	1.083	1.356	1.782	2.179	2.681	3.055	3.930	4.318
13	.2586	.5375	.6938	.8702	1.079	1.350	1.771	2.160	2.650	3.012	3.852	4.221
14	.2582	.5366	.6924	.8681	1.076	1.345	1.761	2.145	2.624	2.977	3.787	4.140
15	.2579	.5357	.6912	.8662	1.074	1.341	1.753	2.131	2.602	2.947	3.733	4.073
16	.2576	.5350	.6901	.8647	1.071	1.337	1.746	2.120	2.583	2.921	3.686	4.015
17	.2573	.5344	.6892	.8633	1.069	1.333	1.740	2.110	2.567	2.898	3.646	3.965
18	.2571	.5338	.6884	.8620	1.067	1.330	1.734	2.101	2.552	2.878	3.610	3.922
19	.2569	.5333	.6876	.8610	1.066	1.328	1.729	2.093	2.539	2.861	3.579	3.883
20	.2567	.5329	.6870	.8600	1.064	1.325	1.725	2.086	2.528	2.845	3.552	3.850
21	.2566	.5325	.6864	.8591	1.063	1.323	1.721	2.080	2.518	2.831	3.527	3.819
22	.2564	.5321	.6858	.8583	1.061	1.321	1.717	2.074	2.508	2.819	3.505	3.792
23	.2563	.5317	.6853	.8575	1.060	1.319	1.714	2.069	2.500	2.807	3.485	3.768
24	.2562	.5314	.6848	.8569	1.059	1.318	1.711	2.064	2.492	2.797	3.467	3.745
25	.2561	.5312	.6844	.8562	1.058	1.316	1.708	2.060	2.485	2.787	3.450	3.725
26	.2560	.5309	.6840	.8557	1.058	1.315	1.706	2.056	2.479	2.779	3.435	3.707
27	.2559	.5306	.6837	.8551	1.057	1.314	1.703	2.052	2.473	2.771	3.421	3.690
28	.2558	.5304	.6834	.8546	1.056	1.313	1.701	2.048	2.467	2.763	3.408	3.674
29	.2557	.5302	.6830	.8542	1.055	1.311	1.699	2.045	2.462	2.756	3.396	3.659
30	.2556	.5300	.6828	.8538	1.055	1.310	1.697	2.042	2.457	2.750	3.385	3.646
32	.2555	.5297	.6822	.8530	1.054	1.309	1.694	2.037	2.449	2.738	3.365	3.622
34	.2553	.5294	.6818	.8523	1.052	1.307	1.691	2.032	2.441	2.728	3.348	3.601
36	.2552	.5291	.6814	.8517	1.052	1.306	1.688	2.028	2.434	2.719	3.333	3.582
38	.2551	.5288	.6810	.8512	1.051	1.304	1.686	2.024	2.429	2.712	3.319	3.566
40	.2550	.5286	.6807	.8507	1.050	1.303	1.684	2.021	2.423	2.704	3.307	3.551
50	.2547	.5278	.6794	.8489	1.047	1.299	1.676	2.009	2.403	2.678	3.261	3.496
60	.2545	.5272	.6786	.8477	1.045	1.296	1.671	2.000	2.390	2.660	3.232	3.460
120	.2539	.5258	.6765	.8446	1.041	1.289	1.658	1.980	2.358	2.617	3.160	3.373
	.2533	.5244	.6745	.8416	1.036	1.282	1.645	1.960	2.326	2.576	3.090	3.291

TABLO 4(a) %10'na GÖRE F DAĞILIMI



v1 =	1	2	3	4	5	6	7	8	10	12	24	
v2 = 1	39.86	49.50	53.59	55.83	57.24	58.20	58.91	59.44	60.19	60.71	62.00	63.33
2	8.526	9.000	9.162	9.243	9.293	9.326	9.349	9.367	9.392	9.408	9.450	9.491
3	5.538	5.462	5.391	5.343	5.309	5.285	5.266	5.252	5.230	5.216	5.176	5.134
4	4.545	4.325	4.191	4.107	4.051	4.010	3.979	3.955	3.920	3.896	3.831	3.761
5	85	3.780	3.619	3.520	3.453	3.405	3.368	3.339	3.297	3.268	3.191	3.105
6	3.776	3.463	3.289	3.181	3.108	3.055	3.014	2.983	2.937	2.905	2.818	2.722
7	3.589	3.257	3.074	2.961	2.883	2.827	2.785	2.752	2.703	2.668	2.575	2.471
8	3.458	3.113	2.924	2.806	2.726	2.668	2.624	2.589	2.538	2.502	2.404	2.923
9	3.360	3.006	2.813	2.693	2.611	2.551	2.505	2.469	2.416	2.379	2.277	2.159
10	3.285	2.924	2.728	2.605	2.522	2.461	2.414	2.377	2.323	2.284	2.178	2.055
11	3.225	2.860	2.660	2.536	2.451	2.389	2.342	2.304	2.248	2.209	2.100	1.972
12	3.177	2.807	2.606	2.480	2.394	2.331	2.283	2.245	2.188	2.147	2.036	1.904
13	3.136	2.763	2.560	2.434	2.347	2.283	2.234	2.195	2.138	2.097	1.983	1.846
14	3.102	2.726	2.522	2.395	2.307	2.243	2.193	2.154	2.095	2.054	1.938	1.797
15	3.073	2.695	2.490	2.361	2.273	2.208	2.158	2.119	2.059	2.017	1.899	1.755
16	3.048	2.668	2.462	2.333	2.244	2.178	2.128	2.088	2.028	1.985	1.866	1.718
17	3.026	2.645	2.437	2.308	2.218	2.152	2.102	2.061	2.001	1.958	1.836	1.686
18	3.007	2.624	2.416	2.286	2.196	2.130	2.079	2.038	1.977	1.933	1.810	1.657
19	2.990	2.606	2.397	2.266	2.176	2.109	2.058	2.017	1.956	1.912	1.787	1.631
20	2.975	2.589	2.380	2.249	2.158	2.091	2.040	1.999	1.937	1.892	1.767	1.607
21	2.961	2.575	2.365	2.233	2.142	2.075	2.023	1.982	1.920	1.875	1.748	1.586
22	2.949	2.561	2.351	2.219	2.128	2.060	2.008	1.967	1.904	1.859	1.731	1.567
23	2.937	2.549	2.339	2.207	2.115	2.047	1.995	1.953	1.890	1.845	1.716	1.549
24	2.927	2.538	2.327	2.195	2.103	2.035	1.983	1.941	1.877	1.832	1.702	1.533
25	2.918	2.528	2.317	2.184	2.092	2.024	1.971	1.929	1.866	1.820	1.689	1.518
26	2.909	2.519	2.307	2.174	2.082	2.014	1.961	1.919	1.855	1.809	1.677	1.504
27	2.901	2.511	2.299	2.165	2.073	2.005	1.952	1.909	1.845	1.799	1.666	1.491
28	2.894	2.503	2.291	2.157	2.064	1.996	1.943	1.900	1.836	1.790	1.656	1.478
29	2.887	2.495	2.283	2.149	2.057	1.988	1.935	1.892	1.827	1.781	1.647	1.467
30	2.881	2.489	2.276	2.142	2.049	1.980	1.927	1.884	1.819	1.773	1.638	1.456
32	2.869	2.477	2.263	2.129	2.036	1.967	1.913	1.870	1.805	1.758	1.622	1.437
34	2.859	2.466	2.252	2.118	2.024	1.955	1.901	1.858	1.793	1.745	1.608	1.419
36	2.850	2.456	2.243	2.108	2.014	1.945	1.891	1.847	1.781	1.734	1.595	1.404
38	2.842	2.448	2.234	2.099	2.005	1.935	1.881	1.838	1.772	1.724	1.584	1.390
40	2.835	2.440	2.226	2.091	1.997	1.927	1.873	1.829	1.763	1.715	1.574	1.377
60	2.791	2.393	2.177	2.041	1.946	1.875	1.819	1.775	1.707	1.657	1.511	1.291
120	2.748	2.347	2.130	1.992	1.896	1.824	1.767	1.722	1.652	1.601	1.447	1.193
	2.706	2.303	2.084	1.945	1.847	1.774	1.717	1.670	1.599	1.546	1.383	1.000

TABLO 4(b) %5'e göre F DAĞILIMI

v1 = v2 = 1	1	2	3	4	5	6	7	8	10	12	24	
2	161.4	199.5	215.7	224.6	230.2	234.0	236.8	238.9	241.9	243.9	249.1	254.3
3	18.51	19.00	19.16	19.25	19.30	19.33	19.35	19.37	19.40	19.41	19.45	19.50
4	10.13	9.552	9.277	9.117	9.013	8.941	8.887	8.845	8.786	8.745	8.639	8.526
5	7.709	6.944	6.591	6.388	6.256	6.163	6.094	6.041	5.964	5.912	5.774	5.628
6	6.608	5.786	5.409	5.192	5.050	4.950	4.876	4.818	4.735	4.678	4.527	4.365
7	5.987	5.143	4.757	4.534	4.387	4.284	4.207	4.147	4.060	4.000	3.841	3.669
8	5.591	4.737	4.347	4.120	3.972	3.866	3.787	3.726	3.637	3.575	3.410	3.230
9	5.318	4.459	4.066	3.838	3.687	3.581	3.500	3.438	3.347	3.284	3.115	2.928
10	5.117	4.256	3.863	3.633	3.482	3.374	3.293	3.230	3.137	3.073	2.900	2.707
11	4.965	4.103	3.708	3.478	3.326	3.217	3.135	3.072	2.978	2.913	2.737	2.538
12	4.844	3.982	3.587	3.357	3.204	3.095	3.012	2.948	2.854	2.788	2.609	2.404
13	4.747	3.885	3.490	3.259	3.106	2.996	2.913	2.849	2.753	2.687	2.505	2.296
14	4.667	3.806	3.411	3.179	3.025	2.915	2.832	2.767	2.671	2.604	2.420	2.206
15	4.600	3.739	3.344	3.112	2.958	2.848	2.764	2.699	2.602	2.534	2.349	2.131
16	4.543	3.682	3.287	3.056	2.901	2.790	2.707	2.641	2.544	2.475	2.288	2.066
17	4.494	3.634	3.239	3.007	2.852	2.741	2.657	2.591	2.494	2.425	2.235	2.010
18	4.451	3.592	3.197	2.965	2.810	2.699	2.614	2.548	2.450	2.381	2.190	1.960
19	4.414	3.555	3.160	2.928	2.773	2.661	2.577	2.510	2.412	2.342	2.150	1.917
20	4.381	3.522	3.127	2.895	2.740	2.628	2.544	2.477	2.378	2.308	2.114	1.878
21	4.351	3.493	3.098	2.866	2.711	2.599	2.514	2.447	2.348	2.278	2.082	1.843
22	4.325	3.467	3.072	2.840	2.685	2.573	2.488	2.420	2.321	2.250	2.054	1.812
23	4.301	3.443	3.049	2.817	2.661	2.549	2.464	2.397	2.297	2.226	2.028	1.783
24	4.279	3.422	3.028	2.796	2.640	2.528	2.442	2.375	2.275	2.204	2.005	1.757
25	4.260	3.403	3.009	2.776	2.621	2.508	2.423	2.355	2.255	2.183	1.984	1.733
26	4.242	3.385	2.991	2.759	2.603	2.490	2.405	2.337	2.236	2.165	1.964	1.711
27	4.225	3.369	2.975	2.743	2.587	2.474	2.388	2.321	2.220	2.148	1.946	1.691
28	4.210	3.354	2.960	2.728	2.572	2.459	2.373	2.305	2.204	2.132	1.930	1.672
29	4.196	3.340	2.947	2.714	2.558	2.445	2.359	2.291	2.190	2.118	1.915	1.654
30	4.183	3.328	2.934	2.701	2.545	2.432	2.346	2.278	2.177	2.104	1.901	1.638
32	4.171	3.316	2.922	2.690	2.534	2.421	2.334	2.266	2.165	2.092	1.887	1.622
34	4.149	3.295	2.901	2.668	2.512	2.399	2.313	2.244	2.142	2.070	1.864	1.594
36	4.130	3.276	2.883	2.650	2.494	2.380	2.294	2.225	2.123	2.050	1.843	1.569
40	4.113	3.259	2.866	2.634	2.477	2.364	2.277	2.209	2.106	2.033	1.824	1.547
60	4.098	3.245	2.852	2.619	2.463	2.349	2.262	2.194	2.091	2.017	1.808	1.527
120	4.085	3.232	2.839	2.606	2.449	2.336	2.249	2.180	2.077	2.003	1.793	1.509
	4.001	3.150	2.758	2.525	2.368	2.254	2.167	2.097	1.993	1.17	1.700	1.389
	3.920	3.072	2.680	2.447	2.290	2.175	2.087	2.016	1.910	1.834	1.608	1.254
	3.841	2.996	2.605	2.372	2.214	2.099	2.010	1.938	1.831	1.752	1.517	1.000

TABLO 4(c) %2.5'a göre F DAĞILIMI

v1 =	1	2	3	4	5	6	7	8	10	12	24	
v2 = 1	647.8	799.5	864.2	899.6	921.8	937.1	948.2	956.7	968.6	976.7	997.2	1018
2	38.51	39.00	39.17	39.25	39.30	39.33	39.36	39.37	39.40	39.41	39.46	39.50
3	17.44	16.04	15.44	15.10	14.88	14.73	14.62	14.54	14.42	14.34	14.12	13.90
4	12.22	10.65	9.979	9.605	9.364	9.197	9.074	8.980	8.844	8.751	8.511	8.257
5	10.01	8.434	7.764	7.388	7.146	6.978	6.853	6.757	6.619	6.525	6.278	6.015
6	8.813	7.260	6.599	6.227	5.988	5.820	5.695	5.600	5.461	5.366	5.117	4.849
7	8.073	6.542	5.890	5.523	5.285	5.119	4.995	4.899	4.761	4.666	4.415	4.142
8	7.571	6.059	5.416	5.053	4.817	4.652	4.529	4.433	4.295	4.200	3.947	3.670
9	7.209	5.715	5.078	4.718	4.484	4.320	4.197	4.102	3.964	3.868	3.614	3.333
10	6.937	5.456	4.826	4.468	4.236	4.072	3.950	3.855	3.717	3.621	3.365	3.080
11	6.724	5.256	4.630	4.275	4.044	3.881	3.759	3.664	3.526	3.430	3.173	2.883
12	6.554	5.096	4.474	4.121	3.891	3.728	3.607	3.512	3.374	3.277	3.019	2.725
13	6.414	4.965	4.347	3.996	3.767	3.604	3.483	3.388	3.250	3.153	2.893	2.595
14	6.298	4.857	4.242	3.892	3.663	3.501	3.380	3.285	3.147	3.050	2.789	2.487
15	6.200	4.765	4.153	3.804	3.576	3.415	3.293	3.199	3.060	2.963	2.701	2.395
16	6.115	4.687	4.077	3.729	3.502	3.341	3.219	3.125	2.986	2.889	2.625	2.316
17	6.042	4.619	4.011	3.665	3.438	3.277	3.156	3.061	2.922	2.825	2.560	2.247
18	5.978	4.560	3.954	3.608	3.382	3.221	3.100	3.005	2.866	2.769	2.503	2.187
19	5.922	4.508	3.903	3.559	3.333	3.172	3.051	2.956	2.817	2.720	2.452	2.133
20	5.871	4.461	3.859	3.515	3.289	3.128	3.007	2.913	2.774	2.676	2.408	2.085
21	5.827	4.420	3.819	3.475	3.250	3.090	2.969	2.874	2.735	2.637	2.368	2.042
22	5.786	4.383	3.783	3.440	3.215	3.055	2.934	2.839	2.700	2.602	2.331	2.003
23	5.750	4.349	3.750	3.408	3.183	3.023	2.902	2.808	2.668	2.570	2.299	1.968
24	5.717	4.319	3.721	3.379	3.155	2.995	2.874	2.779	2.640	2.541	2.269	1.935
25	5.686	4.291	3.694	3.353	3.129	2.968	2.848	2.753	2.613	2.515	2.242	1.906
26	5.659	4.265	3.670	3.329	3.105	2.945	2.824	2.729	2.590	2.491	2.217	1.878
27	5.633	4.242	3.647	3.307	3.083	2.923	2.802	2.707	2.568	2.469	2.195	1.853
28	5.610	4.221	3.626	3.286	3.063	2.903	2.782	2.687	2.547	2.448	2.174	1.829
29	5.588	4.201	3.607	3.267	3.044	2.884	2.763	2.669	2.529	2.430	2.154	1.807
30	5.568	4.182	3.589	3.250	3.026	2.867	2.746	2.651	2.511	2.412	2.136	1.787
32	5.531	4.149	3.557	3.218	2.995	2.836	2.715	2.620	2.480	2.381	2.103	1.750
34	5.499	4.120	3.529	3.191	2.968	2.808	2.688	2.593	2.453	2.353	2.075	1.717
36	5.471	4.094	3.505	3.167	2.944	2.785	2.664	2.569	2.429	2.329	2.049	1.687
38	5.446	4.071	3.483	3.145	2.923	2.763	2.643	2.548	2.407	2.307	2.027	1.661
40	5.424	4.051	3.463	3.126	2.904	2.744	2.624	2.529	2.388	2.288	2.007	1.637
60	5.286	3.925	3.343	3.008	2.786	2.627	2.507	2.412	2.270	2.169	1.882	1.482
120	5.152	3.805	3.227	2.894	2.674	2.515	2.395	2.299	2.157	2.055	1.760	1.310
	5.024	3.689	3.116	2.786	2.567	2.408	2.288	2.192	2.048	1.945	1.640	1.000

TABLO 4(d) %1'e göre F DAĞILIMI

v1 = v2 = 1	1	2	3	4	5	6	7	8	10	12	24
2	4052	4999	5403	5625	5764	5859	5928	5981	6056	6106	6235
3	98.50	99.00	99.17	99.25	99.30	99.33	99.36	99.37	99.40	99.42	99.46
4	34.12	30.82	29.46	28.71	28.24	27.91	27.67	27.49	27.23	27.05	26.60
5	21.20	18.00	16.69	15.98	15.52	15.21	14.98	14.80	14.55	14.37	13.93
6	16.26	13.27	12.06	11.39	10.97	10.67	10.46	10.29	10.05	9.888	9.466
7	13.75	10.92	9.780	9.148	8.746	8.466	8.260	8.102	7.874	7.718	7.313
8	12.25	9.547	8.451	7.847	7.460	7.191	6.93	6.840	6.620	6.469	6.074
9	11.26	8.649	7.591	7.006	6.632	6.371	6.178	6.029	5.814	5.667	5.279
10	10.56	8.022	6.992	6.422	6.057	5.802	5.613	5.467	5.257	5.111	4.729
11	10.04	7.559	6.552	5.994	5.636	5.386	5.200	5.057	4.849	4.706	4.327
12	9.646	7.206	6.217	5.668	5.316	5.069	4.886	4.744	4.539	4.397	4.021
13	9.330	6.927	5.953	5.412	5.064	4.821	4.640	4.499	4.296	4.155	3.780
14	9.074	6.701	5.739	5.205	4.862	4.620	4.441	4.302	4.100	3.960	3.587
15	8.862	6.515	5.564	5.035	4.695	4.456	4.278	4.140	3.939	3.800	3.427
16	8.683	6.359	5.417	4.893	4.556	4.318	4.142	4.004	3.805	3.666	3.294
17	8.531	6.226	5.292	4.773	4.437	4.202	4.026	3.890	3.691	3.553	3.181
18	8.400	6.112	5.185	4.669	4.336	4.102	3.927	3.791	3.593	3.455	3.084
19	8.285	6.013	5.092	4.579	4.248	4.015	3.841	3.705	3.508	3.371	2.999
20	8.185	5.926	5.010	4.500	4.171	3.939	3.765	3.631	3.434	3.297	2.925
21	8.096	5.849	4.938	4.431	4.103	3.871	3.699	3.564	3.368	3.231	2.859
22	8.017	5.780	4.874	4.369	4.042	3.812	3.640	3.506	3.310	3.173	2.801
23	7.945	5.719	4.817	4.313	3.988	3.758	3.587	3.453	3.258	3.121	2.749
24	7.881	5.664	4.765	4.264	3.939	3.710	3.539	3.406	3.211	3.074	2.702
25	7.823	5.614	4.718	4.218	3.895	3.667	3.496	3.363	3.168	3.032	2.659
26	7.770	5.568	4.675	4.177	3.855	3.627	3.457	3.324	3.129	2.993	2.620
27	7.721	5.526	4.637	4.140	3.818	3.591	3.421	3.288	3.094	2.958	2.585
28	7.677	5.488	4.601	4.106	3.785	3.558	3.388	3.256	3.062	2.926	2.552
29	7.636	5.453	4.568	4.074	3.754	3.528	3.358	3.226	3.032	2.896	2.522
30	7.598	5.420	4.538	4.045	3.725	3.499	3.330	3.198	3.005	2.868	2.495
32	7.562	5.390	4.510	4.018	3.699	3.473	3.304	3.173	2.979	2.843	2.469
34	7.499	5.336	4.459	3.969	3.652	3.427	3.258	3.127	2.934	2.798	2.423
36	7.444	5.289	4.416	3.927	3.611	3.386	3.218	3.087	2.894	2.758	2.383
38	7.396	5.248	4.377	3.890	3.574	3.351	3.183	3.052	2.859	2.723	2.347
40	7.353	5.211	4.343	3.858	3.542	3.319	3.152	3.021	2.828	2.692	2.316
60	7.314	5.179	4.313	3.828	3.514	3.291	3.124	2.993	2.801	2.665	2.288
120	7.077	4.977	4.126	3.649	3.339	3.119	2.953	2.823	2.632	2.496	2.115
	6.851	4.787	3.949	3.480	3.174	2.956	2.792	2.663	2.472	2.336	1.950
	6.635	4.605	3.782	3.319	3.017	2.802	2.639	2.511	2.321	2.185	1.791

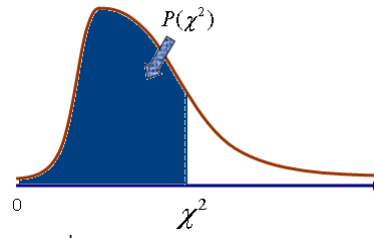
TABLO 4(e) %0.5'e göre F DAĞILIMI

v1 =	1	2	3	4	5	6	7	8	10	12	24	
v2 = 1	16211	20000	21615	22500	23056	23437	23715	32925	24224	24426	24940	25464
2	198.5	199.0	199.2	199.2	199.3	199.3	199.4	199.4	199.4	199.4	199.5	199.5
3	55.55	49.80	47.47	46.19	45.39	44.84	44.43	44.13	43.69	43.39	42.62	41.83
4	31.33	26.28	24.26	23.15	22.46	21.97	21.62	21.35	20.97	20.70	20.03	19.32
5	22.78	18.31	16.53	15.56	14.94	14.51	14.20	13.96	13.62	13.38	12.78	12.14
6	18.63	14.54	12.92	12.03	11.46	11.07	10.79	10.57	10.25	10.03	9.474	8.879
7	16.24	12.40	10.88	10.05	9.522	9.155	8.885	8.678	8.380	8.176	7.645	7.076
8	14.69	11.04	9.596	8.805	8.302	7.952	7.694	7.496	7.211	7.015	6.503	5.951
9	13.61	10.11	8.717	7.956	7.471	7.134	6.885	6.693	6.417	6.227	5.729	5.188
10	12.83	9.427	8.081	7.343	6.872	6.545	6.302	6.116	5.847	5.661	5.173	4.639
11	12.23	8.912	7.600	6.881	6.422	6.102	5.865	5.682	5.418	5.236	4.756	4.226
12	11.75	8.510	7.226	6.521	6.071	5.757	5.525	5.345	5.085	4.906	4.431	3.904
13	11.37	8.186	6.926	6.233	5.791	5.482	5.253	5.076	4.820	4.643	4.173	3.647
14	11.06	7.922	6.680	5.998	5.562	5.257	5.031	4.857	4.603	4.428	3.961	3.436
15	10.80	7.701	6.476	5.803	5.372	5.071	4.847	4.674	4.424	4.250	3.786	3.260
16	10.58	7.514	6.303	5.638	5.212	4.913	4.692	4.521	4.272	4.099	3.638	3.112
17	10.38	7.354	6.156	5.497	5.075	4.779	4.559	4.389	4.142	3.971	3.511	2.984
18	10.22	7.215	6.028	5.375	4.956	4.663	4.445	4.276	4.030	3.860	3.402	2.873
19	10.07	7.093	5.916	5.268	4.853	4.561	4.345	4.177	3.933	3.763	3.306	2.776
20	9.944	6.986	5.818	5.174	4.762	4.472	4.257	4.090	3.847	3.678	3.222	2.690
21	9.830	6.891	5.730	5.091	4.681	4.393	4.179	4.013	3.771	3.602	3.147	2.614
22	9.727	6.806	5.652	5.017	4.609	4.322	4.109	3.944	3.703	3.535	3.081	2.545
23	9.635	6.730	5.582	4.950	4.544	4.259	4.047	3.882	3.642	3.475	3.021	2.484
24	9.551	6.661	5.519	4.890	4.486	4.202	3.991	3.826	3.587	3.420	2.967	2.428
25	9.475	6.598	5.462	4.835	4.433	4.150	3.939	3.776	3.537	3.370	2.918	2.377
26	9.406	6.541	5.409	4.785	4.384	4.103	3.893	3.730	3.492	3.325	2.873	2.330
27	9.342	6.489	5.361	4.740	4.340	4.059	3.850	3.687	3.450	3.284	2.832	2.287
28	9.284	6.440	5.317	4.698	4.300	4.020	3.811	3.649	3.412	3.246	2.794	2.247
29	9.230	6.396	5.276	4.659	4.262	3.983	3.775	3.613	3.377	3.211	2.759	2.210
30	9.180	6.355	5.239	4.623	4.228	3.949	3.742	3.580	3.344	3.179	2.727	2.176
32	9.090	6.281	5.171	4.559	4.166	3.889	3.682	3.521	3.286	3.121	2.670	2.114
34	9.012	6.217	5.113	4.504	4.112	3.836	3.630	3.470	3.235	3.071	2.620	2.060
36	8.943	6.161	5.062	4.455	4.065	3.790	3.585	3.425	3.191	3.027	2.576	2.013
38	8.882	6.111	5.016	4.412	4.023	3.749	3.545	3.385	3.152	2.988	2.537	1.970
40	8.828	6.066	4.976	4.374	3.986	3.713	3.509	3.350	3.117	2.953	2.502	1.932
60	8.495	5.795	4.729	4.140	3.760	3.492	3.291	3.134	2.904	2.742	2.290	1.689
120	8.179	5.539	4.497	3.921	3.548	3.285	3.087	2.933	2.705	2.544	2.089	1.431
	7.879	5.298	4.279	3.715	3.350	3.091	2.897	2.744	2.519	2.358	1.898	1.000

TABLO 4(f) %0.1'e göre F DAĞILIMI

v1 = v2 = 1	1	2	3	4	5	6	7	8	10	12	24
2	998.5	999.0	999.2	999.2	999.3	999.3	999.4	999.4	999.4	999.4	999.5
3	167.0	148.5	141.1	137.1	134.6	132.8	131.6	130.6	129.2	128.3	125.9
4	74.14	61.25	56.18	53.44	51.71	50.53	49.66	49.00	48.05	47.41	45.77
5	47.18	37.12	33.20	31.09	29.75	28.83	28.16	27.65	26.92	26.42	25.13
6	35.51	27.00	23.70	21.92	20.80	20.03	19.46	19.03	18.41	17.99	16.90
7	29.25	21.69	18.77	17.20	16.21	15.52	15.02	14.63	14.08	13.71	12.73
8	25.41	18.49	15.83	14.39	13.48	12.86	12.40	12.05	11.54	11.19	10.30
9	22.86	16.39	13.90	12.96	11.71	11.13	10.70	10.37	9.894	9.570	8.724
10	21.04	14.91	12.55	11.28	10.48	9.926	9.517	9.204	8.754	8.445	7.638
11	19.69	13.81	11.56	10.35	9.578	9.047	8.655	8.355	7.922	7.626	6.847
12	18.64	12.97	10.80	9.633	8.892	8.379	8.001	7.710	7.292	7.005	6.249
13	17.82	12.31	10.21	9.073	8.354	7.856	7.489	7.206	6.799	6.519	5.781
14	17.14	11.78	9.729	8.622	7.922	7.436	7.077	6.802	6.404	6.130	5.407
15	16.59	11.34	9.335	8.253	7.567	7.092	6.741	6.471	6.081	5.812	5.101
16	16.12	10.97	9.006	7.944	7.272	6.805	6.460	6.195	5.812	5.547	4.846
17	15.72	10.66	8.727	7.683	7.022	6.562	6.223	5.962	5.584	5.324	4.631
18	15.38	10.39	8.487	7.459	6.808	6.355	6.021	5.763	5.390	5.132	4.447
19	15.08	10.16	8.280	7.265	6.622	6.175	5.845	5.590	5.222	4.967	4.288
20	14.82	9.953	8.098	7.096	6.461	6.019	5.692	5.440	5.075	4.823	4.149
21	14.59	9.772	7.938	6.947	6.318	5.881	5.557	5.308	4.946	4.696	4.027
22	14.38	9.612	7.796	6.814	6.191	5.758	5.438	5.190	4.832	4.583	3.919
23	14.20	9.469	7.669	6.696	6.078	5.649	5.331	5.085	4.730	4.483	3.822
24	14.03	9.339	7.554	6.589	5.977	5.550	5.235	4.991	4.638	4.393	3.735
25	13.88	9.223	7.451	6.493	5.885	5.462	5.148	4.906	4.555	4.312	3.657
26	13.74	9.116	7.357	6.406	5.802	5.381	5.070	4.829	4.480	4.238	3.586
27	13.61	9.019	7.272	6.326	5.726	5.308	4.998	4.759	4.412	4.171	3.521
28	13.50	8.931	7.193	6.253	5.656	5.241	4.933	4.695	4.349	4.109	3.462
29	13.39	8.849	7.121	6.186	5.593	5.179	4.873	4.636	4.292	4.053	3.407
30	13.29	8.773	7.054	6.125	5.534	5.122	4.817	4.581	4.239	4.001	3.357
32	13.12	8.639	6.936	6.014	5.429	5.021	4.719	4.485	4.145	3.908	3.268
34	12.97	8.522	6.833	5.919	5.339	4.934	4.633	4.401	4.063	3.828	3.191
36	12.83	8.420	6.744	5.836	5.260	4.857	4.559	4.328	3.992	3.758	3.123
38	12.71	8.331	6.665	5.763	5.190	4.790	4.494	4.264	3.930	3.697	3.064
40	12.61	8.251	6.595	5.698	5.128	4.731	4.436	4.207	3.874	3.642	3.011
60	11.97	7.768	6.171	5.307	4.757	4.372	4.086	3.865	3.541	3.315	2.694
120	11.38	7.321	5.781	4.947	4.416	4.044	3.767	3.552	3.237	3.016	2.402
	10.83	6.908	5.422	4.617	4.103	3.743	3.475	3.266	2.959	2.742	2.132

TABLO 5 Kİ-KARE DAĞILIM FONKSİYONU



$\nu =$	1	$\nu =$	1	$\nu =$	2	$\nu =$	2	$\nu =$	3	$\nu =$	3
$x=0,0$	0,0000	$x=4,0$	0,9545	$x=0,0$	0,0000	$x=4,0$	0,8647	$x=0,0$	0,0000	$x=4,0$	0,7385
0,1	0,2482	4,1	0,9571	0,1	0,0488	4,1	0,8713	0,1	0,0082	4,2	0,7593
0,2	0,3453	4,2	0,9596	0,2	0,0952	4,2	0,8775	0,2	0,0224	4,4	0,7786
0,3	0,4161	4,3	0,9619	0,3	0,1393	4,3	0,8835	0,3	0,0400	4,6	0,7965
0,4	0,4729	4,4	0,9641	0,4	0,1813	4,4	0,8892	0,4	0,0598	4,8	0,8130
0,5	0,5205	4,5	0,9661	0,5	0,2212	4,5	0,8946	0,5	0,0811	5,0	0,8282
0,6	0,5614	4,6	0,9680	0,6	0,2592	4,6	0,8997	0,6	0,1036	5,2	0,8423
0,7	0,5972	4,7	0,9698	0,7	0,2953	4,7	0,9046	0,7	0,1268	5,4	0,8553
0,8	0,6289	4,8	0,9715	0,8	0,3297	4,8	0,9093	0,8	0,1505	5,6	0,8672
0,9	0,6572	4,9	0,9731	0,9	0,3624	4,9	0,9137	0,9	0,1746	5,8	0,8782
1,0	0,6827	5,0	0,9747	1,0	0,3935	5,0	0,9179	1,0	0,1987	6,0	0,8884
1,1	0,7057	5,1	0,9761	1,1	0,4231	5,1	0,9219	1,1	0,2229	6,2	0,8977
1,2	0,7267	5,2	0,9774	1,2	0,4512	5,2	0,9257	1,2	0,2470	6,4	0,9063
1,3	0,7458	5,3	0,9787	1,3	0,4780	5,3	0,9293	1,3	0,2709	6,6	0,9142
1,4	0,7633	5,4	0,9799	1,4	0,5034	5,4	0,9328	1,4	0,2945	6,8	0,9214
1,5	0,7793	5,5	0,9810	1,5	0,5276	5,5	0,9361	1,5	0,3177	7,0	0,9281
1,6	0,7941	5,6	0,9820	1,6	0,5507	5,6	0,9392	1,6	0,3406	7,2	0,9342
1,7	0,8077	5,7	0,9830	1,7	0,5726	5,7	0,9422	1,7	0,3631	7,4	0,9398
1,8	0,8203	5,8	0,9840	1,8	0,5934	5,8	0,9450	1,8	0,3851	7,6	0,9450
1,9	0,8319	5,9	0,9849	1,9	0,6133	5,9	0,9477	1,9	0,4066	7,8	0,9497
2,0	0,8427	6,0	0,9857	2,0	0,6321	6,0	0,9502	2,0	0,4276	8,0	0,9540
2,1	0,8527	6,1	0,9865	2,1	0,6501	6,2	0,9550	2,1	0,4481	8,2	0,9579
2,2	0,8620	6,2	0,9872	2,2	0,6671	6,4	0,9592	2,2	0,4681	8,4	0,9616
2,3	0,8706	6,3	0,9879	2,3	0,6834	6,6	0,9631	2,3	0,4875	8,6	0,9649
2,4	0,8787	6,4	0,9886	2,4	0,6988	6,8	0,9666	2,4	0,5064	8,8	0,9679
2,5	0,8862	6,5	0,9892	2,5	0,7135	7,0	0,9698	2,5	0,5247	9,0	0,9707
2,6	0,8931	6,6	0,9898	2,6	0,7275	7,2	0,9727	2,6	0,5425	9,2	0,9733
2,7	0,8997	6,7	0,9904	2,7	0,7408	7,4	0,9753	2,7	0,5598	9,4	0,9756
2,8	0,9057	6,8	0,9909	2,8	0,7534	7,6	0,9776	2,8	0,5765	9,6	0,9777
2,9	0,9114	6,9	0,9914	2,9	0,7654	7,8	0,9798	2,9	0,5927	9,8	0,9797
3,0	0,9167	7,0	0,9918	3,0	0,7769	8,0	0,9817	3,0	0,6084	10,0	0,9814
3,1	0,9217	7,1	0,9923	3,1	0,7878	8,2	0,9834	3,1	0,6235	10,2	0,9831
3,2	0,9264	7,2	0,9927	3,2	0,7981	8,4	0,9850	3,2	0,6382	10,4	0,9845
3,3	0,9307	7,3	0,9931	3,3	0,8080	8,6	0,9864	3,3	0,6524	10,6	0,9859
3,4	0,9348	7,4	0,9935	3,4	0,8173	8,8	0,9877	3,4	0,6660	10,8	0,9871
3,5	0,9386	7,5	0,9938	3,5	0,8262	9,0	0,9889	3,5	0,6792	11,0	0,9883
3,6	0,9422	7,6	0,9942	3,6	0,8347	9,2	0,9899	3,6	0,6920	11,2	0,9893
3,7	0,9456	7,7	0,9945	3,7	0,8428	9,4	0,9909	3,7	0,7043	11,4	0,9903
3,8	0,9487	7,8	0,9948	3,8	0,8504	9,6	0,9918	3,8	0,7161	11,6	0,9911
3,9	0,9517	7,9	0,9951	3,9	0,8577	9,8	0,9926	3,9	0,7275	11,8	0,9919
4,0	0,9545	8,0	0,9953	4,0	0,8647	10,0	0,9933	4,0	0,7385	12,0	0,9926

TABLO 5 Kİ-KARE DAĞILIM FONKSİYONU (Devam)

$\nu =$	4	5	6	7	8	9	10	11	12	13	14
$\alpha=0,5$	0,0265	0,008	0,002	0,0006	0,0001						
1,0	0,09	0,037	0,014	0,0052	0,0018	0,0006	0,0002	0,0001			
1,5	0,173	0,087	0,041	0,0177	0,0073	0,0029	0,0011	0,0004	0,0001		
2,0	0,264	0,1500	0,08	0,0402	0,0190	0,0085	0,0037	0,0015	0,0006	0,0002	0,0001
2,5	0,355	0,224	0,132	0,0729	0,0383	0,0191	0,0091	0,0042	0,0018	0,0008	0,0003
3,0	0,442	0,3000	0,191	0,1150	0,0656	0,0357	0,0186	0,0093	0,0045	0,0021	0,0009
3,5	0,522	0,377	0,2560	0,1648	0,1008	0,0589	0,0329	0,0177	0,0091	0,0046	0,0022
4,0	0,5940	0,451	0,323	0,2202	0,1429	0,0886	0,0527	0,0301	0,0166	0,0088	0,0045
4,5	0,658	0,52	0,391	0,2793	0,1906	0,1245	0,0780	0,0471	0,0274	0,0154	0,0084
5,0	0,713	0,584	0,456	0,3400	0,2424	0,1657	0,1088	0,0688	0,0420	0,0248	0,0142
5,5	0,76	0,642	0,519	0,4008	0,2970	0,2113	0,1446	0,0954	0,0608	0,0375	0,0224
6,0	0,801	0,694	0,577	0,4603	0,3528	0,2601	0,1847	0,1266	0,0839	0,0538	0,0335
6,5	0,835	0,739	0,63	0,5173	0,4086	0,3110	0,2283	0,1620	0,1112	0,0739	0,0477
7,0	0,864	0,779	0,679	0,5711	0,4634	0,3629	0,2746	0,2009	0,1424	0,0978	0,0653
7,5	0,888	0,8140	0,723	0,6213	0,5162	0,4148	0,3225	0,2427	0,1771	0,1254	0,0863
8,0	0,908	0,844	0,762	0,6674	0,5665	0,4659	0,3712	0,2867	0,2149	0,1564	0,1107
8,5	0,925	0,869	0,796	0,7094	0,6138	0,5154	0,4199	0,3321	0,2551	0,1904	0,1383
9,0	0,939	0,891	0,826	0,7473	0,6577	0,5627	0,4679	0,3781	0,2971	0,2271	0,1689
9,5	0,95	0,909	0,853	0,7813	0,6981	0,6075	0,5146	0,4242	0,3403	0,2658	0,2022
10,0	0,96	0,925	0,875	0,8114	0,7350	0,6495	0,5595	0,4696	0,3840	0,3061	0,2378
10,5	0,967	0,938	0,895	0,8380	0,7683	0,6885	0,6022	0,5140	0,4278	0,3474	0,2752
11,0	0,973	0,949	0,912	0,8614	0,7983	0,7243	0,6425	0,5567	0,4711	0,3892	0,3140
11,5	0,979	0,958	0,926	0,8818	0,8251	0,7570	0,6801	0,5976	0,5134	0,4310	0,3536
12,0	0,983	0,965	0,9380	0,8994	0,8488	0,7867	0,7149	0,6364	0,5543	0,4724	0,3937
12,5	0,9860	0,972	0,948	0,9147	0,8697	0,8134	0,7470	0,6727	0,5936	0,5129	0,4338
13,0	0,989	0,977	0,9570	0,9279	0,8882	0,8374	0,7763	0,7067	0,6310	0,5522	0,4735
13,5	0,9909	0,981	0,964	0,9392	0,9042	0,8587	0,8030	0,7381	0,6662	0,5900	0,5124
14,0	0,993	0,984	0,97	0,9488	0,9182	0,8777	0,8270	0,7670	0,6993	0,6262	0,5503
14,5	0,994	0,987	0,976	0,9570	0,9304	0,8944	0,8486	0,7935	0,7301	0,6604	0,5868
15,0	0,995	0,99	0,98	0,9640	0,9409	0,9091	0,8679	0,8175	0,7586	0,6926	0,6218
15,5	0,996	0,992	0,983	0,9699	0,9499	0,9219	0,8851	0,8393	0,7848	0,7228	0,6551
16,0	0,9970	0,993	0,986	0,9749	0,9576	0,9331	0,9004	0,8589	0,8088	0,7509	0,6866
16,5	0,998	0,994	0,989	0,9791	0,9642	0,9429	0,9138	0,8764	0,8306	0,7768	0,7162
17,0	0,998	0,996	0,991	0,9826	0,9699	0,9513	0,9256	0,8921	0,8504	0,8007	0,7438
17,5	0,999	0,996	0,992	0,9856	0,9747	0,9586	0,9360	0,9061	0,8683	0,8226	0,7695
18,0	0,999	0,997	0,994	0,9880	0,9788	0,9648	0,9450	0,9184	0,8843	0,8425	0,7932
18,5	0,9990	0,998	0,995	0,9901	0,9822	0,9702	0,9529	0,9293	0,8987	0,8606	0,8151
19,0	0,999	0,998	0,996	0,9918	0,9851	0,9748	0,9597	0,9389	0,9115	0,8769	0,8351
19,5	0,999	0,998	0,997	0,9932	0,9876	0,9787	0,9656	0,9473	0,9228	0,8916	0,8533
20,0	1	0,999	0,997	0,9944	0,9897	0,9821	0,9707	0,9547	0,9329	0,9048	0,8699
21,0	1	0,999	0,998	0,9962	0,9929	0,9873	0,9789	0,9666	0,9496	0,9271	0,8984
22,0	1	1	0,999	0,9975	0,9951	0,9911	0,9849	0,9756	0,9625	0,9446	0,9214
23,0	1	1	0,999	0,9983	0,9966	0,9938	0,9893	0,9823	0,9723	0,9583	0,9397
24,0	1	1	1	0,9989	0,9977	0,9957	0,9924	0,9873	0,9797	0,9689	0,9542
25,0	1	1	1	0,9992	0,9984	0,9970	0,9947	0,9909	0,9852	0,9769	0,9654
26,0		1	1	0,9995	0,9989	0,9980	0,9963	0,9935	0,9893	0,9830	0,9741
27,0		1	1	0,9997	0,9993	0,9986	0,9974	0,9954	0,9923	0,9876	0,9807
28,0			1	0,9998	0,9995	0,9990	0,9982	0,9968	0,9945	0,9910	0,9858
29,0			1	0,9999	0,9997	0,9994	0,9988	0,9977	0,9961	0,9935	0,9895
30,0				0,9999	0,9998	0,9996	0,9991	0,9984	0,9972	0,9953	0,9924

TABLO 5 Kİ-KARE DAĞILIM FONKSİYONU (Devam)

v =	15	16	17	18	19	20	21	22	23	24	25
n=3	0,0004	0,0002	0,0001								
4	0,0002	0,0011	0,0005	0,0002	0,0001						
5	0,008	0,0042	0,0022	0,0011	0,0006	0,0003	0,0001	0,0001			
6	0,02	0,0119	0,0068	0,0038	0,0021	0,0011	0,0006	0,0003	0,0001	0,0001	
7	0,042	0,0267	0,0165	0,0099	0,0058	0,0033	0,0019	0,0010	0,0005	0,0003	0,0001
8	0,076	0,0511	0,0335	0,0214	0,0133	0,0081	0,0049	0,0028	0,0016	0,0009	0,0005
9	0,123	0,0866	0,0597	0,0403	0,0265	0,0171	0,0108	0,0067	0,0040	0,0024	0,0014
10	0,18	0,1334	0,0964	0,0681	0,0471	0,0318	0,0211	0,0137	0,0087	0,0055	0,0033
11	0,2474	0,1905	0,1434	0,1056	0,0762	0,0538	0,0372	0,0253	0,0168	0,0110	0,0071
12	0,3210	0,2560	0,1999	0,1528	0,1144	0,0839	0,0604	0,0426	0,0295	0,0201	0,0134
13	0,398	0,3272	0,2638	0,2084	0,1641	0,1226	0,0914	0,0668	0,0480	0,0339	0,0235
14	0,475	0,4013	0,3329	0,2709	0,2163	0,1695	0,1304	0,0985	0,0731	0,0533	0,0383
15	0,549	0,4754	0,4045	0,3380	0,2774	0,2236	0,1770	0,1378	0,1054	0,0792	0,0586
16	0,618	0,5470	0,4762	0,4075	0,3427	0,2834	0,2303	0,1841	0,1447	0,1119	0,0852
17	0,681	0,6144	0,5456	0,4769	0,4101	0,3470	0,2889	0,2366	0,1907	0,1513	0,1182
18	0,737	0,6761	0,6112	0,5443	0,4776	0,4126	0,3510	0,2940	0,2425	0,1970	0,1576
19	0,786	0,7313	0,6715	0,6082	0,5432	0,4782	0,4149	0,3547	0,2988	0,2480	0,2029
20	0,828	0,7798	0,7258	0,6672	0,6054	0,5421	0,4787	0,4170	0,3581	0,3032	0,2532
21	0,863	0,8215	0,7737	0,7206	0,6632	0,6029	0,5411	0,4793	0,4189	0,3613	0,3074
22	0,892	0,8568	0,8153	0,7680	0,7157	0,6595	0,6005	0,5401	0,4797	0,4207	0,3643
23	0,916	0,8863	0,8507	0,8094	0,7627	0,7112	0,6560	0,5983	0,5392	0,4802	0,4224
24	0,935	0,9105	0,8806	0,8450	0,8038	0,7576	0,7069	0,6528	0,5962	0,5384	0,4806
25	0,95	0,9302	0,9053	0,8751	0,8395	0,7986	0,7528	0,7029	0,6497	0,5942	0,5376
26	0,9620	0,9460	0,9255	0,9002	0,8698	0,8342	0,7936	0,7483	0,6991	0,6468	0,5924
27	0,971	0,9585	0,9419	0,9210	0,8953	0,8647	0,8291	0,7888	0,7440	0,6955	0,6441
28	0,978	0,9684	0,9551	0,9379	0,9166	0,8906	0,8598	0,8243	0,7842	0,7400	0,6921
29	0,984	0,9761	0,9655	0,9516	0,9340	0,9122	0,8860	0,8551	0,8197	0,7799	0,7361
30	0,988	0,9820	0,9737	0,9626	0,9482	0,9301	0,9080	0,8815	0,8506	0,8152	0,7757
31	0,991	0,9865	0,9800	0,9712	0,9596	0,9448	0,9263	0,9039	0,8772	0,8462	0,8110
32	0,9936	0,9900	0,9850	0,9780	0,9687	0,9576	0,9414	0,9226	0,8999	0,8730	0,8420
33	0,995	0,9926	0,9887	0,9833	0,9760	0,9663	0,9538	0,9381	0,9189	0,8959	0,8689
34	0,9966	0,9946	0,9916	0,9874	0,9816	0,9739	0,9638	0,9509	0,9348	0,9153	0,8921
35	0,998	0,9960	0,9938	0,9905	0,9860	0,9799	0,9718	0,9613	0,9480	0,9316	0,9118
36	0,998	0,9971	0,9954	0,9929	0,9894	0,9846	0,9781	0,9696	0,9587	0,9451	0,9284
37	0,999	0,9979	0,9966	0,9948	0,9921	0,9883	0,9832	0,9763	0,9675	0,9562	0,9423
38	0,999	0,9985	0,9975	0,9961	0,9941	0,9911	0,9871	0,9817	0,9745	0,9653	0,9537
39	0,999	0,9989	0,9982	0,9972	0,9956	0,9933	0,9902	0,9859	0,9802	0,9727	0,9632
40	0,9995	0,9992	0,9987	0,9979	0,9967	0,9950	0,9926	0,9892	0,9846	0,9786	0,9708
41	1	0,9994	0,9991	0,9985	0,9976	0,9963	0,9944	0,9918	0,9882	0,9833	0,9770
42	1	0,9996	0,9993	0,9989	0,9982	0,9972	0,9958	0,9937	0,9909	0,9871	0,9820
43	1	0,9997	0,9995	0,9993	0,9987	0,9980	0,9969	0,9953	0,9931	0,9901	0,9860
44	1	0,9998	0,9997	0,9994	0,9991	0,9985	0,9977	0,9965	0,9947	0,9924	0,9892
45	1	0,9999	0,9998	0,9996	0,9993	0,9989	0,9983	0,9973	0,9960	0,9942	0,9916
46	0,9999	0,9999	0,9998	0,9997	0,9995	0,9992	0,9987	0,9980	0,9970	0,9956	0,9936
47		0,9999	0,9999	0,9998	0,9996	0,9994	0,9991	0,9985	0,9978	0,9967	0,9951
48			0,9999	0,9998	0,9997	0,9996	0,9993	0,9989	0,9983	0,9975	0,9963
49				0,9999	0,9998	0,9997	0,9995	0,9992	0,9988	0,9981	0,9972
50					0,9999	0,9999	0,9998	0,9996	0,9994	0,9991	0,9986
											0,9977

TABLO 6 Kİ-KARE DAĞILIMININ YÜZDELİK DİLİMLERE GÖRE DAĞILIMI

P =	99,95	99,9	99,5	99	97,5	95	90	80	70	60
v=1	0,003927	0,051571	0,039270	0,09171	0,039821	0,0039	0,016	0,064	0,149	0,2750
2	0,00100	0,002001	0,01003	0,02010	0,05064	0,1026	0,211	0,446	0,713	1,022
3	0,01528	0,02430	0,07172	0,1148	0,2158	0,3518	0,584	1,005	1,424	1,869
4	0,06392	0,09080	0,2070	0,2971	0,4844	0,7107	1,064	1,649	2,195	2,753
5	0,1581	0,2102	0,4117	0,5543	0,8312	1,145	1,610	2,343	3,000	3,655
6	0,2994	0,3811	0,6757	0,8721	1,237	1,635	2,204	3,070	3,828	4,570
7	0,4849	0,5985	0,9893	1,239	1,690	2,167	2,833	3,822	4,671	5,493
8	0,7104	0,8571	1,344	1,646	2,180	2,733	3,490	4,594	5,527	6,423
9	0,9717	1,1520	1,735	2,088	2,700	3,325	4,168	5,380	6,393	7,357
10	1,265	1,479	2,156	2,558	3,247	3,940	4,865	6,179	7,267	8,295
11	1,587	1,834	2,603	3,053	3,816	4,575	5,578	6,989	8,148	9,237
12	1,934	2,214	3,074	3,571	4,404	5,226	6,304	7,807	9,034	10,18
13	2,305	2,617	3,565	4,107	5,009	5,892	7,042	8,634	9,926	11,13
14	2,697	3,041	4,075	4,660	5,629	6,571	7,790	9,467	10,82	12,08
15	3,108	3,483	4,601	5,229	6,262	7,261	8,547	10,31	11,72	13,03
16	3,536	3,942	5,142	5,812	6,908	7,962	9,312	11,15	12,62	13,98
17	3,980	4,416	5,697	6,408	7,564	8,672	10,09	12,00	13,35	14,94
18	4,439	4,905	6,265	7,015	8,231	9,390	10,86	12,86	14,44	15,89
19	4,912	5,407	6,844	7,633	8,907	10,12	11,65	13,72	15,35	16,85
20	5,398	5,921	7,434	8,260	9,591	10,85	12,44	14,58	16,27	17,81
21	5,896	6,447	8,034	8,897	10,28	11,59	13,24	15,44	17,18	18,77
22	6,404	6,983	8,643	9,542	10,98	12,34	14,04	16,31	18,10	19,73
23	6,924	7,529	9,260	10,20	11,69	13,09	14,85	17,19	19,02	20,69
24	7,453	8,085	9,886	10,86	12,40	13,85	15,66	18,06	19,94	21,65
25	7,991	8,649	10,52	11,52	13,12	14,61	16,47	18,94	20,87	22,62
26	8,538	9,222	11,16	12,20	13,84	15,38	17,29	19,82	21,79	23,58
27	9,093	9,803	11,81	12,88	14,57	16,15	18,11	20,70	22,72	24,54
28	9,656	10,39	12,46	13,56	15,31	16,93	18,94	21,59	23,65	25,51
29	10,23	10,99	13,12	14,26	16,05	17,71	19,77	22,48	24,58	26,48
30	10,80	11,59	13,79	14,95	16,79	18,49	20,60	23,36	25,51	27,44
32	11,98	12,81	15,13	16,36	18,29	20,07	22,27	25,15	27,37	29,38
34	13,18	14,06	16,50	17,79	19,81	21,66	23,95	26,94	29,24	31,31
36	14,40	15,32	17,89	19,23	21,34	23,27	25,64	28,73	31,12	33,25
38	15,64	16,61	19,29	20,69	22,88	24,88	27,34	30,54	32,99	35,19
40	16,91	17,92	20,71	22,16	24,43	26,51	29,05	32,34	34,87	37,13
50	23,46	24,67	27,99	29,71	32,36	34,76	37,69	41,45	44,31	46,86
60	30,34	31,74	35,53	37,48	40,48	43,19	46,46	50,64	53,81	56,62
70	37,47	39,04	43,28	45,44	48,76	51,74	55,33	59,90	63,35	66,40
80	44,79	46,52	51,17	53,54	57,15	60,39	64,28	69,21	72,92	76,19
90	52,28	54,16	59,20	61,75	65,65	69,13	73,29	78,56	82,51	85,99
100	59,90	61,92	67,33	70,06	74,22	77,93	82,36	87,95	92,13	95,81

TABLO 6 Kİ-KARE DAĞILIMININ YÜZDELİK DİLİMLERE GÖRE DAĞILIMI (Devam)

P =	50	40	30	20	10	5	2,5	1	0,5	0,1	0,05
v=1	0,4549	0,7083	1,074	1,642	2,706	3,841	5,024	6,635	7,879	10,83	12,12
2	1,386	1,833	2,408	3,219	4,605	5,991	7,378	9,210	10,60	13,82	15,20
3	2,366	2,946	3,665	4,642	6,251	7,815	9,348	11,34	12,84	16,27	17,73
4	3,357	4,045	4,878	5,989	7,779	9,488	11,14	13,28	14,86	18,47	20,00
5	4,351	5,132	6,064	7,289	9,236	11,07	12,83	15,09	16,75	20,52	22,11
6	5,348	6,211	7,231	8,558	10,64	12,59	14,45	16,81	18,55	22,46	24,10
7	6,346	7,283	8,383	9,803	12,02	14,07	16,01	18,48	20,28	24,32	26,02
8	7,344	8,351	9,524	11,03	13,36	15,51	17,53	20,09	21,95	26,12	27,87
9	8,343	9,414	10,66	12,24	14,68	16,92	19,02	21,67	23,59	27,88	29,67
10	9,342	10,47	11,78	13,44	15,99	18,31	20,48	23,21	25,19	29,59	31,42
11	10,34	11,53	12,90	14,63	17,28	19,68	21,92	24,72	26,76	31,26	33,14
12	11,34	12,58	14,01	15,81	18,55	21,03	23,34	26,22	28,30	32,91	34,82
13	12,34	13,64	15,12	16,98	19,81	22,36	24,74	27,69	29,82	34,53	36,48
14	13,34	14,69	16,22	18,15	21,06	23,68	26,12	29,14	31,32	36,12	38,11
15	14,34	15,73	17,32	19,31	22,31	25,00	27,49	30,58	32,80	37,70	39,72
16	15,34	16,78	18,42	20,47	23,54	26,30	28,85	32,00	34,27	39,25	41,31
17	16,34	17,82	19,51	21,61	24,77	27,59	30,19	33,41	35,72	40,79	42,88
18	17,34	18,87	20,60	22,76	25,99	28,87	31,53	34,81	37,16	42,31	44,43
19	18,34	19,91	21,69	23,90	27,20	30,14	32,85	36,19	38,58	43,82	45,97
20	19,34	20,95	22,77	25,04	28,41	31,41	34,17	37,57	40,00	45,31	47,50
21	20,34	21,99	23,86	26,17	29,62	32,67	35,48	38,93	41,40	46,80	49,01
22	21,34	23,03	24,94	27,30	30,81	33,92	36,78	40,29	42,80	48,27	50,51
23	22,34	24,07	26,02	28,43	32,01	35,17	38,08	41,64	44,18	49,73	52,00
24	23,34	25,11	27,10	29,55	33,20	36,42	39,36	42,98	45,56	51,18	53,48
25	24,34	26,14	28,17	30,68	34,38	37,65	40,65	44,31	46,93	52,62	54,95
26	25,34	27,18	29,25	31,79	35,56	38,89	41,92	45,64	48,29	54,05	56,41
27	26,34	28,21	30,32	32,91	36,74	40,11	43,19	46,96	49,64	55,48	57,86
28	27,34	29,25	31,39	34,03	37,92	41,34	44,46	48,28	50,99	56,89	59,30
29	28,34	30,28	32,46	35,14	39,09	42,56	45,72	49,59	52,34	58,30	60,73
30	29,34	31,32	33,53	36,25	40,26	43,77	46,98	50,89	53,67	59,70	62,16
32	31,34	33,38	35,66	38,47	42,58	46,19	49,48	53,49	56,33	62,49	65,00
34	33,34	35,44	37,80	40,68	44,90	48,60	51,97	56,06	58,96	65,25	67,80
36	35,34	37,50	39,92	42,88	47,21	51,00	54,44	58,62	61,58	67,99	70,59
38	37,34	39,56	42,05	45,08	49,51	53,38	56,90	61,16	64,18	70,70	73,35
40	39,34	41,62	44,16	47,27	51,81	55,76	59,34	63,69	66,77	73,40	76,09
50	49,34	51,89	54,72	58,16	63,17	67,50	71,42	76,15	79,49	86,66	89,56
60	59,34	62,13	65,23	68,97	74,40	79,08	83,30	88,38	91,95	99,61	102,7
70	69,34	72,36	75,69	79,71	85,53	90,53	95,02	100,4	104,2	112,3	115,6
80	79,34	82,57	86,12	90,41	96,58	101,9	106,6	112,3	116,3	124,8	128,3
90	89,34	92,76	96,52	101,1	107,6	113,1	118,1	124,1	128,3	137,2	140,8
100	99,34	102,9	106,9	111,7	118,5	124,3	129,6	135,8	140,2	149,4	153,2

TABLO 7 BİNOM DAĞILIM FONKSİYONU

$$P(x) = \frac{n!}{x!(n-x)!} p^x (1-p)^{n-x}$$

n : örneklem boyutu

x : başarı sayısı

$p(x)$: n deneme sonucunda x tane başarı elde etme olasılığı

p	x	0,0100	0,0500	0,1000	0,1500	0,2000	0,2500	0,3000	0,3500	0,4000	0,4500	0,5000
2	0	0,9801	0,9025	0,8100	0,7225	0,6400	0,5625	0,4900	0,4225	0,3600	0,3025	0,2500
	1	0,9999	0,9975	0,9900	0,9775	0,9600	0,9375	0,9100	0,8775	0,8400	0,7975	0,7500
	2	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
3	0	0,9703	0,8573	0,7290	0,6141	0,5120	0,4219	0,3430	0,2746	0,2160	0,1664	0,1250
	1	0,9997	0,9927	0,9720	0,9393	0,8960	0,8438	0,7840	0,7183	0,6480	0,5748	0,5000
	2	1,0000	0,9998	0,9990	0,9966	0,9920	0,9844	0,9730	0,9571	0,9360	0,9089	0,8750
4	3	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	0	0,9606	0,8145	0,6561	0,5220	0,4096	0,3164	0,2401	0,1785	0,1296	0,0915	0,0625
	1	0,9994	0,9859	0,9477	0,8905	0,8192	0,7383	0,6517	0,5630	0,4752	0,3910	0,3125
5	2	1,0000	0,9995	0,9963	0,9880	0,9728	0,9492	0,9163	0,8735	0,8208	0,7585	0,6875
	3	1,0000	0,9999	0,9999	0,9995	0,9984	0,9961	0,9919	0,9850	0,9744	0,9590	0,9375
	4	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
5	0	0,9510	0,7737	0,5905	0,4437	0,3277	0,2373	0,1681	0,1160	0,0778	0,0503	0,0313
	1	0,9990	0,9774	0,9185	0,8352	0,7373	0,6328	0,5282	0,4284	0,3370	0,2562	0,1875
	2	1,0000	0,9988	0,9914	0,9734	0,9421	0,8965	0,8369	0,7648	0,6826	0,5931	0,5000
	3	1,0000	0,9999	0,9995	0,9978	0,9933	0,9844	0,9692	0,9460	0,9130	0,8688	0,8125
	4	1,0000	1,0000	1,0000	0,9999	0,9997	0,9990	0,9976	0,9948	0,9898	0,9816	0,9688
6	5	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	0	0,9415	0,7351	0,5314	0,3772	0,2621	0,1780	0,1177	0,0754	0,0467	0,0277	0,0156
	1	0,9985	0,9672	0,8857	0,7765	0,6554	0,5339	0,4202	0,3191	0,2333	0,1636	0,1094
	2	1,0000	0,9978	0,9842	0,9527	0,9011	0,8306	0,7443	0,6471	0,5443	0,4415	0,3438
	3	1,0000	0,9999	0,9987	0,9941	0,9830	0,9624	0,9295	0,8826	0,8208	0,7447	0,6563
	4	1,0000	1,0000	0,9999	0,9996	0,9984	0,9954	0,9891	0,9777	0,9590	0,9308	0,8906
6	5	1,0000	1,0000	1,0000	1,0000	0,9999	0,9998	0,9993	0,9982	0,9959	0,9917	0,9844
	6	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	0	0,9321	0,6983	0,4783	0,3206	0,2097	0,1335	0,0824	0,0490	0,0280	0,0152	0,0078
	1	0,9980	0,9556	0,8503	0,7166	0,5767	0,4450	0,3294	0,2338	0,1586	0,1024	0,0625
	2	1,0000	0,9962	0,9743	0,9262	0,8520	0,7564	0,6471	0,5323	0,4199	0,3164	0,2266
	3	1,0000	0,9998	0,9973	0,9879	0,9667	0,9294	0,8740	0,8002	0,7102	0,6083	0,5000
7	4	1,0000	1,0000	0,9998	0,9988	0,9953	0,9871	0,9712	0,9444	0,9037	0,8471	0,7734
	5	1,0000	1,0000	1,0000	0,9999	0,9996	0,9987	0,9962	0,9910	0,9812	0,9643	0,9375
	6	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9998	0,9994	0,9984	0,9963	0,9922
	7	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	0	0,9227	0,6634	0,4305	0,2725	0,1678	0,1001	0,0577	0,0319	0,0168	0,0084	0,0039
	1	0,9973	0,9428	0,8131	0,6572	0,5033	0,3671	0,2553	0,1691	0,1064	0,0632	0,0352
	2	1,0000	0,9942	0,9619	0,8948	0,7969	0,6785	0,5518	0,4278	0,3154	0,2201	0,1445
8	3	1,0000	0,9996	0,9950	0,9787	0,9437	0,8862	0,8059	0,7064	0,5941	0,4770	0,3633
	4	1,0000	1,0000	0,9996	0,9972	0,9896	0,9727	0,9420	0,8939	0,8263	0,7396	0,6367
	5	1,0000	1,0000	1,0000	0,9998	0,9988	0,9958	0,9887	0,9747	0,9502	0,9115	0,8555
	6	1,0000	1,0000	1,0000	1,0000	0,9999	0,9996	0,9987	0,9964	0,9915	0,9819	0,9648
	7	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9998	0,9998	0,9993	0,9983	0,9961
	8	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000

TABLO 7 BİNOM DAĞILIM FONKSİYONU (Devam)

$$P(x) = \frac{n!}{x!(n-x)!} p^x (1-p)^{n-x}$$

n : örneklem boyutu
x : başarı sayısı
p(x): n deneme sonucunda x tane başarı elde etme olasılığı

p	n	x	0,0100	0,0500	0,1000	0,1500	0,2000	0,2500	0,3000	0,3500	0,4000	0,4500	0,5000
9	0		0,9135	0,6303	0,3874	0,2316	0,1342	0,0751	0,0404	0,0207	0,0101	0,0046	0,0020
	1		0,9966	0,9288	0,7748	0,5995	0,4362	0,3003	0,1960	0,1211	0,0705	0,0385	0,0195
	2		0,9999	0,9916	0,9470	0,8592	0,7382	0,6007	0,4628	0,3373	0,2318	0,1495	0,0898
	3		1,0000	0,9994	0,9917	0,9661	0,9144	0,8343	0,7297	0,6089	0,4826	0,3614	0,2539
	4		1,0000	1,0000	0,9991	0,9944	0,9804	0,9511	0,9012	0,8283	0,7334	0,6214	0,5000
	5		1,0000	1,0000	0,9999	0,9994	0,9969	0,9900	0,9747	0,9464	0,9007	0,8342	0,7461
	6		1,0000	1,0000	1,0000	1,0000	0,9997	0,9987	0,9957	0,9888	0,9750	0,9502	0,9102
	7		1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9996	0,9986	0,9962	0,9909	0,9805
	8		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9997	0,9992	0,9981
	9		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
10	0		0,9044	0,5987	0,3487	0,1969	0,1074	0,0563	0,0283	0,0135	0,0061	0,0025	0,0010
	1		0,9957	0,9139	0,7361	0,5443	0,3758	0,2440	0,1493	0,0860	0,0464	0,0233	0,0107
	2		0,9999	0,9885	0,9298	0,8202	0,6778	0,5256	0,3828	0,2616	0,1673	0,0996	0,0547
	3		1,0000	0,9990	0,9872	0,9500	0,8791	0,7759	0,6496	0,5138	0,3823	0,2660	0,1719
	4		1,0000	0,9999	0,9984	0,9901	0,9672	0,9219	0,8497	0,7515	0,6331	0,5044	0,3770
	5		1,0000	1,0000	0,9999	0,9986	0,9936	0,9803	0,9527	0,9051	0,8338	0,7384	0,6231
	6		1,0000	1,0000	1,0000	0,9999	0,9991	0,9965	0,9894	0,9740	0,9452	0,8980	0,8281
	7		1,0000	1,0000	1,0000	1,0000	0,9999	0,9996	0,9984	0,9952	0,9877	0,9726	0,9453
	8		1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9999	0,9995	0,9983	0,9955	0,9893
	9		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9997	0,9990
	10		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
11	0		0,8953	0,5688	0,3138	0,1673	0,0859	0,0422	0,0198	0,0088	0,0036	0,0014	0,0005
	1		0,9948	0,8981	0,6974	0,4922	0,3221	0,1971	0,1130	0,0606	0,0302	0,0139	0,0059
	2		0,9998	0,9848	0,9104	0,7788	0,6174	0,4552	0,3127	0,2001	0,1189	0,0652	0,0327
	3		1,0000	0,9985	0,9815	0,9306	0,8389	0,7133	0,5696	0,4256	0,2963	0,1911	0,1133
	4		1,0000	0,9999	0,9973	0,9841	0,9496	0,8854	0,7897	0,6683	0,5328	0,3971	0,2744
	5		1,0000	1,0000	0,9997	0,9973	0,9884	0,9657	0,9218	0,8513	0,7535	0,6331	0,5000
	6		1,0000	1,0000	1,0000	0,9997	0,9980	0,9924	0,9784	0,9499	0,9007	0,8262	0,7256
	7		1,0000	1,0000	1,0000	1,0000	0,9998	0,9988	0,9957	0,9878	0,9707	0,9390	0,8867
	8		1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9994	0,9980	0,9941	0,9852	0,9673
	9		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9998	0,9993	0,9978	0,9941
	10		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9995
	11		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
12	0		0,8864	0,5404	0,2824	0,1422	0,0687	0,0317	0,0138	0,0057	0,0022	0,0008	0,0002
	1		0,9938	0,8816	0,6590	0,4435	0,2749	0,1584	0,0850	0,0424	0,0196	0,0083	0,0032
	2		0,9998	0,9804	0,8891	0,7358	0,5584	0,3907	0,2528	0,1513	0,0834	0,0421	0,0193
	3		1,0000	0,9978	0,9744	0,9078	0,7946	0,6488	0,4925	0,3467	0,2253	0,1345	0,0730
	4		1,0000	0,9998	0,9957	0,9761	0,9274	0,8424	0,7237	0,5834	0,4382	0,3044	0,1939
	5		1,0000	1,0000	0,9995	0,9954	0,9806	0,9456	0,8822	0,7873	0,6652	0,5269	0,3872
	6		1,0000	1,0000	1,0000	0,9993	0,9961	0,9858	0,9614	0,9154	0,8418	0,7393	0,6128
	7		1,0000	1,0000	1,0000	0,9999	0,9972	0,9905	0,9745	0,9427	0,8883	0,8062	0,7062
	8		1,0000	1,0000	1,0000	1,0000	0,9999	0,9996	0,9983	0,9944	0,9847	0,9644	0,9270

TABLO 7 BİNOM DAĞILIM FONKSİYONU (Devam)

$$P(x) = \frac{n!}{x!(n-x)!} p^x (1-p)^{n-x}$$

n : örneklem boyutu
x : başarı sayısı
p(x): n deneme sonucunda x tane başarı elde etme olasılığı

p	n	x	0,0100	0,0500	0,1000	0,1500	0,2000	0,2500	0,3000	0,3500	0,4000	0,4500	0,5000
12	9	1	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9998	0,9992	0,9972	0,9921	0,9807
	10	1	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9997	0,9989	0,9968
	11	1	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9998
	12	1	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
13	0	0	0,8775	0,5133	0,2542	0,1209	0,0550	0,0238	0,0097	0,0037	0,0013	0,0004	0,0001
	1	0	0,9928	0,8646	0,6213	0,3983	0,2337	0,1267	0,0637	0,0296	0,0126	0,0049	0,0017
	2	0	0,9997	0,9755	0,8661	0,6920	0,5017	0,3326	0,2025	0,1132	0,0579	0,0269	0,0112
	3	0	1,0000	0,9969	0,9658	0,8820	0,7473	0,5843	0,4206	0,2783	0,1686	0,0929	0,0461
	4	0	1,0000	0,9997	0,9935	0,9658	0,9009	0,7940	0,6543	0,5005	0,3530	0,2280	0,1334
	5	0	1,0000	1,0000	0,9991	0,9925	0,9700	0,9198	0,8346	0,7159	0,5744	0,4268	0,2905
	6	0	1,0000	1,0000	0,9999	0,9987	0,9930	0,9757	0,9376	0,8705	0,7712	0,6437	0,5000
	7	0	1,0000	1,0000	1,0000	0,9998	0,9988	0,9944	0,9818	0,9538	0,9023	0,8212	0,7095
	8	0	1,0000	1,0000	1,0000	1,0000	0,9998	0,9990	0,9960	0,9874	0,9679	0,9302	0,8666
	9	0	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9994	0,9975	0,9922	0,9797	0,9539
	10	0	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9997	0,9987	0,9959	0,9888
	11	0	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9995	0,9983
	12	0	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999
	13	0	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
14	0	0	0,8688	0,4877	0,2288	0,1028	0,0440	0,0178	0,0068	0,0024	0,0008	0,0002	0,0001
	1	0	0,9916	0,8470	0,5846	0,3567	0,1979	0,1010	0,0475	0,0205	0,0081	0,0029	0,0009
	2	0	0,9997	0,9700	0,8416	0,6479	0,4481	0,2811	0,1608	0,0839	0,0398	0,0170	0,0065
	3	0	1,0000	0,9958	0,9559	0,8535	0,6982	0,5213	0,3552	0,2205	0,1243	0,0632	0,0287
	4	0	1,0000	0,9996	0,9908	0,9533	0,8702	0,7415	0,5842	0,4227	0,2793	0,1672	0,0898
	5	0	1,0000	1,0000	0,9985	0,9885	0,9562	0,8883	0,7805	0,6405	0,4859	0,3373	0,2120
	6	0	1,0000	1,0000	0,9998	0,9978	0,9884	0,9617	0,9067	0,8164	0,6925	0,5461	0,3953
	7	0	1,0000	1,0000	1,0000	0,9997	0,9976	0,9897	0,9685	0,9247	0,8499	0,7414	0,6047
	8	0	1,0000	1,0000	1,0000	1,0000	0,9996	0,9979	0,9917	0,9757	0,9417	0,8811	0,7880
	9	0	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9983	0,9940	0,9825	0,9574	0,9102
	10	0	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9998	0,9989	0,9961	0,9886	0,9713
	11	0	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9994	0,9979	0,9935
	12	0	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9998	0,9991
	13	0	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999
	14	0	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000

TABLO 7 BİNOM DAĞILIM FONKSİYONU (Devam)

$$P(x) = \frac{n!}{x!(n-x)!} p^x (1-p)^{n-x}$$

n : örneklem boyutu
x : başarı sayısı
p(x): n deneme sonucunda x tane başarı elde etme olasılığı

p	n	x	0,0100	0,0500	0,1000	0,1500	0,2000	0,2500	0,3000	0,3500	0,4000	0,4500	0,5000
15	0		0,8601	0,4633	0,2059	0,0874	0,0352	0,0134	0,0048	0,0016	0,0005	0,0001	0,0000
	1		0,9904	0,8291	0,5490	0,3186	0,1671	0,0802	0,0353	0,0142	0,0052	0,0017	0,0005
	2		0,9996	0,9638	0,8159	0,6042	0,3980	0,2361	0,1268	0,0617	0,0271	0,0107	0,0037
	3		1,0000	0,9945	0,9444	0,8227	0,6482	0,4613	0,2969	0,1727	0,0905	0,0424	0,0176
	4		1,0000	0,9994	0,9873	0,9383	0,8358	0,6865	0,5155	0,3519	0,2173	0,1204	0,0592
	5		1,0000	1,0000	0,9978	0,9832	0,9390	0,8516	0,7216	0,5643	0,4032	0,2608	0,1509
	6		1,0000	1,0000	0,9997	0,9964	0,9819	0,9434	0,8689	0,7548	0,6098	0,4522	0,3036
	7		1,0000	1,0000	1,0000	0,9994	0,9958	0,9827	0,9500	0,8868	0,7869	0,6535	0,5000
	8		1,0000	1,0000	1,0000	0,9999	0,9992	0,9958	0,9848	0,9578	0,9050	0,8182	0,6964
	9		1,0000	1,0000	1,0000	1,0000	0,9999	0,9992	0,9964	0,9876	0,9662	0,9231	0,8491
	10		1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9993	0,9972	0,9907	0,9745	0,9408
	11		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9995	0,9981	0,9937	0,9824
	12		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9997	0,9989	0,9963
	13		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9995
	14		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	15		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
20	0		0,8179	0,3585	0,1216	0,0388	0,0115	0,0032	0,0008	0,0002	0,0000	0,0000	0,0000
	1		0,9831	0,7358	0,3918	0,1756	0,0692	0,0243	0,0076	0,0021	0,0005	0,0001	0,0000
	2		0,9990	0,9245	0,6769	0,4049	0,2061	0,0913	0,0355	0,0121	0,0036	0,0009	0,0002
	3		1,0000	1,0000	0,9999	0,9987	0,9900	0,9591	0,8867	0,7624	0,5956	0,4143	0,2517
	4		1,0000	1,0000	1,0000	0,9998	0,9974	0,9861	0,9520	0,8782	0,7553	0,5914	0,4119
	5		1,0000	1,0000	1,0000	1,0000	0,9994	0,9961	0,9829	0,9468	0,8725	0,7507	0,5881
	6		1,0000	1,0000	1,0000	1,0000	0,9999	0,9991	0,9949	0,9804	0,9435	0,8692	0,7483
	7		1,0000	1,0000	1,0000	1,0000	1,0000	0,9998	0,9987	0,9940	0,9790	0,9420	0,8684
	8		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9985	0,9935	0,9786	0,9423
	9		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9936	0,9793
	10		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9985	0,9941
	11		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9987
	12		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9998
	13		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	14		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	15		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	16		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	17		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	18		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	19		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	20		1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000

TABLO 7 BİNOM DAĞILIM FONKSİYONU (Devam)

$$P(x) = \frac{n!}{x!(n-x)!} p^x (1-p)^{n-x}$$

n : örneklem boyutu
x : başarı sayısı
p(x): n deneme sonucunda x tane başarı elde etme olasılığı

p	n	x	0,0100	0,0500	0,1000	0,1500	0,2000	0,2500	0,3000	0,3500	0,4000	0,4500	0,5000
25	0	0	0,7778	0,2774	0,0718	0,0172	0,0038	0,0008	0,0001	0,0000	0,0000	0,0000	0,0000
	1	1	0,9742	0,6424	0,2712	0,0931	0,0274	0,0070	0,0016	0,0003	0,0001	0,0000	0,0000
	2	2	0,9981	0,8729	0,5371	0,2537	0,0982	0,0321	0,0090	0,0021	0,0004	0,0001	0,0000
	3	3	0,9999	0,9659	0,7636	0,4711	0,2340	0,0962	0,0332	0,0097	0,0024	0,0005	0,0001
	4	4	1,0000	0,9928	0,9020	0,6821	0,4207	0,2137	0,0905	0,0321	0,0095	0,0023	0,0005
	5	5	1,0000	0,9988	0,9666	0,8385	0,6167	0,3783	0,1935	0,0826	0,0294	0,0086	0,0020
	6	6	1,0000	0,9998	0,9905	0,9305	0,7800	0,5611	0,3407	0,1734	0,0736	0,0258	0,0073
	7	7	1,0000	1,0000	0,9977	0,9745	0,8909	0,7265	0,5119	0,3061	0,1536	0,0639	0,0216
	8	8	1,0000	1,0000	0,9995	0,9920	0,9532	0,8506	0,6769	0,4668	0,2735	0,1340	0,0539
	9	9	1,0000	1,0000	0,9999	0,9979	0,9827	0,9287	0,8106	0,6303	0,4246	0,2424	0,1148
	10	10	1,0000	1,0000	1,0000	0,9995	0,9945	0,9703	0,9022	0,7712	0,5858	0,3843	0,2122
	11	11	1,0000	1,0000	1,0000	0,9999	0,9985	0,9893	0,9558	0,8746	0,7323	0,5426	0,3450
	12	12	1,0000	1,0000	1,0000	1,0000	0,9996	0,9966	0,9825	0,9396	0,8462	0,6937	0,5000
	13	13	1,0000	1,0000	1,0000	1,0000	0,9999	0,9991	0,9940	0,9745	0,9222	0,8173	0,6550
	14	14	1,0000	1,0000	1,0000	1,0000	1,0000	0,9998	0,9982	0,9907	0,9656	0,9040	0,7878
	15	15	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9996	0,9971	0,9868	0,9560	0,8852
	16	16	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9992	0,9957	0,9826	0,9461
	17	17	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9998	0,9988	0,9942	0,9784
	18	18	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9984	0,9927
	19	19	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9996	0,9980
	20	20	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9995
	21	21	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999
	22	22	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	23	23	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	24	24	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	25	25	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
50	0	0	0,6050	0,0769	0,0052	0,0003	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000
	1	1	0,9106	0,2794	0,0338	0,0029	0,0002	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000
	2	2	0,9862	0,5405	0,1117	0,0142	0,0013	0,0001	0,0000	0,0000	0,0000	0,0000	0,0000
	3	3	0,9984	0,7604	0,2503	0,0461	0,0057	0,0005	0,0000	0,0000	0,0000	0,0000	0,0000
	4	4	0,9999	0,8964	0,4312	0,1121	0,0185	0,0021	0,0002	0,0000	0,0000	0,0000	0,0000
	5	5	1,0000	0,9622	0,6161	0,2194	0,0480	0,0071	0,0007	0,0001	0,0000	0,0000	0,0000
	6	6	1,0000	0,9882	0,7702	0,3613	0,1034	0,0194	0,0025	0,0002	0,0000	0,0000	0,0000
	7	7	1,0000	0,9968	0,8779	0,5188	0,1904	0,0453	0,0073	0,0008	0,0001	0,0000	0,0000
	8	8	1,0000	0,9992	0,9421	0,6681	0,3073	0,0916	0,0183	0,0025	0,0002	0,0000	0,0000
	9	9	1,0000	0,9998	0,9755	0,7911	0,4437	0,1637	0,0402	0,0067	0,0008	0,0001	0,0000
	10	10	1,0000	1,0000	0,9907	0,8801	0,5836	0,2622	0,0789	0,0160	0,0022	0,0002	0,0000
	11	11	1,0000	1,0000	0,9968	0,9372	0,7107	0,3816	0,1390	0,0342	0,0057	0,0006	0,0001
	12	12	1,0000	1,0000	0,9990	0,9699	0,8139	0,5110	0,2229	0,0661	0,0133	0,0018	0,0002
	13	13	1,0000	1,0000	0,9997	0,9868	0,8894	0,6370	0,3279	0,1163	0,0280	0,0045	0,0005
	14	14	1,0000	1,0000	0,9999	0,9947	0,9393	0,7481	0,4468	0,1878	0,0540	0,0104	0,0013
	15	15	1,0000	1,0000	1,0000	0,9981	0,9692	0,8369	0,5692	0,2801	0,0955	0,0220	0,0033
	16	16	1,0000	1,0000	1,0000	0,9993	0,9856	0,9017	0,6839	0,3889	0,1561	0,0427	0,0077
	17	17	1,0000	1,0000	1,0000	0,9998	0,9937	0,9449	0,7822	0,5060	0,2369	0,0765	0,0164
	18	18	1,0000	1,0000	1,0000	0,9999	0,9975	0,9713	0,8594	0,6216	0,3356	0,1273	0,0325
	19	19	1,0000	1,0000	1,0000	0,9991	0,9861	0,9152	0,7264	0,4465	0,1974	0,0595	0,0103
	20	20	1,0000	1,0000	1,0000	0,9997	0,9937	0,9522	0,8139	0,5610	0,2862	0,1013	0,0103

TABLO 7 BİNOM DAĞILIM FONKSİYONU (Devam)

$$P(x) = \frac{n!}{x!(n-x)!} p^x (1-p)^{n-x}$$

n : örneklem boyutu
x : başarı sayısı
p(x): n deneme sonucunda x tane başarı elde etme olasılığı

p	n	x	0,0100	0,0500	0,1000	0,1500	0,2000	0,2500	0,3000	0,3500	0,4000	0,4500	0,5000
50	21	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9974	0,9749	0,8813	0,6701	0,3900	0,1611
	22	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9990	0,9877	0,9290	0,7660	0,5019	0,2399
	23	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9996	0,9944	0,9604	0,8438	0,6134	0,3359
	24	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9976	0,9793	0,9022	0,7160	0,4439
	25	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9991	0,9900	0,9427	0,8034	0,5561
	26	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9955	0,9686	0,8721	0,6641
	27	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9981	0,9840	0,9220	0,7601
	28	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9993	0,9924	0,9556	0,8389
	29	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9966	0,9765	0,8987
	30	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9986	0,9884	0,9405
	31	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9995	0,9947	0,9676
	32	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9998	0,9978	0,9836
	33	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9991	0,9923
	34	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9967
	35	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9987
	36	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9995
	37	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999
	38	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	39	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	40	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	41	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	42	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	43	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	44	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	45	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	46	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	47	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	48	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	49	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
	50	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
100	0	0,3660	0,0059	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000
	1	0,7358	0,0371	0,0003	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000
	2	0,9206	0,1183	0,0019	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000
	3	0,9816	0,2578	0,0078	0,0001	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000
	4	0,9966	0,4360	0,0237	0,0004	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000
	5	0,9995	0,6160	0,0576	0,0016	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000
	6	0,9999	0,7660	0,1172	0,0047	0,0001	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000
	7	1,0000	0,8720	0,2061	0,0122	0,0003	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000
	8	1,0000	0,9369	0,3209	0,0275	0,0009	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000
	9	1,0000	0,9718	0,4513	0,0551	0,0023	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000
	10	1,0000	0,9885	0,5832	0,0995	0,0057	0,0001	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000
	11	1,0000	0,9957	0,7030	0,1635	0,0126	0,0004	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000
	12	1,0000	0,9985	0,8018	0,2473	0,0253	0,0010	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000
	13	1,0000	0,9995	0,8761	0,3474	0,0469	0,0025	0,0001	0,0000	0,0000	0,0000	0,0000	0,0000
	14	1,0000	0,9999	0,9274	0,4572	0,0804	0,0054	0,0002	0,0000	0,0000	0,0000	0,0000	0,0000
	15	1,0000	1,0000	0,9601	0,5683	0,1285	0,0111	0,0004	0,0000	0,0000	0,0000	0,0000	0,0000

TABLO 7 BINOM DAĞILIM FONKSİYONU (Devam)

$$P(x) = \frac{n!}{x!(n-x)!} p^x (1-p)^{n-x}$$

n : örneklem boyutu
x : başarı sayısı
p(x): n deneme sonucunda x tane başarı elde etme olasılığı

p	n	x	0,0100	0,0500	0,1000	0,1500	0,2000	0,2500	0,3000	0,3500	0,4000	0,4500	0,5000
100	16	1,0000	1,0000	0,9794	0,6725	0,1923	0,0211	0,0010	0,0000	0,0000	0,0000	0,0000	0,0000
	17	1,0000	1,0000	0,9900	0,7633	0,2712	0,0376	0,0022	0,0001	0,0000	0,0000	0,0000	0,0000
	18	1,0000	1,0000	0,9954	0,8372	0,3621	0,0630	0,0045	0,0001	0,0000	0,0000	0,0000	0,0000
	19	1,0000	1,0000	0,9980	0,8935	0,4602	0,0995	0,0089	0,0003	0,0000	0,0000	0,0000	0,0000
	20	1,0000	1,0000	0,9992	0,9337	0,5595	0,1488	0,0165	0,0008	0,0000	0,0000	0,0000	0,0000
	21	1,0000	1,0000	0,9997	0,9607	0,6540	0,2114	0,0288	0,0017	0,0000	0,0000	0,0000	0,0000
	22	1,0000	1,0000	0,9999	0,9779	0,7389	0,2864	0,0479	0,0034	0,0001	0,0000	0,0000	0,0000
	23	1,0000	1,0000	1,0000	0,9881	0,8109	0,3711	0,0755	0,0066	0,0003	0,0000	0,0000	0,0000
	24	1,0000	1,0000	1,0000	0,9939	0,8687	0,4617	0,1136	0,0121	0,0006	0,0000	0,0000	0,0000
	25	1,0000	1,0000	1,0000	0,9970	0,9125	0,5535	0,1631	0,0211	0,0012	0,0000	0,0000	0,0000
	26	1,0000	1,0000	1,0000	0,9986	0,9442	0,6417	0,2244	0,0351	0,0024	0,0001	0,0000	0,0000
	27	1,0000	1,0000	1,0000	0,9994	0,9659	0,7224	0,2964	0,0558	0,0046	0,0002	0,0000	0,0000
	28	1,0000	1,0000	1,0000	0,9997	0,9800	0,7925	0,3768	0,0848	0,0084	0,0004	0,0000	0,0000
	29	1,0000	1,0000	1,0000	0,9999	0,9888	0,8505	0,4623	0,1236	0,0148	0,0008	0,0000	0,0000
	30	1,0000	1,0000	1,0000	1,0000	0,9939	0,8962	0,5491	0,1730	0,0248	0,0015	0,0000	0,0000
	31	1,0000	1,0000	1,0000	1,0000	0,9969	0,9307	0,6331	0,2331	0,0399	0,0030	0,0001	0,0000
	32	1,0000	1,0000	1,0000	1,0000	0,9985	0,9554	0,7107	0,3029	0,0615	0,0055	0,0002	0,0000
	33	1,0000	1,0000	1,0000	1,0000	0,9993	0,9724	0,7793	0,3803	0,0913	0,0098	0,0004	0,0000
	34	1,0000	1,0000	1,0000	1,0000	0,9997	0,9836	0,8371	0,4624	0,1303	0,0166	0,0009	0,0000
	35	1,0000	1,0000	1,0000	1,0000	0,9999	0,9906	0,8839	0,5458	0,1795	0,0272	0,0018	0,0000
	36	1,0000	1,0000	1,0000	1,0000	0,9999	0,9948	0,9201	0,6269	0,2386	0,0429	0,0033	0,0000
	37	1,0000	1,0000	1,0000	1,0000	1,0000	0,9973	0,9470	0,7025	0,3068	0,0651	0,0060	0,0000
	38	1,0000	1,0000	1,0000	1,0000	1,0000	0,9986	0,9660	0,7699	0,3822	0,0951	0,0105	0,0000
	39	1,0000	1,0000	1,0000	1,0000	1,0000	0,9993	0,9790	0,8276	0,4621	0,1343	0,0176	0,0000
	40	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9875	0,8750	0,5433	0,1831	0,0284	0,0000
	41	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9928	0,9123	0,6225	0,2415	0,0443	0,0000
	42	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9960	0,9406	0,6967	0,3087	0,0666	0,0000
	43	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9979	0,9611	0,7635	0,3828	0,0967	0,0000
	44	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9989	0,9754	0,8211	0,4613	0,1356	0,0000
	45	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9995	0,9850	0,8689	0,5413	0,1841	0,0000
	46	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9997	0,9912	0,9070	0,6196	0,2421	0,0000
	47	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9950	0,9362	0,6931	0,3087	0,0000
	48	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9973	0,9577	0,7596	0,3822	0,0000
	49	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9986	0,9729	0,8173	0,4602	0,0000
	50	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9993	0,9832	0,8654	0,5398	0,0000
	51	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9996	0,9900	0,9041	0,6178	0,0000
	52	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9998	0,9942	0,9338	0,6914	0,0000
	53	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9968	0,9559	0,7579	0,0000
	54	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9983	0,9716	0,8159	0,0000
	55	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9991	0,9824	0,8644	0,0000
	56	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9996	0,9894	0,9033	0,0000
	57	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9998	0,9939	0,9334	0,0000
	58	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9966	0,9557	0,0000
	59	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9982	0,9716	0,0000
	60	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9991	0,9824	0,0000
	61	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9995	0,9895	0,0000
	62	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9998	0,9940	0,0000
	63	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9999	0,9967	0,0000
	64	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9982	0,0000
	65	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000	0,9991	0,0000

TABLO 7 BİNOM DAĞILIM FONKSİYONU (Devam)

n : örneklem boyutu
 x : başarı sayısı
 $p(x)$: n deneme sonucunda x tane başarı elde etme olasılığı

[illegible]